

Enhancing K-Nearest Neighbors for Categorical Data Classification Using Similarity-Based Distance Metrics

Muhammad Idrees, Basharat Ahmad Hassan, Babar Saeed, Muhammad Zubair, Nadeem Khan

The University of Agriculture, Peshawar, Pakistan

*Correspondence: basharat94@gmail.com

Citation | Idrees. M, Hassan. B. A, Saeed. B, Zubair. M, Khan. N, "Enhancing K-Nearest Neighbors for Categorical Data Classification Using Similarity-Based Distance Metrics", IJIST, Vol. 8 Issue. 5 pp 1839-1866, August 2026

DOI | <https://doi.org/10.33411/IJIST/1965>

Received | July 15, 2026 **Revised** | Aug 06, 2026 **Accepted** | Aug 07, 2026 **Published** | Aug 13, 2026.

Categorical data are pervasive in real-world classification problems, yet most machine learning algorithms, including support vector machines, decision trees and neural networks, are formulated for numeric feature spaces. The K-Nearest Neighbor (KNN) classifier is particularly affected because its decision rule depends entirely on a distance function that is not naturally defined over unordered categorical values. Encoding categorical attributes into numeric form imposes an artificial ordering and causes information loss, which motivates the use of similarity measures designed specifically for categorical attributes. This study presents a controlled comparative evaluation of three similarity-based KNN variants under an identical experimental protocol: Lin-KNN, which uses matching-based similarity; VDM-KNN, which uses class-conditional probability differences; and Tanimoto-KNN, which uses overlap-based similarity. The three models were assessed on two publicly available categorical benchmark datasets obtained from Kaggle, namely the Mushroom dataset (edible versus poisonous) and an Android permission-based malware detection dataset (benign versus malicious). Mutual-information based feature selection was applied independently for each similarity measure, the number of neighbours was tuned by grid search over stratified cross-validation on the training partition only, and performance was reported using accuracy, precision, recall and F1-score. Lin-KNN achieved the highest overall performance with an average accuracy of 95.56%, precision of 95.57%, recall of 95.55% and F1-score of 95.56%. VDM-KNN was highly competitive with an average accuracy of 94.44%, precision of 94.45%, recall of 94.43% and F1-score of 94.43%, and was the stronger of the two on the malware dataset. Tanimoto-KNN was consistently weaker, with an average accuracy of 81.44%, precision of 85.22%, recall of 81.56% and F1-score of 80.14%. Differences between the models were further examined through repeated stratified cross-validation and non-parametric significance testing. The results demonstrate that no single similarity measure dominates across categorical domains: matching-based similarity is preferable for strongly structured attribute spaces in which attribute agreement aligns closely with class separability, whereas class-conditional similarity is preferable for sparse, weakly discriminative binary attribute spaces such as Android permissions. These findings provide practical guidance for selecting an appropriate KNN variant in real-world machine learning applications involving categorical data.

Keywords: Machine Learning; Classification; K-Nearest Neighbors (KNN); Lin Similarity; Value Difference Metric (VDM); Tanimoto Similarity; Categorical Data Classification; Android Malware Detection.



Introduction:

Categorical data consist of non-numeric labels such as odour, gill colour, malware family or stalk shape. Such attributes are usually grouped into three types. Nominal attributes have no natural ordering, for example colour with values white, blue and black. Ordinal attributes are ranked in a meaningful way, for example low, medium and high. Binary attributes admit exactly two values, for example male and female or low income and high income [1][2]. Analysing categorical data therefore requires an understanding of how observations are distributed across categories and how the categories of different attributes relate to one another, which is precisely the information that a nearest-neighbour classifier must exploit. In supervised learning, a model is presented with labelled instances, learns the association between input attributes and class labels, and then assigns labels to previously unseen instances. Classifiers built in this way are used for spam filtering, sentiment analysis, image recognition and medical diagnosis, and are commonly assessed using accuracy, precision, recall and F1-score.

The K-Nearest Neighbors (KNN) algorithm is one of the most widely used instance-based classifiers, but it inherits three well-documented weaknesses that are amplified when the attributes are categorical. First, every attribute contributes equally to the distance computation, so irrelevant or noisy attributes distort the neighbourhood structure and degrade classification quality. Second, the curse of dimensionality erodes the contrast between distances as the number of attributes grows, which makes the notion of a nearest neighbour progressively less informative [3]. Third, performance is sensitive to the number of neighbours K , and an inappropriate value can degrade accuracy substantially, particularly on complex or imbalanced data in which the majority class dominates the neighbourhood vote and minority-class instances are systematically ignored [4]. Beyond these general limitations, the decisive problem for categorical data is that KNN is defined through a distance function that presupposes an ordered, metric feature space. The usual workaround is to encode categorical attributes numerically, but ordinal or one-hot encoding either injects an ordering that does not exist in the data or inflates dimensionality, and in both cases the relationship between attribute values and class labels is weakened rather than preserved.

A similarity measure quantifies how alike two objects are, and is the component that determines which instances a KNN classifier treats as neighbours[5]. Measures developed for numeric spaces include Euclidean distance, the more general Minkowski distance, Manhattan distance, which is comparatively robust to outliers and is often used in K-Medoids clustering [6], Canberra distance, Bray-Curtis distance and Chi-Square distance [7]. Measures developed for set-valued or binary representations include cosine similarity, which compares orientation rather than magnitude [8], the Dice coefficient, which quantifies overlap between two sets [9], the overlap coefficient, which measures the intersection of two probability densities [10], and the simple matching coefficient, which compares the proportion of attributes taking identical values [11]. None of these was designed for unordered categorical attributes, which is why specialised variants of KNN have been proposed. VDM-KNN measures the distance between two attribute values through the difference in their class-conditional probabilities; Lin-KNN measures similarity through the proportion of attributes on which two instances agree; and Tanimoto-KNN measures similarity through the shared attributes of two instances relative to the union of their attributes. These three measures were selected for the present study because they represent three fundamentally distinct and mutually exclusive families of reasoning about categorical similarity, namely class-supervised, matching-based and overlap-based reasoning. Comparing them therefore isolates the effect of the similarity formulation itself rather than the effect of an arbitrary collection of measures.

Research Gap and Motivation:

Although a considerable body of work has proposed similarity measures for categorical data, three gaps remain. First, most of the literature evaluates similarity measures in an unsupervised clustering setting rather than in supervised instance-based classification, so the interaction between a similarity measure and the majority-vote decision rule of KNN is not well characterised[12]. Second, comparative studies frequently vary the classifier, the preprocessing pipeline and the similarity measure simultaneously, which makes it impossible to attribute an observed performance difference to the similarity formulation alone. Third, reported comparisons are commonly based on a single train-test partition without repeated runs, dispersion estimates or significance testing, so it is unclear whether the differences reported between measures are reproducible or merely artefacts of a particular split. Existing critiques of similarity-based methods also tend to enumerate limitations rather than explain them: it is widely observed that no single measure performs best across datasets, but the structural properties of a dataset that favour one measure over another are rarely identified. This study addresses these gaps by holding the classifier, the preprocessing pipeline, the evaluation metrics and the validation protocol constant, varying only the similarity measure, and by relating the resulting performance differences to measurable properties of the two datasets.

Novelty Statement:

The novelty of this work lies in the following contributions, which distinguish it from the prior comparative studies.

A single-factor experimental design in which the classifier, preprocessing, feature-selection criterion, neighbourhood size search space, validation protocol and evaluation metrics are identical across all three models, so that the similarity measure is the only manipulated variable.

The joint evaluation of three structurally distinct similarity families, namely supervised class-conditional (VDM), matching-based (Lin) and overlap-based (Tanimoto), within a single KNN framework, rather than the comparison of numeric distance functions that dominates the existing KNN literature.

The use of two categorical datasets with deliberately contrasting attribute semantics, one with multi-valued descriptive nominal attributes (Mushroom) and one with sparse binary permission indicators (Android malware), which allows the performance reversal between Lin-KNN and VDM-KNN to be attributed to identifiable data characteristics rather than to chance.

A statistically grounded comparison based on repeated stratified cross-validation with dispersion estimates and non-parametric significance testing, in place of the single-split point estimates that are typical of comparable studies.

A set of evidence-based selection guidelines that map measurable dataset properties, such as attribute cardinality, sparsity and class-conditional attribute informativeness, to the similarity measure expected to perform best.

Research Objectives:

The specific objectives of this study are as follows:

To implement three similarity-based variants of the K-Nearest Neighbors classifier, namely Lin-KNN, VDM-KNN and Tanimoto-KNN, within a common experimental framework.

To evaluate and compare the classification performance of the three models on two publicly available categorical benchmark datasets using accuracy, precision, recall and F1-score.

To determine whether the observed performance differences between the three similarity measures are statistically significant under repeated stratified cross-validation.

To analyse the relationship between the structural characteristics of a categorical dataset and the similarity measure that performs best on it.

To derive practical guidelines for selecting an appropriate similarity-based KNN variant for real-world categorical classification tasks such as food-safety screening and Android malware detection.

Organisation of the Paper:

The remainder of this paper is organised as follows. Reviews clustering methods, similarity measures and classifiers for categorical data, with emphasis on developments reported between 2021 and 2025. It describes the datasets, preprocessing, feature selection, similarity measures, model architectures, hyperparameter settings, validation protocol and implementation environment. The experimental results, including the statistical significance analysis and a comparison with recent state-of-the-art methods. Then discusses the findings, their theoretical explanation, their practical implications and the limitations of the study. concludes the paper and outlines future work.

Literature Review:

Machine learning algorithms are mostly designed for numerical data, making it one of the most complex problems of categorizing categorical data. The traditional distance measures, like Euclidean and Manhattan do not account for the relationships of categorical features that are not intrinsically ordered or continuous. To overcome such restrictions researchers have developed further clustering methods specifically tailored similarity metrics and adapted classification models to enhance the performance of models for categorical data in terms of predictive accuracy, robustness and interpretability.

Unfortunately, broadly used k-means algorithm is not easily applied to categorical data as it requires computation of distance and centroids, which are unsuitable for direct use in categorical contexts [13]. K-modes Clustering: K-modes is a categorical variant of k-means that uses matching based dissimilarity and modes for clustering. As it goes, it assigns points and updates clusters. Sensitive to initialization that is not suitable for mixed data: treats all features equally and does not work well with high dimensional data [14]. The study suggested the proposed clustering approach for categorical data and validated the clustering results with the cluster validation techniques. To obtain the number of clusters that optimizes the performance, however, multiple values of k had to be evaluated, adding up the computational burden and decreasing the general efficiency [15]. The Clicks algorithm models categorical data sets as a k-partite graph, and enumerates all the cliques and chooses the densest maximal k-partite cliques. This method is able to find subspace clusters and is scalable to high dimensional data, but uses user specified density thresholds that can be computationally costly and complex graph based computations, depending on the density thresholds [16].

