



University of Guilan

Computational Sciences and Engineering

journal homepage: <https://cse.guilan.ac.ir/>

Detection of breast cancer tumours using discriminative features extracted by the dual-objective coral reef algorithm and SVM classifier

Maedeh Kiani Sarkaleh^a, Hossein Azgomi^{a,*}, Azadeh Kiani-Srakaleh^b^a Department of Computer Engineering, Rasht Branch, Islamic Azad University, Rasht, Iran^b Department of Electrical Engineering, Rasht Branch, Islamic Azad University, Rasht, Iran

ARTICLE INFO

Article history:

Received 14 November 2025

Received in revised form 22 December 2025

Accepted 11 April 2026

Available online 11 April 2026

Keywords:

Breast cancer

Image processing

Dual-objective coral reef algorithm

SVM

Feature selection

ABSTRACT

Early detection and diagnosis of breast cancer masses can significantly reduce mortality associated with this disease. This paper introduces a system designed to identify cancerous masses in mammography images. The proposed system is composed of four phases: pre-processing, feature extraction, feature selection, and classification. In the pre-processing phase, the region of interest is isolated from the image, noise is eliminated using a median filter, and contrast is enhanced through histogram adjustment. In the feature extraction phase, the features pertinent to shape, histogram, and texture are extracted from the region of interest. In the third phase, the dual-objective coral reef algorithm is employed for feature selection. Finally, an SVM classifier is utilized in the classification phase. The proposed method was applied to 100 images from the MIAS database and 300 images from the CBIS-DDSM database. Based on the quantitative results, the system demonstrated promising performance in reducing classification errors. Its accuracy, sensitivity, specificity, AUC, MCC, and F1-Score values for the MIAS database were calculated as 95.4%, 99.18%, 91.79%, 0.91, 0.95, and 95.32%, respectively, and for the CBIS-DDSM dataset database were calculated as 96.73%, 98%, 95.46%, 0.934, 0.968, and 96.71%, respectively.

1. Introduction

Cancer is a leading cause of women's and men's mortality globally. Despite significant advances in early diagnosis and treatment of malignant tumors, the global burden of cancer continues to increase. Cancer cells arise from mutations or abnormal alterations in genes that are responsible for cell

* Corresponding author.

E-mail addresses: azgomi@iaurasht.ac.ir (H. Azgomi)

growth and health. These cells can cause tumor formation. Tumors may be classified as either benign or malignant. Benign tumors resemble normal cells visually, grow gradually, do not invade nearby tissues, and do not metastasize. While benign tumors are not cancerous, their presence may elevate the risk of developing cancer. Malignant tumor cells proliferate rapidly and differ in appearance from normal and benign cells. Malignant tumors metastasize by invading adjacent tissues and subsequently disseminating to other body parts [1]. Breast cancer involves the uncontrolled proliferation of breast cells. It is among the most prevalent cancers worldwide. The International Agency for Research on Cancer (IARC) lists breast cancer as the fifth leading cause of death among Iranian women [2]. There is, unfortunately, no effective measure to prevent breast cancer, so early and precise detection of cancerous breast tumors is crucial for appropriate treatment decisions and timely intervention [3].

There are various medical imaging systems for breast examination. Research indicates that MRI, ultrasound, and CT scans have high sensitivity in detecting small tumors, even in dense breasts. However, CT scans are unsuitable for routine screening since they increase cancer risk and the effectiveness of ultrasound imaging depends on the operator's expertise. On the other hand, MRI is expensive and its use is limited because of gadolinium and strong magnetic effects. Therefore, although digital mammography performs moderately in detecting tumors in dense breasts, it has been accepted as a routine screening method worldwide due to its low cost and minimal processing time [4].

Breast cancer prediction is necessary for early diagnosis, the development of treatment plans, and the adoption of preventive measures, which will ultimately contribute to improving patients and reducing mortality rates. The use of classification systems in medical diagnosis is gradually growing. Undoubtedly, the evaluation of patient data and specialist decisions are key factors in diagnosis. However, expert systems and various artificial intelligence techniques for classification significantly assist specialists. Classification systems help minimize potential errors that may arise due to specialists' lack of experience and provide detailed medical data for examination in a shorter time. Breast cancer patients need to be diagnosed and classified automatically, especially when rapid and accurate decision-making is required. Automatic diagnosis of these patients can potentially reduce mortality rates by decreasing the time radiologists spend on examinations [5]. Mammography serves as an effective diagnostic tool for detecting early-stage breast cancer. It is also an effective screening method that contributes to reducing mortality. Numerous innovative studies have been conducted on early detection of breast cancer in recent years. It can be concluded from research that the application of the techniques of image processing and pattern identification to mammographic images significantly aids in the early detection of breast tumors [6]. A fundamental step in diagnosing cancerous tumors is feature extraction. These features measure image characteristics. Since benign and malignant tumors in the breast differ in size and shape, extracting various features can provide us with information about the likelihood of the tumors being cancerous. The texture differences between benign and malignant tumors also cause variations in pixel intensity. Accordingly, the proposed system utilizes histogram, shape, and texture features [4]. Researchers have used various image analysis techniques, such as image classification, segmentation, annotation, and retrieval. Among multiple methods already proposed in computer-aided diagnosis (CAD) systems, diverse techniques of feature selection for obtaining distinctive features for breast tumor classification, alongside suitable classification algorithms, have provoked interest [7, 8]. However, most pattern recognition problems face the escalation of data volume,

leading to the issue of big dimensionality. A potent strategy to address this challenge is feature selection, which involves choosing a subset of initial features and discarding superfluous ones, thereby reducing data dimensions without compromising performance and even enhancing it in many cases [8]. Initially, adding more features boosts learning algorithm efficiency, but beyond a certain threshold, it not only stops improving it but also may degrade it in some cases. As the number of features increases, more data samples will be needed, raising time complexity and the cost of sample supply. Thus, when selecting features, the focus is put on simultaneously minimizing the number of features and enhancing classification accuracy. Feature selection methods eliminate inappropriate and less discriminative features, retaining those with relevant information that can distinguish between classes [9]. There are two approaches to feature selection algorithms. The first approach identifies a feature set independently of their classification performance that can preserve most of the original data's information. This feature selection technique is known as the filter method. The second approach selects a subset of features such that classification performance is maintained or, at least, not degraded. This approach, called the wrapper method, is generally more robust than filter methods but requires extensive computation [10]. Wrapper methods include exact optimization algorithms and heuristic and meta-heuristic algorithms. Exact optimization algorithms strive for fully optimal solutions but are unable to provide suitable solutions for optimization problems in a search space with high dimensions. Classic approximate optimization methods (heuristic algorithms) aim to find near-optimal solutions in a shorter time. However, they may become trapped at local optima, have early convergence, and are problem-specific. Meta-heuristic algorithms also target near-optimal solutions in a shorter time, but they have mechanisms to avoid local optima and can be applied to various issues. Thus, they are extensively used to solve optimization problems. Due to their flexibility, meta-heuristic algorithms are particularly well-suited to reduce computational costs by reducing the number of features and avoiding over-fitting for solving feature selection problems [11].

Classification refers to the process of assigning new cases to distinct classes based on the training cases. Many parametric and non-parametric classifiers are used in pattern recognition problems. Parametric classification is based on a mathematical model that defines the connection between inputs and outputs. These classification approaches are less adaptable but benefit from a faster and easier learning process. Non-parametric classifiers, in contrast, do not depend on a mathematical model and instead learn directly from the training data. The drawbacks of these classifiers include their slower performance, the demand for more training data, and overfitting. Various classifiers have been used in breast cancer diagnosis problems, among which the Support Vector Machine (SVM), K-Nearest Neighbors (KNN), and Artificial Neural Networks (ANN) have been employed extensively and demonstrated promising results [12,13].

This study presents a CAD-based diagnosis system designed to identify malignant masses in mammographic images. This system is capable of discerning whether suspicious masses depicted in mammography are malignant. It encompasses four phases: pre-processing, feature extraction, feature selection, and classification. In the pre-processing phase, after the target area is extracted from the image, noise is eliminated using a median filter, and contrast is enhanced through histogram equalization. In the next phase, the features of the shape, histogram, and texture are extracted from the target region. In the third phase, the dual-objective Coral Reef Optimization (CRO) algorithm is employed to identify distinguishing features. The CRO algorithm is a metaheuristic approach that emulates the growth patterns of corals in reefs. It operates through two

primary processes: reef formation and coral reproduction. In the classification phase, the SVM classifier is utilized due to lower sensitivity to the size of the training database. It is one of the most popular non-parametric techniques for solving binary classification problems. In the CRO algorithm, some initial features are designated as corals within the reefs. Then, new features are generated by regeneration operators through several iterations and positioned on the reefs. CRO's advantage over other metaheuristic feature selection methods lies in the fact that each feature category, such as corals located on the reef, is given an opportunity to prove its efficacy rather than being prematurely discarded once it provides a weak solution [14,15]. After several steps, the corals containing better feature categories grow, while weaker feature categories are eliminated from the reef. The quality of the extracted features is measured using a fitness function. Finally, the best features category is selected as the ultimate features category. The proposed method incorporates a dual-objective CRO for feature selection. The primary goal is to identify a set of features with the maximum value in the fitness function, as the selected features can provide higher accuracy in the classification phase. The secondary goal is to choose a minimal feature set as it contributes to further reducing data size. Given these two objectives concurrently, the features extracted from the dual-objective CRO (DOCRO) include distinguishing features that not only reduce data volume but also improve classification accuracy. Most of the SVM-based metaheuristic approaches introduced so far, such as CSA-SVM and WOAMLP, usually pursue only one objective in the feature selection process and their focus is on increasing the relevance of features to classes. In contrast, the DOCRO is able to select a more optimal and distinct set of features by simultaneously considering two key criteria, namely increasing relevance and reducing redundancy. This dual-objective mechanism improves the performance of the SVM classifier and increases the accuracy of mammographic lesion detection. Since the best features with the highest fitness are selected in DOCRO, not only does computational load and data volume decrease but classification accuracy also increases. By selecting optimal features through DOCRO and utilizing them in the SVM classifier, the hyperplane parameters are chosen to create the largest margin for the classes, ultimately increasing data classification accuracy and the generalizability to test samples.

In general, this paper's contributions can be outlined as follows:

- Detecting and diagnosing breast cancer tumors using the DOCRO algorithm and SVM classifier.
- Using histogram, shape, and texture features to discern malignant masses.
- Reducing data dimensions by selecting pertinent features based on dual-objective optimization.

The rest of the paper is organized in the following manner. Section 2 reviews the related literature. Section 3 delineates the foundational requirements for the research. Section 4 introduces the proposed method, which includes four phases: pre-processing, feature extraction, feature selection, and classification. Section 5 presents the empirical results of implementing the proposed method. Section 6 is dedicated to concluding remarks and finally, recommendations for future works are presented in Section 7.

2. Review of Related Work

In recent years, various methods have been developed for automated image analysis employing meta-heuristic and classification algorithms applied to medical datasets. Many of these computer-based diagnostic systems have proven highly effective and efficient in helping radiologists to diagnose breast cancer. Today, the application of image processing science in the analysis of medical images is garnering significant attention and has a wide range of uses [16,17]. This section reviews numerous machine-learning approaches that have been utilized for breast cancer diagnosis using various meta-heuristic algorithms based on the MIAS dataset since 2019. **Figure 1** displays the categorization of these approaches.

A breast cancer classification method is proposed in [18], utilizing a synthesized Convolutional Neural Network (CNN), an enhanced optimization algorithm, and transfer learning. The primary objective is to assist radiologists in swiftly identifying abnormalities. In this approach, the ant colony optimization (ACO) technique is modified with opposition-based learning. The Enhanced Ant Colony Optimization (EACO) method is then employed to determine the optimal hyperparameters for the CNN architecture. This technique integrates the CNN Residual Network 101 (ResNet101) architecture with the EACO algorithm, resulting in a novel model called EACO–ResNet101. The method is evaluated using the MIAS and DDSM (CBIS-DDSM) mammography datasets. The proposed method was reported to achieve 98.63% accuracy, 98.76% sensitivity, and 98.89% diagnosability on the CBIS-DDSM dataset, and 99.15% accuracy, 97.86% sensitivity, and 98.88% specificity on the MIAS dataset.

Authors in [19] developed a wrapper-based chaotic crow search algorithm (WCCSA) to enhance the diagnosis of breast cancer. The Probabilistic Neural Network (PNN) was employed to classify mammography images as malignant, benign, or normal. The combination of WCCSA and PNN achieved an accuracy of 97% on the MIAS dataset.

A meta-heuristic chaotic-based duck travel optimization (cDTO) algorithm was explored in [20] for classifying images from the MIAS database. Linear discriminant analysis was utilized for feature extraction from mammography images. The selected features were assessed using qualitative criteria, e.g., accuracy, sensitivity, specificity, and error rate. The results revealed that the cDTO classifier achieved a high accuracy of 98.5%.

Reference [21] utilized a block-based discrete wavelet packet transform (BDWPT) for feature extraction, including Shannon entropy, Tsallis entropy, Renyi entropy, and energy. Subsequently, the principal component analysis (PCA) was applied to extract distinguishing features from the primary feature vector. Then, a wrapper-based kernel extreme learning machine (KELM), enhanced by a weighted chaos salp swarm algorithm (WC-SSA), was proposed as a classifier for digital mammograms. The researchers validated the efficacy of the proposed BDWPT + PCA + WC-SSA-KELM scheme using three standard datasets of MIAS, DDSM, and BCDR. The proposed technique attained accuracies of 99.62% and 99.92% on the MIAS and DDSM datasets in the normal-abnormal category, respectively. For benign-malignant classification, accuracies of 99.28%, 99.63%, and 99.60% were achieved on the MIAS, DDSM, and BCDR datasets, respectively.

A novel classification approach was proposed in [22] to address the challenges of digital mammography analysis. Five distinct wavelet families were used for feature extraction from MIAS database mammograms. The extracted features had a statistical nature and served as inputs for the linear discriminant analysis (LDA) and cuckoo search algorithm (CSA) classifiers. The results clearly indicated that the CSA classifier outperformed the LDA classifier with an accuracy of 97.5%.

In [23], a novel training algorithm based on the whale optimization algorithm was proposed for the MLP network. The proposed method utilized MIAS and DDSM databases. The findings demonstrated that it achieved an accuracy of 98.1%. In [3], the initial step involved enhancing the quality of mammography images and the contrast of the abnormal areas by improving image contrast and reducing noise. Then, a color space-based method was used for image segmentation, followed by mathematical morphology. Features were then extracted using the gray-level co-occurrence matrix (GLCM) and discrete wavelet transform (DWT). Finally, an optimized convolutional neural network (CNN) and a new, improved meta-heuristic algorithm, known as the advanced thermal exchange optimizer, were employed for feature classification. The results indicated that the proposed method, when applied to the MIAS database, achieved an accuracy of 93.79%, with sensitivity and specificity rates of 96.89% and 67.7%, respectively.

A hybrid optimization algorithm combining the grasshopper optimization algorithm and the crow search algorithm was proposed in [24] for feature selection and breast mass classification using a multilayer perceptron. The algorithm's performance was compared with the algorithms based on the multilayer perceptron, improved and adaptive genetic algorithm, whale optimization algorithm, and butterfly optimization algorithm. The results revealed a classification accuracy of 97.1%, sensitivity of 98%, and specificity of 95.4% for the MIAS dataset.

In [25], features like texture features, pseudo-Zernike moments, and wavelet features were extracted from the input image. The simulated annealing algorithm was employed to reduce the feature vector size. The masses were finally classified using a meta-heuristic adaptive neuro-based fuzzy inference system in which the frog leap algorithm was employed. The method achieved an accuracy of 97.5% when applied to the Mini-MIAS database.

Authors in [26] explored the best meta-heuristic algorithm for fine-tuning the weights, biases, and hyper parameters of CNN networks, aimed at solving the challenge of characterizing breast image abnormalities. Moreover, hybrid models, including a CNN architecture and five representative meta-heuristic algorithms (GA, WOA, MVO, SBO, and LCBO), were introduced for the efficient detection of breast cancer anomalies. Two distinct tests were conducted to optimize the weights and biases. The first involved training the CNN model, while the second included optimizing the CNN model's training with meta-heuristic algorithms. A comparative analysis was performed to assess the influence of the meta-heuristic algorithms GA, WOA, MVO, SBO, and LCBO on the CNN architecture's performance. The results indicated that MVO, SBO, and LCBO outperformed GA and WOA. The classification accuracies achieved by CNN-GA, CNN-WOA, CNN-MVO, CNN-SBO, and CNN-LCBO in the fifth round were 0.76, 0.75, 0.84, 0.80, and 0.86, respectively. Concurrently, the conventional CNN model attained an accuracy of 0.66. The results suggest that optimization algorithms inspired by physics, biology, and humans adjust better performance parameters for CNNs than those based on evolutionary and swarm-based meta-heuristic algorithms.

In [27], a novel meta-heuristic technique named the separated enemy-driven dragonfly algorithm (SEDDA) was introduced. It is the improvisation of the dragonfly algorithm for selecting an optimal subset of features extracted from digital mammography. Texture features were extracted from the targeted region, and the optimal feature set was determined using the SEDDA algorithm. Utilizing a multilayer perceptron with backpropagation, this method achieved an accuracy of 94.5% and a sensitivity of 87.3% on the MIAS dataset using a multilayer perceptron with backpropagation.

In many optimization algorithms used at the feature selection stage, there is no balance between exploration and exploitation capabilities within a search space. In addition, it is time-consuming to determine the fitness function due to slow convergence rates. Furthermore, the classifier's accuracy diminishes with increasing the number of features, resulting in heightened computational complexity. Evolutionary algorithms can be used in various applications and to solve optimization problems. These algorithms generate a solution set randomly. The selected solutions are propagated using operators to produce offspring. Finally, the subsequent generation is formed by combining the parent and offspring populations. This iterative process continues as long as a predefined termination condition is met.

The Coral Reef Metaheuristic Algorithm, as a category of the nature-inspired algorithms, mimics the behavior and structure of coral reefs. This algorithm is used to solve optimization and search problems across various fields. Recently, this algorithm has been applied to the feature selection problem. In [28], the CRO algorithm was utilized in optimal feature selection for breast cancer diagnosis. An enhanced version of CRO combined with the adaptive b-hill climbing algorithm was employed in [29] to boost CRO's efficiency. In [15], a hybrid binary CRO algorithm combined with Simulated Annealing (SA) was used to identify the best subset of features in a biomedical dataset. There are various problems in the energy sector, which can be expressed as optimization problems. An example is wind speed prediction, to which CRO has been successfully applied [10]. This approach is based on a feature selection problem used to forecast wind speed using meteorological variables. In [30], the CRO algorithm was applied to optimize the performance of multi-reservoir systems for efficient surface water management. A typical NP-hard problem in communication engineering is the Channel Assignment Problem (CAP) for wireless access points (APs), which connect to a certain number of stations. The Coral Reefs Optimization with Substrate Layers (CRO-SL) algorithm in [31] significantly increased throughput while reducing the computational time required for optimal solutions. Another complex optimization problem is Optimal Power Flow (OPF), which involves determining the best operating parameters of a power-related system. In [32], the Coral Reefs Optimization with Substrate Layers algorithm enhanced with the CE method (CE+CRO-SL) demonstrated higher performance in terms of efficiency and accuracy than alternative techniques.

