

Communication aid for specially abled

Sayandeep Patra
Department of ECE
BMS College of Engineering
Bangalore-560019, India.
sayandeep.ec19@bmsce.ac.in

Shravan Achar
Department of ECE
BMS College of Engineering
Bangalore-560019, India.
shravan.ec19@bmsce.ac.in

Supriyo Bid
Department of ECE
BMS College of Engineering
Bangalore-560019, India.
supriyo.ec19@bmsce.ac.in

Sathvik G
Department of ECE
BMS College of Engineering
Bangalore-560019, India.
sathvikg.ec19@bmsce.ac.in

Dr. A Meera
Department of ECE
BMS College of Engineering Bangalore-560019, India.
amira.ece@bmsce.ac.in

Abstract - This project aims to develop a comprehensive and inclusive technology solution for individuals with sensory impairments. It focuses on hand gesture recognition to voice conversion, text identification in images to voice conversion, text to speech conversion, and voice to speech conversion. The implementation involves the use of a Raspberry Pi, a microphone, a speaker, and a camera to enable seamless communication and interaction. The hand gesture recognition feature utilizes machine learning and computer vision techniques to recognize and interpret hand gestures, converting them into corresponding voice commands. The text identification in images feature utilizes image processing algorithms to identify and extract text from images. This text is then converted into voice output. The text to speech conversion feature enables individuals to input text through a keyboard or other input devices. The system converts the entered text into speech output. The voice to speech conversion feature utilizes a microphone to record voice modulations, which are then converted into text using speech recognition algorithms. The converted text is displayed on a screen.

Keywords – Communication Aid, Raspberry Pi, AI/ML

I. INTRODUCTION

One of the most precious gifts to a human being is an ability to see, listen, speak and respond according to the situations. But there are some unfortunate ones who are deprived of this. Making a single compact device for people with Visual, Hearing and Vocal impairment is a tough job. Communication between deaf-dumb and normal person have been always a challenging task. This paper proposes an innovative communication system framework for deaf, dumb and blind people in a single compact device. We provide a technique for a blind person to read a text and it can be achieved by capturing an image through a camera which converts a text to speech (TTS). It provides a way for the deaf people to read a text by speech to text (STT) conversion technology. Also, it provides a technique for dumb people using text to voice conversion. The system is provided with four switches and each switch has a different function. The blind people can be able to read the words using by Tesseract OCR (Online Character Recognition), the dumb people can communicate their message through text which will be read out by espeak, the deaf people can be able to hear others speech from text. of Laptop. This project explores and evaluates the effectiveness of communication aid devices designed specifically to enhance communication for individuals who are deaf, dumb, and blind. Communication is an essential aspect of human interaction, enabling individuals to express thoughts, emotions, and needs, thereby fostering social connections and inclusivity. However, those with combined sensory impairments face significant challenges in effective

communication.

II. METHODOLOGY

To enhance the functionality of the proposed device, we have identified four primary options that allow users to interact with it effectively: text-to-voice, image-to-speech, capture and read image, and speech-to-text. Let's explore each option in more detail:

- **Text-to-Voice:** This feature enables users to input text messages through the keyboard and have them converted into spoken words. The device uses text-to-speech (TTS) conversion algorithms to generate natural-sounding audio output. This functionality empowers vocally impaired individuals to communicate by typing their messages, which are then vocalized by the device's speaker.
- **Image-to-Speech:** With this option, users can capture an image using the integrated camera. The device processes the image and employs optical character recognition (OCR) technology to recognize text within the image. The recognized text is then converted into spoken words using TTS algorithms and played through the device's speaker. This capability allows visually impaired individuals to access written information present in the captured images.
- **Capture and Read Image:** This feature combines the device's camera and TTS capabilities to provide users with the ability to capture an image and have the text within the image read aloud. Once an image is captured, the device uses OCR to extract the text, which is then converted to speech and played back to the user. This functionality assists visually impaired individuals in accessing printed or written content from various sources.
- **Speech-to-Text:** This option allows users to convey their messages verbally using the device's microphone. The speech recognition algorithms within the device process the spoken input and convert it into text format. The recognized text is displayed on the device screen, enabling deaf individuals to understand what others are saying by reading the transcribed speech. This capability promotes effective communication between individuals who are deaf and those who can hear.

