

Vision Seeker: Deep Learning based Missing Person Detection and Recognition in the CCTV and Prerecorded Video

Amruta Tapas Paul¹, Dnyaneshwar Itankar², Ayush Sontakke³, Aaditi Bedarkar⁴, Ayush Patil⁵ and Abhishek Paitode⁶

¹⁻⁶CSE Department, Yeshwantrao Chavan College of Engineering, Nagpur, India

Email: amrutatapas@gmail.com, itankardnyaneshwar@gmail.com, ayushsontakke2004@gmail.com, aaditibedarkar@gmail.com, apayush28@gmail.com, abhishekpaitode0007@gmail.com, lipikshab@gmail.com

Abstract— The increasing number of missing persons cases which is a serious challenge faced by all of the world. Despite of having CCTV in numerous places, there is a lack of technological methods to find them through it in an efficient way. With the ever-increasing population in metropolitan cities, finding a person in a crowded area is also a tedious task to do. This paper proposes Vision seeker- an AI which will help the law and enforcement to find the missing person in a short time. The system integrates three pretrained models i.e. GFP- GAN -which will help to increase the quality of the image, secondly MTCNN for detecting the image even in conditions like poor lighting, crowd areas, etc. from the video by checking it through all frames, and lastly face net with the InceptionResnetV1 architecture, implemented using the facenet_pytorch library will recognize the facial structure by checking the facial embeddings and will compare it with the database of that individual. This model will help to integrate this with the real time access of the CCTV and in return will give us the image and also the timestamps on which the missing person is found. It will also show the location in which that individual has been found. The primary contribution lies in the novel system-level integration and practical deployment pipeline, rather than algorithmic innovation. By implementing this solution, various persons can be located in quick amount of time which will ensure the public safety and will improve the technological implementations of the surveillance and shows us the immense power of the AI in solving real world problems. The system giving the accuracy of 97.63% in pre-recorded video and 87.1% accuracy in real time.

Index Terms— MTCNN, InceptionResNetV1, CCTV, GFP-GAN, face identification, face recognition, missing person finding.

I. INTRODUCTION

The increasing number of missing persons is a significant challenge faced by the humans in today's world especially in the urban environment. Traditional methods for finding the missing person rely heavily on manual checking of CCTV footage - this process is time consuming as well as laborious but also inefficient when video is very long and crowded as very large population finding missing person is a tedious task. This approach frequently results in delayed action, and late recovery of the missing person.

In this paper, we have introduced a novel solution to this problem which integrates various machine models which are trained and will help us to find that particular person in less amount of time with more accuracy and in an efficient way. By integrating image enhancement techniques, string face detection and the good face recognition system, an intelligent surveillance framework can be made to aid law enforcements and rescue organizations. This paper introduces novel approach designed to detect and identify the missing persons in the crowded areas using both live and recorded CCTV footages. The system's novelty resides in its integrated pipeline design, which combines image enhancement, face detection, and face recognition into a cohesive workflow suitable for both offline and real-time analysis. Our proposed architecture enhances low-quality images, detect the faces in the video and generate the precise face recognition and localization of the person in the footage. By creating such architecture, the project seeks to benefit public safety and humanitarian purpose. The created solution can empower law enforcement and rescue organizations with real-time actionable intelligence enhancing missing person recovery speed and accuracy.

