

A Study of Multi Layer Perceptron (MLP)

Neha Kapil

Research Scholar, Glocal University, Saharanpur, U.P.

Dr. Jugnesh Kumar

Supervisor, Asso. Prof. Glocal University, Saharanpur, U.P.

Abstract - Interruptions in the supply of electricity cause numerous losses to consumers, whether residential or industrial and may result in fines being imposed on the regulatory agency's concessionaire. In Brazil, the electrical transmission and distribution systems cover a large territorial area, and because they are usually outdoors, they are exposed to environmental variations. In this context, periodic inspections are carried out on the electrical networks, and ultrasound equipment is widely used, due to non-destructive analysis characteristics. Ultrasonic inspection allows the identification of defective insulators based on the signal interpreted by an operator. This task fundamentally depends on the operator's experience in this interpretation. In this way, it is intended to test machine learning applications to interpret ultrasound signals obtained from electrical grid insulators, distribution, class 25 kV. Currently, research in the area uses several models of artificial intelligence for various types of evaluation. This paper discuss the Multi Layer Perceptron(MLP) and Comparison of single and multiple Multi Layer Perceptron(MLP)

Keywords - Multi Layer Perceptron(MLP), FMIC

INTRODUCTION

Using the characteristics and classifiers discussed in the preceding chapters, the authors conduct a series of experiments. In order to make the most of the zoning method's potential, many zoning strategies are put into practice, and NVDs are computed from four distinct starting points. Both binary and four-dimensional pictures were processed for feature extraction. You may use QDF, MQDF, SQDF, MLP, and MMLP for categorization.

Due of MLP's complexity, a model called Multiple MLP with a structure very similar to MLP is used. The many permutations of character types are tried out. In addition, a comparison was made between the MLP module and the quadratic classifier. By reducing complexity, MQDF helps improve QDF's already impressive output. SQDF is intended to provide improved recognition outcomes.

A basic property We try out RLC, gradient features, and 12 DC. By adding up the DC in each block of 12, we may drastically minimize the number of dimensions involved. As an additional step toward feature extraction, we diagonally cross the photos. Last but not least, a fusion of features approach To improve the accuracy of Malayalam character recognition, a GBF - RLC hybrid is being tested. RLC, gradient characteristics, aspect ratio, centroid, and character code are all combined to see what happens. Experiments are presented in detail in this chapter.

Four data sets were created from the retrieved characteristics of the samples for use in training and testing. Seventy percent of the samples are used for training, while the other thirty percent are used for testing. A new 70:30 split between training and testing data is being evaluated for the next data set, which will be built using a different mix of the identical character samples.

With the four distinct data sets in place, we can be confident that all possible character variants will be used throughout both the training and testing stages.

REIVEW OF LITERATURE

R. Jayadevan et.al (2011) Survey results indicate that improper character segmentation of touched or broken characters is the primary cause of failures in identifying printed Devanagari characters. Along with several page segmentation, word segmentation, and character segmentation techniques, this article also reviews a variety of preprocessing methods, such as skew angle and Shirrekha detection in word, noise reduction, etc. It determined that little work on unconstrained Devanagari handwriting recognition has been documented.

Cheng- Lin Liu et.al (2004) Every day, a quarter of the world, mostly in Asia, uses Chinese characters to communicate. Kanji in both Japan and China are based on the same basic character shape. Among the most often used databases is TUAT (Tokyo University of Agriculture and Technology), which is split into two subsets: the Nakayosi, which is typically used in classifier learning, and the Kuchiblu, which is put to use in evaluation. Half-thinning a character picture and identifying strokes are crucial steps in the pre-processing phase.

Andrew Senior et. al (2015) One author's handwritten page is scanned and sent to be segmented. Height, angle, slope, stroke width, and rotation are some of the factors linked to the word that have been identified as removing. The aforementioned factors have been eliminated thanks to the introduction of normalisation. Noise in the source text may be eliminated using convolution with a 2-dimensional Gaussian filter, followed by iterative application of an erosive thinning method to thin down the strokes, and finally through removal of the noise altogether.

B. Wegmann et. al (2000) compressed picture quality is improved by focusing on feature-specific vector quantization. Because handwriting presents a broad range of challenges for recognition systems. The development of an effective handwriting recognition system requires a significant amount of memory and processing capacity. Due to the finite nature of system resources, optimising the balance between the intended recognition rate and the needed system resources is a crucial consideration when creating a handwriting recognition system.

