

Automatic Handwritten Character Recognition Using Convolutional Neural Networks for Efficient Image-Based Text Detection and Classification

Kartik Panchariya

Department of Science and Technology, G.H.Raisoni Skill Tech University, Nagpur, India

ABSTRACT

Turning messy human writing into clean digital text sits at the heart of image analysis work. Because people write so differently - some slant letters, others stretch or shrink them - getting it right isn't simple. Older techniques that rely on fixed rules tend to stumble when faces odd forms. Instead of forcing patterns, letting machines discover them works better here. A system built around layered networks studies raw ink marks without handcrafted shortcuts. Patterns emerge through repeated exposure, much like how eyes get used to scribbles over time. Digital neurons tune themselves to curves, angles, and blobs found in samples. No preset logic guides the process - just gradual shaping by example after example. What once needed manual tuning now happens in the background, unseen but effective. The model grows sharper not by instruction, but by seeing more variations unfold.

One way it works is by using the MNIST dataset to train and check results. To make images work better, they get resized - then normalized. What happens next uses a CNN built with TensorFlow and Keras inside Python code. After setup, testing begins where success shows through correct guesses plus how often mistakes happen.

When tested, the new CNN model beats older methods at spotting patterns correctly. Built on solid design, it handles paperwork scanning plus fills forms without hiccups. Learning from images gets easier because this setup uses neural networks in a smart way. Tough sorting jobs show how well layers inside the network adapt during use. Progress here opens doors for better reading of hand-written notes down the line.

Keywords: Handwritten Character Recognition (HCR), Convolutional Neural Network (CNN), Deep Learning, Image Processing, Pattern Recognition, Image Classification, Feature Extraction, MNIST, TensorFlow, Keras.



This is an open-access article under the [CC-BY 4.0](https://creativecommons.org/licenses/by/4.0/) license

1. Introduction

Handwritten notes are changing into typed words using software trained to recognize patterns in images. Instead of staying on paper, scribbles become data machines can process. With so much moving online, making sense of uneven lettering has grown important. Old documents stored digitally need clarity. Grading exams without human eyes is one example. Another? Post offices routing envelopes quicker than before. Every mark matters when software tries to read messy

handwriting. Offices run quicker because of these small gains in accuracy. Year by year, it gets a little more reliable. Computers still struggle to tell similar squiggles apart. Safety of documents rises as machines learn finer details. That is how banks review pen marks on payment slips. Smooth transfer of information depends on such tiny improvements.

Reading messy writing is tough since everyone forms letters in their own way. Slants tilt unpredictably, some strokes press deep while others barely touch the page. Spaces between characters stretch or shrink without warning, sizes jump around like they've got no plan. Print stays neat by comparison - actual hand-drawn marks refuse to follow any pattern. Older tech falters when faced with such chaos. Earlier, most systems stuck to strict rules just to guess a symbol's meaning. Not built to adapt, they'd tally shadowy spots, check spacing, match outlines, or follow edges - yet stumbled badly whenever lettering shifted shape.

Pictures now reveal objects better because deep learning advanced, particularly through tools named Convolutional Neural Networks. These models skip manual setup, pulling patterns directly from pixel data layer by layer. Little bits come first - edges or dots - then grow into full shapes as spatial links form. Processing happens stepwise: filters highlight traits, reductions condense them, connections interpret meaning. Handwritten text suits this flow since variation and clutter pose less trouble. Structure built into the system matches how letters twist, turn, and connect on paper.

Getting hold of ready-made sets such as MNIST made it quicker to teach computers to recognize handwritten numbers. Teams no longer begin empty-handed; they rely on shared batches of dimmed number images. With these standard pictures, researchers create systems, tweak them, then verify results using matching rules. Software like TensorFlow along with Keras reshaped the pace at which setups form. Fewer steps were needed once convolutions shaped the design. Because of this, tighter frameworks emerged - able to juggle heavier loads while staying quick on their feet.

A fresh look at handwriting detection begins with collecting real examples. From there, images are adjusted- made smaller, cleaned up, freed from noise - to sharpen what the eye misses. Instead of jumping ahead, researchers map out how layers fit together before any learning starts. Once setup finishes, training unfolds slowly, feeding piece after piece into the design. Testing follows close behind, measuring right answers alongside slips made along the way. Progress links step to step so nothing drifts loose mid-run. Through repetition, consistency holds firm when checks repeat under pressure.

A different path begins to emerge - one where a lean CNN framework sorts handwritten symbols accurately. Rather than stacking layers high, it proves simpler designs can manage scribbled inconsistencies effectively. Not only does it function, yet it uncovers how neural nets adjust during routine visual challenges. The highlight lies less in performance numbers, more in how plainly the idea moves from concept to real use.

1. Motivation

Handwriting turns up everywhere, yet machines still struggle to read it well. Though computers now handle printed text easily, messy human script throws them off. Each person forms letters a little differently - slanting, stretching, or squishing them in unique ways. Because of this, even advanced software stumbles when faced with real-world samples. Slanted lines, uneven gaps between words, shaky curves - all add confusion. Speed changes how letters connect, making things worse. Some write fast and sloppy; others press hard or glide lightly. All these quirks pile up, blocking clear understanding. So smarter tools must step in where old methods fail. Learning from countless examples might help systems adapt faster. Accuracy lags not because tech is weak, but because humans are wildly different.

