

TECHNICAL SCIENCES

AN INTELLIGENT WEB APPLICATION FOR DRUG INFORMATION RETRIEVAL USING OPTICAL CHARACTER RECOGNITION AND DEEP LEARNING

Ghazaryan A.

PhD in Economics, Associate professor at Chair of Business Information Systems and Information Technologies, Armenian State University of Economics, Yerevan, Armenia, ORCID: 0000-0001-6083-5489

Hovhanisyan J.

Graduate student at the Department of Computer Science and Statistics, ASUE, Yerevan Armenia, ORCID: 0009-0000-4398-4263

Kirakosyan A.

Assistant Lecturer at Chair of Business Information Systems and Information Technologies, Armenian State University of Economics, Yerevan, Armenia, ORCID: 0000-0003-4420-2993

Aramyan A.

Lecturer at Chair of Business Information Systems and Information Technologies, Armenian State University of Economics, Yerevan, Armenia, Email: ORCID: 0009-0006-6325-7438

Soomela Moosa Moghadam

*Master's in management (IT in business), ASUE alumni at the Department of Computer Science and Statistics, Yerevan Armenia, ORCID: 0009-0000-1454-581
<https://doi.org/10.5281/zenodo.15024999>*

Abstract

In a period of swift technological progress, access to accurate pharmaceutical information is similarly significant for healthcare professionals and patients. This paper explores the development of a web application that leverages machine learning, particularly optical character recognition (henceforth, OCR) and deep learning, to provide rapid, reliable access to drug information by scanning medication packaging. The application uses neural network models, which include Convolutional Recurrent Neural Networks (hereafter, CRNN) and ResNet architectures, to improve text recognition accuracy under various conditions, such as different lighting and font styles. Python-based frameworks, namely EasyOCR and PaddleOCR, enabled efficient and thorough text extraction and processing. This system streamlines the typically lengthy process of searching for medications, supporting informed decision-making, and mitigating the risk of errors. Additionally, integrated databases and Flask web development tools ensure a simplified and user-friendly interface. This innovative approach to pharmaceutical information retrieval can greatly improve medication management, patient safety, and the overall efficiency of healthcare delivery.

Keywords: Medicine Information, Optical Character Recognition (OCR), ResNet Model, Convolutional Recurrent Neural Network (CRNN), PaddleOCR / KerasOCR, EasyOCR

Introduction

In today's world, where technology is developing rapidly and the requirements for efficient solutions are constantly increasing, the need for quick and convenient access to pharmaceutical information is also growing. Technological advances, particularly in machine learning and healthcare integration, offer new opportunities to improve access to this information. Applying neural networks to recognize text in images is at the forefront of machine learning. It provides more accurate and reliable identification of drug information, even under different lighting conditions and font styles.

Developing applications that meet the needs of modern society is the key requirement of this era. The development of a web application that allows users to search for information about drugs through photos of their packaging greatly simplifies and speeds up the search process. Healthcare systems around the world are facing the challenge of streamlining increasingly pressured processes and improving patient care, which is why the need for innovative solutions that enhance medication management and safety is becoming increasingly apparent. Traditional methods of accessing

drug information - printed drug instructions or consultations with healthcare professionals - are often time-consuming and burdensome. In this context, the development of a user-friendly application that enables users to get comprehensive details about their medication quickly and accurately through scanning is an important advancement of the times. It is particularly beneficial for those who need rapid identification of pharmaceutical products. In the healthcare context, drug information must be accurate and reliable. The development of robust text recognition models and their integration into the web application following quality standards ensures the reliability of the information provided. Creating a web application that allows users to interact with the system through photos increases user-friendliness, which is particularly relevant in situations where users may have limited medical knowledge, providing a more accessible way to access information.

Furthermore, quick access to drug information through a photo-based web application reduces the need for manual searches, facilitating decision-making for healthcare professionals, pharmacists, and patients.

The work aims to develop a web application to search for information about medicines using a photo of the packaging. The search should be performed using the name recognized from the drug directory photo. To complete the task, it is necessary to use neural networks to recognize the drug's name, analyze the measurements, evaluate the results of the detection and text recognition models, and improve these indicators for both models. The information will be retrieved from the drug directory. For the implementation of the mentioned goal, the following tasks were outlined:

- study existing approaches to solve the problem of text recognition from photos, and conduct comparative analysis;
- search for pre-prepared models for solving a given problem, study their working principles, review articles and scientific publications;
- collecting data to test neural networks, and conducting experiments to identify the best models;
- conducting experiments, developing, finalizing, and improving models.

Literature review

In recent years, rapid advancements in information technology have made significant impacts on healthcare, particularly regarding access to reliable data. Digital platforms have become essential for managing vast volumes of information in healthcare. Studies show that these platforms increase accuracy in diagnosis and treatment, allowing healthcare providers and patients alike to access medical information promptly. The manuscript indicates that machine learning-enabled applications are crucial to improving service efficiency and reducing human error, particularly in medication management.

In "Information and Access to Health Care: Is There a Role for Trust?" Michael Thiede examines the crucial role trust plays in healthcare access and the exchange of information between patients and providers (Thiede, 2005). Thiede argues that trust in healthcare institutions and professionals is fundamental for ensuring that patients seek care and adhere to medical guidance. He discusses how trust affects not only patient behavior but also healthcare outcomes, as individuals are more likely to disclose important health information when assured of confidentiality and competence. The paper highlights disparities in healthcare access, noting that communities with higher trust in healthcare systems generally experience better health outcomes. Thiede ultimately suggests that building trust within healthcare systems may mitigate inequalities in access and enhance the overall effectiveness of healthcare services.