In recent years, the combination of metaheuristic algorithms with SVM classifiers has been proposed as an effective approach for mammographic lesion detection. For example, methods such as CSA-SVM and WOAMLP-SVM have selected features using single-objective algorithms and have been able to increase classification accuracy to some extent. However, the main limitation of these approaches is the focus on a single criterion, usually the association of features with classes, which leads to the selection of a larger and sometimes redundant set. In contrast, the DOCRO algorithm, with its dual-objective mechanism, simultaneously optimizes two key criteria, namely increasing relevance and reducing redundancy. This feature allows a more compact and distinct set

of features to be selected, and as a result, the performance of the SVM classifier in mammographic lesion detection is significantly improved.

This article uses the evolutionary CRO algorithm to solve the problem of feature selection in breast cancer tumor diagnosis with better outcomes than those before the feature selection. The SVM classifier, which is less sensitive to the number of training samples, is also employed in the classification phase.

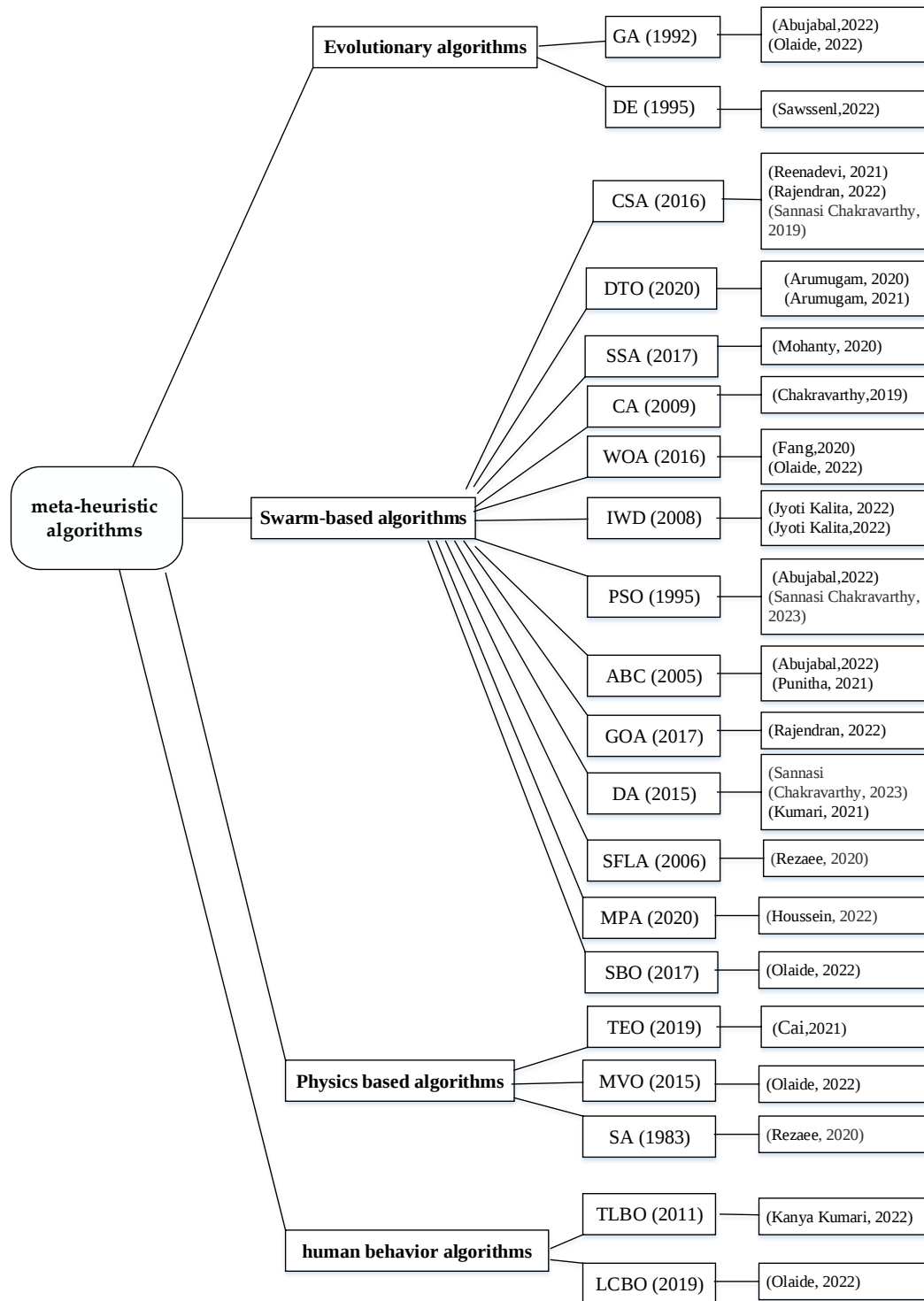


Figure 1. Categorization of machine-learning approaches using various meta-heuristic algorithms based on the MIAS dataset since 2019

3. Research Requirements

Numerous meta-heuristic algorithms have recently been employed to select distinctive features in medical datasets. This study utilized the CRO algorithm for feature selection, which is elaborated upon subsequently. The features selected by CRO were then applied to the SVM classifier to detect cancerous tumors.

3.1. Coral reef optimization (CRO) algorithm

The CRO algorithm [14] is a meta-heuristic algorithm that simulates coral growth behaviors in reefs. This algorithm is composed of two phases: reef formation and coral reproduction. During reef formation, initialization is performed by designating specific spaces within a network. **Figure 2** depicts the reef model, represented by a $N_1 \times N_2$ grid initiated by corals and coral colonies. Various coral larvae grow within the occupied grid squares, while the vacant squares represent potential spaces where new corals can freely settle and grow.

In the second phase, the CRO algorithm seeks optimal solutions via diverse coral reproduction operations. In each phase of CRO, a coral larva is produced, and each larva is labeled with a relevant fitness function $f(\text{REEF}_{ij})$, which is the same optimization objective function. Larvae with higher f values outlive others, and vice versa. It implies that there are superior solutions within the population at higher stages. The comprehensive operational description is provided below [14].

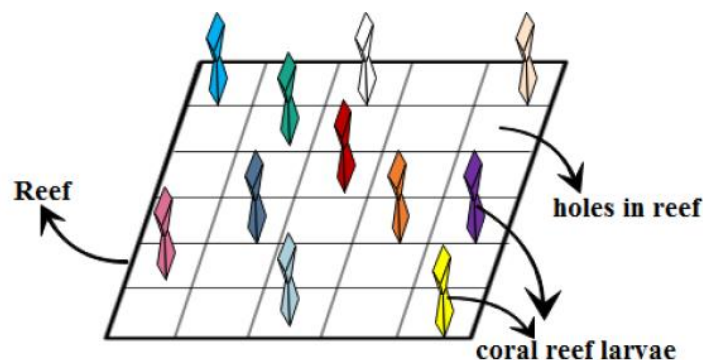


Figure 2. Coral reef model [15]

3.1.1. Broadcast spawning (external sexual reproduction)

Initially, some corals with the probability of F_b are selected for broadcast spawning, and then new larvae are created. Given the probability of p_k , the couples can be selected from the broadcast spawning coral pool at stage k . Once two corals are selected as the larva's parents, they are excluded from selection at stage k (a pair of corals can be selected as parents only in one stage). Couples can be selected either randomly or through a method such as roulette wheel selection. Larvae generated in this phase are not placed directly onto the grid but are instead released into the water.

3.1.2. Brooding (internal sexual reproduction)

In each part of the reef formation phase in the CRO algorithm, brooding is achieved by selecting corals with a probability of F_b-1 . In other words, the remaining larvae, which are not chosen in the prior step, will be utilized. In this stage, new solutions are generated by applying random mutations

to the chosen larvae. Subsequently, these larvae are dispersed into the water, akin to those produced during the broadcast spawning phase.

3.1.3. Larvae setting

The formation of coral reefs involves a struggle to occupy the space. Once all larvae are produced in stage k through broadcast spawning or brooding, they try to anchor on the coral and grow. Initially, the health (fitness function) of each coral is assessed. Then, each larva attempts to occupy a space (i,j) on the coral. If a cell is vacant, the reef will grow there, irrespective of its health. Conversely, if a cell is already occupied with a reef, a new larva will only replace it if it demonstrates superior fitness. All new larvae and corals vie for space, and the healthier ones (those with better fitness function) occupy the reef grid. There are only k opportunities for the larvae to secure a position; otherwise, they will be destroyed by other organisms on the substrate. Success in this phase is contingent upon the coral's robustness (its efficacy in optimizing the solution) or its luck in locating a vacant spot.

3.1.4. Budding (asexual reproduction)

The population is arranged in descending order based on their health functions, and a fraction of the best solutions reproduce themselves, while those with a low F_a for their self-reproduction are selected for asexual reproduction. These new larvae (representing the optimal solutions) seek to colonize other areas of the substrate following the settlement process outlined for larvae.

3.1.5. Depredation

At the end of each step of the algorithm, a depredation of corals is enacted to discard suboptimal solutions and ensure that there is ample vacant space on the reef for subsequent steps. To this end, a minor proportion of the worst solutions with a specific probability is purged. The depredation operator, with a minimal probability of p_d and a low F_d value, expels corals with poor health (lower fitness function) from the reef. For simplicity, we can assume $F_d = F_a$.

A critical difference between the CRO algorithm and the genetic algorithm is that there is no definition of solution neighborhood in the CRO algorithm. Rather, it employs the coral reproduction processes. Indeed, solutions are distributed within the problem space by corals to settle on the reef. Moreover, due to its structural resemblance to cellular evolutionary algorithms, the CRO algorithm is readily amenable to parallel execution [33].

This algorithm repeats until a termination criterion is met. For instance, reaching a number of generations specified by the user may serve as the algorithm's termination condition. **Figure 3** displays the pseudocode of the CRO algorithm.

Pseudo-code of CRO algorithm

Input: Ngen, N1,N2, Fb,Fa,Fd,R0,K,Pd,Opt**Output:** The best coral

```

1: Initialize  $REEF, REEF_{ij}$ , and parameter values ( Fb,Fa,Fd,R0,K,Pd)
2: Calculate fitness value of each coral using Eq. (1)
3: while Stopping criterion is not met do
4: Perform sexual crossover operator using broadcast spawning
5: for Couples of broadcast spawning corals  $P1$  and  $P2$  do
6:  $P1 + P2 \rightarrow L1 + L2$ 
7: end for
8: Perform sexual crossover using brooding
9: for Each brooding coral  $P$  do
10:  $P \rightarrow L$ 
11: end for
12: Settle new larvae (as per the fitness values)
13: Perform asexual reproduction using budding
14: for Each coral which would perform budding do
15:  $Pb \rightarrow Pb$ 
16: end for
17: Perform coral depredation
18: Compute the fitness value of each coral
19: Obtain the current optimal solution with best fitness
20: end while

```

Figure 3. Pseudocode of the CRO algorithm [29]**3.1.6. Computational Complexity**

The computational complexity analysis of any metaheuristic algorithm primarily depends on the search time. The initialization of parameters and the fitness function also influence the time complexity. For the CRO algorithm, the initialization process is estimated to last for $O(n)$ time, where n represents the number of corals, corresponding to the reef size $N_1 \times N_2$ in this study. The fitness function's time complexity is dependent on the classifier used. This study employed the K-Nearest Neighbors (KNN) classifier. For the KNN classifier, if n is the number of data points and d is the number of features per data point, the algorithm's time complexity is $O(n \times d)$. Thus, it can be said that the fitness evaluation generally requires $O(\text{fit})$ time. The main operations involved in the search process include broadcast spawning, brooding, larvae settling, and budding. Broadcast spawning needs $O(d^2)$ time, and brooding requires $O(d)$, where d represents the problem dimension or the number of features, which is 104 features in this study. The larvae settling stage requires $O(n^2)$ time, while budding has a time complexity of $O(n)$. If *maxIter* is considered the maximum number of iterations (generations) for the entire search process, which is 500 in this study, the total time complexity of CRO can be calculated using $O(n + \text{fit} + \text{maxIter} \times (d^2 + d + n^2))$ [29]. If an SVM classifier is used in the fitness function, the time complexity of the algorithm will depend on this classifier. The time complexity of training an SVM model depends on the number of training examples (n) and the number of features (d). In the worst case, the time complexity for training an SVM model is approximately $O(n^2 \times d)$, primarily arising from the need to solve problems with many constraints and dense internal matrices. The time complexity of a classification with a trained SVM model is approximately $O(d)$, indicating that each new example is computed based on its features [34].

3.2. Support vector machine (SVM)

Classification involves assigning new samples to various categories using the training samples. There are various parametric and non-parametric classification methods in pattern recognition problems. SVM is a non-parametric classifier that has achieved notable success in pattern recognition problems in recent years. A linear SVM aims to identify a distinguishing hyperplane that categorizes samples into two classes while maximizing the margin between them. In this classifier, the boundary training data, known as support vectors, are pinpointed using an optimization algorithm and are then used to determine the decision boundary. **Figure 4** illustrates two linearly separable classes and their respective support vectors. As is seen, only the boundary samples are necessary to establish the decision boundary, and selecting these samples so as to maximize the margin between categories minimizes classification error. Hence, this classifier is less sensitive to the number of training samples and is most effective when classifying linearly separable data [35].

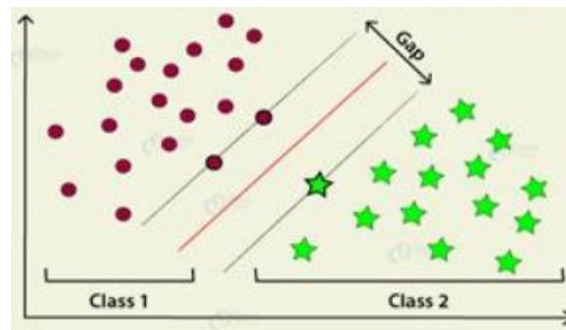


Figure 4. Support Vector Machine [34]

Non-linear SVM is utilized for datasets that are not linearly separable. In a non-linear SVM, samples are mapped into a space with more dimensions using a non-linear kernel function to achieve satisfactory separation. Despite its high computational cost and time-intensiveness, this method has demonstrated considerable effectiveness [36]. **Figure 5** presents the SVM pseudocode.

Pseudo Code of Support Vector Machine
Input: data set Output: C, γ
Normalize the data set 2. For Each C, γ : 2.1. Cross Validate using leave-one-out. 2.1.1. Train and test the SVM. 2.1.2. Store the success rate. 2.2. Compute the average success rate. 2.3. Update the best C and γ if needed. 2.4. Return to 2.1 with next C, γ . 3. Choose C, γ with best average success rate, and perform step (2) using fine scale around the selected parameters.

Figure 5. Pseudocode of Support Vector Machine [37]

3.3. Dataset

The data employed were the images in the MIAS and CBIS-DDSM datasets [38,39]. The MIAS database was created by an organization of research groups in the UK that are interested in

understanding mammography images. The images in this database possessed a digital resolution of 8 bits, and each image had a size of 1024 x 1024 pixels. It encompassed MLO (Mediolateral Oblique)-view images of both left and right breasts from 161 patients. In total, the database contained 322 images, out of which 208 were categorized as normal, 63 as benign, and 51 as malignant. **Figure 6** displays an exemplar from this image collection.

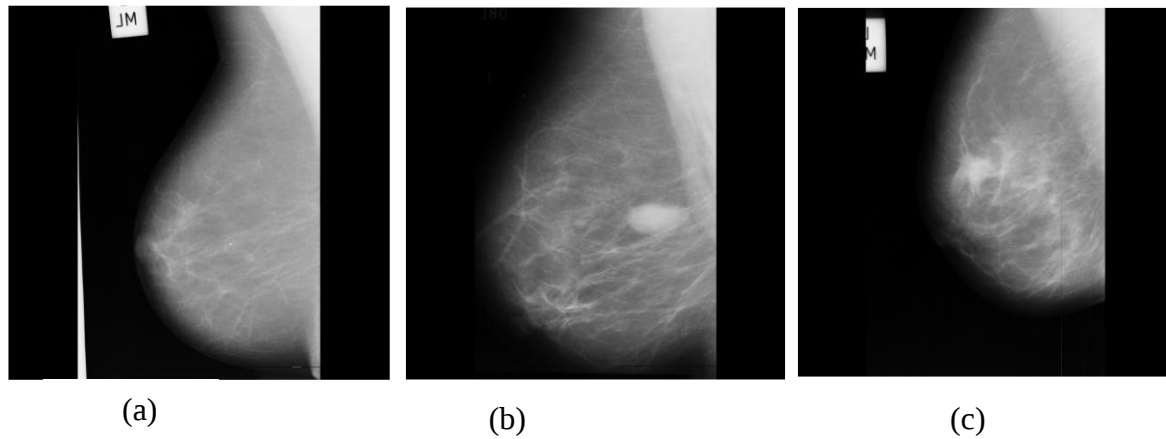


Figure 6. An example of the images in the MIAS database, (a) normal, (b) benign mass in the breast, (c) malignant mass in the breast

A text file contained information on the images in this database, including information on, e.g., the mass's location and radius, the breast's position (left or right), the type of breast tissue (fatty, mixed, or dense), and the nature of the mass (malignant or benign), when available. Based on the information in the dataset's Readme file, the mass center and radius in mammography images were identified, and the ROI was extracted using a cutting method. We used 100 images from the MIAS database, comprising 50 with benign masses and 50 with malignant masses, to implement and assess the proposed method. Out of these, 40 samples from each category were used to train the system, and 10 from each were used to test it. In total, 80 samples were used for training and 20 samples for testing.

The CBIS-DDSM dataset is a standardized and updated version of the DDSM dataset, specifically designed for research in mammography image recognition and analysis. The images in this dataset are available in DICOM format and can be viewed and analyzed using Matlab software. In addition to the overall breast images, ROI images for masses are available in a cropped form. These images have been stored at a 16-bit depth, which can enhance the ability to distinguish between tissue features and lesions. A .csv file is included in the dataset to provide comprehensive information about the mass type and the imaged view. The CBIS-DDSM dataset contains 3,161 mammographic images from patients, including 1,598 with cancerous masses, 1,211 with cancerous microcalcifications, and 1,562 with benign masses and microcalcifications, all captured from two standard views: CC (Cranio-Caudal) and MLO (Mediolateral Oblique) [38]. In the MLO view, the image is recorded obliquely to cover more breast tissue, including areas near the chest wall and axilla. This view provides crucial information about deep breast structures and aids in the detection of lesions that may not be visible in other views. Images from the MIAS dataset, also recorded in the MLO view, were used in this study. For the CBIS-DDSM dataset, images recorded in the MLO view were utilized for both training and testing the system. **Figure 7** illustrates an example of ROI images in the MLO view for both benign and malignant masses.