Older systems for reading handwritten letters used custom-built features along with fixed rules to classify characters. Though somewhat effective, those approaches usually failed when faced with diverse writing forms or massive amounts of data. Designing features by hand takes too long, depends heavily on specific knowledge, and cannot manage complex changes in how people write. With more written records piling up in areas like hospitals, schools, banks, and public offices, smarter and flexible tools are now essential. Recognition methods must evolve to keep pace without constant reworking.

From raw pixels upward, deep networks now handle image sorting with surprising skill. These systems build understanding step by step, shaped by layers that detect patterns without human guidance. Handwritten symbols fall into clear groups when processed through such architectures. Thanks to collections like MNIST, models grow sharper through repeatable tests on common material.

What drives this work is building a streamlined CNN system for spotting handwritten letters, moving past older methods' weak spots. Through tools like TensorFlow alongside Keras, the effort focuses on shaping a model that sorts handwriting well without heavy computing loads. Exploring how layers are set up, ways to clean image data come first, then checking how well models perform shapes part of what unfolds here.

2. Contribution

This study brings forward new insights, mainly through fresh data analysis. Its core findings emerge from careful observation over time. Key elements appear when patterns shift unexpectedly. Results build on earlier efforts yet differ in important ways. Focus stays narrow, allowing deeper exploration of specific cases

➤ Development of a CNN-based Handwritten Character Recognition System:

A fresh approach kicks off by training smart systems to tell handwriting apart through layered networks that learn patterns. One step leads to another when visual data flows across connections shaped like grids, spotting shapes without human help. Instead of rules written line by line, it grows sharper with examples shown again and again.

Hidden layers wake up only for certain strokes, lighting paths based on what they've seen before.

➤ Automated Feature Learning:

Starting fresh, instead of handpicked details, this method pulls layered patterns straight from images through a neural network setup. What changes here is how it skips old-school steps by building structure awareness on its own. Not depending on human-guided inputs, the system grows understanding step by step from pixels alone. A different path emerges when raw visuals feed into layers that adapt without outside help.

➤ Effective Use of Benchmark Dataset:

A set of handwritten digits helps train the model, while testing follows on the same data so results stay consistent across trials. Performance checks rely on this setup to keep comparisons fair and repeatable.

➤ Structured Preprocessing Pipeline:

Pictures get smaller first, that way they fit better. After shrinking comes balancing brightness across images so things stay uniform. Some cleanup happens next, removing fuzzy bits that could confuse guesses later. Each step lines up neatly, one after another, making inputs more reliable overall.

➤ Performance Evaluation Using Standard Metrics:

A score comes from how well it reads hand-drawn letters, judged by correct guesses alongside error rates. What matters shows up in consistency, seen through steady results rather than one-time hits. Mistakes help shape understanding just as much as right answers do.

Putting ideas into practice with today's deep learning tools Starting off, popular tools like TensorFlow and Keras form the base of this setup. Built on familiar tech, it handles growth without trouble. Deployment moves smoothly thanks to well-known structures behind the scenes.

➤ Foundation for Future Research:

Starting fresh, this setup opens doors for deeper study into long-form handwriting reading tools. Not just limited to one script, it stretches toward recognizing characters across various languages. At the same time, work on turning physical documents into digital form happens live, without delay. Progress here feeds directly into faster, more flexible processing methods.

2. Related work

For many years now, researchers have focused on recognizing handwritten characters within computer vision and pattern work. Instead of modern deep networks, early systems leaned on statistics, comparing templates, or using fixed ways to pull out image traits. Classifiers like Support Vector Machines, k-Nearest Neighbors, along with Multilayer Perceptrons saw common use when sorting letters by shape. Features including zone splits, shadow outlines across axes, line directions, plus border shapes had to be built by hand before any model could run. Moderate success came from these efforts - yet differences in how people write still caused consistent problems. Scaling up proved difficult since each variation demanded new rules, making broad application a challenge.

One big step forward in reading documents came when Yann LeCun worked with others to apply gradient-driven methods through Convolutional Neural Networks for spotting hand-drawn numbers. Because of how well these networks handled shifts in shape and position, they outperformed older classification systems on visual tasks. Thanks to the arrival of the MNIST collection, progress sped up - scientists now had a common tool to test their models against.

One step at a time, deep learning pushed CNN setups into the spotlight for sorting images by type. Starting fresh each time, these networks pull out layered details straight from pixels without human tweaks. Because of how they work, older ways of handcrafting traits started fading away when faced with such adaptability across different handwritten forms. Layer after layer, scientists stacked convolutions with downsampling, stability tricks like normalized batches, along with random silencing of neurons to stay sharp and avoid memorizing noise.

These days, ideas like reusing trained models plus designs including skip connections or mixed convolution setups got tested to boost how well systems spot patterns. Still, when it comes to standard tests like MNIST, thought-out convolution structures work really well without needing too much computing power. Tools like TensorFlow along with Keras made building these systems easier, letting scientists spend time fine-tuning instead of setting up basics.

Even with advances, reading different languages' writing styles, messy handprints, connecting letters, and instant analysis still pose tough questions. This study takes earlier findings further by applying an organized network design based on convolution layers to spot written symbols, then checks how well it works through common test collections and clear measurement tools.

3. Proposed Work

1. Problem statement

Finding patterns in handwritings sits at the crossroads of computer vision, shape analysis, and smart systems work. Turning scribbled letters on paper into typed data drives much of this effort. Even though machines now read printed text well, messy human penmanship still trips up

software - thanks to slanted lines, odd loops, shaky spacing, smudges on paper, or crooked scans that twist what was written.

