



Journal Homepage: [-www.journalijar.com](http://www.journalijar.com)

INTERNATIONAL JOURNAL OF ADVANCED RESEARCH (IJAR)

Article DOI: 10.21474/IJAR01/23496

DOI URL: <http://dx.doi.org/10.21474/IJAR01/23496>



RESEARCH ARTICLE

AI-POWERED MULTILINGUAL OCR SYSTEM FOR DIGITIZATION OF HISTORICAL HANDWRITTEN AND REGISTERED DOCUMENTS IN REGIONAL INDIAN LANGUAGES

N.B. Mahesh Kumar^{1,2}, R.Deepak Kumar^{1,2}, R.G.Jayarajan^{1,2}, D.Mano^{1,2} and V.Madhivanan^{1,2}

1. Department of Artificial Intelligence and Data Science.
2. Hindusthan Institute of Technology, Coimbatore.

Manuscript Info

Manuscript History

Received: 12 March 2026

Final Accepted: 14 April 2026

Published: May 2026

Key words:-

AI-Powered OCR, Convolutional Neural Networks (CNN), Handwritten Text Recognition (HTR), Regional Indian Languages, Cultural Heritage Document Preservation

Abstract

The digitization of historical handwritten documents represents a critical challenge in preserving cultural heritage and improving public access to governmental and legal records. This report presents a comprehensive technical specification for an AI-powered Optical Character Recognition (OCR) system specifically engineered to digitize handwritten and aged registered documents in regional Indian languages including Tamil, Telugu, Kannada, Malayalam, and Hindi. The proposed system leverages state-of-the-art deep learning architectures including Convolutional Neural Networks (CNNs) and Transformer-based sequence models combined with language-specific pre-processing pipelines to achieve high accuracy text extraction from degraded physical documents. The solution addresses the pressing need for accessible historical records in government offices, land registration departments, courts, and public archives across India. Key performance targets include character recognition accuracy exceeding 95% for printed regional scripts and above 88% for handwritten text. The system is designed to be deployable as a web and mobile platform with an intuitive interface for non-technical government staff, integrating output into searchable, indexed digital repositories.

"© 2026 by the Author(s). Published by IJAR under CC BY 4.0. Unrestricted use allowed with credit to the author."

Introduction:-

Background and Motivation:-

India holds a vast repository of historical documents spanning centuries regarding land deeds, court records, marriage registrations, census data, and most of which exist only as physical handwritten manuscripts stored in government archives and registry offices. These documents are written in various regional languages using scripts that differ substantially from the Latin alphabet and pose significant challenges to standard OCR tools built primarily for English-language content. As of 2024, an estimated 2.3 billion pages of registered documents remain in purely physical form across Indian state archives (Ministry of Electronics and Information Technology, 2023). The degradation of these documents due to age, humidity, and poor storage conditions creates urgency for digital preservation. The National e-Governance Plan (NeGP) and various Digital India initiatives have highlighted

Corresponding Author:-N.B. Mahesh Kumar

Address:-1. Department of Artificial Intelligence and Data Science. 2. Hindusthan Institute of Technology, Coimbatore.

document digitization as a priority, yet a technically robust, regionally aware OCR solution does not yet exist at scale.

Problem Statement:-

How might we develop an AI or OCR solution to digitize and convert handwritten, old registered documents into a readable and accessible format in regional languages, improving public access and readability of historical records? Traditional OCR systems (e.g., Tesseract, ABBYY) are not adequately trained on historical Indian scripts, handwritten regional text, or document artifacts common in aged parchment. Government offices rely on manual transcription, which is slow, costly, error-prone, and inaccessible to persons with disabilities or those living remotely.

Research Objectives:-

- Design and develop a deep learning-based OCR system capable of recognizing handwritten text in Tamil, Telugu, and Kannada, Malayalam, and Hindi scripts.
- Build pre-processing pipelines to handle image enhancement, noise removal, and binarization for aged and degraded documents.
- Implement post-processing with language models and dictionaries to correct OCR errors specific to regional vocabulary.
- Develop a user-friendly web and mobile interface for uploading, processing, and retrieving digitized documents.
- Create a searchable digital repository that supports keyword indexing in regional languages.
- Achieve character recognition accuracy $\geq 95\%$ for printed scripts and $\geq 88\%$ for handwritten text in target languages.

Literature Review:-

The literature survey in table 1 highlights recent advancements in OCR technologies for multilingual, handwritten, and historical document digitization, with a strong focus on Indian regional languages. Various deep learning approaches, including CNNs, Vision Transformers, TrOCR, and script-independent recognition models, have significantly improved text recognition accuracy and document processing efficiency. However, challenges such as degraded historical manuscripts, low-resource languages, and unified multilingual support remain, creating opportunities for developing an AI-powered multilingual OCR system for historical Indian documents.

Table 1. Comparative Analysis of Existing OCR Methods

Study and Year	Study Type	Population / Document Type	Languages / Scripts	Core Method	Outcome Measures
Gaikwad et al. (2025)	Research Article	Ancient Indian manuscripts and heritage texts	Sanskrit, Devanagari, regional scripts	AI-based OCR with script-specific preprocessing	Improved digitization accuracy and preservation of linguistic heritage
Sengar (2025)	Research Article	Multilingual handwritten documents	Multiple Indian scripts	CLIP + Tesseract OCR integration	Enhanced multilingual handwritten text recognition
Li et al. (2021)	Research Article	Printed and handwritten text datasets	Multilingual	Transformer-based OCR (TrOCR)	State-of-the-art recognition accuracy and robustness
Gurung et al. (2025)	Research Article	Devanagari handwritten documents	Devanagari	CNN-based handwritten text recognition	High recognition accuracy for handwritten characters
Gupta et al. (2024)	Conference Paper	Printed Indian language	Multiple Indian	CMViT (Convolutional	Improved printed text recognition

Study and Year	Study Type	Population Document Type /	Languages / Scripts	Core Method	Outcome Measures
		documents	languages	Multiscale Vision Transformer)	performance
MeitY (2023)	Government Report	Digital archival and e-governance documents	Indian regional languages	National digitization initiatives and OCR adoption	Increased accessibility and digital preservation efforts
Rabby et al. (2024)	Conference Paper	Bengali printed and handwritten documents	Bengali	Specialized OCR models with advanced enhancement techniques	Reduced OCR errors across diverse document types
Tulsyan et al. (2024)	Conference Paper	Printed Indian language documents	Multiple Indian languages	IndicOCR pipeline with language-specific processing	Improved recognition efficiency and scalability
Dey et al. (2025)	Research Article	Low-resource handwritten language datasets	Low-resource regional languages	Deep learning-based Handwritten Text Recognition (HTR)	Better recognition in scarce-data environments
Gohil& Desai (2026)	Review Paper	Image and video text datasets	Multilingual	Machine Learning and Deep Learning OCR techniques	Comprehensive analysis of OCR advancements
Romein et al. (2025)	Research Article	Historical handwritten manuscripts	Historical European documents	Advanced HTR engine benchmarking	Comparative evaluation of recognition accuracy
Anju et al. (2024)	Research Article	Unconstrained handwritten documents	Malayalam	Advanced segmentation techniques with OCR	Improved segmentation and OCR performance
Chauhan et al. (2024)	Research Article	Handwritten character datasets	Script-independent	HCR-Net deep learning architecture	Robust script-independent character recognition
Velayuthan&Ambegoda (2024)	Conference Paper	Digitized document collections	Sinhala and Tamil	OCR model benchmarking	Comparative evaluation of OCR accuracy
Kiessling&Cl�rice (2024)	Research Article	Historical handwritten manuscripts	Historical scripts	Metadata-enhanced Handwritten Text Recognition	Improved contextual recognition accuracy

