

# Enhancement of a Machine Learning Allocation with Video Copy Detection using IoT with Steganography for Raspberry Pi

<sup>1</sup>Karthika P and <sup>2</sup>Balamurali S

<sup>1,2</sup>Department of Computer Applications, Kalasalingam Academy of Research and Education, Krishnan Koil, Viruthunagar, Tamil Nadu, India.

<sup>1</sup>karthikasivamr@gmail.com, <sup>2</sup>sbmurali@gmail.com

Correspondence should be addressed to Karthika P: karthikasivamr@gmail.com

## Article Info

Journal of Machine and Computing (<https://anapub.co.ke/journals/jmc/jmc.html>)

Doi : <https://doi.org/10.53759/7669/jmc202505070>

Received 16 May 2024; Revised from 07 December 2024; Accepted 21 February 2025.

Available online 05 April 2025.

©2025 The Authors. Published by AnaPub Publications.

This is an open access article under the CC BY-NC-ND license. (<https://creativecommons.org/licenses/by-nc-nd/4.0/>)

**Abstract** – The IoT is being created by Internet and billions of physical devices (IoT). These devices are producing more useful or meaningless data. This data must be prepared and transmitted, which is a difficult task. Future research directions and other IoT applications are also discussed. The security sector is undoubtedly one of the IoT framework's application areas. It is essential to a common little effort solution to prevent wrongdoing and guarantee the safety of individuals from the home, military, industry, and other settings. This study covers the depiction driven evolution procedure for AI Security System employing Raspberry Pi IoT setup. It makes note at end of the client's expectation. Such kind of instance is to share the information and protect with the Internet of Things that are more dependent by having the methods of Artificial Intelligence which supports the level of proximity and some of the non-mistaken with an accuracy of 90 percent proof, and the major confirmation to be made here is being a part of this. This method utilizes the help of electric door that are striking the actuator and the USB type web camera with the example of Image-gathering. It is also a kind of application with a standard programming Interfaces to collect similar gaming plans which are to be more related with the Internet of Things which are based on boarding knowledge of Raspberry Pi and also the steganography type distribution.

**Keywords** – Artificial Intelligence, Machine Learning, Internet of Things, Webcam, USB.

## I. INTRODUCTION

The IoT can be considered as arrangement of systems also the methods for exchanging or gathering data. Security [1] accepts a fundamental activity over the necessity care, condition, insurance, and other worries, as the (challenging) dataset is transported over the various supported devices and various clients. When such information is discovered and used abruptly, it may result in extreme harm to the productive structure and the potential assets. Several diagrams which are relating to Internet of Things security, IoT structure security, or unambiguous IoT system components [2]. It suggests a wide-ranging area display for security risk organization in IoT systems. As a result, the IoT structure's information and data are exposed to security risks. As a result of how heavily these systems rely on both the cloud with Internet handling, we wonder how vulnerabilities do the mitigation are handled in the management of high security web application scheme. This can be useful for assessing and controlling the security threats posed by these architectures. It updates for the video steganography assignment for Raspberry Pi type of board would suggests for default for managing security and moreover proposed for the type which are declared in the form of SRM model. To use accessible game plans in verifying [3] web application structure.

As visual and aural gear and software advance quickly and are used widely, the cost of combining, creating, and limiting picture and video data continues to decline. A staggering amount of video data is created and shared on a regular basis. There are enormous numbers of duplicates or almost identical narratives throughout these enormous volumes of chronicles. Overall, according to estimates from social media platforms, Google and other resource video web crawlers, these things would range up to 27% dull accounts which are identical to or closer identical to the most type sort of a video that are most related with the question. An efficient and competent method for video copying has thus become more notable.

Nowadays, the role played by the Internet in people's lives is becoming increasingly important. The unlimited resources on the Web are delivered in a variety of formats, including text, images, sounds, videos, and more. Because of the benefits

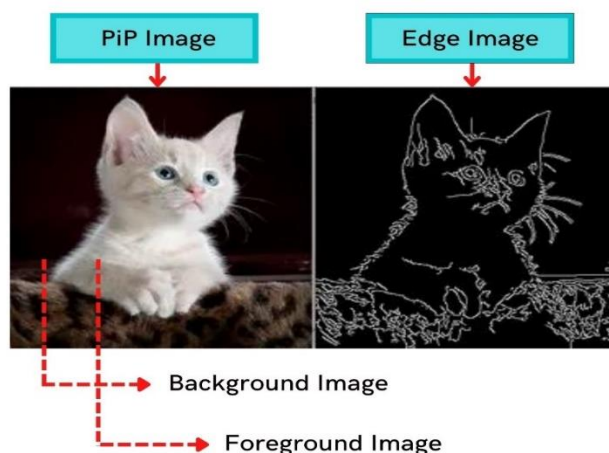
of verbalization and creating a similar interest, video is the type of media that people like. However, there are several histories with the same information, which adds to the burden on the structure. By computing their respective hash estimates, video copy ID determines whether two accounts are from the same material by judging how closely related they are to one another.

The Final output of the essay is too organized as of that. Segment 2 related efforts that improve the security of multimodal biometric data by utilizing an IoT multi-layer framework. A Raspberry Pi-based AI detection framework is shown in segment three. This section presents the framework architecture, several low-effort sensors, a local figuring component based on the Raspberry Pi, and data processing programmes. The fourth segment of the video duplication detection is usually dedicated to identifying the major level management, and the image recognition which can be represented as an image at the edge. Most excellent type of execution to date with a significantly increased time consumption is the differentiating proof precision execution that are similar to method for utilizing all the final edges. Section 5 provides a conclusion.