By offering these four options, our device caters to the unique needs of individuals who are visually impaired, vocally impaired, and deaf. It promotes inclusive communication, access to information, and independence, allowing users to interact with the world more effectively and experience a standard lifestyle similar to that of individuals without such impairments.

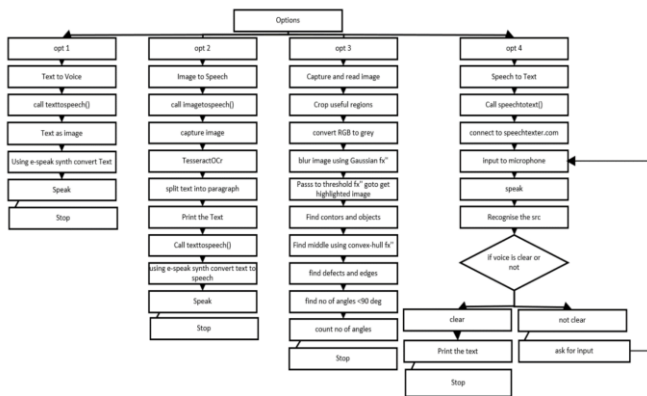


Fig 1. Data flow diagram

- User Input: The function prompts the user to enter text and stores the entered text in the variable "text1".
- Display Entered Text: The entered text is printed to provide a visual confirmation.
- Text-to-Speech Conversion: The function uses the gTTS library to convert the entered text into speech. It instantiates the gTTS class with the entered text, sets the language to English, and the speech speed to normal. The resulting speech is saved as an audio file named "voice.mp3".
- Playback of the Audio: The function loads the saved audio file using the pygame mixer module. It initializes the pygame mixer, loads the audio file for playback, and initiates the playback of the audio.
- Wait for Audio Playback: The function pauses the program execution for the duration of the audio file using the time.sleep() function.
- Cleanup: The function gracefully shuts down the pygame mixer.

To use the text_to_voice() function, you need to have the following dependencies installed: pygame, gTTS, and mutagen. pygame is a library for multimedia applications, gTTS is a library for text-to-speech conversion, and mutagen is a library for working with audio file metadata.

The text_to_voice() function provides a convenient way to convert text to speech and play the resulting audio, making it suitable for applications that require speech synthesis or interactive voice experiences.

IV. IMAGE TO SPEECH

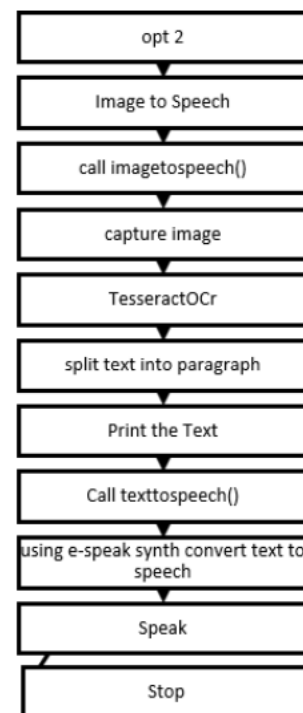


Fig 4. Flow Chart

III. TEXT TO VOICE

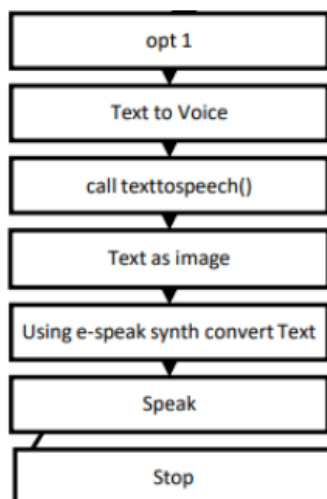


Fig 3. Flow Chart

The text_to_voice() function is a Python function that converts user-inputted text into speech and plays the resulting audio. It utilizes the pygame library for audio playback and the gTTS (Google Text-to-Speech) library for text-to-speech conversion. Here is a summary of the steps performed by the text_to_voice() function:

The provided description outlines the development of an Optical Character Recognition (OCR) image recognition system using machine learning techniques. The system is designed to be deployed on a Raspberry Pi and is capable of capturing images from a connected camera. The captured images are then processed and analyzed using a trained ML/DL model to extract text information.