In *Digital Transformation in Healthcare: Technology Acceptance and Its Applications*, authors (Stoumpos et al., 2020) explore how digital transformation is reshaping healthcare through new technologies and changing patient and provider interactions. The authors emphasize the importance of technology acceptance in healthcare settings, noting that user acceptance is critical for successful implementation and patient satisfaction. The study reviews various digital tools, such as electronic health records, telemedicine, and mobile health applications, which have improved accessibility

and efficiency in healthcare delivery. By analyzing factors influencing technology acceptance, including perceived ease of use and trust in digital tools, the authors provide insights into how healthcare providers can encourage digital tool adoption. The editors also highlight the potential for digital transformation to reduce healthcare costs and improve patient outcomes through data-driven, efficient healthcare practices.

Based on Kayethri, Dharunya, and Harini's study on "Recognition Holic - Medicine Detection Using Deep Learning Techniques" highlights the advancement of deep learning applications in healthcare, particularly in medicine detection (Kayethri et al., 2022). Recent developments in neural network architecture have shown promising results in accurately recognizing pharmaceuticals through image data, which the authors leverage by training their model on a wide array of medication images. Previous studies have underscored the limitations of traditional detection methods, such as barcode scanning and manual checking, which can be both time-consuming and prone to error. This paper contributes to the field by proposing a model capable of adapting to various types of medication, ensuring a high level of accuracy and robustness in diverse pharmacy and healthcare settings. The research paper ultimately emphasizes the potential for deep learning to revolutionize medicine recognition, supporting safer and more efficient pharmaceutical management practices.

In the healthcare industry, especially within pharmacies, ensuring accurate prescription dispensing is essential for patient safety. Previous studies have noted that human errors in medication dispensing can lead to severe health complications, emphasizing the need for automated systems to support pharmacists in verifying prescriptions accurately. Advances in deep learning techniques, particularly in image and text recognition, have paved the way for automated verification systems that reduce these errors. Convolutional Neural Networks (henceforward, CNN) are widely used in image classification tasks, as they can learn intricate patterns within image data, making them suitable for recognizing visual features of medication blister packs. Combining CNNs with Histogram of Oriented Gradients (HOG) for feature extraction has shown promising results in pattern recognition, with high accuracy in image classification models. In addition to visual identification, text extraction techniques such as OCR are applied to read imprints on medication packages. Studies have demonstrated that methods like Character Region Awareness for Text Detection (CRAFT) and Keras-OCR improve text recognition accuracy, particularly when combined with text correction algorithms. By integrating both image and text classification models, an automatic system can leverage majority voting mechanisms to achieve a balanced, high-confidence outcome in prescription verification. Recent work, including the work by Maitrichit and Hnoohom (2020), indicates that this dual-model approach can achieve over 90% accuracy, effectively identifying and verifying medications. These systems have great potential to enhance patient safety and improve

pharmacy operations by minimizing errors in the prescription dispensing process

According to another study (N. Maitrichit and N. Hnoohom, 2020) in the healthcare industry, especially within pharmacies, ensuring accurate prescription dispensing is essential for patient safety. Previous studies have noted that human errors in medication dispensing can lead to severe health complications, emphasizing the need for automated systems to support pharmacists in verifying prescriptions accurately. Advances in deep learning techniques, particularly in image and text recognition, have paved the way for automated verification systems that reduce these errors. CNN is widely used in image classification tasks, as they can learn intricate patterns within image data, making them suitable for recognizing visual features of medication blister packs. Combining CNNs with HOG for feature extraction has shown promising results in pattern recognition, with high accuracy in image classification models. Besides, text extraction techniques such as OCR are applied to read imprints on medication packages. Studies have demonstrated that methods like CRAFT and Keras-OCR improve text recognition accuracy, particularly when combined with text correction algorithms. By integrating both image and text classification models, an automatic system can leverage majority voting mechanisms to achieve a balanced, high-confidence outcome in prescription verification. Recent research, including the work by Maitrichit and Hnoohom (2020), indicates that this dual-model approach can achieve over 90% accuracy, effectively identifying and verifying medications. These systems have great potential to enhance patient safety and improve pharmacy operations by minimizing errors in the prescription dispensing process.

Automating data-filling processes in web systems has gained traction due to the inefficiency and error-prone nature of manual data entry, particularly when dealing with complex form images. Traditional data uploading methods that require manually inputting form data are not only time-intensive but also increase the risk of inaccuracies, especially when handling similar-looking fields or irrelevant edge content. OCR has become an essential technology for converting images into text, allowing web systems to process and digitize form images more effectively. Nonetheless, OCR alone often struggles to accurately parse complex images due to issues such as overlapping fields or low-quality image components. Recent advancements, like the method proposed by Su, Kang, and Fan (2024), combine OCR with multi-text similarity measures to improve data extraction accuracy in complex forms. Techniques such as the Levenshtein distance metric provide a reliable approach for detecting and correcting OCR output by measuring text similarity, enabling better alignment with intended data fields. Studies indicate that integrating OCR with these text similarity measures can significantly enhance the accuracy of data extraction and field mapping. Practical applications of such hybrid methods in web systems have demonstrated data-filling accuracies exceeding 90% in complex forms. These improvements underscore the potential of combining OCR and advanced text-match-

ing methods to automate data-filling processes efficiently. This research area continues to evolve, providing promising solutions for streamlining data processing in web applications.

With the growing complexity of medication regimens, especially for cancer patients and individuals with multiple comorbidities, the risk of drug-drug interactions (DDIs) has become a critical concern. Traditional methods of identifying contraindications are often manual, time-intensive, and prone to human error. Advances in AI and natural language processing have led to innovative tools designed to automate this process and improve medication safety. OCR is a key technology for reading and digitizing text on medication packaging, enabling automated systems to extract necessary information. Nevertheless, OCR alone may not be sufficient for accurately extracting data from complex images, which is where multi-text similarity techniques, as discussed by Su, Kang, and Fan (2024), provide enhanced precision in identifying relevant text fields. Integrating OCR with multi-text similarity and language models offers a promising solution for accurate data extraction in medication management systems. The MEDIC chatbot, as presented by Kim et al. (2024), uses OCR alongside ChatGPT to identify DDIs, employing the Langchain framework to structure and deliver contextually relevant information. This combination allows MEDIC to recognize drug names from packaging images and provide real-time, concise information on potential contraindications, contributing to safer medication practices. Early studies demonstrate that AI-driven systems like MEDIC are highly accurate and user-friendly, making them valuable tools in clinical and home care settings. These advancements underscore the potential of AI to transform medication management, providing healthcare providers and patients with an efficient means of preventing adverse drug interactions.