## II. EXISTING TECHNOLOGIES WITH THE ANALYZED REVIEW

An Internet of Things multi-layer structure which are combines steganography and watermarking [4] techniques to increase the security of multimodal biometric data. It conceals faces in images of fingerprints using a watermarking technique [5][6] Eigen capability. On this stage, the user used a technique called the steganography by hiding data due to most related subjective impression and sometimes the face watermarked biometric or criminological photograph. In each of the three subjective visual channels' three hues, recess mounting sites are randomly placed.

Digital watermarking techniques may be used to link sensitive data, such as the organization's emblem on the host data type, by safeguarding the rights of protected invention of massive dataset [7]. Biometrics-based individual differentiating proof tactics are also sometimes used. These are also used to verify news reports from the media. After enrollment, encryption may be used to improve the security model for biometric data that is stored in the form of database and the second format as of brand (such as the keen type card), or 3) with the biometric device (such as a mobile phone with a unique mark sensor). The result pertaining to biometric information on the web may then be created by looking through and using encrypted formats during validation. The offered encryption models cannot be accessed altered without decoding with the right option.



**Fig 1.** Image Edge Detection Model.

Perception of outline level is the main emphasis of the video differentiating proof. In addition, picture disclosure has the possibility of being displayed as acknowledgement into a square shape that might represent the boundary between a closer view and an establishment picture [8]. The edge picture of **Fig 1** with the square shape having the highest likelihood of being red stepped. Soma of the identifiers is being employed to [9] distinguish added vertical edge location inside those recent pictures when it comes to edge line acknowledgement. If one level parade is maintained throughout the entire film is captured with a defined level line, then just even lines are perceived. The image area can be set up with two or three creative even shapes that might theoretically be used to construct two equal lines in the form of a square [10].

The most likely there are two vertical lines and that are even too in the shape of square [11] is taken into consideration as a kind of image location when level edge and also the vertical lines are observed. In order to get a mean-edge arrangement for image limitation, the edges are differentiated for each major edge. Even more conclusively, shape [12] video image outcomes are related to mix of packaging level picture results. The square shape would be supplied as similar to video image district after a vote form is thrown and the total of square shapes observed from each key packing. Nevertheless, these current approaches have a few shortcomings. If each packing in the movie is taken care of for a certain object, the time utilization is not allowed. Another issue is that if some important housing is examined [13] in order to decipher the

ID, time intelligibility information will be mostly lost and the image area won't be able to be viewed and identified clearly. The suggested approach in the accepting zone is to throw out the two folds of deficiencies.

### III. PROPOSED AREA WHERE BOTH THE AI AND ML ARE APPLIED

The three levels of security that have been suggested can concurrently check everybody and secure the biometric system [14]. An iris with a watermark is used as part of a layout that are being identified with the watermark on the image face, confirming that the face can be recognised, and the security of the biometric (facial) is also maintained. For instance, a computation based on log on1D Gabor was used to create the model iris watermark information. The model is created by convolving Dlog Gabor with default surface of modified iris type of image. Iris Code is one of the model's names [15]. This iris code has a two-fold structure and is unique to each individual. As similar to operation of various multimodal gives 2 degrees of enhancing the security [16] to finalise individual and concurrently ensure those biometric formats, it is still being included with image of substance of a similar human to ensure the actual model. A biometric format is used for making the cross to verify the image with hidden person and the security of biometric, with iris in the face is being watermarked as of similar to confirming the Stomach difference with the face. The proposed computation makes use of three security levels to maintain the accuracy, safety and the personal data. The basic layer of security uses specific watermarking to conceal [17] the client's personal information, including any data that could be most related to the (Aadhaar card, DOB) this might be a single picture and that creates a major retinal fingerprint, while it is being used with the watermark to the image while hiding the mystery from inclusion through the computation of Steganography.

### IV. RESULT AND THE PART OF DISCUSSION

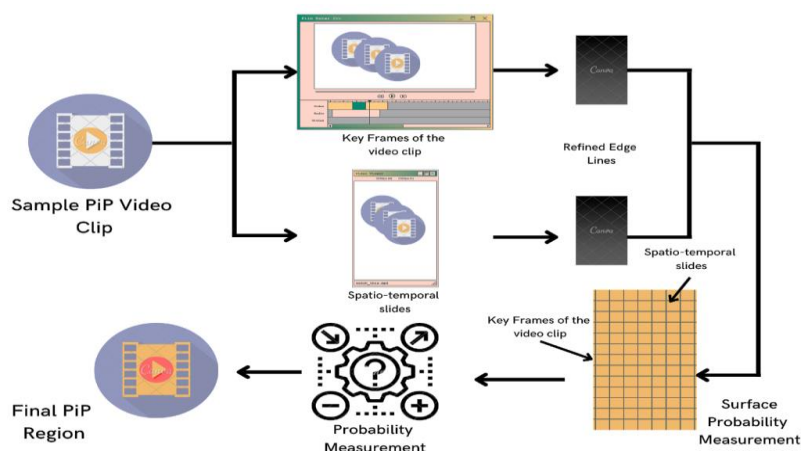
#### *Algorithm in the form of Steganography*

An organization must thoroughly screen traffic that is being managed with time- and processor-intensive, to discern hidden messages. However, those who are concerned with the framework's typical traffic patterns may easily monitor [18] for modifications, such as the continued advancement of large images across the framework, which may call for a more thorough evaluation. It's also shrewd to have - and effectively maintain - a security strategy that lays out in plain sight acceptable use, which data types may and cannot be communicated over the framework, and how it should be ensured. In a similar vein, continue using unapproved applications, forbid using unapproved steganography and encryption at work, and take into consideration forcing the measure [19] of post boxes.