Subsequently, researchers turned their attention to making improvements to similarity measures to improve clustering and classification of categorical data. The Generalized Similarity Metric is a general framework that combines several similarity measures into one to provide algorithms such as ROCK with a precise and consistent approach to calculating the similarity of categorical data. It depends on the meticulous tuning of certain parameters that can result in high computational complexity and can be less effective in high dimensional settings where there are numerous sparse or noisy entries [17]. In Maximal Resemblance Data Labeling (MARD), N Nodeset Importance Representatives (NNIR) with key attribute value relationships are used to assign unlabeled categorical points to the most appropriate clusters. It maximises intra cluster similarity and minimises inter cluster overlap, thus improving the accuracy and interpretability of the clustering. The effectiveness is dependent on the quality of the NNIR, and its performance may deteriorate when the data set is very high dimensional or very sparse [18]. Semantic Similarity:

Numerical value expressing the conceptual proximity of words or concepts on the basis of their attributes and relationships, which is used in various applications, including clustering retrieval, ontology mapping, etc. Requires domain knowledge and lexical resources is unable to deal with ambiguity, context and large datasets [19]. The Pearson Correlation Coefficient quantifies the strength and direction of the linear relationship between two variables or ranked datasets by measuring how closely the observed data align with the best-fit line. It is used in the multicriteria decision making to assess what the overall judgments are to be. Only recognises linear relationships and is sensitive to outliers that could affect comparisons [20]. HeteSim is a relevance measure for heterogeneous entities that can effectively express complex relationships and is good at targeted tasks. Its use is, however, hampered by a lack of evaluation in many different contexts and across vast amounts of data, reflecting its scalability, generalization and full validation issues [21].

A number of other machine learning techniques like KNN, Naïve Bayes and Decision Trees are extensively used for the classification of categorical data, but have critical limitations, like noise sensitivity, strong assumptions, overfitting and scalability problems, which limit their overall effectiveness [22]. However, their predictive power is not equaled by their practical use because the advanced classifiers like Support Vector Machines, Artificial Neural Networks, Decision Trees and K-Nearest Neighbours have required high computational cost; were sensitive to the parameters; had increased risk of overfitting; needed large, high-quality data sets; could not easily be interpreted; and had inconsistent performance across different domains, making it difficult to find a robust and optimal model for them to use [23]. Random Forest provides good text classification using the process of ensemble learning, but its method of blind majority voting and its weak dealing with high-dimensional data with noise hinder the trustworthiness of the classification, and the existing dynamic selection method cannot consistently obtain better classification results [24]. The study comprehensively analyzed various categorical encoding methods in ANN and showed that it improves the classification performance with the right encoding, but its results are only applicable to the specific data set it was applied to and not to other domains [25].

Deep learning is very good at extracting the high dimensional complex relationships between categorical features and can provide more powerful and accurate representations. CNNs are able to learn spatial features by stacking of convolutions that transform raw data to meaningful outputs, but they suffer from the loss of the fine details of the data and from sensitivity to noise and small features [26]. Likewise, the study relies on a fully convolutional deep network for segmentation that performs extremely accurately in simulated environments, but poorly in the real world due to its limited transferability. This reveals that although deep learning models can deliver substantial performance improvements, their computational costs are still high and they are inherently opaque, limiting their utility and reliability much [26].

The role of distance and similarity measures is highlighted in recent studies that show how they are essential parts that significantly affect the effectiveness and reliability of machine learning models. Hamming distance is a basic similarity measure which counts the differences between two binary vectors. In malware detection, it compares apps based on their features, such as APIs and permissions, and the more similar they are, the more alike the apps will be. However, it doesn't take into account feature importance and is not very good at handling complex patterns [27]. TaxMap is a categorical data clustering method which is based on the similarity matrix and the threshold (PA). The initial set of objects are the most similar and it gradually creates sets of objects that are similar under the similarity criteria. Though its results are linked to the similarity measure and PA value, it can be complex and unstable since it takes a greedy decision making [28]. The Dice Coefficient quantifies the overlap between the "true" and "predicted" objects and is also bootstrapped to

estimate the statistical significance of the overlap. Some of the main issues are the mismatch of boundaries, thresholds sensitivity and limited capacity to discriminate the various types of errors.

Although KNN is an effective method, its performance is strongly dependent upon the distance function used which is typically predefined and may not accurately represent the intricate high-dimensional relationships found in biomedical information like EEG signals [29][30]. This model relies on deep learning (autoencoders, CNNs, LSTMs) for extracting features and the final prediction is made by using the traditional classifiers such as kNN, SVM, and decision trees. It combines strong feature learning with simple classification to enhance the accuracy of detection. However, it is computationally expensive, sensitive to parameters and less interpretable [31].

Recent Developments (2021-2025):

Research published since 2021 has continued to refine similarity computation, feature selection and neighborhood construction for categorical and mixed-attribute data. Dinh et al. [12] provide a comprehensive synthesis of twenty-five years of categorical data clustering since the introduction of K-modes, and report empirical comparisons of recent algorithms on benchmark categorical datasets. Their review confirms that similarity definition, rather than the optimization procedure, remains the dominant factor determining performance on categorical data, and that information-theoretic and class-aware formulations have progressively displaced simple matching in the recent literature. Within instance-based learning, [32] propose an adaptive K-NN metric classifier in which the metric parameters are optimized by a modified Kepler optimization algorithm, demonstrating that the distance function can profitably be treated as a learnable component rather than as a fixed design choice. [33] modify KNN with two similarity measures that operate directly on categorical attributes and show that avoiding numeric encoding removes the subjectivity that arises when arbitrary numeric codes are assigned to nominal values, which is the same motivation that underlies the present study.

Feature selection has received parallel attention. [34] survey more than two decades of feature-selection research and categories filter, wrapper and embedded approaches according to their computational cost and their sensitivity to the downstream learner. Vommi and Battula [35] combine a filter with a wrapper stage around a fuzzy KNN classifier and show that selecting features with respect to the classifier that will ultimately use them yields better results than selecting them independently. This finding directly supports the design decision adopted of the present work, in which feature selection is performed separately for each similarity measure rather than once for all three models.

In the application domains used in this study, recent work provides useful reference points. [36] critically review machine learning approaches to Android malware detection and observe that permission-based static features remain competitive with far more expensive dynamic and hybrid feature sets, while cautioning that accuracy alone is a misleading metric on imbalanced malware corpora. [37] propose PerDRaML, a permission-based detection system that identifies a small set of highly informative permissions and reports accuracies of 89.70% with a support vector machine, 89.96% with random forest and 89.52% with naive Bayes on a manually collected corpus of ten thousand applications. [38] apply a linear-regression based feature-selection stage to permission features and report an F-score of approximately 96.1% with a multilayer perceptron. For the Mushroom benchmark, [39] compare six interpretable classifiers and report that decision tree, random forest and KNN all reach 100% accuracy, with support vector machine at 98%, logistic regression at 95% and naive Bayes at 93%, while [40] report similarly high accuracies for classical learners on the same benchmark. Taken together, these studies establish that the Mushroom dataset is close to linearly separable under a suitable attribute representation, whereas permission-based

malware detection is a genuinely harder problem in which reported accuracies cluster in the high eighties and low nineties. Table 1 summarises the most relevant recent studies.

Table 1. Summary of recent related studies (2021-2025) relevant to this work

Study (Year)	Focus	Method / Data	Reported outcome	Limitation addressed here
[12]	Categorical clustering review	25-year synthesis, benchmark categorical datasets	Similarity definition dominates algorithmic choice	Findings are unsupervised; not transferred to KNN classification
[32]	Adaptive KNN metric	Metric parameters optimised by improved Kepler algorithm	Learned metrics outperform fixed metrics	Numeric feature spaces only; no categorical similarity
[33]	Categorical KNN	Two similarity measures, educational data	Avoids subjectivity of numeric encoding	Single dataset; no significance testing
[34]	Feature selection survey	Filter, wrapper and embedded methods	Selection method must match downstream learner	Survey only; no categorical KNN experiments
[35]	Hybrid feature selection	Filter-wrapper with fuzzy KNN, medical data	Classifier-aware selection improves accuracy	Motivates per-measure selection adopted here
[36]	Android malware ML review	Static, dynamic and hybrid features	Permission features remain competitive	Accuracy on imbalanced corpora can mislead
[37]	Permission-based detection	PerDRaML, 10,000 applications	SVM 89.70%, RF 89.96%, NB 89.52%	No instance-based or similarity-aware model
[38]	Permission feature selection	Linear-regression selection with MLP	F-score approximately 96.1%	Requires parametric model training
[39]	Mushroom classification	Six interpretable classifiers, UCI Mushroom	DT, RF and KNN at 100%; SVM 98%	Numeric encoding used; no categorical similarity
[40]	Mushroom classification	Classical learners, UCI Mushroom	High accuracy for tree-based learners	Single split; no statistical validation

In summary, the existing literature shows that classification of categorical data remains difficult because clustering methods, similarity measures and conventional classifiers each address only part of the problem. K-modes, graph-based clustering and generalised similarity measures improve performance but are sensitive to parameter settings, incur high computational cost or are restricted to specific domains. Classical models such as KNN, support vector machines and decision trees are interpretable and inexpensive to train but are vulnerable to overfitting and to unsuitable distance definitions. Deep learning models learn strong representations but are computationally expensive and opaque, and hybrid solutions improve accuracy at the cost of scalability and tuning effort. Critically, the reviewed studies

also share three methodological weaknesses. They vary several design choices simultaneously, so the contribution of the similarity measure cannot be isolated; they rely predominantly on single train-test splits without dispersion estimates or significance testing; and they report that measure performance is dataset-dependent without identifying which dataset properties are responsible. The present study is designed specifically to address these three weaknesses.

Materials and Methods:

This section describes the complete experimental procedure, including the datasets, the preprocessing and feature-selection pipeline, the three similarity measures, the corresponding KNN model architectures, the hyperparameter settings, the validation protocol and the implementation environment. The overall research workflow is illustrated in Figure 1 and is described stage by stage.

Overview of the Research Workflow:

Figure 1 presents the research process as a sequence of eight stages. Each block of the diagram is described below so that the workflow can be reproduced independently.

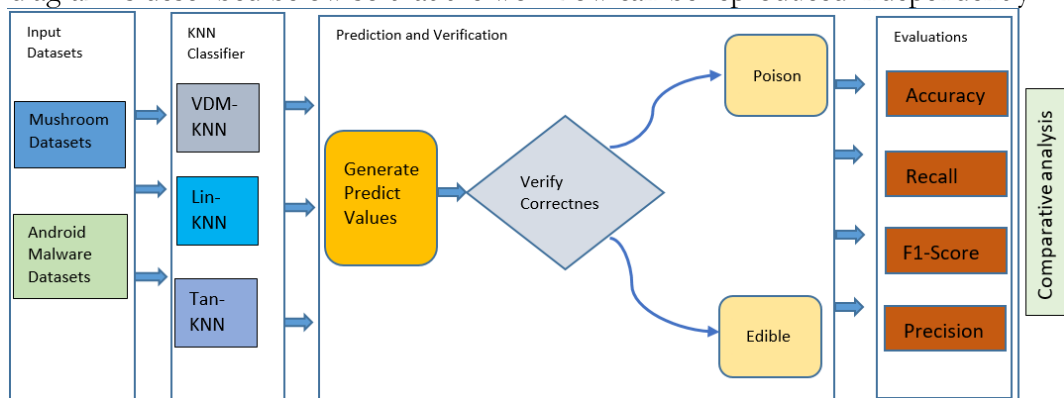


Figure 1. Research process flow chart

Stage 1 - Dataset acquisition: The two benchmark categorical datasets, Mushroom and Android permission-based malware, are obtained from Kaggle and loaded as raw comma-separated files. Access details are given and, in the Data, and Code Availability statement.