II. METHODOLOGY

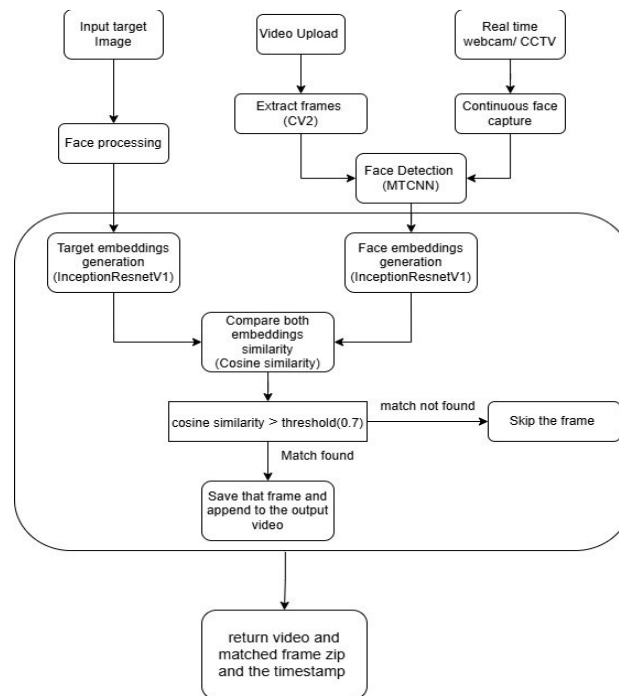


Figure.1 Vision Seeker system architecture

(Fig 2.1) The overall flow of our model which is divided into four robust parts namely Input processing and enhancement, Face detection, Face recognition, Output generation. A. Input Processing and Enhancement In the input section the user will input the image and video then our preprocessing model will enhance the input target image by adjusting its noise, saturation, blurriness using model Generative Facial Prior

- GAN (GFP - GAN)[11] the

task of this model is restoration of face from old, blurry or damaged photos. It has the following components

- Pre trained style GAN feature which provides a high-quality face feature.
- Unet-based generator maps degraded image into a realistic output.
- Degradation aware encoder helps to encode low quality image and degradation.
- Facial component encoder encodes facial feature using style-GAN priori

GFP – GAN requires an input image which is in low quality then it performs degradation encoding where the encoder analyze the noise pattern after this style-GAN prior is used for adjusting to match the degraded input following with U-Net structure uses this input and generate a new high-quality image as an output multiple loss function enhances the result to look realistic and preserve the identity. The final output from this model is enhanced quality, realistic face image with proper identity without loss of details a last sub-step in this step is to

change the format of image using image processing library in this case we are using python PIL/Pillow which takes any image as an input but limited to certain format like PNG, JPEG, BMP, TIFF, GIF. We import `save("output.jpg", format="JPEG")` function which saves the image in JPEG format and to maintain consistency we'll always convert the input image into jpeg format. This will ensure no other image format will disturb the code.

B. Face Detection

Completion of this step newly generated image and video will move to cloud for further processing where our detection model will deal with video. The video is converted into frames with the help of OpenCv.

To improve performance and avoid redundant frame processing, we perform temporal sampling. Given an original video with $fps_{original}$ and target processing frame rate as fps_{target} , we select the frames based on calculated skip interval. The set of sampled frames, $F(\text{sampled})$ can be defined as :

$$F(\text{sampled}) = \{f_i | i \left(\text{mod round} \left(\frac{fps_{original}}{fps_{target}} \right) \right) == 0\} \quad \dots (1)$$

where f_i is the i -th frame of the video. This equation processes the only selected frames and skips the frames. In our system we are processing the video at ~ 7 fps so that we will process the frames after specific interval of frames. Each sampled frame is passed through the MTCNN model for face detection.

If multiple faces are detected, embeddings are computed for all, and the one with the highest similarity to the target is considered, for a good precision and accuracy, we have used Multi-task Cascaded Convolutional Neural Networks (MTCNN) which is used to detect faces and facial landmarks localization. MTCNN is based on 3 CNN which is Proposal Network (P-Net), Refine Network (R-Net) and Output Network (O-Net). MTCNN works in a way where the video is split into multiple frames, P-Net is responsible for detection in each frame in this section the frame is resized and scaled properly. P-Net then scans the images and applies face bounding proposal with confidence score and non-maximum suppression to avoid overlapping detection of objects. Now comes the R-Net which refines the results of P-Net. It takes input from p-net and eliminates boxes which have low confidence score and then refines the boundary by using regression after completion of this the last step is performed which is of landmark localization where it marks the left eye, right eye, nose, left mouth corner and lastly right mouth corner.

At each stage, the network performs the three tasks simultaneously. The overall loss function for each is weighted sum of these tasks:

$$L_k = \alpha_k L_{cls} + \beta_k L_{box} + \gamma_k L_{landmark}$$

Where, L_{cls} is the facial classification loss, L_{box} is the bounding box regression loss, $L_{landmark}$ facial landmark localization loss we'll get multiple frames where faces are detected with high confidence score and this output will work as input for our face recognition model.