Luiz Oliveira et.al. (2002) The method of recognition used is segmentation based. Numbers from Brazilian bank checks and the NIST SD19(hsf 7 series) of number strings have been the primary targets of this work. Within a multihypothesis method, many stages, such as segmentation, recognition, and post- processing, are combined. Under and over segmentation were found to be manageable with effective strategies. Extensive testing using the NIST SD 19 database of numerical values has been performed.

B. B. Chaudhury et. al (2009) Line identification in texts written in Bangla, English, Hindi, and Malayalam are the focus of this project. By drawing a histogram in steps, we may identify individual lines and spaces in a text. It is possible to detect noise and calculate the width of the vertical strip. 18 average standard deviations were used to sketch the first sets of curves. The highest percentage of correct identification was 94.31%, and it was achieved across six distinct scripts on individual lines of text.

U. Pal et. al (2003) An automated method is described for determining whether lines of a document contain text written in various Indian scripts. Flat bed scanners were used to complete the digitization process. Using a histogram-based method, we can transform the data into a grayscale picture.

S. Kompalli et. al (2005) In order to calculate the gradient characteristics, the Sobel operator is used. It does this by gauging the size and direction of intensity shifts in each pixel's immediate vicinity. **Sethi and chatterjee (2011)** outlined a structural method for reading handwritten Devanagari numbers. Line segments in the horizontal, vertical, and diagonal planes are the fundamental building blocks.

U. Pal et. al (2012) The recognition procedure employs four different sets of characteristics: two sets of binary features and two sets of gray-scale picture features. In addition to the computation of the curvature and gradient features, these computations have been completed. Euclidean Distance (ED), Projection Distance (PD), Subspace Method (SM), Modified Quadratic Discriminant Function (MQDF), Modified Projection Distance (MPD), Linear Discriminant Function (LD), Mirror Image Learning (MIL), Support Vector Machine (SVM), etc. are all taken into account for comparisons in this work. The specifics of the outcome reveal that MPD achieves 94.15 percent accuracy for a single character.

Macro Bressan et. al (2003) When characteristics that are statistically independent of one another are used to model classes, classification is complete. The divergence measure for class separability is considerably reduced to the sum of undimensional divergence if class-conditional independence is assumed. We modify divergence and the Bayes decision scheme such that they work with this class-conditioned model. The parameters of the independent component analysis are calculated. Independent Component Analysis is frequently utilised in conjunction with class information. It is clear from the results of the previous section that conditionally independent features considerably reduce the complexity of pattern classification and feature selection issues.

COMPARISON OF SINGLE AND MULTIPLE MLP

A database with 500 samples and 30 classes is used for the preparation procedures in the earlier experiment. Multiple iterations of training and testing were conducted, each time with a new set of modules and character sets. Sixty percent of samples are utilized for training and forty percent for testing across all investigations. Single MLP module (classic design) and six MLP modules (multi-layer perceptron) were used in the experiment (proposed architecture). After doing an experiment, the number of MLP modules was set. It is completely at random how the letters are grouped together.

Table 6.8 shows the outcome using a single MLP module. With FMIC, the average recognition accuracy is 75.01%, but with FMLC, it's 78.35%.

Table 1: Performance of FMIC and FMLC for 16 features on 64 x 64 binarized and thinned pictures using a single MLP module

Data Set	Classification Accuracy (%)	
	FMIC	FMLC
1	74.38	78.12
2	75.12	78.45
3	75.27	77.63
4	75.28	79.18
Average	75.01	78.35

In Table, we see a compilation of the findings from applying six MLP modules to each of the six categories. With FMIC, the average recognition accuracy for the whole three-character set is 92.28 percent, whereas with FMLC, it's 93.30 percent. Using six MMLP modules divided into six classes,

FMLC recognition accuracy is maximized to 93.30%. Recognition accuracy for FMIC is improved by more than 17% when MMLP is used instead of MLP. When comparing MMLP with MLP, FMLC gains over 15% with MMLP.

Table 2: Performance of FMIC and FMLC for character set 1 classification on 64 x 64 binarized and thinned images utilizing six MMLP modules and six groups

Character Set 1	Classification Accuracy (%)	
	FMIC	FMLC
1, 3, 4, 5, 9	94.45	93.78
2, 6, 10, 13, 21	89.35	91.27
7, 18, 19, 25, 28	91.90	93.07
8, 15, 20, 27, 30	92.87	93.95
11, 16, 22, 24, 29	92.68	94.22
12, 14, 17, 23, 26	92.07	93.70
Average	92.22	93.33

Table 3: Performance of FMIC and FMLC for character set 2 classification on 64 x 64 binarized and thinned pictures utilizing six MMLP modules with six groups.