Finally, decide if workers in charge of secret information should direct their attention to vast media archives, notably images, videos, or sounds [20] that will be put on your website. Malignant gatherings employ steganography for transmit data for outcasting method that has access and permitting the site using techniques for such records. Why not utilise automated identification and the type of steganography through the copyright of Web-open recordings, further strengthening your good fortune? Even system passwords or keys may be concealed with it, providing a permanently secure boundary zone.

#### *ANN Algorithm*

Some of the occasion in featuring the vector forms the primary picture database by created picture database are used in the ANN to form up the neural [21] framework. The Back inciting Neural Network is constructed with 300 segment vectors (150 from a library of unique images and 150 from images that have been seen before). The ANN in use consists of 6 hid neurons, 30 info neurons, and 1 yield neuron. The ANN is tested using the remaining 300.



**Fig 2.** Picture to Override and to Highlight Some of the Edges Based on Estimated Sectors.

The position of the video duplication is frequently focused on identifying level, and image recognition that are shown as an edge image. **Fig 2** since the video was clipped, sampled key casings and replacements were extracted. The essential

edges are distinguished and refined using vertical and level edge lines, using the probability of 4 different surfaces that are accessed by the pictures.

Although the video may be viewed in a collection of images and similar to the picture with the actual video that is really chosen, in an average estimation can be used to develop the image model as shown in **Fig 2**. As a result, the focus of video image area shifts from identifying a square form with 4 edge surfaces in a major territory while establishing the accounts to recognising powerful shapes with four edge surfaces in bleeding edge and establishment photographs.

Steps in an ANN algorithm

1. Batch Size = 128, Epochs = 5, Classes = 10, and
2. Measuring the supplied picture to be 28 \*28,
3. Taking input pictures from a data collection.
4. Variable exploration: k=Train data set, X=Test data set (15000,28,28,1). (90000,28,28,1)
5. Create and put together the models
6. Network training.

The algorithm details the main procedures that are being involved with the conditioning and testing the data type for ANN picture categorization.

Artificial Neural Network consider as one of the dataset D and it consists of various k objects  $\{x_1, x_2, \dots, x_k\}$   $x_i$  ( $1 \leq i \leq k$ )

$H(A_j) \rightarrow$  can be considered as image attribution.

$V \rightarrow$  can be denoted as variables

$P \rightarrow$  denotes the Probability of  $A_i, A_j$

$$H(A_j) = -\sum_{v \in \text{domain}(A_j)} P(A_j = v) \log P(A_j = v), \quad (1)$$

$$I(A_i, A_j) = H(A_i) + H(A_j) - H(A_i, A_j) \quad (2)$$

$$1 \leq i, j \leq m (i \neq j), \quad (3)$$

$m \rightarrow$  this can be denoted as an attribute with the n object  $\{x_1, x_2, \dots, x_n\}$ , where  $x_i = \{x_{i1}, x_{i2}, \dots, x_{im}\}$ , approximate differential of  $x_0$

$$\hat{h}(x_0) = \sum_{i=1}^m W_x(y_i) \left( \log a - \frac{a}{b} \log b \right) - a W_x(Y) + a \sum_{i=1}^m OF(x_{0,i}), \quad (4)$$

Were

$$W_x(y_i) = 2 \left( 1 - \frac{1}{1 + \exp(-H_x(y_i))} \right) W_x(Y) = \sum_{i=1}^m W_x(y_i) H_x(y_i) \quad (5)$$

A major advantage is  $W_x(y_i)$  is a collection of real numbers, assuming that the sequence of numbers is real. Equations 3 and 4 use the most common lossy advanced picture pressure framework currently in use for two-dimensional advanced image processing.

$$OF(x_{0,i}) = \begin{cases} 0, & \text{if } n(x_{0,i}) = 1 \\ \delta, [n(x_{0,i})], & \text{otherwise} \end{cases} \quad (6)$$

$$\delta(x) = (x - 1) \log(x - 1) - x \log(x), \quad (7)$$

$$\hat{h}(x_0) = \sum_{i=1}^m \left( \log a - \frac{a}{b} \log b \right) - a H_x(Y) + a \sum_{i=1}^m OF(x_{0,i}), \quad (8)$$

Equation 5 has a calculation that ensures accuracy by combining arbitrary properties with the 2 different coefficients. Higher of x is used to for getting stronger and stronger. Calculations and for implantation and extraction

$$D_s = \frac{D - D_{\min}}{D_{\max} - D_{\min}} \quad (9)$$

$$D_s = \frac{D - D_{mean}}{\sigma_d} \quad (10)$$

The method mentioned above ensures that Equ.6 does not ignore the image's  $D_s$ , since doing so would prevent the coefficients from applying the desired connection without gradually harming the picture data present in the specific picture.

$$y_{min} = \int (W^T \cdot S + C), \quad (11)$$

Basic research concluded that twelve parcels can be a good trade-off where the length of  $y_{min}$ : longtype enough for detect gadget and such kinds and major enough to be fully filled with intriguing bundle from  $y_{min}$ . To get the size of 276 highlights, however, cushioning with 0 quality is used if the  $y_{min}$  does not include enough remarkable bundles to complete it.