To conclude, the reviewed literature underscores the transformative potential of digital technologies in healthcare, particularly in medication management and data processing. The integration of machine learning, OCR, and AI-driven applications has led to significant advancements in error reduction, efficiency, and accessibility. Trust and user acceptance remain fundamental to the success of these systems, as patient engagement and adherence are essential for effective healthcare delivery. Studies consistently show that machine learning, especially deep learning techniques like CNNs combined with OCR and text similarity measures, offers reliable solutions for automated data extraction and medication identification. Recent innovations, such as hybrid systems that integrate image recognition, OCR, and NLP (e.g., the MEDIC chatbot), demonstrate promising results in improving patient safety through efficient drug-drug interaction detection. Collectively, these advancements highlight the growing importance of AI in supporting healthcare providers and patients by ensuring safe, accurate, and accessible medication management, with ongoing analysis aimed at further enhancing these outcomes in clinical and home settings.

In summary, it can be said that currently, machine learning allows us to solve many problems that seemed

unattainable to humanity not so long ago. Solving complex problems thanks to well-known algorithms and models opens many doors to human society, offering people to make their lives easier. Experience shows that in this rapidly developing period, it is necessary to adapt to new technologies as quickly as possible, so as not to lag behind the rhythm of life and not be afraid of new challenges of science, but to accept challenges with honor and diligently develop with the times. If new technologies are used, considering them as an integral part of our lives, we will obtain unlimited new and useful knowledge from science, and by applying them, we will become more competitive and will not give up our place in technology.

Research Methodology

The research methodology for this study is structured around the design, development, and evaluation of a web application that retrieves drug information using EasyOCR and deep learning techniques. The application consists of three main modules: image acquisition, OCR processing, and deep learning-based information retrieval. Users upload or capture images of drug labels, packaging, or prescriptions, which are then processed through EasyOCR for text extraction. To support this system, a comprehensive dataset of high-quality images of drug packaging and labels was compiled from various sources, including publicly available pharmaceutical databases. Each image was preprocessed to optimize OCR performance, including steps like grayscale conversion, noise reduction, and resizing to meet OCR model requirements.

The EasyOCR tool was selected due to its capacity to handle multiple languages and fonts, which is crucial for recognizing text on diverse drug packaging. For optimal performance, EasyOCR was fine-tuned to recognize specific terms in pharmaceutical labeling, helping improve accuracy in challenging cases like blurry or low-resolution images. After text extraction, the data underwent post-processing, applying regular expressions to correct misrecognized characters and ensure text consistency. The next stage involved designing a deep learning model based on a CNN to analyze and classify the extracted text data. The CNN was trained on a labeled dataset that included relevant drug information, such as drug names, dosages, and usage instructions.

To ensure the model's reliability, the dataset was split into an 80/20 ratio for training and validation. Evaluation metrics, including accuracy, precision, recall, and F1-score, were calculated to measure the

model's performance on validation data. Several hyperparameter tuning techniques were applied, adjusting learning rates, batch sizes, and regularization methods to refine model performance. The application was designed as a microservices architecture, with the OCR and deep learning components deployed as separate services for improved scalability. The front end was built to offer users an intuitive interface where they can easily upload images and view retrieved drug information in real time.

For enhanced usability, error handling mechanisms were implemented to log OCR and classification errors, providing users with options to refine or correct search results. Testing was conducted in multiple phases, starting with functional testing to ensure all core features operated as intended. Performance testing evaluated response times and load handling, especially critical for an application designed for quick information retrieval. A usability test phase was conducted with a small group of users to gather feedback on the interface, ease of use, and system response accuracy. The user feedback was analyzed, and iterative improvements were made to enhance the application's functionality and user experience. Moreover, an accuracy test compared the application's retrieval success rate against existing drug information retrieval systems, demonstrating the efficacy of the OCR and deep learning approach. To support error resolution, logs were maintained to monitor OCR extraction and classification accuracy continually. Regular updates were applied to the OCR and deep learning model as new drug information became available, ensuring the system remained up-to-date with current pharmaceutical data.

EasyOCR is an open-source OCR library that aims to provide an easy-to-use interface for text detection and recognition tasks. EasyOCR is designed to recognize text in images and convert it into machine-readable text. EasyOCR typically uses deep learning models, such as CNN or RNN, to perform text detection and recognition tasks. These models are trained on large datasets of annotated images to learn to accurately detect and recognize text in different environments and languages. EasyOCR, as the name suggests, is a Python package that allows for effortless optical character recognition⁹⁷.

A diverse set of text images is fundamental to getting started with EasyOCR. It helps the OCR system handle a wide range of text styles, fonts, and orientations, increasing the system's overall performance.