Character Set 2	Classification Accuracy (%)	
	FMIC	FMLC
3, 6, 15, 18, 26	90.25	91.47
4, 10, 16, 19, 20	93.63	94.37
5, 7, 13, 14, 24	92.95	93.75
1, 12, 22, 25, 27	92.50	93.53
9, 17, 21, 28, 30	92.25	93.37
2, 8, 11, 23, 29	92.75	94.23
Average	92.39	93.45

Table 4: Performance of FMIC and FMLC for character set 3 classification on 64 x 64 binarized and thinned pictures utilizing six MMLP modules with six groups.

Character Set 3	Classification Accuracy (%)	
	FMIC	FMLC
3, 7, 10, 19, 23	90.87	91.30
6, 8, 9, 16, 17	90.55	91.98
5, 11, 13, 15, 22	93.30	94.02
18, 20, 21, 24, 27	91.60	93.52
2, 4, 25, 26, 30	93.80	93.25
1, 12, 14, 28, 29	93.20	94.57
Average	92.22	93.11

In Table, we see a compilation of the findings from applying six MLP modules to each of the six categories. With FMIC, the average recognition accuracy for the whole three-character set is 92.28 percent, whereas with FMLC, it's 93.30 percent. Using six MMLP modules divided into six classes, FMLC recognition accuracy is maximized to 93.30%. Recognition accuracy for FMIC is improved by more than 17% when MMLP is used instead of MLP. When comparing MMLP with

MLP, FMLC gains over 15% with MMLP.

Table 5: Performance of FMIC and FMLC for character set 1 classification on 64×64 binarized and thinned pictures utilizing five MMLP modules with seven groups.

Character Set 1	Classification Accuracy (%)	
	FMIC	FMLC
1,3,4,5,9,12	91.77	91.65
2,6,10,13,14,21	89.15	91.80
7,17,18,19,25,28	90.77	92.33
8,15,20,23,27,30	91.00	91.77
11,16,22,24,26,29	90.77	92.55
Average	90.69	92.02

Table 6: Performance of FMIC and FMLC for character set 2 classification on 64×64 binarized and thinned pictures utilizing five MMLP modules with seven groups.

Character Set 2	Classification Accuracy (%)	
	FMIC	FMLC
2,3,6,15,18,26	88.33	88.47
4,8,10,16,19,20	92.53	92.93
5,7,11,13,14,24	92.77	92.53
1,12,22,23,25,27	90.85	91.83
9,17,21,28,29,30	91.82	92.87
Average	91.26	91.73

Table 7: Performance of FMIC and FMLC for character set 3 classification on 64×64 binarized and thinned pictures utilizing five MMLP modules with seven groups

Character Set 3	Classification Accuracy (%)	
	FMIC	FMLC
1,7,10,19,23,2	88.70	90.70
6,8,9,12,16,1	88.43	90.33
5,11,13,14,15,22	93.20	93.75
3,18,20,21,24,27	89.87	91.22
2,4,25,26,29,30	92.85	92.63
Average	90.61	91.73

Five MLP modules were used to analyze data from eight different categories, and their combined findings are shown in Tables. FMIC achieves an average recognition accuracy of 90.83 percent for the three- character set, whereas FMLC achieves an accuracy of 92.03%. The optimal recognition rate for FMLC is 92.03% while employing five MMLP modules divided into eight classes. As with FMIC, the recognition accuracy may be improved by more than 15% with MMLP as opposed to MLP. More than 13% more FMLC is produced by MMLP than by MLP.

Table 8: Performance of FMIC and FMLC for character set 1 classification on 64×64 binarized and thinned pictures utilizing five MMLP modules with eight groups.

Character Set 1	Classification Accuracy (%)	
	FMIC	FMLC
1,3,4,5,9,12,29	91.37	91.22
2,6,10,13,14,21,26	88.53	90.87
7,11,17,18,19,25,28	88.05	91.48
8,15,16,20,23,27,30	89.02	91.40
22,24	97.10	96.98
Average	90.81	92.39

Table 9: Performance of FMIC and FMLC for character set 2 classification on 64×64 binarized and thinned pictures utilizing five MMLP modules with eight groups

Character Set 2	Classification Accuracy (%)	
	FMIC	FMLC
2,3,6,15,18,26,30	87.37	88.00
4,8,10,16,19,20,29	91.85	91.82
5,7,11,13,14,21,24	89.50	91.00
1,12,17,22,23,25,27	88.15	90.63
9,28	97.75	97.80
Average	90.92	91.85

Table 10: Performance of FMIC and FMLC for character set 3 classification on 64 x 64 binarized and thinned pictures utilizing five MMLP modules with eight groups