The methodology involves the following steps:

V. GESTURE TO VOICE

- **Data Collection:** A diverse dataset of labeled characters is collected for training the ML/DL model. Synthetic data generation techniques and publicly available OCR datasets are used to create a comprehensive training dataset.
- **Preprocessing:** Captured images undergo preprocessing steps such as resizing, grayscale conversion, noise reduction, and normalization to enhance their quality and prepare them for OCR analysis.
- **ML/DL Model Selection:** A Convolutional Neural Network (CNN) model is chosen for the OCR task due to its effectiveness in image recognition. The CNN architecture includes convolutional layers, pooling layers, fully connected layers, and an output layer for character classification.
- **Training the Model:** The model is trained using the preprocessed image data. Stochastic Gradient Descent (SGD) is used as the optimization algorithm, and performance metrics like accuracy, precision, recall, and F1-score are used to evaluate the model.
- **Integration with Raspberry Pi:** The trained OCR model is deployed on a Raspberry Pi along with the necessary software components. The Raspberry Pi camera module is used to capture real-time images, which are processed by the ML model to extract text. The extracted text can be displayed or stored for further analysis.

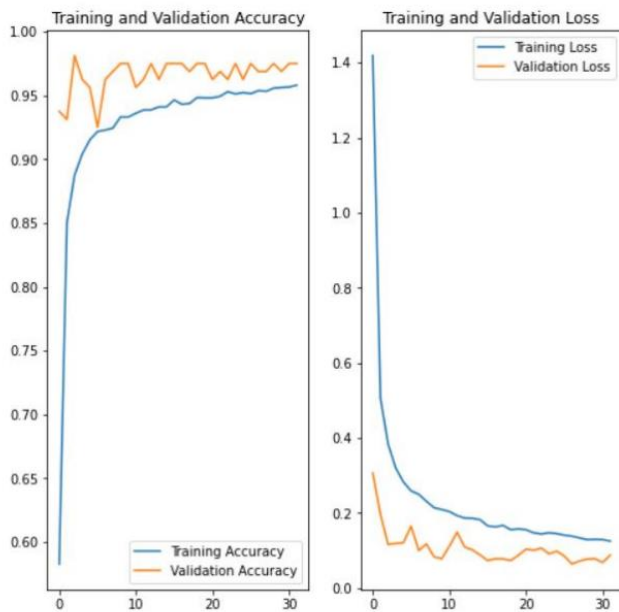


Fig 5. Subplots of the training and validation model

This demonstrates various aspects of the implementation, including data preprocessing, model architecture definition, model training, evaluation, and usage of libraries like TensorFlow, OpenCV, and gTTS for text-to-speech conversion.

In summary, the OCR system utilizes machine learning techniques and a trained model to recognize and extract text from images captured by a Raspberry Pi camera. The system can be extended for real-world applications such as document scanning, text recognition in images, or assistive technologies for the visually impaired.

It follows the following steps:

- **Imports necessary libraries:** The code imports libraries such as imutils, OpenCV, numpy, mediapipe, tensorflow.keras.models, pygame, time, gtts, and mutagen.mp3.
- **Sets up hand detection:** The code initializes hand detection functionality using the mpHands module from Mediapipe. It defines parameters like the maximum number of hands to detect and the minimum detection confidence threshold.
- **Loads a pre-trained gesture recognizer model:** The code loads a pre-trained model for gesture recognition using the load_model function from TensorFlow's Keras API. The model is loaded from a file.
- **Loads class names:** The code reads class names from a file.
- **Initializes the webcam:** The code initializes the webcam for capturing live video frames.
- **Enters a video frame processing loop:** The code enters a continuous loop to process each video frame captured by the webcam.
- **Preprocesses the frame:** Inside the loop, the code reads each frame, flips it horizontally, and performs other necessary preprocessing steps.

