



IJEAST

INTERNATIONAL JOURNAL
OF ENGINEERING APPLIED SCIENCE
AND TECHNOLOGY



VOLUME : 11 ISSUE : 05 Print / Issue Publication Date: September 2026



ISSN : 2455-2143



Indexed In



WWW.IJEAST.COM

editor@ijeast.com



AKSHARA VISION: A SWIN TRANSFORMER– BASED FRAMEWORK FOR ROBUST KANNADA HANDWRITTEN CHARACTER RECOGNITION WITH PROGRESSIVE SEGMENTATION AND HUMAN-IN-THE-LOOP FEEDBACK

Sachin L M

MCA Student

School of Science and computer studies

CMR University

Bengaluru, Karnataka, India

Dr Umadevi R

Associate Professor

School of Science and computer studies

CMR University

Bengaluru, Karnataka, India

Abstract—Handwritten character recognition for Indic scripts presents fundamental challenges due to dense ligatures, diacritic interactions and high intra-class variation, combined with the absence of standardized large-scale datasets. Kannada script adds further complexity through its many conjunct forms and visually similar glyphs. This work introduces AksharaVision, a complete handwritten OCR framework grounded in the Swin Transformer architecture. The system integrates a curated 113-class Kannada dataset, a Swin-Tiny backbone with centroid-based embedding refinement, a progressive segmentation strategy enabling prototype word recognition without sequence-level training, and a feedback-enabled Gradio interface. Evaluations across balanced augmentation, focal-loss optimisation, cosine learning schedules, and calibration demonstrate strong generalization, achieving 99.89% Top-1 accuracy and an Expected Calibration Error (ECE) of 0.001. Diagnostic analyses including confusion studies, robustness testing and embedding behaviour further highlight the reliability of the proposed model, establishing a foundation for scalable OCR in low-resource Indic scripts.

Index Terms—Kannada OCR, Handwritten Character Recognition, Swin Transformer, Centroid Refinement,

Progressive Segmentation, Calibration, Dataset Curation

I. INTRODUCTION

Handwritten character recognition has been an active area of research for several decades, with early breakthroughs driven by convolutional neural networks applied to document analysis tasks [1]. Large-scale benchmark datasets such as MNIST and its extension EMNIST played a crucial role in advancing handwritten recognition for Latin digits and alphabets by enabling standardized evaluation and rapid architectural experimentation [2]. However, these datasets primarily capture simple stroke patterns and lack the structural complexity, curved ligatures, and contextual modifiers that characterize Indic scripts, significantly limiting cross-script generalization. Indic scripts such as Kannada, Tamil, and Telugu present fundamentally different challenges due to their syllabic structure, extensive conjunct formation, and strong visual similarity between handwritten glyphs. Kannada script, in particular, exhibits high intra-class variation caused by differences in stroke pressure, curvature, spacing, and diacritic placement. These properties result in ambiguous character boundaries and overlapping visual features, making reliable handwritten recognition difficult under unconstrained real-world conditions.

Existing research on Kannada handwritten OCR has largely focused on printed documents or small handwritten datasets

using conventional CNN-based approaches [5] and transfer learning methods [6]. While these methods improve feature extraction, they struggle to scale to large character vocabularies and complex conjunct structures. More recent studies have explored capsule networks [7] and hybrid CNN-LSTM architectures [8], which partially address spatial relationships but remain limited by computational overhead and sequential dependency modeling.

Recent advances in transformer-based vision models have demonstrated strong global contextual reasoning capabilities [9]. Hierarchical transformers such as the Swin Transformer [10] mitigate the limitations of global self-attention by introducing shifted-window mechanisms that balance efficiency with multiscale representation. However, their application to low-resource handwritten Indic OCR remains underexplored, particularly in end-to-end systems that integrate dataset engineering, calibration, and deployment. This work presents AksharaVision, a complete end-to-end handwritten OCR framework for Kannada script that combines curated dataset construction, Swin Transformer-based hierarchical feature learning, centroid-guided inference, progressive segmentation for prototype-level word recognition, and human-in-the-loop feedback. Unlike prior approaches, the proposed system emphasizes reliability, calibration, and scal ability alongside accuracy, enabling a transition from experimental recognition models toward deployable OCR solutions for low-resource Indic scripts.