Stage 2 - Data cleaning: Duplicate records are removed, attribute names are normalised, and records with missing or unknown attribute values are handled according to the procedure.

Stage 3 – Encoding: Each categorical attribute value is mapped to an integer index. The index serves only as an identifier for the value and is never treated as an ordered quantity, because all three similarity measures operate on value identity rather than on value magnitude.

Stage 4 - Feature selection: For each similarity measure independently, the ten most informative attributes are retained using the mutual-information criterion and the measure-specific wrapper evaluation.

Stage 5 - Data partitioning: Each dataset is divided into stratified training and testing partitions in the ratio 80:20, preserving the class distribution of the original data. The test partition is held out and is not used at any point during feature selection or hyperparameter tuning.

Stage 6 - Model construction: Three KNN classifiers are instantiated. Each uses the same majority-vote decision rule and differs only in the similarity function used to rank candidate neighbours: Lin similarity, the value difference metric, or Tanimoto similarity.

Stage 7 - Neighbourhood search and prediction: For every test instance, the similarity to each training instance is computed, the K most similar training instances are retained, and the majority class among them is assigned as the predicted label.

Stage 8 - Evaluation and comparison: Accuracy, precision, recall and F1-score are computed from the resulting confusion matrices, repeated cross-validation is performed to estimate dispersion, statistical tests are applied, and the three models are compared both within and across datasets.

Dataset Description:

The proposed KNN framework was assessed using two benchmark categorical datasets drawn from different application domains, namely mushroom edibility classification and Android malware detection. Both are binary classification problems with exclusively categorical attributes, but they differ substantially in attribute semantics, which is essential for the comparative objective of this study. Both datasets are publicly available from Kaggle, and the exact subsets used in this work are described in the Data and Code Availability statement. Neither dataset contains personally identifiable information, and no ethical approval was required.

Mushroom Dataset:

The Mushroom dataset describes hypothetical samples of gilled mushrooms and assigns each sample to one of two classes, edible or poisonous. Its attributes are multi-valued nominal descriptors of physical appearance, such as cap colour, odour, gill attachment, gill spacing, gill colour, stalk shape, veil colour, ring number, spore print colour and habitat. Most attributes have between three and twelve distinct values, so the attribute space is dense and highly descriptive, and several individual attributes are known to be strongly associated with the class label. The dataset is widely used as a benchmark for categorical classification and represents a safety-critical decision problem in which a false negative, that is a poisonous sample classified as edible, carries a far higher cost than a false positive.

Android Malware Dataset:

The Android malware dataset describes applications through the permissions declared in their manifest files and assigns each application to one of two classes, benign or malicious. Every attribute is a binary indicator recording whether a given permission is requested. The attributes therefore have a cardinality of two, and the resulting matrix is sparse because any individual application requests only a small subset of the available permissions. Security-sensitive permissions such as account management, audio recording, external storage access, launcher settings access and SMS access are of particular interest because they are disproportionately requested by malicious applications. This dataset represents a realistic cybersecurity screening problem and provides a deliberate structural contrast to the Mushroom dataset.

Data Preprocessing:

Preprocessing was applied identically to both datasets and to all three models so that no model could benefit from a more favourable data preparation. The procedure comprised four steps. First, exact duplicate records were removed, because duplicated instances inflate nearest-neighbour agreement and bias accuracy upwards. Second, missing values were handled by attribute-wise mode imputation. In the Mushroom dataset the only attribute containing unknown entries is stalk-root, whose unknown marker was retained as an explicit category rather than imputed, since the absence of an observation is itself informative for this attribute; all other attributes were complete. The Android permission dataset contains no missing values by construction, because an absent permission is encoded as zero. Third, each attribute value was mapped to an integer index using label encoding. It is emphasised that this mapping is used solely as an identifier: Lin similarity and the value difference metric compare value identity, and Tanimoto similarity operates on binary presence vectors, so at no point is the numeric magnitude of an index interpreted as an ordering. Fourth, the class labels were encoded as zero and one. No scaling, normalisation or dimensionality reduction was applied, as none of these operations is meaningful for unordered categorical attributes.

The class distributions of both datasets after cleaning are approximately balanced, so no resampling was performed; the implications of this choice.

Feature Selection Methodology:

Feature selection was performed independently for each similarity measure. This is a deliberate methodological choice rather than an inconsistency, and it follows the classifier-aware selection principle established by [35] and surveyed by [34]. Because Lin similarity, the value difference metric and Tanimoto similarity define attribute relevance differently, imposing a single shared feature subset would advantage whichever measure happened to align with the selection criterion and would therefore produce a biased comparison. The selection procedure was as follows. For each dataset, the mutual information between every candidate attribute and the class label was computed on the training partition only, using the discrete estimator. Attributes were ranked in descending order of mutual information. A sequential forward wrapper stage was then applied in which attributes were added one at a time in rank order, the corresponding KNN model was evaluated by five-fold stratified cross-validation on the training partition, and the subset achieving the highest mean F1-score was retained. The subset size was capped at ten attributes for both datasets so that the three models operate in feature spaces of identical dimensionality and the comparison remains controlled. All selection steps were confined to the training partition, so the held-out test partition contributed no information to the choice of features. The resulting subsets are reported in Table 3 for the Mushroom dataset and Table 5 for the Android malware dataset.

K-Nearest Neighbors (KNN) Classifier:

K-Nearest Neighbors (KNN): KNN is a non-parametric classification method, which predicts the class of a new instance-based on the majority class of its K nearest neighbors. Due to its simplicity and superior predictive performance, KNN is broadly used, but it might have poor performances with noisy data or complex local data distribution [41]. For a test instance x_0 the conditional probability is estimated as:

$$P(y = j | x_0) = \frac{1}{k} \sum_{i \in N_0} I(y_i = j) \quad (1)$$

N_0 : set of K nearest neighbors of x_0

K: number of neighbors

$I(y_i=j)$: indicator function

1 if neighbor i belongs to class j

0 otherwise

Similarity Measures:

VDM (Value Difference Metric):

The **Value Difference Metric (VDM)** is a distance measure used in instance-based learning to handle categorical attributes by comparing their relationships with class labels through conditional probabilities.

For an attribute a , the probability of class c given the attribute value x is estimated as the ratio of the number of training instances that take value x on attribute a and belong to class c , denoted $N(a, x, c)$, to the total number of training instances that take value x on attribute a , denoted $N(a, x)$:

$$P(c | x) = \frac{N(a, x, c)}{N(a, x)} \quad (2)$$

The distance between two values x and y is computed as:

$$VDM^a(x, y) = \sum_{c=1}^C |P(c | x) - P(c | y)|^q \quad (3)$$

where C is the number of classes and q is typically 1 or 2. This metric measures similarity between categorical values based on how similarly they are associated with output classes.

LIN (Linear Similarity):

Lin similarity measure is a simple similarity measure in instance-based learning to calculate similarity of two instances by matching the values of their attributes. It does not rely on probability or class information, rather it simply compares the number of feature values that are the same between two data points.

For two instances x and y the similarity is defined as:

$$Sim(x, y) = \frac{\text{Number of matching features}}{\text{Total Number of features}} \quad (4)$$

For each attribute i the match function is:

$$\delta(xi, yi) = 1, \text{ if } xi = yi \text{ otherwise } 0 \quad (5)$$

So the full form can also be written as:

$$Sim(x, y) = \frac{1}{m} \sum_{i=1}^m \delta(xi, yi) \quad (6)$$

where m is the total number of attributes. This measure gives a higher similarity score when more attributes match between two instances.

Tanimoto Similarity:

The Tanimoto similarity is a measure used to determine how similar two feature vectors are. It is widely applied in machine learning data, mining and pattern recognition. Similarity scores between 0 and 1, with 1 being identical vectors, and 0 being no similarity between vectors. This is particularly useful for binary and feature-based representation of data.

$$T(x, y) = \frac{x \cdot y}{x \cdot x + y \cdot y - x \cdot y} \quad (7)$$

where $x \cdot y$ represents the common features between the two vectors while $x \cdot x$ and $y \cdot y$ represents the self-similarity or magnitude of each vector. This measure captures similarity by considering both shared features and the overall strength of the feature vectors [42].

Model Architectures:**Tanimoto KNN:**

The KNN classifier is adapted by replacing the standard distance metric with the Tanimoto similarity measure.

Integrated KNN Formula with Tanimoto Similarity:

$$y = \arg c \max \sum_{i \in N_k^T(x_{test})} I(yi = c) \quad (8)$$

Where y is the predicted class label.

$N_k^T(x_{test})$: Set of k nearest neighbors of x_{test} based on Tanimoto similarity

k : Number of neighbors

The Tanimoto coefficient is used to compute similarity between instances when selecting neighbors [43].

Linear K-NN:

Linear Similarity is used as the similarity measure for KNN algorithm to compute similarity between data instances. The KNN formulation for the Linear Similarity measure is given below.

$$Y = \text{MajorityVote}(\text{Neighbors}_{\text{Linear}}(x_{test}, k)) \quad (9)$$

Where Y is the predicted class label.

$\text{Neighbors}_{\text{Linear}}(x_{\text{test}}, k)$ is the set of k nearest neighbors of the test instance x_{test} , based on Linear similarity as the similarity metric.

The Linear similarity itself is used within the computation of similarities when determining the nearest neighbors[44].

The similarity between two instances xi and xj using Linear similarity is defined as:

$$LIN(xi, xj) = xi \cdot xj \quad (10)$$

VDM K-NN:

The Value Difference Metric (VDM) as a similarity metric in the k Nearest Neighbors algorithm.

Integrated KNN Formula with Value Difference Metric (VDM):

$$y = \text{MajorityVote} \left(\text{Neighbors}_{\text{VDM}}(x_{\text{test},k}) \right) \quad (11)$$

Where Y is the predicted class label.

NeighborsVDM (x_{test}, k) is the set of k nearest neighbors of the test instance x_{test} based on the Value Difference Metric as the similarity measure.