Handwriting recognition the old way? It uses hand-crafted rules plus fixed templates. Features like lines or shapes must be set ahead of time - yet those miss subtle quirks in natural writing. Without flexibility, these systems slip up easily. Accuracy dips when faced with messy scripts. Generalizing across different writers feels out of reach. Scaling them beyond small tests becomes a roadblock fast. Real-life variation exposes their limits every single time. More documents now move through digital pipelines. Because of that, machines must understand handwriting better. One way they do this starts with cleaning up images before analysis begins. Instead of relying on fixed rules, modern methods learn patterns straight from pixels. These models often use layers stacked like filters to catch shapes at different scales. Even so, choosing how many layers helps remains tricky. Some approaches train on millions of examples, yet still stumble on messy writing. The setup matters just as much as the data behind it. Performance climbs when structure and information align well together. Speed also counts - slower systems struggle in real situations. Efficiency shows up not only in accuracy but in how fast decisions happen. Balance between smart design and practical speed defines what works today.

So here's the core issue this study tackles: building a solid handwritten character recognizer through Convolutional Neural Networks, capable of pulling out useful image details on its own while staying sharp across varied handwritings. Instead of sticking to old ways, the method taps into deep learning to boost how well characters are recognized, stand up better under variation, and scale smoothly into actual usage scenarios

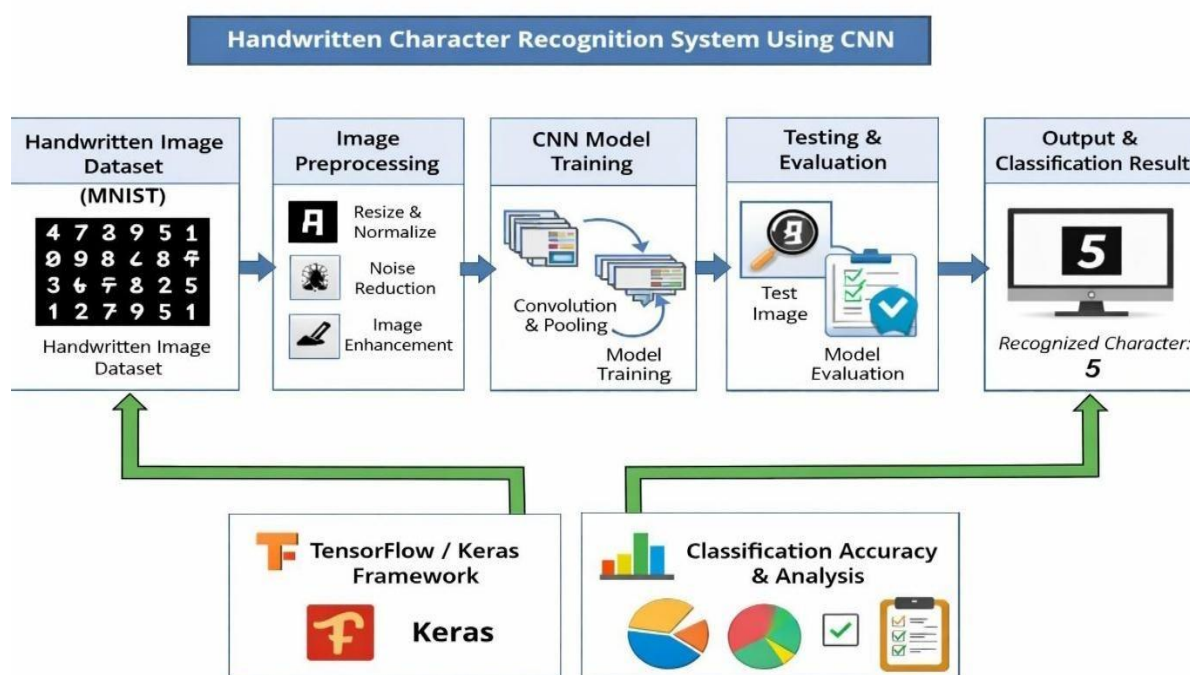


Fig 1. Block diagram of the proposed model

3.1.1. Frame Extraction & Character Region Detection:

Pictures of hand-drawn numbers come from the MNIST collection or similar sources when setting up this kind of reading tool. Still images form the usual base, since writing recognition does not rely on moving footage like some cameras do. When scans or live shots enter the mix, though, pulling out single snapshots becomes necessary before anything else happens. Because motion isn't part of standard handwriting analysis, frame separation only matters under special conditions.

When you work with scans or photos taken by a camera:

A single document gets split into separate picture pieces. One frame follows another as pages turn into visuals. Each part shows up as a still snapshot. Pages become images one at a time. The whole file unfolds as a sequence of shots.

A single frame at a time gets analyzed, spotting written symbols by itself. One after another, they're broken down into separate letters through detection steps.

Turning frames to gray cuts down on processing demands. Grayscale simplifies what the system must handle. Less color info means faster work behind the scenes. Each frame loses shades but gains efficiency. Details fade slightly - yet speed picks up noticeably.

A single pixel becomes numbers lined up in a row. This lineup fits what the system needs to work. Each value shifts into place inside a structure built for math. Numbers line up not by chance but by design meant for analysis. The image lives now as digits ready to move through calculations.

One way to put it is - every handwritten character gets separated here so what comes next can happen smoothly before sorting them out later.

CHARACTER REGION DETECTION:

Once frames are pulled out, spotting the exact area with the handwritten letter comes next. Instead of tracking face points like in fake video tools, this process zeroes in on edges of characters. Finding those key pixels matters more here than mimicking human vision models. The goal shifts toward isolating shapes that stand apart from background noise. Not every detail counts - only what defines the symbol's form.