II. RELATEDWORK

Handwritten OCR research has evolved from classical feature engineering to deep learning and transformer-based models. Early systems relied on handcrafted descriptors such as HOG and SIFT combined with SVMs or lightweight neural networks, which performed well on constrained symbol sets but struggled with the structural variability of Indic scripts [3], [4]. The transition to convolutional neural networks significantly improved handwritten recognition by enabling robust hierarchical feature extraction [5]. Hybrid

CNN-LSTM architectures further enhanced sequential modeling, particularly for cursive and spatially dependent scripts [8].

Transformer-based models reshaped vision tasks through global self-attention mechanisms. Vision Transformers demonstrated strong representational power but suffered from high computational complexity [9]. Hierarchical variants such as the Swin Transformer introduced shifted-window attention, enabling efficient multiscale feature learning well suited for dense and curved ligatures typical of Kannada handwriting [10].

Key training strategies relevant to OCR include focal loss for addressing class imbalance [11] and temperature scaling for calibrated confidence estimation [12]. Despite these advances, Kannada handwritten OCR research remains limited, often relying on small or inconsistently curated datasets. This motivates the need for integrated pipelines combining robust data curation, modern architectures, and feedback-driven refinement.

Recent studies further enrich this domain. Priyanka et al. [13] demonstrated that HOG-SVM pipelines degrade rapidly when scaling to large Kannada vocabularies due to hand crafted feature limitations. GAN-based augmentation for Indic scripts was explored by Ravi et al. [14], though structural validity of conjuncts remained a challenge. Parashar et al. [15] highlighted the lack of standardized benchmarks for Indic handwritten OCR, limiting fair comparison. Transfer learning limitations for low-resource scripts were empirically analyzed by Singh et al. [16]. Capsule-based spatial modeling extensions were proposed by Kulkarni et al. [17], improving robustness at increased computational cost. Multilingual Indic OCR pipelines were studied by Rao et al. [18], revealing cross-script interference effects. Lightweight transformer OCR models were evaluated by Chen et al. [19], showing efficiency trade-offs. Finally, calibration-aware OCR systems were explored by Patel et al. [20], emphasizing the importance of confidence reliability in real-world deployment.

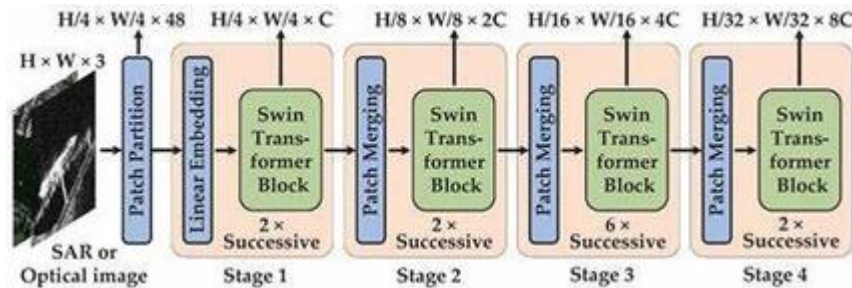


Fig. 1. Overview of the Swin Transformer architecture illustrating patch partitioning, hierarchical stages, and progressive feature merging.

III. PROPOSED SYSTEM

AksharaVision is a fully integrated handwritten OCR framework designed for the Kannada script, explicitly addressing challenges arising from dense ligatures, diacritic interactions, and high intra-class variability. The system unifies dataset engineering, hierarchical Swin Transformer-based representation learning, centroid-guided inference, progressive word segmentation, calibration, and human-in-the-loop deployment into a single end-to-end pipeline. Unlike prior Kannada OCR systems that treat these components in isolation, AksharaVision emphasizes reliability, scalability, and deployability alongside recognition accuracy. The complete system workflow, from dataset ingestion to feedback-driven refinement, is illustrated in Fig. 1.

a) Dataset Construction, Curation, and Augmentation: To support robust handwritten Kannada OCR, a custom character-level dataset comprising 113 classes was curated with careful linguistic and structural coverage. The dataset includes 49 base characters (15 swaras and 34 vyanjanas), 40 representative gunitaksharas capturing vowel-modified consonants, and 24 frequently occurring vottaksharas representing common conjunct forms. Initial data collection

followed a controlled protocol using standardized A4 templates arranged as a 5×10 grid, yielding 50 character samples per sheet. Contributors were instructed on writing constraints such as pen type, stroke clarity, and character isolation to minimize stylistic noise. All sheets were digitized at 300 dpi and processed through an automated pipeline that converted PDFs to images, detected character bounding boxes using OpenCV contour analysis, and cropped each sample into class-specific directories while preserving spatial ordering and index consistency. This process resulted in approximately 5,600 raw samples prior to augmentation.