$$P_{copy} = P(y_{min} = 1 | S) = \frac{\exp(W^T \cdot S + C)}{1 + \exp(W^T \cdot S + C)} \quad (12)$$

We suggest a two-overlap  $P_{copy}$ -system approach to be flexible and pertinent for a growing variety of gadget kinds. First, we used to train a single type of classifier for similar device type in Equ. 8. Whether the inputs that are given is unique finger impression by coordinating the type of gadgets or not is a paired choice provided by each classifier.

$$\text{logit}(P_{copy}) = \log\left(\frac{P_{copy}}{1 - P_{copy}}\right) = W^T \cdot S + C. \quad (13)$$

An Equ.9 Some classifiers can recognise an obscure unique finger print, which allows them to classify various device types. In these situations, logit ( $P_{copy}$ ) is used to break ties between various matches by using an alternative separation measure. Although alter deliberate alone might be distinguish across gadget kinds, this approach is more laborious than order.

$$y_{min} = \begin{cases} 1, & \text{logit}(P_{copy}) > 0.5 \\ 0, & \text{logit}(P_{copy}) < 0.5. \end{cases} \quad (14)$$

In Equation 10, the important edges—both the level and vertical lines that are cornered—are distinct and complicated, and the probability of 4 surfaces are determined by the images' depictions of edges. While a video will be seen as a collection of all type of images, the image area through which a video cut is made is picked.

$$H_o^{(i)}: \omega_i = 0 \quad (i = t, \omega, h, p, r). \quad (15)$$

As a result, the central subject of the video image area becomes less apparent as having a form of square.  $\omega_i = 0$  ( $i = t, \omega, h, p, r$ ) using the four edge lines between cutting-edge which are established images by identifying solid objects with 4 edge surfaces in the frontal area and established images.

$$\omega = \frac{B_{\omega}^2}{SE_{\omega}^2}, \quad (16)$$

Key edges may be noticed in the attempted edges, and each term which describes the edge tends to have 1 estimation line. Equations 12 and 13 are mostly responsible for incorrectly estimating the probability of surface of the edges with the default surface would be, especially in stories with a constant sceneexchanging surface,

$$P = P\{\omega \geq \omega_{\alpha}\}. \quad (17)$$

$$k(S_i, S_j) = f(\|S_i\|, \|S_j\|). \quad (18)$$

Movie that is available for viewing as a series of images proceed via spatial estimate( $S_i, S_j$ ) many images with estimation, the spatio-common ready to assessed, and the brief estimation  $k, f(\|S_i\|, \|S_j\|)$ , in Equ.14

$$k(S_i, S_j) = \|S_i\| \|S_j\|. \quad (19)$$

For example, Equ.fifteen a comparative line of an  $y$  that is equivalent to (observing in red red in  $k(S_i, S_j)$ ) is picked from all of the video lines, and if necessary, an estimation of a spatio-transient cut is made.  $\|S_i\| \|S_j\|$  of the entire video.

$$k(S_i, S_j) = \frac{S_i \cdot S_j}{\|S_i\| \cdot \|S_j\|}. \quad (20)$$

When choosing the probability function of the edge surface in Equation 16, stays in vertical lines that are located between the frontal territory and establishment accounts are employed concurrently through level lines acting as key housings.

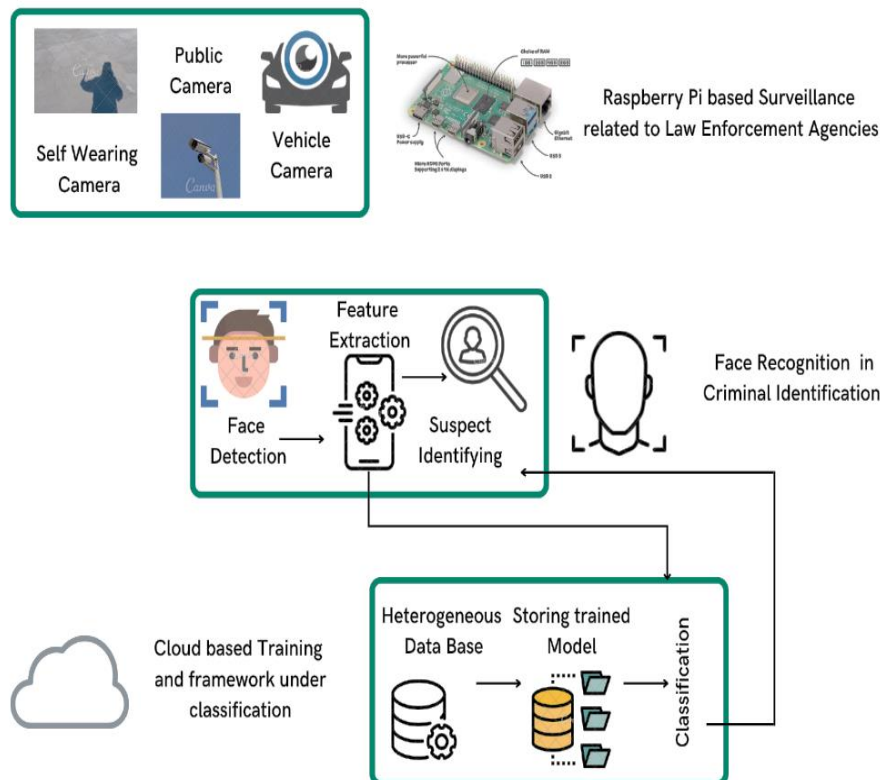
$$k(S_i, S_j) = \left( \frac{S_i \cdot S_j}{\dim(S)} \right)^2. \quad (21)$$