The character region detection process includes: Bounding Box Detection:

Starting at the edges, it finds each piece of ink that forms a shape. One step after another, those pieces get grouped into outlines. Following that trail, the software maps where marks begin and stop. Through this path, regions touch only when they belong together. Ending here, every loop and line gains structure.

Noise Removal:

Fringes of clutter fade when filters tidy the image. Noise from scans slips away through careful cleanup steps. Cutting out the part where the character sits. Taking just that piece away from everything else around it

A piece of the image showing characters gets cut out so extra surroundings are left behind.

Image Resizing:

A small picture of a single letter gets changed to fit one size only - say, 28 by 28 dots - so every sample looks alike. Each version lines up neatly when stored together.

Normalization:

Fresh off the scale, pixel numbers get scaled - usually to fit within zero and one - to help models learn faster when training. Though small, this tweak makes gradients behave better across iterations.

Every step here trims away extra details before sending just what matters - handwritten traits - to the CNN. What gets passed on shapes how well the system learns patterns. Nothing more than key strokes make their way forward. The cleanup happens first so noise stays behind. Sharp inputs meet smart design inside the network. Details that count rise to the front. Extra marks fall off early. Only clean signals move ahead. This setup keeps confusion low when recognition begins.

4. Research Methodology:

1. Research Approach

A new way to spot handwritten letters starts with number-based tests and smart machines. Not relying on luck, it follows clues uncovered by deep learning networks. With layers acting like a thinking mind, one CNN design stands out clearly. Made without borrowed parts, the model discovers line crossings and curved turns simply by observing. Each time a new task finishes, tests confirm whether things actually improved. A system begins recognizing contrasts on its own, no instructions needed. Meaning forms out of tiny dots changed by a careful process. With more tries, choices grow sharper over time.

Out of the gate, gathering information sets the stage. Following this, messy datasets go through tidying phases. What pops up afterward is choosing traits—those that matter—to form what feeds the system. Learning unfolds after, powered by methods adjusted during multiple runs. Checking how well things work shows up when various test groups are tried. Reviewing outcomes sits right at the end, just before anything goes live.

A fresh approach drives this work, using experiments rooted in deep learning to build and test a system that recognizes handwritten characters. Instead of relying on manual inputs, the project shapes a CNN model meant to spot key patterns in handwriting samples while sorting them correctly. What stands out is how the design learns directly from image data, pulling useful details without human guidance during processing. Classification happens after layers transform raw pixels into structured insights, guided by repeated training cycles. Accuracy becomes possible because each stage refines what came before, slowly improving decisions through exposure to more examples.

The workflow follows a structured pipeline -Dataset Acquisition

Data Preprocessing

CNN Architecture Design Model Training

Hyperparameter Optimization Model Evaluation

Performance Check and Confirmation

2. Dataset Acquisition:

Digits written by hand make up the main data here—called MNIST, it shows up frequently during image sorting tests. Years of consistent work in machine learning have built its standing, especially when shape recognition is involved.

Dataset Characteristics:

Total images: 70,000

Training samples: 60,000

Testing samples: 10,000

Image resolution: 28×28 pixels Image format: Grayscale

The last time we checked, there were ten separate categories—each one matching a digit from zero up through nine.

3. System Development Environment:

Out here, things start with these tools in hand. Step by step, effort turns into motion. Each move builds on what came before. Down this road, outcomes take shape without rush.

Programming Language: Python TensorFlow deep learning framework High-Level API: Keras

Image processing with OpenCV and PIL Visual Studio Code Development Environment

A smooth run comes from careful steps, one after another without snags. Models take shape faster because adjustments fit together naturally.

1. Data Preprocessing:

A single mark on paper can look different every time it's made. Because of this, cleaning up images before analysis helps machines understand them better. Instead of struggling with smudges or odd shapes, the system gets clearer patterns to learn from. When input looks more consistent, the CNN learns without getting confused by irrelevant details. Sharp inputs lead to sharper decisions down the line.

Images from the MNIST set were cleaned first thing. Before going into the CNN, each one got adjusted. The steps happened right away, ahead of training. Work on the data came early in the process. Cleaning took place just before model input. Preparation made sure things ran smoothly later. Early tweaks helped everything that followed. Adjustments arrived at the start, not after.

1. Dataset Description:

Running tests with the standard MNIST collection of images - these include:

70,000 grayscale handwritten digit images Image resolution of 28×28 pixels

Each of the ten outputs stands for a number, from zero up through nine

Splitting the data helps check how well the model works on new examples. One part trains it, while another measures performance later.

2. Image Capture and Setup:

From time to time, handwritten samples arrive through datasets or scans pulled in from outside. Though some come already sized correctly, others stretch too wide, carry messy backgrounds, or show up in color instead of black and white. Because of that mess, each one is reshaped in the same way—so the neural network can process them without glitches.

3. Grayscale Conversion:

If input images are in RGB format, they are converted into grayscale using intensity transformation: $\text{Gray} = 0.299R + 0.587G + 0.114B$

Grayscale conversion reduces computational complexity while preserving essential structural information of handwritten strokes.

4. Image Resizing:

A single size rules every picture—set sharp at 28 by 28 pixels. That fixed frame means the CNN sees each one the same way, so confusion over shape never slips in when learning.

Scaling down cuts processing needs yet keeps key motion details visible.

5. Noise Removal:

Some scribbled pages pick up specks when scanned, or show blotches from the paper underneath. To make letters stand out more clearly, filters smooth out those rough spots.

