

- **Title:** Summarizes the main idea of your project. Clever/catchy acronyms and/or puns are welcome. Deep thoughts also welcome.

Title Ideas:

Low Light Image Enhancement

- **Who:** Names and logins of all your group members. Take credit for your awesome work!

Group Details:

Jeremy Lutz: jlutz

Juan Mina: jmina

Jose Celaya-Alcala:

Truong Cai: tcai7

- **Introduction**

Modern digital cameras have issues producing high quality images in extreme low light conditions, often resulting in high noise and color loss. Conventional techniques that employ denoising algorithms or burst imaging help with regard to somewhat low-light conditions but often fail to produce quality images in extreme low light conditions (<0.1 lux). We aim to apply the concepts presented in [this paper](#) in order to demonstrate how CNNs can be applied to outperform these conventional algorithms. We aim to extend this concept further to also resolve details within TEM images.

We chose this paper because of its particular focus on an end-to-end CNN aimed at resolving low quality images and because of the impressive results that came out of it. Additionally, the architecture in the paper appears to be rather versatile with regard to processing low-quality images and therefore capable of applications beyond conventional photography.

This is best described as an image processing problem with supervised learning. We aim to predictively enhance the quality of an image with extremely poor lighting by using a convolutional neural net with supervised learning.

- **Related Work**

The primary architecture used in the paper we are re-implementing is [U-net](#). In this paper, the researchers develop an architecture to classify pixel data in images (image segmentation) for microscopic biomedical images. This architecture is composed of a series of convolution and max-pooling layers in which the resolution is downsampled, followed by a symmetric series of what they call “up-convolution” layers in which the resolution is upsampled. This leads to effective localization and rapid execution on images of considerable size.

Existing implementations:

1. <https://github.com/cchen156/Learning-to-See-in-the-Dark>

- **Data**

The data will be sourced from the See-in-the-Dark (SID) dataset which contains 5094 short-exposure images, each paired with an associated long-exposure image. These images are either 4240x2832 or 6000x4000 in size. These short-exposure images are taken from two different cameras (Sony **a**7S II and Fujifilm X-T2) with three different exposure times (1/10, 1/25, 1/30 seconds). The associated long-exposure images are taken with the same

cameras of the same scene with the same lighting, only with much longer exposure times (therefore greater low-light quality). If we consider the raw, un-compressed data, then this dataset comes out to be ~270 GB in size, so it is a fairly large dataset and will likely require some significant preprocessing; possibly involving batched data loading or other clever techniques.

- **Methodology**

The architecture of our model is almost entirely a fully-convolutional neural network, specifically the architecture is the [U-net](#) architecture. We will be training our model using the SID dataset. We will compare the output of the model with the ground truth image in order to calculate a loss and then adjust the parameters of the model using typical SGD. The hardest part about implementing this model may be the scale of the model (~4000 epochs for training) and the size of the data we are implementing. This will likely require a great deal of time and computation to train and test extensively.

- **Metrics**

Ultimately our project will have succeeded if we can provide a model that can brighten a low light image and denoise it. We will test the model on poorly lit images taken by a standard smartphone and compare the result of the model with the initial smartphone images.

In our case, the conventional notion of raw pixel accuracy is not entirely well correlated with the effectiveness of the model. Therefore, it is necessary to define some other units of measure that are relevant to what we are trying to solve, i.e. image noise and general image structure.

In the paper we are sourcing, the authors primarily made use of two measurements to define the effectiveness of the model, specifically Peak Signal to Noise Ratio (PSNR) and Structural Similarity Index (SSIM). The authors of this paper were aiming to outperform

conventional techniques used to enhance low light images and so these units of measure were relevant in determining how effectively the model reduced noise and conserved image structure. For our project, we define accuracy using the same two values so we can compare our results with that of the paper and other noise reducing algorithms.

Our base goal is to get a functioning network that will replicate the results that are provided in the paper we are basing our project on. Our target goal is to completely replicate the results that are provided in the paper we are basing our project on. Our stretch goal is to replicate the results of the paper and extend the functionality of the model to TEM images.

- **Ethics:**

Fortunately for our case, we're trying to train our denoising algorithms using images of non-human objects. We escape the difficulty of racial bias. But the price we pay is that perhaps our model might actually underperform when trying to resolve the human images in the dark. In essence we recognize that we prevent racial bias by potentially weakening the scope of our algorithms.

I could see the result of this project being applied to self-driving cars. In that case all people who commute would be at stake. This actually poses some serious considerations because the consequences for a mistake could be fatal. On this end, we would have to think carefully whether our metric of success could possibly provide some bound on the error that the algorithm will make in a real-life setting. If so, what is the relationship between them. For example, currently we elect to use Peak-to-Noise ratio type of metric to define our accuracy or success in processing the low-light images to get back the low-noise bright image. If the images later would be used for other applications like object recognition, then perhaps it would be worthwhile to consider the metric of successes instead to be given an object recognition architecture, what the accuracy of classification our result images give.

- **Division of labor:** Briefly outline who will be responsible for which part(s) of the project.

Jeremy

- Preprocessing
- Implementing the architecture
- Training the model
- Write up / Reflection

Jose

- Digital Poster
- Presentation
- Measuring accuracy / Testing model performance

Juan

- Digital Poster
- Write up / Reflection
- Preprocessing
- Implementing the architecture

Wesley

- Implementing the architecture
- Training the model
- Measuring accuracy / testing model performance
- Presentation
- Write up/reflection