Equation 17 removes the level and vertical edge images separately ( $S_i, S_j$ ), which substitute the current two headings for the prior eight headings, lowering the figure of point course. **Fig 3.** Shows Detailed View of The Proposed Face Recognition Framework for Suspect and Missing Person Through Identification. The chance of real edges and upheaval edges is now increased by choosing a truly small edge for acute edge assurance. In addition, picture extraction is intended to link a little nearby edge.  $k(S_i, S_j)$  divisions while cutting short, detached ones. The raucous edges are intended to reduce the number of edge lines used for extra frameworks. The Raspberry Pi utilised a method for video copy acknowledgement.  $k(S_i, S_j)$  that is enlarged in every requested video, the image is detected and restricted (wherever distinct). The following figure, **Fig 3** represents the detailed view of the proposed face recognition framework for suspect and missing person through identification and **Fig 4.** represents the Region-of-Interest selection related to ORB.

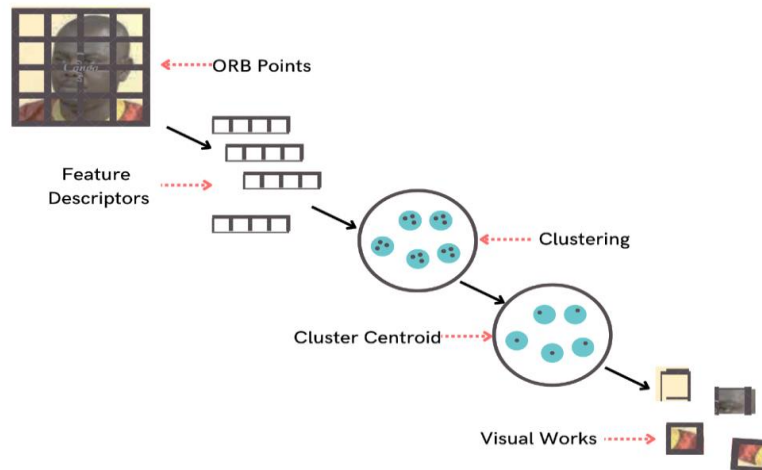
**Input:** Informative and correlated subsets  $L = \{A_1, A_2, \dots, A_i\}$  of  $D$   
**Output:** Label matrix label of  $L$   
 $L \leftarrow \emptyset$   
 Compute and store the Refined Approximate Differential Holoentropy  $\widehat{h}_r(x_0)$  and  $OF(x_{o_d})$  for all  $x_{o_d}$   
 For all  $x_{o_d} \in D$   
 If  $\widehat{h}_r(x_0) > 0$  or  $OF(x_{o_d})$  is the minimum label  $l_{i,j} = 1$   
 Else label  $l_{i,j} = 0$   
 Return label

ALGORITHM 1: Outlier detection on informative dataset.

Yrain  $\rightarrow \{(0, 1) 1 \rightarrow \text{Rainfall} > 50 \text{ mm per day or } 30 \text{ mm per } 12 \text{ hours and } 0 \rightarrow \text{otherwise}$



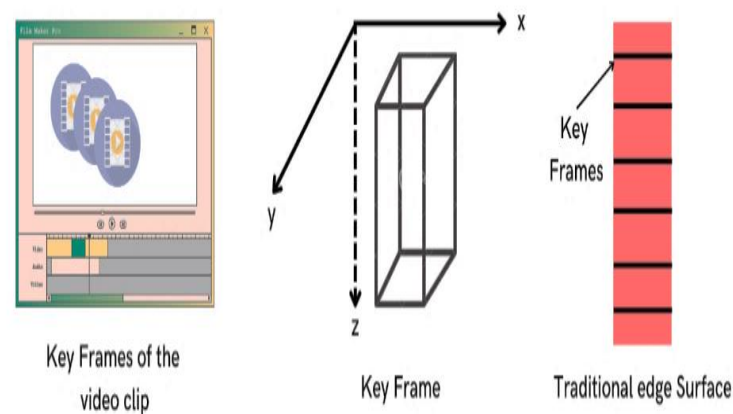
**Fig 3.** Detailed View of The Proposed Face Recognition Framework for Suspect and Missing Person Through Identification.



**Fig 4.** Feature Extraction That are Related To ORB.

**Table 1.** Raspberry Pi Managing Related to The Proposed Model

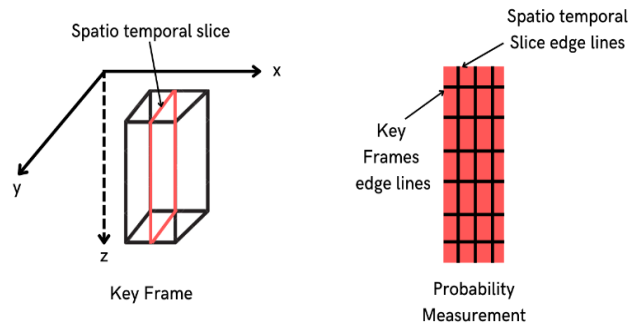
| Name of the models                            | Detailed about the configuration                                                                  |
|-----------------------------------------------|---------------------------------------------------------------------------------------------------|
| Operating System                              | Noobs type of board (Raspbian)                                                                    |
| Coding Experience                             | Python version 2.7                                                                                |
| Library functions                             | Numpy, Scientific Python, PythonLab, Matplotlib.pyplot, GPIO and RPI                              |
| Imaging Tools                                 | OpenCV version 3.1.0, Scikit-Learn, Scikit-Image, MATLAB packages that are matched with the Board |
| Analyzing the performance with its utilities. | BCMStat, CFML, CommandBox, ContentBox, Profiler in the base of line                               |