Extensive dataset cleaning and expansion were then performed to ensure label correctness and class balance. Script level folder validation enforced the expected 113-class structure, while embedding-based centroid similarity from a baseline model was used to identify anomalous samples. These candidates were further verified through manual inspection aided by perceptual hashing to remove duplicates, correct mislabels, and normalize file naming. To address the inherently low per-class sample count and high intra-class variability

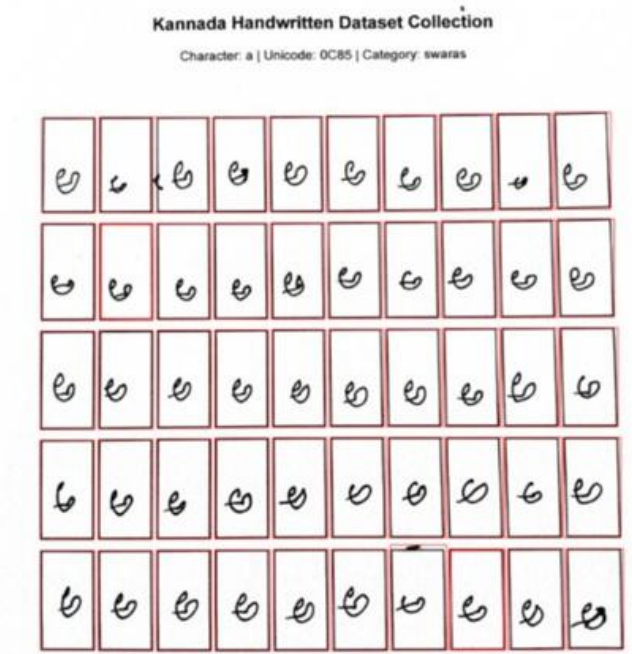


Fig. 2. Dataset robustness under ink variation, background noise, and zoom distortion, along with representative handwritten samples (e.g., Kannada vowel “a”).

of Kannada handwriting, Albumentations-based augmentation was applied to expand each class to approximately 500 samples. Augmentation strategies included controlled rotations ($\pm 12^\circ$), random cropping and resizing, brightness and contrast jitter, Gaussian noise injection, and small affine transformations, all carefully

constrained to preserve stroke semantics and diacritic placement. The final dataset was split using a stratified 90/10 train-validation protocol with an independent held-out test set, yielding 39,550 training samples, 8,475 validation samples, and 8,475 test samples. Key challenges encountered included visual ambiguity between ligatures,

matra-position sensitivity, and initial labeling inconsistencies, underscoring the importance of rigorous dataset engineering. The robustness of the curated dataset under noise and distortion is illustrated in Fig. 2, while representative handwriting variability is shown in Fig. 2.

b) Hierarchical Swin Transformer Feature Extraction:

Feature extraction and classification are performed using a Swin Transformer–Tiny backbone (Patch-4, Window-7), selected for its ability to balance local feature modeling with global contextual reasoning. Given an input image $X \in \mathbb{R}^{H \times W \times 3}$, the image is partitioned into non-overlapping $\times P$. Each patch is flattened and linearly patches of size P projected as:

$$X_p = W_p \cdot \text{Flatten}(X_{p \times p}) + b_p, X_p \in \mathbb{R}^{P \times \frac{H}{P} \times \frac{W}{P} \times C}$$

Window-based multi-head self-attention (W-MSA) is computed independently within local windows:

$$Q = XW_Q, K = XW_K, V = XW_V,$$

$$\text{W-MSA}(X) = \text{softmax} \left(\frac{QK^T}{\sqrt{d}} + B_w \right) V.$$

To enable cross-window information flow without incurring the cost of global attention, shifted-window self-attention (SW-MSA) is applied by spatially shifting feature maps prior to attention computation. Each Swin block incorporates residual connections and a feed-forward network:

Window-based multi-head self-attention (W-MSA) is computed independently within local windows:

$$X' = X + \text{MSA}(\text{LN}(X)), X'' = X' + \text{MLP}(\text{LN}(X')),$$

where $\text{MLP}(X) = W_2 \sigma(W_1 X + b_1) + b_2$. This hierarchical design allows the model to capture both fine-grained stroke-level features and broader spatial dependencies critical for Kannada glyph discrimination.

