



Kurdish Handwritten Word Recognition with CRNN Using Multi-Backbone Evaluation

Israa Mahdi Alsaqi^{1*} , Polla Fattah² 

^{1,2}Software Engineering and Informatics Department, College of Engineering, Salahaddin University, Erbil 44001, IRAQ

DOI: <https://doi.org/10.63841/iue24539>

Received 24 May 2025; Accepted 21 Jun 2025; Available online 25 Oct 2025

ABSTRACT

We introduce the Kurdish Handwritten Word Dataset (KHWD), this is the first large-scale collection of handwritten word images created from printed forms by 8,000 native Sorani Kurdish writers. The dataset contains around 400,000-word samples covering 10,000 unique words, captured via high-resolution scanning (1200 DPI) and processed through an automated segmentation pipeline to produce aligned image–transcription pairs (with CSV mappings). This work specifically highlights the challenges of the Sorani script (cursive joins, diacritics, dot-position distinctions, and right-to-left writing) that motivate our dataset design. We employ a Convolutional Recurrent Neural Network (CRNN) for recognition, evaluating multiple CNN backbones (ResNet18, ResNet50, GoogleNet, and MobileNetV2). The best results were achieved by MobileNetV2 pretrained on isolated Kurdish characters (achieving 97% training and 96% validation accuracy) and then fine-tuned on a combination of KHWD and synthetically generated word images, this augmentation yields significantly better performance. The improved MobileNetV2 achieves approximately 17.98% CER, 54.31% WER, and 45.69% word-level accuracy, reducing CER by 1.47%, WER by 3.28%, and increasing word-level accuracy by 3.28% compared to the baseline MobileNetV2.

Keywords: Kurdish OCR, handwritten text recognition, CRNN, Mobile Net, synthetic data, low-resource languages



1 INTRODUCTION

handwritten text recognition (HTR) for low-resource languages like Kurdish is still an open challenge. Because of the complexity and the cursive nature of the Kurdish script. Segmentation and recognition become more difficult by the contextual character forms, diacritical markings, dot-position differences, and cursive writing observed in the Sorani Kurdish script, deep learning-based OCR has also shown promise for digitization of handwritten documents. For example [1] achieved high recognition accuracy on official records by using a Mask R-CNN-based system to identify handwritten Kurdish forms within an e-government framework. Because of a lack of large word-level datasets for evaluation, current OCR and HTR systems for Kurdish have mainly focused on printed text or isolated character classification [2]. In addition to current developments in printed Kurdish OCR, attempts have started to make historical publications machine-readable. [3] suggested enhancing Tesseract and other OCR systems to handle damaged Kurdish books from the early 20th century, with an emphasis on problems like bleed-through, scratch, and irregular fonts. Their study focuses on printed text, highlighting the significant need for OCR solutions across various Kurdish scripts and formats. To address this gap, the Kurdish Handwritten Word Dataset (KHWD) is presented and evaluated various deep learning models on it. Our contributions in this paper are as follows:

- We introduce KHWD, a large-scale scanned corpus of around 400,000 handwritten word images from 8,000 writers, covering around 10,000 unique Sorani words along with their transcriptions.
- We develop a Convolutional Recurrent Neural Network (CRNN) for whole word recognition and assess its efficacy employing multiple convolutional backbones (ResNet18, ResNet50, GoogleNet, and MobileNetV2).

*Corresponding author: israamahdi82@gmail.com
<https://ojs.cihanrtv.com/index.php/public>

- We demonstrate that pretraining MobileNetV2 on isolated Kurdish characters and augmenting training with synthetic handwritten-word images yields significant improvements in recognition (lower CER/WER, higher word accuracy) on this dataset.

2 RELATED WORKS

Even recent expansion, research on Kurdish handwritten text recognition is still in the early stages when compared to other scripts. Using traditional feature-based techniques, early attempts focused on identifying single Kurdish characters or printed text[4]. These methods don't scale efficiently to cursive, word-level recognition and often use small-scale datasets. To our knowledge, no previous work has addressed large-scale Sorani Kurdish word-level HTR until now. While several recent studies have proposed character-level datasets and recognition benchmarks [4, 5], there remains a clear need for robust word-level benchmarks. Recent attempts to address this gap include the work by [6] introduced an EMNIST-style dataset with 58 characters (letters and digits) captured from 3,800 native Sorani writers, establishing baseline benchmarks for isolated recognition. However, both studies focused exclusively on character-level analysis. For OCR tasks, Convolutional Recurrent Neural Networks (CRNN) work effectively because of combining CNNs for feature extraction with LSTMs to model the dependencies between sequential characters. High accuracy and low latency have been demonstrated by a CRNN model with seven CNN layers and two LSTM layers on challenging OCR benchmarks. These findings demonstrate the potential of CRNN for low-resource scripts, such as Kurdish, where word-level HTR is still not well understood[7].