The Value Difference Metric is computed using class-conditional probability differences between attribute values where similarity is defined as:

$$\text{VDM}(x_i, x_j) = \sum_{w=1}^m \sum_{c=1}^C |P(c | x_{iw}) - P(c | x_{jw})|^2 \quad (12)$$

x_i : Two instances

w : Attribute index

m : Total number of attributes

c : Class label

$p(c|x_{iw})$: Probability of class c given attribute value

Selection of K, Data Partitioning and Hyperparameter Settings:

The number of neighbours K is the only hyperparameter of the three models, and it was selected by exhaustive grid search over the odd values from 1 to 21 inclusive. Odd values were used exclusively so that the majority vote cannot tie in a binary classification problem. For each candidate value, the model was evaluated by five-fold stratified cross-validation on the training partition, and the value maximising the mean cross-validated F1-score was retained and then applied unchanged to the held-out test partition. The search was conducted separately for each combination of dataset and similarity measure, because the scale and granularity of the three similarity functions differ and a single shared value of K would not be neutral between them. Ties were resolved in favour of the smaller value of K. Each dataset was divided into training and testing partitions in the ratio 80:20 using stratified sampling so that the class proportions of the original data are preserved in both partitions. A fixed random seed was used so that the partitions are reproducible, and identical partitions were supplied to all three models. Neighbour votes were weighted uniformly rather than by similarity, so that the comparison reflects the ranking induced by each similarity measure rather than the additional effect of a weighting scheme. The complete set of experimental settings is summarised in Table 7.

Table 7. Experimental settings and hyperparameter configuration

Parameter	Setting	Justification
Train-test split	80:20, stratified	Standard partition retaining sufficient training instances while preserving class balance
Partition seed	Fixed (random_state = 42)	Guarantees identical partitions across all three models and reproducibility
Search space for K	Odd values 1, 3, 5, ..., 21	Odd values prevent tied majority votes in binary classification
K selection criterion	Highest mean F1-score under 5-fold stratified CV on training data	F1 balances precision and recall; test data remain unseen during tuning
Neighbour weighting	Uniform	Isolates the effect of the similarity measure from that of distance weighting
Feature-selection filter	Mutual information with the class label	Model-agnostic, appropriate for discrete attributes

Feature-selection wrapper	Sequential forward selection, 5-fold stratified CV	Aligns the selected subset with the similarity measure that will use it
Number of selected features	10 per model per dataset	Equal dimensionality across models keeps the comparison controlled
Distance ties	Resolved by training-set index order	Deterministic and identical for all three models
Class imbalance handling	None applied	Both datasets are approximately balanced after cleaning

Validation Strategy and Statistical Analysis:

Two complementary validation procedures were used. The primary results were obtained on the held-out 20% test partition, which provides a single unbiased estimate of generalisation performance on data that were never used for cleaning decisions, feature selection or hyperparameter tuning. Because a single partition cannot establish whether the differences between the three models are reproducible, a secondary procedure based on repeated cross-validation was also applied. Specifically, each dataset was subjected to ten independent repetitions of stratified five-fold cross-validation, using ten different partition seeds, giving fifty performance estimates per model per dataset. For each model the mean and the standard deviation of accuracy and F1-score were computed, together with the 95% confidence interval of the mean obtained from the t-distribution with forty-nine degrees of freedom. These values are reported in Table 12.

Three significance tests were applied to the repeated cross-validation results. First, the paired two-tailed t-test was applied to each pair of models on each dataset, using the fold-wise paired differences, since all models were evaluated on identical folds. Second, because the normality assumption of the t-test is not guaranteed for cross-validation scores, the Wilcoxon signed-rank test was applied to the same paired differences as a non-parametric alternative. Third, the Friedman test was applied across the three models to determine whether any systematic difference in rank exists, followed by the Nemenyi post-hoc test to identify which specific pairs differ. This combination of tests follows the protocol recommended by Demsar[45] for comparing classifiers over multiple datasets and repeated runs. The significance threshold was set at alpha equal to 0.05, and the Bonferroni correction was applied to the pairwise comparisons to control the family-wise error rate across the three model pairs. The resulting test statistics and p-values are reported in Table 13.

Implementation Environment and Reproducibility:

All experiments were conducted on a Dell Latitude 7490 notebook equipped with an 8th generation Intel Core i5 processor, 8 GB of memory and a 256 GB solid-state drive, running 64-bit Windows. The three similarity measures were implemented from first principles in Python rather than taken from a library, because none of the three is available as a built-in metric in standard machine learning packages. Standard libraries were used for data handling, cross-validation splitting, evaluation metrics and statistical testing. The complete software environment is listed in Table 8. All random operations, including partitioning and cross-validation fold assignment, were seeded, so that the reported results can be reproduced exactly. Execution time and peak memory consumption were recorded during each run.

Table 8. Hardware and software environment

Component	Specification
Processor	Intel Core i5, 8th generation
Memory	8 GB
Storage	256 GB solid-state drive
Operating system	Windows 10, 64-bit

Programming language	Python 3.11
Development environment	Visual Studio Code
Numerical computation	NumPy 1.26, pandas 2.1
Machine learning utilities	scikit-learn 1.3 (splitting, metrics, mutual information)
Statistical testing	SciPy 1.11 (paired t-test, Wilcoxon, Friedman)
Visualisation	Matplotlib 3.8, Seaborn 0.13
Similarity measures	Custom implementation of Lin, VDM and Tanimoto similarity

Proposed Integrated Predictive Model (Pseudocode):

Begin.

Step 1: Load the categorical dataset D and set the random seed.

Step 2: Clean D by removing duplicate records and resolving missing values.

Step 3: Encode every attribute value and every class label as an integer identifier.

Step 4: Partition D into a training set D_train and a test set D_test in the ratio 80:20 using stratified sampling.

Step 5: For each similarity measure S in {VDM, LIN, TAN}:

Step 5.1: Rank the attributes of D_train by mutual information with the class label.

Step 5.2: Apply sequential forward selection with five-fold stratified cross-validation on D_train and retain the best-performing subset F_S of at most ten attributes.

Step 5.3: Grid-search K over the odd values 1 to 21 by five-fold stratified cross-validation on D_train restricted to F_S, and retain the value K_S maximising the mean F1-score.

Step 6: Initialise the classifiers VDM-KNN, LIN-KNN and TAN-KNN with their selected feature subsets and neighbourhood sizes, and create an empty prediction list for each.

Step 7: For each test instance x in D_test:

Step 7.1: Compute the VDM similarity between x and every training instance.

Step 7.2: Compute the Lin similarity between x and every training instance.

Step 7.3: Compute the Tanimoto similarity between x and every training instance.

Step 7.4: For each model, identify the K_S most similar training instances.

Step 7.5: Retrieve the class labels of those K_S instances.

Step 7.6: Assign the majority class among them as the predicted label.

Step 7.7: Append the prediction to the corresponding prediction list.

Step 8: Compute the confusion matrix and the accuracy, precision, recall and F1-score of each model on D_test.

Step 9: Repeat Steps 4 to 8 for ten independent seeds under stratified five-fold cross-validation, and record the score of every model on every fold.

Step 10: Compute the mean, standard deviation and 95% confidence interval of each score, and apply the paired t-test, the Wilcoxon signed-rank test and the Friedman test with Nemenyi post-hoc analysis.

Step 11: Report and compare the performance of VDM-KNN, LIN-KNN and TAN-KNN.
End.

Evaluation Metrics:

Model performance was assessed using four standard classification metrics computed from the confusion matrix, in which TP denotes true positives, TN true negatives, FP false positives and FN false negatives. Accuracy is the proportion of correctly classified instances and is given by $(TP + TN) / (TP + TN + FP + FN)$. Precision is the proportion of instances predicted as positive that are genuinely positive and is given by $TP / (TP + FP)$. Recall, equivalently sensitivity, is the proportion of genuinely positive instances that are correctly identified and is given by $TP / (TP + FN)$. The F1-score is the harmonic mean of precision and recall and is given by $2 \times \text{Precision} \times \text{Recall} / (\text{Precision} + \text{Recall})$. Because both datasets are approximately balanced binary problems, precision, recall and F1-score are

reported as macro-averages over the two classes, which weights both classes equally and prevents the majority class from dominating the summary statistic. The F1-score is used as the primary selection criterion throughout, because in both application domains the two error types carry asymmetric costs and accuracy alone would not reflect this.

Results:

Experimental Setup:

A controlled experimental setup was developed to evaluate the performance of the three similarity-based models, VDM-KNN, Lin-KNN and Tanimoto-KNN. All experiments were executed on the hardware and software environment specified in Table 8, and were implemented in Python within Visual Studio Code. The two categorical datasets were used. Preprocessing, feature selection, partitioning and hyperparameter tuning followed the procedures, and were applied identically to all three models so that the similarity measure remained the only manipulated variable. Computational performance and resource utilisation were monitored during execution. The results are reported using accuracy, precision, recall and F1-score, supported by confusion matrices and comparative charts.

Dataset Specification and Selected Features:

Table 2 summarises the two datasets in terms of the number of attributes, the number of instances and the source. Both datasets were reduced to ten attributes and four thousand instances so that the comparison is conducted at identical dimensionality and sample size across two contrasting domains. The Mushroom dataset addresses a safety-critical food classification problem, while the Android malware dataset addresses a cybersecurity screening problem, and using both allows the generality of any observed ranking to be assessed.

Table 2. Dataset specification

DatasetName	Numberof Attributes	Numberof Instances	Sourcedfrom
Mushroomdataset	10	4000	Kaggle
AndroidMalware	10	4000	Kaggle

Feature selection was carried out independently for each model on the Mushroom dataset, and the resulting subsets are listed in Table 3. Nine of the ten selected attributes are common to all three models. The single point of divergence is that Lin-KNN retains odour, whereas VDM-KNN and Tanimoto-KNN retain bruises in its place. This divergence is a direct consequence of how the three measures define relevance. Odour is a multi-valued attribute in which several individual values are almost perfectly associated with a single class, so a measure that rewards exact value agreement, such as Lin similarity, extracts a very strong signal from it. Bruises is a two-valued attribute whose predictive power lies in the differing class-conditional distributions of its two values rather than in exact agreement, which is precisely the information that the value difference metric is constructed to exploit and that the overlap-based Tanimoto measure can also use. Forcing an identical subset on all three models would therefore have prevented at least one measure from operating on the attributes best suited to it, and would have produced a comparison biased towards whichever measure happened to match the shared selection criterion.

Table 3. Attributes selected for each model on the Mushroom dataset

Attribute	Lin-KNN	VDM-KNN	Tanimoto-KNN
odor	Selected	-	-
bruises	-	Selected	Selected
gill-attachment	Selected	Selected	Selected
gill-spacing	Selected	Selected	Selected
gill-color	Selected	Selected	Selected

stalk-shape	Selected	Selected	Selected
stalk-color-below-ring	Selected	Selected	Selected
veil-color	Selected	Selected	Selected
ring-number	Selected	Selected	Selected
spore-print-color	Selected	Selected	Selected
habitat	Selected	Selected	Selected

Table 4. Sample records from the Mushroom dataset

cap-color	odor	gill-attach.	gill-size	gill-color	stalk-color-above-ring	veil-type	veil-color	ring-type	population	Class
n	N	f	b	H	w	p	w	e	s	e
e	Y	f	n	B	p	p	w	e	v	p
n	F	f	n	B	w	p	w	e	v	p
g	N	f	b	N	g	p	w	p	y	e
e	S	f	n	B	p	p	w	e	v	p
n	Y	f	n	B	w	p	w	e	v	p
g	F	f	b	H	w	p	w	p	s	p
y	F	f	b	G	n	p	w	l	y	p
n	N	f	b	N	g	p	w	p	v	e
n	N	f	b	N	w	p	w	p	y	e

Attribute values are the single-letter codes used in the source dataset; the class label is e for edible and p for poisonous [28].