⁹⁷ <https://pyimagesearch.com/2020/09/14/getting-started-with-easyocr-for-optical-character-recognition>

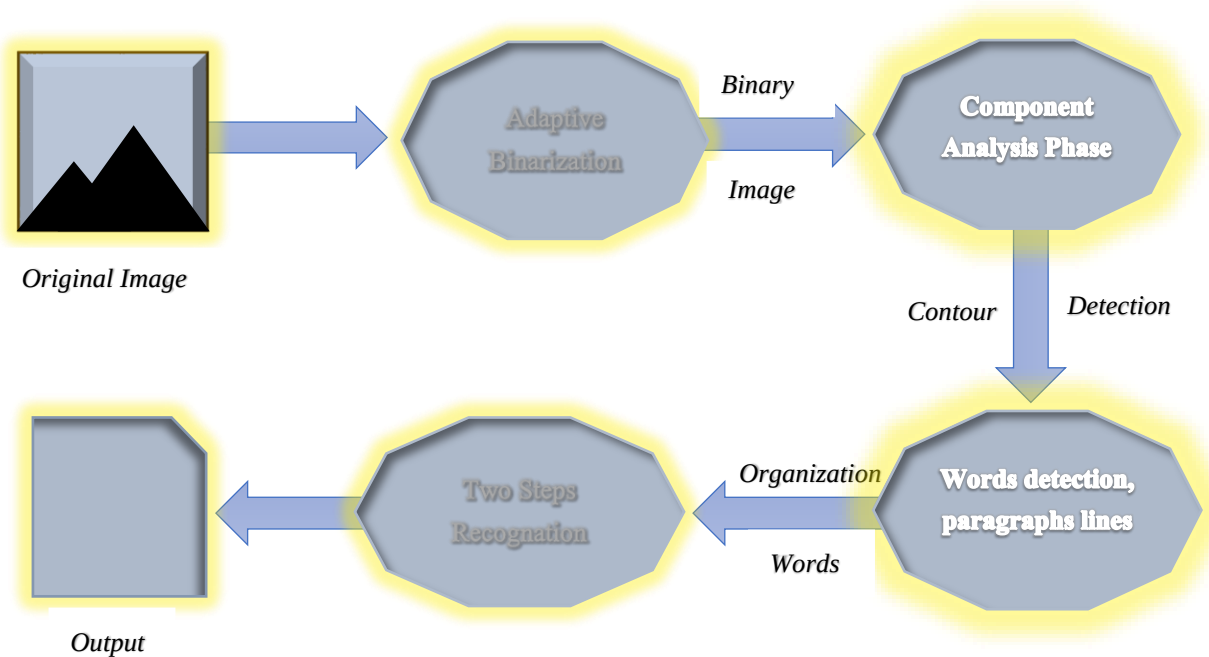


Figure 1: Steps of EasyOCR

Email validation is a method that determines whether an email address is available and valid. This implements a quick process that detects typos, whether they are "accidental" errors or deliberate misdirection. This tool checks email addresses at the point of database entry.

In a nutshell, the research methodology employed in this study successfully integrated image processing, OCR, and deep learning to create an innovative web application for drug information retrieval. The system leveraged EasyOCR for robust text extraction across a diverse range of drug packaging, augmented by preprocessing techniques that optimized OCR accuracy. By training a CNN on a well-curated dataset, the application achieved reliable classification of essential drug information, such as drug names and dosages. Careful dataset splitting, model evaluation, and hyperparameter tuning contributed to a refined, high-performing deep-learning model.

The modular design, encapsulated within a micro-services architecture, allowed for scalable deployment and seamless interaction between OCR and classification components. A user-friendly front end, complemented by thorough error handling and logging mechanisms, provided a reliable and accessible platform for users. Rigorous testing phases, including functional,

performance, and usability assessments, further validated the application's robustness and responsiveness.

Comparative accuracy tests demonstrated the application's efficacy relative to existing drug information retrieval systems, showcasing the potential of integrating OCR and deep learning in the pharmaceutical domain. The methodology also emphasized continual improvement through user feedback analysis and model updates, ensuring long-term relevance and alignment with evolving drug data. Overall, this methodology lays a strong foundation for further advancements in digital health applications and drug information retrieval.

Analysis

In practical scenarios, predefined form templates are frequently used to streamline data collection processes. For instance, standardized templates are commonly employed in activities such as opening a bank account, processing insurance claims, or administering school examinations. These templates maintain consistent field positions and dimensions, enabling predefined formats to guide the data input and recognition process⁹⁸. Since the layout remains uniform, the form structure can be defined beforehand, allowing for direct mapping and recognition of specific fields using the template. The procedure for identifying and matching these fields typically involves the following steps:

⁹⁸ Kumar, P.; Revathy, S. An Automated Invoice Handling Method Using OCR. In Proceedings of the Data Intelligence

and Cognitive Informatics: Proceedings of ICDICI 2020, Tirunelveli, India, 8–9 July 2020; pp. 243–254