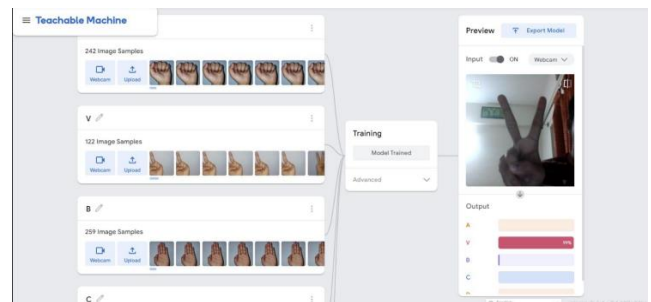


Fig 6. Image Dataset Training

- **Predicts hand landmarks:** The code uses the hand detection module to predict hand landmarks in the frame.
- **Extracts and visualizes landmarks:** If hand landmarks are detected, the code extracts the coordinates of the landmarks and visualizes them on the frame.
- **Predicts the gesture and generates voice output:** The code predicts the gesture by passing the landmarks to the pre-trained model. It converts the predicted gesture into speech using the gTTS library and saves it as an audio file. The speech is played using the pygame library.
- **Shows the prediction on the frame:** The code overlays the predicted gesture on the frame using OpenCV's putText function.
- **Displays the frame:** The code displays the frame with the predicted gesture and voice output.
- **Releases resources:** After exiting the loop, the code releases the webcam and destroys active windows.

The implemented system continuously captures frames from the webcam, detects hand landmarks, predicts the gesture, converts the gesture to speech, and displays the predicted gesture and voice output in real-time.

The code utilizes various libraries and modules to handle video processing, hand detection, deep learning model loading, speech generation, and audio playback. These libraries and modules provide the necessary functionalities for implementing the real-time gesture recognition system.

VI. VOICE TO TEXT

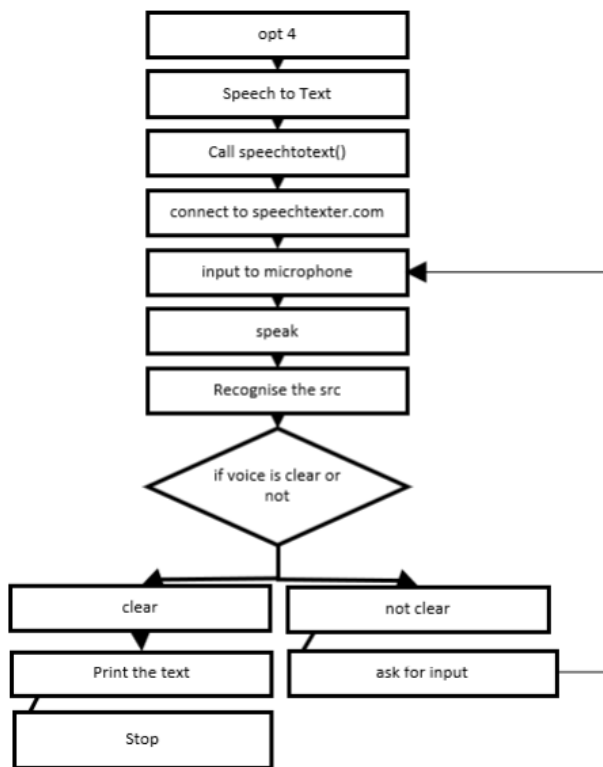


Fig 7. Flow Chart

This project uses SpeechTexter.com, a web-based platform that offers speech recognition services to convert spoken language into written text. It highlights key features such as advanced speech recognition, real-time transcription, a user-friendly interface, language support, compatibility across devices, and application integration options.

