

HYBRID DRL-CTO: OPTIMIZATION DRIVEN DEEP LEARNING FRAMEWORK FOR CEREBELLAR ATAXIA BASED ON HUMAN GAIT ANALYSIS

EDARA SREENIVASA REDDY¹, SUNIL PREM KUMAR PRATHIPATI²

¹Professor, VIT AP, SCOPE, India

²Research Scholar, ACHARYA NAGARJUNA UNIVERSITY, Department of CSE, India

E-mail: ¹Sreenivasareddy.e@vitap.ac.in, ²Sunilpremnelson@gmail.com

ABSTRACT

This study presents a Hybrid Deep Repeated Learning with Competitive Tuning Optimization (Hybrid DRL-CTO) framework for detecting and classifying Cerebellar Ataxia (CA) using human gait analysis. The proposed approach integrates a repeated learning mechanism with the CTO algorithm to enhance parameter optimization, improve convergence, and reduce computational complexity. Deep S3P features extracted from regions of interest in gait frames enable the model to capture discriminative representations while minimizing feature dimensionality, thereby improving classification accuracy. Experimental evaluations demonstrate that Hybrid DRL-CTO achieves 96.8% accuracy, 95.1% precision, and 97.1% recall, consistently outperforming existing techniques in k-fold validation. The combination of CNN-LSTM with the CTO algorithm further prevents overfitting and ensures reliable performance, offering a robust solution for CA classification with high efficiency and reduced computational effort.

Keywords: *Cerebellar Ataxia, Repeated Learning, Deep Learning, Gait Analysis, Competitive Tuning Optimization.*

1. INTRODUCTION

Neurodegenerative Diseases (NDD) are growing rapidly in worldwide with a huge number of patients in the last decades. World Health Organization (WHO) predicted that NDD disorder will represent the second leading disease within 2030, which causes death globally [1]. CA is a neurological disorder that exposes symptoms, like coordination and balance issues this affects the cerebellum part of the brain, and individuals with CA frequently experience unclear speech, clumsiness, and instable gait [2]. Human gait analysis supports the detecting system to recognize the movement patterns of humans that seem uncommon and associate them with illness [3]. The aim of gait analysis in individuals who are affected by CA is to capture movement differences. These differences like postural unsteadiness or slow movements are very significant for assessing the evolution of CA [1]. Scale for Assessment and Rating of Ataxia (SARA) is a scale of clinical rating, which is used to gauge the severeness of ataxia by revealing the sign of disordered motion. This scale measures the severity of ataxia in various areas, including speech, posture, gait, and lower and

upper limbs, which limits from a score of 0 to 40, that is no ataxia to severe ataxia. This valuation scheme is affected severely by the experience of clinicians and the intrinsic subjectiveness of human-related processes [4].

Recent gait research shows that people with CA had shorter strides, slower gait speeds, shorter swinging phases, and longer supporting phases [5]. Generally, human gait analysis employed quantitative approaches to support in diagnosing NDD like CA that cause movement and postural inabilities [3]. Numerous applications exist for gait assessments are the investigation of motion disorders, including early neurology diagnosis, physical therapy, rehabilitation, and physical activity measurement. Observing patients with multiple sclerosis for ataxias diagnosis therapy and disease progression tracking in gait detection equipment are the major troubles in the ataxic gait field [6]. Studies of gait in traumatic brain injury, stroke, and other diseases have been calculated as gait detection, and are seen to be a crucial pre-operative evaluation technique in the field of pediatric orthopedic diseases. Nevertheless, the conventional movement capture approaches are based on the camera arrangements or wearable

sensor, which are suggested for gait valuation, additionally that are possess some drawbacks such as high cost, enclosure in laboratory setting, and the struggle of application on individuals [7] [3].

Neurodegenerative diseases, including CA, are rising globally and predicted by WHO to become the second leading cause of death by 2030. Early and accurate detection is urgently needed because misdiagnosis or delayed diagnosis leads to poor patient outcomes, higher healthcare costs, and significant burden on families. Current clinical assessment methods are highly subjective, vary by clinician experience, and frequently fail to detect early or subtle symptoms, denying timely intervention and rehabilitation. Machine learning (ML) with manual feature extraction has been used in the majority of CA experiments [3]. The quantity and quality of selected features have an important impact on the ML model, out of over 100 reported signal processing features. Despite in other domains, there are few deep learning applications in the CA field [8]. Together with the Decision Tree (DT), Support Vector Machine (SVM), Bayesian method, K-Nearest Neighbour (k-NN), and the two-layer Neural Network (NN) algorithms were frequently utilized in Computational and standard classification techniques. With all these techniques, the selected features significant with time, frequency, and scale domains utilize Deep Neural Networks (DNN) and Artificial Intelligence (AI) to optimize multilayer mathematical systems and their coefficients. These techniques are frequently applied to the examination of motion disorders, and analysis of body motion, and the recognition of normal kinematics in human activity [6].

To overcome the existing drawbacks, the Hybrid DRL-CTO method is proposed to accomplish the precise level of CA categorization based on the gait analysis of humans. The frame selection using adaptive quality metrics such as Signal-to-Interference-plus-Noise Ratio (SINR) Structural Similarity Index Measure (SSIM), and Peak Signal to Noise Ratio (PSNR), helps to reduce the load of computation and memory. Further, the extraction of Deep Structural Spatio Skeleton-based Pose estimated (Deep S3P) feature smakes the analysis easier, which reduces the data complexity. The key contributions of the research include:

Competitive Tuning Optimization Algorithm: The CTO algorithm is the combination of competitive characteristics of the Imperialist Competitive Algorithm (ICA) and tuning characteristics of Teaching Learning based Optimization (TLBO). The high prevalence and

disabling nature of CA, combined with diagnostic delays and low accuracy in conventional tools, necessitate development of advanced AI-driven methods for improved detection and classification.

- **Hybrid Deep Repeated Learning based Competitive Tuning Optimization Model:** The incorporation of repeated learning process in the CNN-LSTM model enables the proposed method to precisely classify the level of CA with high learning rate and improved accuracy due to the repeated working of the same input data. Additionally, the incorporation of the CTO algorithm aids in reducing the over fitting issues, avoids trapping to local optimal points, and boosts global convergence.

The research article is arranged as follows: Section 2 covers literature survey and challenges, the system model for CA classification is presents in section 3, and section 4 elaborates the methodology of the Hybrid DRL-CTO method, the result sreachdover the proposed research is explains in Section 5, and the research ends with the conclusion and future works in section 6

2. LITERATURE SURVEY

Existing approaches in the research of CA detection and classification schemes are explained in this section with its advantages as well as limitations.

Stoean, C. et al. [9] suggested a method, that employed Monte Carlo Dropout within the deep learning (DL-MCD-DT) to repeatedly differentiate presymptomatic symbols of type 2 of CA in saccadic trials, which is attained from electrocardiograms. The suggested model aided in offering graphic explainability of the topological model. The suggested model is more robust and trains the decision tree model with additional information, however, it failed to resolve the complex issue of categorizing the registers precisely, particularly in identifying the presymptomatic characteristic. Shanmuga Sundari, M., and Jadala, V.C. [10] developed a model based on machine learning for Analysis of Gait (AoG) prediction by the occurrence of worst gait patterns before AoG. The weights of the Recurrent Neural Network (RNN) and NN were flawlessly calibrated by the developed models to progress the capability of the system to predict diseases. The personalized labelling and weak gait forms were used to forecast AoG rapidly and precisely. The main benefits of the developed model were systems accessibility and cost efficiency; however, low sensitivity was the

drawback of the suggested method, which was caused by the inherent variances across patients.

Procházka, A. et al. [6] suggested an optimized deep learning convolutional neural network (ODLCNN) system to differentiate normal and ataxic gait using accelerometric information. The suggested method was based on the occurrence component analysis of signals and the body locations were recorded with 60Hz of frequency. The lack of feature specifications allowed the suggested method to simplify the entire classification process, however, its accuracy was affected by various extra components of the signal. Ngo, T., et al. [8] suggested an image transformation-based federated learning method to diagnose CA. The federated learning provides privacy defense for the participant hospitals for execution and the suggested method saved the time of deployment and analysis by removing the feature extraction of laborious work. The suggested altered recurrence plot reduced the stored data; thus, the processing time was decreased, however, due to the presence of composite and heavier network like DenseNet in the suggested model, it failed to produce high accuracy.

Rahman, W. et al. [11] developed a ML model to detect the person with the characteristics of ataxia-impacted gait, and the rigidity of CA was measured from the gait of SARA. The participants were separated from their backgrounds and created various features by a developed method to capture characteristics of gait like speed, stability, and step width. The dataset was gathered on individuals either identified or at-danger of CA, therefore the outcomes failed to translate to another kind of ataxia. Dostál, O., et al. [15] presented a method to assess and categorize accelerometric information attained over the gait. Frequency domain features were utilized to categorize the people with the syndrome of neurology and the discrimination capability of sensors on the position of the body. The suggested methodology was comparatively simple and reasonable for discriminating the patients of ataxic from controls, however, it failed to look into concerns with mobility and development.

