

Reconstruction of Images with RPCA and Augmenting 3D model

by

Tasnia Jasim Tahiti

18101482

Talha Yusuf

18101170

Chowdhury Abrar Shakhawat

18101497

A thesis submitted to the Department of Computer Science and Engineering
in partial fulfillment of the requirements for the degree of
B.Sc. in Computer Science

Department of Computer Science and Engineering
Brac University
January 2022

© 2022. Brac University
All rights reserved.

Declaration

It is hereby declared that

1. The thesis submitted is my/our own original work while completing degree at Brac University.
2. The thesis does not contain material previously published or written by a third party, except where this is appropriately cited through full and accurate referencing.
3. The thesis does not contain material which has been accepted, or submitted, for any other degree or diploma at a university or other institution.
4. We have acknowledged all main sources of help.

Student's Full Name & Signature:



Tasnia Jasim Tahiti
18101482



Talha Yusuf
18101170



Chowdhury Abrar Shakhawat
18101497

Approval

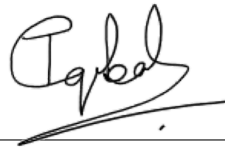
The thesis/project titled “Reconstruction of Images with RPCA and Augmenting 3D model” submitted by

1. Tasnia Jasim Tahiti (18101482)
2. Talha Yusuf(18101170)
3. Chowdhury Abrar Shakhawat (18101497)

Of Fall, 2021 has been accepted as satisfactory in partial fulfillment of the requirement for the degree of B.Sc. in Computer Science on January 20, 2022.

Examining Committee:

Supervisor:
(Member)



Dr. Muhammad Iqbal Hossain
Assistant Professor
Department of Computer Science and Engineering
Brac University

Program Coordinator:
(Member)

Dr. Md. Golam Rabiul Alam
Assistant Professor
Department of Computer Science and Engineering
Brac University

Head of Department:
(Chair)

Sadia Hamid Kazi, PhD
Chairperson and Associate Professor
Department of Computer Science and Engineering
Brac University

Abstract

Crime scene photography is one of the most essential parts of documentation for all physical crime scenes in the field of forensic science. It is an organized and structured operation, having a goal of helping the investigators with the hidden information about the place of crime and finally recognizing the criminal. Photography and videography of the crime scene is a key procedure as this will further help to analyze and find viable indications about the crime. This obtained footages lead to the idea of creating a 3D model of the crime site. By using sophisticated methods in computer vision, the output obtained can be realistic reconstructions. The article further goes over a few essential designs such as crucial elements that should be included during a sophisticated Crime Scene Investigation (CSI) study. Our proposal here is to generate a realistic formation of the crime site with the multiple sophisticated designs. This also concerns the estimation of image properties such as pixel orientation, frequency, noise level etc. Moreover, we present a set of 3D pictures that have been generated and compared to state-of-the-art view synthesis methods. Furthermore, we demonstrate why using RPCA provides us with the best outcome possible.

Keywords: Image Processing; Machine Learning; Crime scene; 3D; inpainting; PCA

Acknowledgement

First and foremost, we give thanks to the Almighty Allah for allowing us to finish our thesis without any serious setbacks.

Secondly, we would like to express our gratitude to our supervisor, Dr. Muhammad Iqbal Hossain sir, for his unwavering support and guidance throughout our project. He never failed to motivate us to study more. He assisted us with unlimited help whenever we needed it.

Thirdly, Mr. Sabbir Rahman sir, for his unwavering assistance with our paper. We collaborated on new research subjects in this sector. All of the feedback he provided was quite beneficial to us in our subsequent efforts.

Finally, without our parents' unwavering support, it may not be conceivable. We are currently on the verge of graduating thanks to their kind assistance and prayers.

Table of Contents

Declaration	i
Approval	ii
Abstract	iii
Acknowledgment	iv
Table of Contents	v
List of Figures	vii
List of Tables	viii
Nomenclature	ix
1 Introduction	1
1.1 Introduction	1
1.2 Problem Statement	2
1.3 Research Objectives	3
1.4 Thesis Structure	4
2 Background	5
2.1 Literature Review	5
2.1.1 Iterative Reconstruction	5
2.1.2 Image Completion (IC)	5
2.1.3 3D Reconstruction through learning-based in-painting model	7
2.2 Related Works	7
3 Methodology	8
3.1 Input data	9
3.1.1 Datasets	9
3.2 Data pre-processing	9
3.2.1 Preparing for processing	9
3.2.2 Measuring multiple variants	11
3.3 Subsampling	12
3.4 Image Completion	13
3.4.1 The matrix completion case	13
3.4.2 Image retrieval using direct Convolutional Neural Network . .	14
3.5 Noise Removal	15

3.5.1	Stabilization	15
3.5.2	Detail extraction	16
3.5.3	De-blurring	16
3.6	Edge Classification	17
3.7	3D Reconstruction	17
3.7.1	Layered depth image	17
3.7.2	Training the model	17
3.7.3	3D textured mesh conversion	17
4	Implementation and Results	18
4.1	Implementation of proposed method	18
4.1.1	Image reconstruction using RPCA	18
4.1.2	Data Training	20
4.1.3	3D Implementation	21
4.2	Results	22
4.3	Comparison with current related fields	24
5	Conclusion and Future work	26
5.1	Conclusion	26
5.2	Future work	26
	Bibliography	29

List of Figures

2.1	Feature extraction and max pooling	6
2.2	One convolutional layer with 32x32 image of 3 channels	6
3.1	The flow chart of the proposed RPCA image reconstruction and 3D implementation model	8
3.2	input image	10
3.3	grayscale image	10
3.4	contour plot	10
3.5	linear relationship between the x and y variable of 200 points	11
3.6	the contrast of the black background has increased after histogram equalization	11
3.7	Histogram before equalization	12
3.8	Histogram after equalization	12
3.9	(a) output of original object sequence, (b) output after being sub- sampled, (c) magnified image of (a), (d) magnified image of (b).[24]	13
3.10	stabilization	15
3.11	extraction	16
4.1	A noisy image	19
4.2	De-noised image using RPCA	19
4.3	Block diagram for PCA implementation	20
4.4	Model training procedure	21
4.5	inpainting process	22
4.6	before subsampling	22
4.7	after subsampling	23
4.8	Crime scene extracted from www.history.com	23
4.9	After executing the 3D inpainting	24
4.10	Missing pixel retriever with swath filter	24
4.11	Missing pixel retrieval using direct CNN with 20 epoch training . . .	25

List of Tables

3.1		11
3.2	Direct Convolution Networks Summary	15

Nomenclature

The next list describes several symbols & abbreviation that will be later used within the body of the document

ART Algebraic Reconstruction Technique

EM Expectation Maximization

MRF Markov Random Field

PCA Principal Component Analysis

RPCA Robust Principal component analysis

SSIM structural similarity

Chapter 1

Introduction

1.1 Introduction

Crime scene photography or forensic photography is one of the most essential techniques used in Investigation Department. It is an indispensable method in conducting investigation and collecting a permanent visual record of the crime scene that can be later utilized by Crime Investigation professionals during formal trials. It helps the professionals to demonstrate a clearer image of the crime scene by reconstructing the events that took place. This helps to link suspects to scenes and thus detect the criminal in lesser time and in much efficient manner. It is used in crime scenes like traffic collisions, burglaries, homicides etc. However, an unclear image can be misleading with inaccurate information and lack of detail. Moreover, many evidents become useless for missing some portion of it. This thesis will help build the lost parts of an image, with cleansing it with better details, making it easier to detect and match patterns, providing evidence at a faster rate.

The first use of forensic photography was in the nineteenth century in Belgium. In 1870s, it became an advanced technology and a significant stepping stone in the world of forensic science. The main objective of forensic photography is limited to capture evidence that is admissible in court. Where evidents cannot be collected physically, forensic photographs come in handy. Such photographs include body of the victim, bullet shell casings, broken glass, fingerprints, tire tracks, blood stains, bruise of victim etc. While detecting specific pattern or rebuilding crime scene, highly enhanced images are needed. Such image must include full 180° or 360° view in some cases. Image reconstruction through PCA and augmenting 3D model can help enhancing these forensic images providing accurate points, proportions, with bit rates up to less than 0.5 bit/pixel.