Although [8] achieved 97% accuracy on isolated characters, their dataset lacks cursive or contextual character connections essential for word-level recognition which released a comprehensive dataset of over 40,000 handwritten Kurdish isolated characters. Additionally,[9] introduced the KurdSet dataset, a Kurdish handwritten digit's dataset consisting of 15,600 images from 1,560 participants and evaluated it using CNN-based models such as ResNet50, DenseNet121, MobileNet, and a custom CNN. Their models best-performance showed the dataset's effectiveness for digit-level recognition tasks in Kurdish, achieving test accuracy of up to 99.73%. Similarly, [10] proposed an ensemble transfer learning model combining DenseNet201, InceptionV3, and Xception for Kurdish character recognition, attaining robust performance on limited resources. However, these studies were restricted to non-cursive, character-level tasks and lacked large-scale word-level benchmarks. Recently, [11]evaluated advanced HTR engines across historical document corpora and noted that deep learning models with layout-awareness and multi-stage recognition pipelines significantly outperformed traditional OCR, particularly in low-resource contexts like Sorani Kurdish.

In contrast, Arabic and Urdu scripts—structurally similar to Kurdish—have seen significant improvements in full-word HTR through hybrid and transformer-based models. [12] a CNN-BiLSTM-CTC architecture enhanced with attention and adaptive augmentation for Arabic handwritten word recognition. Their model achieved a Character Error Rate (CER) of 15.2% and a Word Accuracy Rate (WAR) of 82.4% on the IFN/ENIT benchmark, showing the critical role of attention mechanisms in modeling cursive script dependencies.

For Urdu HTR .The TrOCR transformer architecture, pretrained on large-scale synthetic Urdu data and fine-tuned on real handwritten words, demonstrated strong generalization on noisy, unconstrained handwriting samples, with significantly reduced CER compared to CRNN baselines [13]. Similarly, [14]proposed ET-Network, a lightweight transformer model for Urdu handwritten recognition, achieving high accuracy while maintaining computational efficiency for low-resource applications.

The proposed KHWD dataset in the current study directly addresses the lack of large-scale word-level Kurdish datasets. It includes 400,000 handwritten word images from 8,000 participants, and to our knowledge, it is the most comprehensive corpus of its kind for Sorani Kurdish. By using Connectionist Temporal Classification (CTC), we benchmark several types of CNN backbones (ResNet18, ResNet50, GoogLeNet, and MobileNetV2) within a CRNN architecture. We show that MobileNetV2 performs significantly better than deeper models when pretrained on isolated Kurdish characters and fine-tuned with synthetic word images. Our improved model gets a CER of 17.98% and WER of 54.31%, closely aligning with results reported in Arabic and Urdu OCR literature. Furthermore, our use of on-the-fly data augmentation—including rotation, contrast jitter, and Gaussian blur—is inspired by[15], who showed these techniques boost HTR robustness. This structure guarantees generalization across handwriting styles and prepares the model for real-world OCR deployment in low-resource settings. [16] highlighted those preprocessing techniques like noise reduction, slant correction, and binarization have a quantifiable effect on CNN-based HTR systems, especially in settings with historical or degraded documents.

To date, KHWD is the first and only Kurdish dataset to support comparative benchmarking at the word level. Our work bridges a critical research gap and provides a foundation for future studies to explore transformer-based Kurdish OCR.

3 METHODOLOGY

In this section we explained the entire pipeline for recognizing Kurdish handwritten words using CRNN models. The process of creating a dataset, preprocessing methods, data augmentation techniques, dataset splitting, model architecture, and training configuration are all covered. Each step is intended to guarantee strong model generalization across various handwriting styles and to address the difficulties of identifying Sorani Kurdish script.