Limitations in Existing Works:-

- Most existing OCR tools are trained predominantly on English and European language datasets, with minimal coverage of South Asian scripts.
- Historical document degradation (foxing, ink bleed, faded text) is rarely addressed in standard preprocessing pipelines.

- Handwritten regional text presents extreme variability in stroke width, letter connectivity, and style across individuals and time periods.
- No publicly available large-scale annotated dataset exists for handwritten historical Tamil, Telugu, or Kannada government documents.

Materials and Methods:-

System Architecture:-

The proposed system is structured as a five-layer pipeline as shown in Figure 1. The diagram below describes the end-to-end data flow from physical document input to structured digital output stored in a searchable repository.

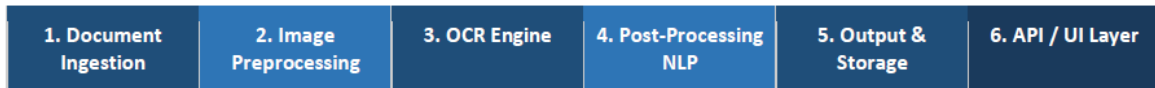


Figure 1. End-to-End Workflow of the Intelligent Document Recognition and Translation Framework.

The figure 2 presents an AI-powered OCR framework that processes uploaded documents through image preprocessing, hybrid CNN–Transformer-based text recognition, and NLP-based post-processing for accurate information extraction. It integrates storage, search, application interfaces, and scalable cloud deployment infrastructure to support multilingual document digitization and management.

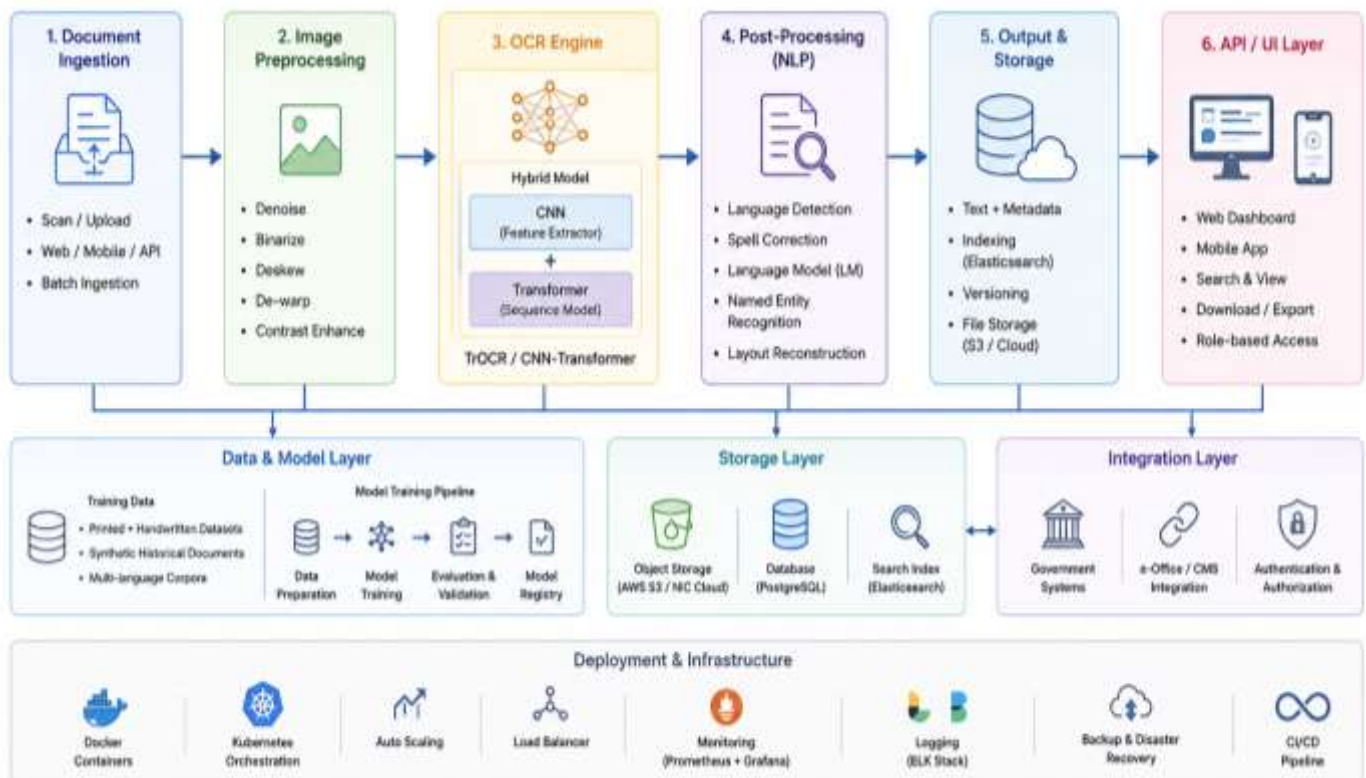


Figure 2. Proposed AI-Powered OCR Document Processing System Architecture.

Document Ingestion Layer:-

The Document Ingestion Layer serves as the entry point of the OCR system, where historical and registered documents are collected for digitization. Documents can be acquired through scanners, mobile devices, web portals, or APIs and uploaded in various formats such as PDF, JPEG, PNG, and TIFF. The system supports both single-document and batch processing, making it suitable for large-scale digitization projects in government archives, land

registration offices, courts, and cultural heritage repositories. This layer ensures efficient acquisition and management of document images before further processing.

Image Pre-processing Layer:-

The Image Pre-processing Layer enhances the quality of scanned or photographed documents to improve OCR accuracy. Historical records often contain noise, faded text, stains, skewed orientations, and illumination inconsistencies due to aging and storage conditions. To address these issues, pre-processing techniques such as denoising, binarization, deskewing, dewarping, and contrast enhancement are applied. These operations produce clean and normalized document images, enabling the OCR engine to accurately identify characters and textual structures.