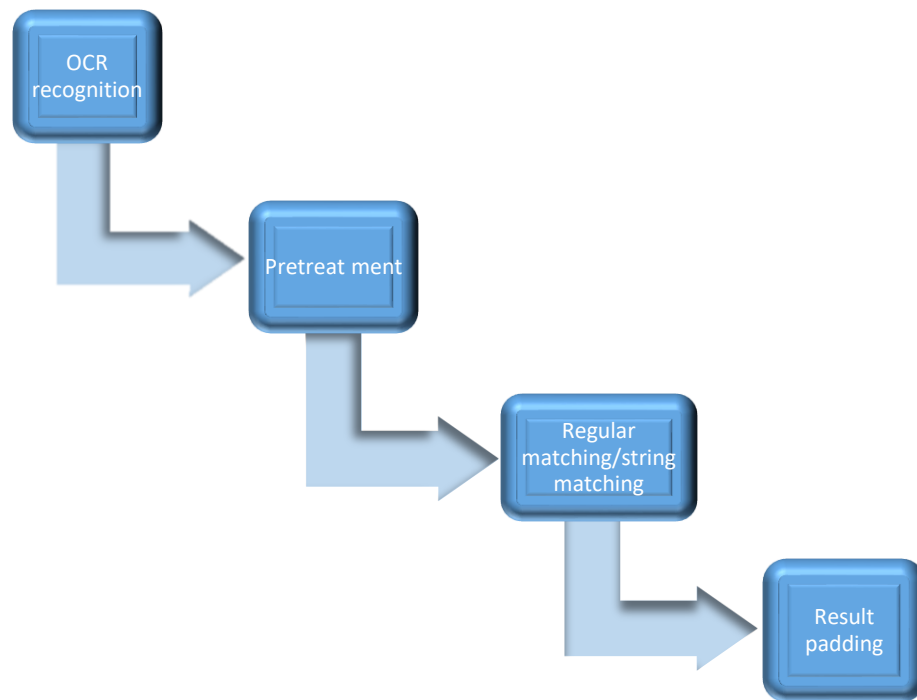


Figure 2. Field matching step diagram

After performing OCR-based image recognition, a set of extracted text data is generated, which includes both task-relevant information and extraneous characters such as spaces and line breaks. To enhance usability, this text data must undergo preprocessing. The preprocessing involves eliminating unnecessary elements like spaces and line breaks, as well as reformatting specially structured data. Importantly, the preprocessing step also preserves the positional information of the corresponding fields, which represents their relative arrangement within the form. This positional data is crucial for understanding the form's structure and context. For tasks involving predefined templates, this approach significantly simplifies the complexity of field matching. By utilizing the predefined field positions, the matching process can be confined to a narrower scope. Since predefined templates often contain fields with

similar structures and attributes, field matching can be efficiently executed on a smaller scale. This process relies on field names defined within the template and employs filtering techniques, such as regular expressions or string-matching algorithms, to locate relevant data.

Experimental Results and Analysis

Based on the existing web system development framework, integrating OCR technology and text similarity algorithms, and improving the similarity comparison process of corresponding fields according to the actual functional needs of the system, we can achieve our proposed goal of automatically filling in images for different frames of data forms on corresponding web pages containing data content. As shown in Figure 4 below, this is the overall structure of the system functions.

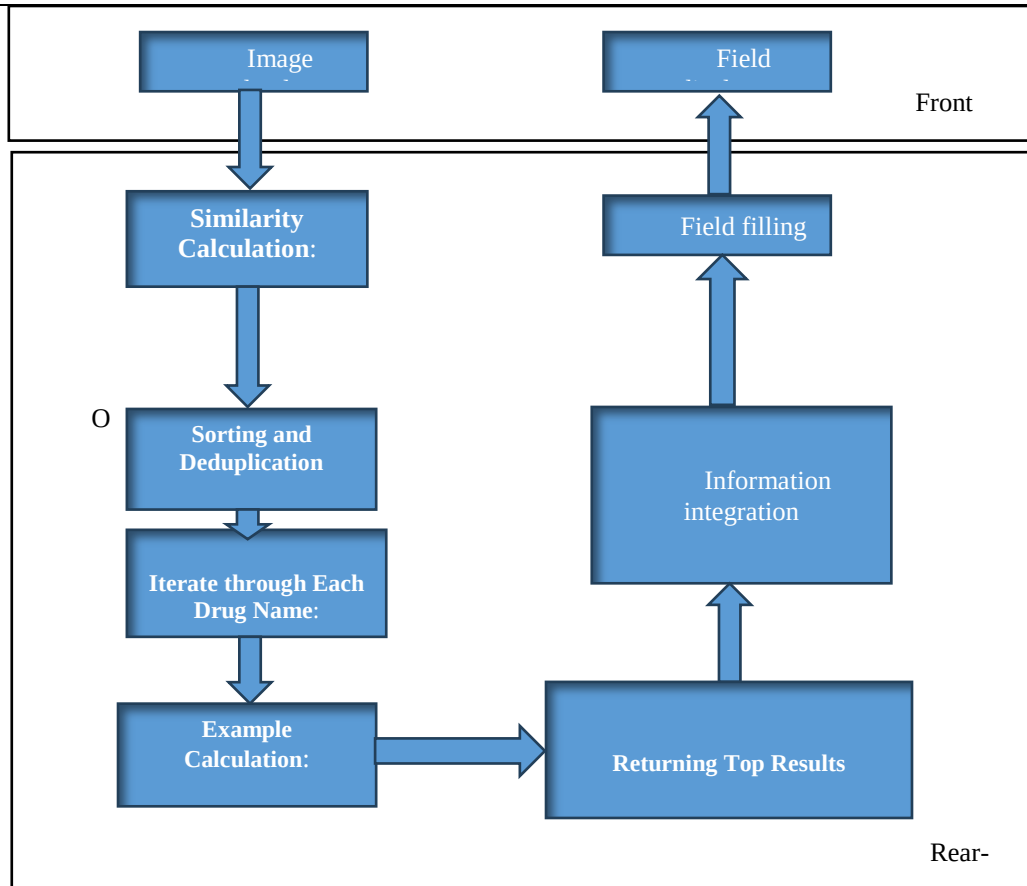


Figure 3. Functional structure diagram.

Let the input image I contain visual information about a drug label (name of the drug, active

1. Similarity Calculation:

For each substring found in the drug_text (e.g., the drug name "**Ibuprofen**", converted to lowercase), the similarity score is calculated as:

$$\text{Similarity Score} = \frac{\text{len}(\text{substring})}{\text{len}(\text{drugname})} \quad (1)$$

This formula gives a ratio of how much of the drug name matches the search text. A longer matching substring relative to the drug name's length will result in a higher similarity score, indicating a closer match.

2. Sorting and Deduplication:

- After calculating similarity scores, the code sorts the matches in descending order based on the similarity score, prioritizing matches with longer substrings within the drug name.

- The code then removes duplicate drug names to avoid listing the same drug multiple times with different substring matches.

Walkthrough

• Iterate through Each Drug Name:

- For the drug name "**Ibuprofen**" in the dataset, all possible substrings are generated and checked against the search text to see if they are present.