Additionally, we provide a breakdown of how to convert voice to text using a Raspberry Pi and a connected serial device. The process involves importing necessary libraries (such as time, serial, and RPi.GPIO), configuring the serial port settings, and continuously listening for voice input by reading data from the serial port. When new voice data is received, it is decoded and printed as recognized text.

In summary, SpeechTexter.com provides an accessible online tool for converting speech to text, while the Raspberry Pi instructions outline the process of utilizing a connected serial device to convert voice data into text on the Raspberry Pi.

VII. RESULTS AND DISCUSSIONS

This project focused on improving accessibility and communication for individuals with sensory impairments, particularly those who are visually impaired, vocally impaired, and deaf. It incorporated four key features: hand gesture recognition to voice conversion, text identification in images to voice conversion, text to speech conversion, and voice to speech conversion. The system utilized a Raspberry Pi, microphone, speaker, and camera to enable these functionalities.

The hand gesture recognition feature used the MediaPipe library and a pre-trained deep learning model to detect and recognize hand gestures in real-time. This allowed users to interact with the system using intuitive hand gestures, enhancing communication and control.

The text identification feature utilized Optical Character Recognition

(OCR) with the Tesseract library. By capturing images with the camera, the system extracted text from the images and converted it into voice output. This enabled users to access textual information in their surroundings, promoting independence and information accessibility.

The text to speech conversion feature used the GTTS library to synthesize speech from text input. Users could input text through a keyboard or other input methods, and the system generated human-like speech output in real-time. This facilitated communication and interaction using text-based input.

The voice to speech conversion feature involved capturing voice input through a microphone and converting it into text using speech recognition techniques. The recognized text was then converted into voice output using text-to-speech functionality. This allowed users to express their thoughts and commands through voice input, facilitating communication.

By integrating these features into a portable Raspberry Pi-based system, the project provided a comprehensive solution for individuals with sensory impairments, enhancing their communication, information access, and interaction capabilities. The project demonstrated the potential of technology in improving the quality of life for individuals with sensory challenges.

Future improvements could include refining the gesture recognition model, optimizing text identification algorithms, improving speech synthesis quality, and exploring additional accessibility features based on user feedback. The project has paved the way for further research and innovation in assistive technologies, promoting inclusivity and accessibility in society.

VIII. CONCLUSION AND FUTURE WORK

As technology continues to advance, there are several potential future trends and advancements that can further enhance the capabilities and impact of this project:

1. **Advanced Gesture Recognition:** The field of gesture recognition is continuously evolving, with advancements in machine learning and computer vision. Future trends may include more robust and accurate hand gesture recognition models that can recognize complex gestures and movements with higher precision. This could enable a wider range of gestures to be translated into voice commands, enhancing the system's usability and versatility.
2. **Improved Image Processing Techniques:** Image processing algorithms for text identification and extraction are constantly being refined. Future trends may include the integration of advanced image processing techniques, such as image enhancement, denoising, and adaptive thresholding, to improve the accuracy and reliability of text identification in images. This could result in better performance in reading text from various sources and challenging environments.
3. **Natural Language Understanding:** Enhancing the system's ability to understand and interpret natural language can significantly improve its usability. Future trends may focus on incorporating natural language understanding techniques and technologies, such as natural language processing (NLP) and machine learning, to enable more interactive and context-aware communication. This could enable users to engage in more dynamic and meaningful conversations with the system.
4. **Multimodal Interaction:** Combining multiple modes of interaction, such as gestures, speech, and touch, can enhance the user experience and accessibility. Future trends may involve integrating additional input modalities, such as touchscreens or haptic interfaces, to complement the existing gesture and voice-based interactions. This would provide users with more options

and flexibility in how they interact with the system.

5. Cloud-Based Processing and Services: The project's current implementation relies on the processing capabilities of the Raspberry Pi. However, future trends may involve leveraging cloud computing resources to offload computationally intensive tasks, such as image processing or speech recognition. This could enable more complex algorithms and models to be utilized, expanding the system's capabilities without being limited by the device's hardware constraints.