3.1 DATASET CREATION

The KHWD dataset was created by generating dedicated data entry forms containing lists of Kurdish words. Participants (8,000 native speakers) were asked to handwrite the given words on these forms. The forms were collected in multiple sessions Manually checked for completeness or any mistakes in the spelling then scanned at 1200 DPI resolution to preserve fine details. We applied an automated segmentation pipeline that detected word bounding boxes and extracted individual word images. Each image has links to a ground-truth transcription using a CSV mapping file. The final dataset contains around 400,000-word images from 10,000 distinct words. Additionally, current Kurdish character-level datasets, such as the one developed by[8], Provide over 40,000 isolated character samples and show standardized data collection through printed grids and Photoshop-based segmentation. Their work, while focused on isolated characters, highlights the necessity for large word-level databases such as KHWD. This approach is like previous dataset efforts: for example,[17] collected forms from 1,560 participants, resulting in over 45,000 isolated Kurdish character images. Because of the distinct set of characters and ambiguous writing style standards, including contextual shapes and diacritical signs, Kurdish script faces unique challenges for OCR. Letters are distinguished by their specific arrangements and number of dots, and so different letters may look the same with different dot patterns. Sorani script's combination of cursive writing and right-to-left direction complicates character segmentation and model training. All these factors necessitate the design of our complex and diverse dataset. In parallel to our work, the CKSD project by[18] presented an accurate dataset of 27 fonts containing isolated Kurdish characters, common words, and dates collected from various document sources. CKSD's contribution validates the current movement towards structured, scalable Kurdish text datasets for NLP and OCR applications, while it focusses more on font diversity and OCR-related font recognition. Figure 1 outlines our data collection pipeline and Figure2 illustrating File and Folder structure of the dataset.

3.2 DATA PREPROCESSING

Before training, each word image was preprocessed to improve model robustness. All images were first converted to grayscale, as color information is irrelevant in handwritten text. We normalized pixel intensities to the $[0,1]$ range using PyTorch's `ToTensor()` transformation. Each image was then resized to a fixed height of 64 pixels, and a fixed width of 256 pixels, maintaining general shape without requiring character-level segmentation. This resizing ensures uniform input dimensions for batch training with convolutional models. The images were not padded or filtered. We confirmed that the preprocessing preserved both legibility and ground truth label integrity for all samples. A recent study by[19] It was proved that applying a structured pipeline of preprocessing techniques, such as binarization, edge detection, and skeletonization, significantly enhances CNN-based Kurdish character recognition , with over 97% accuracy in testing on a dataset of 40,000 images. Their findings highlight that the amount of preprocessing directly affects the model's performance and significantly impacts final recognition results.

3.3 DATA AUGMENTATION

To improve generalization and mimic real-world handwriting variation, we applied on-the-fly data augmentation during training. Inspired by prior OCR practices, the following transformations were used: random rotation ($\pm 5^\circ$), small translations (up to 3% of image width and height), brightness and contrast jitter ($\pm 30\%$), and slight Gaussian blur to simulate image noise. These augmentations help the model become robust to variations in writing angle, brightness, stroke continuity, and style. All transformations were applied using PyTorch's `transforms.Compose()` with carefully selected parameter ranges.

3.4 DATASET SPLITTING

The dataset was randomly split into training, validation, and test sets, which 80% of images were used for training, 10% for validation (hyperparameter tuning), and 10% for final testing. We ensured that all 40 samples of a given

word from a single form were placed in the same split to avoid leakage. This split strategy is consistent with prior work; for example, [20] used an 80:20 train/test split for a similar Kurdish character dataset.

3.5 MODEL ARCHITECTURE

The current study implements a Convolutional Recurrent Neural Network (CRNN) architecture for complete Kurdish handwritten word recognition, as shown in figure 3. This architecture combines the features extraction capabilities of Convolutional Neural Networks (CNNs) with the ability of Recurrent Neural Networks (RNNs) for processing sequence model, making it appropriate for sequence-to-sequence text tasks without requiring character-level segmentation.

THE CRNN ARCHITECTURE CONSISTS OF THREE MAIN COMPONENTS:

CNN Backbone (Feature Extractor): We used four different pretrained CNN backbones — ResNet18, ResNet50, GoogLeNet (Inception v1), and MobileNetV2 — to extract spatial features from input word images. These backbones were initialized with ImageNet-pretrained weights to enable transfer learning. Each CNN outputs a feature map with width and channel dimensions dependent on the backbone. A 1×1 convolution layer is then applied to standardize the feature dimension to 256 channels, reducing computational complexity for the sequence model.

Sequence Modeling (BiLSTM): The 2D feature map from the CNN is converted to a sequence of vectors, after that are passed to a two-layer bidirectional Long Short-Term Memory (BiLSTM) network. This allows the model to capture both forward and backward context in the word image. Each LSTM layer contains 256 hidden units per direction, totaling 512 output features per time step.

Transcription Layer (CTC): The BiLSTM outputs are passed over a fully connected (linear) layer followed by a softmax to compute character probabilities at each timestep. Connectionist Temporal Classification (CTC) loss is used to align the predicted sequence with the target transcription. CTC enables training without requiring character boundaries in the image.