The method can improve the quality of noisy images where signal-noise ratio is very low. Desired object can be analyzed by taking 2D image as input. Image reconstructing and augmenting 2D picture might require investment in data storage and processing power, which can be easily handled by the rapid development of the related field. Moreover, optimum slice thickness and resolution can be achieved through the method. This will help matching profiles where extreme details are important, such as fingerprints, tire marks and so on.

Enhanced and augmented 3D images can come in handy when capturing evidence that possesses time restraints. For example, bloodstain evidence, which can determine the position of a victim at the time of a bloody blow, can often change over

time. However, generating a 3D model can recreate the crime scene to analyze and re-evaluate certain factors. This will help figuring out details that have been missed in initial attempts. Where some crucial part of an image is missing, for example, digit of a car plate, models like MRF (Markov Random Field) can be used to enable missing texture of an image by copying pixels of a target image itself.

1.2 Problem Statement

As the time passes by, photography is becoming one of the most integral part of documenting crime scenes, homicides, stealing, traffic collisions, etc. Images have become one of the most important pieces of the puzzle while solving crimes. These images give the investigators the inch perfect explanations of what has happened, and sometimes also giving the officers the perspective of how it happened. The main focus of crime scene photography is to provide the investigators the precise image of the crime site. Its' further focus is to preserve the physical evidence present on the scene.[18] This is a sensible and important factor of solving a case, and thus it requires the need of a skilled photographer. Forensic photography is an indispensable tool and an essential technique that is used in the field of forensic odontology which plays an important role in crime investigations.[26]

With the increasing popularity of image processing, more and more fields are taking interest in containing and developing their techniques through the help of it and merging into a secured database for further use. While working with digital images in investigations, professionals create a hash value for the CD-ROM that has been procured earlier. [18] A Chain of evidence (COE) procedure is used to document origin, usage, storage, and processing of evidence with its integrity clearly [26].

Crime scene photographs are captured from multiple viewpoints in order to have a wider-angle view of the environment. This is where augmentation of 3D image can play huge role. In many cases, evident gathering time is a constraint and limited. For such cases, creating a 3D model can give unlimited time of accessing the realistic crime scene. This can further help to relocate the varieties of viewpoints and collect any data that has been missed out due to shortage of time or limitation of access (both time wise or physically). Moreover, in times, wide angled images can help detecting many important aspects of the crime scene, such as, 3D image of bullet wound can help detecting the standing position of the shooter. The exact formation of multi-dimensional room can help visualizing the situation from any possible places for the forensic professionals.

Image processing covers vast areas of work such as reconstructing missing image data, removal of unnecessary objects, error concealment, image enlargement etc. [15]. Structural and textural reconstruction can be done to accurately rebuild image edges. In many investigation cases, such image processing techniques are necessary. For example, detecting missing car plate number from a still image, enlarging CCTV footage and minimize noise, building missing part of portrait sketches and so on. 3D model is referred to use in a case where building spatial structure of the whole environment is helpful to analyze in much more sophisticated way. [21] It helps to reconstruct a more hyper realistic environment and helps more than just still images or sketches.

Successful image reconstruction can be implemented through Adaptive Inverse Projection, which uses sparse projection. [15] Expectation Maximization (EM) and

Algebraic Reconstruction Technique (ART) can be utilized for iterative approach in reconstructing blurry image [3]. Iterative image restoration can be also conducted with ML-EM algorithm [14]. Other kinds of technique can restore any part of the evident based on fragment-based completion or copying pixel from a non-parametric sample image [15].

In many cases, crime scene photographs can be misleading and confusing for the viewers. Even shooting in the most ideal situation, taking both wide-angle shots and close-up shots can be challenging to develop the precise perspective. Many details can be missed or changed due to time constraint as well. In many cases artificial lighting and colored filters should be added. In order to find the evidences' relation with respect to other objects, certain techniques are needed. Curvature minimization is used to reconstruct images with the help of geometric analysis. [33] Both peak signal-to-noise ratio (PSNR) and structural similarity (SSIM) are used on synthetic and real image to estimate discrete mean curvature and Gaussian curvature. This can be further used in image segmentation and image super-resolution. Robust Principal Component Analysis (RPCA) is examined [25]. It can be used in vast areas of processing evidence such as hyper spectral image denoising, image completion, inpainting and so on. Many times, evidence is collected during the night time resulting in blurry and noisy images. In such case PCA can separate noisy low-rank components, as well as handle global illumination conditions that can be very helpful to servile night time images.

The growing significance of crime scene photography is due to the stretch of newer and more advanced technologies for the time being. In many cases, crime scene photographs have furnished desired explanation. Image reconstruction can help when mechanical aspects fall short in capturing clear and focused evidence. 3D model of the crime scene can accentuate panoramic images as well [21]. It can generate precise point clouds with a range of 120m.

Therefore, as stated, our approach can help in investigative procedures, maintenance of archival data though iterative image reconstruction, completion and augmentation of 3D model. SO, the question could be phrased as follows:

Can PCA and 3D image augmentation facilitate with proper transparency of an image?

This research will answer this through exploring different methods and algorithms to maximize the transparency of 2D images and provide 3D image with much higher quality in the fastest rate possible.

1.3 Research Objectives

This research helps to develop a 3D image from 2D image using iterative reconstruction. Moreover, this research establishes the development of the quality of the images. To make a crime scene more informative and helpful and to specify a crime suspect more accurately is the goal of this research. The main objectives of this research are:

1. To decrease the amount of noise in an image.
2. To recreate the 3D version of the crime scene from the 2D images.
3. To complete an incomplete image.
4. To reconstruct a damaged image.
5. To be able to make an image clearer.

1.4 Thesis Structure

Photographs is one of the main ways through which people gather information. People know about a place through different pictures, get to know about unknowns through photographs. As the days went by, the usage of photographs had spread to different work sectors and it has mostly helped in the sector of crime investigation. But the job is made hard when the image is distorted or damaged. Thus, comes the requirement of image retrieving. The research of ours would help to lessen the amount of noise in the picture and also helps to complete an incomplete image. The research also helped us to develop a better 3D image which will help to give the viewer more information about the image.

In our research, the data is firstly sent to pre-processing state where the data is processed through various steps to get ready for detection. Later in the processing state the noise and distortion of image is detected and removed with the help of RPCA algorithm. Further, the processed data is compared and then clustered which results in a sharpened image. Moreover, these pictures are gathered and an 3D image is formed by using the learning-based in painting model

Chapter 2

Background

2.1 Literature Review

Nowadays, people have a perspective from watching tv shows that how Crime Scene Photography is done. But many don't know, this is a useful step in solving the crime. The crime scene acts as a silent witness and this silent witness, indeed, proves to be very helpful in finding crucial evidences which are related to the story of the crime and allows detectives to establish the correct reasoning and therefore to gain a valuable insight in basically how the crime was committed [17], [2].

In other fields, such as Medical, the use of 3D imaging is rapidly increasing. For instance, Image reconstruction is widely used in computed tomography (CT) and also has been extensively studied from various angles, to search methods as well as algorithms for more accurate execution of the reconstruction tasks.[16] A detailed 3D image will help in knowing and understanding the crime scene better.

2.1.1 Iterative Reconstruction

Iterative reconstruction is basically a technique in image reconstruction which refers to iterative algorithms used to reconstruct 2D and 3D images in certain imaging techniques. In experiments, using with synthetic noise-free and additive noisy projection data of dental phantoms, after experiments it is found that both the simultaneous iterative algorithms produce superior image quality as well as compared to filtered back projection after linearly fitting projection gaps.[3]

For a wide range of scans, metal deletion technique yields reduced metal streak artifacts and better-quality images than does filtered back projection, linear interpolation, or selective algebraic reconstruction technique.[14] We are going to use this technique mainly as we are not sure when we are going to perform the experiment, if we will have a large set of projections available or not, and also when the projections are not equally distributed uniformly in an angle, or even when the projections are missing at different orientations.

2.1.2 Image Completion (IC)

The primarily goal of IC, is to complete the missing portion from the surrounding structure but first restore, in a non detectable way for any-one who is not aware of the original image.[20] While reconstructing an image it is not impossible that there

might be some missing portion of the image. Here comes image completion, and that is to generate multiple and diverse plausible results but only when presented with a masked image.[30]

Retrieving missing pixels using direct Neural Network

For retrieving missing part of an image, we have used the convolution process, which uses a multidimensional array of input data and a multidimensional array of parameters as a kernel to execute a finite summation over a finite number of array members. To demonstrate a convolution operation, we may use a two-dimensional kernel K and a two-dimensional input image I :

$$S(i, j) = (I * K)(i, j) = \sum_m \sum_n I(m, n) K(i - m, j - n) \quad (2.1)$$

By picking the values for m and n based on the input image's dimensionality, we can perform neural network on various images. It is not essential to perform linear operations on the entire image due to the convolution's architecture, but rather to choose specific parts of the image to condense and discover patterns.