C. Face recognition

To perform recognition, we've used InceptionResnetV1 a face embedding and recognition model used to create embedding from faces and used for verification. It integrates residual connection and inception modules and trained using arc face, triplet Loss for large scale dataset. This model will take the frame as an input like aligned face and output a 512-dimensional feature vector which is compact numerical representation of the face. In this section one thing to keep in mind the embeddings of same person are very closer in terms of cosine space and then we will compare embeddings with cosine similarity checking. Cosine similarity is computed between target person embeddings (E_{target}) and the embeddings of the detected face ($E_{detected}$).

Cosine similarity measures the cosine of angle between the two vectors, making it robust to difference in vector magnitude.

$$\text{similarity}(E_{target}, E_{detected}) = \cos(\theta)$$

$$\text{similarity}(E_{target}, E_{detected}) = \frac{E_{target} \cdot E_{detected}}{\|E_{target}\| \cdot \|E_{detected}\|}$$

$$\text{similarity}(E_{\text{target}}, E_{\text{detected}}) = \frac{\sum_{i=1}^{512} E_{\text{target}}^i \cdot E_{\text{detected}}^i}{\sqrt{\sum_{i=1}^{512} E_{\text{target}}^2} \sqrt{\sum_{i=1}^{512} E_{\text{detected}}^2}}$$

The threshold is calculated based on try and evaluate basis, two most important metrics in this context are: False Acceptance Rate (FAR): The probability that the system incorrectly accepts an unauthorized person (an imposter). This is a critical security metric.

$$FAR = \frac{\text{total number of negative pairs}}{\text{Number of false positives}}$$

False Rejection Rate (FRR): The probability that the system incorrectly rejects an authorized person. This is a critical usability or convenience metric.

$$FRR = \frac{\text{total number of positive pair}}{\text{Number of false negatives}}$$

D. Output Generation

We have used the custom CCTV-style videos (40–60 seconds each, 720p resolution) recorded in varied lighting and crowd conditions for mimicking the CCTV video and laptop camera used for live testing in various environment & lighting conditions. The final stage, the system compiles the result as follows.:

For each frame where the match is found i.e. similarity > 0.7, a bounding box is drawn around the face. The corresponding similarity score is written on the box to maintain the transparency. All such matched frames are collated into the final output video and zip file is created, if no matching person found after processing the entire video “person not found message” is displayed.

E. Real Time Module

The real time modules work similarly as the recorded module but the main difference is that the frames are captured concurrently if the match is found message is displayed with person similarity score and compiled zip file with video.

F. Evaluation of Results

The results are evaluated based on the three important basis accuracy, precision and recall.

Accuracy: It measures the overall correctness of the model. It is the ratio of total correct predictions by the total number of predictions made.

$$\text{Accuracy} = \frac{TP + TN}{TP + TN + FP + FN}$$

Precision: It measures the quality of positive predictions, focuses on minimizing false positive

$$\text{Precision} = \frac{TP}{TP + FP}$$

Recall: It measures the model’s ability to find all the positive samples, focusing on minimizing the false negatives.

$$\text{Recall} = \frac{TP}{TP + FN}$$

Where, TP (True Positive): model correctly predicted the positive classes, TN (True Negative): model correctly predicted negative classes, FP (False Positive): model incorrectly predicted positive classes, FN (False Negative): model incorrectly predicted negative classes.

III. RELATED WORK

- A Real-Time Framework for Human Face Detection and Recognition in CCTV Images suggests an efficient and cost-saving system that can automatically detect and recognize human faces in real-time CCTV footage

by using machine learning and deep learning algorithms. The authors have created a framework having four main steps: image acquisition using an IP camera, image preprocessing (grayscale conversion and Canny edge detection), face detection using the Viola-Jones algorithm, and then face recognition through feature extraction with Principal Component Analysis (PCA) and classification using K-Nearest Neighbours, Decision Tree, Random Forest (RF), and Convolutional Neural Networks. Among all technique, CNN achieved the highest accuracy of 97.5%, showing that it's the most reliable, while KNN also performed well with accuracy over 94%.