Dropout regularization is implemented with a probability of 0.5 to both the CNN feature output and the BiLSTM hidden states to reduce overfitting. All models process grayscale images reshaped to 64×256 and operate end-to-end in PyTorch.

3.6 TRAINING CONFIGURATION

All models were trained using the Adam optimizer, with a learning rate of $1e-4$, weight decay of $1e-4$ and a batch size of 128 for 30 epochs. To prevent stagnation in optimization, a learning rate scheduler (ReduceLROnPlateau) was used, reducing the learning rate by an amount of 0.5 if the validation loss failed to improve for 3 sequential epochs. A powerful Windows 11 laptop with an Intel® Core™ i7-14650HX processor (2.20 GHz, 14 cores), 40 GB of DDR5 RAM (5600 MHz), and an NVIDIA GeForce RTX 4070 GPU with 8 GB of VRAM was used for the training. PyTorch 2.x with CUDA 11.8 support was implemented to perform out the experiments.

All models were trained from scratch using the KHWD dataset. for the MobileNetV2-based CRNN, we implemented a two-phase training technique:

1. Pretraining on Isolated Characters: The CNN backbone and BiLSTM layers were first trained on a dataset of isolated Sorani Kurdish handwritten characters. This phase helped the model to learn each individual character shapes, and common writing variations.

2. Fine-tuning on Combination of KHWD and synthetic data: The pretrained model was then fine-tuned on the combination of KHWD word-level dataset and synthetic data generated by using six different cursive and print-style Kurdish fonts. This phase helped the model to get better generalization.

This two-phase approach to implementation enabled better weight initialization and enhanced convergence speed, which led to reduced CER/WER and increased accuracy relative to direct training on KHWD.

4 RESULTS AND ANALYSIS

The performance of each model on the KHWD test set is summarized in Table 1 and illustrated in Figure 1 Which we report Character Error Rate (CER), Word Error Rate (WER), and word-level accuracy (exact match). Lower CER and WER are better, while higher accuracy is better.

The evaluation metrics are computed as follows:

Character Error Rate (CER)

$$CER = \frac{(S + D + I)}{N}$$

Where:

- S is the number of substitutions,
- D is the number of deletions,
- I is the number of insertions, and
- N is the total number of characters in the ground truth.

Word Error Rate (WER)

$$WER = \frac{(S + D + I)}{N}$$

Where:

- S Number of incorrect words (substitutions)
- D Number of missing words (deletions)
- I Number of extra words (insertions)
- N Total number of words in the ground truth

Word-level Accuracy

$$Accuracy = \frac{1}{N} \sum_{i=1}^N 1 [\hat{y}_i = y_i]$$

Where:

- $\hat{y}_i$ is the predicted word,
- y_i is the ground truth word, and
- $1[\cdot]$ Indicator function (1 if correct, 0 otherwise)
- N Total number of words in the ground truth

These formulas are commonly used in sequence-level OCR evaluation. CER and WER are normalized Levenstein distances at the character and word level respectively, while word accuracy measures the proportion of samples where the predicted word exactly matches the ground truth.

The improved MobileNetV2 model achieves the best results, that shows a CER of 0.179, a WER of 0.543, and a word accuracy of 0.456. The baseline MobileNetV2 (without pretraining or synthetic data) gets the best performance also (CER 0.194, WER 0.575, acc 0.424), exceeding deeper networks like GoogLeNet (CER 0.202, WER 0.583, acc 0.416) and ResNet50 (CER 0.233, WER 0.614, acc 0.383). The shallowest network, ResNet18, had the worst performance (CER 0.250, WER 0.678, acc 0.321), likely due to overfitting or optimization difficulty. Overall, MobileNetV2 benefits most from the training strategy, delivering significant error reductions over the baseline as shown in Table 1.

In Figure 4: Bar chart comparing the Character Error Rate (CER), Word Error Rate (WER), and word-level accuracy of the five CRNN models on the KHWD test set. The improved MobileNetV2 model (far right) achieves the lowest CER and WER, and the highest word-level accuracy.

Figure 5 shows the Character Error Rate (CER) per epoch during training and validation for each CNN backbone integrated with the CRNN model. All models showed a rapid initial reduction in CER, referring to efficient early learning. However, the convergence behavior is different: ResNet18 and ResNet50 showed slower improvement and higher final CER, may be because to overfit or fail in feature extraction for this task. GoogLeNet and MobileNetV2 showed better generalization and lower CER over epochs. The improved MobileNetV2, pretrained on unique Kurdish characters and fine-tuned using synthetic word samples, consistently achieves the lowest validation Character Error Rate (CER) and indicates stable learning dynamics. This highlights the effectiveness of our pretraining and augmentation techniques in improving recognition accuracy.

