

AN IDEAL REGIONAL LEARNING MODEL FOR DISABILITY PREDICTION USING PROVIDED INPUT SAMPLES

R. Santhosh

Abstract

One type of specific learning disability that might cause academic dissatisfaction is dyslexia. Detection is crucial yet challenging, particularly for languages with simple orthographies. Finding dyslexia remains a major goal for research efforts across many fields, as seen by numerous studies published in numerous scientific journals. These activities, like classifying medical datasets, benefit greatly from deep learning (DL) and other statistical techniques. To detect dyslexia, we have integrated fundamental DL algorithms into this paper. However, a significant amount of manual labor is still required to properly classify the large dataset utilized for training. Recent research has demonstrated that DL models are useful in reducing the workload associated with feature engineering. However, applying the Regional Convolutional Network Model (RCNM) directly to the classification issue does not result in a statistically significant gain in performance. We offer an efficient learning model built on genetic algorithms for better DL prediction. In the initial step of the model, dyslexia features are extracted, and then DL algorithms are used. Next, we merge many DL models by choosing the best weights using an adaptive genetic method. To demonstrate that our suggested method greatly improves accuracy, we ran experiments on the Dyslexia dataset.

Keywords : dyslexia, prediction, machine learning, feature, accuracy

I. INTRODUCTION

Moreover, 10% of people worldwide suffer from dyslexia, a unique learning disorder with neurobiological causes. Poor spelling and decoding skills, as well as challenges with accurate and/or fluent word recognition, are characteristics of dyslexia. These problems can arise from phonological component impairments, which may seem shocking when considering other cognitive skills and the

provision of efficient classroom education [1]. The first step in overcoming the difficulties caused by dyslexia is acknowledging that one has it. However, in the absence of appropriate accommodations, dyslexic students are more likely to struggle in the classroom; according to one study, 35% of dyslexic students drop out of high school, and another estimates that only 2% of them will complete an eight-year bachelor's degree [2]. Symptoms of dyslexia usually appear after a child reaches a specific age and literacy level. These days, screening (pre-) readers requires a professional therapist or costly tools such as fMRI scans. Other than reading and writing, dyslexia has been studied in the context of visual perception, short-term memory, auditory perception and executive functioning [3]. Our work shows an efficient way to screen for suspected dyslexia in pre-readers using ML with data from language-independent content integrated into a web-based game. According to theory, a game may be publicly accessible, warning concerned parents that their child may have dyslexia and urging them to get help from a professional. It is quite difficult to diagnose dyslexia in languages with distinct orthographies. Because of the more solid connection between grapheme (letter) and phoneme (sound), people with dyslexia find it easier to learn and read in languages with shallow orthographies, like English [4]. People with dyslexia are sometimes misdiagnosed and labeled as having a "hidden disability" because their symptoms are less severe in languages with transparent orthographies. Unfortunately, when speakers discover that they might have dyslexia as a result of academic failure, it is sometimes too late to intervene successfully. In the current diagnostic and screening process, in-person, time-consuming examinations [5] are used to measure reading errors, reading speed (words per minute), reading words, writing errors, reading fluency, text comprehension, and reading pseudo-words.

Automatic methods for diagnosing and classifying dyslexia have been used in recent decades to increase the accuracy of illness recognition and help doctors make better decisions. Earlier approaches that mainly used ML techniques for automatic heartbeat categorization required feature selection, feature extraction, and classification [6]. A subfield of artificial intelligence called machine learning builds mathematical models using training data. The system aims to make decisions on its own and provide an accurate diagnosis without the need for human input. In [7], Mitchell

Professor & Head,
Department of Computer Science and Engineering
Karpagam Academy of Higher Education
Coimbatore, India
santhoshrd@gmail.com

* Corresponding Author

offers a traditional yet practical definition. Because ML can be used to improve the highly beneficial computer software that is created and utilized in many applications, including the medical profession, as well as to help researchers answer significant scientific and technical concerns, researchers are interested in DL [8] – [10]. For automatic dyslexia identification, we propose a hybridized and optimized DL method in this study. To improve the DL classifier's performance, we offer a learning model and selection techniques that improve the classifier model using Regional Convolutional Network Model (RCNM). The proposed model gives promising outcomes.