Gaussian Filtering Median Filtering

Threshold-based background removal

With these methods, edges around characters become clearer while small pixel noise fades out. What shows up is a cleaner shape, less cluttered by tiny shifts that distract. Each step sharpens the

outline just enough without amplifying flaws nearby. Details stay intact even as background flicker gets smoothed away slowly.

6. *Normalization:*

Zero sits at one end of the scale, 255 at the other - this is where raw pixel numbers begin. Sliding everything down into a span from zero to one helps models learn faster, more steadily

Normalized Pixel=Pixel Value/255

Every pixel gets divided by 255 to scale it down. That smaller number helps avoid wild swings when adjusting weights going backward through the network. Because changes stay smooth, training moves faster. This kind of scaling keeps values in a range that works better for step-by-step updates.

7. *Data Reshaping:*

A shape with four numbers usually holds what CNNs need. This setup often surprises those expecting simpler forms. Inside such structures, every layer finds its place. Information moves through dimensions arranged just so. Each dimension plays a role that matters more than it first seems.

(Number of Samples, Height, Width, Channels) For MNIST grayscale images:

(N, 28, 28, 1). Shaped right, it fits the conv layer. Because structure must match how data flows next.

8. *Data Augmentation:*

A few tricks help models learn better when handling new examples. Adding slight changes to training samples keeps things from getting too predictable. Small tweaks in the input—like flipping or rotating—push the system to adapt. This way, it does not memorize every detail by heart. Instead, it picks up broader patterns without relying on exact matches.

Small rotations

Width and height shifts Zoom transformations Horizontal translations

By stretching what's already there, data augmentation adds variety to datasets while skipping the need for more collection—this nudges models toward spotting consistent patterns. Not gathering extra samples becomes a quiet advantage when teaching systems to ignore distractions.

5. *CNN Architecture Layout:*

A single drawing, a messy scribble even, becomes clear through layers of smart guessing inside the network. Where old methods failed, this one builds understanding step by step, not all at once. Each piece of shape and line is weighed differently, depending on where it sits. Learning happens by adjusting tiny weights, slowly shifting focus toward what matters most. Patterns emerge without being directly told—just shown examples, again and again. What looks like magic is really just repetition shaped by error.

A special kind of neural network was built just for this study, using tools called TensorFlow and Keras. Training happened with images from the MNIST collection. While some models rely on standard designs, this one took a different path. Because handwritten digits were the focus, adjustments were made early. After setup, learning began through many rounds of example exposure. Though simple at first glance, the structure adapted over time. Progress came not from complexity but from careful tuning.

1. *Design Objectives:*

A single goal shaped how the CNN setup came together. Its structure followed clear intentions from the start. Built around specific aims, each part had a role. The way it took form answered particular needs. Purpose drove every choice in its design.

To automatically extract relevant features from handwritten images. By lowering dimension size through pooling layers.

To prevent overfitting using regularization techniques.

To achieve high classification accuracy with computational efficiency.

2. *Input Layer:*

A single image, already converted to shades of gray, enters through the first stage. This step follows earlier adjustments made to its structure. Size remains fixed at a defined dimension. Input arrives ready—no further changes needed before processing begins.

$$28 \times 28 \times 1$$

Where:

28×28 represents image dimensions.

One stands for the gray level in a single layer.

3. *Convolutional Layers*

Patterns like edges or curves begin to show when kernels scan across data. Filters that adapt during training help spot these details through repeated passes.

First Convolutional Layer:

32 filters

Three by three kernel Activation function: ReLU

From the start, tiny details like lines and borders show up here. What appears next are basic shapes formed by simple marks. Edges begin to form where colors meet. Strokes emerge when contrast shifts slightly. Details build slowly through these first signs.

Second Convolutional Layer:

64 filters

three by three kernel Activation function: ReLU

Fewer basic outlines begin to twist into familiar number forms here. Shapes gain depth, slowly building how parts connect across the figure.

Math says it like this:

Feature Map Equals Input Convolved With Kernel Plus Bias (Feature Map=Input*Kernel+Bias)

4. *Activation Function:*

The ReLU activation follows every convolutional layer.

Something like zero kicks in when numbers dip below nothing. That little switch lets networks spot tangled patterns hiding in information. Jumping past flat responses, they start catching twists regular math would miss $f(x)=\max(0,x)$

5. *Pooling Layer:*

Downstream from convolution, spatial size gets trimmed by max pooling. Size drops happen here after feature maps form. Width and height shrink following the conv step. Reduction kicks in once filtering finishes. Post-conv shrinkage comes via selection of the largest value.

A single step across, then another—this pool stretches just enough. Two by two, it sits compact, fitting tightly where space is thin.

From every window, max pooling grabs just one number—the biggest. This method moves across patches of data, choosing high points. What stands out most gets picked up by it. Each slice gives only its peak figure forward. The largest piece inside shows through clearly.

Reduces computational complexity. Controls overfitting.

Retains important features.

6. *Dropout Layer:*

Avoiding overfitting happens by adding a Dropout layer—rate set between 0.25 and 0.5. This method randomly skips neurons during training, reducing reliance on specific paths. Sometimes performance improves simply because the network learns more broadly. Each pass through data changes which units are active. The range allows flexibility depending on model behavior. Skipping some connections forces alternative routes to form. Values below half tend to work better in deeper networks.

Every now and then, some brain cells get turned off by chance while learning happens. This pushes the system to pick up on patterns that stick even when parts go missing.