**Fig 5.** Estimation That Is Most Related to The Edge Surface and Casing with The Conventional Tool.

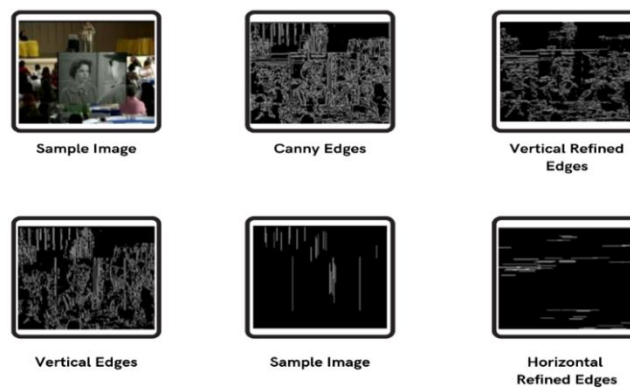
Although it is simple to estimate the chance of four surfaces in, it is tedious to eliminate edges from every edge of the video. Similar to this, certain Artificial Intelligence techniques are mostly related to animate video image disclosure. **Table 1.** Shows Raspberry Pi Managing Related to The Proposed Model There are several key edges visible in the edges, and each of the key type edge corresponds to a certain measurement line on the edge surface that are being (appeared in **Fig 5**). Probability of edge formed surface may be obtained by identifying the actual probabilities of the all-similar lines. However, the video's sequential order will not be preserved according to the these isolated even lines alone, which frequently leads to inaccurate edge surface estimate, particularly in stories that manages with the continuous scene swapping. A video that is ready for analysis while a series of images with spatial type of estimation (X, Y) with the transitory estimation T, the spatio-common that is ready for analysis while many images with estimation (X, T), or (Y, T). Here the kind of spatio-transient part is most similarly transmitted with estimate (X, T) of the entire type of video without the picture representation in accordance with criteria, for instance, is being compared to the identical y which are always being noted with the form of all set of video samples.





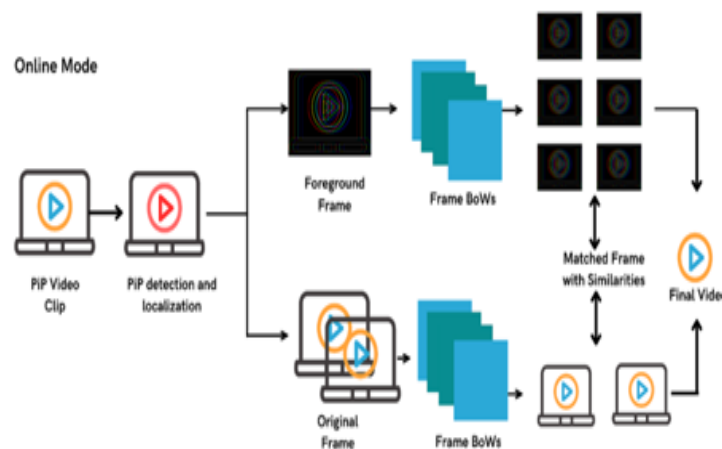
**Fig 6.** Model of Spatio Based on The Temporal Slice.

**Fig 6** outlines the spatially common incisions are ready to reveal the vertical lines. These vertical line types, which are in the frontal zone and establishing chronicles' edge surfaces, are employed concurrently by the lines according to the key where the housings while determining the likelihood at the management of edge surface.



**Fig 7.** Board Like Raspberry PI Utilized for Video Type of Identification.

The first eight headings are dropped in favour of the present two orientations, which lowers the calculation of point course by eliminating the vertical and level edge photos separately among the principal cautious edges shown in the figure, **Fig 7**. Currently, a truly small edge is chosen for survey rate's smart edge guarantee, increasing the likelihood of both genuine and disruptive edges. In order to select fewer edge lines for new frameworks, image extraction is then designed to combine tiny adjacent corners are being considered as parcel where the empty shot can be identified as the disengaging reject with the active edges.



**Fig 8.** Video Copy Identification System with The Help of Raspberry PI.

**Fig 8.** describes the expanded Raspberry Pi's employed video copy acknowledgment mechanism. Every request video recognises and limits (wherever distinct) the image. As of right now and going front, are eliminated from nearby see housings and unique those corners or the edges in the form of separately. In order to quickly scan the most relevant reference



diagrams surveyed using question plots, altered records are employed. Finally, the choice of the request file whereas to whether like a copied reference of video with it replicated will be directed by these linked reference edges. It will be assessed how well the video copy acknowledgment structure is being used.