II. RELATED WORKS

Given that dyslexia is among the most complex neurological conditions affecting the brain, researchers in current neuroscience are becoming increasingly interested in it. The International Dyslexia Association defines dyslexia as a disorder with challenges in processing languages, spelling, and recognizing accurate words. Auditory, phonological problems, and Visual are the primary causes of dyslexia. These disorders appear as anomalies that start to form at birth. Apps and games are among the many resources available for diagnosing and assisting people with dyslexia [11]. Gamification has been incorporated into several frameworks, applications and use cases. Gamified game designers utilize elements of other games to increase user motivation and engagement. With a focus on the gaming experience, games' linguistic content is modified for screen readers and screen pre-readers. The author uses the only one to report an 89.2% success rate using phonological processing features. However, games weren't included, and accumulating features necessitates extensive testing, which costs money and effort. The classification is carried out on a very small sample size ($n = 56$), and over-fitting is not addressed nor verified.

The study presented in [12] used k-means clustering, support vector machines, and convolutional neural networks. MRI scans were used to distinguish between people with and without dyslexia. At 94.8%, ANN showed the highest accuracy on a dataset of 56 samples. In [12], it is shown that MRI scans can be used to differentiate between individuals with and without dyslexia. SVM was the ML approach employed in this study, which employed a dataset of 236 people. In this investigation, an 83% success rate was found. It was possible to identify dyslexia in 80 people using EEG scans at the beginning phase with an average accuracy of 89.6%, 89.7%, and 85.7%, respectively, by applying ML techniques such as the ANN, fuzzy logic classifier, and k means. Numerous current studies employing nonlinear classifiers for diagnosing dyslexia are trained with almost

identical proportions, even though an increased risk of dyslexia in children is only 5%–12%. The effectiveness of learning algorithms is severely impacted by this direct avoidance of the class imbalance issue, making it difficult to apply trained models to the general population.

The authors of [13] divide 925 participants' MRI ++-scan data into dyslexic and non 7/dyslexic groups using linear SVMs. The study found that accuracy is almost 80%. The author identified dyslexia in a dataset comprising 267 people by using eye tracker features in a linguistic computer game. In this study, researchers reported an 85% accuracy with MATLAB's SVM from the LIBSVM Toolbox. In a different study, researchers examined data from 185 participants to determine dyslexia using eye-tracking equipment. Using the SVM with automatic recursive feature elimination, a 96% accuracy rate was attained. Using an eye-tracking feature, researchers have also employed the SVM to diagnose dyslexia, according to [14]. On a dataset of size 97, an accuracy of 80% is achieved in this work. Eighty-eight proficient readers (scoring in the 80 or 90%) and 97 low-skilled readers (scoring in the 5% for word decoding) were among the 185 Swedish children (ages 9–10) examined in another study [14]. Unsurprisingly, the SVM classifier achieved a 10-fold cross-validation accuracy of 95.6%, considering the substantial variations in reading proficiency between the groups and their approximately comparable sizes. Numerous aspects of eye movement were investigated in the study, such as the duration of forward and backward saccades and the duration of each fixation. The same 95% classification accuracy was achieved when the same dataset was reanalyzed in [15] utilizing a Hybrid Kernel SVM-PSO approach according to Particle Swarm Optimization (PSO). Information obtained from the eye movements of the participants was used to determine the primary components of this analysis.

III. METHODOLOGY

With ages ranging from 7 to 17, 392 (10.7%) of the 3,644 participants in this data set were diagnosed with dyslexia. The data set for each participant had 196 features in total. While the following 196 aspects represent the participant's performance as a result of their interaction with the test's 32 questions, the first four characteristics relate to the participant's personal information.

The full system and the elements of the suggested model are depicted. Initially, feature maps are produced by the image parser. Two branches are then fed these feature maps. We used a two-stage strategy similar to the left branch of our study. The Region Network processes feature maps to produce region proposals (RoIs). The box offsets, binary masks, and class

labels for every ROI are predicted using the grading network in the next step. Furthermore, we include a right branch that generates an epithelial cell score to determine whether epithelial cells are present in the image. The outcomes of the network and the grading network determine the network's final prediction. A post-processing step employing a conditional random field is applied after the initial predictions. Our model's main goals are to detect the existence of epithelial cells and provide a segmentation mask that matches the Gleason grade. These two tasks are supposed to be carried out independently. Different networks make up the grading network. Classification loss L_{obj} .