Seetharama, P.D. and Math, S. [16] introduced a weight feature optimization based extreme gradient boost classifier (FWO-XGB) to categorize ataxic individuals and healthy individuals through the features of gait. FWO-XGB method was effective in decreasing misclassification using the misclassification-aware weight optimization method. The effectiveness of the weight

optimization technique guaranteed improved specificity and sensitivity, however, FWO-XGB failed to test on the dataset of multi-label classification. Seetharama, P.D. and Math, S. [6] suggested feature selection and ranking-based extreme gradient boost (FSR-XGB) methods to categorize the severity of ataxia. An FSR-XGB method was effective in differentiating various levels of ataxia's severity, this is owing to the adoption of the original selection of features and grading method through the method of two-level cross-authentication.

Hypothesis: A hybrid deep repeated learning framework with competitive tuning optimization, leveraging discriminative gait features, will achieve more accurate and efficient cerebellar ataxia detection/classification than existing models.

Scope: This paper presents a methodology for CA detection and classification using gait analysis and the READISCA dataset only. It does not address other neurodegenerative diseases, non-gait-based assessments, nor generalized movement disorder diagnosis.

3. CHALLENGES

The challenges originate from the existing approaches of CA classification are explained below,

- In [6], the suggested method had no feature specification, which made the whole classification process simple, however, the gait analysis of wearable sensors had affected the accuracy levels of signal processing in the limited spatial area.
- The suggested model [4] was more convenient and low-cost effective, however, it possessed low sensitivity, which was caused due to the characteristics variance across the patients.
- In [10], the suggested method reduced workload and prevented the data leakage, however, the absence of matrix transformation techniques affected the base model and resulted in poor final prediction and performance.
- The suggested method [15] is moderately simple and cost-effective for differentiating ataxic patients, but the dataset was trained with less amount of data and focused on fewer neurological problems.

The FWO-XGB method [16] reduced the false positive rate and worked well with the imbalanced

data, however, it failed to validate the FWO-XGB method by testing it on the multiple label classification dataset

4. PROBLEM STATEMENT

Impairment function of the cerebellum causes the disease CA, which exposes symptoms like uncoordinated movement and balance issues. If CA is left untreated, it leads to several complications including, psychological illness, and some forms of ataxia can lead to death. Numerous methods are developed to detect CA accurately, yet they are ended up with certain drawbacks. In the existing model, there is no stacking of two or more distinct models with other pattern matrix transformation techniques. A research methodology is designed to produce accurate results in detecting cerebral ataxia by solving the disadvantages of existing approaches like low model performance, low sensitivity, poor final detection, high complexity, and low accuracy. Here, the proposed method used the dataset of 150 videos with 6 seconds to detect CA. Each video contains m number of frames, which is represented as,

$$W = \left\{ \sum w \right\} \quad (1)$$

Where W denotes the video data, w represents the j frame of the video. Within the video, the frames are selected to focus on the particular moment in the video by the frame selection measure, such as PSNR, SSIM, and SINR. The selected frames from one video are represented as,

$$W = \left\{ \sum wf \right\} \quad (2)$$

Where, W is the selected frames from one video, wf indicate the j selected frame, n denotes the total amount of frames in one video. The selected frames from all the videos are represented as,

$$S = \left\{ \sum W \right\} \quad (3)$$

Here S signifies the total amount of selected frames from all the input videos, k indicates the total amount of videos. The selected frames (images) may be in poor quality and contrast, which may distress the model's accuracy. Hence, these issues are overcome in the pre-processing phase

and image quality is improved. The pre-processed data is denoted as wf . Pre-Processed image is given as input to the Region-of-Interest (ROI) extraction phase, which decreases the complexity of the data. Then the ROI extracted image undergoes the feature extraction process for identification and representation of distinctive within the image. The extracted features are given as input to the proposed model for predicting the classes of CA, which is represented as

$$X = \begin{cases} 0; \text{if } "wf" \text{ is normal} \\ 1; \text{else if } "wf" \text{ is abnormal with slight CA} \\ 2; \text{else if } "wf" \text{ is abnormal with moderate CA} \\ 3; \text{otherwise } "wf" \text{ is abnormal with advanced CA} \end{cases} \quad (4)$$

where X represents the type of CA classes. This is done through the model by reducing the error using the loss of categorical cross-entropy measure. Then, the categorical cross-entropy ρ is calculated as

$$\rho = - \sum x \cdot \log(x) \quad (5)$$

Here x represents the true probability distribution, x indicates the predicted probability distribution. Thus, the proposed method attains high effectiveness in CA detection and classification.

5. SYSTEM MODEL

The CA classification system is significant because it aids in determining the exact diagnosis and treatment. The input videos are gathered and key frames are selected to avoid unnecessary information. After extracting the features, it is given to the DL model, which is capable of identifying the patterns in the image and providing accurate classified output. The outcome is assessed using the evaluation measures, such as recall, precision, and accuracy. This system helps medical experts to make precise and timely diagnoses. Figure 1 shows the architecture of system model for classifying CA.

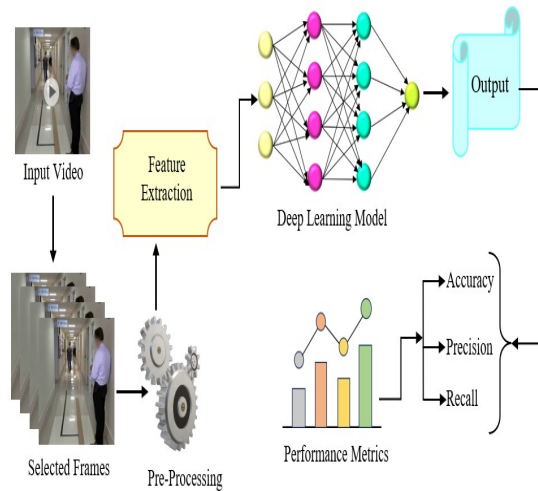


Figure 1: System Model for Cerebellar Ataxia

Proposed Hybrid Deep Repeated Learning based Competitive Tuning Optimization Model for Cerebellar Ataxia Detection and Classification

The primary intention of the research is to detect and classify the level of CA based on the human gait analysis. The proposed Hybrid DRL-CTO is composed of different stages, like input video collection, frame selection, pre-processing and ROI extraction, features extraction, and model training phase. At first, the input video is collected from the dataset, then it is further processed to select the frames using adaptive quality metrics such as PSNR, SSIM, and SINR to perform the detection mechanism more optimally. Thereafter, pre-processing and ROI extraction mechanisms are carried out to increase the quality of video frames by improving the brightness and contrast of the image. The pre-processed and ROI-extracted frames are fed to the features extraction module, where features, like Skeleton-based gait features, pose estimation, kinematic Spatio temporal features, and deep flow map-based structural features are effectively extracted. After extracting the features, a repeated learning-based CNN-LSTM model is used to detect the behaviors in such a way that training of this model is performed using the CTO algorithm that is obtained through the integration of ICA and TLBO. Finally, the proposed classifier classifies CA into four different classes such as normal, abnormal with slight CA, abnormal with moderate CA, and abnormal with advanced CA. The diagrammatic illustration of the proposed framework is shown in Figure 2.

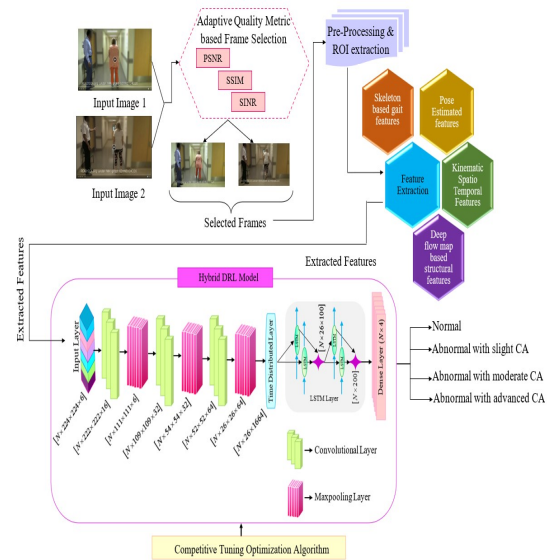


Figure 2: Schematic illustration of proposed Hybrid DRL-CTO Approach

This research adopts the following design:

1. Collection of 150 videos from the READISCA dataset.
2. Adaptive frame selection using PSNR, SSIM, and SINR
3. Pre-processing and ROI extraction to enhance image quality.
4. Feature extraction via Deep S3P (skeleton-based gait, pose estimation, kinematic spatio-temporal, deep flow-based features).
5. Training a repeated learning CNN-LSTM model.
6. Parameter tuning with Competitive Tuning Optimization (CTO).
7. Evaluation using k-fold validation, with accuracy, precision, and recall metrics reported.