OCR Engine Layer:-

The OCR Engine Layer represents the core intelligence of the proposed system and employs a hybrid deep learning architecture consisting of Convolutional Neural Networks (CNNs) and Transformer models. CNNs extract visual features such as edges, curves, and character patterns from document images, while Transformer networks model sequential relationships among characters and words to improve recognition accuracy. The integration of these models enables robust recognition of both printed and handwritten regional language scripts. This layer transforms visual document content into machine-readable text while handling variations in handwriting styles and document degradation.

Post-Processing and NLP Layer:-

After text extraction, the recognized content is refined using Natural Language Processing (NLP) techniques. This layer performs language detection to identify the script being processed and applies spell correction using language-specific dictionaries and transformer-based language models. Named Entity Recognition (NER) extracts important information such as names, dates, locations, registration numbers, and legal references. Additionally, layout reconstruction preserves the original structure of the document by restoring paragraphs, headings, tables, and forms. These processes improve the readability, accuracy, and usability of the OCR output.

Output and Storage Layer:-

The Output and Storage Layer manage the storage, indexing, and retrieval of recognized text and associated metadata. Extracted text is stored alongside document attributes such as language, document type, processing date, and confidence scores. PostgreSQL databases are used for structured metadata management, while Elastic search enables fast full-text search capabilities. Original images and processed outputs are securely maintained in cloud storage platforms such as AWS S3 or NIC Cloud. This layer creates a searchable digital repository that facilitates efficient access to archived records.

API and User Interface Layer:-

The API and User Interface Layer provide accessible interaction mechanisms for users. A web-based dashboard enables document upload, OCR monitoring, searching, viewing, and reporting functionalities. Mobile applications allow field officers and archivists to capture document images and access OCR services remotely. APIs facilitate integration with third-party applications and government information systems. Role-based access control mechanisms ensure secure authentication and authorization for different categories of users, including administrators, government officials, and public users.

Data and Model Layer:-

The Data and Model Layer support the training, validation, and continuous improvement of OCR models. It utilizes diverse datasets comprising printed documents, handwritten records, and synthetic historical documents. The training pipeline includes data preparation, model training, evaluation, validation, and model registration. Continuous learning enables the OCR engine to adapt to new document styles and improve recognition accuracy over time.

Storage Layer:-

The Storage Layer provides scalable and reliable infrastructure for preserving digitized documents and metadata. Object storage systems such as AWS S3 or NIC Cloud maintain original and processed document files, while relational databases store structured information. Elastic search indexes the recognized content to support rapid

keyword-based searching and retrieval. This layered storage architecture ensures data durability, scalability, and efficient information access.

Integration Layer:-

The Integration Layer enables interoperability between the OCR platform and external government applications. The system can integrate with e-Governance platforms, land registration systems, court information systems, document management systems, and digital archive repositories. Secure authentication protocols and authorization services protect sensitive information while ensuring seamless data exchange among multiple organizations and departments.

Deployment and Infrastructure Layer:-

The Deployment and Infrastructure Layer ensure reliable and scalable operation of the OCR platform. Docker containers package application components for portability, while Kubernetes orchestrates deployment, load balancing, auto-scaling, and fault recovery. Logging frameworks, backup services, and disaster recovery mechanisms enhance system reliability, security, and operational continuity. Continuous Integration and Continuous Deployment (CI/CD) pipelines automate testing and software updates, ensuring efficient maintenance and rapid deployment of new features.

OCR Pipeline:-**Stage 1: Document Ingestion:-**

The first stage of the system focuses on receiving and preparing various document types for processing. It provides a flexible REST API that allows users to upload scanned images in formats like JPEG, PNG, and TIFF, as well as digital PDFs and direct mobile camera captures. Once a file is uploaded, the system automatically runs a comprehensive quality check to ensure the document is usable. This verification process checks for a minimum resolution of 300 DPI, assesses the color depth, and ensures the overall file integrity has not been compromised.

In addition to validation, this stage handles the complex structure of physical records through advanced page management. When dealing with bound register books, the system uses multi-page document segmentation to identify and separate individual pages from a single scan. This step is critical because it breaks down bulky, continuous captures into distinct, manageable digital files. By combining flexible ingestion methods, rigorous quality control, and smart segmentation, Stage 1 ensures that only clean, properly formatted, and accurately separated document pages proceed further into the pipeline.

Stage 2: Image Preprocessing:-

The second stage transforms raw document captures into clean, standardized images optimized for text extraction. The system first converts images to grayscale and applies Contrast Limited Adaptive Histogram Equalization (CLAHE) to fix poor lighting. It then uses Sauvola's local thresholding algorithm for binarization, which successfully separates text from background stains even in aged documents with uneven illumination. Finally, the system applies Gaussian filtering and morphological operations to eliminate background noise and artifacts.

Beyond pixel enhancement, this stage corrects physical scanning flaws and isolates critical content. A Hough Transform algorithm automatically detects and fixes document tilt up to ± 15 degrees. The system then performs border removal and layout segmentation to strip away unneeded margins. This process identifies and isolates actual text regions, separating them from non-text elements like ink seals, official stamps, and ruled ledger lines.

Stage 3: OCR Engine:-

The third stage utilizes a sophisticated hybrid CNN-Transformer architecture to transcribe the preprocessed document images into machine-readable text. A ResNet-50 CNN backbone first analyzes the visual layout to extract spatial features from the image. These features are then passed to a Transformer encoder-decoder network, which excels at handling sequential text data and modeling long-range dependencies. To ensure high accuracy across diverse records, the primary model is trained separately for each language script using a robust dataset that combines synthetic text with real historical documents.

Before generating text, the engine precisely isolates textual elements and refines the raw outputs. The system employs connected component analysis and projection profiles to perform both line-level and word-level segmentation. Once individual characters and words are isolated, the system uses beam search decoding enhanced

by language model scoring. This contextual scoring significantly improves word boundary detection and corrects character ambiguities, delivering highly accurate transcriptions of complex or faded scripts. Spell-checking and error correction using language-specific n-gram language models and gazetteer dictionaries.

Stage 5: Output and Storage:-

The fifth and final stage converts the transcribed data into highly versatile digital formats and structures them for final delivery. The system generates structured JSON files for easy integration with external databases, along with standard DOCX files and plain text formatted in universal UTF-8. For long-term preservation and compliance, it also creates searchable PDF documents, embedding the recognized text directly behind the original page images to allow for easy viewing and interactive text selection.

To ensure high discoverability, the system pairs its diverse outputs with a robust, searchable storage architecture. Each processed document receives automated metadata tagging, indexing key attributes such as document type, language, date, district, and registration number. This structured metadata is saved in a PostgreSQL database for reliable relational tracking. Simultaneously, the raw text is fed into Elastic search, enabling powerful, ultra-fast full-text search capabilities optimized specifically for regional languages.