- The paper demonstrates an excellent framework to enhance the performance of face recognition (FR) systems for face misalignment-a prevalent issue when faces are not aligned correctly between training and testing stages. To address this, the authors propose a face shaped-guided deep feature alignment approach that leverages facial key point heatmaps (face shape priors) at training time to direct robust face feature learning. The architectures incorporate pixel alignment and feature alignment. Through extensive experiments on benchmarks such as LFW, CALFW, YTF, and Mega Face, the method demonstrates better robustness to different misalignment cases and competitive or better performance than state-of-the-art FR and face alignment learning methods, and all while decreasing testing computational complexity.
- a very effective and lightweight system for real-time masked face recognition on mobile devices. Inspired by the heavy use of face masks during the COVID-19 pandemic, the authors address the challenge of facial recognition occlusion, which significantly reduces the performance of traditional face recognition systems. The system makes use of a Convolutional Neural Network (CNN) architecture optimized for mobile devices to ensure minimal computational overhead while maintaining high accuracy. The proposed model integrates modules combines modules for face detection, mask detection, and identify recognition in a streamlined pipeline. The solution keeps balances of performance with efficiency making it good choice for deployment in mobile applications, surveillance, and security systems where facial recognition under occlusion is critical
- A Multi-task Cascaded Algorithm with Optimized Convolutional Neural Network for Face Detection (MTOCNN) Methodology- Expands the Multi-task Cascaded Convolutional Neural Network (MTCNN) by improving architecture with optimizations and adding improvements. Eliminates landmark detection from initial stages (P-Net, R-Net) for simpler models and training Adds ELUs (Exponential Linear Units) in place of PReLU for enhancing convergence speed. Enforces Mosaic Data Augmentation (from YOLOv4) to enhance robustness by mosaicking four images for enhancing background diversity. Authors' Claims- 97.6% accurate on Fddb — a 2.7% improvement over MTCNN, and earlier models such as Faceness-Net (88%) and Cascade CNN (85.1%). Decreases training time by 20%, and model size by 1%. ELU activation outperforms PReLU in reducing loss on CIFAR-10 and LFW datasets.
- Blind Image Restoration and Data Augmentation Using GFP-GAN and StyleGAN Methodology- GFP-GAN employs a VGG-19 based encoder-decoder network for facial feature extraction from low-resolution images through identity-aware restoration. Integrated StyleGAN for high-level feature (style) and low-level content (structure) processing in the latent space to perform data augmentation. Real facial datasets augmented by joint training enhance structure realism and image fidelity. GFP-GAN. StyleGAN employs a number of facial image references in order to augment dataset size and scope. Contributes improved performance in low- quality input scenarios like forensic application, surveillance, and medical imaging.Limitations- The GFP- GAN can hallucinate details or produce artificial colouring when grayscale or severely degraded input is given. Persistently struggles with extreme pose variations or severe blur.
- Smart Face Recognition System using Hybridization on of MTCNN and Facenet Algorithms IEEE International Conference paper methodology used is MTCNN FaceNet based recognition CNN powered surveillance detection Limitations in Practical scenarios needs improvement for Deployment.
- Identity Verification in Real Time Proctoring: An Integrated Approach with Face Recognition and Eye Tracking Siamese Network with Eye Tracking VGGFace-2 is used for pretraining and eye tracking GBM model. Accurate Face and Attention Tracking. Limited Generalization for Underrepresented Groups Focus on Cross Demographic Calibration.
- Comparative study of Face Detection using Cascaded Haar Hog and MTCNN algorithms Algorithm Performance Under Various Conditions. Training Classifiers with Face/Non -Face Images MTCNN Out performance Haar, HOG Haar and HOG Lightweight, Limited Performance balancing model complexity and performance.
- The paper demonstrating the framework to enhance the performance of face recognition ;a prevalent issue when faces are not aligned correctly between training and testing stages. The authors propose a face shaped-guided deep feature alignment approach that leverages face shape priors at training time to direct robust face feature learning. This incorporate feature and pixel alignment. Through extensive experiments on

benchmarks the method demonstrates better robustness to different misalignment cases and competitive or better performance than state-of-the-art face recognition and face alignment learning methods, and all while decreasing testing computational complexity

