

IMAGE SYNTHESIS AND LIGHT CORRECTION USING MACHINE LEARNING APPROACH

Jyoti Ranjan Labh^{1*}, Rakesh Kumar Dwivedi²

^{1*}Teerthanker Mahaveer University, Moradabad, India (jrlabh@gmail.com)

²Teerthanker Mahaveer University, Moradabad, India (dwivedi.rakesh02@gmail.com)

Abstract:

In the context of visual effects and computer graphics, image generation and image enhancement are one of the basic problems, Same situation in other areas like space science and medicinal science. The picture quality is impacted by the presence of light in the image. This study introduces a unique method for generating and enhancing low-illumination photographs and generation of brighter time image it implies the usage of a deep mastering set of rules, particularly a deep convolutional Wasserstein generative adversarial network (DC-WGAN). The method involves changing snap shots from RGB to CIELAB shade space, which aligns extra intently with human visible belief, taking into account specific illumination estimation and mitigation of uneven lighting fixtures outcomes. By employing DC-WGAN to decorate the brightness aspect via an extended era community, the algorithm captures and amplifies important photograph features. The stronger LAB photographs are then converted back to RGB area to produce the final improved photos. The effectiveness of this approach is tested thru experiments under well known, special, and realistic conditions, demonstrating advanced overall performance as compared to 4 typically used algorithms. This study affords an important technological advancement for enhancing robot target reputation and renovation operations in area environments.

Keywords: Recurrent, Low-illumination, photograph enhancement, deep studying, deep convolutional Wasserstein generative adversarial Network (DC-WGAN), CIELAB coloration area, picture processing, space surroundings, illumination estimation, sensible robots, RGB

INTRODUCTION

A large amount of the information we consume is visual in nature. This visual data is primarily stored in images, and as electronic technology advances, so too does the variety of electronic imaging equipment used to record this data. The specific application environment often dictates the choice of imaging equipment, and the hardware capabilities of the device have a direct impact on the final image quality. In many images, low visibility conditions such as environmental disturbance like snow, rain or other environments problems can lead to color distortion, grainy images, and other image quality issues. Although advanced imaging technology can improve image quality, environmental conditions still affect its effectiveness. Various environmental factors still affect the quality of captured images, even with the best devices, and technology cannot fully compensate for these fundamental limitations. In general, standard image processing software is designed to work best with images taken under ambient light. Image classification, object recognition, and image analysis are among the applications that can significantly degrade in performance when used with low-light images. Image enhancement in low light has become the core of computer vision. The purpose of this technique is to restore the missing details in unpleasant images taken in harsh lighting conditions. The goal is to produce realistic-looking images that preserve important structural information by maximizing contrast, reducing noise and by adjusting the graphical elements and. Improving low-light image quality to

generate the image of different timings and it is essential for both human vision and computer vision, including object recognition and image recognition. Therefore, research on those images have low-light conditions image enhancement methods is of great value. Photograph processing technology is good technique to overcome those challenges, play a crucial role in improving the first-rate of visible records collected by using robotic structures. One of the most important issues is the processing of low-illumination pics, wherein traditional enhancement methods frequently fall quick. Techniques together with histogram equalization (HE) and gamma correction have been hired to adjust photo comparison and brightness. However, those strategies may be inconsistent, particularly in low-illumination eventualities wherein the visual statistics is drastically compromised. Retinex theory gives every other technique by means of decomposing photos into reflectance and illumination additives. This principle has been prolonged through various algorithms consisting of multiscale Retinex (MSR), single-scale Retinex (SSR) and multiscale Retinex with colour recuperation (MSRCR), which intention to improve image fidelity and decorate color representation. Despite their effectiveness, those methods can nonetheless struggle with complicated interference in area environments. Recent advancements in deep learning have brought new opportunities for picture enhancement. Techniques like convolutional neural networks (CNNs) and generative Adversarial networks (GANs) have validated big upgrades in photo type, popularity, and enhancement. Deep learning techniques leverage hierarchical function extraction and gaining knowledge of techniques to achieve better results in low-illumination conditions.