For word-level recognition, we evaluated the MobileNetV2-based CRNN model on a separate Kurdish handwritten character classification task. Figure 6 shows the training and validation accuracy over 30 epochs. The model converged quickly and achieved over 97% training accuracy and approximately 96% validation accuracy, indicating strong generalization on character-level recognition, which supports the effectiveness of pretraining strategies for future word recognition.

Table 2 shows ten sample predictions generated by our improved MobileNetV2-based CRNN model on the KHWD test set. This model was initially pretrained on isolated handwritten Kurdish characters to capture basic character-level patterns. It was then fine-tuned using the full KHWD word-level dataset combined with synthetic handwritten words.

The samples show both successful recognition (e.g., ناوهوه → ناوهوه، چهنديهتي → چهنديهتي) and common OCR issues, such as character-level replacements (e.g., خالي → خالي) and word-level distortions (e.g., نهفلاتون → نهفالتون). These demonstrate the kinds of mistakes that resulted in the observed Word Error Rate (WER) of 54.31% and Character Error Rate (CER) of 17.98%.

Table 1. Test set performance of CRNN models on KHWD. Metrics include Character Error Rate (CER), Word Error Rate (WER), and word-level accuracy (%)

Model	CER (%)	WER (%)	Accuracy on word (%)
ResNet18	25.02	67.87	32.13
ResNet50	23.36	61.46	38.34
GoogLeNet (InceptionV1)	20.25	58.34	41.66
MobileNetV2	19.45	57.59	42.41
MobileNetV2+(improved)	17.98	54.31	45.68

Table 2. Sample predictions from the improved MobileNetV2 + BiLSTM + CTC model

Ground Truth	Prediction
چهنديهتي	چهنديهتي
نهفلاتون	نهفالتون
ناوهوه	ناوهوه
سهريهشت	سهريهشت
سزادان	سزادان
خالي	خالي
سترانوس	سترانوس
شنيينو	شنيينو
ميعماري	ميعاري
لهبي	لهبي

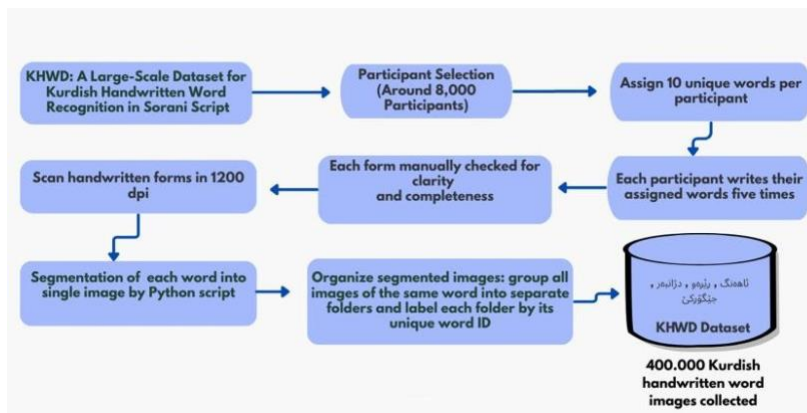


FIGURE 1. Outlines our data collection pipeline, illustrating how the dataset was systematically compiled from participant forms to ensure high-quality and diverse handwriting samples for OCR training

```

KHWD_Dataset/
|-- words.csv
|-- WA-0014/
|   |-- WA-0014_W1.png
|   |-- ...
|   |-- WA-0014_W40.png
|-- WA-0015/
|   |-- WA-0015_W1.png
|   |-- ...
|   |-- WA-0015_W40.png
    
```

FIGURE 2. Illustrates the organized file and folder structure of KHWD

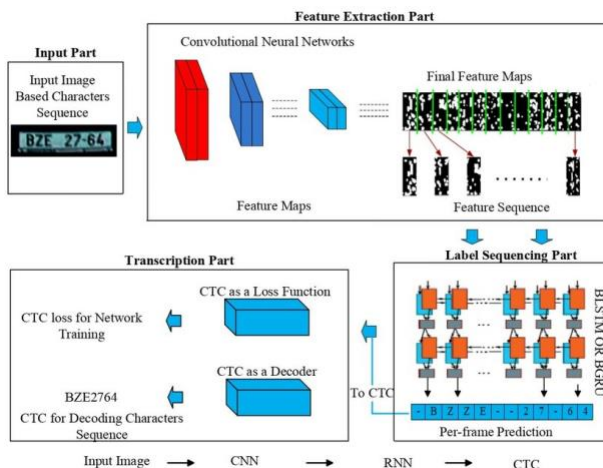


FIGURE 3. Overview of the CRNN architecture using CNN for feature extraction, BiLSTM for sequence modeling, and CTC for transcription [17]