We used 300 ROI images from the CBIS-DDSM database, comprising 150 with benign masses and 150 with malignant masses, to implement and assess the proposed method. Out of these, 120 samples from each category were used to train the system, and 30 from each were used to test it. In total, 240 samples were used for training and 60 samples for testing.

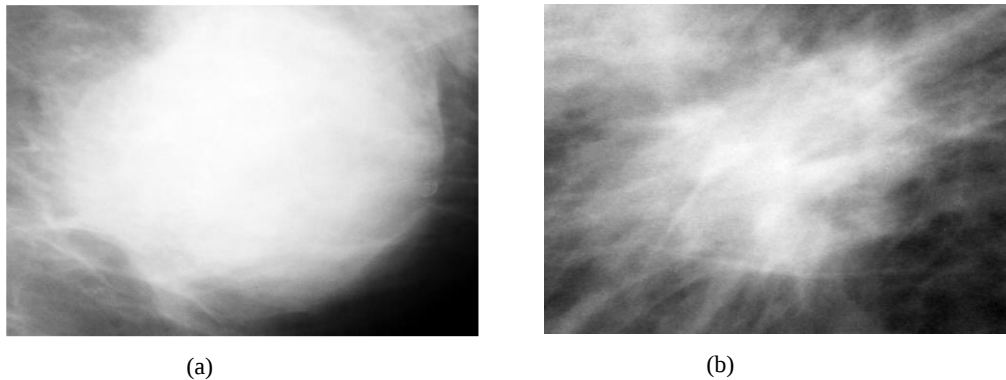


Figure 7. An example of the images in the CBIS-DDSM database, (a) ROI of benign mass, (b) ROI of malignant mass

For both datasets, the training and test samples were selected by the cross-validation method in which when there are n samples, t samples are used to train the classifier, and the remaining $n - t$ samples are used to test it. The classifier was constructed using the training samples, and its performance was evaluated with the test samples. This process was iterated for various combinations of training samples drawn from the entire dataset. The average outcomes for the evaluation parameters of the classifier were then reported. Owing to the variation in training and test samples across iterations and the inclusion of all samples as both training and test samples, the cross-validation method was deemed appropriate for assessing the proposed system. Additionally, the system was evaluated using all the samples.

3.4. Parameters Configuration

In the pre-processing phase, the image first underwent a median filter followed by histogram adjustment. Preserving the fine details in mammography images, including tumor margins and microclassifications, is essential for accurate analysis. The median filter effectively reduces noise while maintaining the original image structure by replacing pixel values with the median of their neighboring values without blurring the edges or altering major peaks in the image histogram. The filter mask size was selected as 3×3 by experimentation to achieve optimal results.

Histogram equalization enhances image contrast and improves the clarity of low-contrast areas by expanding the histogram and evenly distributing brightness across the entire range. This technique highlights brightness differences between different textures and clarifies otherwise indistinguishable details. Subsequently, 104 features related to texture, shape, and histogram were extracted from the image.

The DOCRO method was employed to select pertinent features. The DOCRO Algorithm is an improved version of the CRO Algorithm. The CRO algorithm initiates feature selection by designating some features as corals, regenerating and optimizing them iteratively. Stronger features thrive, while weaker ones are removed, measured by a fitness function. DOCRO extends this with

a dual-objective approach: maximizing classification accuracy and minimizing feature set size. This ensures efficient data reduction and improved performance. The parameters for the method, which are listed in **Table 1**, were obtained through experimentation to achieve the best outcomes.

Table 1. CRO parameters values

The parameter	Value
Ngen	500
N1	5
N2	5
Fb	0.9
Fa	0.01
Fd	Fa
R0	0.6
k	3
Pd	0.1
Opt	max

In a metaheuristic algorithm, the population size and the maximum number of generations, which represent the number of iterations applied, are two crucial parameters that significantly impact the algorithm's performance and efficiency. A larger population size can provide more diversity. In this study, a population size of 5×5 was considered. *Ngen*, or the number of generations, specifies the number of iterations the algorithm must run to reach an optimal solution, set to 500 in this study. The values of *Fb* and *Fa* are related to coral reproduction. To ensure efficient exploration of the search space, a high value of corals that use sexual reproduction is required. Conversely, a small amount of asexual reproduction is recommended. Thus, $Fb = 0.9$ and $Fa = 0.01$ were considered appropriate. The depredation process is performed to remove inappropriate solutions with the parameter $Fd = Fa$ and a probability $Pd = 0.1$. The KNN classifier was used for the fitness function. By searching from the set $\{1,2,3,4\}$, a value of $K = 2$ was chosen to achieve the best result.

A non-linear SVM with a radial basis function was applied for mass classification. The optimal value for the width parameter was determined by experimenting from the set $\{0.0001, 0.0005, 0.001, 0.005, 0.01, 0.05, 0.1, 0.5, 1, 5, 10, 50, 100\}$ to attain the best classification performance. If the training data were placed on the boundary of the separating hyperplane, the margin would decrease. To counteract this issue, a smooth margin was used, permitting a certain number of misclassifications during training to enhance the margin, thereby improving generalizability and performance on test data. This adjustment involved incorporating a penalty parameter into the classifier equations. In the proposed system, the penalty parameter was determined by exploring values from the set $\{0.001, 0.1, 1, 10, 100, 1000, 10000\}$ to achieve the most favorable outcome.

3.5. System configuration and evaluation parameters

All evaluations were conducted with the MATLAB software suite in a PC with Intel Core i5-1035G1 CPU 1.00GHz 1.19 GHz, 8 GB DDR3 main memory, and 1TB SATA hard drive.

The proposed system's efficacy in diagnosing breast cancer tumors was assessed by parameters like accuracy, sensitivity, specificity, area under curve (AUC), Matthews correlation coefficient (MCC),

and F1 Score derived from the confusion matrix [40, 41]. The confusion matrix was obtained using the sample labels calculated by the SVM classifier. **Table 2** presents an example of this matrix.

Table 2. The confusion matrix

Category	The label obtained from the classifier	
	malignant	Benign
malignant	TP	FN
Benign	FP	TN

$$Accuracy = \frac{TP+TN}{TP+FP+TN+FN} \times 100 \quad (1)$$

$$Sensitivity = \frac{TP}{TP+FN} \times 100 \quad (2)$$

$$Specificity = \frac{TN}{TN+FP} \times 100 \quad (3)$$

$$Precision = \frac{TP}{TP+FP} \times 100 \quad (4)$$

$$Recall = \frac{TP}{TP+FN} \times 100 \quad (5)$$

$$F1 - Score = 2 \times \frac{Precision \times Recall}{Precision + Recall} \quad (6)$$

$$MCC = \frac{TP \times TN - FP \times FN}{\sqrt{(TP+FP)(TP+FN)(TN+FP)(TN+FN)}} \quad (7)$$

$$AUC = \int_0^1 \frac{TP}{FP+TP} d \frac{FP}{FN+TN} = \frac{1}{P(FN+TN)} \int_0^1 TP d FP \quad (8)$$

In these equations, TP represents the number of malignant tumors accurately identified as malignant, FP represents the number of benign tumors misclassified as malignant, TN represents the number of benign tumors correctly identified as benign, and FN represents the number of malignant tumors misclassified as benign.

3.6. Potential Problem Limitations

Due to the complications in the diagnosis of breast cancer tumors in mammographic images, this research has been associated with challenges that are identified as potential problem limitations. Also, in order to provide a more complete view of this study, existing limitations are clarified to pave the way for future research. These limitations include the following:

- Lack of data diversity: The samples used may lack sufficient cultural, racial, and geographic diversity, potentially impacting the accuracy and generalizability of the results.
- High time and resource consumption: The preprocessing, feature extraction, and feature selection processes can be time-consuming and require significant computational resources. Also, the DOCRO algorithm may require more computational resources on very large feature sets, and its scalability should be investigated in the future.
- Need for labeled data: Training the proposed model requires labeled data whose preparation and careful examination can be both time-consuming and costly.
- Sensitivity to algorithm parameter settings: The research results may be highly sensitive to the settings of the algorithms used, meaning that changes in values and parameters can significantly affect the model's performance.

4. Proposed Method

This paper introduces a computer-based system designed to detect cancerous masses in mammographic images. The proposed method is composed of four phases: pre-processing, feature extraction, feature selection, and classification. The flowchart of the proposed method is depicted in *Figure 8*.

In the pre-processing phase, the region of interest (ROI) is extracted, and techniques like median filtering and histogram expansion are used to enhance the image quality. In the feature extraction phase, the features related to shape, histogram, and texture are derived from the ROI. Subsequently, the dual-objective CRO method is employed to select distinguishing features, and the SVM classifier is used to make the final decision.

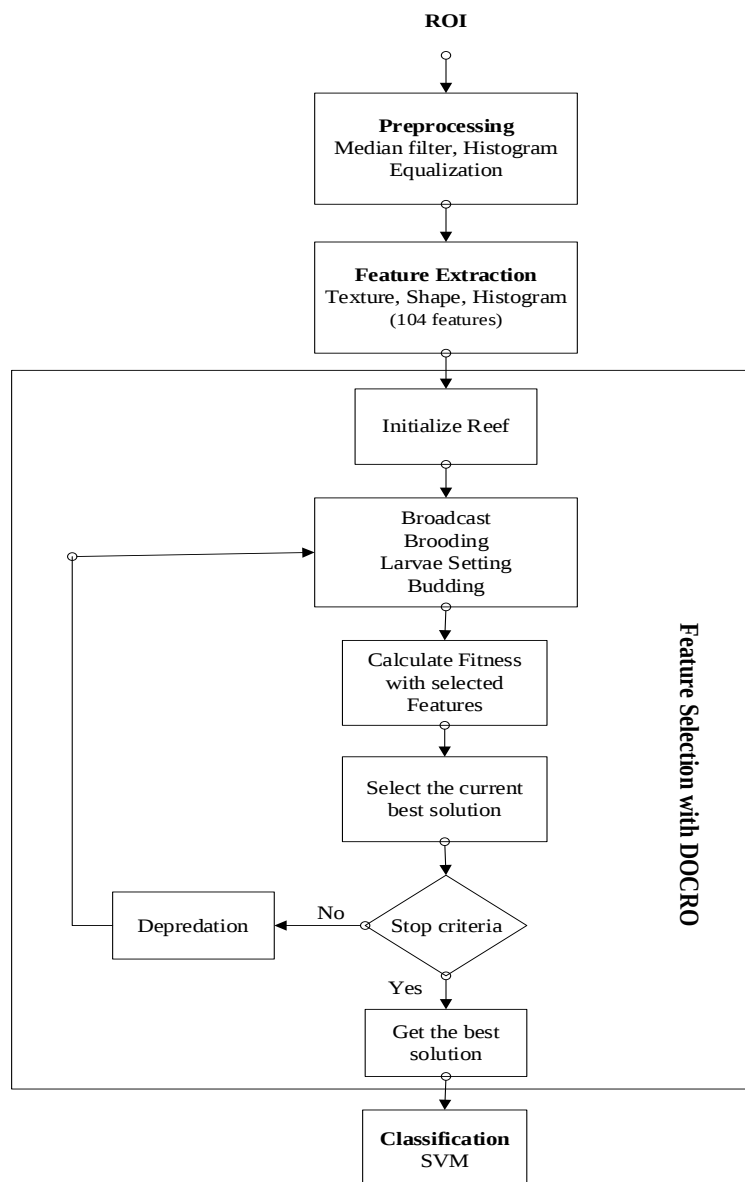


Figure 8. Flowchart of the proposed method

4.1. Extracting the region of interest (ROI)

Masses in mammographic images are the ROI for analysis and can be identified either automatically or manually. Automatic ROI determination is performed using clustering methods and image processing techniques. In the manual determination method, the mass area is cut from the breast image using the mass location information in the dataset. For the MIAS dataset, a text file contained information on the images, including the mass's location and radius information. Based on the information in the dataset's Readme file, the mass center and radius in mammography images were identified, and the ROI was extracted using a cutting method. For the CBIS-DDSM dataset, ROI images for masses can be downloaded and obtained in a cut form and these images were used. A .csv file containing full information about the mass type and imaging view is also included in the dataset.

4.2. Pre-processing

Following the extraction of the ROI from the breast image, the median filter is first applied, followed by histogram equalization to enhance it. The median filter is a non-linear spatial filter that substitutes the pixel value with the median of the intensity values within its vicinity. It causes less image blurriness than the smoothing filter with a similar size. Moreover, this filter proves efficient in eliminating various types of random noise. The median filter can be applied using an $n \times n$ square mask with different sizes. Histogram adjustment is applied to bolster contrast. This technique redistributes the image's histogram so that the intensity levels of the processed image cover a broader range, thereby amplifying the image contrast [42,43].

Median filter functions are used to clean noise in mammography images, while histogram equalization is used to increase image contrast so that tumor details are more clearly visible. The combination of these two techniques in pre-processing ensures that the data used for feature extraction is clean and has optimal contrast, which will increase the accuracy and effectiveness of breast cancer tumor detection in the next stage. **Figures 9** and **10** present the median filter and histogram equalization pseudocode.

Preserving fine details, such as tumor margins and microclassifications, is crucial in mammography image analysis. The median filter reduces noise and preserves edges and important features without blurring, by replacing each pixel's value with the median value of its neighbors. This filter keeps the main peaks of the histogram intact and does not alter the image's brightness distribution. Histogram equalization increases brightness contrast by expanding the image histogram, thereby highlighting brightness differences between tissues. This method improves the clarity of low-contrast areas and enables more accurate analysis. The combination of these two methods is effective, particularly for low-contrast images.

Pseudo-code of Median filter

Input:

Image: A 2D array representing the grayscale image

WindowSize: Size of the square filter window (must be odd, e.g., 3, 5, 7)

Output:

FilteredImage: A 2D array representing the image after applying the median filter

1: Begin

2: Initialize FilteredImage as a copy of Image.

3: Calculate the padding size: PadSize = Floor (WindowSize / 2).

4: Pad the input Image with PadSize pixels using the "replicate" or "zero" padding method.

5: **For** each pixel (i, j) in the original Image:

6: Extract the WindowSize x WindowSize neighborhood around the pixel from the padded image.

7: Flatten the neighborhood values into a 1D array.

8: Compute the median of the array.

9: Assign the median value to FilteredImage[i, j].

10: **end for**

11: Return FilteredImage.

9: End

Figure 9. Pseudocode of the median filter algorithm

Pseudo-code of Histogram equalization

Input: Gray levels of digital images as [0, K-1]**Output:** Histogram of the enhanced image**1: Begin**2: **for** $l \leftarrow 0$ to $K-1$ 3: Calculate distribution probability $P(r_l) \leftarrow \frac{n_l}{N}$ 4: Compute cumulative probability $C(r_l) \leftarrow \sum_{j=0}^{l-1} P(r_j)$, $0 \leq C(r_l) \leq 1$ 5: Compute equalization vector $U_l \leftarrow (K-1) * C(r_l)$ 6: Calculate distance between U_l and $U_l + 1$ as $\Delta U_l \leftarrow (K-1) * P(r_l)$ 7: **end for****8: End**

Figure 10. Pseudocode of the histogram equalization

4.3. Feature extraction

Feature extraction is a critical step in computer-aided systems for medical image diagnosis. The features extracted from images include all discernible structures, considering their nature. They quantify image properties and are used for classification. Extracted features must ensure optimal classification performance by maximizing similarity within the same category and dissimilarity between different categories. Benign and malignant breast masses exhibit distinct sizes and shapes. Masses that are round or oval with smooth, well-defined margins are typically benign. Conversely, masses with irregular shapes and indistinct, ragged margins are often malignant. Shape features describe the geometric structure of objects or regions under analysis. In mammographic images, the shape of a tumor is a critical factor for distinguishing between benign and malignant tumors.

Typically, malignant tumors have irregular shapes and poorly defined boundaries, whereas benign tumors tend to have smoother and more regular outlines. Histogram features are derived from the distribution of pixel intensities in an image. An intensity histogram represents the number of pixels corresponding to each intensity value, offering insights into the brightness and contrast distribution across the image. These features are particularly useful in analyzing medical images like mammograms, as they help identify variations in intensity that might indicate abnormalities or unusual tissue structures.

Moreover, malignant masses, unlike benign ones, display irregular textures, so texture features derived from the gray-level co-occurrence matrix (GLCM) are precious [4]. Texture features provide valuable information about the distribution patterns of pixel intensities within an image, playing a significant role in distinguishing tumor regions from healthy tissues. These features capture local intensity variations, the degree of regularity or irregularity, and the roughness or smoothness of specific image regions. In mammographic images, texture analysis can help identify abnormal tissues, such as tumors, which often exhibit more complex and irregular texture patterns compared to the uniform patterns of healthy or benign tissues.

The GLCM provides information about the relationships between adjacent pixel values in an image. Four angles, i.e., 0, 45, 90, and 135°, are used to determine pixel relationships, yielding a GLCM for each angle. The 22 features listed in **Table 3** are computed for each matrix. The extracted features comprise 88 texture features, 6 histogram features, and 10 shape features (see Table 3).

Table 3. Extracted features

Texture (GLCM)	Autocorrelation, Contrast, Correlation, Correlation: Matlab, Cluster Prominence, Cluster Shade, Dissimilarity, Energy, Entropy, Homogeneity, Maximum probability, Sum of squares: Variance, Sum average, Sum variance, Sum entropy, Difference variance, Difference entropy, Information measure of correlation1, Information measure of correlation2, Inverse difference (INV) is homom, Inverse difference normalized (INN), Inverse difference moment normalized
Histogram	Mean, Variance, Skewness, Kurtosis, Energy, Entropy
Shape	Solidity, Area, MajorAxisLength, MinorAxisLength, ConvexArea, Eccentricity, Orientation, Perimeter, Extent, EquivDiameter

4.4. Feature selection with dual-objective CRO (DOCRO)

The high number of features extracted at the feature extraction phase may lead to category overlap, increase data volume, and slow down processing. By selecting distinctive features, category overlap can be minimized with fewer features, which will reduce data volume. Feature selection is an NP-hard problem. In recent years, various meta-heuristic algorithms have been employed to solve the feature selection problem in medical datasets. The CRO algorithm is one such algorithm that can select an optimal subset of features to enhance classification performance [28].