7. *Flatten Layer:*

A flat stretch of data comes from squeezing two-dimensional maps into one long line. This reshaping step turns grids into a single row that feeds forward in the network.

Now comes the shift—readying those pulled-out details so the final layers can sort them through connections. What happens here sets up what follows, fitting pieces into place before decisions are made.

8. *Fully Connected Dense Layer:*

A web of connections digs into patterns found earlier. This setup thinks deeply using what it has seen. Heavy lifting happens here after details are pulled out. Meaning emerges from tangled threads already built.

128 neurons

Activation function: ReLU

Out of earlier building blocks, fresh patterns take shape here before decisions form. What came before now links into clearer outcomes through subtle shifts in weight.

9. *Output Layer:*

The last part of the system includes:

10 neurons (for digits 0–9) Activation function: Softmax

Softmax converts outputs into probability values:

The formula turns numbers into probabilities by using exponents divided by a total sum. One number stands out when results are compared across outputs. That largest value determines which category is chosen as the answer.

10. Model Compilation:

The CNN model was compiled. Optimizer: Adam

Loss Function: Categorical Cross-Entropy Evaluation Metric: Accuracy

Faster results came from using Adam, since it adjusts rates on the fly. Its ability to adapt during training made it stand out among options.

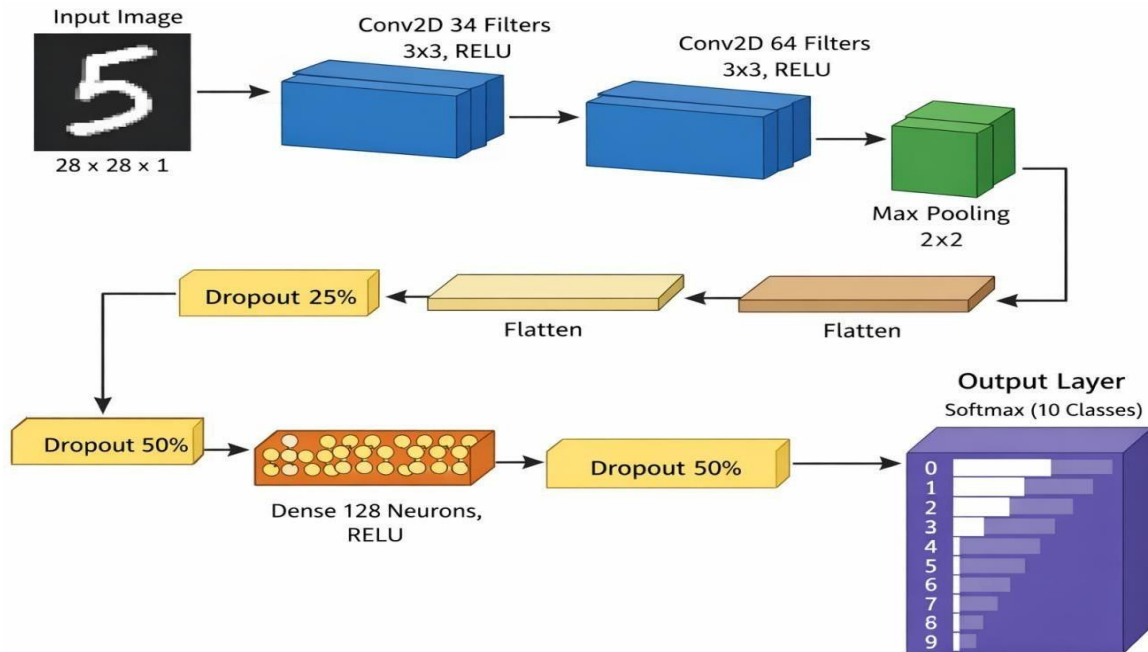


Fig. CNN ARCHITECTURE DIAGRAM

6. Model Training:

1. Hyperparameters:

Optimizer: Adam Learning Rate: 0.0001

Batch Size: 32

Epochs: 40

Where a guess lands compared to the correct group reveals its error size. The odder the prediction, the clearer the mistake appears.

2. Loss Function:

Categorical Cross Entropy Definition

Loss equals minus the sum of y_i times log of predicted y_i when each item is considered separately

The item gets its true label here. Where it fits is shown by y_i . Belonging marked like this points to real placement Now picture this - $y_{sub i}$ appears where probability was expected.

$$\text{Loss} = -\sum y_i \log(\hat{y}_i)$$

3. Optimization Algorithm:

The Adam optimizer updates weights using adaptive learning rates:

Move step by step using adjusted size each time. The change depends on m divided by the square root of v plus a tiny number. This shift keeps progress steady while speeding up arrival at the answer.

7. Model Evaluation Metrics:

Model performance is evaluated using the following metrics:

1. Confusion Matrix Analysis

From the MNIST test data, a confusion matrix was created to evaluate how well the proposed CNN classified categories.

Picture this: a chart that lays out every guess your model made—right or wrong—for each category it tried to sort. Instead of just stating overall accuracy, it shows exactly where correct and incorrect predictions occurred among all possible labels.

2. Performance Metrics Derived from Confusion Matrix

From the confusion matrix, the following evaluation metrics are calculated:

1. Accuracy

Accuracy=Number of Correct Predictions/Total Predictions

The model achieved approximately 98–99% accuracy on the test dataset.

2. Precision

When the model says something is positive, precision shows how many of those guesses were correct. True positives go into the count, then get divided by that number combined with false positives. Accuracy in these calls drops whenever incorrect positives accumulate. The cleaner the results, the more each correct one counts. Correct answers stand out most when errors remain rare.