- If a substring is found, the similarity score is calculated for that substring using the formula:

$$\text{Similarity Score} = \frac{\text{len}(\text{substring})}{\text{len}(\text{drugname})} \quad (2)$$

• Example Calculation:

- Let's assume the drug name is "**Ibuprofen**" and the search text is "**pro**".

- Substrings of "**Ibuprofen**":

"I", "Ib", "Ibu", "Ibup", "pro", ..., "Ibuprofen"

- When the substring "**pro**" is found in the search text, its similarity score is calculated as:

$$\text{Similarity Score} = \frac{\text{len}(\text{"pro"})}{\text{len}(\text{Ibuprofen})} = \frac{3}{9} = 0.3333(3)$$

• Returning Top Results:

- The code sorts the results by similarity and returns the top 5 unique matches, ensuring the most relevant matches appear at the top.

Summary

The formula:

$$\text{Similarity Score} = \frac{\text{len}(\text{substring})}{\text{len}(\text{drugname})} \quad (4)$$

is used as a similarity measure, making it a proportional score of how much of the drug name matches the search text. In this case, the substring "**pro**" within "**Ibuprofen**" achieves a score of **0.3333**, indicating a partial match. This technique ensures the search effectively identifies relevant drug names using substring matching.

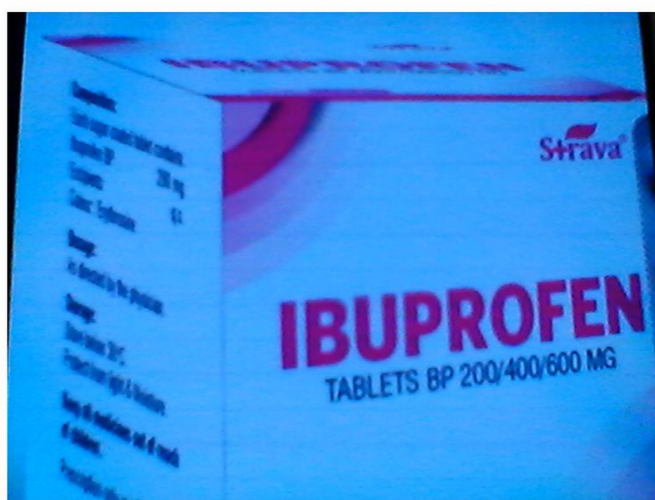
EasyOCR applies image processing and text extraction steps to identify text elements T from I . The extracted text T can be expressed as:

$$T = OCR(I) = CNN(I) \quad (5)$$

where $CNN(I)$ represents the convolutional neural network-based feature extraction within EasyOCR to identify individual characters and sequences in multiple languages and fonts. Let's assume the input image I contains the following details.

The main part of the web application is the section for scanning the medicine box. You need to click the "Open Camera" button and in the opened window, take

a picture of the medicine box by clicking the "Take a photo of the medicine box" button.



Take a photo of the medicine box

Image1. The scanning of the medicine box

The image I represents the drug label, which EasyOCR will process to extract the necessary text

It contains the following details:

- Name of the drug: Ibuprofen
- Active ingredient of the drug: Ibuprofen
- Form of the drug: Capsule

This process would result in structured data that can be stored or used in an application to display information to the user.

EasyOCR applies its OCR function to identify and extract text elements T from I . Here is how the text extraction process can be expressed:

$$T = OCR(I) = CNN(I) \quad (6)$$

where:

- T is the extracted text, containing the relevant information about the drug label.
- $CNN(I)$ is the convolutional neural network (CNN) model in EasyOCR, which identifies and processes individual characters and text sequences in the image.

Example output T for ibuprofen Label

Drug recognition

Open camera Remove Result

Detected text: Tva IBUPROFEN TABLETS BP 20841u

 Ibuprofen Active ingredient: Ibuprofen Dosage form: Capsule Manufacturer: Pfizer Price: 20000 Expiration Date: 31.12.2026 Quantity: 50 Dosage: 400mg Barcode: 2345678901234 Add to Database	 Ketoprofen Active ingredient: Ketoprofen Dosage form: Capsule Manufacturer: Sanofi Price: 39000 Expiration Date: 31.07.2026 Quantity: 1550 Dosage: 100mg Barcode: 1234567890152 Add to Database	 Bupropion Active ingredient: Bupropion Dosage form: Tablet Manufacturer: Mylan Price: 33000 Expiration Date: 28.02.2026 Quantity: 290 Dosage: 150mg Barcode: 1234567890026 Add to Database	 Baclofen Active ingredient: Baclofen Dosage form: Tablet Manufacturer: Merck Price: 25000 Expiration Date: 30.04.2026 Quantity: 980 Dosage: 10mg Barcode: 1234567890095 Add to Database
---	---	--	---

Image2. The scanning of the medicine box

Assuming EasyOCR processes the input image correctly, the output T could be:

$T = \{\text{ibuprofen, capsule, Pfizer, 20000, 31.12.2026, 50, 400mg, 2345678901234}\}$

This extracted text can then be further analyzed to categorize details:

Image Preprocessing for Enhanced OCR Accu-

To enhance the OCR accuracy for extracting text from the input image I containing drug information, we apply a series of image preprocessing steps.

The goal is to prepare the image for optimal OCR extraction of specific text elements T from I .

Given the image containing the following text elements:

T

$= \{\text{ibuprofen, capsule, Pfizer, 2000}\eta\eta., 31.122026, 50, 4\}$

1. **Grayscale Conversion:** Convert the image to grayscale to improve text contrast by removing color, making it easier to distinguish characters.

$$I_{gray} = \text{Grayscale}(I)$$

2. **Noise Reduction:** Apply filters to reduce background noise, such as shadows or textures, which can interfere with text recognition.

$$I_{clean} = \text{Denoise}(I_{gray})$$

3. **Image Binarization:** Convert the cleaned image to black and white, isolating text from the background by enhancing the contrast between text and non-text areas.

$$I_{binary} = \text{Threshold}(I_{clean})$$

4. **Resizing:** Scale the image to the optimal resolution for OCR processing, ensuring that text size meets OCR model requirements

$$I' = \text{Resize}(I_{binary}, d_w, d_h)$$

Where d_w and d_h are the target dimensions for best OCR results.