Feature selection is a crucial issue in classification systems. Limiting the number of input features in a classifier helps create a good and computationally more compressed predictive model. In the field of medical diagnosis, a smaller feature subset means lower testing and diagnosis costs. The number of features used for classification is particularly important in classification problems, as some features in a high-dimensional feature vector may be inappropriate and redundant, have lower discrimination power, do not contribute to the classifier's learning process, and may even lead to class entanglement, which reduces classification accuracy and increases the learning model's

training time. Various feature selection methods can address this issue. Research shows that metaheuristic algorithms are effective in solving feature selection problems. The CRO is a cellular-type metaheuristic algorithm with excellent self-tuning exploration and exploitation capabilities. This motivated us to present and use an improved version of CRO, called DOCRO, in this study. DOCRO selects the best features, leading to increased classifier accuracy. In this method, the original CRO approach is used and only the larval adjustment process is modified. Once all larvae are produced through broadcast spawning and brooding, the larval adjustment process begins. This process enhances population diversity and helps avoid local optima. During the larval adjustment phase, the best larvae settle on the reef, and a battle starts for space. In DOCRO, the larva with the better fitness value wins the fight and occupies the place of the coral currently residing on the reef. Otherwise, the coral, which has occupied a certain space on the reef, wins, and the new larva either relocates on the reef or dies, depending on its current k value [44].

The binary-coded version of the DOCRO algorithm, designed to solve the feature selection problem, is outlined below. Then, the proposed fitness function used in this study is presented. Finally, the fundamental steps of the CRO algorithm, as applied to feature selection, are described and the selected features for both datasets are reported.

4.4.1. Encoding

In this section, an encoding scheme is initially used for feature selection. The coral reefs, REEF, are composed of an $N_1 \times N_2$ square grid, as depicted in **Figure 11(a)**. Squares occupied by coral larvae are marked as 1, while empty squares are marked as 0. In **Figure 11(a)**, each filled square, REEF_{ij}, in the REEF represents a coral larva from the coral population. As shown in **Figure 11(b)**, the initial population of each coral larva is modeled as a one-dimensional binary vector with a length of L . This vector is composed of a sequence of binary values, 0 and 1, denoting a subset of features. The binary values within the coral population represent the subset of selected features, where m signifies the number of squares occupied by coral larvae. Each coral larva encompasses a variety of features. If the j th feature is chosen, the value of the feature f_j is set to 1; otherwise, it remains 0.

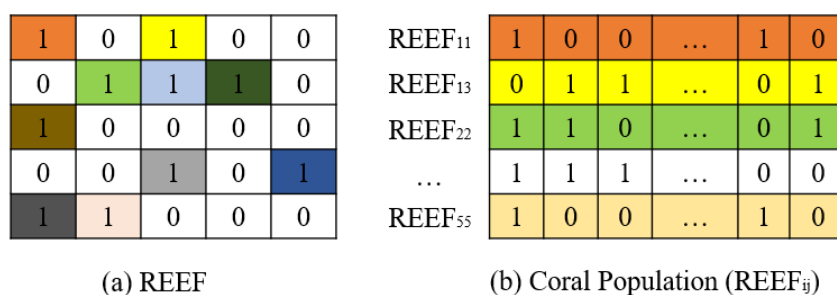


Figure 11. Representation of a reef and initial population of corals

4.4.2. Dual-objective fitness function

Two primary objectives in feature selection are to minimize data volume to enhance processing speed and to select suitable features to improve classification accuracy. In light of these objectives, the fitness function in CRO is formulated as **Eq. (9)**. This function is composed of two sections. The first section aims to augment classification accuracy by selecting features by which the maximum number of samples can be classified correctly. Samples are classified using the KNN method, and NC_i denotes the number of accurately classified samples based on the i th solution. The second

section aims to satisfy the objective of reducing data volume by choosing the smallest possible set of features. Employing the proposed dual-objective fitness function enables a simultaneous increase in classification accuracy and a significant reduction in data volume.

Objective 1: $\max(NC_i)$

Objective 2: $\min \left(\sum_{j=1}^{N2} REFF(i, j) \right)$

$fitness(i) = \{NC_i / \sum_{j=1}^{N2} REFF(i, j)\}$ (9)

4.4.3. Steps of running the CRO algorithm

Each coral larva represents a solution that is solved by competing in the growth space. The new solution is perpetually updated through the previously described coral reproduction process.

- **Initialization:** The coral population is initialized randomly. Each cell can contain only a single larva. The total number of population members must not be equal to or more than the number of cells in the reef matrix because if all matrix cells are filled, the algorithm will offer minimal opportunity for new solutions to be placed and proliferate within the matrix in the subsequent steps.
- **External reproduction operator:** The broadcast spawning process is executed by applying the crossover operator. Two corals, P1 and P2, are selected randomly to generate offspring, L1 and L2, through external sexual reproduction. It can be regarded as a two-point crossover operator in GA. **Figure 12** illustrates the CRO broadcast spawning in which the blue segment of P1 is exchanged with the yellow segment of P2, and the new results will be their offspring.

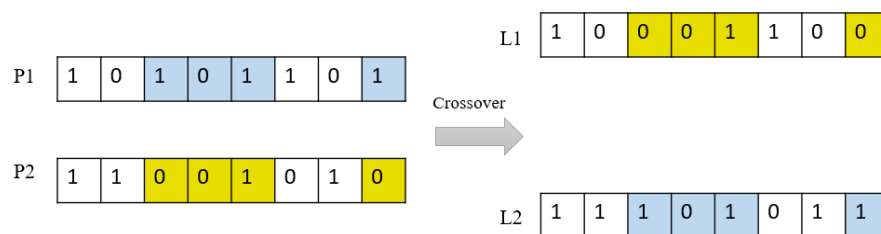


Figure 12. Example of Broadcast spawning

- **Internal reproduction operator:** The brooding process is executed through the random mutation of coral larvae, as illustrated in **Figure 13**. Initially, a coral larva, P, is randomly chosen and reversed into a new solution, L. **Figure 13** presents an instance of the internal reproduction operation. In this instance, bit #4 is selected at random, its value is altered from 0 to 1, and a new solution is created. This operator ensures that solutions do not converge toward a local minimum. The solution generated in this stage is added to the list of new solutions, which may subsequently have the chance to be set on the reef.

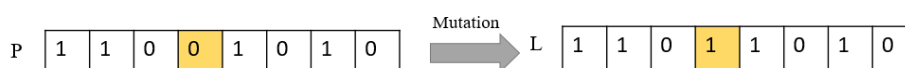


Figure 13. Example of brooding

- *Setting new solutions*: The best larvae are settled onto the reef to adjust the larvae, while those with the lowest fitness values are eliminated. The new solution generated during sexual reproduction searches for a grid to settle on. Only the healthiest corals with the highest fitness levels can successfully settle.
- *Asexual reproduction operator*: In the asexual reproduction step, all reef corals are ranked based on their health function (fitness function) values. The better corals are then copied to create new solutions.
- *Depredation*: Finally, a portion of worse solutions are purged in this process.

To assess the quality of each solution, **Eq. (9)** is employed to calculate the fitness value of each coral in the reef, indicating their health status. The algorithm continues to iterate until the termination conditions are satisfied.

4.4.4. Selected features

After applying the DOCRO method to the extracted features, 21 features were selected for the MIAS dataset, and 18 features were selected for the CBIS-DDSM dataset. **Table 4** presents the features selected by DOCRO for both datasets. The symbol ($^{\circ}$) above the numbers indicates the degree. The 21 features selected for the MIAS data included GLCM features at all four calculated angles. In addition, two histogram features and one shape feature were selected. Similarly, the 18 features selected for the CBIS-DDSM data included GLCM features at all four angles, two histogram features, and one shape feature. Some of the features selected were common to both datasets, making them dataset-resistant features for breast tumor detection.

Table 4. Selected features

21 Selected features for Mias dataset	Texture (GLCM)	Autocorrelation (45°,0°), Contrast (45°,0°), Cluster Prominence (135°,0°), Energy (90°,45°), Homogeneity (0°), Sum of squares: Variance (45°), Difference variance (0°), Difference entropy (90°), Information measure of correlation1(90°,45°), Inverse difference (INV) is homom (135°,90°), Inverse difference normalized (INN) (135°,90°)
	Histogram	Variance, Entropy
	Shape	Convex Area
18 Selected features for CBIS-DDSM dataset	Texture (GLCM)	Correlation: Matlab (0°), Cluster Prominence (45°,0°), Cluster Shade (90°), Dissimilarity (135°), Energy (45°), Entropy (90°), Sum average (135°), Sum entropy (0°), Difference variance (0°), Information measure of correlation1(45°), Information measure of correlation2 (90°,0°), Inverse difference (INV) is homom (0°), Inverse difference normalized (INN) (90°)
	Histogram	Kurtosis, Entropy
	Shape	Convex Area

4.5. Classification with SVM

After feature selection by the DO-CRO algorithm, cancer tumors are classified using SVM as a classical classifier that has shown remarkable results in the classification of two-class data. In linear SVM, the objective is to find a separating hyperplane that classifies samples into two classes by

maximizing the margin between them. In this classifier, boundary training samples, called support vectors, are calculated using an optimization algorithm, and the decision boundary is determined based on these support vectors. These support vectors for classes are the closest samples to the boundary selected to maximize the margin between the two classes. Consequently, only support vectors are needed to determine the boundary, which increases the margin and reduces the classification error. Therefore, SVM performs excellently in classifying linearly separable samples and is less sensitive to the number of training samples. In this study, SVM was used to detect whether tumors were cancerous or not, which was a two-class problem. The features selected by DOCRO do not allow for linear separation of the samples of the two classes, making linear SVM inadequate for detecting masses accurately. To increase detection accuracy, nonlinear SVM was used as it has been successfully applied to nonlinearly separable data in many cases.

To solve linearly inseparable problems, SVM maps data by non-linear kernels to a space with a higher dimension and finds the optimal linear boundary within that space. Polynomial and radial basis functions are the kernels commonly used for non-linear SVMs. The SVM model used in this study is the Gaussian radial basis function, which has demonstrated high efficacy in classification tasks, as per *Eq. (10)*, in which γ represents the width parameter, and its optimal value is determined through experimentation.

$$k(x, u) = \exp\left(-\frac{\|x-u\|^2}{\gamma^2}\right) \quad (10)$$

In some cases, even with a nonlinear kernel, samples may be located on the boundary of the separating hyperplane, thereby reducing the margin between the two classes. This can decrease the generalizability of the classifier to test samples and increase the classification error. This study used a soft margin to address this issue for which a penalty parameter (C) was added to the classifier relations. Using a soft margin allows some samples to be misclassified during training. By increasing the margin between the two classes, generalizability will be enhanced and ultimately the accuracy of test sample detection will be improved. By optimally adjusting the classifier parameters through experiments and considering the desired evaluation criteria according to *Eq. (1) to (8)*, the SVM classifier is expected to detect benign and malignant tumors with appropriate accuracy using the features selected by DOCRO.

5. Experimental Results

In order to demonstrate the performance of the proposed method in identifying cancerous tumors in mammographic images, the evaluations were conducted on the MIAS and CBIS-DDSM datasets. We utilized 100 images from the MIAS database, consisting of 50 images with benign masses and 50 with malignant masses, along with 300 images from the CBIS-DDSM database, consisting of 150 images with benign masses and 150 with malignant masses, to implement and evaluate the proposed method. Since the number of images with mass in the MIAS dataset is limited and in order to maintain a balance between classes, 50 samples were used for each class. The number of images with mass in the CBIS-DDSM dataset is much higher than in the MIAS dataset, but in order to fairly compare the performance of the method on the two datasets and evaluate the proposed method in a situation where we are faced with a limited number of samples, 150 samples were used for each class. This research also focused on providing a proof-of-concept for combining the dual-objective coral reef algorithm with the SVM classifier. In addition, the simultaneous use of two different

datasets provides a kind of cross-validation that reduces the risk of overfitting to a specific set and shows that the proposed method can be generalized to different data recording conditions.

For both datasets, the training and test samples were selected by 5-fold cross-validation. The classifier was constructed using the training samples, and its performance was evaluated with the test samples. This process was iterated for various combinations of training samples drawn from the entire dataset. The average outcomes for the evaluation parameters of the classifier were then reported. In the first stage of implementing the proposed method, the image of the ROI, which includes a benign or malignant mass, is entered into the pre-processing phase. At this phase, median filter functions are used to clean noise, and histogram equalization is used to increase image contrast so that tumor details become more clearly visible.

Figure 14 shows the preprocessing of ROI images for a benign mass (**Figure 14a**) and a malignant mass (**Figure 14g**) from the MIAS dataset. The changes made to the ROIs using the median filter and histogram adjustment are depicted, along with the corresponding histograms. Preserving fine details, such as tumor margins, microclassification, and abnormal tissue, is crucial in mammographic image analysis. Unlike linear filters like the Gaussian filter, which smooths edges, the median filter better preserves sharp edges and minor features by replacing each pixel's value with the median value of neighboring pixels. This reduces image noise without blurring important medical information and removes very bright (white dots) or very dark (black dots) pixels from the histogram. As shown in **Figure 14**, the median filter preserves brightness variations that typically represent edges or important structures in medical images as it replaces the pixel's value with the median value. Consequently, the main peaks in the image histogram remain intact. Unlike linear filters such as the Gaussian filter, which redistributes pixel values across the entire brightness range by blurring the image, the median filter does not introduce significant changes to the overall brightness distribution and preserves the original structure of the histogram.

A challenge in mammographic image analysis is their low contrast. Histogram equalization broadens the brightness distribution of the image, spreading brightness values more evenly across the entire possible range. This results in more prominent brightness differences between different tissues. As shown in **Figure 14**, adjusting the histogram makes minor brightness differences that may be undetectable in the original image more visible. Before adjustment, the image histogram is limited to a narrow range of brightness. After adjustment, the histogram expands, evenly distributing brightness across the entire range and enhancing the clarity of low-contrast areas.

In the feature extraction phase, 104 features related to shape, histogram, and texture are derived from the ROI. Subsequently, the dual-objective CRO method is employed to select 21 and 18 discriminative features for MIAS and CBIS-DDSM images, respectively. Finally, the SVM classifier is employed to make the final decision. In order to evaluate the effectiveness of the proposed method in classifying masses in mammography images, parameters of accuracy, sensitivity, specificity, area under the curve (AUC), Matthews correlation coefficient (MCC), and F1 score obtained from the confusion matrix were used. These standard evaluation parameters provide comprehensive insight into overall accuracy.

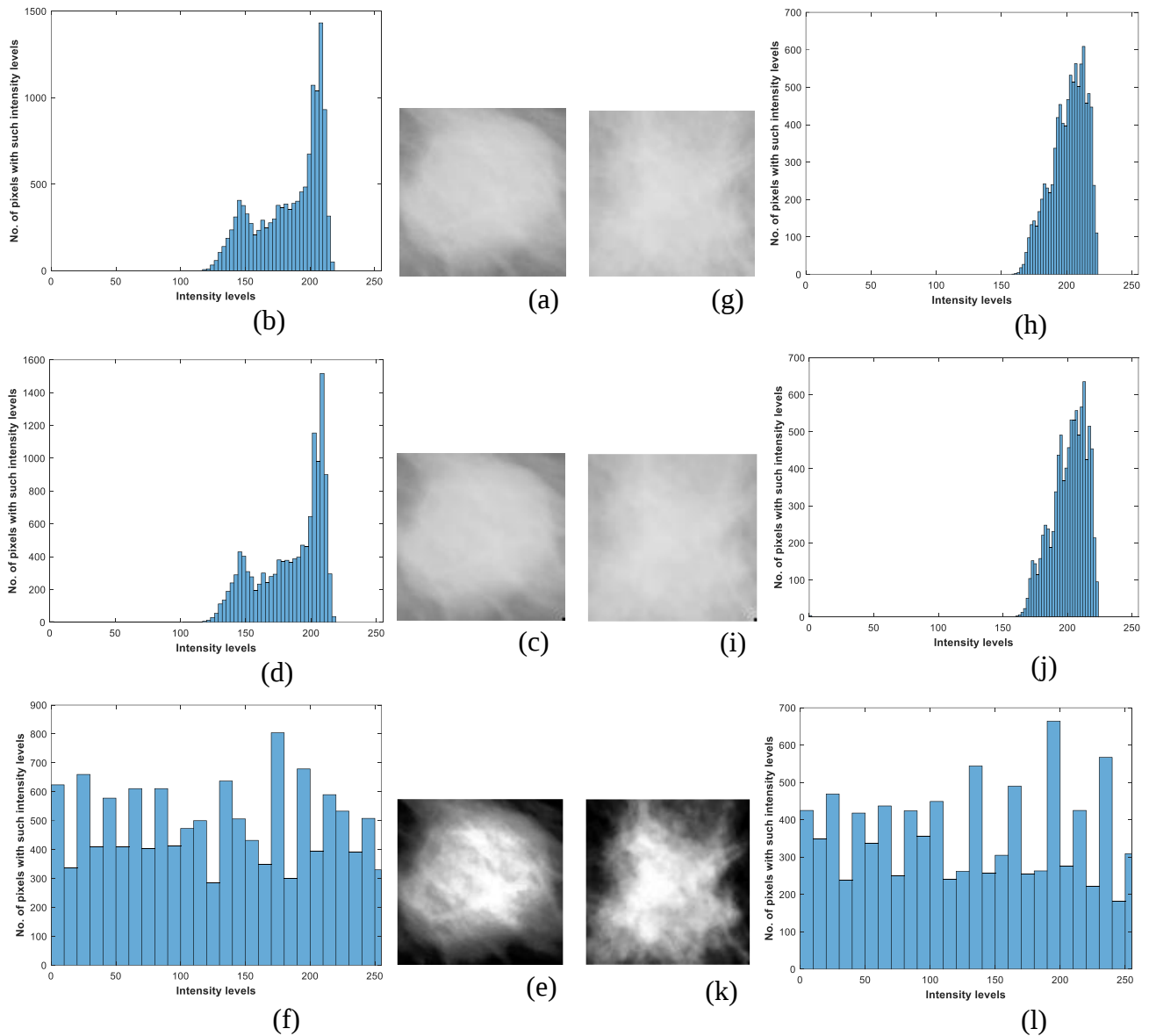


Figure 14. Preprocessing results: (a) original ROI image with benign mass, (b) histogram of image “a”, (c) result of applying the median filter on image “a”, (d) histogram of image “c”, (e) result of histogram equalization on image “c”, (f) histogram of image “e”, (g) original ROI image with malignant mass, and (h) histogram of image “g”

The confusion matrix is an important tool for evaluating cancer detection systems. **Table 5** reports the confusion matrices calculated by applying the test samples to each test iteration for both the MIAS and CBIS-DDSM datasets. According to **Table 5**, for the MIAS dataset using 21 features selected by DOCRO, all test samples were correctly identified in the best case. In the worst case, with $FN = 1$ and $FP = 3$, four samples were incorrectly labeled as other classes. There were two iterations in the best-case scenario and one iteration in the worst-case scenario. By applying 104 features to the classifier for the MIAS dataset, too, two samples were incorrectly identified with $FN = 2$ in the best case. In the worst case, with $FN = 1$ and $FP = 4$, five samples were incorrectly identified. For the CBIS-DDSM dataset, using 18 features selected by DOCRO, all test samples were correctly identified in the best case. In the worst case, with $FN = 3$ and $FP = 6$, nine samples were incorrectly labeled. There were three iterations in the best-case scenario and one iteration in the worst-case scenario. By applying 104 features to the classifier for the CBIS-DDSM dataset, in

the best case, with FN = 3 and FP = 2, five samples were misidentified. In the worst case, with FN = 9 and FP = 4, 13 samples were misidentified.