5.1 Input Video Collection and Frame Selection

Initially, input videos are collected from the source, which comprises 150 videos and each video contains 6-second clips. Each video contains m number of frames, which can be mathematically represented as,

$$W = \{w, w, \dots, w, \dots, w\} \quad (6)$$

Where, w represents the j frame of the video, W indicates the video data. The frames are selected from the collected video using the adaptive quality metric-based frame selection. The quality metrics for selecting frames from the videos are PSNR, SSIM, and SINR. The main aim of the

frame selection process is to adaptively select the frames from the input video clips for the reconstruction process depending on the movement happening among two successive frames. Selecting frames can reduce the computational load, memory, and processing time required to process frames. PSNR depends on the approximation of signal-to-noise proportion. The PSNR rate proceeds towards infinity as the Mean Squared Error (MSE) rate proceeds towards zero. This displays that a smaller value of PSNR indicates more arithmetical variations among the frames. The relation among the two frames is calculated in decibel form. The decibel scale's logarithm term is used to compute the PSNR value because of the extensive dynamic range of the signals. This dynamic range differs among the biggest and least values that are unstable by its quality. The PSNR value is mathematically represented as,

$$PSNR = 20 \times \log \left(\frac{MaxPixel}{\sqrt{MSE}} \right) \quad (7)$$

Where, MSE is the value, which used to compute the variation among two successive frames. The SSIM is one of the familiar quality measures, which is utilized to calculate the similarity among two frames. It is considered to be interrelated with the perception of quality of the Human Visual System. The SSIM is considered a replacement for the outdated fault summation approach, which is designed using image distortion modeling as per the incorporation of three aspects such as luminance distortion, loss of correlation, and contrast distortion. The SSIM is mathematically represented as,

$$SSIM(w, w) = \frac{(2\lambda\lambda + r)(2\mu + r)}{(\lambda + \lambda + r)(\mu + \mu + r)} \quad (8)$$

Where, λ and λ represents the average of w and w likewise, μ and μ indicates the variance of w and w respectively, μ denotes the covariance of both w and w . Here, the term w and w are the two consecutive frames of the video W . r and r are the variables, which utilize to balance the division with feeble denominator. SINR is also used to measure the superiority of frames by associating the power of the desired frames with the power of undesirable noise and interference. The SINR is mathematically represented as,

$$SINR = \frac{P}{I + G} \quad (9)$$

Where, P is the power of the i frame, I denotes the power of interference of another a frame, G represents the Gaussian noise. Each frame is compared with the consecutive frames and compute the average value of the three adaptive quality metrics, which is represented as,

$$Average = \frac{PSNR + SSIM(w, w) + SINR}{3} \quad (10)$$

Where, $Average$ denotes the average value of the three adaptive quality metrics. The frames with the minimum values are considered to have high diversity, so the frames with low comparison values are selected, and that is considered as the selected frames. The selected frames from one video are represented as,

$$W = \{wf, wf, \dots, wf, \dots, wf\} \quad (11)$$

Where, W is the selected frames from one video, wf indicate the j selected frame, n represents the total number of frames from one video. The selected frames from all the videos are represented as,

$$S = \{W, W, \dots, W\} \quad (12)$$

Where S represents the total number of selected frames from all the input videos, k indicates the total number of videos. Each selected frame has a dimension of $[1 \times 1920 \times 1080 \times 3]$

5.2 Pre-Processing and ROI Extraction of Selected Frames

The original selected frames (images) are not suitable for the further process, due to the occurrence of various noise, low quality, and contrast in the selected frames that may distress the model's accuracy. So, the pre-processing phase is significant to improve the quality of the selected images. Pre-processing with ROI extraction reduces model's complexity. The pre-processing stage comprises noise removal, contrast quality, and enhancement. In contrast enhancement, it upgrades the visual quality by improving the brightness and contrast of the image. This can pave the way to improve the effectiveness of CA

detection and classification. The resultant pre-processed image is represented as,

$$P = wf \quad (13)$$

Where, P represents the pre-processed image. The appropriate ROI extraction enhances the contradiction by finding the significant regions and eradicating the unnecessary regions from the pre-processed image.

In ROI extraction, initially, the morphological function is applied to eliminate the imperfection, which mostly affects the texture and structure of the images. Morphology functions operate based on set theory and depend on pixel ordering relation instead of its numerical implications. It is very beneficial in ROI extraction because it directly deals with shape extraction of the input images. Then, find the contour region, finding contours is similar to discovering front objects from the background. It can be described just by joining the curve in all constant points, which have identical intensity or color. After finding the contour region, the size of the image is resized from the dimension of $[1 \times 1920 \times 1080 \times 3]$ to $[N \times 224 \times 224 \times 3]$. The image resizing has the advantages of two-fold, that is each images are considered as the same size and the computation complexity is diminished efficiently. Finally, the ROI extracted images are represented as,

$$G = ROI(P) \quad (14)$$

Where the ROI extracted image is denoted as G . Then, the ROI extracted image is given to the feature extraction process with the dimension of $[N \times 224 \times 224 \times 3]$ for further extraction.

5.3 Feature Extraction based on Deep Structural Spatio Skeleton-based Pose Estimated Features (Deep S3P)

Feature extraction is the process of extracting the relevant features to create the feature vector from the ROI-extracted image. The process of feature extraction is accomplished through Deep Structural Spatio Skeleton based Pose estimated (Deep S3P) Features, which include four various types of features, namely skeleton-based gait features, pose estimation, kinematic Spatio temporal features, deep flow-based structural features. The Deep S3P Features aid the classifier to easily classify among various classes of CA and

extracting Deep S3P Features improves the efficacy and performance of the Hybrid DRL-CTO model, also it makes the analysis process easier and reduces the complexity of data. Let feature extraction process input be G with the dimension of $[1 \times 224 \times 224 \times 3]$.

Skeleton-based gait features are the type of features, which represent the walking pattern of the person by examining their skeleton data. These types of features are extracted using the pre-trained neural network model, like Open Pose and the Deep-CNN model. Open Pose offers a bottom-up method to assess the real-time gestures of various people without any requirement of character finders and enhance the processing speed. The ROI-extracted image G is given as input to the pre-trained Open Pose model, then the output of the Open Pose model is given to the Deep-CNN to extract the skeleton-based gait features. Deep-CNN is developed by heaping various layers of convolution and pooling to form a deep structure. The low-level features are extracted using the first convolutional layer, though the composite features are extracted based on the extra convolutional layer. The skeleton-based gait features are represented as,

$$E = \sum G \otimes K + b \quad (15)$$

Where E represents the extracted skeleton-based gait features with the dimension of $[1 \times 224 \times 224]$, G denotes the input feature map, which is the ROI extracted image, K is the convolutional filter, b indicates the bias, both convolutional filter and bias are the learnable parameter in Deep CNN, and $\otimes$ represents the 2-dimensional convolution operation.

Pose estimated features are extracted using Open Pose, a pre-trained neural network model. It traces a person's body, face, joints of hands, and legs to estimate the pose of the person. The existing pose estimation approach used a top-down method to estimate the pose, which failed to estimate the correct pose. Here, Open Pose uses a bottom-up approach that initially identifies the body part of each individual in the image and then accumulates the joints depending on the linkable pairs among the parts of the body, this enables the model to estimate the pose features effectively. The dimension of extracted pose estimated features is $[1 \times 224 \times 224 \times 3]$ and it is represented as E .

Kinematic Spatio-Temporal features are extracted using the Kinematic Spatio-Temporal model, which is based on MLP, that are recurrently applied across the axes of temporal or spatial. The

input of this model is G and the input vector is processed and transposed through the block of MLP. The block of the second MLP is a patch-wise feature extraction block. Each block consists of an activation function and two fully connected layers. It extracts the features in the multiple blocks and the features are united with the feature of insertion length. Combined features are processed by Fully connected (FC) layers to produce accurate Kinematic Spatiotemporal features. The operation for extracting these features is represented as,

$$V = G + h\phi(hN(G)) \quad (16)$$

$$E = V + h\phi(hN(V)) \quad (17)$$

Where, h, h, h, h are the parameter of fully connected layer, V indicates the data of insertion length, N represents the function of normalization, ϕ signifies the GELU activation function, E denotes the kinematic spatio-temporal features with the dimension of $[1 \times 224 \times 224]$.

The pre-trained ResNet-101 model is a residual network with deep 101 layers. The ROI extracted image G is given as input with the dimension of is $[1 \times 224 \times 224 \times 3]$ to the pre-trained ResNet-101 model. Then the processed output is given to the Local Ternary Pattern (LTP) to extract the deep flow map-based structural features. LTP works by evaluating the relationship among pixel's intensity rate within the neighborhood of an image. It considers the intensity of the neighbour to be greater or smaller or equal than the central pixels. LTP converts the relationship into a code of ternary, that is utilized to represent the characteristics of the local texture. It concentrates on the relationships of relative intensity rather than the absolute values which makes the LTP robust. Moreover, LTP can able to capture the rough and well texture effectively. The extracted deep flow map-based structural features are represented as,

$$E = \sum I(Q - Q).3 \quad (18)$$

$$I = \begin{cases} 1, & \text{if } Q \geq th \\ 0, & \text{if } Q = th \\ -1, & \text{if } Q \leq th \end{cases} \quad (19)$$

where Q represents the neighbor pixel's intensity, th denotes the threshold value, Q indicates the current position pixel's intensity, I is the gradient size, and E represents extracted deep flow map-based structural features with the dimension of $[1 \times 224 \times 224]$. Finally, the extracted Deep S3P features are concatenated to gain the final feature vector α with the dimension of $[1 \times 224 \times 224 \times 6]$, which is represented as,

$$\alpha = \{E \| E \| E \| E\} \quad (20)$$

Where α represents the extracted final feature vector, and it is further fed into the Hybrid DRL model for accurate CA detection and classification.