- **Face Recognition from Blurred Images Using Deep Learning: Methodology-** Describes and contrasts conventional vs deep learning-based methods of face recognition for blurring situations. GFP-GAN used for face restoration before recognition. Examines several deep learning-based models of recognition: FaceNet, VGG-Face, DeepFace, and a self-programmed ResNet50. Dataset Utilized Surveillance Dataset 8,000 blurred faces with different illumination and poses. MSCeleb-1M: Massive celebrity dataset for training. Chicago Face. FaceNet, because of its magnitude and 128-D embeddings, gives the best results with 83.72% accuracy on blurred, recovered faces. GFP-GAN facilitates identity-preserving deblurring, improving recognition outcomes. Classical models such as DeepFace overfit and perform poorly with blurred data (accuracy only 50%). Limitations- Blurry, occluded, and far-off views continue to damage recognition accuracy. Restoration networks may inject artifacts or bias facial features. Scope Restoration with GFP-GAN, along with a strong recognition model such as FaceNet, greatly enhances face recognition for difficult conditions. Progresses a pipeline strategy: Restoration -Embedding - Matching as a strong technique.
- **GFPGAN vs. PULSE – PSNR-Based Image Restoration Comparison: Methodology-** gfpGAN uses a pre-trained facial prior model to condition multi-level restoration, with the aim of ensuring realism and preserving identity. PULSE: Leverages latent space exploration of StyleGAN for restoring high resolution images by optimization of latent codes to align the degraded input. This is an iterative, computationally intensive process. GFP-GAN performs better or on par with PULSE in PSNR on evaluated datasets. GFP-GAN provides a more efficient, faster, and identity preservation solution compared to PULSE, which is slow and occasionally creates artifacts. GFP-GAN achieves a superior visual realism-structural faithfulness trade-off in the recovered images. Limitations- PULSE and other latent-optimization methods are computationally expensive, hence real-time use is not forthcoming. Although GFPGAN solves the identity preservation, its quality can deteriorate for very severe corruption, like low-resolution or occlusion-dense inputs. GFP-GAN is presently a candidate for useful practical real-world face-associated image restoration usage where speed and retention of integrity are critical issues. Illustrates how the utilization of semantic facial priors in GAN models is able to surpass the shortcomings of regular latent optimization techniques. [12]

Existing systems often focus on individual components enhancement, detection, or recognition without offering an integrated, deployable solution for missing person search. Vision Seeker fills this gap by combining robust pretrained models into an end-to-end pipeline optimized for real-world CCTV footage.

IV. RESULTS AND DISCUSSION

To validate the effectiveness of the missing person detection, the testing is conducted on the recorded offline video and by the real time webcam feed. As explained in the methodology (FIG. 2.1), the system evaluates the video on the basis of cosine similarity between target and candidate embeddings. In case of pre-recorded video instead of checking the, every other frame in the second of the video we will be doing at ~7 fps and we are batch processing the images face comparison to reduce the time of processing and make it faster.

At original fps value the processing requiring the more time than the duration of the video and sometime same time as video based on the video and crowd and various condition. To mitigate this issue we have decided to perform processing at the 7fps.

Based on the empirical evaluation, the similarity threshold of 0.7 was selected to balance the false acceptance rate (FAR) and False Rejection Rate (FRR). In pre-recorded tests, the system achieved a FAR of around 2-3% and a FRR of 3-4% at this threshold, achieving an effective trade-off suitable for surveillance applications.

The prerecorded video takes nearly 20.3 seconds to process the 40 second video at 3 fps and 30 second at 7fps. If the matched frame cosine similarity > threshold (0.7 in our case) each matched frame, if identified, was saved, timestamped, and compiled into a ZIP archive and annotated video output for further verification. If the match not found it returns "Target not found in the video". The target image processing is done to make it proper, if it's blurred, low quality or in low light and the output is shown in the fig. 1.

The target image after embedding generation is compared with the embedding of the faces in the frames of video or the real time frame capture. It returns the video and the zip file of matched faces frame with the target image with a rectangle on the target image and similarity score. The processed target image by GFP-GAN. A side-by-

side target and matched face frame is highlighted in the (fig. 2). Highlighting how a face match looks like in the targeted image and the corresponding video frame.



Figure 1. Target face optimization

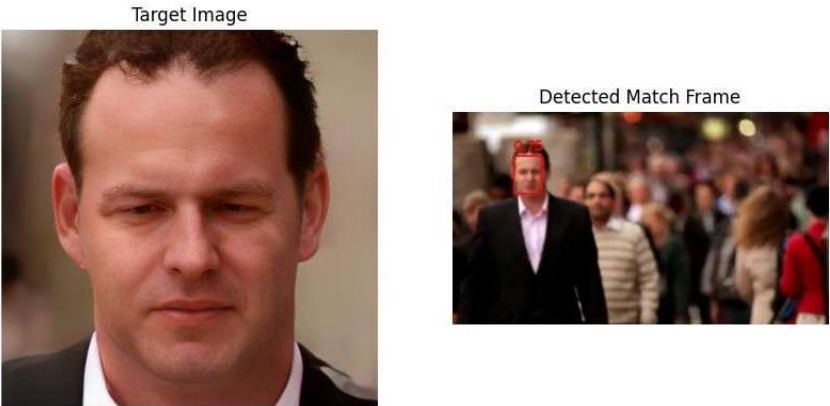


Figure 2 Example of target image and matched frame with bounding box and similarity score