$$L_{obj} = \sum_{i=1}^N (-y_i \log(p_i) - (1 - y_i) \log(1 - p_i)) \quad (1)$$

One popular loss function for binary classification is L_{obj} . The ground truth representation for the given image is $y_i \in \{0, 1\}$, where $y_i = 0$ denotes the absence of an ROI and $y_i = 1$ shows the presence of at least one ROI. If N represents all the images in the training datasets, Our model's sigmoid layer output, represented by $p_i \in (0, 1)$ should be defined as the probability that an ROI would be present in the image. Consequently, our model's overall loss is provided by:

$$L = L_{obj} + L_{cls} + L_{box} + L_{mask} \quad (2)$$

Upon generating predictions for each patch using our model, we reconstructed the original tiles by seamlessly integrating the individual image patches. Fig 1 final two rows show that each patch's margins may exhibit artifacts due to this stitching step. We employed a fully coupled conditional random field (CRF) model to tackle this issue. The author originally developed this method to efficiently compute picture segmentations, showing that it could utilize long-range dependencies and capture precise edge features. Subsequently, the model incorporated this methodology as a post-processing step within a convolutional neural network (CNN). In the context of a conditional random field (I, X) , where X encompasses the entire image represented as $\{x_1, x_2, \dots, x_N\}$, the probability $P(X/I)$ is expressed as $1/Z(I) \exp(-E(X/I))$. Here, the label corresponding to the i th pixel is represented by x_i and N denotes the total number of pixels in the image. In the energy function, the model identifies the first term on the right as the unary potential and the second as the pairwise potential. The segmentation head in the grading network calculates the unary potential and the probability of label assignment at pixel I using the formula $\theta_i(x_i) = -\log P(x_i)$. The pairwise potential is defined as $\theta_{i,j} = \mu(x_i, x_j) K_m = 1_{\omega, k_m(f, f_j)}$, where $\mu(x_i, x_j)$ equals one if x_i is equal to x_j and 0 otherwise. Each k_m represents the Gaussian kernel, which

relies on features (denoted as f) derived from pixels i and j and is influenced by a learnable parameter ω_m . As demonstrated, we incorporate bilateral position and color components within the kernels.

$$\omega_1 \exp\left(-\frac{\|p_i - p_j\|^2}{2\sigma_\alpha^2}\right) - \frac{\|I_i - I_j\|^2}{2\sigma_\beta^2} + \omega_2 \exp\left(-\frac{\|p_i - p_j\|^2}{2\sigma_\gamma^2}\right) \quad (3)$$

Where, I stands for pixel color intensity and P for pixel position, the initial kernel term compels neighboring pixels of identical color to be categorized within the same class. In contrast, the subsequent kernel term eliminates small, isolated regions. The experiment empirically established the "scale" of the Gaussian kernels, which are governed by the hyper-parameters $\sigma_\alpha, \sigma_\beta$, and σ_γ . Henceforth, this study will refer to the fully connected Conditional Random Field (CRF).

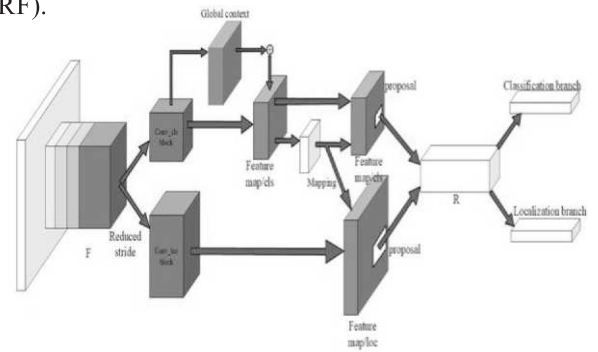


Fig 1: Regional Convolutional Network model

IV. NUMERICAL RESULTS

We assessed the performance of our classification approach using popular metrics such as F1-score, accuracy, precision, and Recall. Below is a description of the equations for the various metrics:

F1-score: integrates Recall and precision into a single value.

$$F1 = 2 * \frac{\text{precision} * \text{recall}}{\text{precision} + \text{recall}} \quad (4)$$

In this study, the terms True Positive, False Positive, True Negative, and False Negative are denoted by the abbreviations TP, TN, FP, and FN, respectively. The model, developed using MATLAB 2020 was evaluated utilizing a graphics card with 16 gigabytes of random-access memory (RAM) and an RTX 2060 resolution.

Accuracy: The calculation can be expressed as the ratio of the sum of TN and TP representing accurate predictions, to the overall number of predictions made.

$$Acc = \frac{TP + TN}{TP + FP + TN + FN} \quad (5)$$

Precision : This can be characterized as the proportion of all positive predictions, which includes both TP and FP.