5.4 Repeated Learning based Convolutional Neural Network-Long Short-Term Memory Model for Cerebellar Ataxia Detection

The Hybrid Deep Repeated Learning (Hybrid DRL) model is the fusion of both the Convolutional Neural Network (CNN) and Long Short-Term Memory (LSTM) Model with a repeated learning process, which is proposed to accomplish an exactCA classification and detection by resolving the issues of vanishing gradient, long-term dependencies. The Hybrid DRL model can acquire the features from both time and spatial dimensions and learn long-range dependencies among time steps and process series of data by iterating beyond time steps. The learning rate of the proposed Hybrid DRL model is improved by the process of repeated learning, which repeatedly learns with the data. The detection accuracy of the Hybrid DRL model is improved due to the incorporation of repeated learning in the model, it works repeatedly for 3 times for the same data, so the accuracy for the detection of CA is high. Additionally, the Hybrid DRL model is more flexible and computationally efficient and it also improved the performance of memory. Figure 3 displays the architecture of the hybrid DRL model.

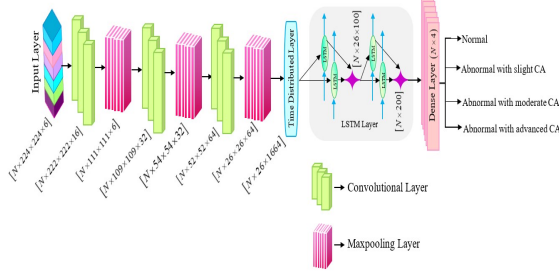


Figure 3: Architecture of Hybrid DRL model

The Hybrid DRL model consists of convolutional blocks and LSTM layers. Initially, the input data α , with the dimension $[1 \times 224 \times 224 \times 6]$ is fed into the convolutional layer. The convolutional process that takes place in the convolutional layer, which is represented as,

$$J = \sum \alpha \otimes K + b \quad (21)$$

Where, J indicates the output feature map, α represents the input feature vector, b indicates the bias, K denotes the convolutional filter, and $\otimes$ represents the 2-dimensional convolution operation. Every convolutional layer is accompanied by a pooling layer to diminish the dimension of the features and eliminate the repeated features to avoid overfitting issues additionally, it decreases the amount of training constraints, which allows the model to concentrate on the prominent feature and diminish the issues of computational complexity. The output of max pooling layer is represented as,

$$p = \max(J) \quad (22)$$

The output of the convolutional block P , with the dimension of, $[1 \times 26 \times 26 \times 64]$ is fed into the Time Distributed Layer (TDL), which is used to change the dimensions of the data. The TDL can convert a layer to a sequential slice of the input, that obtains more historical period steps of time sequence data and gains the long-range features of input variables. The TDL changes the dimension of the output from the convolutional block to $[1 \times 26 \times 1664]$, which is represented as P . The output of TDL P is given to the LSTM block of the Hybrid DRL model to detect and classify CA. It comprises three mechanisms such as output gate, a forget gate, and an input gate. The output of time distributed layer is given as input to

the forget gate. The outcome of the forget gate is computed using Eqn (23), which is represented as.

$$O = \sigma(v \cdot [hd, p] + b) \quad (23)$$

Where O indicates the outcome of the forget gate, v represents the forget gate's weight, p denotes the input, which is the output generated from the time distributed layer, σ signifies the sigmoid function, hd represents the preceding hidden state, and b denotes the bias of the forget gate. The input gate and candidate input gate are computed using Eqn. (24) and (25), which are represented as,

$$O = \sigma(v \cdot [hd, p] + b) \quad (24)$$

$$\tilde{O} = \tanh(v \cdot [hd, p] + b) \quad (25)$$

Where, O and $\tilde{O}$ indicates the output of the input gate and the candidate input gate respectively, $\tanh$ is the tanh activation function, v denotes the input gate's weight, candidate input gate's weight is represented as v , and the bias of both input and candidate input gate is denoted as b and b respectively. The state of the current cell is updated using Eqn. (26), which is represented as,

$$O = O \times O + O \times \tilde{O} \quad (26)$$

Where, O represents the state of the current cell. The outcome of the Output gate is computed using Eqn (27), which is represented as.

$$O = \sigma(v \cdot [hd, p] + b) \quad (27)$$

Where O indicates the outcome of the output gate, v represents the output gate's weight, and b denotes the bias of the forget gate. The output of the LSTM layer of the Hybrid DRL model is obtained by calculating the output of the state of the current cell and the output of the output gate, which is represented as,

$$Z = O \times \tanh(O) \quad (28)$$

Where, Z indicates the output of the LSTM layer, which is given to the layer of dropout, that aids in reducing the complexity. Followed by a layer of dropout, the dense layer is established in

the Hybrid DRL model which offers the data acquired from the previous layer and effectively detects and classifies the CA in four diverse classes such as normal, abnormal with slight CA, abnormal with moderate CA and abnormal with advanced CA. The parameters such as weights and biases of the Hybrid DRL model are tuned by the Competitive Tuning Optimization (CTO) Algorithm, which is described elaborately in the below section.

5.5 Competitive Tuning Optimization Algorithm

CTO algorithm is obtained by merging the competitive characteristics of the Imperialist Competitive Algorithm (ICA) [15] and tuning characteristics of Teaching learning based optimization (TLBO) [16]. Strong competing strategy and the effect of the influence of an educator make the proposed CTO algorithm more ideal and feasible. A good educator creates a good means for the outcomes of the students and the students also study from communication among themselves, which aids in their grades and the teacher's quality affects the results of the learners. Likewise, the competition of imperialistic will progressively result in the power of the powerful empires increasing and the power of the weaker ones decreasing. The proposed CTO algorithm overcomes limitations such as low convergence speed, high computational complexity, and sensitivity to local optima, and attained the best global solutions for constant non-linear function with low computational effort and high reliability.

Motivation: The proposed CTO algorithm is developed to expand the search space and improve the capability of global search based on assuring the necessary convergence speed, also the proposed CTO algorithm reveals powerful competitiveness in terms of accuracy of solution, robustness, and convergence speed and avoids trapping to local optimal points and boosts global convergence. Furthermore, the CTO algorithm can offer high performance in resolving issues of high complexity in small and enormous dimensions and achieve effective results.

Inspiration: The imperialistic competition as well as the teaching-learning procedure in the classroom are inspired to use CTO algorithm. It is based on the behavior of socio-political and results of the influence of the educator on the learner's results in the classroom. The teacher is mostly considered an extremely educated person, who distributes their knowledge with the students. It is

apparent that an excellent teacher teaches learners so that they can able to get better outcomes in terms of their grades or marks. The imperialistic competition initiates between all the competitive groups. Any competitive group that is not able to learn or triumph in the competition is collected and stored for exploiting the experiences. These hybrid characteristics in the optimization aid in attaining optimum results, that enable to achieve of accurate detection and classification of CA. The steps of the CTO algorithm are described as follows,

5.5.1 Initialization

Let us assume the initial solution of the CTO algorithm H , which is mathematically represented as,

$$H = \{H, \dots, H, \dots, H\} \quad (29)$$

Here, $H \in (b, v)$, b and v represents the weight and bias of the model, each solution H represents a competitive group and H denotes the total number of solutions. The significance of each group supporting the global convergence, that is based on the learning factor of each group. The competitive group with a high learning factor is capable of learning from the real environment and experience and it always retains the high learning factor.

On the other hand, the competitive group with a lower learning factor exhibits poor learning capacity, which needs a constant teaching environment. It is significant to note that the poorly learned competitive group from the teaching environment is reported as the worst learner and the competitive group is stated as the worst solution. Generally, the standard competitive algorithm ignores the worst solution, whereas, in the proposed CTO algorithm, the worst groups are gathered and stored for exploiting the experiences during the learning scenarios

5.5.2 Fitness Evaluation

The fitness function of the solution is measured based on accuracy, which is represented as in the Equation. (30).

$$Ay = \left(\frac{t + t}{t + t + f + f} \right) \quad (30)$$

Where t signifies the true positive rate, f shows the rate of false positive, f is the false negative rate, and t indicates the value of true

negative. Ay indicates the fitness measure, which is computed based on accuracy, by which the best solution is considered by the maximum fitness value.

5.5.3 Evaluation of Learning Factor

For all M solutions, the learning factor T is calculated by using the Eqn 30, which is represented as

$$T = \left(\frac{Ay(H) + Ay(H^-)}{2} \right) \quad (31)$$

Where H denotes the current solution of the competitive group in the current iteration, H^- denotes the current solution of the competitive group in the previous iteration, and Ay indicates the fitness of the learning model. At the first iteration, a common best solution H^- is declared.

5.5.4 Learning and Teaching Phase:

The obtained solutions in the phases of learning and teaching are renewed and updated until t .

Case (i)-Exploitation Phase: If $A \leq 1$, when the fitness of the model A is greater than equal to 1, the CTO algorithm passes into the exploitation phase.

Sub Case (i)- Global Learner Phase: If $T(H) \geq T \leq T^-$, when the $T(H)$ learning factor of the current solution of the competitive group in the new iteration is greater than or equal to the global learning factor, T^- and less than or equal to the learning threshold T , then the global learner phase is employed for the solution renewal.