6. User Customization and Personalization: Recognizing the diverse needs and preferences of individuals with sensory impairments, future trends may involve incorporating customization and personalization features. This could include user profiles, adaptive settings, and personalized voice output characteristics. Allowing users to tailor the system to their specific requirements would enhance the overall user experience and promote inclusivity.

7. Integration with Smart Home and IoT Devices: As smart home technology and Internet of Things (IoT) devices become more prevalent, integrating the project's functionalities with these systems opens up new possibilities. Future trends may involve seamless integration with smart home devices, allowing users to control various aspects of their Communication aid for the specially abled.

REFERENCES

- [1] L. Anusha and Y. U. Devi, "Implementation of gesture based voice and language translator for dumb people," 2016 International Conference on Communication and Electronics Systems (ICES), Coimbatore, India, 2016, pp. 1-4, doi: 10.1109/CESYS.2016.7889990.
- [2] S. S. Shinde, R. M. Autee and V. K. Bhosale, "Real time two way communication approach for hearing impaired and dumb person based on image processing," 2016 IEEE International Conference on Computational Intelligence and Computing Research (ICCIC), Chennai, India, 2016, pp. 1-5, doi: 10.1109/ICCIC.2016.7919572.
- [3] G. I. Hapsari, G. A. Mutiara and D. T. Kusumah, "Smart cane location guide for blind using GPS," 2017 5th International Conference on Information and Communication Technology (ICoICT), Melaka, Malaysia, 2017, pp. 1-6, doi: 10.1109/ICoICT.2017.8074697.
- [4] S. He, "Research of a Sign Language Translation System Based on Deep Learning," 2019 International Conference on Artificial Intelligence and Advanced Manufacturing (AIAM), Dublin, Ireland, 2019, pp. 392-396, doi: 10.1109/AIAM48774.2019.00083
- [5] Vishal, D., Aishwarya, H.M., Nishkala, K., Royan, B.T. and Ramesh, T.K., 2017. Sign Language to Speech Conversion: Using Smart Band and Wearable Computer Technology. In IEEE International Conference on Computational Intelligence and Computing Research (ICCIC).
- [6] K. Tiku, J. Maloo, A. Ramesh and I. R., "Real-time Conversion of Sign Language to Text and Speech," 2020 Second International Conference on Inventive Research in Computing Applications (ICIRCA), Coimbatore, India, 2020, pp. 346-351, doi: 10.1109/ICIRCA48905.2020.9182877.
- [7] P. Vijayalakshmi and M. Aarthi, "Sign language to speech conversion," 2016 International Conference on Recent Trends in Information Technology (ICRTIT), Chennai, India, 2016, pp. 1-6, doi: 10.1109/ICRTIT.2016.7569545.
- [8] Ce Zhang, Xin Pan, Huapeng Li, Andy Gardiner, Isabel Sargent, Jonathon Hare, Peter M. Atkinson, A hybrid MLP-CNN classifier for very fine resolution remotely sensed image classification, ISPRS Journal of Photogrammetry and Remote Sensing, Volume 140, 2018, Pages 133-144, ISSN 0924-2716
- [9] L. Lu, Y. Yi, F. Huang, K. Wang and Q. Wang, "Integrating Local CNN and Global CNN for Script Identification in Natural Scene Images," in IEEE Access, vol. 7, pp. 52669-52679, 2019, doi: 10.1109/ACCESS.2019.2911964.