If target not found it returns the message target not found as shown in fig 3 Not found comparison of performance of system in real time and the recorded video example is shown in the Table I

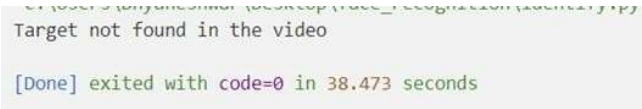


Figure 3 not found message

TABLE I. PERFORMANCE COMPARISON

Measures (Average)	Pre-record video (%)	Real time webcam feed (%)
Accuracy	97.62	87.1
Precision	99	98.23
Recall	85.71	84.3
F1-score	91.87	90.73

The system performed well on the prerecorded video achieving the high accuracy the high precision indicates that when face is detected it is almost always correct. The system misses the actual occurrences of target face; this might be due to blur images, lighting conditions or motion blur in the frames. Real time introduces more variability like lighting fluctuation, frame dropping and hardware latency. in all of this the system has maintained high precision. It's less as compared to pre-recorded this might be because of the

real-time operational variability. We have processed the 5 recorded videos with various environmental conditions.

The comparison of pre-recorded and real time accuracy, precision and recall is given in the bar graph in the figure 4.

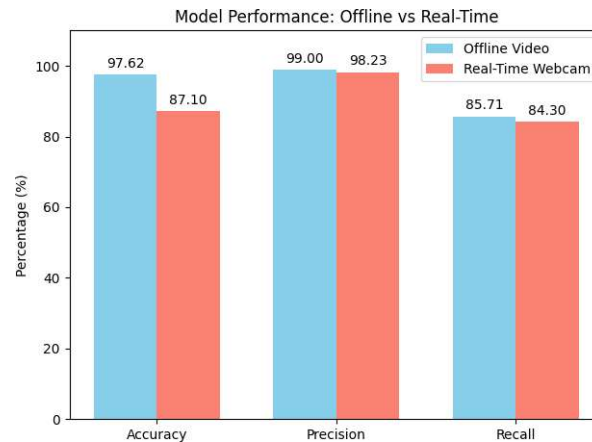


Figure 4. comparative bar graph prerecorded vs real time

V. CONCLUSION AND SUMMARY

A. Conclusion

Through this project we aim to focus on finding the missing person through real time video captured through CCTV. Using various technologies blended together, this results in reducing the time and effort to locate a particular person from a large volume of people. Faster finding of a missing person enabling fast analysing and identifying of the missing person through integrating our system to CCTV and ultimately reducing the searching time to a less amount. support for law enforcement Helping the police and different investigative agencies to help them find and locate missing person which will help the official people use it in a very efficient. This project is a solution to a critical problem of finding missing person. By integrating and leveraging firstly GFP-GAN for image enhancement, MTCNN for face detection, and Face_net for facial recognition, it automatically reduces the time and efforts one needs to find and locate a missing person. It also provides many types of social benefits which will aid our society in the time of crisis. It is an important step towards a more secured technological world which will help attain various good attributes to the boon of our society.

B. Summary

This project aims to deliver the rapid and best use of technologies mainly MTCNN for face detection, and Face Net (InceptionResNetV1) for facial recognition through embeddings and cosine similarity, which will detect the image of the missing person in the best possible way. It will enhance public safety and also the impact of it towards law and enforcement will be a good one.

Software Components

- GFP-GAN (Generative Facial Prior - Generative Adversarial Network): It is a machine model which is used to enhance the image quality which will help to identify the missing person in a more enhanced way. It has various preprocessing modules which helps to image to be cleaned and blur-free.
- MTCNN (Multi-task Cascaded Convolutional Neural Network): This model helps us to find out that particular image from the video. It splits the video into multiple frames. The modules detect the key features of the image and detects it and makes a box around the resulted image.
- InceptionResNetV1: It is a model used for image recognition and used for face embeddings generation. The libraries used here is facenet_pytorch for InceptionResNetV1, cosine similarity is used to matching the embedding.