The competitive group in the global learner phase exhibits optimal learning performance. The influence of the teacher in this phase is found optimal and the learner groups exhibit linear learning experiences, which not only contribute towards the global point but also bring a quality H^- , that boosts the accuracy of the classifier. The global learner's solution update is represented as follows,

$$R = u - \frac{t \times (u - u^-)}{t} \quad (32)$$

$$R = u + (u - u^-) \times \left(\frac{t - t^-}{t - 1} \right) \quad (33)$$

$$R = \sum C(H) \times \|H^- - H\| \quad (34)$$

$$C(H) = \begin{cases} 1 & ; \text{otherwise} \\ 0 & ; \|H^- - H\| \leq T \end{cases} \quad (35)$$

$$R = \sum U(H) \times \|H^- - H\| \quad (36)$$

$$U(H) = \begin{cases} 1 & ; \|H^- - H\| \leq T \\ 0 & ; \text{otherwise} \end{cases} \quad (37)$$

$$R = q - (q - q^-) \cdot \cos(e - 1) \quad (38)$$

$$(Hs - Ht)(1 + 4R) + R \cdot \text{round}[1 - \text{round}(Qij)] \quad (39)$$

Where R, R, R , and R represents the random value, u, u^- are the inertia factors, which avoids trapping to local optimal points and boosts global convergence, and also it changes linearly with the learning ability of the learner. $C(H)$ represents the global factor of the L learner, and $U(H)$ is the local factor of the L learner. q represents the initial learning value and q^- denotes the minimum learning value. H indicates the current global solution of the competitive group at the current iteration, H^- is the previous solution of the competitive group at the current iteration, H^- denotes the best solution of the competitive group in the current iteration, B represents the interactive experience, and R signifies the interactive degree. $\text{rand}(0,1)$ is the random numeral, which is dispersed consistently between (0,1). If a new learner is produced with an improved fitness value, it will exchange the current one. This phase leads the population to speeden the convergence speed.

Sub Case (ii)- Personal Learner Phase: If $T(H) \geq T \geq T^-$, when the $T(H)$ learning factor of the current solution of competitive group in the next iteration is greater than or equal to the global learning factor, T^- and greater than or equal to the learning threshold T , the personal learner phase occurs.

In the global learner phase, group learning and group performance are significant, where the individual learners interact with each other,

supporting knowledge sharing between the individual learners and supporting overall group performance. The personal learner's solution updation is represented as,

$$H_{\text{new}} = H_{\text{old}} + R (H_{\text{best}} - H_{\text{old}}) + R (H_{\text{best}} - H_{\text{old}}) \quad (40)$$

Where, H_{new} represents the updated solution of the global learner, R and R indicates random value, and $R, R \in (0,1)$. The Learning threshold is represented as,

$$T = \frac{T(H_{\text{new}}) + T(H_{\text{old}}) + T(H_{\text{best}}) + T(H_{\text{best}}) + \dots + T(H_{\text{best}})}{t} \quad (41)$$

Where, $T(H_{\text{new}})$ represents the learning factor of the competitive group's solution in the current iteration, $T(H_{\text{old}})$, $T(H_{\text{best}})$, and $T(H_{\text{best}})$ indicates the learning factor of the competitive group's solution in the previous iterations, and $T(H_{\text{best}})$ represents the learning factor of the competitive group's solution in the maximum iteration. The learning threshold evaluates the linear progression of competitive groups, which ensures the continuous learning ability of groups in successive iterations. The random factor is represented as,

$$\text{Ran}() = \frac{T(H_{\text{new}}) + T(H_{\text{old}})}{2} \quad (42)$$

Where $\text{Ran}()$ represents the random factor, H_{new} denotes the solution of the global competitive group, H_{old} represents the solution of the local competitive group. The random factor can be obtained by taking the average of the solution of both the global group and the subgroup.

In the global learner phase, every learner exhibits good potential, which is acquired through teaching, while weak learners remain hidden due to the overall group learning. Likewise, in the personal learner phase, all weak learners come into the picture, and there is a necessity to learn individually from the teachers. Here, knowledge sharing between individuals is not identified while individual learners utilize the teacher's ability to update their self-knowledge.

Sub Case (iii)- Weak Learner Phase: If $T(H_{\text{new}}) \leq T$, when the $T(H_{\text{new}})$ learning factor of the current solution of the competitive group in the

next iteration is less than equal to the global learning factor, the weak learner phase occurs.

Through proper teaching, individual learning and group learning putsome learners in no-learn mode, where the teacher seeks special attention and special teaching scenarios for the weak learners. The teacher's efforts to distribute knowledge between the weak and best learners, and raise the intelligence level of the entire class, which aids the learner to score good grades or marks. Thus, according to the learner's ability, the teacher improves the mean of the class. The weak learner's solution updation is represented as,

$$H_{\text{new}} = H_{\text{old}} + R (H_{\text{best}} - D \times H_{\text{old}}) + R (H_{\text{best}} - H_{\text{old}}) + R (H_{\text{best}} - H_{\text{old}}) \quad (43)$$

$$R = t \left(1 - \frac{\max(T(H_{\text{new}}), T(H_{\text{old}}))}{\max[T(H_{\text{best}})]} \right) \quad (44)$$

$$R = t \left(1 - \frac{\max(T(H_{\text{new}}), T(H_{\text{old}}))}{\max[T(H_{\text{best}})]} \right) \quad (45)$$

$$R = t \left(1 - \frac{\max(T(H_{\text{new}}), T(H_{\text{old}}))}{\max[T(H_{\text{best}})]} \right) \quad (46)$$

Where D represents the teaching factor, H_{new} indicates the solution of weak learner group in the current iteration, H_{old} denotes the average learner knowledge, H_{best} is the common best solution. The average of learner knowledge acquired through group learning and individual learning is represented as,

$$H_{\text{new}} = \left(\frac{H_{\text{new}} + H_{\text{old}}}{2} \right) - R (H_{\text{new}} - H_{\text{old}}) - R (H_{\text{new}} - H_{\text{old}}) \quad (47)$$

$$R = q - (q - q_{\text{min}}) \cdot \cos(e - 1) \quad (48)$$

$$R = y - \frac{(y - y_{\text{min}}) \times [T(H_{\text{new}}) - T(H_{\text{old}})]}{T(H_{\text{best}}) - T(H_{\text{old}})} \quad (49)$$

Where, R and R are the random values, H_{new} represent the updated solution of the global learner's group, H_{old} indicates the updated solution of the global learner's group in the next iteration.

$y, y_{\min}$ represents the maximum and minimum weight value, $q_{\max}$ is the maximum learning value, and the initial and maximum values lie between 0 to 2. In $R, \cos(e^{2-u} - 1)$ changes in cos form and explains the self-learning nature of the learners.

Case (ii)-Exploration Phase: If $A > 1$, when the fitness of the model A is greater than 1, the CTO algorithm passes into the exploration phase.

Sub Case (i)- Adaptive Learning Phase: If $T(H) \geq \text{Ran}()$, when the $T(H)$ learning factor of the current solution of the competitive group in the next iteration is greater than equal to the random factor $\text{Ran}()$, the adaptive learning phase occurs.

The learning finds the scheduled learning through observation and prolonged experiences and there no specific teacher is present. The adaptive learning solution updation is represented as,

$$H_i = H_i + 21[(D+1+\beta)H_i + (D+\beta(1-\beta)(3-\beta))H_t] \quad (50)$$

$$\beta = t \times \left(1 - \frac{t}{T}\right) \times (H) \quad (51)$$

Where H indicates the solution of random learner and outside the search boundary, β represents the adaptive fusion degree, H^+ , H^- , and H are denotes the solution of competitive group in the previous iteration, D is the teaching factor, H represents the updated solution of the adaptive learning.

Sub Case (ii)- Random Learning Phase: If $T(H) < \text{Ran}()$, when the $T(H)$ learning factor of the current solution of the competitive group in the next iteration is less than the random factor $\text{Ran}()$, a random learning phase occurs.

The random learning of the learners is exhibited in this phase. Any learner can communicate with another learner randomly to transfer the acquired knowledge from the teaching mechanism. The solution updation of random learning is represented as,

$$H = H + \text{rand} \times (H - H) \quad (52)$$

Where, rand is the random numeral, which is dispersed consistently between (0,1),

H represents the solution of the group with the random learner, and H denotes the current solution of the competitive group in the current iteration. The last updated solution is stated as the best solution. The CTO algorithm improves the performance of the model by tuning the parameters of the classifier for the detection and classification of CA. Figure 4. Shows the overall flowchart of the CTO algorithm.

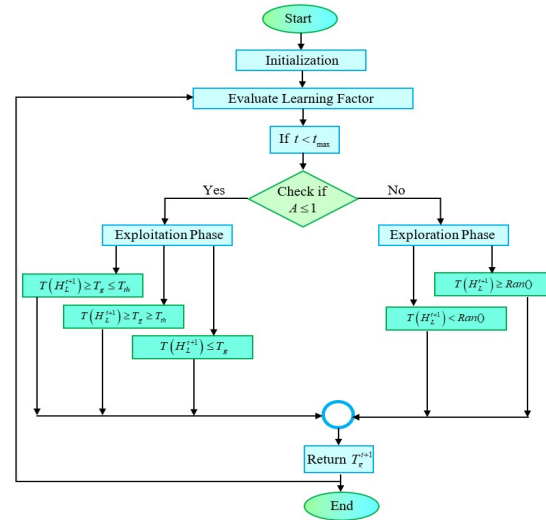


Figure 4: Flowchart of CTO Algorithm

6. RESULTS AND DISCUSSION

This section describes the experimental outcomes attained through the proposed Hybrid DRL- CTO model over the analysis of comparative and performance.