Character Set 3	Classification Accuracy (%)	
	FMIC	FMLC
1,7,10,13,19,25,28	88.75	89.45
2,6,8,9,12,16,17	88.15	89.78
3,11,14,18,22,23,30	88.02	89.33
5,15,20,21,24,26,27	90.92	93.20
4,29	97.90	97.47
Average	90.75	91.85

Tables describe the outcomes of using four MLP modules over nine different categories. For the three- character set, FMIC achieves an average recognition accuracy of 88.1%, whereas FMLC achieves an accuracy of 89.81%. Using four MMLP modules divided into nine groups, FMLC recognition accuracy is maximized to 89.81%. When compared to MLP, MMLP improves FMIC identification accuracy by over 13%. When comparing MMLP to MLP, FMLC shows a greater than 11% improvement.

CONCLUSION

Diagonal-based feature1 is determined to be the simplest, most efficient, and best (97.6% for 54 (k = 11) features) of the features employed in the works for Malayalam HCR. All the approaches utilizing a single feature (RLC / Gradient) are observed to have a lower recognition accuracy with a fewer number of blocks. On the other hand, by fusing these traits together, we may improve identification accuracy even with a reduced number of blocks. Using a combination of characteristics from many sources is shown to be the most effective strategy for improving recognition accuracy. We found that SQDF for GBF - RLC with 147 (k = 58) features had the highest recognition rate (99.66%). Based on our experimental findings using MLP, we can conclude that the Malayalam HCR system has the maximum recognition accuracy (99.78%) while employing the GBF - RLC. Using the same collection of characteristics and the MMLP classifier, we can get a 99.86% recognition rate. QDF (98.81%) and MQDF (99.37%) classifiers achieve high rates of recognition accuracy. Out of all the published works.

REFERENCES

1. Arica, N. and F.T. Yarman-Vural, An overview of character recognition focused on off-line handwriting. Systems, Man, and Cybernetics, Part C: Applications and Reviews, IEEE Transactions on, 2001. 31(2): p. 216-233.
2. B V Dhandra et al., (2010) B. V. Dhandra, Mallikarjun Hangarge, Gururaj Mukarambi, —Spatial Features for Handwritten Kannada

3. C Leila, (2011) Chergui Leila, Kef Maâmar, Chikhi Salim, —Combining Neural Networks for Arabic Handwriting Recognition, Programming and Systems (ISPS), 2011 10th International Symposium on, pp.74-79.
4. Das, N., et al. A Novel GA-SVM Based Multistage Approach for Recognition of Handwritten Bangla Compound Characters. in Proceedings of the International Conference on Information Systems Design and Intelligent Applications (INDIA 2012) held in Visakhapatnam, India, January 2012. Springer.
5. G L Prajapati and A Patle, (2010) Gend Lal Prajapati , Arti Patle , —On Performing Classification using SVM with Radial Basis and Polynomial Kernel Functions, Third International Conference on Emerging Trends in Engineering and Technology , 978-0-7695-4246-1/10, 2010 IEEE, doi: 10.1109/ICETET.2010.134.
6. H Bunke and P S P Wang, (1997) H Bunke, P S P Wang, —Handbook of Character Recognition and Document Image Analysis, World Scientific Publishing Company, pp 49 — 78, 1997.
7. R. Jayadevan, S. Kolhe, P. Patil and U. Pal, "Offline Recognition of Devanagari Script: A Survey", IEEE Transactions on systems, man and cybernetics - part - c: Applications and reviews, Vol.41, No.6, November 2011.
8. Cheng- Lin Liu, S. Jaeger and M. Nakagawa, —Online Recognition of Chinese Characters: The State — Of — the Art, IEEE Transactions on Patten Analysis and Intelligence, Vol.26, No.2, pp. 198-213, February 2004.
9. A.W. Senior and J. Robinson, —An Offline Cursive Handwriting Recognition System, IEEE Transactions on Patten Analysis and Machine Intelligence, Vol.20, No.3, pp. 309- 321, March 2015.
10. B. Wegmann and C. Zetzshe, —Feature — Specific Vector Quantization Of Images, IEEE Transactions on Image Processing, Vol.5, No.2, pp. 274-288, February 2000.
11. Luiz Oliveria, F. Bortolozzi and R. Sabourn, —Automatic Recognition of Handwritten Numeral Strings: A Recognition and Verification Strategy, IEEE Transactions on Patten Analysis and Machine Intelligence, Vol.24, No.11, pp. 1438-1453, November 2002.
12. B.B. Chaudhuri and Sumedha Bera, "Handwritten Text Line Identification In Indian Scripts", Proceedings of the tenth international conference on Document Analysis and Recognition, Vol. No 2, pp.636- 640, October 2009.