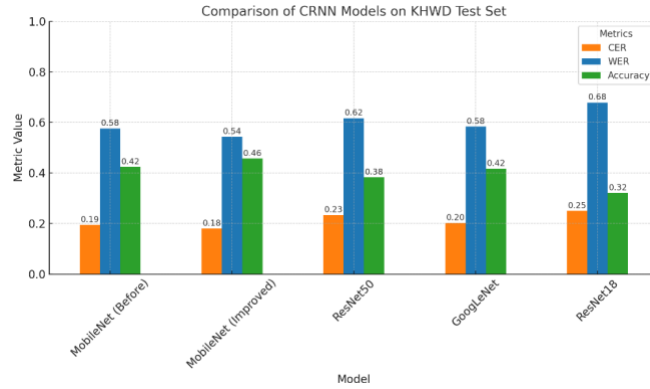


FIGURE 4. Illustrates that the improved MobileNetV2 outperforms other models in all three metrics, confirming its robustness for low-resource Kurdish OCR tasks.

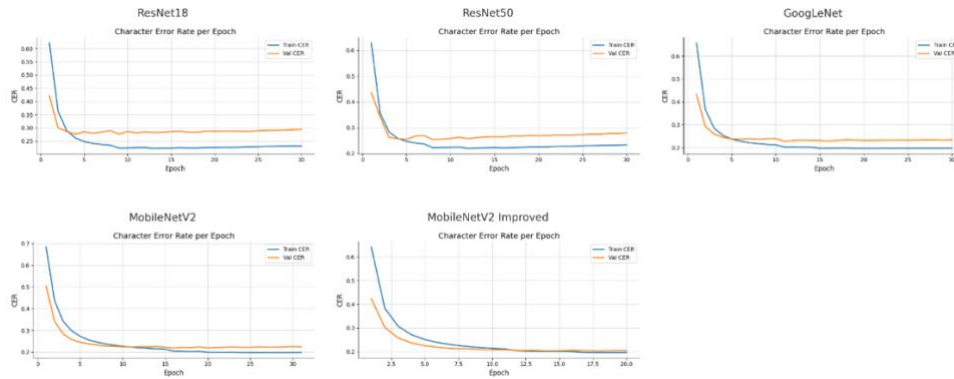


FIGURE 5. Shows CER per epoch for all models, revealing that the improved MobileNetV2 not only converges faster but also maintains the lowest CER throughout, emphasizing its training stability and generalization

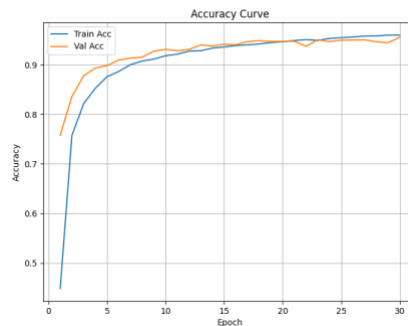


FIGURE 6. Provides empirical support for the effectiveness of the two-phase training approach, demonstrating how character-level pretraining accelerates convergence and boosts final accuracy in word-level tasks.

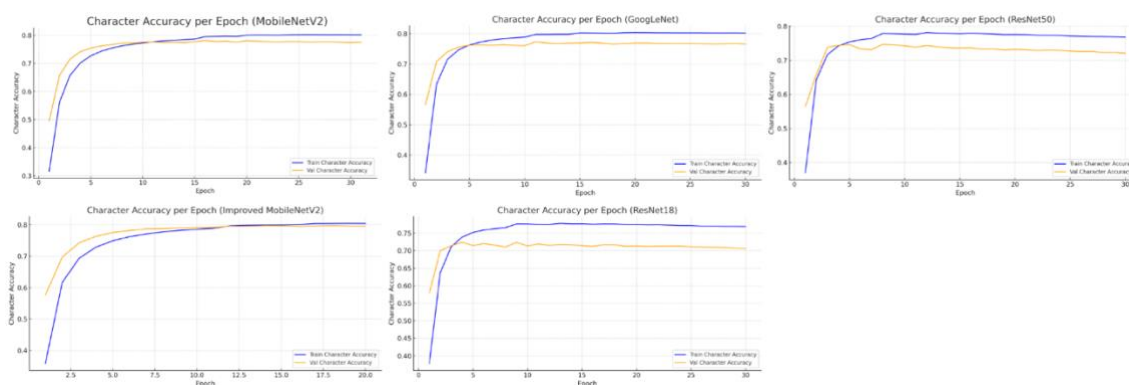


FIGURE 7. Character accuracy per epoch for CRNN models using different CNN backbones. Improved MobileNetV2 shows the highest validation accuracy, highlighting the benefits of pretraining and synthetic augmentation