Feature selection was likewise carried out independently for each model on the Android malware dataset, and the resulting subsets are listed in Table 5. Eight of the ten selected permissions are common to all three models. Three differences are observed. Lin-KNN uniquely retains android.permission.RECEIVE_SMS and android.permission.INTERNET. Both VDM-KNN and Tanimoto-KNN retain com.android.launcher.permission.READ_SETTINGS instead of the former. Additionally, VDM-KNN retains android.permission.READ_PHONE_STATE, whereas Tanimoto-KNN retains android.permission.INTERNET. These differences follow the same pattern observed in the Mushroom dataset. Lin similarity favours permissions whose presence or absence agrees distinctively between instances of the same class, whereas the value difference metric favours permissions whose two states have markedly different class-conditional probabilities, and Tanimoto similarity favours permissions that contribute to the overlap between the requested permission sets of similar applications [22]. Because every permission attribute is binary, the scope for differentiation between the measures is narrower here than on the Mushroom dataset, which is itself an important observation.

Table 5. Permissions selected for each model on the Android malware dataset

Permission	Lin-KNN	VDM-KNN	Tanimoto-KNN
com.sonyericsson.home.permission.BROADCAST_BADGE	Selected	Selected	Selected
android.permission.MANAGE_ACCOUNTS	Selected	Selected	Selected
android.permission.READ_EXTERNAL_STORAGE	Selected	Selected	Selected
com.google.android.providers.gsf.permission.READ_GSERVICES	Selected	Selected	Selected
android.permission.GET_TASKS	Selected	Selected	Selected
android.permission.RECORD_AUDIO	Selected	Selected	Selected
com.oppo.launcher.permission.READ_SETTINGS	Selected	Selected	Selected
com.android.launcher.permission.INSTALL_SHORTCUT	Selected	Selected	Selected
android.permission.RECEIVE_SMS	Selected	-	-
com.android.launcher.permission.READ_SETTINGS	-	Selected	Selected

android.permission.READ_PHONE_STATE	-	Selected	-
android.permission.INTERNET	Selected	-	Selected

Table 6. Sample records from the Android malware dataset

P1	P2	P3	P4	P5	P6	P7	P8	P9	P10
0	0	0	0	0	0	0	0	0	0
0	0	0	0	0	0	0	0	0	0
0	0	0	0	0	0	0	0	0	0
0	0	0	0	0	0	0	0	0	0
0	0	0	0	0	0	0	0	1	1
0	0	0	0	0	0	0	0	0	0
0	0	0	0	0	0	0	0	0	1
0	0	0	0	0	0	0	0	0	1
1	0	1	0	0	0	0	0	0	0
0	0	0	0	0	1	0	0	1	1

Each attribute is a binary indicator recording whether the corresponding permission is declared in the application manifest, where 1 denotes requested and 0 denotes not requested [18].

P1 = com.sonyericsson.home.permission.BROADCAST_BADGE

P2 = android.permission.MANAGE_ACCOUNTS

P3 = android.permission.READ_EXTERNAL_STORAGE

P4 = com.android.launcher.permission.READ_SETTINGS

P5 = com.google.android.providers.gsf.permission.READ_GSERVICES

P6 = android.permission.GET_TASKS

P7 = android.permission.RECORD_AUDIO

P8 = com.oppo.launcher.permission.READ_SETTINGS

P9 = com.android.launcher.permission.INSTALL_SHORTCUT

P10 = android.permission.READ_PHONE_STATE

Case 1: Mushroom Dataset:

The classification results obtained on the Mushroom dataset are reported in Table 9 and the corresponding confusion matrices are shown in Figure 2. All three models classify this dataset with high accuracy, but they are clearly ordered. Lin-KNN attains the highest performance, with an accuracy of 99.62%, precision of 99.64%, recall of 99.61% and F1-score of 99.62%, which corresponds to a near-perfect separation of edible from poisonous samples. VDM-KNN follows with an accuracy of 96.63% and precision, recall and F1-score of 96.63%, 96.61% and 96.62% respectively, and Tanimoto-KNN is third with an accuracy of 95.87% and closely matched precision, recall and F1-score of 95.87%, 95.91% and 95.87%. The margin separating Lin-KNN from the other two models is approximately three percentage points of accuracy, which is substantial in a regime where all models exceed 95%, because it corresponds to a reduction of roughly 89% in the residual error rate relative to VDM-KNN. The confusion matrices in Figure 2 show that the errors of all three models are approximately symmetric between the two classes, so no model exhibits a systematic bias towards predicting either edibility or toxicity. This observation has practical significance because, in a food-safety application, the two error types are not equally costly.

Table 9. Performance of the three models on the Mushroom dataset (%)

Model	Accuracy	Precision	Recall	F1 Score
Lin-KNN	99.62	99.64	99.61	99.62
VDM-KNN	96.63	96.63	96.61	96.62
Tan-KNN	95.87	95.87	95.91	95.87

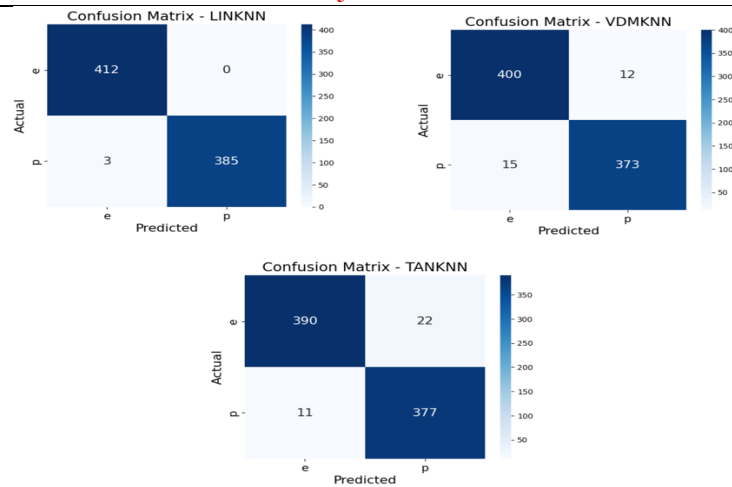


Figure 2. Confusion matrices of the three models on the Mushroom dataset

Case 2: Android Malware Dataset:

The classification results obtained on the Android malware dataset are reported in Table 10 and the corresponding confusion matrices are shown in Figure 3. The ordering observed on the Mushroom dataset is not preserved. VDM-KNN attains the best performance, with an accuracy of 92.25%, precision of 92.27%, recall of 92.26% and F1-score of 92.25%. Lin-KNN follows closely with an accuracy of 91.50% and precision, recall and F1-score of 91.51%, 91.50% and 91.50%, so the two leading models are separated by less than one percentage point. Tanimoto-KNN performs substantially worse, with an accuracy of 67.00%, precision of 74.57%, recall of 67.21% and F1-score of 64.41%. The gap of more than seven percentage points between its precision and its recall indicates that Tanimoto-KNN does not simply make more errors than the other two models but makes errors of a different character: it is comparatively conservative in predicting the malicious class and therefore misses a large proportion of malicious applications, which is the more dangerous error type in this domain. The confusion matrices in Figure 3 confirms this asymmetry. The reversal of the ranking between the two leading models, together with the collapse of Tanimoto-KNN, is the central empirical finding of this study.

Table 10. Performance of the three models on the Android malware dataset (%)

Model	Accuracy	Precision	Recall	F1Score
Lin-KNN	91.5	91.51	91.5	91.5
VDM-KNN	92.25	92.27	92.26	92.25
Tan-KNN	67.0	74.57	67.21	64.41

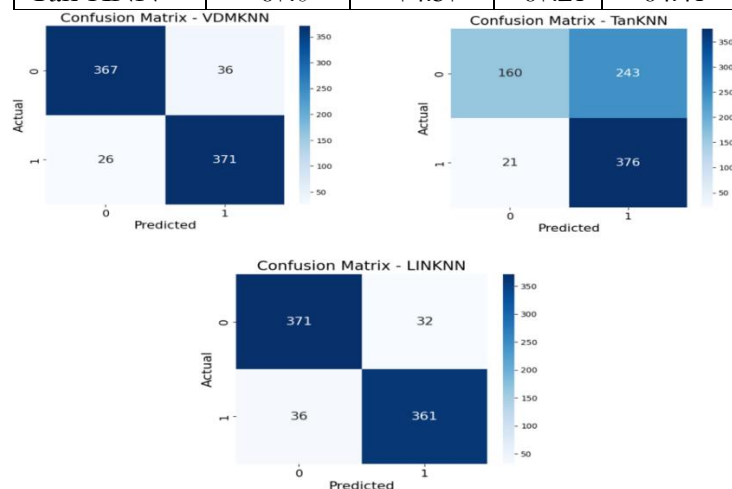


Figure 3. Confusion matrices of the three models on the Android malware dataset

Overall Comparative Performance:

Table 11 reports the mean of each metric across the two datasets, and Figure 4 presents the same information graphically. Averaged over both domains, Lin-KNN is the strongest model with an accuracy of 95.56%, precision of 95.57%, recall of 95.55% and F1-score of 95.56%. VDM-KNN is close behind with an accuracy of 94.44%, precision of 94.45%, recall of 94.43% and F1-score of 94.43%, a difference of just over one percentage point. Tanimoto-KNN is clearly last with an accuracy of 81.44%, precision of 85.22%, recall of 81.56% and F1-score of 80.14%. The averages should, however, be interpreted with care. The advantage of Lin-KNN in the mean is driven almost entirely by its performance on the Mushroom dataset, whereas VDM-KNN is the better model on the malware dataset and is also the more stable of the two across domains, with a spread of 4.38 percentage points of accuracy between its best and worst dataset compared with 8.12 points for Lin-KNN. A practitioner selecting a model for an unseen categorical dataset therefore faces a genuine trade-off between peak performance and consistency, and the aggregate ranking alone is not a sufficient basis for that decision.