5. **Rotation Correction:** If the text is not horizontally aligned, apply a rotation transformation to align it horizontally.

$$I' = \text{Rotate}(I')$$

Segmentation of Text Areas: Isolate each section of text (e.g., "ibuprofen," "capsule," "Pfizer," "2000 AMD/ $\eta\eta.$," etc.) to help the OCR engine read each part individually and avoid misinterpretation of adjacent text areas. The information content in self-made images includes useful information and irrelevant information. The useful information is further divided into similar field information and dissimilar field information. Our testing dataset contains 203 images.

The evaluation criteria for the sim mentioned above are the final results obtained through multiple experiments. For the evaluation indicators of results, we use the following accuracy formula:

$$P = \frac{TP}{TP+FP} \quad (7)$$

For the accuracy of the final filling of an image, we divide the number of correctly filled fields (TP) by the number of fields it fills (TP + FP).

$$\text{Top-1 Accuracy} = \frac{\text{Number of correct Top-1 predictions}}{\text{Total number of predictions}} \quad (8)$$

The Top-1 Accuracy calculated on our testing dataset is 94.1%, while the Top-5 Accuracy is 98%. Top-1 accuracy is a metric commonly used to evaluate the performance of classification models. It measures the proportion of instances where the highest probability prediction made by the model (i.e., the top prediction) matches the actual label of the data point. In simpler terms, it's the percentage of times the model's most confident guess is correct.

Top-5 accuracy is another metric used to evaluate the performance of classification models, especially when there are a large number of classes. It measures the proportion of instances where the true label of a data point is among the top five most probable predictions made by the model.

This article focuses on the problem of data uploading and filling in web systems. Based on the existing advanced OCR technology and text similarity algorithms, combined and improved, the goal of filling fields in complex form images was effectively achieved. According to the test results, it was found that the accuracy of image recognition and filling in practical applications can reach over 90%. However, this method also has some limitations, as its proposal is based on practical engineering problems, so the field information considered is also related to this project. If fields need to be changed, they may need to be readjusted. In the future, further optimization can be based on this method, such as considering real-time updates

of database tables corresponding to fields, in order to better adapt to different image forms.

Discussion

Our study highlights how effective OCR and deep learning techniques, especially CNNs, can be in creating a reliable web application for retrieving drug information. The system depicted impressive accuracy, with a Top-1 score of 94.1% and a Top-5 score of 98%. These results underscore how advanced technology can tackle challenges in managing medications and facilitating access to pharmaceutical details. Our research builds on the work of other researchers and moves the field forward. For instance, Maitrichit and Hnoohom (2020) also achieved over 90% accuracy with CNNs and OCR; however, we have taken it a step further. By adding advanced text-matching algorithms, we have made data extraction and field mapping even more precise. Our system also stands out due to the use of Easy-OCR, fine-tuned specifically for pharmaceutical labels, which performs better than traditional methods like manual data entry or barcode scanning.

Comparing our work to the "Recognition Holic" study by Kayethri et al. (2022), it can be seen that their approach focused on accuracy in medicine detection. We have built on that by ensuring our system works well with a variety of packaging styles, even low-resolution images. This adaptability is thanks to preprocessing steps like noise reduction, grayscale conversion, and rotation correction, which align with Su,

Kang, and Fan's (2024) recommendations on optimizing OCR accuracy for complex tasks. Our web application brings real and unique value. It is simple: take a photo of a drug package, and the app retrieves detailed information in seconds. For users, this means less time searching and fewer chances of medication mistakes. It is a win for patients and healthcare professionals alike. Whether it is verifying prescriptions in a pharmacy or reconciling medications in a hospital, this system streamlines the process and makes it safer.

The app's design also ensures it is scalable and adaptable, which means it can be integrated into various healthcare environments. Besides, since the database can be updated, the system will remain current and reliable over time. Undoubtedly, no system is perfect. One of the main challenges is the need for high-quality images. Blurry or poorly lit photos might not work as well, which could limit the app's usability in some situations. Another hurdle is processing multilingual or heavily stylized fonts, which can be tricky for OCR. And then there is the need to keep the database updated—a critical task to ensure users always receive accurate information. Eventually, the computational demands of the deep learning models may be a barrier for users with less powerful devices.

Enhancing the OCR model to handle multiple languages and improving its ability to process more challenging images are top priorities. Expanding the training dataset to include more packaging styles and resolutions will make the system even more versatile. We are also thrilled about the potential of incorporating Natural Language Processing (NLP) tools. Imagine the app not only identifying medications but also flagging potential drug interactions or contraindications. Collaborations with pharmaceutical companies could help us build an even more comprehensive and reliable database. And by optimizing the app for mobile use, we can make it more accessible to people everywhere.

This work is part of a larger movement to bring digital innovation to healthcare. By using OCR and deep learning, we are making medication management safer, faster, and easier. The app streamlines workflows, reduces errors, and, most importantly, helps protect patients. By continuing to refine and expand this technology, we can make a lasting impact on global health, improving lives and transforming how we manage medications.

Conclusion

In summary, the development of a web application for obtaining drug information represents a significant advancement in pharmaceutical technology with the potential to improve medication management. In this research, various innovative solution frameworks were explored, examining their advantages and challenges. The primary objective of this web application is to provide comprehensive yet accurate and reliable information to every user. Following this, it is equally important to ensure easy access to the process of obtaining information, enabling users to make informed decisions regarding their health.