C. Future Work

Several avenues can be explored to extend and enhance the current system

- Adding the functionality to track the multiple persons in the database in the prerecorded video and at real time.
- Creating the embeddings from the multiple face angles such that can get the more accurate results.
- Incorporate the ethical guidelines like face blurring the non-target individuals and blockchain based face embeddings storage.
- Integration of Temporal Face Tracking: Incorporating tracking algorithms such as SORT or Deep SORT can allow the system to maintain temporal consistency.
- reduce false positives, and more robustly follow the same person across frames, especially in crowded or dynamic scenes.

REFERENCES

- [1] Ullah R, Hayat H, Siddiqui AA, et al. A Real-Time Framework for Human Face Detection and Recognition in CCTV Images. Ali N, ed. Mathematical Problems in Engineering. 2022;2022:1-12. doi:<https://doi.org/10.1155/2022/3276704>
- [2] Kim HI, Yun K, Ro YM. Face Shape-Guided Deep Feature Alignment for Face Recognition Robust to Face Misalignment. IEEE Transactions on Biometrics, Behavior, and Identity Science. Published online 2022:1-1. doi:<https://doi.org/10.1109/tbiom.2022.3213845>
- [3] Kocacinar B, Tas B, Akbulut FP, Catal C, Mishra D. A Real-Time CNN-Based Lightweight Mobile Masked Face Recognition System. IEEE Access. 2022;10:63496-63507. doi:<https://doi.org/10.1109/access.2022.3182055>
- [4] Zhang L, Wang H, Chen Z. A Multi-task Cascaded Algorithm with Optimized Convolution Neural Network for Face Detection. Published online January 1, 2021:242-245. doi:<https://doi.org/10.1109/acctcs52002.2021.00054>
- [5] Harsh Khatter, Tyagi N, Avishi Tayal, Gupta P. Blind Image Restoration and Data Augmentation. Published online March 15, 2024:32-37. doi:<https://doi.org/10.1109/icdt61202.2024.10489715>
- [6] Raj N, Agarwal V, Gupta N. Smart Face Recognition System Using Hybridization of MTCNN and Facenet Algorithms. Published online December 7, 2024:1-6. doi:<https://doi.org/10.1109/inspect63485.2024.10896125>
- [7] Chatterjee P, Jayanti Dansana, Swain S, Mahendra Kumar Gourisaria, Bandyopadhyay A. Identity Verification in Real Time Proctoring: An Integrated Approach with Face Recognition and Eye Tracking. Published online August 23, 2024:1-6. doi:<https://doi.org/10.1109/iacis61494.2024.10721819>
- [8] Thilinda Edirisooriya, Eranda Jayatunga. Comparative Study of Face Detection Methods for Robust Face Recognition Systems. Published online December 6, 2021. doi:<https://doi.org/10.1109/slaai-icai54477.2021.9664689>
- [9] Kim HI, Yun K, Ro YM. Face Shape-Guided Deep Feature Alignment for Face Recognition Robust to Face Misalignment. IEEE Transactions on Biometrics, Behavior, and Identity Science. Published online 2022:1-1. doi:<https://doi.org/10.1109/tbiom.2022.3213845>
- [10] Nukala S, Yuan X, Roy K, Odeyomi OT. Face Recognition for Blurry Images Using Deep Learning. Published online May 24, 2024:46-52. doi:<https://doi.org/10.1109/ccai61966.2024.10603301>
- [11] Shravan D, Ramkumar G, Meenakshisundaram N. Generative Facial Prior Generative Adversarial Networks based Restoration of Degraded Facial images in Comparison of PSNR with Photo Upsampling via Latent Space Exploration. 2024 International Conference on Advances in Data Engineering and Intelligent Computing Systems (ADICS). Published online April 18, 2024:1-5. doi:<https://doi.org/10.1109/adics58448.2024.10533635>
- [12] Kaur S, Sharma D. Comparative Study of Face Detection Using Cascaded Haar, Hog and MTCNN Algorithms. Published online November 23, 2023. doi:<https://doi.org/10.1109/aece59614.2023.10428242>