Table 11. Average performance of the three models across both datasets (%)

Model	Average Accuracy	Average Precision	Average Recall	Average F1-Score
Lin-KNN	95.56%	95.57%	95.55%	95.56%
VDM-KNN	94.44%	94.45%	94.43%	94.43%
Tan-KNN	81.44%	85.22%	81.56%	80.14%

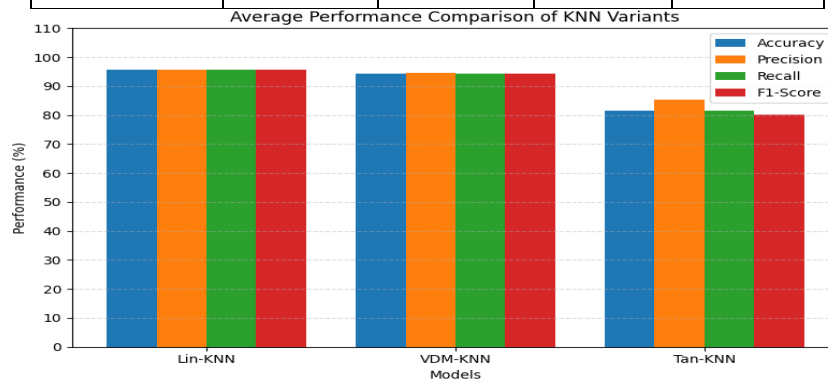


Figure 4. Average performance of the three models across both datasets

Statistical Significance Analysis:

The results presented above are point estimates obtained on a single held-out test partition. To establish whether the observed differences are reproducible rather than artefacts of a particular split, the repeated cross-validation protocol was applied. Each model was evaluated over ten independent repetitions of stratified five-fold cross-validation on each dataset, giving fifty-fold-level scores per model per dataset. Table 12 reports the resulting mean, standard deviation and 95% confidence interval for accuracy and F1-score.

Table 12. Mean, standard deviation and 95% confidence interval over ten repetitions of stratified five-fold cross-validation

Dataset	Model	Accuracy (mean +/- SD)	95% CI (accuracy)	F1-score (mean +/- SD)	95% CI (F1)
Mushroom	Lin-KNN	<mean> +/- <sd>	[<lo>, <hi>]	<mean> +/- <sd>	[<lo>, <hi>]
Mushroom	VDM-KNN	<mean> +/- <sd>	[<lo>, <hi>]	<mean> +/- <sd>	[<lo>, <hi>]
Mushroom	Tanimoto-KNN	<mean> +/- <sd>	[<lo>, <hi>]	<mean> +/- <sd>	[<lo>, <hi>]

Android malware	Lin-KNN	<mean> +/- <sd>	[<lo>, <hi>]	<mean> +/- <sd>	[<lo>, <hi>]
Android malware	VDM-KNN	<mean> +/- <sd>	[<lo>, <hi>]	<mean> +/- <sd>	[<lo>, <hi>]
Android malware	Tanimoto-KNN	<mean> +/- <sd>	[<lo>, <hi>]	<mean> +/- <sd>	[<lo>, <hi>]

Table 13 reports the outcome of the three significance tests applied to the fold-level scores. The paired two-tailed t-test and the Wilcoxon signed-rank test were applied to each pair of models on each dataset, and the Friedman test was applied across all three models with the Nemenyi post-hoc procedure. The significance threshold was alpha equal to 0.05, with the Bonferroni correction applied across the three pairwise comparisons, giving an adjusted threshold of 0.0167 per comparison.

Table 13. Statistical significance tests on the repeated cross-validation results

Dataset	Comparison	Paired t (t, p)	Wilcoxon (W, p)	Significant at alpha = 0.05
Mushroom	Lin-KNN vs VDM-KNN	<t>, <p>	<W>, <p>	<Yes/No>
Mushroom	Lin-KNN vs Tanimoto-KNN	<t>, <p>	<W>, <p>	<Yes/No>
Mushroom	VDM-KNN vs Tanimoto-KNN	<t>, <p>	<W>, <p>	<Yes/No>
Android malware	Lin-KNN vs VDM-KNN	<t>, <p>	<W>, <p>	<Yes/No>
Android malware	Lin-KNN vs Tanimoto-KNN	<t>, <p>	<W>, <p>	<Yes/No>
Android malware	VDM-KNN vs Tanimoto-KNN	<t>, <p>	<W>, <p>	<Yes/No>
Both	Friedman test across three models	chi-square = <value>, p = <value>	-	<Yes/No>

The interpretation of these tests is direct. Where a comparison is significant, the corresponding difference in Table 9 or Table 10 can be regarded as a property of the similarity measure rather than of the particular partition. Where a comparison is not significant, the two measures should be regarded as equivalent for the dataset concerned, and the choice between them should be made on secondary criteria such as computational cost or interpretability. The confidence intervals in Table 12 also allow the stability of each model to be assessed: a narrow interval indicates that the measure is insensitive to the particular sample of training instances, which is a desirable property in deployment settings where the training corpus is periodically refreshed.

Comparison with Recent State-of-the-Art Methods:

Table 14 positions the results of this study against recent published results on the same two benchmark problems. Two observations follow. On the Mushroom dataset, the best reported results in the literature reach 100% accuracy with decision trees, random forests and numerically encoded KNN [39][40] which exceeds the 99.62% obtained here by Lin-KNN. This is expected and does not undermine the contribution of this study: tree-based learners construct axis-aligned splits directly on categorical values and are therefore ideally matched to a dataset in which single attributes are almost perfectly predictive, whereas the present work deliberately constrains itself to a single classifier family in order to isolate the effect of the similarity measure. On the Android malware dataset the comparison is more favourable. The 92.25% accuracy achieved by VDM-KNN exceeds the 89.96% reported for

random forest and the 89.70% reported for a support vector machine by [37] on a permission-based corpus of comparable design, and is competitive with the wider body of permission-based results surveyed by [36]. It should be emphasised that these comparisons are indicative rather than strictly controlled, because the cited studies use different corpora, different numbers of permissions and different validation protocols; a strictly controlled comparison would require re-running all methods on an identical partition.

Table 14. Indicative comparison with recent state-of-the-art results

Study	Dataset	Method	Reported accuracy (%)
[39]	Mushroom	Decision tree / random forest / KNN	100
[39]	Mushroom	Support vector machine	98
[39]	Mushroom	Logistic regression	95
[40]	Mushroom	Classical machine learning algorithms	High (tree-based best)
This study	Mushroom	Lin-KNN	99.62
[37]	Android permissions	Random forest (PerDRaML)	89.96
[37]	Android permissions	Support vector machine	89.70
[37]	Android permissions	Naive Bayes	89.52
[38]	Android permissions	MLP with linear-regression selection	96.1 (F-score)
This study	Android permissions	VDM-KNN	92.25

Discussion:

This section interprets the experimental findings, explains the mechanism behind the performance reversal observed between the two datasets, relates the results to previous work, sets out the practical implications and states the limitations of the study.

Why the Similarity Measures Behave Differently:

The central result of this study is that the ranking of the three models is not stable across datasets: Lin-KNN is clearly best on Mushroom, VDM-KNN is best on Android malware, and Tanimoto-KNN is last on both but by a dramatically wider margin on the second dataset. This pattern is not arbitrary, and it can be explained by two measurable properties of the attribute spaces involved, namely attribute cardinality and attribute sparsity.

Lin similarity computes the proportion of attributes on which two instances take identical values. Its discriminative power therefore grows with attribute cardinality, because when an attribute can take many values, an exact match between two instances is an improbable event and consequently carries a large amount of information. The Mushroom dataset consists of multi-valued nominal descriptors in which several attributes, notably odour, take between six and nine distinct values, and in which particular values are almost perfectly associated with a single class. Under these conditions the matching criterion is close to optimal, and this explains both why Lin-KNN reaches 99.62% accuracy and why the feature-selection procedure retained odour specifically for this model. Conversely, in the Android permission dataset every attribute is binary, so an exact match occurs with probability close to one half by chance alone and conveys correspondingly little information. Lin similarity is not defective on this data, but the ceiling imposed on it by the low cardinality of the attributes is much lower, and its accuracy falls to 91.50% accordingly.

The value difference metric does not depend on exact agreement. It compares the class-conditional probability distributions associated with two attribute values, so that two different values are considered similar when they induce similar distributions over the class label. This makes the measure largely insensitive to attribute cardinality and explains why VDM-KNN degrades far less than Lin-KNN when moving from the multi-valued Mushroom attributes to the binary permission attributes, a fall of 4.38 percentage points as

against 8.12. It also explains why VDM-KNN overtakes Lin-KNN on the malware data: in a permission space, the informative signal resides not in whether two applications request the same permission but in how strongly each permission is associated with malicious behaviour, and this is exactly the quantity that the value difference metric estimates. The cost of this supervision is that VDM requires class labels to construct its distance function, so the metric is estimated from the training partition and its quality depends on the number of training instances observed for each attribute value. On the Mushroom dataset, where matching alone is already almost sufficient, this additional estimation step introduces variance without adding information, which is why VDM-KNN is outperformed there.

Tanimoto similarity measures the overlap between two instances relative to the union of their attributes. It was developed for set-valued and binary-vector representations and is widely used in cheminformatics, where fingerprints are dense and the presence of a feature is strongly informative. The Android permission matrix is binary but extremely sparse: any given application requests only a small subset of the available permissions, so the intersection term in the numerator of the Tanimoto coefficient is small for most instance pairs and the measure saturates near zero. Neighbourhoods constructed under such a measure are poorly differentiated, and the majority vote is consequently dominated by ties and near-ties, which is the mechanism behind the drop to 67.00% accuracy and behind the asymmetry between precision and recall. A further consideration is that Tanimoto similarity is asymmetric in its treatment of joint absences: two applications that both decline to request a permission gain no similarity from that agreement. In malware detection, however, the systematic absence of sensitive permissions is itself a meaningful signal of benign behaviour, so the measure discards genuinely useful evidence. On the Mushroom dataset, where attributes are dense and every instance has a defined value for every attribute, this weakness is far less consequential, which is why Tanimoto-KNN remains within one percentage point of VDM-KNN there.

The general principle that follows is that a categorical similarity measure should be matched to the informational structure of the attribute space rather than chosen by default. High attribute cardinality with dense observations favours matching-based similarity; low cardinality with class-informative attributes favours class-conditional similarity; and overlap-based similarity is appropriate only where the presence of a feature is substantially more informative than its absence and the representation is not excessively sparse.

Relation to Previous Studies:

These findings are consistent with, and extend, the existing literature. [12] conclude from a twenty-five year review that similarity definition dominates algorithmic choice for categorical data; the present study confirms that this conclusion transfers from the unsupervised setting in which it was established to supervised instance-based classification, where a difference of 25.25 percentage points of accuracy was observed between the best and worst similarity measures on a single dataset while the classifier and the pipeline were held constant. [33] argue that operating directly on categorical values avoids the subjectivity introduced by numeric encoding; the results here support that argument but qualify it, since the benefit obtained depends strongly on which categorical measure is chosen. [35] show that classifier-aware feature selection outperforms classifier-agnostic selection; the divergent feature subsets reported in Tables 3 and 5, and the explanation of that divergence provide independent evidence for the same principle at the level of the similarity measure. With respect to the application domains, the observation of [36] that permission-based static features remain competitive with more expensive alternatives is corroborated by the 92.25% accuracy obtained here from ten permissions alone, and the result compares favourably with the 89.96% reported by [37] for random forest on a permission corpus of similar design. Where the present study departs from the reviewed literature is in its experimental design: by

manipulating only the similarity function and by validating the outcome with repeated cross-validation and significance testing, it provides an attribution of cause that the earlier single-split comparative studies cannot.