The confusion matrix allows for analyzing the system performance with more details. It is an instrument to identify the system's strengths and weaknesses. In the confusion matrix, FP indicates cases where the system mistakenly identifies non-cancerous samples as cancerous. A high FP rate suggests that the diagnostic system is overly cautious, mistakenly identifying many non-cancerous samples as cancerous. This can arise from the system's inability to distinguish subtle differences between cancerous and non-cancerous patterns. This error is not life-threatening and can be corrected by performing a biopsy. A higher FN shows that the diagnostic system is less cautious and mistakenly identifies cancerous samples as non-cancerous. This error is associated with the system's inability to extract appropriate features related to cancerous tissues and is very dangerous, as it can lead to disease progression and potentially the patient's death. TP represents the system's ability to correctly identify cancer patients. An increase in TP is essential to avoid missing cancer cases. TN indicates the system's ability to correctly identify non-cancer patients, reducing false alarms and unnecessary costs and stress. High TP and TN rates signify the system's high accuracy in correctly identifying masses. Based on the results in **Table 5**, the proposed system using features extracted by DOCRO demonstrates a good ability to recognize subtle differences between cancerous and non-cancerous patterns, reduce false alarms, and extract appropriate features related to cancer tissues.

Table 5. The confusion matrix by applying test samples for each experiment

		number of experiment									
		1		2		3		4		5	
MIAS dataset	21 selected features	10	0	10	0	10	0	10	0	9	1
		3	7	1	9	0	10	0	10	3	7
MIAS dataset	104 features	6	4	6	4	8	2	8	2	8	2
		1	9	1	9	0	10	0	10	2	8
CBIS-DDSM dataset	18 selected features	27	3	30	0	30	0	28	2	30	0
		6	24	0	30	0	30	3	27	0	30
CBIS-DDSM dataset	104 features	21	9	28	2	27	3	23	7	26	4
		4	26	5	25	2	28	1	29	4	26

Figures 15 to 20 depict the accuracy, sensitivity, specificity, F1-Score, MCC, and AUC in percentage for each iteration of the test for both datasets. These figures display the results of testing the proposed system using both the test dataset and the complete dataset. The TEST FIT and ALL FIT graphs show the results from the test and complete datasets, respectively, with features selected via the DOCRO algorithm. The TEST RAW and ALL RAW graphs illustrate the results derived from the datasets with the 104 originally extracted features.

Figure 15 displays the accuracy values in five test iterations for both the test dataset and the complete dataset. It is observed that all samples were accurately classified in two tests on the test datasets with the features selected by DOCRO. For the MIAS dataset, the minimum accuracy observed for the test samples with 21 selected features was 80%. The highest and lowest accuracy for all samples with 21 features were 96% and 94%, respectively. They were 87% and 85% for all samples with

104 features, respectively. For the CBIS-DDSM dataset, the minimum accuracy observed for the test samples with 18 selected features was 85%. The highest and lowest accuracy for all samples with 18 features were 97.66% and 96.33%, respectively. They were 90.66% and 88% for all samples with 104 features, respectively.

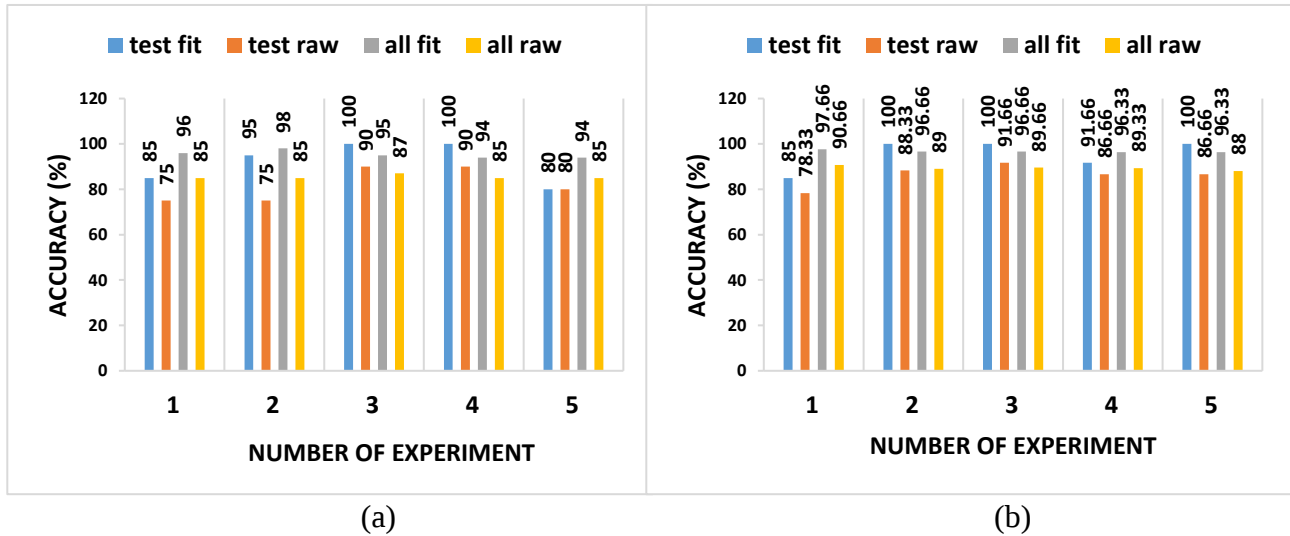


Figure 15. Accuracy values for each test iteration: (a) the MIAS dataset and (b) the CBIS-DDSM dataset

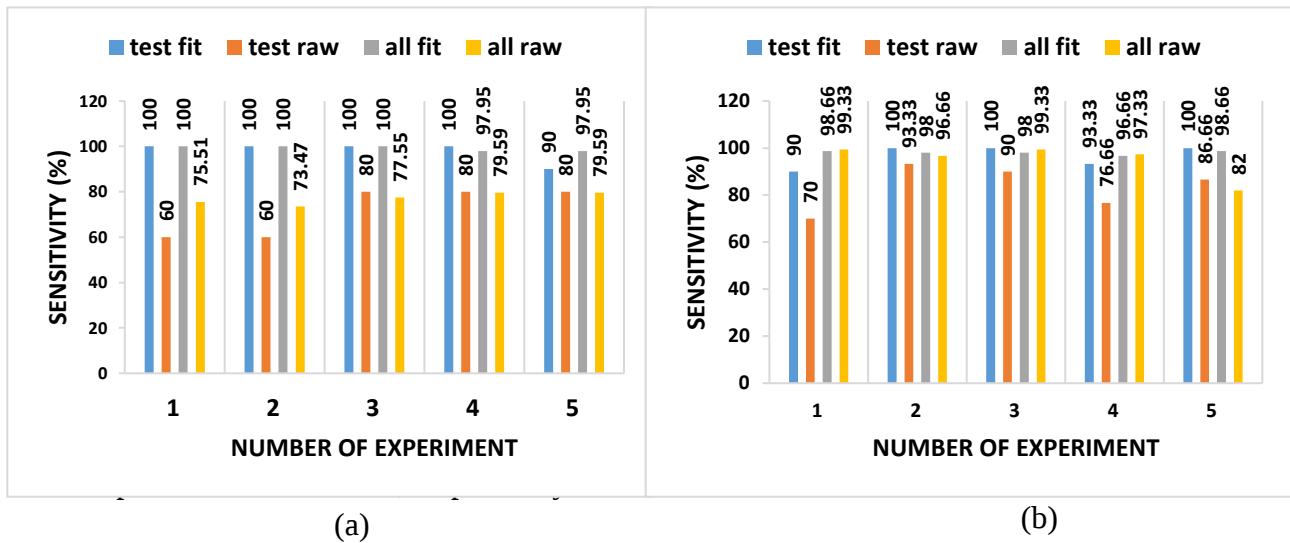


Figure 16. Sensitivity values for each test iteration: (a) the MIAS dataset and (b) the CBIS-DDSM dataset

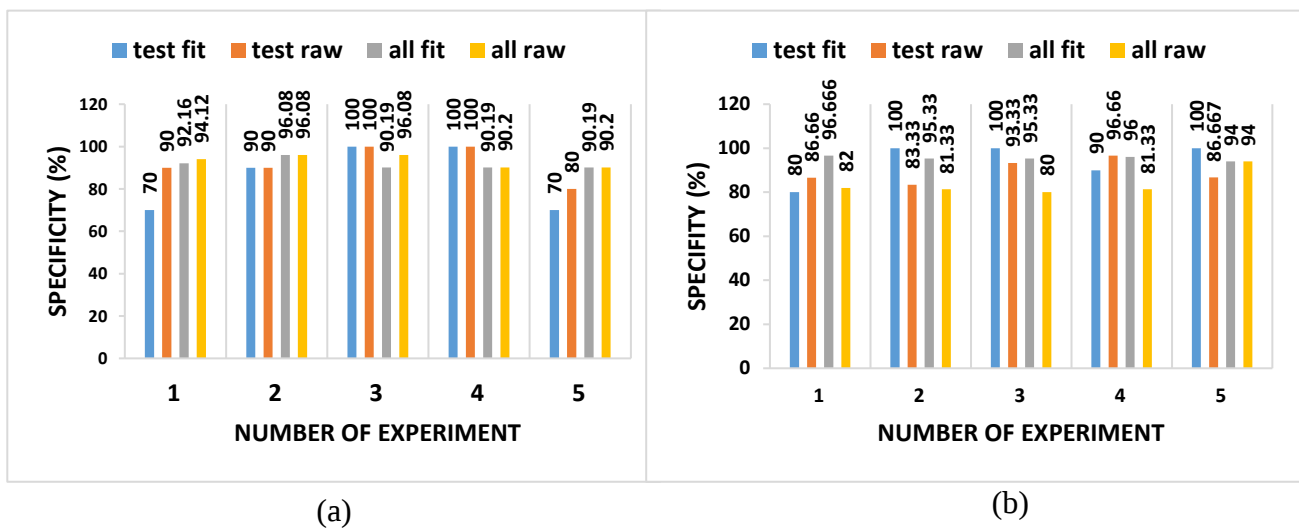


Figure 17. Specificity values for each test iteration: (a) the MIAS dataset and (b) the CBIS-DDSM dataset

Figure 16 shows that the system achieved 100% sensitivity in seven tests for the MIAS dataset. The lowest recorded sensitivity for the test samples for the MIAS dataset with 104 features was 60%, and the lowest recorded sensitivity for the test samples for the CBIS-DDSM dataset with 104 features was 70%. **Figure 17** displays the specificity values in five iterations for both the test dataset and the complete dataset. The system attained 100% specificity in four experiments for the MIAS dataset and in three experiments for the CBIS-DDSM dataset. For the MAS dataset, the maximum and minimum specificity were 96.08% and 90.19% for all samples with 21 features and 96.08% and 90.2% for all samples with 104 features, respectively. For the CBIS-DDSM, the maximum and minimum specificity were 96.66% and 94% for all samples with 18 features and 94% and 80% for all samples with 104 features, respectively.

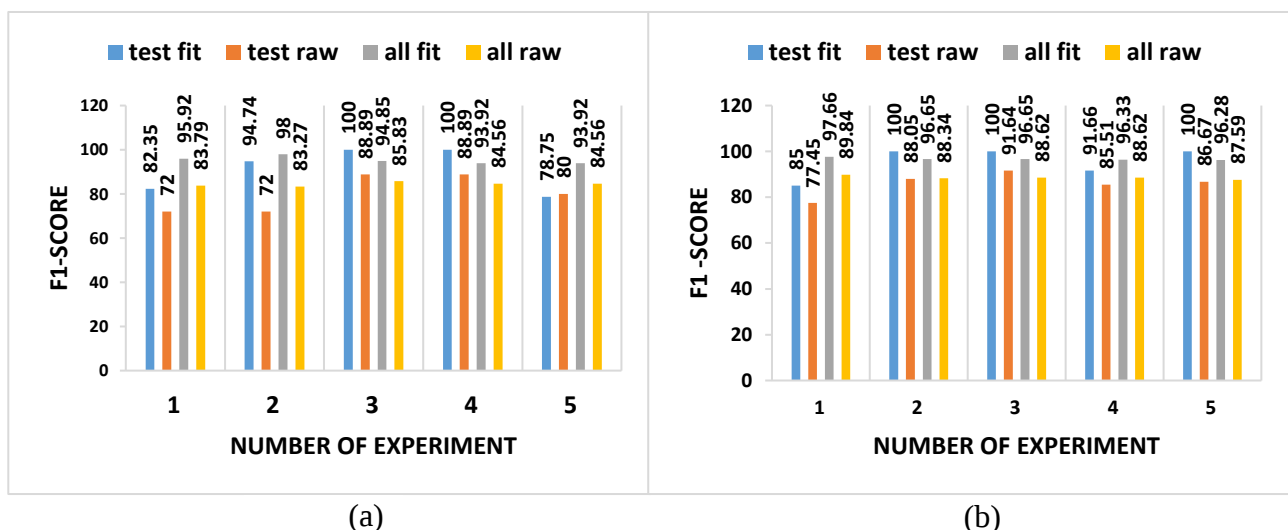


Figure 18. F1-Score values for each test iteration: (a) the MIAS dataset and (b) the CBIS-DDSM dataset

Figure 18 depicts the F1-Score values for five test iterations for both the test dataset and the complete dataset. The test results for the MIAS dataset indicate that in two iterations for test samples with 21 features, the system achieved an F1-Score of 100%. The lowest F1-Score for the test samples with 104 features was 72%. For the CBIS-DDSM dataset in three iterations for test samples with 18

features, the system achieved an F1-Score of 100%. The lowest F1-Score for the test samples with 104 features was 87.59%.

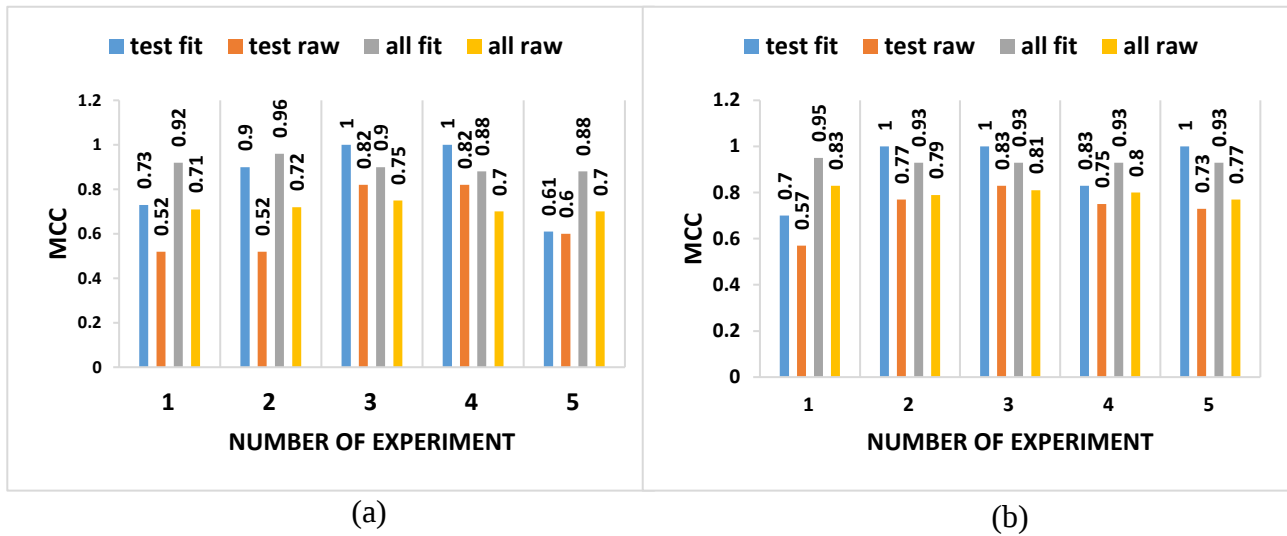


Figure 19. MCC values for each test iteration: (a) the MIAS dataset and (b) the CBIS-DDSM dataset

The MCC values for five iterations of testing on both the test data and the total data are depicted in **Figure 19**. For the MIAS dataset, the tests showed that in two iterations for the test data with 21 selected features, the system's MCC value achieved the optimal score of one. The highest and lowest MCC values for all samples with 21 features were 0.96 and 0.88, respectively. For all samples with 104 features, these values were 0.75 and 0.70, respectively. For the CBIS-DDSM dataset, the experiments showed that in three iterations for the test data with 18 selected features, the system's MCC value achieved the optimal score of one. The highest and lowest MCC values for all samples with 18 features were 0.95 and 0.93, respectively. For all samples with 104 features, these values were 0.83 and 0.77, respectively.

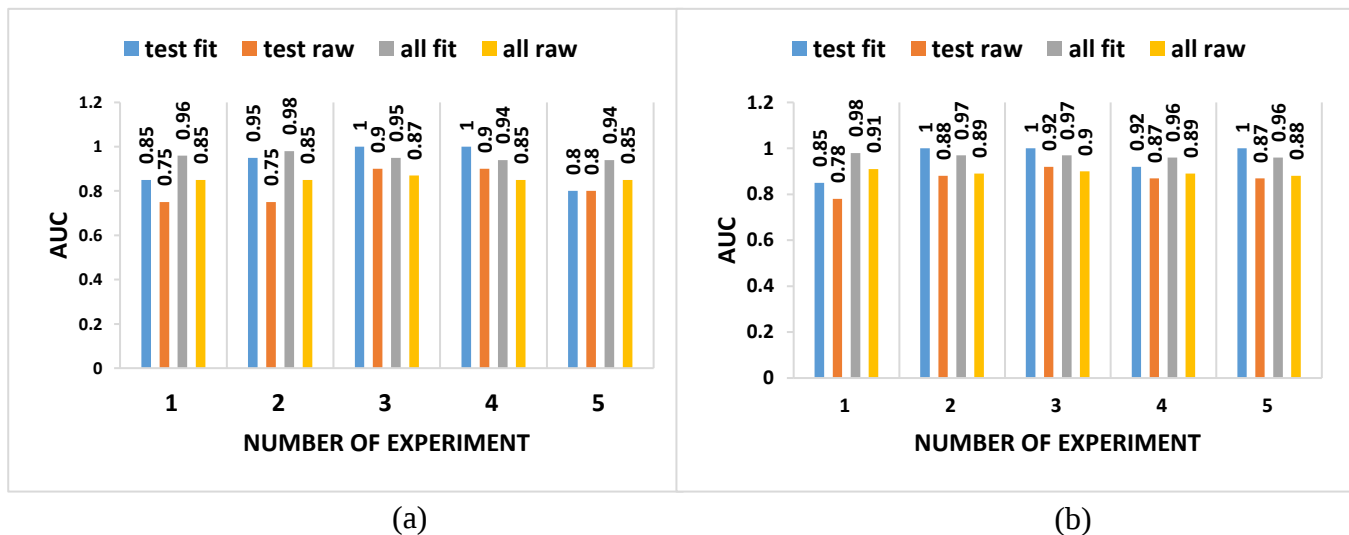


Figure 20. AUC values for each test iteration: (a) the MIAS dataset and (b) the CBIS-DDSM dataset

As is evident in **Figure 20**, according to the experiments, the AUC value reached one in two iterations for the test data with 21 features for the MIAS dataset, and the minimum AUC value for test data with 104 features was recorded at 0.75. For the CBIS-DDSM dataset, the AUC value reached one

in three iterations for the test data with 18 features, and the minimum AUC value for test data with 104 features was recorded at 0.78.