Recall : This metric assesses the ratio of true positive predictions, which refers to the accurately identified positive instances, to the total number of positive occurrences within the dataset.

$$Recall = \frac{TP}{precision} \quad (6)$$

To determine how well our approach worked and how it compared to other models, we conducted a series of trials using various ML techniques. Some classification techniques were experimentally assessed. The exact Dyslexia data set was used to test each of these classifiers. The objective was to test various classifiers for detecting the dyslexic handicap and select the most effective one. Every experiment used the same data split: 70% for training and 30% for testing. The F1-score, Accuracy, Precision, and Recall were calculated to compare the classifiers. Table 1 uses a unified strategy to compare the outcomes of multiple different categorization ML models on the Dyslexia dataset.

Table 1: Classification outcomes

Models	Acc	Pre	Recall	F1
DT	78	80	79	79
Boost	82	83	82	82
LR	85	85	86	85
SGD	66	62	67	63
RF	88	89	89	89
GBM	84	84	84	84
ETC	83	83	83	83
GNB	86	86	87	86
SVM	88	87	87	87
RCNM	96	97	97	96

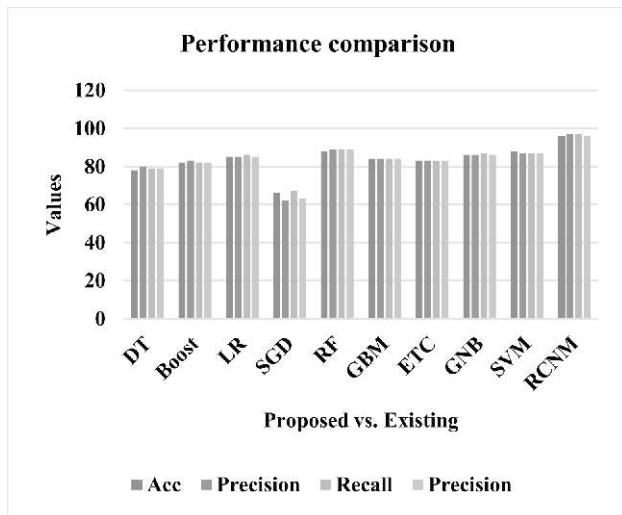


Fig 2 : Classification results

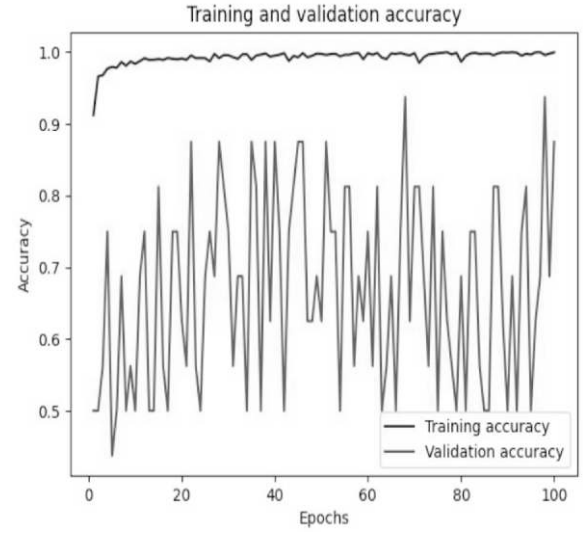


Fig 3: Accuracy analysis

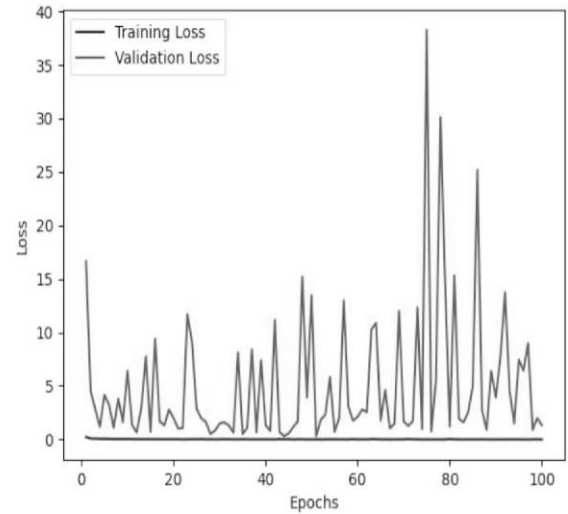


Fig 4: Loss analysis

Table 1 provides a summary of the information acquired throughout the studies. This kind of data analysis reveals that the proposed model performed noticeably better in accuracy rate than the other classifiers as in Fig 2 to Fig 4.