Depending on the kernel we select, different patterns will be identified when the kernel travels around the picture. A 3×3 kernel, for example, with only one row of 1 and all other indices set to 0, would identify a short horizontal bar in the picture. And these patterns will be saved in a number of different limitless filters.

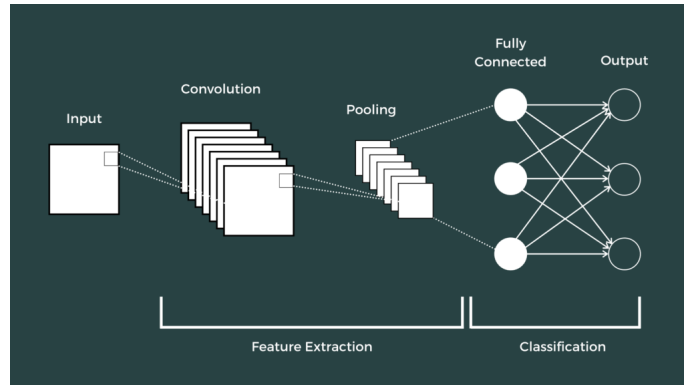


Figure 2.1: Feature extraction and max pooling

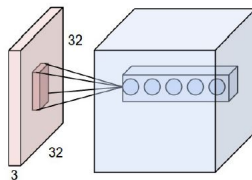


Figure 2.2: One convolutional layer with 32x32 image of 3 channels

2.1.3 3D Reconstruction through learning-based in-painting model

The 3D picture will be produced in a mesh by adding up the results, i.e., inpainted depth where there will be color values. There are numerous rendering representations available, and mesh is ideal in this situation. As a result of employing this, we can obtain results more quickly.

A Layered Depth Image with explicit pixel connection as the underlying representation and present a learning-based inpainting model can be used that repeatedly synthesizes fresh local color-and-depth content into the occluded region in a spatial context-aware fashion. Using conventional graphics engines, the resultant 3D images may be displayed effectively with motion parallax. Such multi-layer representation is very helpful to render the whole crime scene to visualize at any moment possible.

2.2 Related Works

This part aims to critically review previous relevant works in the field of image reconstruction and image processing in any context of crime investigation. We analyze the different techniques used for the main results achieved and we show how image reconstruction and 3D imaging has its own specific challenges due to the lack of information, the limited evidence and the massive number of errors and noise that makes the image to deteriorate.

Image reconstruction has elemental impacts on image quality. Most methods for image reconstruction aim at reconstructing the image or forming the whole image. To determine the depth of the image camera calibration is done. After the depth maps are obtained, using the image registration the final mesh is created.[1]

The problem during image reconstruction is shortage of information to recreate the image. The environment or the crime scene is one of the most valuable factors for the image reconstruction. It helps the investigators to establish their first thoughts on the case and an important sight of how the crime was done.[17] The author claims that by sequencing the video footages and by keyframing it the crime scene can be given a 3D look which results in the improvement of the judgement of the case. [21] The research work [8] shows us a different way to detect human and a different approach to image reconstruction. The application they used in the paper is outdoor surveillance. The advantage of reconstructing in this manner is that the desired result can also be obtained in outdoor uncontrolled environment. The goal of this paper is to observe human motion. However, the motion can be different in front of the camera, it is less dependent on the environmental factors. [6], [7]

Another research work shows that by reducing the curvatures on image surfaces we can reconstruct images [33]. Their claim is by integrating over Ω we get a more efficient way to reconstruct image. [4],[5]. Crime sketches and photographs of the crime area are one of the most well organized and competent way to solve a crime. [18] These are the roots to solving the case and it gives the most significant clues and measurements of the crime scene. Motion detection and reducing noise in an image has resulted in a better reconstruction of the image. Our approach extends the quality of image reconstruction when 3D augmentation is added to the process at the same time. Crime detection with the help of this processed image with reported result better and be more fruitful in solving the crime case.

Chapter 3

Methodology

The proposed system takes an image of a related object of interest and finds any defects or distortion in order to remove them and turn into fresh picture for detection and comparison purposes in the forensic lab. For this, input data needs to be first pre-processed; which is Grayscale Transform and Histogram Equalization in order to flatten grayscale histogram to create highest similar intensities. The processed data is then sent to be analyzed through RPCA and formed into a 3D structure for further clearer depiction.

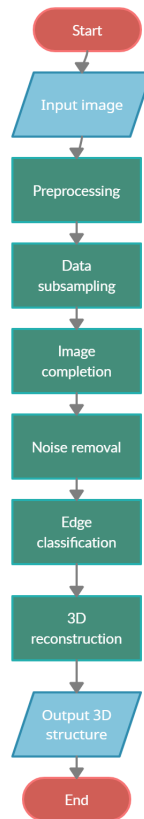


Figure 3.1: The flow chart of the proposed RPCA image reconstruction and 3D implementation model

So, the whole process can be described in the 3 stages below:

1. Data preprocessing: In this stage, the processing of input image takes place to

get ready for detecting noise and distortion.

2. Processing: In this stage, distortion and noise is removed through RPCA algorithm and redefined image is sent to Layered Depth In-painting and formed into a 3D image.

3. Comparison: In this stage, processed data is compared to previous data set in order to compare the noise removal and multi-layer representation of the object of interest.

After preprocessing, the image is clustered and the less noisy image is sent to remove any low frequency component and creates sharpened image containing the edges. Then the 3D picture is formed through learning-based in-painting model.

3.1 Input data

Huge amounts of data are now more prevalent in any IT sector, yet they might be challenging to analyze. Robust Principal component analysis (RPCA) helps to deduct the dimensionality of datasets. It improves readability while also reducing information loss. The system creates advanced uncorrelated variables and also maximizes variance. Thus, it reduces the dimensionality of a dataset, preserving as much ‘variability’ as possible.

3.1.1 Datasets

The CIFAR-10 database is a group of images commonly used to train machine learning and computer visualization algorithms. It is a chief dataset which is operated in machine learning. The dataset is made up of 60000 32*32 colour images which is divided into 10 classes. 50000 training images and 10000 test images are observed there. The database is divided into five training groups and one test set, each containing 10000 images. The test collection contains exactly 1000 randomly selected images for each class. Training collections contain random remaining images, but some training groups may contain more images of one class than another. Among them, training collections contain exactly 5,000 images from each class. Different classes have completely different characteristics. Each and every class have distinct images and none of the classes overlap.

3.2 Data pre-processing

Preprocessing dataset can fall into two categories: 1) Preparing data for further processing, 2) Measuring multiple variants of the image through plots.

3.2.1 Preparing for processing

If the input image is not grayscale or binary, it needs to be transformed into a 2D array grayscale image and thresholding is applied. Then the grayscale image’s contour is detected as Binarized images are easier to process when finding contours with this algorithm. All the contours shapes by calling shape contours in the console which will help to extract information from the image such as how many dots are present

[27]. PCA (principal component analysis) determines a list of the principal axes, which is then used to quantify relationship between x and y . The coefficients multiplying with each element are the principal components and serve as the dataset's low-dimensional representation. The transformed data will be reduced to a single dimension. What it does is, to mainly reconstruct different components which are in dataset by only attaching the first few of them together and it appreciates the optimal basis functions only [23]. For pixel wise representation, if the vector x is the collection of pixels in the dataset, then-

$$Image(x) = mean + x_1(basis\ 1) + x_2(basis\ 2) + x_3(basis\ 3)$$



Figure 3.2: input image



Figure 3.3: grayscale image

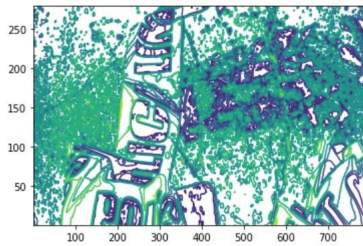


Figure 3.4: contour plot



Figure 3.5: linear relationship between the x and y variable of 200 points

PCA components of the above graph:

-0.94446029	-0.32862557
-0.32862557	0.94446029

Table 3.1

3.2.2 Measuring multiple variants

For noise filtering, total number of PCA component is needed. The number of components is then used in the inverse transformation in rebuilding the filtered image. Then plotting histogram and through equalizing the gray-scale histogram is flattened to create most similar intensities. Finally, CDF or Cumulative Distributive Function is used to map the pixels to a desired range. The comparison before equalization and after equalization can be shown below:



Figure 3.6: the contrast of the black background has increased after histogram equalization

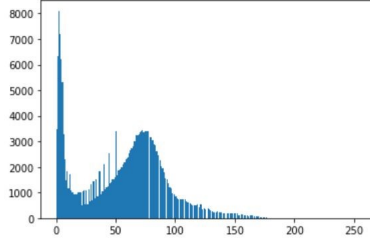


Figure 3.7: Histogram before equalization

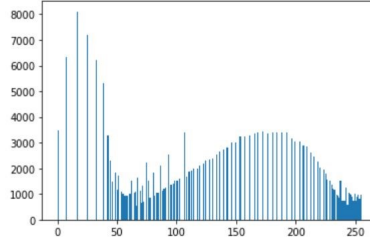


Figure 3.8: Histogram after equalization

3.3 Subsampling

According to Schroeder and martin subsampling is a procedure that decreases the size of the data by choosing a subset of the actual data [9]. According to Fabien, the main idea of image subsampling is creating a half-sized imaged image whilst throwing away every other row and column [29]. Subsampled images are sharper and has more clearer edges. If we observe the figure 3.9 carefully, the first image is blurry and lacks information. Whereas, in the second image, the photo is more informative and clearer. The alignments or the registered frames are more accurate, and thus the image produced results much more similar to our reference image. The final outcome image is less distorted which is achieved by making the energy better and efficient and the subsampled video is achieved.

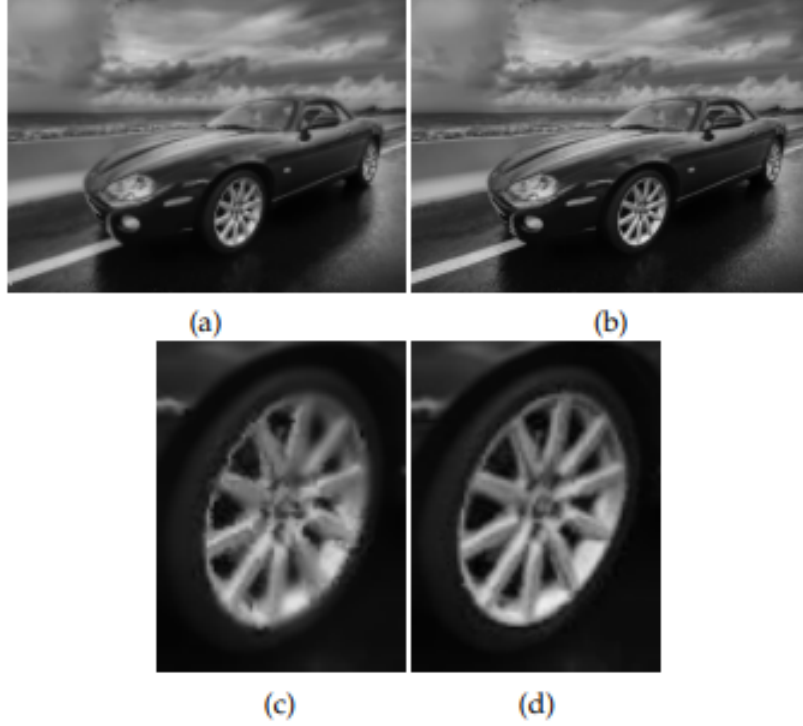


Figure 3.9: (a) output of original object sequence, (b) output after being subsampled, (c) magnified image of (a), (d) magnified image of (b).[24]

3.4 Image Completion

From section 2.1.2, we get to know that, the goal of Image completion is primarily to restore. The missing part reconstruction is mainly addressed by matrix completion (MC). First, we start with a matrix concept of the completion process in section 2.1.2. [25]

3.4.1 The matrix completion case

The main job of MC is, to just recover a low rank matrix and from its entries (partial observations). Basically, they are just distributed amongst two categories: Matrix completion by rank minimization and Matrix completion by matrix factorization [22]. We will discuss the first approach. The first approach or Matrix completion by rank minimization is to decrease the rank of input matrix to a negligible level. [25]

Matrix completion by rank minimization

Now, the main purpose here is a matrix given $L \in R^{m \times n}$ and find the missing portion with minimum rank which leads to the matrix $A \in R^{m \times n}$ (approximately) [25]. According to Candes and Recht, the problem is re-formulated as [10]:

$$\begin{aligned} & \underset{L}{\text{minimize}} && \text{rank}(L), \\ & \text{subject to} && P_{\Omega}(L) = P_{\Omega}(A) \end{aligned}$$

Here, In the equation (1) the values of P_{Ω} are denoting the operator which is limited to the materials of Ω (different entry which are given), i.e., The values which

are given in $\mathbf{A}$ are also same for the $P_{\Omega}(\mathbf{A})$ and for all the entry given in Ω and for the items outside of Ω and $\text{rank}(\mathbf{L})$ is indicating matrix $\mathbf{L}$'s rank [25]. A proposal was given by Candes and Recht to by the nuclear norm with the rank(.) [10]: $\min_{\mathbf{L}} \|\mathbf{L}\|_*$ Subject to $P_{\Omega}(\mathbf{L}) = P_{\Omega}(\mathbf{A})$ (2) Here, r is the rank of $\mathbf{L}$ and also $\|\mathbf{L}\|_* = \sum_{i=1}^r \sigma_i \mathbf{L}$ has some values here which are so, $\sigma_1, \sigma_2, \dots, \sigma_r$. These are mainly singular. It was theoretically proved by Candes and Recht that, the solution can be recovered (in exact) and also with a high probability [25] i.e., the nuclear norm is making the problem tractable. Also, to mention, the entries which are observed is expected to be enriched in noise in real world applications. So, we need to reduce the noise now, and which means making the problem (eqn. 2), a matrix completion approach which is relatively stable while relaxing the equality constraint was proposed [12]:

$$\underset{\mathbf{L}}{\text{Minimize}} \quad 1/2 \|\mathbf{P}_{\Omega}(\mathbf{L}) - \mathbf{P}_{\Omega}(\mathbf{A})\|_F^2 + \lambda \|\mathbf{L}\|_*$$

Here, the parameter λ is controlling the rank of $\mathbf{L}$ and $\|\cdot\|_F$ is denoting the Frobenious norm. The noise level should be the factor on which selection of λ should depend on [12].

3.4.2 Image retrieval using direct Convolutional Neural Network

The loss function for this approach is the mean squared error among the reconstructed awed car snap shots and the original pictures within the training database. Through many studies, efficient encoding of the information of the images by the strength of the power of convolution networks are confirmed. To construct the other remaining networks, convolution layers were chosen. For that reason, i would be able to find out the effective performance of direct convolution network. The enter records may be a combination of the three-channel masked schooling snap shots with a vicinity map indicating in which the mask are. The vicinity map is a matrix with 1s at the unmasked region and 0s at the masked part. adding the region map as the fourth channel of the input photos ought to supply the model a more specific target to work on. additionally, the detail of the networks I used is supplied inside the table underneath.

We have used keras for implementing direct convolutional network for retrieving the missing pixels. Total 9 convolution layers have been added fo every 32*32 images having RGB channel. The layers are conducted with relu activation function and output layer has sigmoid as activation function. Max pooling is done in all the layers.

Layers	Padding	Pool Layer	Output Shape
CONV1 + RELU1	same	input	32x32x32
CONV2 + RELU2	same	pool1	32x32x64
CONV3 + RELU3	same	pool2	32x32x128
CONV4 + RELU4	same	pool3	32x32x256
CONV5 + RELU5	same	pool4	32x32x256x512
CONV6 + RELU6	same	up6	32x32x128x256
CONV7 + RELU7	same	up7	32x32x64x128
CONV8 + RELU8	same	up8	32x32x32x64
CONV9 + RELU9	same	up9	32x32x32
Output + Sigmoid	same	-	32x32x3

Table 3.2: Direct Convolution Networks Summary

3.5 Noise Removal

Noise removal refers to fixing any level of distortion from a blurry image where pixels are spurious. Distortion can be of three types: 1) mildly distorted, 2) strongly distorted and 3) severely distorted. Depending on level of distortion, noise removal can be staged into three stages: 1) stabilization, 2) detail extraction and 3) de-blurring.