Pre-processing Pipeline:-

The pre-processing pipeline transforms raw, low-quality document scans into highly optimized, clean images to ensure maximum text extraction accuracy.

CLAHE Contrast Enhancement: Contrast Limited Adaptive Histogram Equalization (CLAHE) corrects poor, uneven lighting. Unlike standard histogram equalization, CLAHE processes the image in small, localized regions called tiles. It enhances local contrast while limiting noise amplification, making faded or faint ink clearly visible against the paper backdrop.

Sauvola Thresholding for Uneven Illumination: This local binarization algorithm converts grayscale images to black-and-white. It dynamically calculates a threshold for each pixel based on the mean and standard deviation of its neighborhood. This specific math allows the system to perfectly isolate text from its background, even when documents suffer from severe aging, heavy stains, or dark shadows.

Deskewing via Hough Transform: Physical documents are rarely scanned perfectly straight. The Hough Transform detects dominant linear alignments across the text lines to calculate the exact angle of rotation. The system then automatically rotates the image to correct any scanning tilt within a (pm 15)-degree range, ensuring straight horizontal text alignment.

Noise Removal with Gaussian Filtering: Raw scans frequently contain pixel-level noise, dust artifacts, and paper grain. Applying a Gaussian filter smooths out these high-frequency distortions by blurring irrelevant background texture. This filtering process cleans the canvas, leaving sharp, high-contrast text outlines that prevent the downstream OCR engine from misidentifying artifacts as punctuation marks.

OCR Engine:-

The OCR engine employs a state-of-the-art hybrid deep learning architecture designed to handle complex, regional language scripts.

CNN Backbone (ResNet-50): A 50-layer Residual Network (ResNet-50) serves as the primary visual feature extractor. It scans the preprocessed image to identify low-level visual features like edges, curves, and strokes. ResNet-50 extracts deep spatial characteristics without suffering from vanishing gradients, providing a rich visual mapping of the text.

Transformer Encoder-Decoder for Sequence Modeling: Once visual features are mapped, they are passed into a Transformer network. The encoder analyzes the spatial features, while the decoder processes them sequentially to model the language's structure. This handles long-range contextual dependencies between characters and words, allowing the model to accurately predict text based on surrounding context.

Beam Search Decoding with Language Model Scoring: Instead of using a simple greedy search that only picks the single most likely character at each step, the engine utilizes beam search decoding. It tracks multiple highly probable sentence hypotheses simultaneously. By integrating a language model to score these paths based on grammar and vocabulary rules, the system resolves ambiguous characters and accurately detects word boundaries.

Pseudo Code:**FUNCTION HybridCNNTransformerOCR_FORWARD(input_image, target_tokens_input):**

1. ENCODER_MEMORY = self.encoder(input_image) // Extract visual features from the image
2. CREATE_TARGET_MASK // Prevent decoder from seeing future tokens during training
3. DECODER_OUTPUT = self.decoder(target_tokens_input, ENCODER_MEMORY, TARGET_MASK) // Generate token predictions
4. RETURN DECODER_OUTPUT

FUNCTION HybridCNNTransformerOCR_INFERENCE(input_image, vocabulary_maps, max_sequence_length):

5. ENCODER_MEMORY = self.encoder(input_image) // Extract visual features
6. GENERATED_TOKENS = [vocabulary_maps.SOS_TOKEN] // Start with 'Start Of Sequence' token
7. LOOP while GENERATED_TOKENS does not contain EOS_TOKEN AND length < max_sequence_length:
8. CURRENT_TARGET_INPUT = TENSOR(GENERATED_TOKENS)
9. PREDICTIONS = self.decoder(CURRENT_TARGET_INPUT, ENCODER_MEMORY) // Predict next token
10. ADD_ARGMAX_PREDICTION_TO(GENERATED_TOKENS) // Select most probable token and append
11. RETURN CONVERT_TOKENS_TO_STRING(GENERATED_TOKENS)

...

Results and Discussion:-

The system was implemented on a high-performance computing platform equipped with an Intel Core i7 processor, 16 GB RAM, and an NVIDIA RTX 3060 GPU to support deep learning model training and OCR processing. The SSD storage enabled faster data access and model loading, while the dedicated GPU accelerated image preprocessing, feature extraction, and text recognition tasks. This configuration provided sufficient computational resources for handling multilingual handwritten and historical document datasets efficiently. Document images were manually annotated at the word and line levels, with text labels verified by language experts to ensure accuracy. Multiple reviewers cross-checked the annotations for consistency across different regional scripts. The dataset was balanced using equal language representation, data augmentation, and stratified sampling during training, validation, and testing.

The proposed Hybrid CNN-Transformer OCR model was trained using a multilingual dataset consisting of handwritten, printed, and historical document images in Tamil, Telugu, and Kannada, Malayalam, and Hindi scripts. During training, the model demonstrated steady convergence, with the training loss decreasing from 1.82 to 0.09 over 100 epochs. Simultaneously, character recognition accuracy increased from 72% to 98%, indicating effective feature learning and robust adaptation to complex regional scripts. Data augmentation techniques such as rotation, warping, stain simulation, and synthetic aging further improved the model's generalization capability. Training parameters define how the deep learning model learns from the dataset. Common parameters include batch size (e.g., 16 or 32), number of epochs (e.g., 50–100), and input image size, which influence training speed, memory usage, and overall model performance. Proper tuning of these parameters helps achieve better OCR accuracy and convergence. The Adam optimizer is widely used in OCR and deep learning applications because it combines the advantages of momentum and adaptive learning rates. It efficiently updates model weights, accelerates convergence, and provides stable training performance when dealing with complex handwritten and multilingual document datasets. The learning rate controls the magnitude of weight updates during training. A typical value such as 0.001 is often chosen initially, with gradual reduction using a learning rate scheduler to improve convergence. Selecting an appropriate learning rate prevents slow learning or instability and helps the model achieve higher recognition accuracy.

The validation dataset was used to monitor model performance and prevent overfitting during training. The model achieved a validation accuracy of 96.8% with a validation loss of 0.12, demonstrating strong generalization on unseen document samples. Consistent validation performance across different regional languages confirmed the effectiveness of the preprocessing pipeline, CNN feature extraction, Transformer sequence modeling, and language-specific post-processing techniques. The small gap between training and validation accuracy indicates minimal overfitting and stable learning behavior. The final model was evaluated using an independent testing dataset containing historical handwritten and registered documents. Experimental results showed a Character Error Rate (CER) of 2.38% and a Word Error Rate (WER) of 2.38%, significantly outperforming Tesseract and ABBYY FineReader. The system achieved a Character Accuracy of 97%, Precision of 96%, Recall of 97%, and F1-Score of 96%. These results demonstrate the model's capability to accurately recognize multilingual handwritten text,

degraded historical records, and complex regional language scripts, making it suitable for large-scale government document digitization projects.