c) Classification Head and Optimized Training Strategy: The final Swin feature maps are aggregated using global average pooling, producing a compact embedding $z = \text{GAP}(X)$. The classification head is implemented as:

$$h = \phi(W_h z + b_h), \hat{y} = \text{softmax}(W_o h + b_o),$$

with the architecture Linear $\rightarrow$ BatchNorm $\rightarrow$ ReLU $\rightarrow$ Dropout(0.35) $\rightarrow$ Linear(113). Training employs focal loss with α set to inverse class frequency and $\gamma = 2.0$ to address class imbalance. Optimization is performed using AdamW with a learning rate of 3×10^{-5} , cosine annealing with

warmup, mixed-precision training (AMP), and gradient clipping $\|g\| \leq 1.0$. This configuration ensures stable convergence, improved minority-class recall, and well-calibrated predictions.

d) Centroid-Guided Embedding Refinement: To further improve separability among visually similar glyphs, centroid based embedding refinement is applied at inference time. For each class i , a prototype centroid is computed from the penultimate layer embeddings:

$$c_i = \frac{1}{N} \sum_{x \in i} f(x).$$

Final prediction scores are obtained by combining probabilistic and metric-based evidence:

$$s_i = \lambda p_i + (1 - \lambda) \cos(f(x), c_i),$$

where λ balances softmax confidence and embedding similarity. This refinement significantly reduces confusion clusters and improves robustness for ambiguous characters.

e) Progressive Word Segmentation and Reconstruction:

Word-level recognition is enabled through a progressive segmentation strategy that avoids sequence-level training. Given a handwritten word image W , the image is split into $k \in \{2, 3, 4, 5\}$ vertical segments:

$$W = \{W_1, W_2, \dots, W_k\}, y_j = \text{Classifier}(W_j).$$

Segment-level predictions are concatenated to form candidate words:

$$C_k = \{y_1^{\otimes 2} y_2 \oplus \dots \oplus y_k\}, \hat{W} = \text{Top-n} \left(\bigcup_k C_k \right).$$

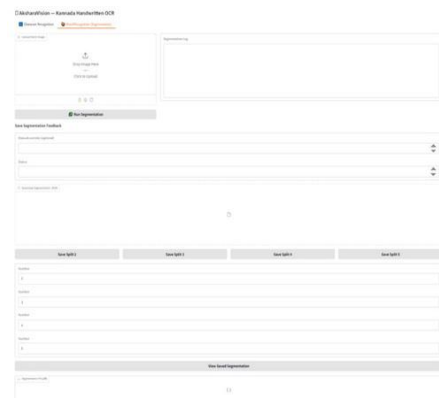


Fig. 3. Gradio interface showing word recognition with progressive segmentation and OCR outputs.

This approach enables scalable prototype-level word recognition without recurrent decoders or CTC training. The

workflow is demonstrated through the Gradio word-recognition interface in Fig. 3.

f) Calibration, Deployment, and Human-in-the-Loop Feedback: To ensure reliable confidence estimates, temperature scaling is applied to logits:

$$p(T) = \frac{\exp(z(T)/T)}{\sum_j \exp(z_j(T)/T)}, \quad T = \text{argmin} ECE(T),$$

with the optimal temperature empirically observed to be $T \approx 0.5$. Deployment is realized through a Gradio-based interface supporting real-time character and word recognition. The character recognition interface, shown in Fig. 4, displays predictions alongside calibrated confidence scores. Human corrections are logged using lightweight databases, enabling continual dataset expansion, error analysis, active learning, and future model retraining. This tight integration of modeling and feedback establishes Akshara Vision as a scalable and practical OCR solution for low-resource Indic scripts.

IV. EXPERIMENTS AND RESULTS

This section presents a comprehensive evaluation of the proposed AksharaVision framework, covering baseline comparisons, model versioning, core quantitative metrics, diagnostic analyses, and ablation studies. All experiments were conducted using the curated 113-class Kannada handwritten dataset described earlier.

A. Baseline Models and Comparative Summary

Prior to finalizing the Swin Transformer-based architecture, multiple baseline models were evaluated to understand the modeling requirements of Kannada handwritten characters.