V. CONCLUSION

This work offered a solution to the identifying challenge of dyslexia. The model was built using RCNM. The main goal of the proposed model was to construct a robust classifier that can accurately identify dyslexia using learning technique. The second objective was to find the best weights for the classifier in our algorithm. The Regional Convolutional Network Model (RCNM) were used to finish this study. Furthermore, the outcome of the learning module was predicted using the proposed technique. The results of the experimental validation of our model demonstrated its good performance, with an accuracy of 95%. Future projects could benefit from several enhancements. To improve the findings, we

specifically intend to target hybrid deep learning-based dyslexia diagnosis. Additionally, we want to employ multi-modal and other types of datasets.

REFERENCES

- [1] Goswami, U., Barnes, L., Mead, N., Power, A. J., & Leong, V. (2016). Prosodic similarity effects in short-term memory in developmental dyslexia. *Dyslexia*, 22(4), 287–304.
- [2] Al-Haija, Q. A., Krichen, M., & Elhaija, W. A. (2022). Machine-learning-based darknet traffic detection system for IoT applications. *Electronics*, 11(4), 556.
- [3] Qaisar, S., Mihoub, A., Krichen, M., & Nisar, H. (2021). Multirate processing with selective subbands and machine learning for efficient arrhythmia classification. *Sensors*, 21(4), 1511.
- [4] Mihoub, A., Snoun, H., Krichen, M., Salah, R. B. H., & Kahia, M. (2020). Predicting COVID-19 spread level using socio-economic indicators and machine learning techniques. In *Proceedings of the 1st International Conference on Smart Systems and Emerging Technologies (SMARTTECH)* (pp. 128–133).
- [5] Rauschenberger, M., Baeza-Yates, R., & Rello, L. (2019). Technologies for dyslexia. In *Web Accessibility: A Foundation for Research* (pp. 603–627). Springer.
- [6] Mora, A., Riera, D., Gonzalez, C., & Arnedo-Moreno, J. (2015). A literature review of gamification design frameworks. In *Proceedings of the 7th International Conference on Games and Virtual Worlds for Serious Applications (VS-Games)* (pp. 1–8).
- [7] Seaborn, K., & Fels, D. I. (2015). Gamification in theory and action: A survey. *International Journal of Human-Computer Studies*, 74, 14–31.
- [8] Rauschenberger, M., Willems, A., Ternieden, M., & Thomaschewski, J. (2019). Towards the use of gamification frameworks in learning environments. *Journal of Interactive Learning Research*, 30(2), 147–165.
- [9] Rello, L., Romero, E., Rauschenberger, M., Ali, A., Williams, K., Bigham, J. P., & White, N. C. (2018). Screening dyslexia for English using HCI measures and machine learning. In *Proceedings of the International Conference on Digital Health* (pp. 80–84).
- [10] Poole, A., Zulkernine, F., & Aylward, C. (2017). Lexa: A tool for detecting dyslexia through auditory processing. In *Proceedings of the IEEE Symposium Series on Computational Intelligence (SSCI)* (pp. 1–5).
- [11] Al-Barhamtoshy, H., & Motaweh, D. M. (2017). Diagnosis of dyslexia using computation analysis. In *Proceedings of the International Conference on Informatics, Health & Technology (ICIHT)* (pp. 1–7).
- [12] Tamboer, P., Vorst, H. C. M., Ghebreab, S., & Scholte, H. S. (2016). Machine learning and dyslexia: Classification of individual structural neuroimaging scans of students with and without dyslexia. *NeuroImage: Clinical*, 11, 508–514.
- [13] Asvestopoulou, T., Manousaki, V., Psistakis, A., Smyrnakis, I., Andreadakis, V., Aslanides, I. M., & Papadopoulou, M. (2019). DysLexML: Screening tool for dyslexia using machine learning. *arXiv:1903.06274*.
- [14] Rello, L., Baeza-Yates, R., Ali, A., Bigham, J. P., & Serra, M. (2020). Predicting risk of dyslexia with an online gamified test. *PLoS ONE*, 15(12), e0241687.
- [15] Kumari, P., Kumar, D., & Mittal, M. (2021). An ensemble approach for classification and prediction of diabetes mellitus using soft voting classifier. *International Journal of Cognitive Computing in Engineering*, 2, 40–46.