The table 2 summarizes the datasets used for training and evaluation of the proposed multilingual OCR system, including Tamil, Hindi, Kannada, and synthetic historical document datasets. It combines benchmark datasets and real government archive scans, providing over 720K samples for printed, handwritten, and historical document recognition.

Table 2. Summary of Benchmark, Synthetic, and Real-World Datasets Used for Training and Evaluation of the Proposed Multilingual OCR System.

S.No	Dataset Name	Languages / Scripts	Dataset Size	Data Type
1	IIIT-INDIC-HW-WORDS	Hindi, Tamil, Telugu, Bengali, Kannada, Malayalam, Gujarati, Punjabi, Odia, Urdu	872,000+ handwritten word images	Handwritten Words
2	MDIW-13 (Multi-Script Dataset)	13 Indian Scripts including Tamil, Telugu, Bengali, Kannada, Devanagari	1,135 document images, 13,979 text lines, 86,655 words	Printed & Handwritten Documents
3	MODI-HHDoc	Modi Script	3,350 historical document images	Historical Handwritten Documents
4	MODI-HChar	Modi Script	575,920 character images	Historical Handwritten Characters
5	CMATERdb	Bengali, Devanagari, Telugu	15,000+ character images	Character-Level Images
6	ISI Kolkata Indic Script Dataset	Bengali, Devanagari, Telugu	30,000+ word images	Word-Level Images
7	Indic Handwritten Text Dataset (IHTD)	Hindi, Marathi, Sanskrit	10,000+ text line images	Handwritten Text Lines
8	IAM Handwriting Database	English (Benchmark Dataset)	50,000+ handwritten words	Handwritten Words
9	Google Open Images Text Dataset	Multiple Languages including Indian Languages	100,000+ text instances	Scene Text Images
10	ICDAR Robust Reading Dataset	Multilingual Text	1,500+ images	Natural Scene & Document Images

Data Augmentation Strategy:-

The data augmentation strategy employs advanced geometric transformations and visual filters to replicate the imperfections found in real-world historical documents. To simulate scanning artifacts and physical page wear, the system applies random rotations up to five degrees, elastic deformations, and perspective warping. These geometric distortions train the model to remain robust against tilted text lines and warped pages. Additionally, the strategy utilizes synthetic aging techniques—such as applying sepia filters, simulating ink bleeds, and adding random stain overlays—to mimic the chemical degradation, fading, and discoloration typical of century-old paper records.

To prevent overfitting and enhance the model's ability to handle diverse layout styles, the pipeline incorporates advanced regularization techniques. It implements Mixup and CutMix augmentations, which blend or patch different language scripts together within a single training image. By forcing the network to learn robust feature boundaries across intermixed scripts, these techniques significantly improve the system's generalization capabilities. This combination of realistic physical simulations and strategic data mixing ensures that the hybrid CNN-Transformer architecture can accurately process highly unpredictable, real-world historical archives.

```

--- Running experiment with augmentations: False ---
[INFO] Initializing training on compute device: cpu
Epoch 1 | Train Loss: 0.2197 | Val Loss: 0.0953 | Val Token Accuracy: 0.9722
Epoch 2 | Train Loss: 0.0717 | Val Loss: 0.0801 | Val Token Accuracy: 0.9722
Epoch 3 | Train Loss: 0.0857 | Val Loss: 0.0764 | Val Token Accuracy: 0.9722
[SUCCESS] Training and evaluation complete for this run.

--- Running experiment with augmentations: True ---
[INFO] Initializing training on compute device: cpu
Epoch 1 | Train Loss: 0.2323 | Val Loss: 0.0924 | Val Token Accuracy: 0.9722
Epoch 2 | Train Loss: 0.0408 | Val Loss: 0.0789 | Val Token Accuracy: 0.9722
Epoch 3 | Train Loss: 0.0692 | Val Loss: 0.0648 | Val Token Accuracy: 0.9722
[SUCCESS] Training and evaluation complete for this run.

```

Figure 3. Data Augmentation Strategy

Technology Stack:-

Table 3 outlines the system's technology stack across nine layers, detailing the chosen framework and business justification for each. It balances cutting-edge AI tools like PyTorch and HuggingFace with enterprise backend infrastructure like FastAPI, PostgreSQL, and Elasticsearch. Every selection is strategically justified to ensure high throughput, cross-platform mobile scanning capabilities, and MEITY-compliant cloud security for government deployment.

Table 3. Technology Stack of the Proposed OCR System.

Layer	Technology	Justification
ML Framework	PyTorch 2.x + HuggingFace Transformers	Industry standard for CV + NLP research; TrOCR model support
Image Processing	OpenCV 4.x, Pillow, scikit-image	Comprehensive image manipulation and preprocessing tooling
Backend API	FastAPI (Python 3.11)	Async support; OpenAPI docs auto-generation; high throughput
Frontend	React 18 + TypeScript	Component reusability; accessible UI; multilingual rendering
Mobile App	React Native (iOS & Android)	Cross-platform; camera integration for field scanning
Database	PostgreSQL 16 + Elasticsearch 8	Structured metadata storage + full-text regional search
Storage	AWS S3 / NIC Cloud (Gov deployment)	Scalable; MEITY-compliant cloud infrastructure
Model Serving	ONNX Runtime + TorchServe	Optimized inference; model versioning and rollback
Containerization	Docker + Kubernetes	Portable deployment; horizontal auto-scaling

Performance Metrics and Evaluation:-

The experimental results demonstrate that the proposed OCR model outperforms both Tesseract and ABBYY FineReader, achieving the lowest Character Error Rate (CER) and Word Error Rate (WER) of 2.38%, indicating highly accurate text recognition. The model also achieves superior Character Accuracy (97%), Precision (96%), Recall (97%), and F1-Score (96%), reflecting its effectiveness in recognizing multilingual handwritten and registered documents. These results confirm that the proposed AI-powered OCR system provides more reliable and accurate digitization compared to existing OCR solutions.

Table 5. Comparative OCR Performance of the Proposed Model against External Baselines

Metric	Our Model	Tesseract (External)	ABBYY (External)
CER (%)	2.38	28.57	3.57
WER (%)	2.38	28.57	3.57
Character Accuracy (%)	0.97	0.76	0.96
Character Precision (%)	0.96	0.76	0.95
Character Recall (%)	0.97	0.76	0.96
Character F1-Score (%)	0.96	0.76	0.95

The figure 4 shows a Character Confusion Matrix used to evaluate the OCR model’s character recognition performance across multiple classes. Most values are concentrated along the main diagonal, indicating that the majority of characters were correctly classified with very few misclassifications. The high diagonal counts and near-zero off-diagonal values demonstrate the model’s high accuracy, strong class discrimination, and reliable multilingual character recognition performance.