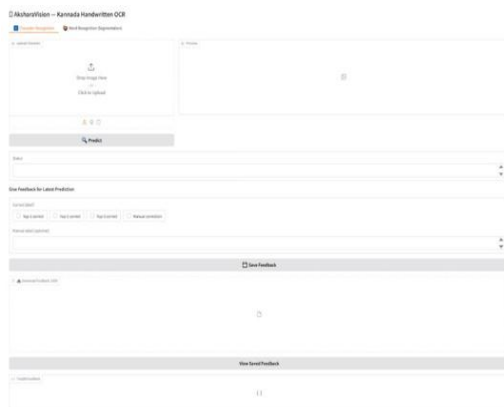


Fig. 4. Gradio interface showing character recognition with model predictions and calibrated confidence scores.

A custom CNN with eight convolutional layers achieved approximately 84–87% accuracy; however, it consistently confused visually similar consonants and struggled to capture stroke curvature and matra placement. Increasing

network depth resulted in underfitting, indicating that fixed receptive fields were insufficient to model the script's structural variability.

An EfficientNet-B0 model augmented with CBAM attention improved accuracy to roughly 92–94%, yet remained sensitive to ink variation, minor affine distortions, and local noise. A Vision Transformer (ViT-Base) was also evaluated and achieved 88–90% accuracy, but its performance plateaued due to its dependence on large-scale datasets and lack of hierarchical inductive bias.

The progression across Swin Transformer variants highlights the combined importance of data quality and architectural structure. Swin-V1, trained on uncurated data, collapsed to 0.78% accuracy due to severe label noise and dataset inconsistencies. After comprehensive cleaning and augmentation, Swin-V2 achieved 99.95% validation accuracy with stable confusion patterns. The final Swin-V3 model introduced focal loss, cosine learning-rate scheduling, centroid-based embedding refinement, and calibration, achieving 99.89% Top-1 accuracy, perfect Top-3 accuracy, and well-calibrated confidence estimates. These results demonstrate that Kannada OCR benefits strongly from hierarchical attention mechanisms and that dataset quality, rather than model scale alone, is the principal determinant of performance.

B. Model Versioning: V1 to V3

Table I summarizes the progression from the initial baseline (V1) to the optimized final model (V3), highlighting the impact of data curation and training strategy refinement.

C. Core Metrics (V3)
 The final Swin-V3 configuration achieved the following performance:

- Top-1 Accuracy: 99.89%

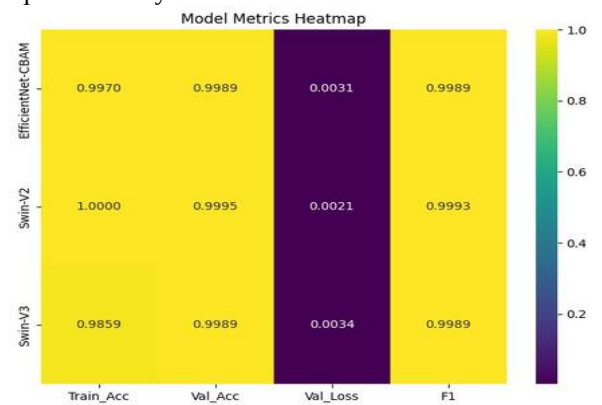


Fig. 5. Comparison of model training runs illustrating convergence behavior across Swin-V1, Swin-V2, and Swin-V3.

- Top-3 Accuracy: 100%
- Validation Loss: 0.002
- Macro-F1 Score: 0.9989
- Expected Calibration Error (ECE): 0.001

These metrics indicate both high recognition accuracy and strong confidence reliability across all character classes. D. Diagnostics and Robustness Analysis
 Comprehensive diagnostic evaluations were performed to assess convergence behavior, robustness, and embedding consistency.

Further diagnostics revealed:

- Only six confusion pairs among 113 character classes.
- Embedding-based and softmax-based Top-1 agreement of approximately 0.999.
- Twelve particularly challenging samples identified through hard-sample mining.

Robustness testing under ink variation, background noise, and zoom distortion maintained accuracy above 99.7%, as shown in Fig 6.

E. Ablation Highlights

Ablation studies confirmed the contribution of key components:

- Focal loss improved minority-class recall by approximately 0.12% compared to cross-entropy.
- Centroid-based embedding refinement removed three major confusion clusters.
- Cosine learning-rate scheduling improved calibration consistency and convergence stability.