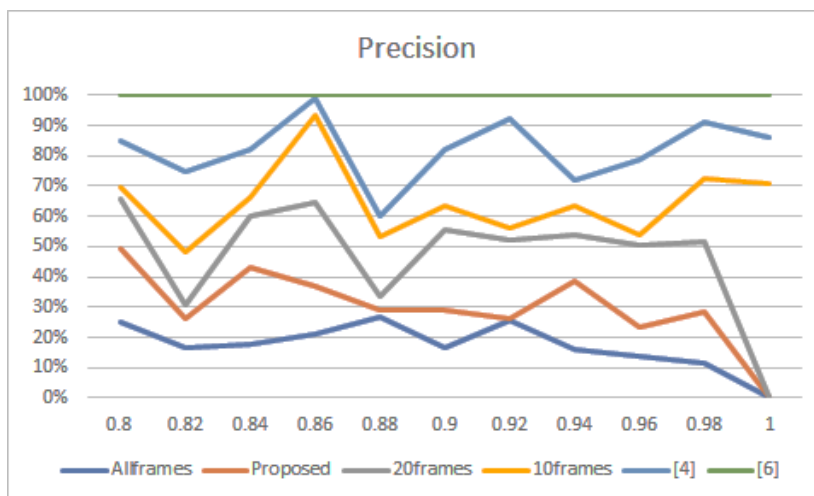


Fig 9. Precision Analysis.

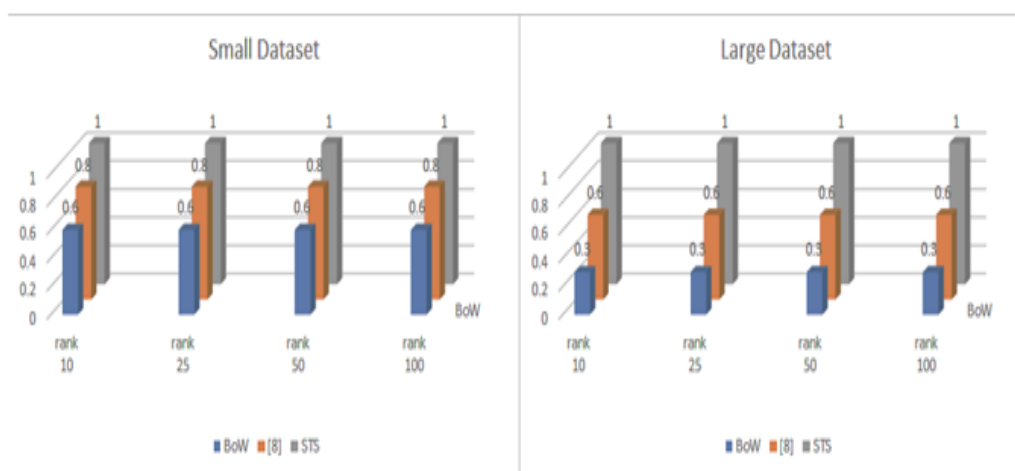


Fig 10. Copy Detection of The Video Retrieval for Accuracy from Small and Large Datasets.

The most wonderful execution to date with significantly increased time utilisation is the unmistakable proof type of execution with method for utilising all corners. Expected technique is depicted in **Fig 9** by the way of showing its precision, which is highlighted in light blue, and the incorrect restriction, which is highlighted in red. Here reason is due to because these types of square-shaped items are being harder to see. **Fig 10** displays the survey accuracy curves for various methods, along with the handling time that is evaluated using distinct image region tensile level; copying identifiable evidence throughout the whole video has improved precision and used less time. Additionally, increased precision and decreased time usage for video may be confirmed by edge line modification.

## V. CONCLUSION

In this work, the Internet of Things and Artificial Intelligence system components for the RaspberryPi board have been balanced. After release of the Raspberry Pi, vulnerabilities and mitigation strategies were brought to light. These potential outcomes from a reference can be used to check IoT systems. To address the emergence of the IoT security risk idea, it applies this basic reference model to the discovered IoT security risks. Finally, it has shown how the model may apply while forming the certified Internet of Things system by providing the secure customer assistance.

### CRedit Author Statement

The authors confirm contribution to the paper as follows:

**Conceptualization:** Karthika P and Balamurali S; **Methodology:** Balamurali S; **Software:** Karthika P; **Data Curation:** Karthika P and Balamurali S; **Writing- Original Draft Preparation:** Balamurali S; **Visualization:** Karthika P and Balamurali S; **Investigation:** Karthika P and Balamurali S; **Supervision:** Balamurali S; **Validation:** Balamurali S;

**Writing- Reviewing and Editing:** Karthika P and Balamurali S; All authors reviewed the results and approved the final version of the manuscript.

### Data Availability

No data was used to support this study.

### Conflicts of Interests

The author(s) declare(s) that they have no conflicts of interest.

### Funding

No funding agency is associated with this research.