5 DISCUSSION

Our results demonstrate that a lightweight MobileNetV2 backbone outperforms deeper networks on the KHWD task. We hypothesize that MobileNetV2's architectural efficiency and the benefit of targeted pretraining make it better suited for this low-resource scenario than larger models that may overfit. Pretraining isolated characters helps the model learn basic character shapes and stroke variations unique to Kurdish script, providing an effective basis for word recognition. The addition of synthetic word images further improves generalization by presenting the model to a greater diversity of writing styles and distortions. Figure 7 presents a comparative view of the character-level accuracy progression during training and validation across all five CRNN configurations. These plots reveal notable differences in convergence behavior and generalization ability among the CNN backbones.

To evaluate the performance of the model and the accuracy of predictions for each word across the 40 handwritten samples, we used a majority-vote analysis. The results show that 56.83% of words were having partial agreement ($\geq 50\%$ of samples matched the same prediction), while only 1.04% of words achieved full agreement across all samples. However, 42.13% of words have low agreement, because the difference in the handwriting styles indicate a large variation in recognition. The challenges that result from intra-word variability are highlighted by this analysis, which also shows how important it is to use a large, varied dataset to train an OCR model well. While words with low agreement seemed more complicated, those with near-perfect agreement were typically written more simply or consistently.

FUTURE WORK

Future research will focus on enhancing this work by exploring Transformer-based architectures like TrOCR and ViT-CTC hybrids. These models have shown higher accuracy in recognizing complex, cursive scripts due to their ability to capture long-range dependencies and contextual information, which are especially relevant to the Sorani Kurdish script. Additionally, we aim to include attention mechanisms within the CRNN framework to enable the model to focus on the most relevant parts in the input image. Such integration is expected to enhance the recognition accuracy, particularly for words with diacritical variations and ambiguous strokes. Finally, to lower semantic-level errors and increase OCR reliability in practical deployment scenarios, future research may also use language modelling or lexicon-based decoding.

CONCLUSION

The current study presented KHWD, a novel large-scale dataset for Kurdish handwritten word recognition, and used it to evaluate benchmark CRNN models with various convolutional backbones. The dataset, consisting of around 400,000 images covering 10,000 Sorani Kurdish words collected from thousands of writers, captures the diversity of handwritten styles and the complexity of the script. Our experiments show that a MobileNetV2-based CRNN, when pretrained on isolated characters and augmented with synthetic word images, achieves the best recognition performance (lowest CER and WER, highest word accuracy) on KHWD. This highlights the effectiveness of model

pretraining and synthetic data in environments with limited resources. We encourage releasing the KHWD dataset to support further research on Kurdish HTR and other low-resource handwriting recognition tasks.