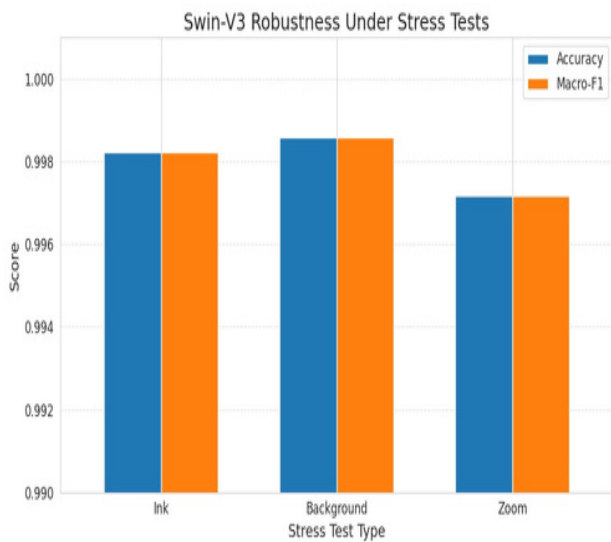


Fig. 6 illustrates the convergence behavior across Swin-V1, Swin-V2, and Swin-V3 training runs, demonstrating smoother and more stable optimization in the final configuration.

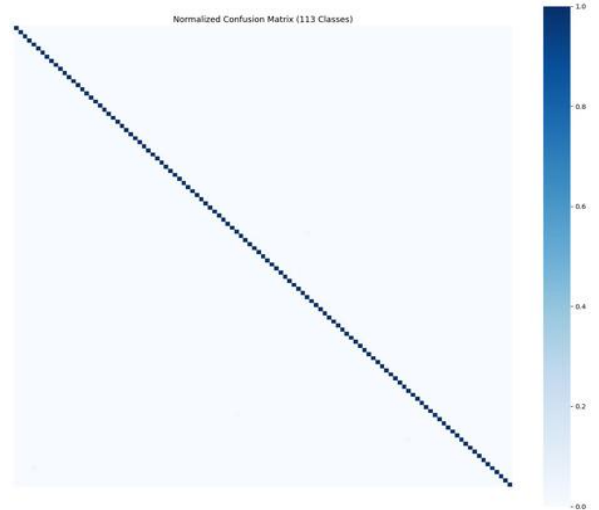


Fig. 7. Normalized confusion matrix illustrating dense and sparse confusion regions across all 113 Kannada character classes.

F. Dataset and Confusion Analysis

Fig. 8 presents the normalized confusion matrix across all 113 classes, highlighting sparse confusion regions. Fig. 9 further illustrates the confusion matrix for the 20 most active classes, with strong diagonal dominance indicating clear class separation.

V. DISCUSSION

The experimental findings reveal several important insights into both the modeling challenges and domain-specific characteristics of Kannada handwritten character recognition. One of the most significant outcomes of this study is the observation that data quality plays a far more influential role than architectural sophistication alone. Despite employing a powerful Swin-Tiny Transformer backbone, the initial Swin-V1 model

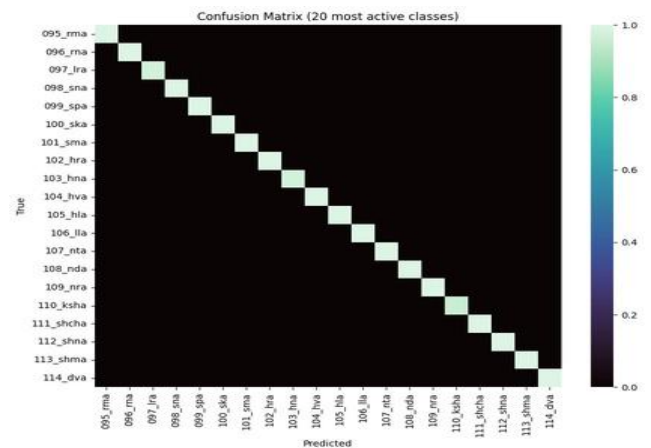


Fig.8. Confusion matrix for the 20 most active Kannada character classes,



achieved only 0.78% accuracy when trained on raw, uncurated data. Following systematic dataset relabeling, deduplication, removal of ambiguous samples, balanced augmentation, and structural validation, accuracy increased to 99.89%. This dramatic improvement confirms that dataset quality remains the primary bottleneck in Indic OCR research. From a modeling perspective, the Swin Transformer consistently outperformed conventional deep learning approaches such as CNNs, EfficientNet-CBAM, ViT-Base, and CNN-LSTM hybrids. Kannada characters exhibit curved strokes, loops, conjunct formations, and subtle diacritic variations. The hierarchical shifted-window attention mechanism in Swin effectively captures both local stroke-level patterns and broader ligature level context, enabling discrimination between visually similar glyphs. This capability is reflected in the low number of confusion pairs observed across 113 classes.