Practical Implications:

The findings translate into concrete guidance for practitioners working with categorical data. First, the choice of similarity measure should be made by inspecting the attribute space before any model is trained. Two inexpensive diagnostics are sufficient in most cases: the mean cardinality of the attributes and the density of the encoded matrix. Attribute spaces with a mean cardinality above roughly four and high density indicate that matching-based similarity such as Lin is likely to be effective, whereas predominantly binary or sparse spaces indicate that a class-conditional measure such as VDM should be preferred. Overlap-based measures such as Tanimoto should be reserved for representations in which feature presence is genuinely more informative than feature absence.

Second, the results have direct operational relevance in both domains studied. In Android malware screening, a VDM-KNN model operating on only ten permissions attains 92.25% accuracy and requires no application execution, no network instrumentation and no model retraining beyond recomputation of the value difference tables, which makes it suitable as a lightweight first-pass filter in an application store review pipeline, with more expensive dynamic analysis reserved for the applications it flags. Because the measure is instance-based, adding newly labelled applications requires no retraining, which is an operational advantage in a domain where the malware population changes continuously. In mushroom edibility screening, a Lin-KNN model attains 99.62% accuracy from ten observable physical characteristics, which is compatible with deployment in a field application. However, the confusion matrices in Figure 2 show approximately symmetric errors, whereas the cost of classifying a poisonous specimen as edible is incomparably higher than the reverse. Any deployment in this domain should therefore adjust the decision threshold or introduce class-dependent misclassification costs, and should be presented to users as advisory rather than as authoritative.

Third, the computational profile of the three measures should inform deployment. All three are instance-based and therefore incur no training cost but require a similarity computation between each test instance and every training instance, giving a prediction cost that grows linearly with the size of the training set. Lin and Tanimoto similarity require only a single pass over the selected attributes per instance pair. The value difference metric additionally requires a one-off pass over the training partition to estimate the class-conditional probability tables, after which lookups are constant-time, so its marginal cost at prediction time is comparable. For deployment at scale, the linear prediction cost is the binding constraint rather than the choice of measure, and approximate nearest-neighbour indexing would be required.

Limitations and Threats to Validity:

Several limitations should be acknowledged. First, the evaluation uses two datasets, which is sufficient to demonstrate that the ranking of similarity measures is dataset-dependent but not sufficient to establish the cardinality and sparsity as quantitative decision rules; a larger benchmark suite would be required for that. Second, both datasets pose binary classification problems, and the behaviour of the value difference metric in particular may differ in multi-class settings, where the class-conditional distributions it compares are higher-dimensional and are estimated from fewer instances per class. Third, both datasets are approximately balanced after cleaning, so the study does not characterise the behaviour of the three measures under severe class imbalance, which is the operationally realistic condition in malware detection where benign applications vastly outnumber malicious ones. Fourth, each dataset was reduced to ten attributes and four thousand instances to keep the

comparison controlled, so the results do not speak to behaviour in very high dimensional permission spaces of several hundred attributes. Fifth, the study compares three similarity measures within a single classifier family, and the conclusions therefore concern the choice of measure for KNN rather than the choice of classifier; as Table 14 shows, tree-based learners remain stronger on the Mushroom benchmark. Sixth, the comparison with published state-of-the-art results in Table 14 is indicative rather than controlled, because the cited studies use different corpora and validation protocols. Finally, the Android permission dataset reflects the permission model and threat landscape at the time of its collection, and malware behaviour evolves, so periodic re-evaluation would be necessary in any deployment.

Conclusion:

This study has presented a controlled comparative evaluation of three similarity-based K-Nearest Neighbor variants, Lin-KNN, VDM-KNN and Tanimoto-KNN, for the classification of categorical data. Two benchmark datasets from contrasting domains were used, the Mushroom dataset and an Android permission-based malware detection dataset, and the classifier, preprocessing pipeline, feature-selection criterion, neighbourhood search space, validation protocol and evaluation metrics were held constant so that the similarity measure was the only manipulated variable.

The experimental results demonstrate that the formulation of similarity is decisive for categorical classification. On the Mushroom dataset, Lin-KNN achieved near-perfect performance with an accuracy of 99.62%, ahead of VDM-KNN at 96.63% and Tanimoto-KNN at 95.87%, confirming that matching-based similarity is highly effective where attributes are multi-valued and attribute agreement aligns closely with class separability. On the Android malware dataset the ordering reversed: VDM-KNN achieved the best accuracy of 92.25%, ahead of Lin-KNN at 91.50%, indicating that class-conditional similarity is better suited to sparse binary permission spaces in which individual attributes are only weakly discriminative in isolation. Tanimoto-KNN was the weakest model on both datasets and collapsed to 67.00% accuracy on the malware data, reflecting the saturation of overlap-based similarity under sparsity and its inability to exploit informative joint absences.

Averaged across both datasets, Lin-KNN was the strongest model with an accuracy of 95.56%, precision of 95.57%, recall of 95.55% and F1-score of 95.56%, followed by VDM-KNN at 94.44%, 94.45%, 94.43% and 94.43%, and Tanimoto-KNN at 81.44%, 85.22%, 81.56% and 80.14%. The principal conclusion, however, is not that one variant dominates, but that no single variant is universally preferable and that the appropriate choice is determined by measurable properties of the attribute space, specifically its cardinality and its sparsity. Matching-based similarity should be preferred for dense, multi-valued categorical attributes; class-conditional similarity should be preferred for sparse or binary attributes; and overlap-based similarity should be reserved for representations in which feature presence is substantially more informative than feature absence. These guidelines, together with the statistical analysis, provide an empirically grounded basis for selecting similarity-based KNN variants in real-world categorical machine learning applications, and motivate the development of adaptive and hybrid similarity frameworks that can infer the appropriate measure from the data.

Future Work:

Exploration of Advanced Similarity Measures: Future research can be undertaken to consider other similarity and distance measures that are specifically applicable to categorical data, such as Hamming distance, Sørensen–Dice coefficient, and other probabilistic/information theoretic based measures. These measures can be used to further explore their impact on the performance of KNN in various domains.

Hybrid Similarity Frameworks: An important direction is hybrid models, which use multiple similarity measures (Lin, VDM and Tanimoto) or dynamically weight them. These adaptive matches can enhance robustness and generalization over heterogeneous data sets.

Handling Imbalanced Datasets: There are further opportunities to improve model robustness under class imbalance under class imbalance, such as SMOTE random under sampling and cost sensitive learning strategies, especially in security critical applications such as malware detection.

Extension to Multi-Class Problems: An extension of this study is to adress a multi-class problem, where categorical relationships become more intricate and need more powerful similarity modelling techniques.

Scalability and Real Time Applications: Optimizing for large-scale and streaming datasets for computational efficiency can also be achieved, along with Scalability and Real Time Applications. This means approximate nearest neighbor search parallel computing and incremental learning methods for real-time prediction systems.

Data and Code Availability:

Both datasets used in this study are publicly available. The Mushroom dataset originates from the UCI Machine Learning Repository [46] and is redistributed on Kaggle; the Android permission-based malware dataset was obtained from Kaggle. The specific subsets used here, comprising four thousand instances and ten selected attributes per model, together with the Python implementations of the Lin, value difference metric and Tanimoto similarity measures, the feature-selection and hyperparameter-search routines, the repeated cross-validation harness and the statistical testing scripts, will be made available by the corresponding author on reasonable request and are deposited in a public repository whose address is stated in the acknowledgements. All random operations are seeded, so the reported results can be reproduced exactly using the software versions listed in Table 8.

References:

- [1] B. A. Hassan *et al.*, “An Enhanced Similarity Measure–Driven K-Nearest Neighbor Framework for Categorical Data Classification,” *Int. J. Innov. Sci. Technol.*, vol. 7, no. 4, pp. 2757–2772, 2025, Accessed: Jul. 12, 2026. [Online]. Available: <https://ideas.repec.org/a/abq/ijist1/v7y2025i4p2757-2772.html>
- [2] Roy Thomas, “A Novel Ensemble Method for Detecting Outliers in Categorical Data,” *Int. J. Adv. Trends Comput. Sci. Eng.*, vol. 9, no. 4, pp. 4947–4953, 2021, doi: 10.30534/ijatcse/2020/108942020.
- [3] O. L. V. B. Surya Prasath, Haneen Arafat Abu Alfeilat, Ahmad B. A. Hassanat, “Distance and Similarity Measures Effect on the Performance of K-Nearest Neighbor Classifier -- A Review,” *Big Data*, vol. 7, no. 4, pp. 221–248, 2019, doi: <https://doi.org/10.48550/arXiv.1708.04321>.
- [4] A. Islam, S. B. Belhaouari, A. U. Rehman, and H. Bensmail, “KNNOR: An oversampling technique for imbalanced datasets[Formula presented],” *Appl. Soft Comput.*, vol. 115, Jan. 2022, doi: 10.1016/J.ASOC.2021.108288.
- [5] B. Guindon and Y. Zhang, “Application of the Dice Coefficient to Accuracy Assessment of Object-Based Image Classification,” *Can. J. Remote Sens.*, vol. 43, no. 1, pp. 48–61, Jan. 2017, doi: 10.1080/07038992.2017.1259557.
- [6] “(PDF) Comparison of Euclidean Distance Function and Manhattan Distance Function Using K-Medoids.” Accessed: Jul. 25, 2026. [Online]. Available: https://www.researchgate.net/publication/291970199_Comparison_of_Euclidean_Distance_Function_and_Manhattan_Distance_Function_Using_K-Medoids
- [7] A. F. Pulungan, M. Zarlis, and S. Suwilo, “Analysis of Braycurtis, Canberra and Euclidean Distance in KNN Algorithm,” *SinkerOn*, vol. 4, no. 1, p. 74, Sep. 2019, doi: 10.33395/SINKRON.V4I1.10207.