- [10] S. A. Sabab and M. H. Ashmafee, "Blind Reader: An intelligent assistant for blind," 2016 19th International Conference on Computer and Information Technology (ICCIT), Dhaka, Bangladesh, 2016, pp. 229-234, doi: 10.1109/ICCITECHN.2016.7860200.
- [11] B. Rajapandian, V. Harini, D. Raksha and V. Sangeetha, "A novel approach as an AID for blind, deaf and dumb people," 2017 Third International Conference on Sensing, Signal Processing and Security (ICSSS), Chennai, India, 2017, pp. 403-408, doi: 10.1109/SSPS.2017.8071628.
- [12] S. S. Shinde, R. M. Autee and V. K. Bhosale, "Real time two way communication approach for hearing impaired and dumb person based on image processing," 2016 IEEE International Conference on Computational Intelligence and Computing Research (ICCIC), Chennai, India, 2016, pp. 1-5, doi: 10.1109/ICCIC.2016.7919572.
- [13] G. I. Hapsari, G. A. Mutiara and D. T. Kusumah, "Smart cane location guide for blind using GPS," 2017 5th International Conference on Information and Communication Technology (ICoICT), Melaka, Malaysia, 2017, pp. 1-6, doi: 10.1109/ICoICT.2017.8074697.
- [14] S. Chatteraj, K. Vishwakarma and T. Paul, "Assistive system for physically disabled people using gesture recognition," 2017 IEEE 2nd International Conference on Signal and Image Processing (ICSIP), Singapore, 2017, pp. 60-65, doi: 10.1109/SIPROCESS.2017.8124506.
- [15] S. Mirri, C. Prandi, P. Salomoni and L. Monti, "Fitting like a GlovePi: A wearable device for deaf-blind people," 2017 14th IEEE Annual Consumer Communications & Networking Conference (CCNC), Las Vegas, NV, USA, 2017, pp. 1057-1062, doi: 10.1109/CCNC.2017.7983285.
- [16] K. Amrutha and P. Prabu, "ML Based Sign Language Recognition System," 2021 International Conference on Innovative Trends in Information Technology (ICITIIT), Kottayam, India, 2021, pp. 1-6, doi: 10.1109/ICITIIT51526.2021.9399594
- [17] K. Nimisha and A. Jacob, "A Brief Review of the Recent Trends in Sign Language Recognition," 2020 International Conference on Communication and Signal Processing (ICCSP), Chennai, India, 2020, pp. 186-190, doi: 10.1109/ICCSP48568.2020.9182351.
- [18] P. R. Sanmitra, V. V. S. Sowmya, and K. Lalithanjana, "Machine Learning Based Real Time Sign Language Detection", IJRESM, vol. 4, no. 6, pp. 137-141, Jun. 2021.
- [19] M. Jaiswal, V. Sharma, A. Sharma, S. Saini and R. Tomar, "An Efficient Binarized Neural Network for Recognizing Two Hands Indian Sign Language Gestures in Real-time Environment," 2020 IEEE 17th India Council International Conference (INDICON), New Delhi, India, 2020, pp. 1-6, doi: 10.1109/INDICON49873.2020.9342454.
- [20] P.K. Athira, C.J. Sruthi, A. Lijiya, A Signer Independent Sign Language Recognition with Co-articulation Elimination from Live Videos: An Indian Scenario, Journal of King Saud University - Computer and Information Sciences, Volume 34, Issue 3, 2022, Pages 771-781, ISSN 1319-1578, https://doi.org/10.1016/j.jksuci.2019.05.002.
- [21] Á. Z. Kaló and M. L. Sipos, "Key-Value Pair Searching System via Tesseract OCR and Post Processing," 2021 IEEE 19th World Symposium on Applied Machine Intelligence and Informatics (SAMI), Herl'any, Slovakia, 2021, pp. 000461-000464, doi: 10.1109/SAMI50585.2021.9378680.