3.5.1 Stabilization

A severely distorted image influences the fusion of an image. Stabilization can be effective for sharp and mildly distorted images. Through stabilization, oscillations of a mildly distorted image can be reduced

Subsample of input image = $\{I^{\text{samp}}_j\}_{j=1}^{T_{\text{samp}}}$

And stabilized sequence of subsample = $\{I^{\text{stable}}_j\}_{j=1}^{T_{\text{samp}}}$

Each sample I^{samp}_j can be stabilized by mapping $D^{\text{r}}_{\text{ref}} : \Omega \rightarrow \Omega$ from I^{ref} to I^{samp}_j [where I^{ref} represents the reference image closest to actual image I]

Therefore, $I^{\text{stable}}_j = I^{\text{samp}}_j \circ D^j_{\text{ref}}$ If the deformation field is $\vec{V}^{\text{ref}}_1$, then- $D^{\text{ref}}_i(p) =$



Figure 3.10: stabilization

$$p + \vec{V}^{\text{reg}}_i(p) \forall p \in \Omega$$

Here the numbers which are zero or not, at first are remarked in the sparse matrix in this way RPCA identifies the attendance of general pattern Then centroid method is

$$\text{implemented. } \vec{V}^{\text{stable}}_i = \begin{cases} \frac{1}{T_{\text{samp}}} \sum_{j=1}^{T_{\text{samp}}} L^j_{i,p^z} & \text{if } \|S^*_{i,p}\|_0 < \frac{XYT_{\text{samp}}}{2} \\ \frac{1}{T_{\text{samp}}} \sum_{j=1}^{T_{\text{samp}}} L^j_{i,p^z} & \text{Otherwise} \end{cases}$$

If I_i^{samp} is a frame that contains deformation, then I_i^{samp} is deformed in accordance with $\vec{V}_i^{\text{stable}}$, $I_i^{\text{stable}} = I_i^{\text{samp}} \vec{V}_i^{\text{stable}}$, $1 \leq I \leq T_{\text{samp}}$. Then the stabilized image $\{I_i^{\text{stable}}\}_{i=1}^{T_{\text{samp}}}$ is obtained.

3.5.2 Detail extraction

In some distorted images, we tend to extract some specific portion of it to gain highest texture in the output. Detail extraction works in handy for such scenario. In this, firstly a detail layer is formed where weights are produced by smoothing corresponding binary patch of the subsample.

Weight, $W^i = H_G(\Omega^{iB}, \hat{\Omega}^i)$ [H_G is guided filter]

Then adding the weights, we get the detail layer- $L^D = \sum_{k=1}^M W^i \Omega^i$

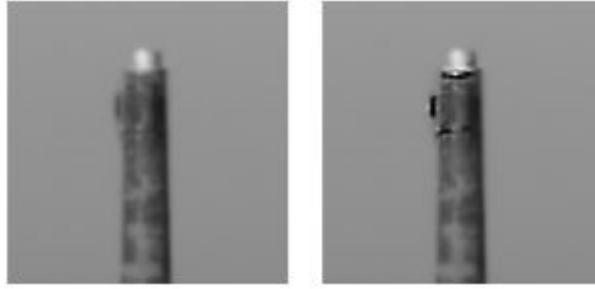


Figure 3.11: extraction

3.5.3 De-blurring

After detail extraction, de-blurring is needed to remove blurring artifacts (overly sharpened layer details). De-blurring is implemented on low rank parts of the image. If I^{LR} is a low rank part, then it is absorbed with texture detail layer.

$$I^{\text{final}} = \text{deblur}(I^{\text{LR}}) + \beta L^D$$

If

G = a blurred image

F = latent sharp image

h = blur kernel,

$$G = F \otimes h + n$$

Then to restrain F and H , regularization terms R_f and R_h are used, $R_f(F) =$

$$\|\rho(Fx) + \rho(Fy)\|_1 \quad \left[\text{where } \rho(x) = \begin{cases} -\theta_1|x| & x \leq l_t \\ -(\theta_2 x^2 + \theta_3) & x > l_t \end{cases} \right] \quad R_h(h) = \|h\|_1$$

Through L+S decomposition, observed image can be viewed as the sum of a low-rank clean image which then can be transformed into stabilized image; that further is formed into detail extraction and de-blurred for noise removal and have a clear vision of any detail needed in the field.

3.6 Edge Classification

Edges are the boundaries between objects and the background. High pass filtering is a process that tends to keep the high frequency data within an image while decreasing the low frequency information [34]. High pass filter enlarges the brightness of the center pixel [34]. The image is sharpened through High pass filtering.

Image processing is an integral part to ensure the quality of the image [13]. Noise reduction in image processing helps vastly during analysis or usage of the images. Image noise reduces the precision of image details that can be recognized by OCR (optical character recognition) software [13]. Blurred edges results in the deterioration of the information.

3.7 3D Reconstruction

3.7.1 Layered depth image

For this method to work, a RGBD image needs to be used. RGBD which is a depth image pair with aligned color, as an input and an image is generated (LDIs, [11]) with inpainted color. Some of the details of such an image is it can hold multiple pixels, from zero to several. Here, the value stored of LDI pixel by each color and value of depth. Here, the local connectivity of pixels is represented. We explicitly represent, each pixel stores tips that could either push the neighbor in any of the Compass directions (east, west, north, south) or zero. The pixels of LDIs can be defined as normal image pixels (4-connected like) which are situated within smooth regions (no neighbors available across depth discontinuities.). For 3D photography, LDIs can be a useful representation for some reasons:

- (1) an arbitrary number of layers can be handled naturally by them, and
- (2) also, thinly scattered, (it is efficient in handling memory and also converted in a textured mesh whose weight is negligible and it can also render fast.) [32]

3.7.2 Training the model

Secondly, to construct the layered image, the need of training the model arises. The hyper-parameters are used here in order to classify edge and generate them. [28]. This hyper-parameter can be used with ADAM optimizer [31] which also has a value of $\beta = 0.9$. One of the mentionable things is that it has a learning rate consisting of 0.0001 that too initially. So, it can be said that this is a depth-generator model. Now, while training the model, it is required to train both of them (depth generator and edge) by giving images as input in the code and then we can find desired output.

3.7.3 3D textured mesh conversion

Lastly, for the final product the 3D constructed image will be printed in a mesh by adding up the results i.e., inpainted depth where there will be color values. For rendering, there are many representations and amongst them mesh is perfect for this case. As by using this we can gain the results faster.

Chapter 4

Implementation and Results

The projected system is implemented in laboratory platform in an ipynb file. Implementation has total 2 parts- 1) image reconstructing using RPCA and 2) 3D implementation. A noisy image file containing is used and subsampled to compute PCA algorithm and filter noise from it in lower-dimensional representation. Further, 2D image is transformed into 3D mesh through inpainting method.

4.1 Implementation of proposed method

For reconstructing and 3D formation, our proposed model contains 2 files. These two files are described below:

1. `pca_image_reconstruction.ipynb`: This ipynb file implements RPCA and computes principle components and variance. It takes an image as input and preserves 50 percent of the variance and use inverse of the transform to construct de-noised image

2. `3d-photo-inpainting.ipynb`: This ipynb file implements 3D and computes the 3D depth inpainting. It takes an image as input merges all the synthesized pixels to give us a 3D picture.

3. `pixel_retriever.ipynb`: The ipynb implements direct neural network on Cifar10 dataset and retrieves missing pixels from the images

Here, ipynb file has been used as it is easier to implement Python algorithms and preserves any changes that have been applied in any cell of the code. The data axes are transferred into principal axes through affine transformation. A simple noisy image is used as input that contains blurry pixels. By preserving a specific percentage of variance, the number of principle components is determined and detected. Then using sklearn sub module from PCA module, inverse transform is used to filter the noise.

4.1.1 Image reconstruction using RPCA

Dataset Processing

Input consists of blurry and noisy image that is transferred into greyscale image in order to Pixel segmentation is done by using Simple Linear Iterative Clustering (CLIC). It takes the pixel values of the image separates them into a predefined number of sub-regions or sub samples.