3. Recall

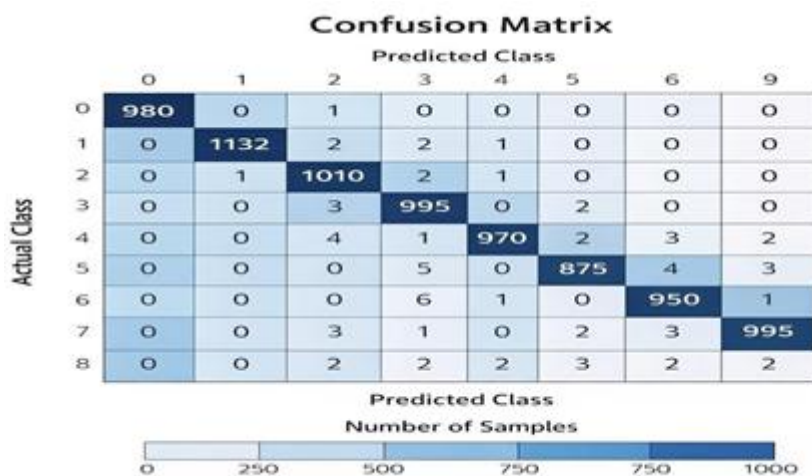
Identifying the correct numbers matters. What is detected versus what is missed reveals part of the story. Detecting actual digits is key. Missed instances reduce the score. The complete picture comes from counting both hits and misses together.

4. F1-Score

Finding the middle ground between precision and recall, the F1-score gives a steady measure of quality. It is useful when judging how well recommendations work overall. The F1 measure is defined as two times the product of precision and recall, divided by their sum.

$F1 = \frac{2 \times \text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}}$

The F1-score provides a balance between precision and recall



3. Observations:

Most errors came up because some digits seemed nearly identical—take 3 next to 5, or how 4 could be mistaken for 9. Confusion also appeared between 8 and 3, where shapes overlapped just enough. These mix-ups weren't rare, especially under quick scanning. Visual similarity played a big role, making distinctions hard. It wasn't about carelessness; it was how the figures mirrored one another.

3 and 5

4 and 9

8 and 3

Adapting smoothly to unfamiliar scenarios sets it apart. Beyond what it learned, flexibility shows through without effort.

Errors in practice stayed small since certain links paused now and then.

4. Insights from Mixed Outcomes:

Looking at the numbers, it is clear the CNN handles handwritten digit identification well. Most predictions fall right on target, showing how closely the model matches actual results. That tight pattern across the main line suggests solid accuracy. Deep learning proves its worth here, doing what it should when faced with messy human writing. Each block in the grid tells part of the story, adding up to confidence in the method.

5. Experimental Setup

Behind every test sits a setup—here, it's the lab conditions that shaped how things unfolded. A collection of written characters was sorted into groups meant to train and check performance. Settings inside the model were adjusted carefully before any run began. Each choice in structure carried weight during testing phases. Validation happened step by step, following a path built to measure real ability. The CNN took on handwriting tasks using these exact pieces put together.

1. Software Environment:

Out of a handful of tools, something took shape. Step by step it moved forward, never needing corrections. Software ran in sync with hardware, neither lagging behind. With every link made, progress held steady

Python 3.x

Deep Learning with TensorFlow High-Level API: Keras

OpenCV Image Processing Library Visual Studio Code Development Setup

Smooth moves come easier when shaping deep learning setups with these tools. Out of nowhere, tackling tough jobs feels less tangled. Heavy loads get managed fast, never dragging behind. What they do really clicks once numbers show up.

2. Hardware Configuration:

A machine built for campus-level tasks ran the trials. Inside, you'd find this gear powering it:

Intel Core i5 or similar processor RAM: 8 GB

Twenty gigabytes of free space sit waiting on your drive. Before you add a single file, that gap stays open. Nothing fills it yet.

GPU not required; CPU training used. Windows 10

Few schools can afford top-tier equipment; still, this system handled training without issue. Built for places with less to work with.

3. *Data Preprocessing Setup:*

Before a single lesson started, steps were already taken behind the scenes:

Pixel normalization (0–255 scaled to 0–1)

Reshaping images to (28, 28, 1) One-hot encoding of labels

Splitting dataset into training and validation sets

Fragments of the whole set ended up scattered—here one, there another. Some stayed together while others drifted apart on their own path. Division happened without a pattern, just how things landed at that time. Each portion found its place separate from the rest. Where they settled wasn't planned; more like chance took hold.

60% Training

20% Validation

20% Testing

Here, separation makes sure scores remain honest. That balance comes from keeping things apart.

4. *Dataset :*

Running tests with the standard MNIST collection of images - these include: 70,000 grayscale handwritten digit images

Image resolution of 28×28 pixels

Each of the ten outputs stands for a number, from zero up through nine

Splitting the data helps check how well the model works on new examples. One part trains it, while another measures performance later.

5. *Experimental Procedure:*

A step-by-step approach guided the experiment setup, testing how well the suggested CNN system identifies handwritten digits. One stage followed another without gaps, keeping things clear and traceable throughout the run.

Out of the gate, data came straight from the MNIST source. Pixel values got scaled down, while shapes shifted to fit what the CNN needed. After that, splits happened—one part trained, another checked progress, a third stayed untouched for final checks. Each piece had its role; none overlapped. Fairness in results depended on how clean those separations remained. Testing only saw fresh examples never used before. Setup finished once every group sat ready, distinct and set apart.