REFERENCES

- [1] S. M. Shareef and A. M. Ali, "Deep learning-based digitization of Kurdish text handwritten in the e-government system," *Indonesian Journal of Electrical Engineering and Computer Science*, vol. 35, no. 3, pp. 1865-1875, 2024, doi: 10.11591/ijeecs.v35.i3.pp1865-1875.
- [2] R. M. Ahmed *et al.*, "Kurdish Handwritten character recognition using deep learning techniques," *Gene Expression Patterns*, vol. 46, p. 119278, 2022/12/01/ 2022, doi: <https://doi.org/10.1016/j.gexp.2022.119278>.
- [3] B. Yaseen and H. Hassani, "Making Old Kurdish publications processable by augmenting available optical character recognition engines," *arXiv preprint arXiv:2404.06101*, 2024.
- [4] I. A. Qader and T. A. Rashid, "Kurdish Handwritten Character Recognition System: Review of Methods and Progress," *Authorea Preprints*, 2024.
- [5] P. A. Abdalla *et al.*, "A vast dataset for Kurdish handwritten digits and isolated characters recognition," *Data in Brief*, vol. 47, p. 109014, 2023. [Online]. Available: <https://pmc.ncbi.nlm.nih.gov/articles/PMC10018436/>.
- [6] H. D. Majeed, G. S. Nariman, R. S. Azeez, and B. B. Abdulqadir, "Kurdish standard EMNIST-like character dataset," *Data in Brief*, vol. 52, p. 110038, 2024/02/01/ 2024, doi: <https://doi.org/10.1016/j.dib.2024.110038>.
- [7] A. Yadav, S. Singh, M. Siddique, N. Mehta, and A. Kotangale, "OCR using CRNN: A deep learning approach for text recognition," in *2023 4th International Conference for Emerging Technology (INCET)*, 2023: IEEE, pp. 1-6.
- [8] R. M. Ahmed, T. A. Rashid, P. Fatah, A. Alsadoon, and S. Mirjalili, "An extensive dataset of handwritten central Kurdish isolated characters," *Data in Brief*, vol. 39, p. 107479, 2022.
- [9] S. H. Ali and M. B. Abdulrazzaq, "KurdSet Handwritten Digits Recognition Based on Different Convolutional Neural Networks Models," *TEM Journal*, vol. 13, no. 1, pp. 221-233, 2024. [Online]. Available: <https://doi.org/10.18421/TEM131-23>.
- [10] A. M. Qadir, P. A. Abdalla, M. I. Ghareb, D. F. Abd, and K. M. HamaKarim, "An Ensemble Transfer Learning Model for the Automatic Handwriting Recognition of Kurdish Letters," *Iraqi Journal for Electrical and Electronic Engineering*, vol. 21, no. 2, pp. 54-63, 2025. [Online]. Available: <https://doi.org/10.37917/ijeec.21.2.6>.
- [11] C. A. Romein, A. Rabus, G. Leifert, and P. B. Ströbel, "Assessing advanced handwritten text recognition engines for digitizing historical documents," *International Journal of Digital Humanities*, 2025/05/12 2025, doi: 10.1007/s42803-025-00100-0.
- [12] B. Imane, A. Alae, K. Ghizlane, and M. Mrabti, "Enhancing Arabic handwritten word recognition: a CNN-BiLSTM-CTC architecture with attention mechanism and adaptive augmentation," *Discover Applied Sciences*, vol. 7, no. 5, p. 460, 2025/05/06 2025, doi: 10.1007/s42452-025-06952-z.
- [13] M. D. A. Cheema, M. D. Shaiq, F. Mirza, A. Kamal, and M. A. Naeem, "Adapting multilingual vision language transformers for low-resource Urdu optical character recognition (OCR)," *PeerJ Computer Science*, vol. 10, p. e1964, 2024.
- [14] A. Hamza, S. Ren, and U. Saeed, "ET-Network: A novel efficient transformer deep learning model for automated Urdu handwritten text recognition," *PLOS ONE*, vol. 19, no. 5, p. e0302590, 2024, doi: 10.1371/journal.pone.0302590.
- [15] S. Minz, R. Kanojia, T. Yadav, and N. Jayanthi, "Enhancing Accuracy in Handwritten Text Recognition with Convolutional Recurrent Neural Network and Data Augmentation Techniques," in *2023 Third International Conference on Secure Cyber Computing and Communication (ICSCCC)*, 2023: IEEE, pp. 803-808.
- [16] R. Dey, R. C. Balabantaray, S. Mohanty, D. Singh, M. Karuppiyah, and D. Samanta, "Approach for preprocessing in offline optical character recognition (ocr)," in *2022 Interdisciplinary Research in Technology and Management (IRTM)*, 2022: IEEE, pp. 1-6.
- [17] S. H. Ali and M.-B. Abdulrazzaq, "KurdSet: A Kurdish Handwritten Characters Recognition Dataset Using Convolutional Neural Network," *Computers, Materials & Continua*, vol. 79, no. 1, pp. 429-448, 2024. [Online]. Available: <http://www.techscience.com/cmc/v79n1/56307>.
- [18] J. A. Qadir, S. K. Jameel, W. O. Khudhur, and K. H. Manguri, "CKSD: Comprehensive Kurdish-Sorani Database," *IAPGOS*, vol. 1, pp. 153-156, 31.03.2025 2025. [Online]. Available: <http://doi.org/10.35784/iapgos.6521>.

- [19] H. Majeed and G. S. Nariman, "Offline Kurdish Character Handwritten Recognition (OKCHR) Using CNN With Various Preprocessing Techniques," *Research Square*, 2023 2023, doi: 10.21203/rs.3.rs-2565015/v1.
- [20] T. Mohammed and A. Ahmed, "Offline Writer Recognition for Kurdish Handwritten Text Document Based on Proposed Codebook," *UHD Journal of Science and Technology*, vol. 5, p. 21, 03/31 2021, doi: 10.21928/uhdjst.v5n2y2021.pp21-27.