RPCA algorithm implementation

PCA algorithm for noise filtering

```
noisy ← numpy.random.normal(digits.data,4)
pca ← PCA(0.50).fit(noisy)
pca.n_components_
components ← pca.transform(noisy)
filtered ← pca.inverse_transform(components)
```

RPCA (Robust Principle Component Analysis) is implanted on the dataset of turbulence -distorted sequence of digits. The method is implemented on ipynb file where sklearn. decomposition submodule has been used. This sub-module has RPCA variants such as RandomizedPCA and SparsePCA. These two components help in approximating principle components and enforce sparsity on the components. Variance larger than noise does not affect the input image. Then the data has been reconstructed using the largest subset of principal components. The input digits



Figure 4.1: A noisy image

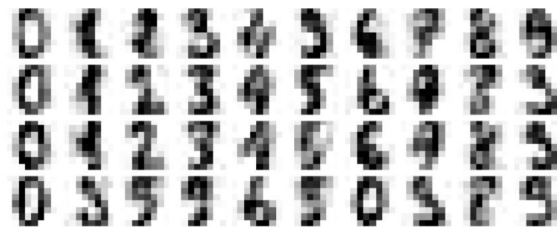


Figure 4.2: De-noised image using RPCA

here contain spurious pixels. Then components have been trained such as projection preserve 50 percent of the variance. 50 percent variance results into total of 12 principle components in the principal axes. Finally inverse transformation has been used to reconstruct image with less noise.

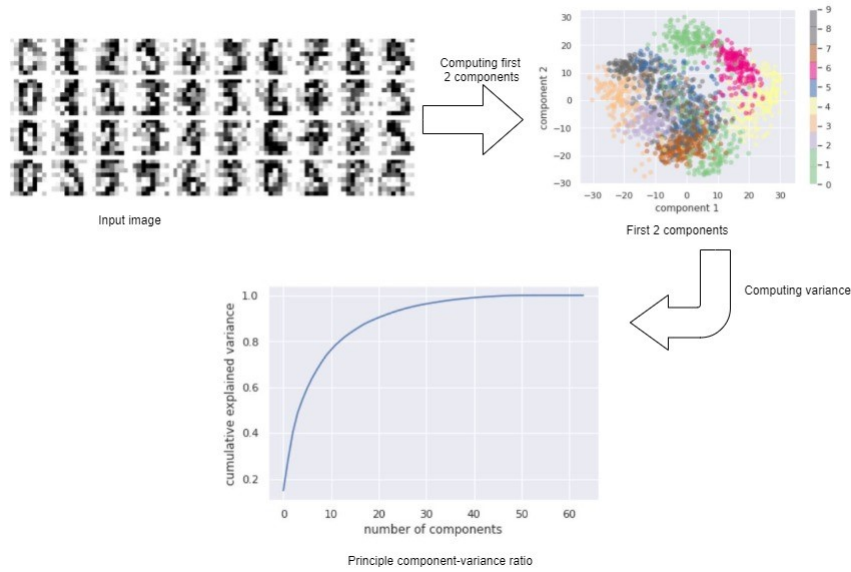


Figure 4.3: Block diagram for PCA implementation

4.1.2 Data Training

We split our data into a train and test set using keras library to load data from Cifar10 dataset. Apple 60000 images of Cifar10, 50000 images are for training samples and 10000 is for test samples. Then using CV2 interface random pixels of the images are being masked with the help of wandb tool. Random thickness of wandb tool and random pixel of XY coordinate from the input image has been chosen. Then cv2.line has been used to concatenate or conduct bitwise and operation for creating random mask on images. The inpainting algorithm is as below:

Inpainting with wandb Algorithm

```

for each epoch: do
    labels = []
    images = []
    predictions = []
    sample_idx → 0 ~ len(testgen)
    sample_images, sample_labels → testgen [sample_idx]
    for i in range (32): do
        inpainted_image self.model.predict(np.exand_dims(sample_images[i],axis=0)
        sample_images[i] → images
        sample_labels[i] → labels
        inpainted_imagereshape(inpainted_image.shape[1:]) → predictions
        wandb.log({"images": [wandb.Image(image) for image in images]})
        wandb.log({"labels": [wandb.Image(label) for label in labels]})
        wandb.log({"predictions": [wandb.Image(inpainted_image) for inpainted_image
        in predictions]})

```

Then the model has been trained with fit method. Training has been done for 20 epochs and for each epoch there are steps equal to the length of traingen, which is a numpy array. The array is an augmentation of original image and masked image. 20 epochs training takes around 55 seconds.

```

Train for 1562 steps, validate for 312 steps
Epoch 1/20
1562/1562 [=====] - 47s 30ms/step - loss: 0.0127 - dice_coef: 0.6042 - val_loss: 0.0124 - val_dice_coef: 0.6059
Epoch 2/20
1562/1562 [=====] - 50s 32ms/step - loss: 0.0127 - dice_coef: 0.6042 - val_loss: 0.0128 - val_dice_coef: 0.6058
Epoch 3/20
1562/1562 [=====] - 52s 33ms/step - loss: 0.0125 - dice_coef: 0.6042 - val_loss: 0.0139 - val_dice_coef: 0.6038
Epoch 4/20
1562/1562 [=====] - 51s 32ms/step - loss: 0.0126 - dice_coef: 0.6043 - val_loss: 0.0121 - val_dice_coef: 0.6074
Epoch 5/20
1562/1562 [=====] - 52s 33ms/step - loss: 0.0126 - dice_coef: 0.6042 - val_loss: 0.0115 - val_dice_coef: 0.6058
Epoch 6/20
1562/1562 [=====] - 51s 33ms/step - loss: 0.0123 - dice_coef: 0.6043 - val_loss: 0.0139 - val_dice_coef: 0.6039
Epoch 7/20
1562/1562 [=====] - 52s 33ms/step - loss: 0.0124 - dice_coef: 0.6043 - val_loss: 0.0118 - val_dice_coef: 0.6071
Epoch 8/20
1562/1562 [=====] - 52s 34ms/step - loss: 0.0120 - dice_coef: 0.6043 - val_loss: 0.0137 - val_dice_coef: 0.6038
Epoch 9/20
1562/1562 [=====] - 52s 33ms/step - loss: 0.0122 - dice_coef: 0.6043 - val_loss: 0.0139 - val_dice_coef: 0.6085
Epoch 10/20
1562/1562 [=====] - 53s 34ms/step - loss: 0.0120 - dice_coef: 0.6043 - val_loss: 0.0107 - val_dice_coef: 0.6066
Epoch 11/20
1562/1562 [=====] - 52s 33ms/step - loss: 0.0119 - dice_coef: 0.6043 - val_loss: 0.0113 - val_dice_coef: 0.6049
Epoch 12/20
1562/1562 [=====] - 51s 33ms/step - loss: 0.0120 - dice_coef: 0.6043 - val_loss: 0.0115 - val_dice_coef: 0.6072
Epoch 13/20
1562/1562 [=====] - 53s 34ms/step - loss: 0.0119 - dice_coef: 0.6043 - val_loss: 0.0124 - val_dice_coef: 0.6043
Epoch 14/20
1562/1562 [=====] - 51s 33ms/step - loss: 0.0116 - dice_coef: 0.6043 - val_loss: 0.0126 - val_dice_coef: 0.6064
Epoch 15/20
1562/1562 [=====] - 52s 33ms/step - loss: 0.0119 - dice_coef: 0.6043 - val_loss: 0.0114 - val_dice_coef: 0.6066
Epoch 16/20
1562/1562 [=====] - 52s 33ms/step - loss: 0.0116 - dice_coef: 0.6044 - val_loss: 0.0127 - val_dice_coef: 0.6094
Epoch 17/20
1562/1562 [=====] - 53s 34ms/step - loss: 0.0116 - dice_coef: 0.6044 - val_loss: 0.0109 - val_dice_coef: 0.6077
Epoch 18/20
1562/1562 [=====] - 52s 33ms/step - loss: 0.0116 - dice_coef: 0.6044 - val_loss: 0.0122 - val_dice_coef: 0.6063
Epoch 19/20
1562/1562 [=====] - 54s 35ms/step - loss: 0.0115 - dice_coef: 0.6044 - val_loss: 0.0106 - val_dice_coef: 0.6069
Epoch 20/20
1562/1562 [=====] - 51s 33ms/step - loss: 0.0115 - dice_coef: 0.6044 - val_loss: 0.0107 - val_dice_coef: 0.6069

```

Figure 4.4: Model training procedure

Create Mask Algorithm

```

for i in images: do
    img=images[i].copy()
    mask=np.full(32x32x3),255,np.uint8
    for i in range (random 1~10): do
        X1,X2 = random 1~32
        Y1,Y2 = random 1~32
        thickness = random 1~3
        CV2.line(mask,(X1,Y1),(X2,Y2),(1,1,1),thickness)
    masked_image=CV2.bitwise_and(img,mask)

```

4.1.3 3D Implementation

LDI inpainting Algorithm

For running the program on Google Colab, there are some prerequisites. Firstly, we need to prepare the environment at python (installing torch, opencv-python, vispy, moviepy, transforms3d, network etc.). Secondly, we need to download script and also the pretrained model of the 3d-photo-inpainting. Lastly uploading the .jpeg file and executing the 3d photo in painting to get desired results.