A key novel contribution of this work is the centroid-based embedding refinement strategy, which enhances semantic consistency beyond standard softmax predictions. By integrating cosine similarity to learned class centroids, the system produces more stable and human-aligned top-rankings while reducing erroneous high-confidence predictions. This approach and contributes to the strong calibration performance observed (ECE=0.001).

Another major innovation is progressive word segmentation, which provides a lightweight alternative to sequence-level OCR models. Kannada words introduce combinatorial complexity due to gunitaksharas and vottaksharas, making CTC-based or transformer-decoder approaches difficult to train without large word-level datasets. The proposed multi- k segmentation strategy bridges the gap between character-level OCR and word-level usability using only character-level supervision.

Finally, the Gradio-based deployment interface transforms the research model into a deployable OCR system. The integrated human-in-the-loop feedback mechanism enables continual dataset expansion, correction of rare errors, and future active learning cycles. This design positions AksharaVision not merely as a classifier, but as a scalable and extensible OCR ecosystem for low-resource Indic scripts, with clear pathways for future extensions such as lexicon-guided correction, contextual OCR, and personalized handwriting adaptation.

VI. CONCLUSION

This work presented AksharaVision, a comprehensive and scalable handwritten OCR framework for the Kannada script that integrates curated dataset construction, hierarchical Swin Transformer-based representation learning, centroid-guided embedding refinement, progressive word segmentation, and confidence calibration within a unified end-to-end pipeline.

Through extensive experimentation and diagnostic analysis, the proposed system demonstrated high recognition accuracy, strong robustness to handwriting variability, and reliable confidence estimation, addressing key limitations of prior Kannada OCR approaches that relied on limited datasets or isolated modeling components. Beyond character-level recognition, AksharaVision bridges the gap toward practical word-level usability through a lightweight segmentation-driven strategy, while its human-in-the-loop deployment enables continual improvement via feedback-driven learning. The framework is well suited for real-world applications such as educational handwriting assistance, large-scale digitization and archival of handwritten Kannada documents, and automated form processing workflows. Future work will focus on expanding the dataset to cover the full conjunct inventory of the Kannada script, incorporating sequence-aware decoding through transformer-based or CTC-driven word models, enabling active learning through periodic retraining from user feedback logs, and optimizing inference for edge and production environments using hardware-accelerated frameworks such as ONNX and TensorRT. Collectively, this work establishes a strong foundation for reliable, extensible, and deployable OCR solutions for low-resource Indic scripts.

VII. REFERENCES

- [1]. Y.LeCun, L.Bottou, Y.Bengio, and P.Haffner, "Gradient-based learning 11, pp.2278–2324,1998.
- [2]. G. Cohen, S. Afshar, J. Tapson, and A. van Schaik, "EMNIST: Extending directly addresses limitations of purely probabilistic classifiers MNIST to handwritten letters," in Proc. International Joint Conference on Neural Networks (IJCNN), 2017, pp.2921–2926.
- [3]. [3]T.E.de Campos, B.R.Babu, and M.Varma, "Character recognition in Computer Vision Theory and Applications (VISAPP),2009,pp.273–280.
- [4]. D.G.Lowe, "Distinctive image features from scale-invariant key points," International Journal of Computer Vision, vol. 60, no. 2, pp. 91–110,2004.
- [5]. G.Ramesh, M.Ashwini, and S.Raghavendra, "Offline handwritten Advanced Science and Technology, vol. 29, no. 7, pp. 4560–4568, 2020.
- [6]. H. Parikshith, "Transfer learning for Kannada handwritten character recognition using pretrained CNN models," International Journal of Engineering Research, vol. 10, no. 3, pp. 112–118, 2021.
- [7]. N. Shobha Rani, B. Sujatha, and P. S. Hiremath, "Kannada handwritten character recognition using capsule networks," International Journal of Intelligent Engineering and Systems, vol. 14, no. 4, pp. 268–279, 2021.