6.1 Experimental Setup

The complete experimentation is performed in the PYTHON tool through Windows 10 OS, an Intel Core processor, and 16 GB storage of RAM.

6.2 Dataset Description

The data is collected from the READISCA (NIH-funded international clinical trial readiness study). It has the critical data to correctly design the medical trails and correctly understand the outcomes in early ataxic carriers pre-ataxic of SCA1 and SCA3 mutations.

6.3 Evaluation Metrics

The metrics of evaluation, which is used to assess the Hybrid DRL-based CTO model are accuracy, recall, and precision.

Accuracy: It computes the percentage of accurate prediction between complete predictions

formed by the model and it is scientifically shown in Equation (30),

Recall: It is demarcated as the number of true positives divided through the whole amount of elements that truly fit the class of positives. Recall is mathematically represented as,

$$recall = \frac{t}{t + f} \quad (53)$$

Where, t represents the true positive value, and f denotes the value of false negative.

Precision: The amount of true positives divided through the whole amount of elements that truly fit the class of positives, that is summation of true positives and false positives. Precision is mathematically represented as,

$$precision = \frac{t}{t + f} \quad (54)$$

6.4 Experimental Analysis

The experimental analysis section includes the results of each step in CA detection and classification using the proposed Hybrid DRL-based CTO model, which is shown in Figure 5.

Input Image					
Pre-Processed Image					
Extracted Features	Skeleton based gait features				
	Pose Estimated Features				
	Kinematic Spatio Temporal Features				
	Deep Flow map based structural Features				
Classified Output					

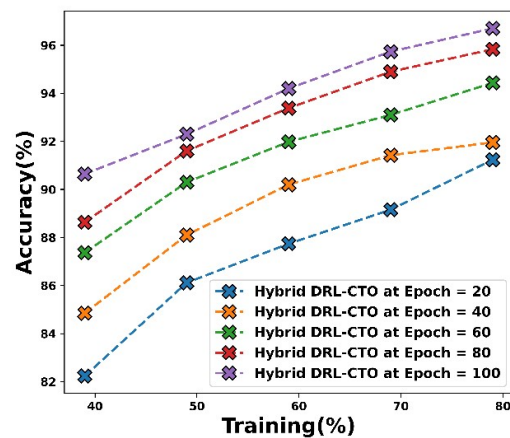
Figure 5: Experimental Results of CA classification

6.5 Performance Analysis

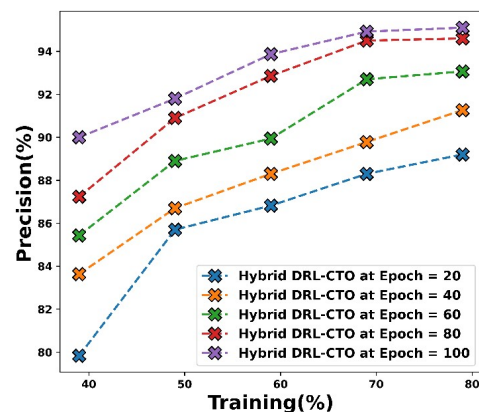
Analysis of Performance is accomplished with the dataset READISCA [17] and it is executed depending on the K-Fold and training percentage.

6.5.1 Performance Analysis with Training Percentage

Figure 6 shows the performance of the proposed Hybrid DRL-CTO model by evaluating the metrics such as precision, accuracy, and recall. These values were recorded across various epochs with a stable training percentage of 80%. The accuracy of the Hybrid DRL-CTO model for the training percentage 80% begins with 91.2% at epoch 20, further, it increases at 40 is 91.9%, at epoch 60 it reaches 94.4%, 95.8% for epoch 80 and at epoch 100 it is 96.7%. The precision of the Hybrid DRL-CTO model for the training percentage 80% attains 89.2% at epoch 20, then it increases to 91.2% at epoch 40, 93.06% at epoch 60, 94.6% at epoch 80, and at epoch 100 it upsurges to 95.1%. The Hybrid DRL-CTO model achieves the recall of 91.5%, 92.6%, 94.9%, 96.1%, and 97.7% for the training percentage 80, with the epochs 20, 40, 60, 80, and 100 respectively. This consistent improvement across metrics with epochs indicates the Hybrid DRL-CTO model's effectiveness and enhanced performance over time with the set training percentage.



Accuracy



Precision

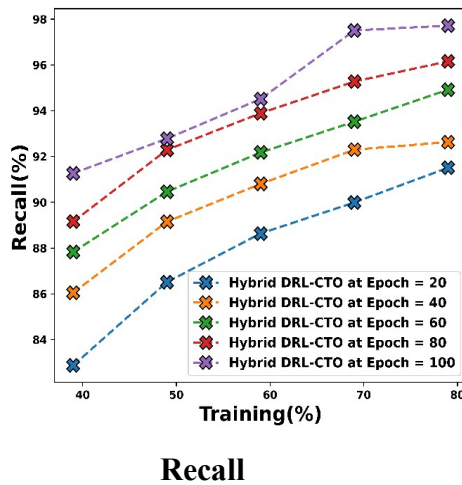


Figure 6: Performance Analysis with Training Percentage

6.5.2 Performance Analysis with K-Fold

Figure 7 shows the performance of the proposed Hybrid DRL-CTO model by measuring the metrics such as precision, accuracy, and recall. With a reliable K-Fold 10, the parameter values metric values vary over epochs 20, 40, 60, 80, and 100. The Hybrid DRL-CTO model accuracy is achieved as 87.9%, 89.4%, 91.2%, 94.6%, and 96.8% for epochs 20, 40, 60, 80, and 100 respectively. The precision of the Hybrid DRL-CTO model at epoch 20 is 86.9%, 88.2% for epoch 40, at epoch 60 it attains 89.8%, at epoch 80 it is 93.7%, and 95.02% for epoch 100. At epoch 20, recall is 88.9%, then it rises to 90.5% at epoch 40, 92.5% at epoch 60, 95.5% at epoch 80 and finally accomplishing 97.1% at epoch 100. This steady improvement across metrics with each epoch highlights the Hybrid DRL-CTO model's improved performance over time.

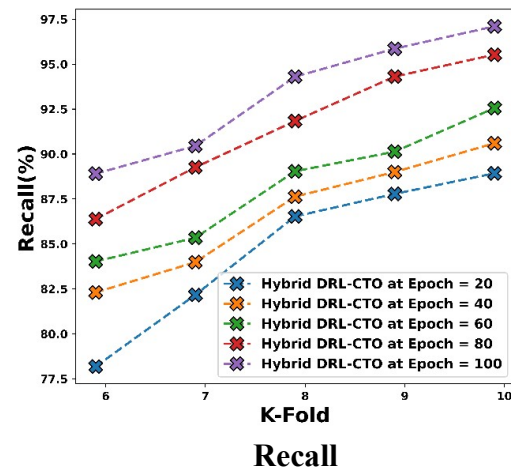
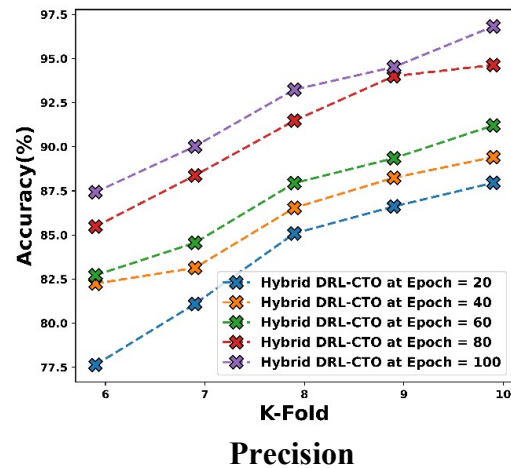
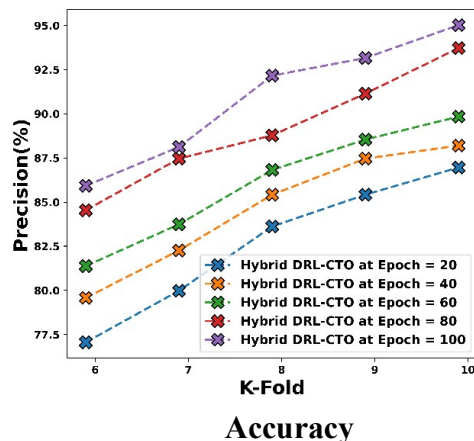


Figure 7: Performance Analysis with K-Fold

6.6 Comparative Analysis

The efficiency of the proposed Hybrid DRL-CTO model is compared with the existing DL-MCD-DT [9], FWO-XGB [13], FSR-XGB [14], ODLCNN [6], Hybrid DRL [18], ICA-Hybrid DRL [15], and TLO-Hybrid DRL [16] approaches in terms of Precision, accuracy, and Recall.