From section 3.7.1, after we take an image, the next processes are discussed below. In fig.(a) Firstly, we see a fully connected LDI. In gray, the depth edge is marked which is discontinuity. In fig.(b) Secondly, a foreground silhouette (color: green) will be formed also a background silhouette (red) after the pixel connections been divided in fig 4.4 it is seen that it attacks depth. In fig.(c) Now, region (berry colored) with a region that is in synthesis (color in pinkish red) is spawned for the

background silhouette of new LDI pixels. In fig.(d) Lastly, in the final LDI product, all the mentioned synthesized pixels have been merged and we get our inpainted product. [32]

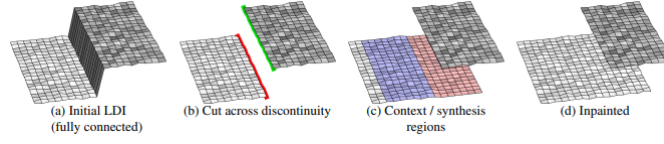


Figure 4.5: inpainting process

4.2 Results

The figures given below show us the results after all the process done above. By using Google COLAB, the results are obtained.

The vectors in figure.4.5, are representing the principal axes of the data, and the it is also to be mentioned, while describing the distribution of the data length of the vector is an indication of how "important" that axis is. Before 3D construction we have used rPCA to reconstruct and while using the model, it is in figure 4.5. A valuable feature selection routine is PCA's signal noise filtering quality—for example, if we train a classifier on a lower-dimensional representation rather than high-dimensional data, the results will be such that the noise will be filtered automatically remaining around the inputs. In figure 4.5, a 64-dimensional point

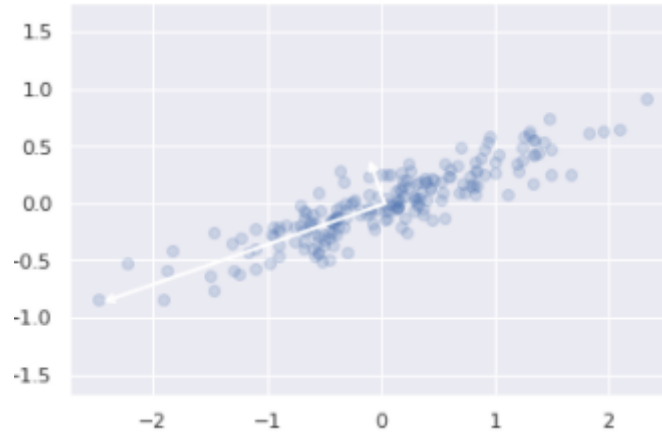


Figure 4.6: before subsampling

cloud is visible, in the points, the largest variance is the projection of different data. Principally, the rotation in 64-dimensional space and the optimal stretch have been found which allows to check out the layout of the digits in different dimensions, did without reference to the labels.[23]

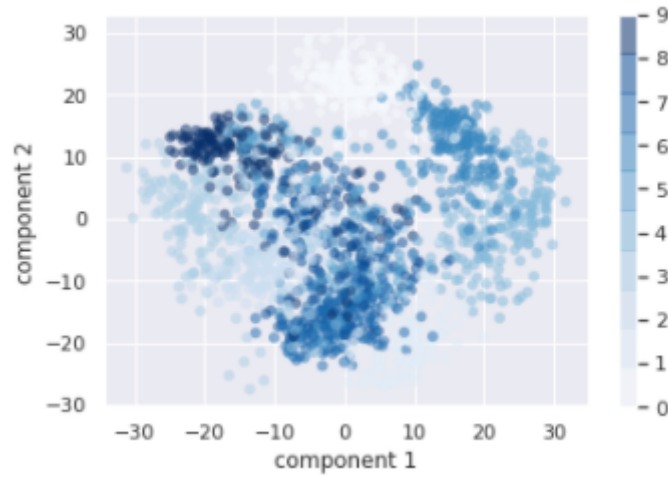


Figure 4.7: after subsampling

Again, a random image of a crime scene is selected in figure 4.10, where the 3D inpainting is executed. Now, starting with Camera calibration (for depth). Multiple depth maps are obtained. After getting them, a final mash is created which is projected in figure 4.11 As the scenes where not visible properly, multiple frames



Figure 4.8: Crime scene extracted from www.history.com

are given here. The original 3D rendered file, figure.4.8 can be viewed at this [link](#).



Figure 4.9: After executing the 3D inpainting

4.3 Comparison with current related fields

There are multiple inpainting methods and in-built libraries for retrieving missing pixels and denoising images. Our method is more efficient, fast and more accurate compared to such state of the art methods. For example, missing-pixel-filler python package is used to retrieve missing pixels of an image by filling missing parts with non-null neighboring pixels. By using swath filling, this type of result can be achieved:

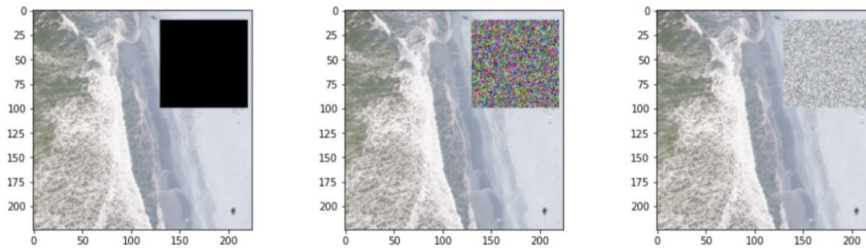


Figure 4.10: Missing pixel retriever with swath filter

On the contrary, our direct neural network with keras achieves more accuracy with much less complications and easier structure:

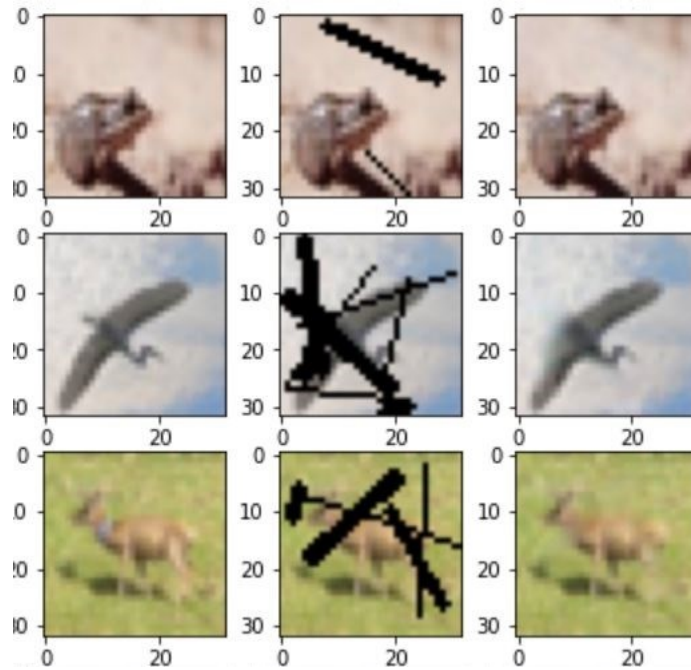


Figure 4.11: Missing pixel retrieval using direct CNN with 20 epoch training

Chapter 5

Conclusion and Future work

5.1 Conclusion

The conditions of the crime scene, the types of the evidence, differ drastically from one to the next. To successfully capture evidence, one needs to have expertise and understanding in both photography and forensics, as well as eager to experiment with lighting, solve problems, and be inventive.[19] Investigation departments are still following the same traditional way of analyzing crime scenes. But truly, the crime scenes are the only witnesses if no other witness were present. As a result, analyzing them with a full effort is a need. There has been a shortage of research that how can this work be done more effectively. Thus, this puts up an idea that crime scenes can be analyzed more perfectly using 3D image reconstruction of the scenes. And this research makes an attempt to implement this idea by using the ideas mentioned above. From the concurrent methods our method RPCA is relatively faster.