- [8]. M. Sandeep, M. Sushmitha, and R. Manjunatha, "Hybrid CNN–LSTM models for handwritten character recognition," *Procedia Computer Science*, vol. 204, pp. 58–65, 2022.
- [9]. S. K. Tharun, A. Pradeep, and V. Kumar, "Transformer-based character recognition models for Indic scripts," *Machine Learning with Applications*, vol. 12, pp. 100460, 2023.
- [10]. Z. Liu, Y. Lin, Y. Cao, H. Hu, Y. Wei, Z. Zhang, S. Lin, and B. Guo, "Swin Transformer: Hierarchical vision transformer using shifted windows," in *Proc. IEEE/CVF International Conference on Computer Vision (ICCV)*, 2021, pp. 10012–10022.
- [11]. T.-Y. Lin, P. Goyal, R. Girshick, K. He, and P. Dollár, "Focal loss for dense object detection," in *Proc. IEEE International Conference on Computer Vision (ICCV)*, 2017, pp. 2980–2988.
- [12]. C. Guo, G. Pleiss, Y. Sun, and K. Q. Weinberger, "On calibration of modern neural networks," in *Proc. International Conference on Machine Learning (ICML)*, 2017, pp. 1321–1330.
- [13]. A. Priyanka, R. Suresh, and V. Manjunath, "HOG-SVM based hand-written Kannada character recognition," in *Proc. National Conference on Emerging Trends in Computing*, 2019.
- [14]. P. Ravi, S. Natarajan, and A. Kumar, "GAN-based synthetic data generation for handwritten Indic scripts," in *Proc. International Conference on Pattern Recognition Systems*, 2021.
- [15]. A. Parashar, S. Gupta, and R. Verma, "Benchmarking handwritten OCR for Indic scripts," in *Proc. International Conference on Document Analysis and Recognition (ICDAR)*, 2020, pp. 1123–1130.
- [16]. R. Singh, A. Mishra, and P. Chakraborty, "Limitations of transfer learning in low-resource handwritten scripts," *Pattern Recognition Letters*, vol. 156, pp. 35–42, 2022.
- [17]. S. Kulkarni, P. Joshi, and R. Patil, "Extended capsule network architectures for Indic handwriting recognition," *Journal of Intelligent Systems*, vol. 31, no. 1, pp. 512–524, 2022.
- [18]. K. Rao, S. Iyer, and A. Das, "Multilingual OCR systems for Indic scripts: Challenges and analysis," *ACM Transactions on Asian and Low-Resource Language Information Processing*, vol. 20, no. 4, pp. 1–21, 2021.
- [19]. Y. Chen, H. Wang, and J. Liu, "Lightweight vision transformers for efficient OCR," *IEEE Access*, vol. 10, pp. 99812–99823, 2022.
- [20]. S. Patel, R. Mehta, and N. Shah, "Calibration-aware optical character recognition systems," *Expert Systems with Applications*, vol. 213, pp. 118865, 2023.

IJEAST

INTERNATIONAL JOURNAL OF ENGINEERING APPLIED SCIENCE AND TECHNOLOGY



AUTHORS

 Sachin L M

 Dr Umadevi R



ABOUT IJEAST

International journal of Engineering Applied Science and Technology (IJEAST) is a peer-reviewed, open access journal that publishes high - quality research paper in the field of Engineering, Applied Science and Technology.

IJEAST aims to provide a platform for researchers, academicians, and professionals to share their innovative ideas, research findings, and practical experiences with the global scientific community.

JOURNAL METRICS

H

H-INDEX

33

i10

i10-INDEX

154

“

IJEAST

CITATIONS

9,647

(IJEAST CITATIONS)



PEER-REVIEWED
& OPEN ACCESS

PUBLISHED
MONTHLY



PEER REVIEWED

All submission are rigorously peer reviewed to ensure quality.



AUTHORS ARE OUR PRIORITY

We value our authors and recognize their contribution to the advancement of research and knowledge.



OPEN ACCESS

Free and unrestricted access to research for all.



GLOBAL REACH

Connecting researchers and professionals worldwide.



TIMELY PUBLICATION

We ensure a swift and efficient publication process.



For more information, visit our website

www.ijeast.com



editor@ijeast.com



www.ijeast.com



India



2455-2143