The results for accuracy, sensitivity, and specificity on data with 21 and 18 features selected by DOCRO for the MIAS and CBIS-DDSM datasets, respectively were obtained by changing the number of features in each iteration of the experiment. The average values are displayed in **Figures 21 to 23** and **24 to 26** for the MIAS and CBIS-DDSM datasets, respectively. According to these figures, the values for the three evaluation parameters had an upward trend as the number of features increased, peaking at 21 and 18 features for the MIAS and CBIS-DDSM datasets, respectively. Similar results were observed for the data with 104 features, as depicted in **Figures 27 to 32**, according to which an increase in the number of features did not result in an upward trend in the graph. The fluctuating patterns observed in these graphs suggest that certain features led to an overlap between classes, so they reduced the classifier's efficiency when they were included in the feature set. The use of a feature selection method can not only enhance the classifier's performance by eliminating redundant features but also reduce data volume. These figures confirm the positive impact of using a feature selection method on classification outcomes and demonstrate the effectiveness of the DOCRO method in selecting discriminative features for diagnosing breast cancer tumors.

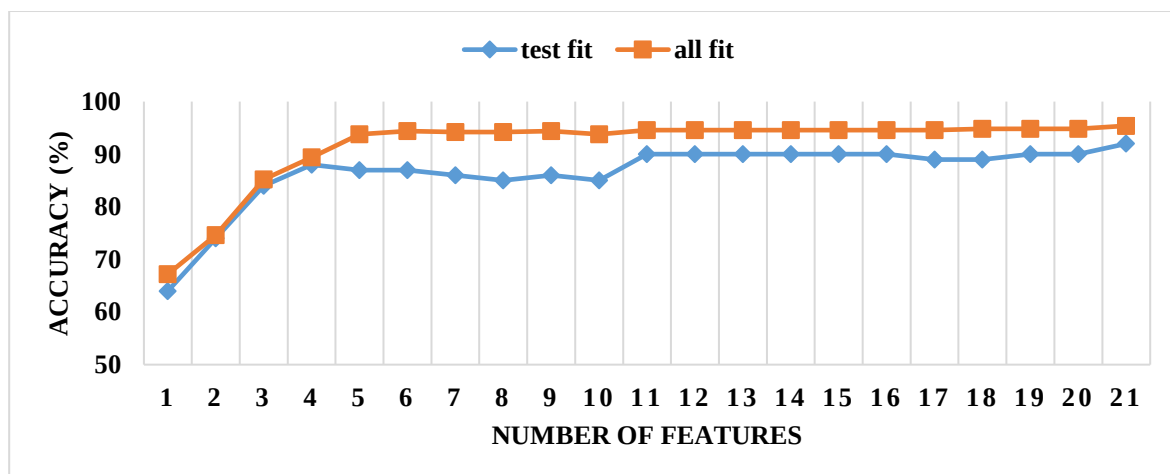


Figure 21. Accuracy values with changing the number of features for the MIAS dataset

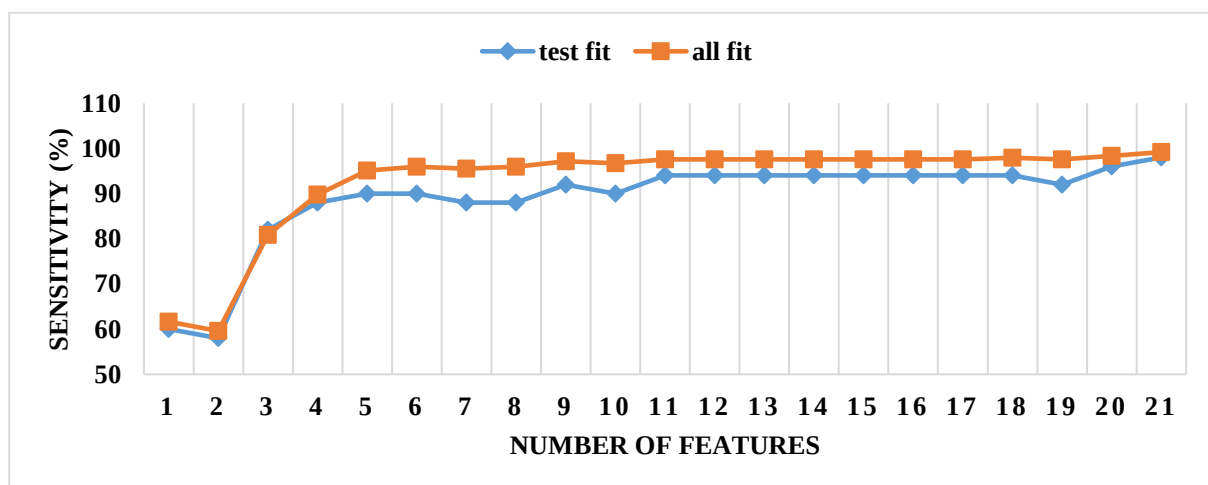


Figure 22. Sensitivity values with changing the number of features for the MIAS dataset

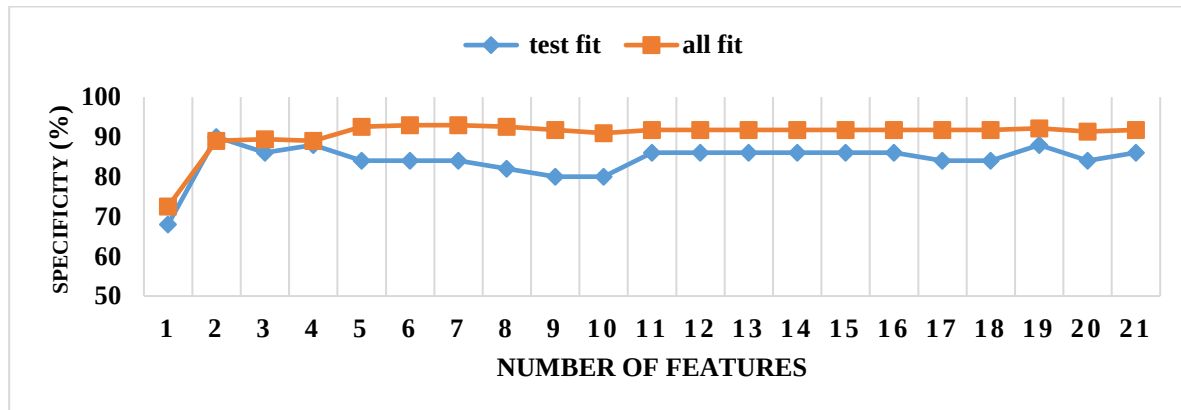


Figure 23. Specificity values with changing the number of features for the MIAS dataset

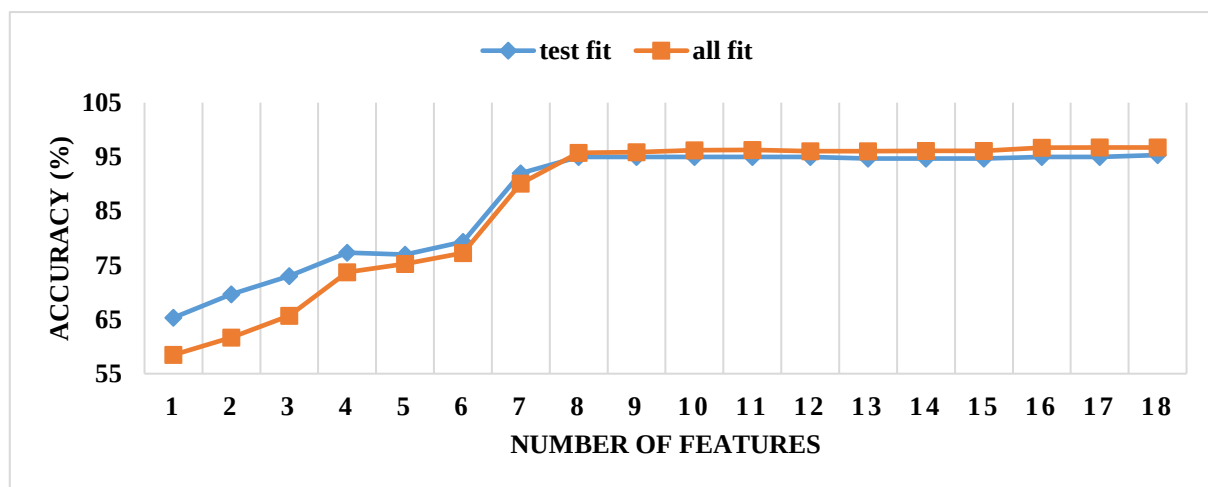


Figure 24. Accuracy values with changing the number of features for the CBIS-DDSM dataset

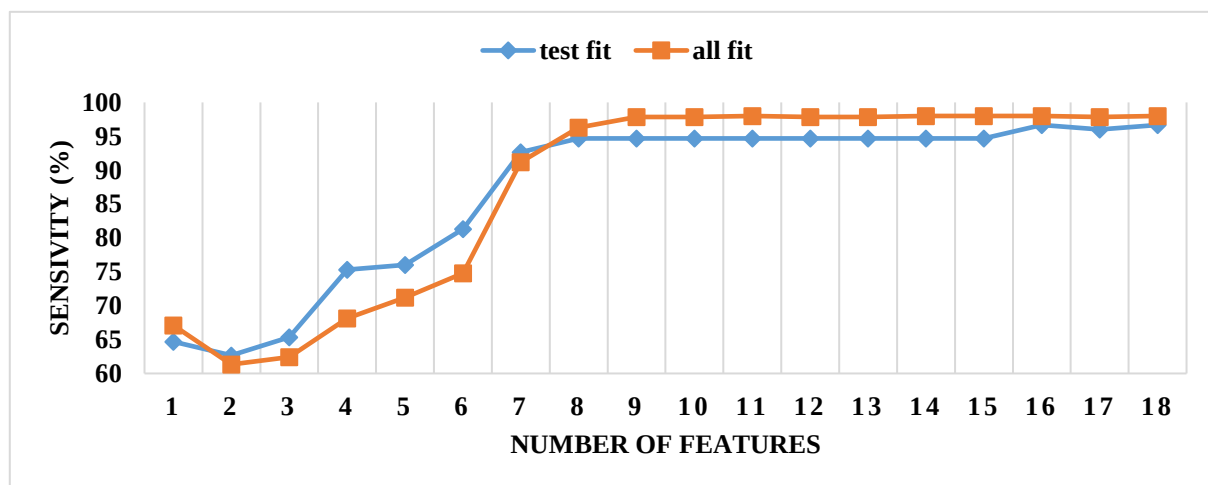


Figure 25. Sensitivity values with changing the number of features for the CBIS-DDSM dataset

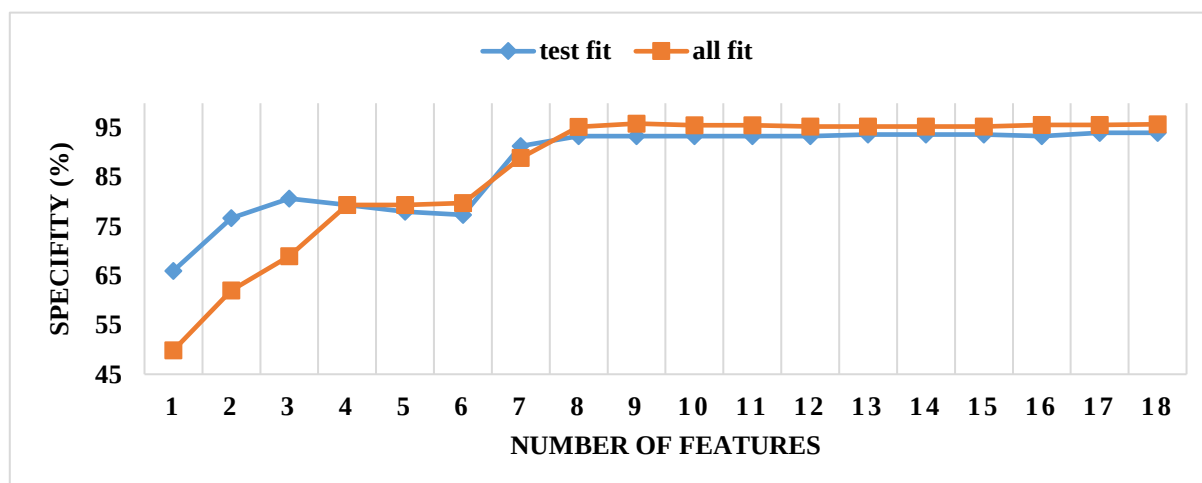


Figure 26. Specificity values with changing the number of features for the CBIS-DDSM dataset

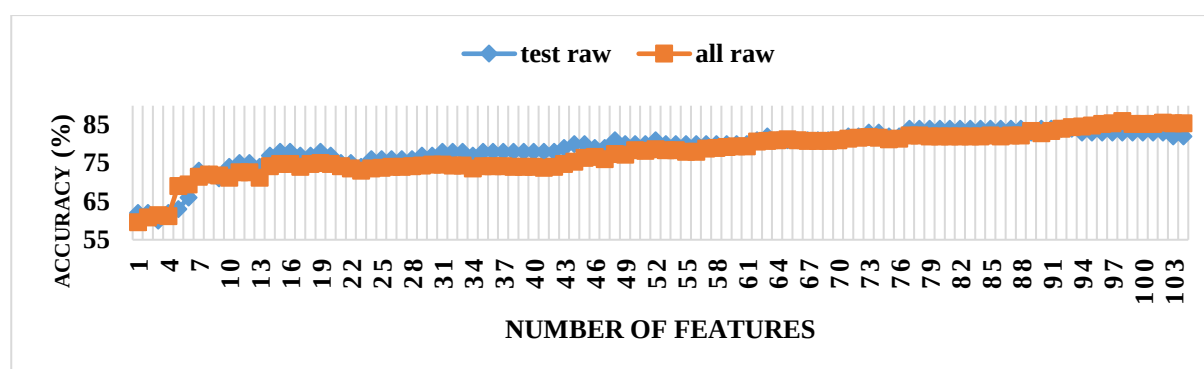


Figure 27. Accuracy values with changing the number of features for the MIAS dataset

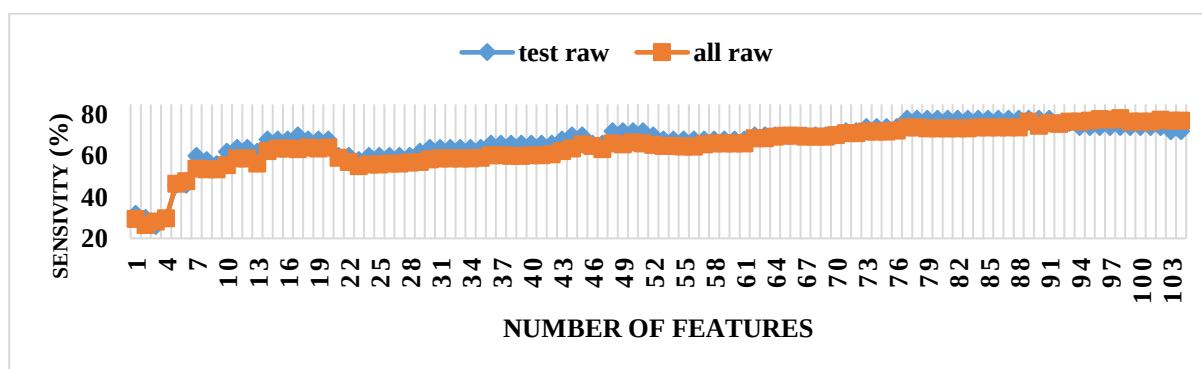


Figure 28. Sensitivity values with changing the number of features for the MIAS dataset

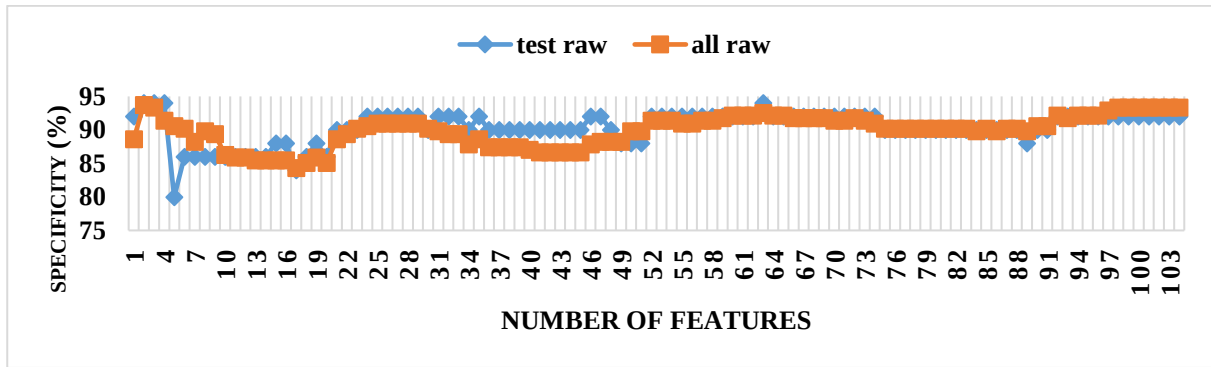


Figure. 29 Specificity values with changing the number of features for the MIAS dataset