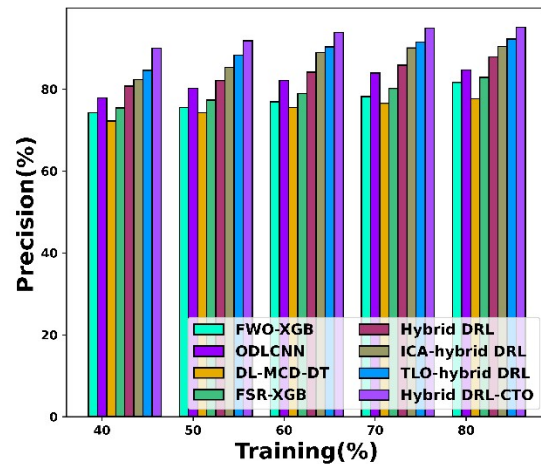
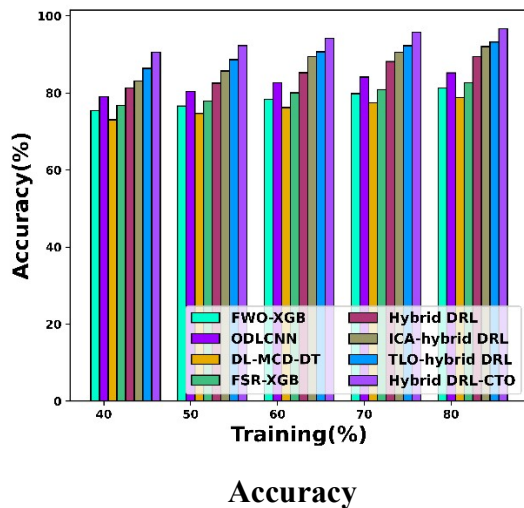
The proposed Hybrid DRL-CTO model's results were critiqued against prior literature. Compared to ODLCNN, FWO-XGB, DL-MCD-DT, FSR-XGB, Hybrid DRL, ICA-Hybrid DRL, and TLO-Hybrid DRL, the Hybrid DRL-CTO showed consistently higher accuracy, recall, and precision. For example, accuracy improved by 15.8% over FWO-XGB and 18.4% over DL-MCD-DT, demonstrating robustness and efficiency.



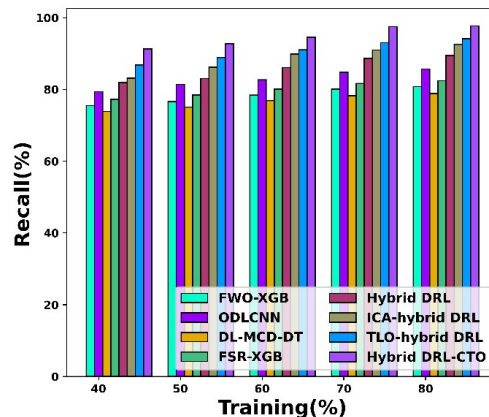
However, limitations remain, including need for validation on more diverse datasets (see Table 1)

6.6.1 Comparative Analysis with Training Percentage

Figure 8 displays the comparative analysis of the proposed Hybrid DRL-CTO model with existing FWO-XGB, ODLCNN, DL-MCD-DT, FSR-XGB, Hybrid DRL, ICA-Hybrid DRL, and TLO-Hybrid DRL approaches in terms training percentage with metrics like precision, accuracy, and recall. The accuracy of the Hybrid DRL-CTO model achieves 96.7%, which is higher than the existing FWO-XGB by 15.8%, ODLCNN by 11.8%, DL-MCD-DT by 18.4%, FSR-XGB by 14.5%, Hybrid DRL by 7.5%, ICA-Hybrid DRL by 4.8%, and TLO-Hybrid DRL by 3.6% for the training percentage of 80. The Hybrid DRL-CTO model reaches a precision of 98.51%, which shows the precision of the Hybrid DRL-CTO model is more progressive than the existing FWO-XGB, ODLCNN, DL-MCD-DT, FSR-XGB, Hybrid DRL, ICA-Hybrid DRL, and TLO-Hybrid DRL by 14.1%, 10.9%, 18.3%, 12.9%, 7.6%, 4.9%, 3.02% respectively in training percentage of 80% and the recall of the Hybrid DRL-CTO model model shows 17.2%, 12.2%, 19.3%, 15.6%, 8.4%, 5.2%, and 3.6% enhancement than the existing FWO-XGB, ODLCNN, DL-MCD-DT, FSR-XGB, Hybrid DRL, ICA-Hybrid DRL, and TLO-Hybrid DRL approaches.



Precision



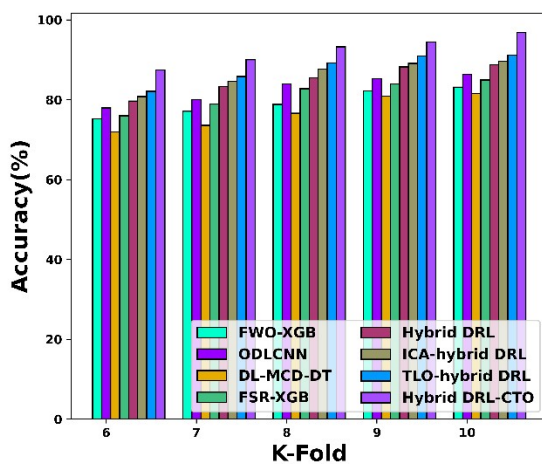
Recall

Figure 8: Comparative Analysis with Training Percentage

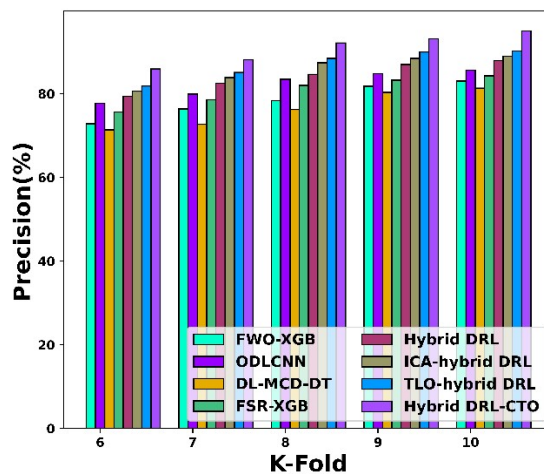
6.6.2 Comparative Analysis with K-Fold

Figure 9 shows the comparative outcomes of the Hybrid DRL-CTO model with existing FWO-XGB, ODLCNN, DL-MCD-DT, FSR-XGB, Hybrid DRL, ICA-Hybrid DRL, and TLO-Hybrid DRL approaches in terms of K-Fold. The Hybrid DRL-CTO model reaches an accuracy of 96.82% for K-Fold 10, which shows 14.11%, 10.81%, 15.74%, 12.26%, 8.36%, 7.41%, and 5.83% advanced than the existing FWO-XGB, ODLCNN, DL-MCD-DT, FSR-XGB, Hybrid DRL, ICA-Hybrid DRL, and TLO-Hybrid DRL methods respectively. The precision achieved by the Hybrid DRL-CTO model for K-Fold 10 is 95.02% which is relatively higher

than the existing FWO-XGB, ODLCNN, DL-MCD-DT, FSR-XGB, Hybrid DRL, ICA-Hybrid DRL, and TLO-Hybrid DRL approaches by 12.58%, 9.89%, 14.39%, 11.28%, 7.49%, 6.35%, and 5.05% respectively. For the K-Fold of 10, the Hybrid DRL-CTO model achieves a recall of 97.10%, that shows the recall is improved over the existing FWO-XGB by 14.29%, ODLCNN by 10.34%, DL-MCD-DT by 15.76%, FSR-XGB by 11.88%, Hybrid DRL by 8.58%, ICA-Hybrid DRL by 7.02%, and TLO-Hybrid DRL by 5.15% respectively.



Accuracy



Precision

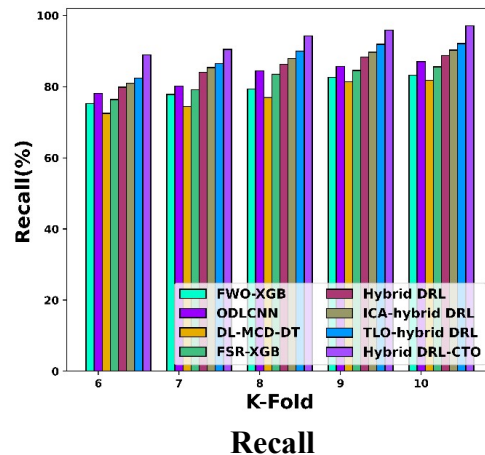


Figure 9; Comparative Analysis with K-Fold

6.7 ROC Analysis

Figure 10 shows the ROC analysis of the proposed Hybrid DRL-CTO model with existing approaches. When the false positive rate is 0.9%, the true positive rate of the existing FWO-XGB, ODLCNN, DL-MCD-DT, FSR-XGB, Hybrid DRL, ICA-Hybrid DRL, TLO-Hybrid DRL and the proposed Hybrid DRL-CTO model are 0.84%, 0.85%, 0.86%, 0.90%, 0.92%, 0.93%, 0.96% and 0.97% respectively. Therefore, from the results, it is evident that the proposed Hybrid DRL-CTO model attained higher accuracy in CA classification when related to the existing FWO-XGB, ODLCNN, DL-MCD-DT, FSR-XGB, Hybrid DRL, ICA-Hybrid DRL, and TLO-Hybrid DRL approaches.