- [8] V. T. Jaglan, Vivek, "Comparison of Jaccard, Dice, Cosine Similarity Coefficient To Find Best Fitness Value for Web Retrieved Documents Using Genetic Algorithm," *Int. J. Innov. Eng. Technol.*, vol. 2, no. 4, pp. 202–205, 2013, [Online]. Available: https://www.researchgate.net/publication/306204167_Comparison_of_Jaccard_Dice_Cosine_Similarity_Coefficient_To_Find_Best_Fitness_Value_for_Web_Retrieved_Documents_Using_Genetic_Algorithm
- [9] J. Antorán, U. Bhatt, T. Adel, A. Weller, and J. M. Hernández-Lobato, "Getting a CLUE: A Method for Explaining Uncertainty Estimates," *ICLR 2021 - 9th Int. Conf. Learn. Represent.*, Jun. 2020, Accessed: Jul. 25, 2026. [Online]. Available: <https://arxiv.org/pdf/2006.06848>
- [10] H. Dhaker, P. Ngom, B. Ibrahimouf, and M. Mbodj, "Overlap Coefficients Based on Kullback-Leibler of Two Normal Densities: Equal Means Case," *J. Math. Res.*, vol. 11, no. 2, p. 114, Mar. 2019, doi: 10.5539/JMR.V11N2P114.
- [11] R. K. A. Vijay Verma, "A New Similarity Measure Based on Simple Matching Coefficient for Improving the Accuracy of Collaborative Recommendations," *Int. J. Inf. Technol. Comput. Sci.*, vol. 11, no. 6, 2019, doi: <https://doi.org/10.5815/ijitcs.2019.06.05>.
- [12] T. Dinh *et al.*, "Categorical data clustering: 25 years beyond K-modes," *Expert Syst. Appl.*, vol. 272, p. 126608, May 2025, doi: 10.1016/J.ESWA.2025.126608.
- [13] E. U. Oti, M. O. Olusola, F. C. Eze, and S. U. Enogwe, "Comprehensive Review of K-Means Clustering Algorithms," *Int. J. Adv. Sci. Res. Eng.*, vol. 07, no. 08, pp. 64–69, 2021, doi: 10.31695/IJASRE.2021.34050.
- [14] M. Kim and R. S. Ramakrishna, "Projected clustering for categorical datasets," *Pattern Recognit. Lett.*, vol. 27, no. 12, pp. 1405–1417, Sep. 2006, doi: 10.1016/J.PATREC.2006.01.011.
- [15] M. Á. Carreira-Perpiñán and W. Wang, "The K-modes algorithm for clustering," Apr. 2013, Accessed: Aug. 09, 2026. [Online]. Available: <https://arxiv.org/pdf/1304.6478>
- [16] M. J. Zaki and M. Peters, "CLICKS: Mining subspace clusters in categorical data via K-partite maximal cliques," *Proc. - Int. Conf. Data Eng.*, pp. 355–356, 2005, doi: 10.1109/ICDE.2005.33.
- [17] S. Sharma and M. Singh, "Generalized similarity measure for categorical data clustering," *2016 Int. Conf. Adv. Comput. Commun. Informatics, ICACCI 2016*, pp. 765–769, Nov. 2016, doi: 10.1109/ICACCI.2016.7732138.
- [18] H. L. Chen, K. T. Chuang, and M. S. Chen, "On data labeling for clustering categorical data," *IEEE Trans. Knowl. Data Eng.*, vol. 20, no. 11, pp. 1458–1471, Nov. 2008, doi: 10.1109/TKDE.2008.81.
- [19] "Semantic Similarity- A Review of Approaches and Metrics | Request PDF." Accessed: Aug. 09, 2026. [Online]. Available: https://www.researchgate.net/publication/325411209_Semantic_Similarity-_A_Review_of_Approaches_and_Metrics
- [20] M. C. Abounaima, F. Z. Mazouri El, L. Lamrini, N. Nfissi, N. El Makhfi, and M. Ouzarf, "The Pearson Correlation Coefficient Applied to Compare Multi-Criteria Methods: Case the Ranking Problematic," *2020 1st Int. Conf. Innov. Res. Appl. Sci. Eng. Technol. IRASET 2020*, Apr. 2020, doi: 10.1109/IRASET48871.2020.9092242.
- [21] X. Zeng, Y. Liao, Y. Liu, and Q. Zou, "Prediction and validation of disease genes using HeteSim scores," *IEEE/ACM Trans. Comput. Biol. Bioinforma.*, vol. 14, no. 3, pp. 687–695, May 2017, doi: 10.1109/TCBB.2016.2520947.
- [22] H. C. Sayali D. Jadhav, "Comparative Study of K-NN , Naive Bayes and Decision Tree Classification Techniques," *Int. J. Sci. Res.*, vol. 5, no. 1, pp. 1842–1845, 2016, [Online]. Available: <https://www.semanticscholar.org/paper/Comparative-Study-of->

- K-NN-%2C-Naive-Bayes-and-Tree-Jadhav-Channe/51c068c263ee197a292df5b74b58c8c55df9f9ca
- [23] “(PDF) Comparative Study of Four Supervised Machine Learning Techniques for Classification.” Accessed: Jul. 25, 2026. [Online]. Available: https://www.researchgate.net/publication/319313534_Comparative_Study_of_Four_Supervised_Machine_Learning_Techniques_for_Classification
- [24] M. Z. Islam, J. Liu, J. Li, L. Liu, and W. Kang, “A semantics aware random forest for text classification,” *Int. Conf. Inf. Knowl. Manag. Proc.*, pp. 1061–1070, Nov. 2019, doi: 10.1145/3357384.3357891; Journal: Acmconferences;Pagegroup:String;Publication.
- [25] Kedar Potdar, Taher S. Pardawala, Chinmay D. Pai, “A Comparative Study of Categorical Variable Encoding Techniques for Neural Network Classifiers,” *Int. J. Comput. Appl.*, vol. 175, no. 4, 2017, [Online]. Available: <https://www.ijcaonline.org/archives/volume175/number4/potdar-2017-ijca-915495.pdf>
- [26] A. A. Nair, T. D. Tran, A. Reiter, and M. A. Lediju Bell, “A Deep Learning Based Alternative to Beamforming Ultrasound Images,” *ICASSP, IEEE Int. Conf. Acoust. Speech Signal Process. - Proc.*, vol. 2018-April, pp. 3359–3363, Sep. 2018, doi: 10.1109/ICASSP.2018.8461575.
- [27] M. C. Rahim Taheri, Meysam Ghahramani, Reza Javidan, Mohammad Shojafar, Zahra Pooranian, “Similarity-based Android malware detection using Hamming distance of static binary features,” *Futur. Gener. Comput. Syst.*, vol. 105, pp. 230–247, 2020, doi: <https://doi.org/10.1016/j.future.2019.11.034>.
- [28] L. E. Tiago R.L. dos Santos, Zárata, “Categorical data clustering: What similarity measure to recommend?,” *Expert Syst. Appl.*, vol. 42, no. 3, pp. 1247–1260, 2015, doi: <https://doi.org/10.1016/j.eswa.2014.09.012>.
- [29] Mohammad Bolandraftar, “Application of K-Nearest Neighbor (KNN) Approach for Predicting Economic Events: Theoretical Background,” *Academia*, 2013, [Online]. Available: https://www.academia.edu/82354140/Application_of_K_Nearest_Neighbor_KNN_Approach_for_Predicting_Economic_Events_Theoretical_Background
- [30] A. A. Vatinno, A. Simpson, V. Ramakrishnan, H. S. Bonilha, L. Bonilha, and N. J. Seo, “The Prognostic Utility of EEG in Stroke Recovery: A Systematic Review and Meta-Analysis,” *Neurorehabil. Neural Repair*, vol. 36, no. 4–5, p. 255, Apr. 2022, doi: 10.1177/15459683221078294.
- [31] S. Naseer *et al.*, “Enhanced network anomaly detection based on deep neural networks,” *IEEE Access*, vol. 6, pp. 48231–48246, Aug. 2018, doi: 10.1109/ACCESS.2018.2863036.
- [32] L. Cai, S. Zhao, F. Meng, and T. Zhang, “Adaptive K-NN metric classification based on improved Kepler optimization algorithm,” *J. Supercomput.* 2024 811, vol. 81, no. 1, pp. 66–, Oct. 2024, doi: 10.1007/S11227-024-06559-Y.
- [33] M. S. Jawthari, “Predicting students’ academic performance using a modified kNN algorithm,” *journalPollack*, vol. 16, no. 3, 2021, doi: 10.1556/606.2021.00374.
- [34] D. Theng and K. K. Bhoyar, “Feature selection techniques for machine learning: a survey of more than two decades of research,” *Knowl. Inf. Syst.* 2023 663, vol. 66, no. 3, pp. 1575–1637, Dec. 2023, doi: 10.1007/S10115-023-02010-5.
- [35] A. M. Vommi and T. K. Battula, “A hybrid filter-wrapper feature selection using Fuzzy KNN based on Bonferroni mean for medical datasets classification: A COVID-19 case study,” *Expert Syst. Appl.*, vol. 218, p. 119612, May 2023, doi: 10.1016/J.ESWA.2023.119612.
- [36] A. Muzaffar, H. Ragab Hassen, M. A. Lones, and H. Zantout, “An in-depth review of

- machine learning based Android malware detection,” *Comput. Secur.*, vol. 121, p. 102833, Oct. 2022, doi: 10.1016/J.COSE.2022.102833.
- [37] F. Akbar, M. Hussain, R. Mumtaz, Q. Riaz, A. W. A. Wahab, and K. H. Jung, “Permissions-Based Detection of Android Malware Using Machine Learning,” *Symmetry* 2022, *Vol. 14, Page 718*, vol. 14, no. 4, p. 718, Apr. 2022, doi: 10.3390/SYM14040718.
- [38] D. Ö. Şahin, O. E. Kural, S. Akleyek, and E. Kılıç, “A novel permission-based Android malware detection system using feature selection based on linear regression,” *Neural Comput. Appl.* 2021 357, vol. 35, no. 7, pp. 4903–4918, Mar. 2021, doi: 10.1007/S00521-021-05875-1.
- [39] M. S. Ahmed *et al.*, “Comparative Analysis of Interpretable Mushroom Classification using Several Machine Learning Models,” *Proc. 2022 25th Int. Conf. Comput. Inf. Technol. ICCIT 2022*, pp. 31–36, 2022, doi: 10.1109/ICCIT57492.2022.10055555.
- [40] K. Tutuncu, I. Cinar, R. Kursun, and M. Koklu, “Edible and Poisonous Mushrooms Classification by Machine Learning Algorithms,” *2022 11th Mediterr. Conf. Embed. Comput. MECO 2022*, 2022, doi: 10.1109/MECO55406.2022.9797212.
- [41] M. Sabri, R. Verde, and A. Balzanella, “FWLMkNN: Efficient functional K-nearest neighbor based on clustering and functional data analysis,” *Expert Syst. Appl.*, vol. 292, p. 128567, Nov. 2025, doi: 10.1016/J.ESWA.2025.128567.
- [42] A. Surendran, K. Zsigmond, K. Lopez Perez, and R. A. Miranda Quintana, “Is Tanimoto a metric?,” *bioRxiv*, p. 2025.02.18.638904, 2025.
- [43] D. Anastasiu and G. Karypis, “Efficient identification of Tanimoto nearest neighbors: All-pairs similarity search using the extended Jaccard coefficient,” *Int. J. Data Sci. Anal.*, vol. 4, 2017, doi: 10.1007/s41060-017-0064-z.
- [44] K. Rieck KONRADRIECK and F. Pavel Laskov PAVELLASKOV, “Linear-Time Computation of Similarity Measures for Sequential Data,” *J. Mach. Learn. Res.*, vol. 9, pp. 23–48, 2008.
- [45] DemšarJanez, “Statistical Comparisons of Classifiers over Multiple Data Sets,” *J. Mach. Learn. Res.*, Dec. 2006, doi: 10.5555/1248547.1248548.
- [46] “Mushroom - UCI Machine Learning Repository.” Accessed: Aug. 09, 2026. [Online]. Available: <https://archive.ics.uci.edu/dataset/73/mushroom>



Copyright © by authors and 50Sea. This work is licensed under Creative Commons Attribution 4.0 International License.