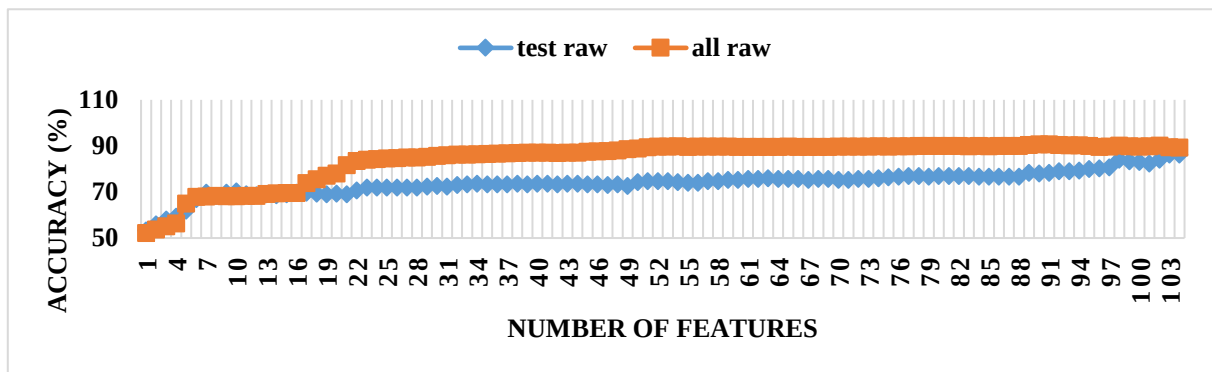


Figure 30. Accuracy values with changing the number of features for the CBIS-DDSM dataset

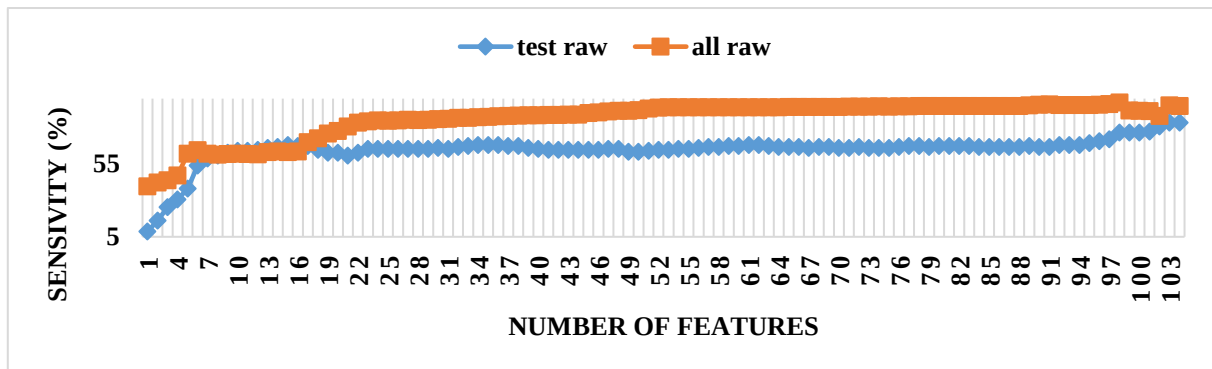


Figure 31. Sensitivity values with changing the number of features for the CBIS-DDSM dataset

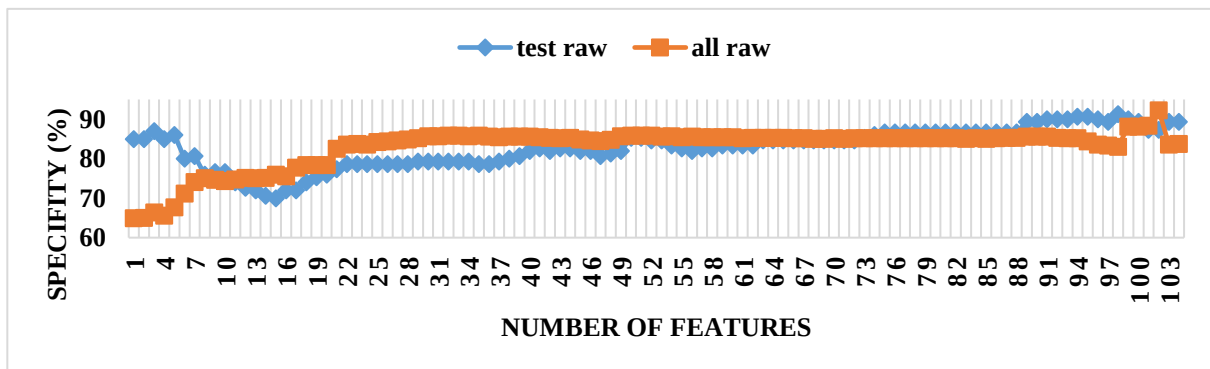


Figure 32. Specificity values with changing the number of features for the CBIS-DDSM dataset

Figure 33 displays the learning curve for the DOCRO feature extraction algorithm using the MIAS and CBIS-DDSM databases. **Figure 33(a)**, derived from the MIAS data, reveals a rapid convergence towards the optimal value. The blue line (best fitness) and the red line (mean fitness) stabilize swiftly. Minor fluctuations in this curve are attributed to random search efforts, reflecting the algorithm's continuous attempts to enhance values. The narrow gap between the best fitness and the mean fitness in the final generations signifies a reduction in population diversity and an overall convergence towards the optimal features. **Figure 33(b)**, derived from the CBIS-DDSM data, demonstrates rapid convergence towards the optimum. Both the blue line (best fitness) and the red line (mean fitness) reach stability around generation 50. Initial fluctuations in the red line highlight greater population diversity and efforts to explore the overall space, which diminish as the algorithm advances. The decreasing difference between the two lines in the final generations, signifying population convergence towards the optimal features, is evident. From generation 150 onwards, both lines are almost completely stable, reflecting the algorithm's success in identifying and maintaining the optimal solution over time. **Figure 33** illustrates that the DOCRO algorithm reached optimal values for both datasets in the early generations, successfully identifying effective features. Additionally, the minimal difference between the best fitness and the mean fitness, as well as their stability in the final stages, indicates that the majority of selected features are of high quality. This suggests that the algorithm successfully identified a set of optimal features that could enhance the performance of the diagnostic system. The achievement of optimal features by the DOCRO algorithm, and its effect on increasing the classification accuracy of both training and testing samples, can play a crucial role in preventing overfitting.

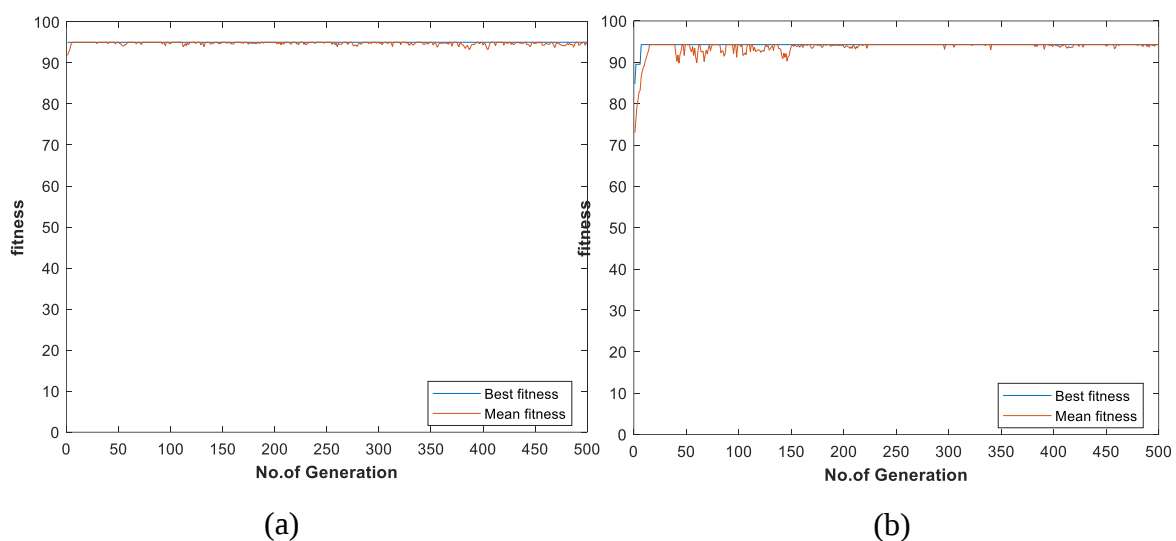


Figure 33. The learning curve of the DOCRO by (a) the MIAS dataset and (b) the CBIS-DDSM dataset

The method proposed for diagnosing breast cancer tumors was evaluated by the ROC curve. This curve serves as an assessment parameter for binary classification problems. The AUC, calculated by *Eq. (10)*, quantifies the classifier's capability to distinguish between classes. A higher AUC signifies the system's superior performance in identifying both positive and negative classes [45]. The ROC curve can indirectly indicate signs of overfitting. There is the possibility of overfitting if the AUC value is very high on the training data but significantly lower on the test data. This difference indicates that while the model has performed well on the training data, it cannot be generalized to new data.

As depicted in **Figures 34** and **35**, the ROC curve values range from 0 to 1. In these figures, the AUC value approached 1 for both test and total samples, indicating the proposed system's high efficiency in detecting breast cancer tumors. Also, there is no sign of overfitting by comparing the ROC graphs of the test and total data.

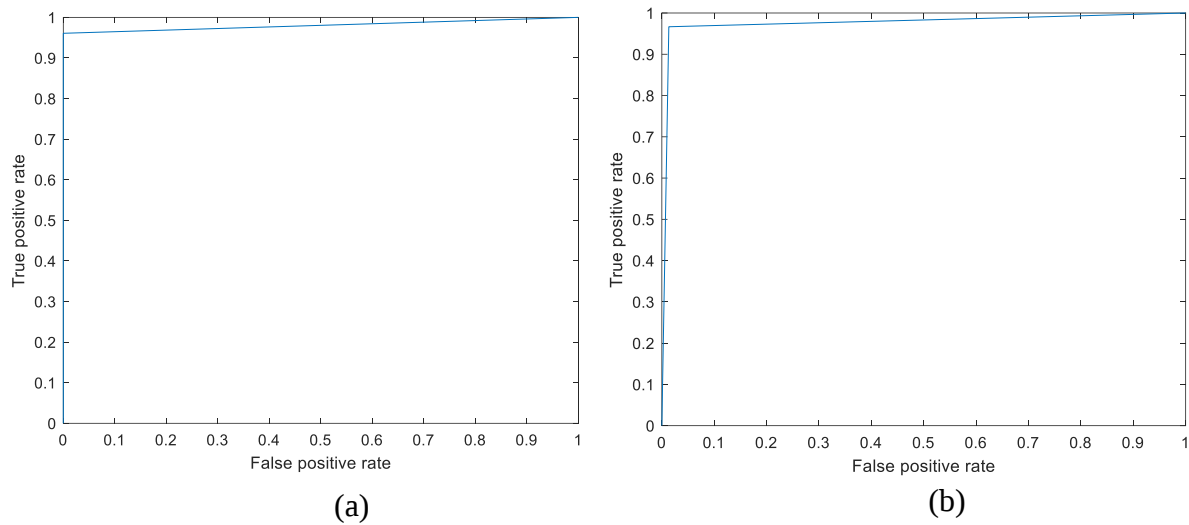


Figure 34. ROC curve of the proposed system by applying total samples for the best result, (a) the MIAS dataset and (b) the CBIS-DDSM dataset

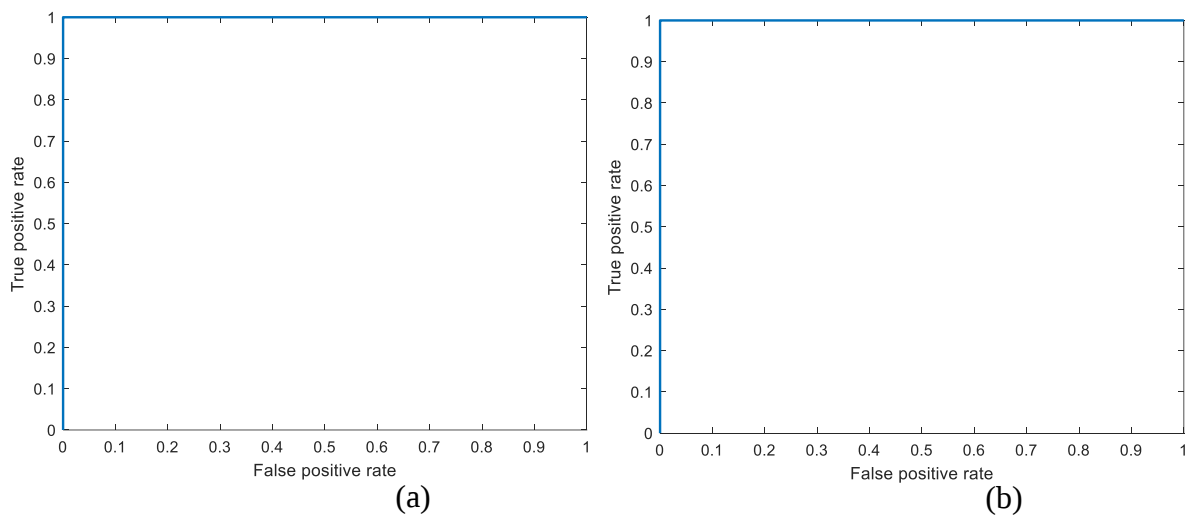


Figure 35. ROC curve of the proposed system by applying test samples for the best result, (a) the MIAS dataset and (b) the CBIS-DDSM dataset

Tables 6 and **7** report the mean and standard deviation of the evaluation parameters for the overall performance of the proposed system using samples with 21 and 104 features for the MIAS dataset and 18 and 104 features for the CBIS-DDSM dataset. According to the results, all evaluation parameters obtained for the samples with selected features by the DOCRO method showed a significant improvement over those for the samples with the initially extracted 104 features. The only exception was the specificity, for which a decrease was observed in the results for the MIAS dataset. A proper system must have high accuracy for the early diagnosis of breast cancer, which is crucial for saving lives. It is also of paramount importance to determine whether the tumor is malignant at the mammography stage. Accurate diagnosis of a cancerous tumor as malignant can

be life-saving. Hence, the proposed system must have a high TP. Also, detecting a non-cancerous mass in the proposed system is not life-threatening. Such a misdiagnosis can be corrected during biopsy, but misdiagnosing a cancerous mass as healthy (a false negative) can lead to disease progression and potentially the patient's death. Therefore, a low FP rate is critically important. Given the significance of FN and TP rates, it is evident that accuracy and sensitivity are more critical than specificity in evaluating a cancerous tumor detection system. For the MIAS dataset, the average accuracy and sensitivity for classification with 21 features selected by the DOCRO method were 10-15% and 22-25% higher than those obtained with 104 features, respectively. For the CBIS-DDSM dataset, the average accuracy for classification with 18 features selected by the DOCRO method was 7-9% higher than those obtained with 104 features. Other evaluation parameters also showed significant improvement for both datasets. The results obtained on the CBIS-DDSM dataset are better than those on the MIAS dataset, which could be due to the higher intensity resolution of the CBIS-DDSM images, potentially enhancing the ability to distinguish between tissue features and lesions. Although the significant results obtained for both datasets demonstrate the effectiveness of the proposed method for detecting cancerous tumors in mammography images. The proposed system utilizes shape, histogram, and texture features, acknowledging that cancerous tumors exhibit changes in shape, border, and texture compared to benign tumors. Texture features are extracted from multiple angles to ensure comprehensive mass characterization. However, some extracted features are redundant, so they increase data volume and negatively impact classification performance. Thus, the DOCRO method was employed for feature selection to optimize classifier performance. A key advantage of DOCRO over many meta-heuristic feature selection methods is that it does not immediately discard features that initially perform poorly, instead giving them a second or even more chances. Over several iterations, features that exhibit better performance grow, while poorer ones are eliminated from the reef. The quality of the extracted features was assessed using the proposed dual-objective fitness function, and the best category with 21 and 18 features was selected for the MIAS and CBIS-DDSM datasets, respectively. With this dual-objective fitness function, discriminative features were obtained that produced better results at the classification stage while significantly reducing data volume.

Table 6. The mean and standard deviation for the evaluation parameters of the proposed system for the MIAS dataset

The results obtained from SVM classification for the MIAS dataset		Test samples with 104 features	Total samples with 104 features	Test samples with 21 selected features	Total samples with 21 selected features
Accuracy (%)	Average	82	85.4	92	95.4
	Standard deviation	7.58	0.89	9.08	1.67
Sensitivity (%)	Average	72	77.14	98	99.18
	Standard deviation	10.95	2.66	4.47	1.12
Specificity (%)	Average	92	93.33	86	91.79
	Standard deviation	8.37	2.97	15.17	2.56
MCC	Average	0.66	0.72	0.85	0.91
	Standard deviation	0.15	0.02	0.17	0.03
AUC	Average	0.82	0.854	0.92	0.95
	Standard deviation	0.0758	0.0089	0.0908	0.016
F1-score	Average	80.36	84.4	91.17	95.32
	Standard deviation	8.45	0.97	10.01	1.71

Table 7. The mean and standard deviation for the evaluation parameters of the proposed system for the CBIS-DDSM dataset

The results obtained from SVM classification for the CBIS-DDSM dataset		Test samples with 104 features	Total samples with 104 features	Test samples with 18 selected features	Total samples with 18 selected features
Accuracy (%)	Average	86.33	89.33	95.33	96.73
	Standard deviation	4.91	0.97	6.81	0.54
Sensitivity (%)	Average	83.33	94.93	96.67	98
	Standard deviation	9.72	7.33	4.71	0.81
Specificity (%)	Average	89.33	83.73	94	95.46
	Standard deviation	5.48	5.79	8.94	0.98
MCC	Average	0.73	0.8	0.906	0.934
	Standard deviation	0.1	0.02	0.136	0.0089
AUC	Average	0.864	0.894	0.954	0.968
	Standard deviation	0.0513	0.0114	0.0677	0.0084
F1-score	Average	85.86	88.6	95.27	96.71
	Standard deviation	5.24	0.81	6.9248	0.56

In recent years, various methods have been proposed for analyzing abnormalities in medical images using metaheuristic and classification algorithms. Many of these methods have proven to be highly effective and efficient in assisting radiologists. This study used shape, histogram, and texture features for feature extraction. Given that not all features obtained from these methods may be effective or efficient and there may be redundancy or correlation between them, the DOCRO algorithm was proposed to select the most distinctive features. This algorithm pursues two important objectives: selecting a set of effective features to increase diagnosis accuracy, and reducing data dimensionality to minimize the computational burden. The results obtained from the proposed method for both the MIAS and CBIS-DDSM datasets, compared to the results of other methods (as shown in **Table 8**), demonstrate the remarkable ability of the proposed system to detect breast cancer masses, which can provide considerable confidence to physicians in using this method for diagnosing breast cancer. The results obtained show that the proposed method based on DOCRO-SVM has a very favorable and significant performance compared to classical classification methods; in such a way that accuracy, sensitivity and identifiability show significant improvements compared to traditional algorithms. In addition, choosing a smaller and more distinct set of features has reduced computational complexity and increased system efficiency. On the other hand, compared to deep learning-based approaches, although in some cases the final accuracy has been reported to be slightly lower, it should be noted that deep learning methods require a very large volume of data and significant computational power. While our proposed method has been able to provide competitive results with more limited data and less complexity, which is an important advantage for practical applications in resource-constrained environments.