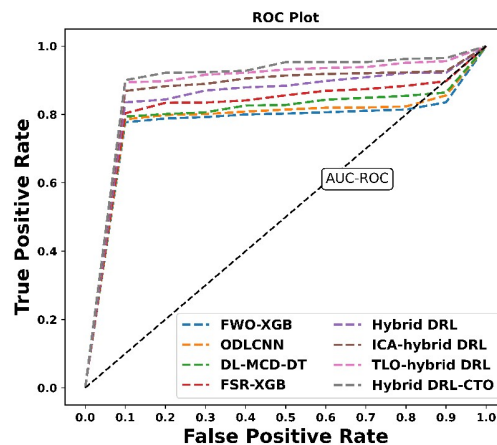


Figure 10; ROC Analysis

6.8 Comparative Discussion

Table 1. Comparative Discussion of the Hybrid DRL-CTO model

The evaluation of the proposed Hybrid DRL-CTO model with the existing approaches is deliberated in this section. The current approaches considered in the research are DL-MCD-DT [9], FWO-XGB [13], FSR-XGB [14], ODLCNN [6], Hybrid DRL [18], ICA-Hybrid DRL [15], and TLO-Hybrid DRL [16]. These approaches have numerous advantages; though, they also have some limitations. So, these existing approaches are unsuccessful in achieving an accurate CA classification. The DL-MCD-DT model has intricate issues in classifying the registers and recognizing the presymptomatic features. The ODLCNN attained a low accuracy rate due to the lack of specification of features in the model and that made the entire categorization process simple. FWO-XGB method failed to validate the model by testing it on the dataset of several label classifications. Due to the lack of consideration of extra classes of ataxia, robustness validation of the FSR-XGB method is difficult. These limitations of the existing methods are overcome by the Hybrid DRL-CTO model, it makes the process of analysis easier using the extraction of Deep S3P Features and also reduces the complexity of data. The extracted Deep S3P Features help the Hybrid DRL model to classify the exact output between the various classes of CA. The Hybrid DRL model is more flexible and computationally efficient and it also improved the performance of memory. The classification accuracy is enhanced due to the incorporation of the CTO algorithm in the Hybrid DRL model. The CTO algorithm diminishes the complexity of computation and enhances the accuracy by eradicating the issues of over fitting, which allows the proposed Hybrid DRL-CTO model to attain better outcomes in CA classification, which are presented in Table 1.

Method s		F W O- X G B	OD LC NN	D L- M C D- D T	F S R - X G B	Hy bri d D R L	IC A- Hy bri d D R L	T L O- Hy bri d D R L	Hy bri d D R L- C T O
TP=80%	Acc ura cy (%)	81 .3	85.2	78 .8	8 2. 6	89. 4	92. 02	93. 1	96. 7
	Pre cisi on (%)	81 .6	84.6	77 .6	8 2. 8	87. 8	90. 4	92. 2	95. 1
	Rec all (%)	80 .8	85.7	78 .8	8 2. 4	89. 4	92. 5	94. 1	97. 71
K-fold=10	Acc ura cy (%)	83 .1	86.3	81 .5	8 4. 9	88. 7	89. 6	91. 1	96. 8
	Pre cisi on (%)	83 .0 6	85.6	81 .3	8 4. 3	87. 9	88. 9	90. 2	95. 02
	Rec all (%)	83 .2	87.0 6	81 .8 0	8 5. 5 6	88. 77	90. 29	92. 10	97. 1

7. CONCLUSION

In this research, the Hybrid DRL-CTO approach is proposed to exactly detect and classify CA based on human gait analysis. The adaptive quality metric-based frame selection aids in reducing the computational load, memory, and processing time required to process a frame. The extraction of Deep S3P features from the ROI extracted frame, makes the Hybrid DRL model to extract more important features, which produce a great impact on the classification correctness by reducing the feature dimensionality. The CTO algorithm tunes the parameters of the Hybrid DRL model to enhance the efficiency and accuracy of the model by resolving the issues of high complexity in small and enormous dimensions. Therefore, the proposed Hybrid DRL-CTO model classifies the CA accurately with enhanced performance and reduced computational complexity. The Hybrid DRL-CTO model achieves higher accuracy and it is assessed by the evaluation metrics, such as

precision, accuracy, and recall that attain the values of 95.02%, 96.8%, and 97.1% respectively. Unanswered questions include the generalizability of results to other movement disorders, performance on additional datasets, and practical integration within real-time clinical systems. Further research should address these open issues. The future work will be focused on designing a hybrid mechanism for detecting and classifying CA and concentrating on assessing the performance with extra datasets rather than evaluating with the single dataset.

REFERENCES:

- [1] Cicirelli G, Impedovo D, Dentamaro V, Marani R, Pirlo G, D'Orazio TR. Human gait analysis in neurodegenerative diseases: a review. *IEEE Journal of Biomedical and Health Informatics*. 2021; 26(1):229-42.
- [2] Vattis K, Oubre B, Luddy AC, Ouillon JS, Eklund NM, Stephen CD, Schmahmann JD, Nunes AS, Gupta AS. Sensitive quantification of cerebellar speech abnormalities using deep learning models. *IEEE Access*. 2024.
- [3] Khalil H, Saad AM, Khairuddin U. Diagnosis of Cerebellar Ataxia Based on Gait Analysis Using Human Pose Estimation: A Deep Learning Approach. *In2022 IEEE-EMBS Conference on Biomedical Engineering and Sciences (IECBES)*. 2022; 201-206.
- [4] Kashyap B, Pathirana PN, Horne M, Power L, Szmulewicz DJ. Modeling the progression of speech deficits in cerebellar ataxia using a mixture mixed-effect machine learning framework. *IEEE Access*. 2021; 9:135343-53.
- [5] Jin L, Lv W, Han G, Ni L, Hu X, Cai H. Gait characteristics and clinical relevance of hereditary spinocerebellar ataxia on deep learning. *Artificial Intelligence in Medicine*. 2020;103:101794.
- [6] Procházka A, Dostál O, Cejnar P, Mohamed HI, Pavelek Z, Vališ M, Vyšata O. Deep learning for accelerometric data assessment and ataxic gait monitoring. *IEEE Transactions on Neural Systems and Rehabilitation Engineering*. 2021; 29:360-7.
- [7] Khalil H, Coronado E, Venture G. Human motion retargeting to Pepper humanoid robot from uncalibrated videos using human pose estimation. *In2021 30th IEEE International Conference on Robot & Human Interactive Communication (RO-MAN)*. 2021; 1145-1152.
- [8] Ngo T, Nguyen DC, Pathirana PN, Corben LA, Delatycki MB, Horne M, Szmulewicz DJ, Roberts M. Federated deep learning for the diagnosis of cerebellar ataxia: Privacy preservation and auto-crafted feature extractor. *IEEE Transactions on Neural Systems and Rehabilitation Engineering*. 2022; 30:803-11.
- [9] Stoean C, Stoean R, Atencia M, Abdar M, Velázquez-Pérez L, Khosravi A, Nahavandi S, Acharya UR, Joya G. Automated detection of presymptomatic conditions in Spinocerebellar Ataxia type 2 using Monte Carlo dropout and deep neural network techniques with electrooculogram signals. *Sensors*. 2020; 20(11):3032.
- [10] Shanmuga Sundari M, Jadala VC. Neurological disease prediction using impaired gait analysis for foot position in cerebellar ataxia by ensemble approach. *Automatika: časopis za automatiku, mjerenje, elektroniku, računarstvo i komunikacije*. 2023; 64(3):540-9.
- [11] Rahman W, Hasan M, Islam MS, Olubajo T, Thaker J, Abdelkader AR, Yang P, Paulson H, Oz G, Durr A, Klockgether T. Auto-gait: Automatic ataxia risk assessment with computer vision from gait task videos. *Proceedings of the ACM on Interactive, Mobile, Wearable and Ubiquitous Technologies*. 2023; 7(1):1-9.
- [12] Dostál O, Procházka A, Vyšata O, Ťupa O, Cejnar P, Vališ M. Recognition of motion patterns using accelerometers for ataxic gait assessment. *Neural Computing and Applications*. 2021; 33:2207-15.
- [13] Seetharama PD, Math S. Ataxic person prediction using feature optimized based on machine learning model. *International Journal of Electrical and Computer Engineering (IJECE)*. 2024; 14(2):2100-9.
- [14] Seetharama PD, Math S. Ataxia severity classification using enhanced feature selection and ranking optimization through machine learning model. *Indonesian Journal of Electrical Engineering and Computer Science*. 2023; 32(3):1605-13.

-
- [15] Ayar M, Isazadeh A, Gharehchopogh FS, Seyedi M. NSICA: Multi-objective imperialist competitive algorithm for feature selection in arrhythmia diagnosis. *Computers in Biology and Medicine*. 2023; 161:107025.
 - [16] Toopshekan A, Abedian A, Azizi A, Ahmadi E, Rad MA. Optimization of a CHP system using a forecasting dispatch and teaching-learning-based optimization algorithm. *Energy*. 2023; 285:128671.
 - [17] READISCA dataset, “<https://rochester.app.box.com/v/AtaxiaDataset>”, Accessed on October 2024.
 - [18] Musfeld P, Souza AS, Oberauer K. Repetition learning is neither a continuous nor an implicit process. *Proceedings of the National Academy of Sciences*. 2023; 120(16):e2218042120.