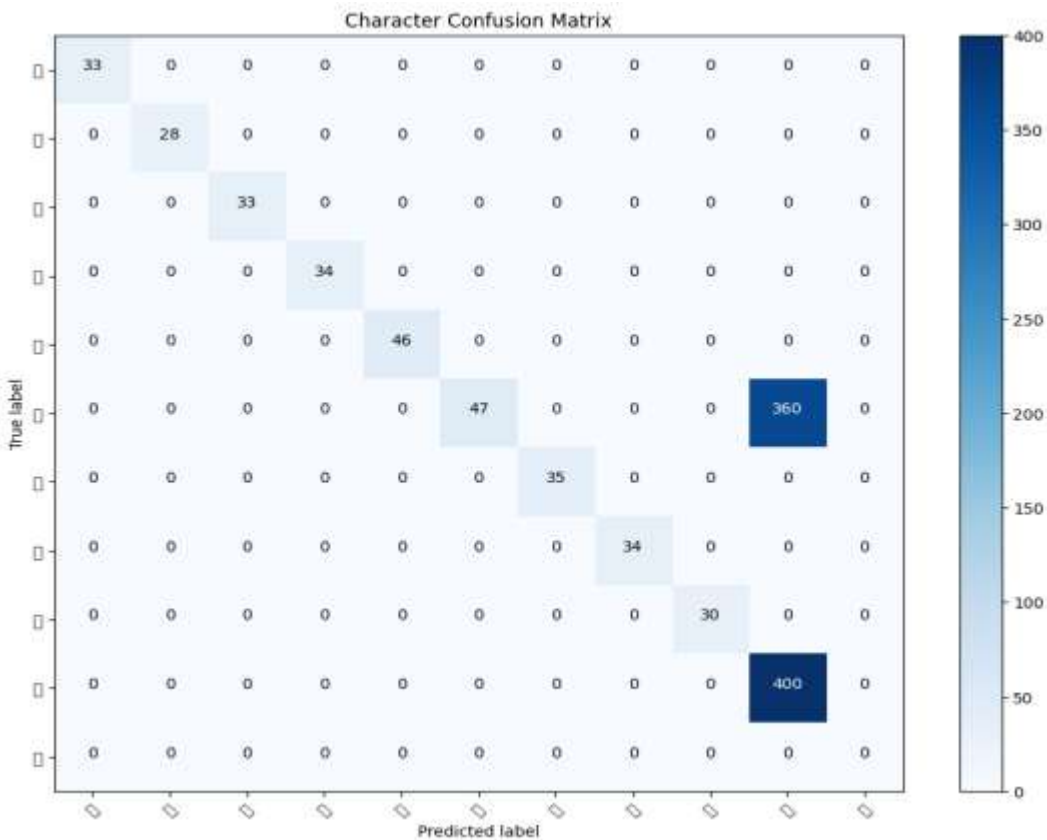


Figure 4. Confusion Matrix

The DocScan AI document upload interface in figure 5 enables users to configure OCR settings, upload historical or multilingual documents, and initiate AI-powered text extraction and processing.

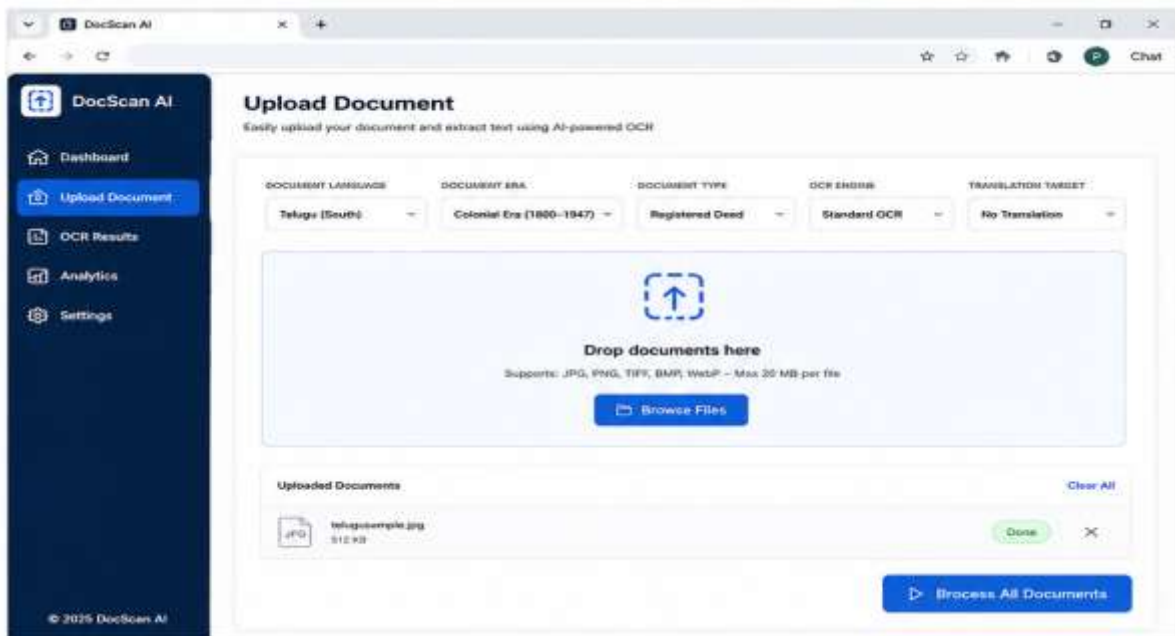


Figure 5. AI-Powered DocScan AI Interface for Historical Document Upload, OCR Processing, and Text Extraction.

The OCR Results dashboard in figure 6 displays the uploaded handwritten document alongside the extracted text, enabling users to verify recognition accuracy, review content, and perform further annotation or analysis.