In this paper, we used machine learning algorithm PCA(Principal Component Analysis) and a deep learning algorithm namely CNN(Convolutional Neural Network). These helped us to reduce the dimensionality of large data sets and to differentiate different objects. We successfully retrieved a damaged image using the listed methods. Moreover, we used enlist to gather a numeric dataset.

5.2 Future work

Currently, we looked into different types of pictures and was able to retrieve them to a certain extent. Though the retrieval of number plates from different vehicles were seen near perfectly after the retrieval, some of the landscape images lacked that proficiency. We would further like to work more on this details of the landscape to make it more perfect and give the viewer more opportunity to get more information from the image.

Bibliography

- [1] L. G. Brown, “A survey of image registration techniques,” vol. 24, no. 4, 1992, ISSN: 0360-0300. DOI: 10.1145/146370.146374. [Online]. Available: <https://doi.org/10.1145/146370.146374>.
- [2] *Criminalistics Handbook*. U.S. Army Military, 1993.
- [3] G. Wang, D. Snyder, J. O’Sullivan, and M. Vannier, “Iterative deblurring for ct metal artifact reduction,” *IEEE Transactions on Medical Imaging*, vol. 15, no. 5, pp. 657–664, 1996. DOI: 10.1109/42.538943.
- [4] N. Sochen, R. Kimmel, and R. Malladi, “A general framework for low level vision,” *IEEE transactions on image processing : a publication of the IEEE Signal Processing Society*, vol. 7, pp. 310–8, Feb. 1998. DOI: 10.1109/83.661181.
- [5] A. Yezzi, “Modified curvature motion for image smoothing and enhancement,” *IEEE transactions on image processing : a publication of the IEEE Signal Processing Society*, vol. 7, pp. 345–52, Feb. 1998. DOI: 10.1109/83.661184.
- [6] J. Davis and A. Bobick, “The representation and recognition of action using temporal templates,” Aug. 2000.
- [7] R. Fablet and M. Black, “Automatic detection and tracking of human motion with a view-based representation,” Jul. 2002, ISBN: 978-3-540-43745-1. DOI: 10.1007/3-540-47969-4_32.
- [8] H. Kjellström, “Detecting human motion with support vector machines,” vol. 2, Sep. 2004, 188–191 Vol.2, ISBN: 0-7695-2128-2. DOI: 10.1109/ICPR.2004.1334092.
- [9] W. Schroeder, “Overview of visualization,” in. Dec. 2005, pp. 3–, ISBN: 978-0123875822. DOI: 10.1016/B978-012387582-2/50003-4.
- [10] E. Candès and B. Recht, “Exact matrix completion via convex optimization,” *Communications of the ACM*, vol. 9, pp. 717–772, Nov. 2008. DOI: 10.1007/s10208-009-9045-5.
- [11] Y. Pritch, E. Kav-Venaki, and S. Peleg, “Shift-map image editing,” Nov. 2009, pp. 151–158. DOI: 10.1109/ICCV.2009.5459159.
- [12] E. Candes and Y. Plan, “Matrix completion with noise,” *Proceedings of the IEEE*, vol. 98, pp. 925–936, Jul. 2010. DOI: 10.1109/JPROC.2009.2035722.
- [13] C.-Y. Hsu, H.-Y. Huang, and L.-T. Lee, “An interactive procedure to preserve the desired edges during the image processing of noise reduction,” *EURASIP J. Adv. Sig. Proc.*, vol. 2010, Dec. 2010. DOI: 10.1155/2010/923748.

- [14] F. E. Boas and D. Fleischmann, "Evaluation of two iterative techniques for reducing metal artifacts in computed tomography," *Radiology*, vol. 259 3, pp. 894–902, 2011.
- [15] T. Ogawa and M. Haseyama, "Missing image data reconstruction based on adaptive inverse projection via sparse representation," *Multimedia, IEEE Transactions on*, vol. 13, pp. 974–992, Nov. 2011. DOI: 10.1109/TMM.2011.2161760.
- [16] M. Soyapi, M. Yusoff, R. Sulaiman, and K. Shafinah, "Image reconstruction for ct scanner by using filtered back projection approach," vol. 88, pp. 797–803, Jan. 2012.
- [17] *Criminalistics Course Book*. Turkish Gendarmerie Schools Command, 2013.
- [18] R. Rohatgi and A. Kapoor, "Importance of still photography at scene of crime: A forensic vs. judicial perspective," *Journal of Harmonized Research ISSN: 2321- 7456*, vol. 11, pp. 271–274, Oct. 2014.
- [19] S. Staggs, *Crime Scene and Evidence Photography (2nd Edition)*. Staggs Publishing, 2014.
- [20] S. Zarif, I. Faye, and D. Awang Rambli, "Image completion: Survey and comparative study," *International Journal of Pattern Recognition and Artificial Intelligence*, vol. 29, p. 1 554 001, Dec. 2014. DOI: 10.1142/S0218001415540014.
- [21] G. E. Bostanci, "3d reconstruction of crime scenes and design considerations for an interactive investigation tool," Dec. 2015.
- [22] R. Sun and Z.-Q. Luo, "Guaranteed matrix completion via non-convex factorization," *IEEE Transactions on Information Theory*, vol. 62, pp. 6535–6579, Nov. 2016. DOI: 10.1109/TIT.2016.2598574.
- [23] J. T. Vanderplas, *Python Data Science Handbook: Tools and techniques for developers*. O'Reilly, 2016.
- [24] C. P. Lau, Y. Lai, and L. M. Lui, "Restoration of atmospheric turbulence-distorted images via rpca and quasiconformal maps," *Inverse Problems*, vol. 35, Apr. 2017. DOI: 10.1088/1361-6420/ab0e4b.
- [25] A. Sobral, "Robust low-rank and sparse decomposition for moving object detection: From matrices to tensors," Ph.D. dissertation, May 2017. DOI: 10.13140/RG.2.2.33578.82884.
- [26] S. Gouse, S. Karnam, G. H C, and S. Murgod, "Forensic photography: Prospect through the lens," *Journal of Forensic Dental Sciences*, vol. 10, p. 2, Jan. 2018. DOI: 10.4103/jfo.jfds_2_16.
- [27] C. Liu, J. Yang, D. Ceylan, E. Yumer, and Y. Furukawa, "Planenet: Piece-wise planar reconstruction from a single rgb image," Apr. 2018.
- [28] T. Miyato, T. Kataoka, M. Koyama, and Y. Yoshida, "Spectral normalization for generative adversarial networks," Feb. 2018.
- [29] M. Fabien, *Image subsampling and downsampling*, Mar. 2019. [Online]. Available: https://maelfabien.github.io/computervision/cv_3/#.
- [30] C. Zheng, T.-J. Cham, and J. Cai, "Pluralistic image completion," Jun. 2019, pp. 1438–1447. DOI: 10.1109/CVPR.2019.00153.

- [31] R. Ranftl, K. Lasinger, D. Hafner, and V. Koltun, “Towards robust monocular depth estimation: Mixing datasets for zero-shot cross-dataset transfer,” *IEEE Transactions on Pattern Analysis and Machine Intelligence*, vol. PP, pp. 1–1, Aug. 2020. DOI: 10.1109/TPAMI.2020.3019967.
- [32] M.-L. Shih, S.-Y. Su, J. Kopf, and J.-B. Huang, “3d photography using context-aware layered depth inpainting,” Jun. 2020, pp. 8025–8035. DOI: 10.1109/CVPR42600.2020.00805.
- [33] Z. Qiuxiang, K. Yin, and Y. Duan, “Image reconstruction by minimizing curvatures on image surface,” *Journal of Mathematical Imaging and Vision*, vol. 63, Jan. 2021. DOI: 10.1007/s10851-020-00992-3.
- [34] *High pass filtering*. [Online]. Available: <https://www.l3harrisgeospatial.com/docs/highpassfilter.html>.