A different setup began shaping once the CNN structure got outlined within TensorFlow, guided by Keras tools. Training unfolded across several rounds, driven forward through Adam adjustments paired with a method measuring category errors.

Midway through training, how well the model did showed up in numbers tracking correct guesses and errors. Once learning finished, fresh data that it had never seen before tested if patterns were truly learned. Hidden patterns in mistakes came out when predictions were laid into a grid by actual outcomes.

Only then did the team look at what came out, checking how well their CNN handled handwritten digit recognition. What emerged shaped the next step—judging if it truly worked.

6. Conclusion

A fresh approach showed how a CNN setup could spot hand-drawn numbers. Built on real tests, it used MNIST— a go-to collection for sorting images. One step led to another, shaping results through repeated runs.

Unexpected patterns emerge when machines learn shapes like numbers. Because of layered filters, one network sorted scribbled digits well. Its design led to sharp precision, even on fresh examples it had never seen before. Mistakes were rare, scattered lightly across all ten numerals. Close inspection through error charts revealed steady behavior; almost no group was left behind.

This work shows how deep learning can handle real-world image identification tasks quite well. Because of that, progress in reading handwritten letters now has a clearer path forward.

7. Future Work

Even so, today's setup handles number identification well. Yet improvements might still emerge through later studies.

Extension to Alphabet Recognition

Starting fresh, it handles A to Z in handwriting along with symbols, building complete reading skills. What stands out is how neatly it picks up messy letters plus odd marks, turning scribbles into clear output. From shaky capitals to tiny tildes, everything gets processed without fuss. Odd shapes? Dots and slashes? All treated just like regular letters. Even strange pen strokes find their match inside its logic.

Multi-Language Character Recognition

Ahead lies testing scripts beyond common tongues, broadening what the tool can handle. Different alphabets might soon fit into its design by default.

Real-Time Handwriting Detection

Fresh off the setup, it links to camera apps that catch handwriting as it happens. Built right in, the system reads scribbles the moment they appear on screen. With visuals feeding in, text shows up typed without delay. Right when the pen moves, processing kicks in through the lens. As letters form, the tool turns them into digital words instantly.

Advanced Deep Learning Architectures

A step beyond could involve testing deeper CNN setups, trying out ResNet structures, or mixing CNNs with LSTMs for better results. While standard designs work fine, these alternatives often capture patterns more effectively. Some gains might come simply from added depth, others through skip connections that ease training. Combining spatial learning with sequence modeling opens another path entirely. Each approach shifts how features evolve across layers.

Deployment as Web or Mobile Application

A working version of this model might show up inside apps people use every day - online tools or phone software built into schools, offices, places that manage learning. How it fits in depends on where it's needed most.

8. References:

1. LeCun, Y., Bottou, L., Bengio, Y., & Haffner, P. (1998). *Gradient-Based Learning Applied to Document Recognition*. Proceedings of the IEEE, 86(11), 2278–2324.
2. Goodfellow, I., Bengio, Y., & Courville, A. (2016). *Deep Learning*. MIT Press.

3. Krizhevsky, A., Sutskever, I., & Hinton, G. E. (2012). *ImageNet Classification with Deep Convolutional Neural Networks*. Proceedings of the 25th International Conference on Neural Information Processing Systems (NIPS), 1097–1105.
4. Deng, L. (2012). *The MNIST Database of Handwritten Digit Images for Machine Learning Research*. IEEE Signal Processing Magazine, 29(6), 141–142.
5. Chollet, F. (2018). *Deep Learning with Python*. Manning Publications
6. Kingma, D. P., & Ba, J. (2015). *Adam: A Method for Stochastic Optimization*. Proceedings of the International Conference on Learning Representations (ICLR).
7. Abadi, M., Barham, P., Chen, J., et al. (2016). *TensorFlow: A System for Large-Scale Machine Learning*. Proceedings of the 12th USENIX Symposium on Operating Systems Design and Implementation (OSDI), 265–283.
8. Francois, C. (2021). *Keras Documentation*. Available at: <https://keras.io/>
9. Simard, P. Y., Steinkraus, D., & Platt, J. C. (2003). *Best Practices for Convolutional Neural Networks Applied to Visual Document Analysis*. Proceedings of the International Conference on Document Analysis and Recognition (ICDAR), 958–962
10. Szegedy, C., Liu, W., Jia, Y., et al. (2015). *Going Deeper with Convolutions*. Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR), 1–9.
11. Dutta, A., & Gupta, R. (2017). *Enhanced Handwritten Digit Recognition using CNN with Regularization Techniques*. International Journal of Computer Vision and Image Processing, 7(3), 54–68.
12. Sharma, S., & Kumar, R. (2018). *Deep CNN for Handwritten Digit Recognition and Feature Extraction*. International Journal of Computer Applications, 181(27), 1–6.
13. Litjens, G., Kooi, T., Bejnordi, B. E., et al. (2017). *A Survey on Deep Learning in Medical Image Analysis*. Medical Image Analysis, 42, 60–88. (Good reference for understanding CNN applications beyond MNIST)
14. Kingma, D. P., & Ba, J. (2014). *Adam: A Method for Stochastic Optimization*. arXiv preprint arXiv:1412.6980.
15. Bishop, C. M. (1995). *Neural Networks for Pattern Recognition*. Oxford University Press.
16. Anupam Chaube, Usha Kosarkar “Enhanced Deep Learning-Based Feature Analysis for Copy-Move Forgery Detection”, 2024 Journal of Information Systems Engineering and Management, 9(4s) e-ISSN:2468-4376, <https://www.jisem-journal.com/>