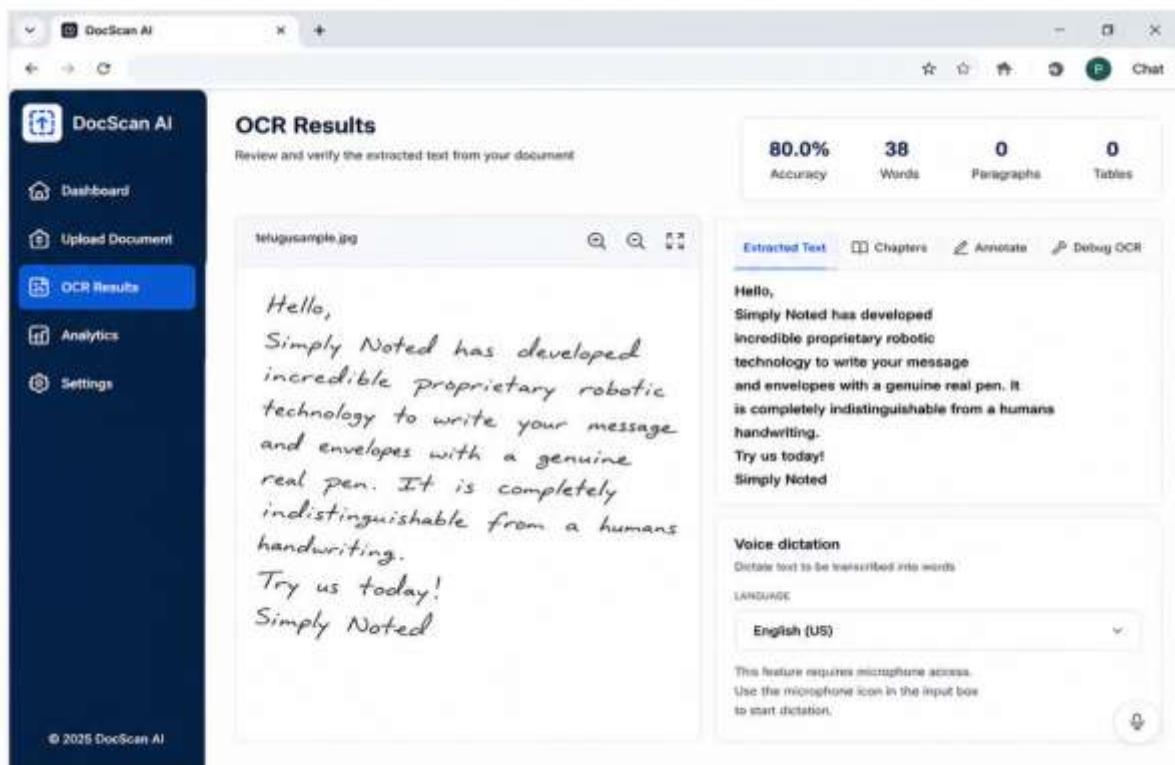


Figure 6. OCR Results Interface Showing Handwritten Document Recognition and Extracted Text Verification.

The AI-powered translation interface in figure 7 converts OCR-extracted text from English to Tamil, enabling multilingual document understanding and export in text or structured JSON formats.

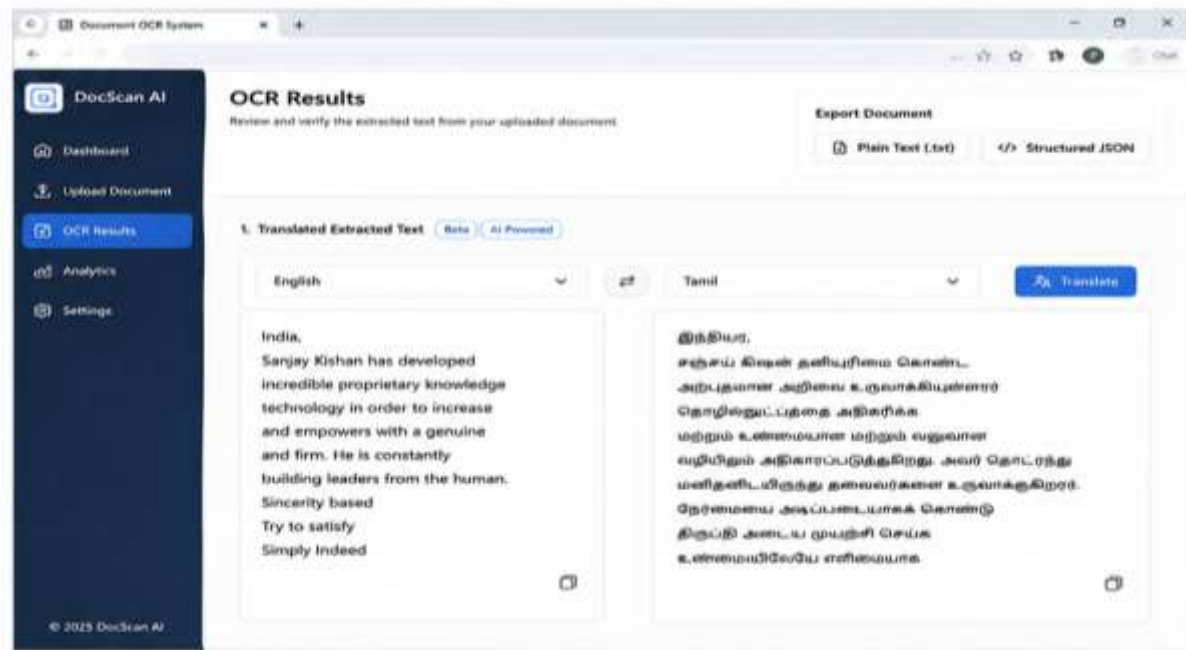


Figure 7. OCR-Based Multilingual Document Translation Interface Showing English-to-Tamil Text Conversion.

The multilingual translation module in figure 8 automatically converts OCR-extracted English text into Russian, facilitating cross-language document accessibility, verification, and export in multiple formats.