- [22] Chawla, Muskan and Jain, Rachna and Nagrath, Preeti, Implementation of Tesseract Algorithm to Extract Text from Different Images (May 1, 2020). Proceedings of the International Conference on Innovative Computing & Communications (ICICC) 2020, SSRN: <https://ssrn.com/abstract=3589972> or <http://dx.doi.org/10.2139/ssrn.3589972>.
- [23] Pawar, N., Shaikh, Z., Shinde, P. and Warke, Y.P., 2019. Image to text conversion using tesseract. Image, 6(02).
- [24] Clausner, C., Antonacopoulos, A. and Pletschacher, S., 2020. Efficient and effective OCR engine training. International Journal on Document Analysis and Recognition (IJDAR), 23, pp.73-88.
- [25] R. Mittal and A. Garg, "Text extraction using OCR: A Systematic Review," 2020 Second International Conference on Inventive Research in Computing Applications (ICIRCA), Coimbatore, India, 2020, pp. 357-362, doi: 10.1109/ICIRCA48905.2020.9183326.
- [26] Y. -M. Su, H. -W. Peng, K. -W. Huang and C. -S. Yang, "Image processing technology for text recognition," 2019 International Conference on Technologies and Applications of Artificial Intelligence (TAAI), Kaohsiung, Taiwan, 2019, pp. 1-5, doi: 10.1109/TAAI48200.2019.8959877.
- [27] Akinbade, D., Ogunde, A.O., Odim, M.O. and Oguntunde, B.O., 2020. An adaptive thresholding algorithm-based optical character recognition system for information extraction in complex images. Journal of Computer Science, 16(6), pp.784-801.
- [28] J. Kunjumon and R. K. Megalingam, "Hand Gesture Recognition System For Translating Indian Sign Language Into Text And Speech," 2019 International Conference on Smart Systems and Inventive Technology (ICSSIT), Tirunelveli, India, 2019, pp. 14-18, doi: 10.1109/ICSSIT46314.2019.8987762.
- [29] H. -T. Luong and J. Yamagishi, "Bootstrapping Non-Parallel Voice Conversion from Speaker-Adaptive Text-to-Speech," 2019 IEEE Automatic Speech Recognition and Understanding Workshop (ASRU), Singapore, 2019, pp. 200-207, doi: 10.1109/ASRU46091.2019.9004008.
- [30] T. D. Chung, M. Drieberg, M. F. Bin Hassan and A. Khalyasmaa, "End-to-end Conversion Speed Analysis of an FPT.AI-based Text-to-Speech Application," 2020 IEEE 2nd Global Conference on Life Sciences and Technologies (LifeTech), Kyoto, Japan, 2020, pp. 136-139, doi: 10.1109/LifeTech48969.2020.1570620448.
- [31] B. Sudharsan, S. P. Kumar and R. Dhakshinamurthy, "AI Vision: Smart speaker design and implementation with object detection custom skill and advanced voice interaction capability," 2019 11th International Conference on Advanced Computing (ICoAC), Chennai, India, 2019, pp. 97-102, doi: 10.1109/ICoAC48765.2019.247125.
- [32] Nandalal, V. A. Kumar, T. Manikandan, K. Sindhuprika, R. Sujitha and G. Suvathy, "Raspberry Pi based Smart Library for Visually Challenged Citizens," 2020 Fourth International Conference on I-SMAC (IoT in Social, Mobile, Analytics and Cloud) (I-SMAC), Palladam, India, 2020, pp. 449-453, doi: 10.1109/I-SMAC49090.2020.9243467.
- [33] F. Barkani, H. Satori, M. Hamidi, O. Zealouk and N. Laaidi, "Amazigh Speech Recognition Embedded System," 2020 1st International Conference on Innovative Research in Applied Science, Engineering and Technology (IRASET), Meknes, Morocco, 2020, pp. 1-5, doi: 10.1109/IRASET48871.2020.9092014.
- [34] A. Ravi, S. Khasimbee, T. Asha, T. N. S. Joshna and P. G. Jyothirmai, "Raspberry pi based Smart Reader for Blind People," 2020 International Conference on Electronics and Sustainable Communication Systems (ICESC), Coimbatore, India, 2020, pp. 445-450, doi: 10.1109/ICESC48915.2020.9155941.
- [35] S. Banthia and V. Priya, "Improving reality perception for the visually impaired," 2020 International Conference on Emerging Trends in Information Technology and Engineering (icETITE), Vellore, India, 2020, pp. 1-7, doi: 10.1109/ic-ETITE47903.2020.093