### Competing Interests

There are no competing interests

### References

- [1]. M. M. Esmaceli, M. Fatourehchi, and R. K. Ward, "A Robust and Fast Video Copy Detection System Using Content-Based Fingerprinting," *IEEE Transactions on Information Forensics and Security*, vol. 6, no. 1, pp. 213–226, Mar. 2011, doi: 10.1109/tifs.2010.2097593.
- [2]. J. M. Barrios and B. Bustos, "Competitive content-based video copy detection using global descriptors," *Multimedia Tools and Applications*, vol. 62, no. 1, pp. 75–110, Nov. 2011, doi: 10.1007/s11042-011-0915-x.
- [3]. Jaap Haitisma and Ton Kalke. 2012. "A highly robust audio fingerprinting system". In *Proceedings of the International Symposium on Music Information Retrieval*. 107–115.
- [4]. M. Jiang, Y. Tian, and T. Huang, "Video Copy Detection Using a Soft Cascade of Multimodal Features," *2012 IEEE International Conference on Multimedia and Expo*, pp. 374–379, Jul. 2012, doi: 10.1109/icme.2012.189.
- [5]. Y. Lei, W. Luo, Y. Wang, and J. Huang, "Video Sequence Matching Based on the Invariance of Color Correlation," *IEEE Transactions on Circuits and Systems for Video Technology*, vol. 22, no. 9, pp. 1332–1343, Sep. 2012, doi: 10.1109/tcsvt.2012.2201670.
- [6]. H. Liu, H. Lu, and X. Xue, "A Segmentation and Graph-Based Video Sequence Matching Method for Video Copy Detection," *IEEE Transactions on Knowledge and Data Engineering*, vol. 25, no. 8, pp. 1706–1718, Aug. 2013, doi: 10.1109/tkde.2012.92.
- [7]. J. Song, Y. Yang, Z. Huang, H. T. Shen, and J. Luo, "Effective Multiple Feature Hashing for Large-Scale Near-Duplicate Video Retrieval," *IEEE Transactions on Multimedia*, vol. 15, no. 8, pp. 1997–2008, Dec. 2013, doi: 10.1109/tmm.2013.2271746.
- [8]. K. Taşdemir and A. E. Çetin, "Content-based video copy detection based on motion vectors estimated using a lower frame rate," *Signal, Image and Video Processing*, vol. 8, no. 6, pp. 1049–1057, Mar. 2014, doi: 10.1007/s11760-014-0627-6.
- [9]. P. Karthika and P. Vidhyasaraswathi, "Digital Video Copy Detection Using Steganography Frame Based Fusion Techniques," *Proceedings of the International Conference on ISMAC in Computational Vision and Bio-Engineering 2018 (ISMAC-CVB)*, pp. 61–68, 2019, doi: 10.1007/978-3-030-00665-5\_7.
- [10]. A. BenHajjoussef, T. Ezzedine, and A. Bouallègue, "Gradient-based pre-processing for intra prediction in High Efficiency Video Coding," *EURASIP Journal on Image and Video Processing*, vol. 2017, no. 1, Jan. 2017, doi: 10.1186/s13640-016-0159-9.
- [11]. P.-Y. Lin, B. You, and X. Lu, "Video exhibition with adjustable augmented reality system based on temporal psycho-visual modulation," *EURASIP Journal on Image and Video Processing*, vol. 2017, no. 1, Jan. 2017, doi: 10.1186/s13640-016-0160-3.
- [12]. I. Batioua, R. Benouini, K. Zenkour, and H. E. Fadili, "Image analysis using new set of separable two-dimensional discrete orthogonal moments based on Racah polynomials," *EURASIP Journal on Image and Video Processing*, vol. 2017, no. 1, Mar. 2017, doi: 10.1186/s13640-017-0172-7.
- [13]. B.-Y. Sung and C.-H. Lin, "A fast 3D scene reconstructing method using continuous video," *EURASIP Journal on Image and Video Processing*, vol. 2017, no. 1, Feb. 2017, doi: 10.1186/s13640-017-0168-3.
- [14]. Y. Cai, Y. Lu, S. H. Kim, L. Nocera, and C. Shahabi, "Querying geo-tagged videos for vision applications using spatial metadata," *EURASIP Journal on Image and Video Processing*, vol. 2017, no. 1, Feb. 2017, doi: 10.1186/s13640-017-0165-6.
- [15]. N. Nan and G. Liu, "Video Copy Detection Based on Path Merging and Query Content Prediction," *IEEE Transactions on Circuits and Systems for Video Technology*, vol. 25, no. 10, pp. 1682–1695, Oct. 2015, doi: 10.1109/tcsvt.2015.2395771.
- [16]. VidhyaSaraswathi.P and M. Venkatesulu, "A Secure Image Content Transmission using Discrete chaotic maps" *Jokull Journal*, Vol.63, No.9, pp.404-418, September-2013.
- [17]. Y. A. V. Phamila and R. Amutha, "Discrete Cosine Transform based fusion of multi-focus images for visual sensor networks," *Signal Processing*, vol. 95, pp. 161–170, Feb. 2014, doi: 10.1016/j.sigpro.2013.09.001.
- [18]. O. Prakash, R. Srivastava, and A. Khare, "Biorthogonal wavelet transform based image fusion using absolute maximum fusion rule," *2013 IEEE Conference on Information And Communication Technologies*, pp. 577–582, Apr. 2013, doi: 10.1109/cict.2013.6558161.
- [19]. K. Sharmila, S. Rajkumar, and V. Vijayarajan, "Hybrid method for multimodality medical image fusion using Discrete Wavelet Transform and Entropy concepts with quantitative analysis," *2013 International Conference on Communication and Signal Processing*, pp. 489–493, Apr. 2013, doi: 10.1109/iccsp.2013.6577102.
- [20]. L. Liu, H. Bian, and G. Shao, "An effective wavelet-based scheme for multi-focus image fusion," *2013 IEEE International Conference on Mechatronics and Automation*, pp. 1720–1725, Aug. 2013, doi: 10.1109/icma.2013.6618175.
- [21]. R. Mahmoud, T. Yousuf, F. Aloul, and I. Zuolkernan, "Internet of things (IoT) security: Current status, challenges and prospective measures," *2015 10th International Conference for Internet Technology and Secured Transactions (ICITST)*, pp. 336–341, Dec. 2015, doi: 10.1109/icitst.2015.7412116.