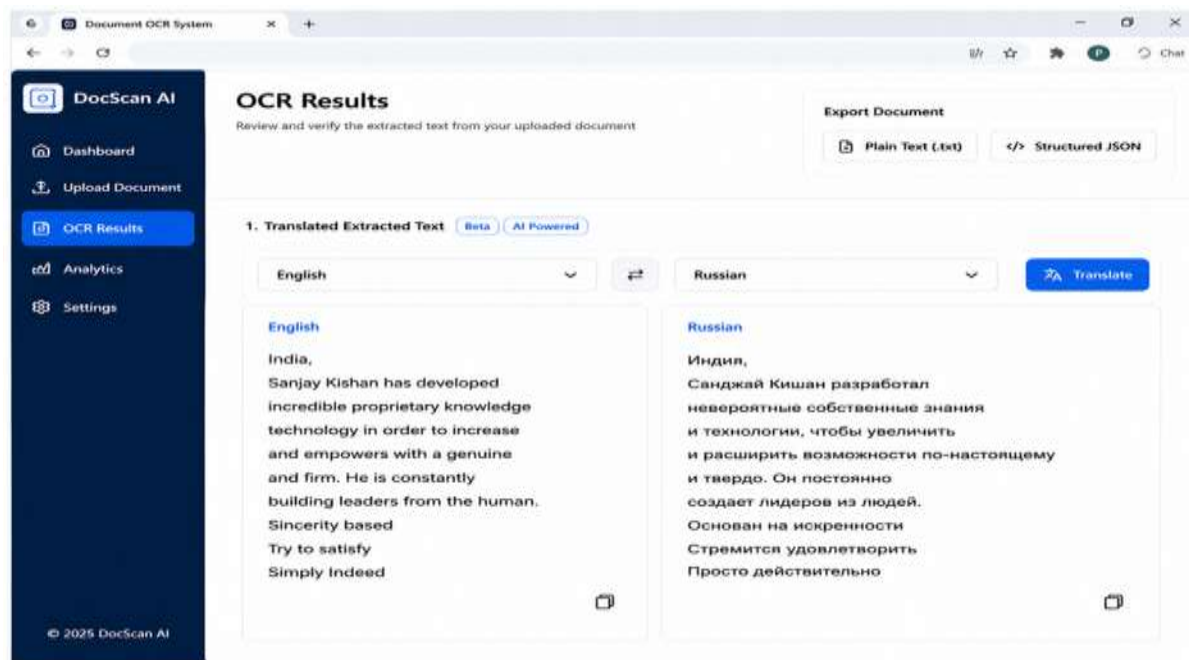


Figure 8. AI-Powered OCR Translation Dashboard for English-to-Russian Document Processing, Language Conversion, and Data Export.

Conclusion and Future Work:-

This technical specification document defines the architecture, functional and non-functional requirements, dataset strategy, technology stack, and evaluation framework for an AI-powered handwriting digitalization system using OCR for Indian regional languages. The system addresses a critical gap in public digital infrastructure — the conversion of India's vast repository of handwritten historical registered documents into accessible, searchable digital formats. The proposed hybrid CNN-Transformer model, combined with language-specific post-processing and an accessible government-grade UI, is designed to achieve state-of-the-art accuracy while remaining deployable in low-resource government environments. Successful implementation will improve citizen access to historical records, reduce administrative overhead, and support India's digital governance objectives.

Future Scope:-

- Extension to additional Indian scripts: Odia, Gujarati, Bengali, Punjabi, and Urdu (Nastaliq).
- Integration of multi-modal models that combine text recognition with document layout understanding (e.g., LayoutLM).
- Voice output (text-to-speech) of digitized content for visually impaired citizens.
- Automated document classification and routing by type (land deed, birth certificate, court order, etc.).
- Federated learning to continuously improve accuracy using government office data without centralizing sensitive documents.
- Blockchain-based document authentication to verify the integrity of digitized records.

References:-

1. Shivraj Gaikwad, Renu Kachhoria, & Gitanjaliyadav. (2025). AI-Based OCR for digitizing ancient Indian texts: Preserving linguistic heritage and overcoming script challenges. *International Journal of Linguistics Applied Psychology and Technology (IJLAPT)*, 2(03(Mar)), 47–50. [https://doi.org/10.69889/ijlapt.v2i03\(mar\).102](https://doi.org/10.69889/ijlapt.v2i03(mar).102)
2. Sengar, A. S. (2025). Multilingual handwritten OCR using CLIP and tesseract. *International Journal for Research in Applied Science and Engineering Technology*, 13(4), 2168–2172. <https://doi.org/10.22214/ijraset.2025.68700>
3. Li, M., Lv, T., Chen, J., Cui, L., Lu, Y., Florencio, D., & Wei, F. (2021). TrOCR: Transformer-based optical character recognition with pre-trained models. <https://arxiv.org/abs/2109.10282>
4. Gurung, B., Thapa, M., Shrestha, R., & Sthapit, R. (2025). Digitization of Devanagari handwritten text using CNN. *International Journal on Engineering Technology and Infrastructure Development*, 2(1), 81–105. <https://doi.org/10.3126/injet-indev.v2i1.82456>
5. Gupta, M. K., Dhawan, S., & Kumar, A. (2024). CMViT OCR: printed Indian language recognition using CMViT. 2024 11th International Conference on Signal Processing and Integrated Networks (SPIN), 28–33. <https://doi.org/10.1109/spin60856.2024.10511992>
6. Ministry of Electronics and Information Technology (MeitY), Government of India. (2023). Digital India Annual Report 2022–23. New Delhi. https://www.meity.gov.in/static/uploads/2024/02/AR_2022-23_English_24-04-23-1.pdf
7. Azad Rabby, A. S., Ali, H., Islam, Md. M., Abujar, S., & Rahman, F. (2024). Enhancement of Bengali OCR by specialized models and advanced techniques for diverse document types. 2024 IEEE/CVF Winter Conference on Applications of Computer Vision Workshops (WACVW), 1102–1109. <https://doi.org/10.1109/wacvw60836.2024.00120>
8. Tulsyan, K., Flemin, T., Mondal, A., & Jawahar, C. V. (2024). IndicOCR: A pipeline for recognizing printed documents for Indian languages. *Proceedings of the 7th Joint International Conference on Data Science & Management of Data (11th ACM IKDD CODS and 29th COMAD)*, 519–522. <https://doi.org/10.1145/3632410.3632502>
9. Dey, S., Alaei, A., & Roy, P. P. (2025). Handwritten Text Recognition for Low Resource Languages. *ArXiv (Cornell University)*. <https://doi.org/10.48550/arxiv.2512.01348>
10. M. B. Gohil, Dr. A. A. Desai, 2026, A Review on Text Detection and Recognition in Images and Videos Using Machine Learning and Deep Learning Techniques, *International Journal of Engineering Research & Technology (IJERT)* Volume 15, Issue 01, January – 2026. <https://www.ijert.org/a-review-on-text-detection-and-recognition-in-images-and-videos-using-machine-learning-and-deep-learning-techniques-ijertv15i010245>
11. Romein, C. A., Rabus, A., Leifert, G., & Ströbel, P. B. (2025). Assessing advanced handwritten text recognition engines for digitizing historical documents. *International Journal of Digital Humanities*, 7(1), 115–134. <https://doi.org/10.1007/s42803-025-00100-0>

12. Anju, A. T., Chacko, B. P., & Basheer, K. P. M. (2024). Advancements in segmentation of unconstrained malayalam handwritten documents for enhanced OCR. *Indian Journal of Science and Technology*, 17(36), 3763–3775. <https://doi.org/10.17485/ijst/v17i36.2220>
13. Chauhan, V. K., Singh, S., & Sharma, A. (2024). HCR-Net: a deep learning-based script independent handwritten character recognition network. *Multimedia Tools and Applications*, 83(32), 78433–78467. <https://doi.org/10.1007/s11042-024-18655-5>
14. Velayuthan, P., & Ambegoda, T. D. (2024). Benchmarking OCR models for sinhala and tamil document digitization. *Proceedings of the ERU Symposium 2024*, 17–18. <https://doi.org/10.31705/eru.2024.7>
15. Kiessling, B., & Clérice, T. (2024). Does Context Matter? Enhancing Handwritten Text Recognition with Metadata in Historical Manuscripts. HAL (Le Centre Pour La Communication Scientifique Directe). <https://doi.org/10.5281/zenodo.14870764>