To organize the smooth operation of the web application, the attachment of the necessary databases plays an important role in the application design. Almost every operation performed in the application is

recorded in tables in the database. Examples include registering, logging in, and adding medications to the web application. An administrator is registered in the application, who can add medicines to the home page of the web application by attaching certain information to them. The information contains comprehensive characteristics of the drug. They are: the name of the drug, the active ingredient contained in the drug, the pharmaceutical company that produces the drug, the average price of the drug in the market, its expiry date, dosage, the barcode of the drug, and the most important point is the description of the drug, what kind of drug it is, for what purpose it is used, with a prescription or without a prescription. is achieved. The information also indicates what diseases it is used for.

After adding this drug in the case of an admin account, he can attach the drug to the base linked to the web application. To have a regular user account, individuals must first register in the web application. Following the sequence of filling in the necessary data, the user can register and then log in. The user can see some medicines on the site. However, if he cannot find the information about the drug from the drug list presented in the web application, he can find the information by using the drug recognition functionality. For this, he merely needs to take a picture of the drug box in the appropriate section, namely "Medicine recognition", in addition to pressing the camera, the "Open camera" button, and then wait for the result. By scanning the medicine box, the system will find the necessary information from the database and display it on the user's device. The user can add the drug he is looking for to the basket in order not to "lose" the drug. If the drug is no longer needed, the user can delete the drug from his cart.

All the mentioned actions, including adding the drug to the cart and later deleting it from there, are recorded in the database in the corresponding table. Those tables, in turn, have key values and are connected by appropriate links. It ensures the smooth operation of the web application.

One of the main goals of this research is to contribute to having an informed society. By providing users with the knowledge, they need to understand and correctly take the medication they are about to use, the web application reduces the risk of medication errors, adverse reactions, and non-adherence. All of this can lead to improved health outcomes, reduced healthcare costs, and a higher quality of life for users. Moreover, the integration of additional features in terms of information, such as touching on the topics most important to users in a web application, is primarily in the interests of the users themselves. By including reliable information and details about sensitive topics, the web application helps increase user awareness and enrich the knowledge base.

While this app offers several benefits to users, it poses just as many challenges from a developer's perspective. One of the primary challenges is the need to ensure the accuracy and reliability of drug information to users. Accepting this challenge with dignity and looking for ways to solve it, reliable sources were the best solution. Examples of reliable sources are pharmaceutical and pharmaceutical companies. After researching and identifying the necessary information, efforts

were made to include it in the drug information retrieval web application. In addition to its innovative functionality, the web application provides ample information.

Considering the needs, wishes, and interests of society, different sections of the website contain equally valuable information. Surveys conducted to understand the public's expectations of the web application were instrumental in the development and design of the web application. Targeting the wants and needs of each survey participant, the web application included several sections containing very different, yet important information. Seeing the need to present such a format, it was considered important to include information from the simplest information to the most sensitive topics.

The public always needs a reliable source of information. Even though, they always tend to take the easy way out and choose shortcuts to achieve their goals. That is why the web application has been designed to be simple enough and easy to use for the user, and at the same time

Comparing the traditional methods of obtaining information about medicines and the advanced solution offered by the web application, we can say that the public, striving to make their daily life even easier, will appreciate this innovative solution and include it in their daily life.

As a result of this research, the set problems were solved and the corresponding expected results were obtained. Therefore, the development of a web application for obtaining information about a drug by scanning its box holds great promise in the process of changing medication management. A web application can have a significant impact on the delivery of health care and improve the well-being of individuals worldwide. Due to further updates of the web application, it is planned to include several more features to diversify and further improve the application.

References:

1. DataCamp. (n.d.). SQLAlchemy tutorial: Using SQLAlchemy with Python. Retrieved November 8, 2024, from <https://www.datacamp.com/tutorial/sqlalchemy-tutorial-examples>

2. Fetchify. (n.d.). What is email validation, and why is it important? Fetchify. Retrieved November 8, 2024, from <https://www.fetchify.com/what-is-email-validation-and-why-is-it-important>

3. Kayethri, et al (2022). Recognition Holic: Medicine Detection Using Deep Learning Techniques Computer Science & Engineering: An International Journal (CSEIJ), vol 12, N06, December 2022, DOI:10.121/cseij.2022.12614

4. Kim, W.-T., Shin, J.-G., Yoo, I.-S., Lee, J.-W., Jeon, H.-J., Yoo, H.-S., Kim, Y., Jo, J.-M., Hwang, S.-J., Lee, W.-J., Park, S., & Kim, Y.-J. (2024). MEDIC: GPT-Powered Chatbot for Drug-Drug Interaction, Mayo Clinic Proceedings: Digital Health, 2024, ISSN 2949-7612, <https://doi.org/10.1016/j.mcpdig.2024.09.001>

5. Maitrichit, N., & Hnoohom, N. (2020). Intelligent Medicine Identification System Using a Combination of Image Recognition and Optical Character Recognition. 2020 15th International Joint Symposium on Artificial Intelligence and Natural Language Processing (ISAI-NLP), Bangkok, Thailand, 2020, pp. 1-5, doi: 10.1109/ISAI-NLP51646.2020.9376816

6. SQLite. (n.d.). About SQLite. Retrieved November 8, 2024, from <https://www.sqlite.org/about.html>

7. Stoumpos, et al (2023). Digital Transformation in Healthcare: Technology Acceptance and Its Applications. International Journal of Environmental Research and Public Health, 20(4), 3407; <https://doi.org/10.3390/ijerph20043407>

8. Su, H. Kang, et al (2024). Research on a Web System Data-Filling Method Based on Optical Character Recognition and Multi-Text Similarity. Applied Sciences, 14(3), 1034. <https://doi.org/10.3390/app14031034>

9. Thiede, M. (2005). Information and access to health care: is there a role for trust?. Social Science & Medicine, 61(7), 1452-1462. <https://doi.org/10.1016/j.socscimed.2004.11.076>