Table 8. Comparison of the results obtained from the proposed method and a number of other methods

Author	Utilized dataset	Feature selection / Optimizing Hyperparameters	Classifier	Accuracy
Reenadevi, 2021[19]	MIAS	CSA	PNN	97%
			ANN	93.6%
			KNN	92.7%
			BPNN	93.67%
			SVM	91.6%
			RF	92.85%
Arumugam, 2021[20]	MIAS	-	LDA	82.5%
			cDTO	98.2%
Krishnavenia, 2020 [46]	MIAS	-	DTO	92%
			NBR	84%
Mohanty, 2020 [21]	MIAS	PCA	WC-SSA	99.28%
	DDSM			99.63%
	BCDR			99.60%
Sannasi Chakravarthy, 2019 [22]	MIAS		CSA	97.5%
			LDA	85%
Fang,2021[23]	MIAS, DDSM	WOA	MLP	98.1
Jyoti Kalita, 2022[13]	MIAS	2-way threshold based IWD	SVM	99%
			NB	98%
			k-NN	97.5%
			DT	95.5%
			RF	95%
Cai,2021[3]	MIAS	TEO	CNN	93.79%
			MLP	81.99
			TL	82.91
			MI	82.91
Jyoti Kalita,2022 [47]	MIAS	IWD	SVM	97.9%
Rajendran, 2022 [24]	MIAS	GOCSA	MLP	97%
		EAGA		90%
		TSWOA		88%
		BOA		84%
		WOA		83%
		GOA		86%
Chakravarthy, 2023 [48]	MIAS, INbreast, WDBC	PSO DFOA CSOA	NB, SVM	84.35%, 83.19%, 97.36%
Rezaee, 2020 [25]	MIAS	SA	MM-ANFIS (SFLAANFIS)	97.5%
Olaide, 2022[26]	MIAS	GA	CNN	76%
		WOA		75%
		MVO		84%
		SBO		80%
		LCBO		86%
kumari, 2021[27]	MIAS	SEDDA	MLP	94.5%
Punitha, 2021[49]	MIAS, WBCD, WDBC, WPBC, DDSM	ABC	ANN	99.25%
		WOA		
Houssein, 2022 [50]	MIAS,CBIS-DDSM	MPA	CNN	98.88%
				98.32%
Thirumalaisamy, 2023 [18]	MIAS,CBIS-DDSM	ACO	CNN	99.15
		OBL		98.63
Proposed method	MIAS	CRO	SVM	95.4%
	CBIS-DDSM	CRO	SVM	96.73

6. Conclusion

Breast cancer is characterized by the uncontrolled proliferation of breast cells. While there is, unfortunately, no definitive method to prevent breast cancer, its early and precise detection is crucial for appropriate treatment decisions. The CAD system aims to facilitate the early detection of breast cancer tumors and minimize diagnostic errors, thereby enabling timely treatment and reducing mortality rates. This paper introduced a method for identifying cancerous tumors in mammography images. The proposed method is composed of four phases: pre-processing, feature extraction, feature selection, and classification. Initially, the target mass area is delineated. Then, a median filter is used to remove the noise, and the histogram adjustment is used to increase the contrast. Subsequently, shape, histogram, and texture features are extracted from the image. The DOCRO method is, at the next step, applied to minimize the number of features by extracting the smallest yet most distinctive set for classification using a non-linear SVM. The proposed system was evaluated using 100 images from the MIAS and 300 images from the CBIS-DDSM databases. For the MIAS dataset, accuracy, sensitivity, specificity, AUC, MCC, and F1-Score values were estimated for the proposed method at 95.4%, 99.18%, 91.79%, 0.91, 0.95, and 95.32%, respectively. For the CBIS-DDSM database, these values were estimated for the proposed method at 96.73%, 98%, 95.46%, 0.934, 0.968, and 96.71%, respectively. The current results are based solely on quantitative measures and no statistical tests have been performed.

The results demonstrate that the application of the dual-objective fitness function in DOCRO allowed for selecting distinguishing features, which not only significantly reduced data volume but also yielded notable classification results, but to finally confirm its performance, testing on larger and more diverse datasets is needed. The proposed method is currently a proof of concept and needs to be tested in real clinical settings to finally confirm its performance.

7. Future work

The method proposed in this study was used to classify mammographic images using features extracted by the DOCRO algorithm. The evaluation data included 100 samples from the MIAS dataset and 300 samples from the CBIS-DDSM dataset. However, to further improve and develop this method, the following future research goals and strategies are suggested:

1. Evaluation with larger and more diverse datasets: Evaluating the proposed method with different and larger datasets can help check the robustness of the extracted features against dataset variations. The racial and geographical diversity in such datasets will contribute to the accuracy and stability of the system.
2. Application in other diagnostic areas: The proposed system can be examined in a wide range of other diagnostic areas to measure its efficiency and generalizability.
3. Combination with deep learning methods: Combining the proposed method with deep learning algorithms and networks can improve accuracy and efficiency. This combination is likely to lead to the extraction of more features and enhance the medical image classification process.

4. Using new metaheuristic algorithms for feature selection: Implementing new metaheuristic algorithms on the dataset to select the best features and comparing the results with the current method can help assess their ability to improve system performance.
5. Studying the effect of optimization methods and machine learning algorithms: Utilizing advanced optimization methods and machine learning algorithms to improve the system's accuracy and efficiency and evaluating their impact on diagnostic accuracy and reducing human errors can be effective.
6. Creating a comprehensive evaluation framework: It can be useful to develop a comprehensive evaluation framework that includes various criteria for accuracy, stability, and generalizability of the system under different conditions.
7. Developing and testing in real clinical environments: Implementing and testing the system in real clinical environments to evaluate its performance and reliability in practical conditions is advisable.
8. Use of statistical tests: Using statistical tests such as t-test and ANOVA to examine the significance of the results and ensure that the observed improvements are not simply random helps in improving the system's performance.

Although the results of the experiments on the MIAS and CBIS-DDSM datasets indicate the favorable performance of the proposed method, it should be noted that the experiments were conducted under limited sample conditions, which could affect the generalizability of the results. However, the simultaneous use of two different datasets provided a kind of cross-validation and reduced the risk of overfitting to a specific set. In the future, we plan to test the proposed method on larger and more diverse datasets to further verify its generalizability to different clinical and practical scenarios.

By pursuing these goals and strategies, it is hoped that the proposed method can become an effective and powerful tool for detecting and classifying medical images, thereby improving the accuracy and efficiency of medical diagnoses.

References

- [1] Tsai, C. C., Wang, C. Y., Chang, H. H., Chang, P. T. S., Chang, C. H., Chu, T. Y., ... Kuo, C. Y. (2024). Diagnostics and therapy for malignant tumors. *Biomedicines*, 12(12), 2659.
- [2] Ebrahimi, A., Arian, A., Sari, A. A., & Ahmadinejad, N. (2022). Diagnostic accuracy of opportunistic breast cancer screening based on mammography in Iran. *Iranian Journal of Radiology*, 19(3).
- [3] Cai, X., Li, X., Razmjoo, N., & Ghadimi, N. (2021). Breast cancer diagnosis by convolutional neural network and advanced thermal exchange optimization algorithm. *Computational and Mathematical Methods in Medicine*, 2021, 5595180.
- [4] Bagchi, S., Tay, K. G., Huong, A., & Debnath, S. K. (2020). Image processing and machine learning techniques used in computer-aided detection system for mammogram screening—A review. *International Journal of Electrical and Computer Engineering (IJECE)*, 10(3), 2336–2348.
- [5] Shaban, W. M. (2023). Insight into breast cancer detection: New hybrid feature selection method. *Neural Computing and Applications*, 35(9), 6831–6853.
- [6] Jiang, Y., Nishikawa, R. M., Schmidt, R. A., Metz, C. E., Giger, M. L., & Doi, K. (1999). Improving breast cancer diagnosis with computer-aided diagnosis. *Academic Radiology*, 6(1), 22–33.
- [7] Yeh, J. Y., & Chan, S. (2017, July). Population-based metaheuristic approaches for feature selection on mammograms. In *2017 IEEE International Conference on Agents (ICA)* (pp. 140–144). IEEE.

- [8] Bolon-Canedo, V., & Remeseiro, B. (2020). Feature selection in image analysis: A survey. *Artificial Intelligence Review*, 53(4), 2905–2931.
- [9] Pudjihartono, N., Fadason, T., Kempa-Liehr, A. W., & O'Sullivan, J. M. (2022). A review of feature selection methods for machine learning-based disease risk prediction. *Frontiers in Bioinformatics*, 2, 927312.
- [10] Salcedo-Sanz, S., Pastor-Sánchez, A., Prieto, L., Blanco-Aguilera, A., & García-Herrera, R. (2014). Feature selection in wind speed prediction systems based on a hybrid coral reefs optimization–Extreme learning machine approach. *Energy Conversion and Management*, 87, 10–18.
- [11] Beheshti, Z., & Shamsuddin, S. M. H. (2013). A review of population-based meta-heuristic algorithms. *International Journal of Advanced Soft Computing Applications*, 5(1), 1–35.
- [12] Sahoo, G., & Kumar, Y. (2012). Analysis of parametric & non-parametric classifiers for classification technique using WEKA. *International Journal of Information Technology and Computer Science (IJITCS)*, 4(7), 43.
- [13] Kalita, D. J., Singh, V. P., & Kumar, V. (2022). Two-way threshold-based intelligent water drops feature selection algorithm for accurate detection of breast cancer. *Soft Computing*, 26(5), 2277–2305.
- [14] Salcedo-Sanz, S., Del Ser, J., Landa-Torres, I., Gil-López, S., & Portilla-Figueras, J. A. (2014). The coral reefs optimization algorithm: A novel metaheuristic for efficiently solving optimization problems. *The Scientific World Journal*, 2014, 739768.
- [15] Yan, C., Ma, J., Luo, H., & Patel, A. (2019). Hybrid binary coral reefs optimization algorithm with simulated annealing for feature selection in high-dimensional biomedical datasets. *Chemometrics and Intelligent Laboratory Systems*, 184, 102–111.
- [16] Veeranjanyulu, K., Lakshmi, M., & Janakiraman, S. (2025). Swarm intelligent metaheuristic optimization algorithms-based artificial neural network models for breast cancer diagnosis: Emerging trends, challenges and future research directions. *Archives of Computational Methods in Engineering*, 32(1), 381–398.
- [17] Bourouis, S., Band, S. S., Mosavi, A., Agrawal, S., & Hamdi, M. (2022). Meta-heuristic algorithm-tuned neural network for breast cancer diagnosis using ultrasound images. *Frontiers in Oncology*, 12, 834028.
- [18] Thirumalaisamy, S., Thangavilou, K., Rajadurai, H., Saidani, O., Alturki, N., Mathivanan, S. K., ... Gochhait, S. (2023). Breast cancer classification using synthesized deep learning model with metaheuristic optimization algorithm. *Diagnostics*, 13(18), 2925.
- [19] Reenadevi, R., Sathiya, T., & Sathiyabhama, B. (2021). Classification of digital mammogram images using wrapper based chaotic crow search optimization algorithm. *Annals of the Romanian Society for Cell Biology*, 25(5), 2970–2979.
- [20] Arumugam, K., Ramasamy, S., & Subramani, D. (2021). Chaotic duck traveler optimization (cDTO) algorithm for feature selection in breast cancer dataset problem. *Turkish Journal of Computer and Mathematics Education*, 12(4), 250–262.
- [21] Mohanty, F., Rup, S., Dash, B., Majhi, B., & Swamy, M. N. S. (2020). An improved scheme for digital mammogram classification using weighted chaotic salp swarm algorithm-based kernel extreme learning machine. *Applied Soft Computing*, 91, 106266.
- [22] SR, S. C., & Rajaguru, H. (2019). Comparison analysis of linear discriminant analysis and cuckoo-search algorithm in the classification of breast cancer from digital mammograms. *Asian Pacific Journal of Cancer Prevention*, 20(8), 2333.
- [23] Fang, H., Fan, H., Lin, S., Qing, Z., & Sheykhahmad, F. R. (2021). Automatic breast cancer detection based on optimized neural network using whale optimization algorithm. *International Journal of Imaging Systems and Technology*, 31(1), 425–438.
- [24] Rajendran, R., Balasubramaniam, S., Ravi, V., & Sennan, S. (2022). Hybrid optimization algorithm based feature selection for mammogram images and detecting the breast mass using multilayer perceptron classifier. *Computational Intelligence*, 38(4), 1559–1593.
- [25] Rezaee, K., Rezaee, A., Shaikhi, N., & Haddadnia, J. (2020). Multi-mass breast cancer classification based on hybrid descriptors and memetic meta-heuristic learning. *SN Applied Sciences*, 2(7), 1297.
- [26] Oyelade, O. N., & Ezugwu, A. E. (2022). Characterization of abnormalities in breast cancer images using nature-inspired metaheuristic optimized convolutional neural networks model. *Concurrency and Computation: Practice and Experience*, 34(4)*, e6629.
- [27] Kumari, S. V., & Rani, K. U. (2021). SEDDA—An optimal feature selection approach for breast cancer diagnosis. *Webology*, 18(5).

- [28] Abouelmagd, L. M., Shams, M. Y., El-Attar, N. E., & Hassanien, A. E. (2021). Feature selection based coral reefs optimization for breast cancer classification. In *Medical informatics and bioimaging using artificial intelligence: Challenges, issues, innovations and recent developments* (pp. 53–72). Springer International Publishing.
- [29] Ahmed, S., Ghosh, K. K., Garcia-Hernandez, L., Abraham, A., & Sarkar, R. (2021). Improved coral reefs optimization with adaptive β -hill climbing for feature selection. *Neural Computing and Applications*, 33(12), 6467–6486.
- [30] Emami, M., Nazif, S., Mousavi, S. F., Karami, H., & Daccache, A. (2021). A hybrid constrained coral reefs optimization algorithm with machine learning for optimizing multi-reservoir systems operation. *Journal of Environmental Management*, 286, 112250.
- [31] Romero-Barrera, A. J., Cruz-Piris, L., Tejedor-Romero, M., Gimenez-Guzman, J. M., & Marsa-Maestre, I. (2024). An environmentally adaptive CRO-SL algorithm based on dynamic agents for the channel assignment problem in wireless networks. *IEEE Access*.
- [32] Marcelino, C. G., Pérez-Aracil, J., Wanner, E. F., Jiménez-Fernández, S., Leite, G. M. C., & Salcedo-Sanz, S. (2023). Cross-entropy boosted CRO-SL for optimal power flow in smart grids. *Soft Computing*, 27(10), 6549–6572.
- [33] Hossein, A., & Karim, S. M. (2019). A novel coral reefs optimization algorithm for materialized view selection in data warehouse environments. *Applied Intelligence*, 49(11), 3965–3989.
- [34] Abdiansah, A., & Wardoyo, R. (2015). Time complexity analysis of support vector machines (SVM) in LibSVM. *International Journal of Computer Applications*, 128(3), 28–34.
- [35] Buongiorno, R., Caudai, C., Colantonio, S., & Germanese, D. (2023). Introduction to machine learning in medicine. In *Introduction to artificial intelligence* (pp. 39–68). Springer International Publishing.
- [36] Awad, M., & Khanna, R. (2015). Support vector machines for classification. In *Efficient learning machines: Theories, concepts, and applications for engineers and system designers* (pp. 39–66). Apress.
- [37] Dror, B., Yanai, E., Frid, A., Peleg, N., Goldenthal, N., Schlesinger, I., ... Raz, S. (2014, December). Automatic assessment of Parkinson's disease from natural hands movements using 3D depth sensor. In *2014 IEEE 28th Convention of Electrical & Electronics Engineers in Israel (IEEEI)* (pp. 1–5). IEEE.
- [38] Suckling, J., Parker, J., Dance, D., et al. (2015). *Mammographic Image Analysis Society (MIAS) database v1.21* [Data set]. Apollo—University of Cambridge Repository.
- [39] Sawyer-Lee, R., Gimenez, F., Hoogi, A., & Rubin, D. (2016). *Curated breast imaging subset of digital database for screening mammography (CBIS-DDSM)* [Data set]. The Cancer Imaging Archive.
- [40] Metz, C. E. (1978, October). Basic principles of ROC analysis. In *Seminars in Nuclear Medicine*, 8(4), 283–298. WB Saunders.
- [41] Sadad, T., Munir, A., Saba, T., & Hussain, A. (2018). Fuzzy C-means and region growing based classification of tumor from mammograms using hybrid texture feature. *Journal of Computational Science*, 29, 34–45.
- [42] Gonzalez, R. C., & Woods, R. E. (2002). *Digital image processing* (2nd ed.). Prentice-Hall.
- [43] Rajendran, R., Balasubramaniam, S., Ravi, V., & Sennan, S. (2022). Hybrid optimization algorithm based feature selection for mammogram images and detecting the breast mass using multilayer perceptron classifier. *Computational Intelligence*, 38(4), 1559–1593.
- [44] Salcedo-Sanz, S. (2017). A review on the coral reefs optimization algorithm: New development lines and current applications. *Progress in Artificial Intelligence*, 6(1), 1–15.
- [45] Raiesdana, S. (2021). Breast cancer detection using optimization-based feature pruning and classification algorithms. *Middle East Journal of Cancer*, 12(1), 48–68.
- [46] Krishnavenia, A., Shankar, R., & Duraisamy, S. (2020). An efficient methodology for breast tumor segmentation using Duck traveler optimization algorithm. *PalArch's Journal of Archaeology of Egypt/Egyptology*, 17(9), 7747–7759.
- [47] Kalita, D. J., Singh, V. P., & Kumar, V. (2022). Detection of breast cancer through mammogram using wavelet-based LBP features and IWD feature selection technique. *SN Computer Science*, 3(2), 175.
- [48] Chakravarthy, S. S., Bharanidharan, N., & Rajaguru, H. (2023). Deep learning-based metaheuristic weighted k-nearest neighbor algorithm for the severity classification of breast cancer. *IRBM*, 44(3), 100749.
- [49] Stephan, P., Stephan, T., Kannan, R., & Abraham, A. (2021). A hybrid artificial bee colony with whale optimization algorithm for improved breast cancer diagnosis. *Neural Computing and Applications*, 33(20), 13667–13691.

- [50] Houssein, E. H., Emam, M. M., & Ali, A. A. (2022). An optimized deep learning architecture for breast cancer diagnosis based on improved marine predators' algorithm. *Neural Computing and Applications*, 34(20), 18015–18033.