Theoretical Basis:

Image fusion is an important technique in image digitization, in which information from multiple images of the same scene is combined. To create a single, higher-quality image, these features are extracted, rearranged, and combined [1]. This technique mostly uses multiple exposures to improve the quality of re-light the images. Multi-exposure image fusion creates a composite image that extracts specific features from multiple images exposed at different positions and also combines them. To improve the contrast. consider a fusion model and an optimization method. In order to achieve better accuracy in low-resolution regions, this method first uses the illumination estimation to generate a weight matrix to blend the images and then implement the camera response model to calculate the exposure. In continuation to improve the quality, the input images are fused according to the weight matrix. In this study, another optimization technique based on the camera system is proposed. Recent advances in deep learning have had a significant impact on computer vision research, especially in object recognition, image interpretation and in terms of knowledge. The most used deep learning model is the convolutional neural network, whose fundamental properties such as energy distribution, clustering, and local receptive fields reduce learning accuracy are reduced and make training easier. These systems also allow for widespread use of complex objects, as they enable image manipulation such as translation, rotation, and scaling. The researchers used convolutional neural networks to improve the quality of low-light images. For adaptation and noise reduction, used a stacked sparse denoising autoencoder trained on low-light features [21] with noise generated by do it and on. MSRNet, a low-light image enhancement network, by extending the MSR method and simulating multidimensional Retinex wu with Gaussian convolution kernels [2].

GAN: Generative Adversarial Network

To simulate the distributions, Goodfellow [22]. introduced the Generative Adversarial Network (GAN) [3]. The GAN consists of a discriminator and a generator. The discriminator seeks to distinguish between synthetic and real images, and the generator seeks to produce images that may confuse the two producers. The results showed the distribution of information in this zero-sum game with the discriminator. , Figure 1 shows how the Gan model working. The parameters to guide the design

$$\begin{aligned} \min_G \max_D L_{cGAN}(G, D) \\ = E_{I,J}[\log(D(I, J)) + E_I[\log(D(I, G(I)))] \\ \times E_I[\log(1 - D(I, G(I)))] \end{aligned} \quad \text{---[1]}$$

However, GAN-based methods cannot be overcome because they do not pre-model the data. Conditional GANs, such as DE-GAN, have been proposed as a solution to this problem. In contrast to traditional GANs, DE-GANs improve controllability and learning by using necessary parameters to guide the design.

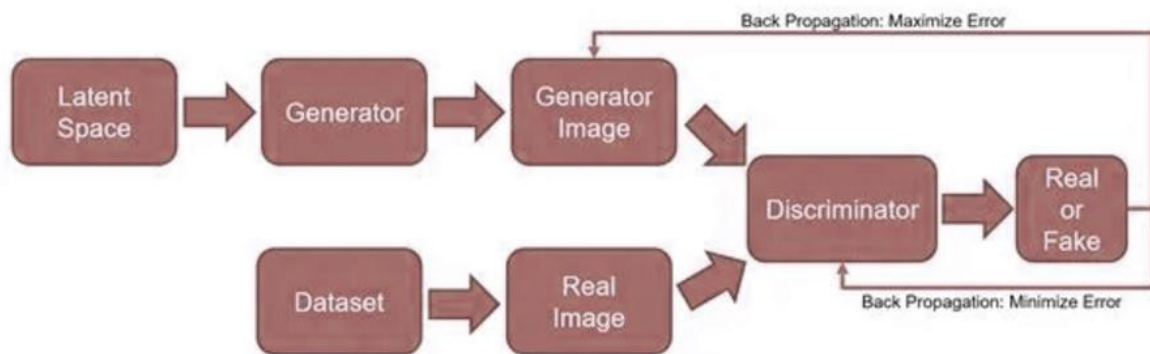


Fig:1, Generative Adversarial Networks (GANs)

G denotes the standard deviation, D denotes the deviation, the first index denotes I, the original index denotes J, and the corresponding (false) index denotes G(I). In the DE-GAN model, it learns to form a $G:I \rightarrow J$ graph, and thus learns to transform from the original graph I to a graph J. Figure 2 shows how DE-GAN was developed [4]. Here, the original image J is the corresponding ground truth, the generated image G(I) is the amplified image, and the original image integrated system. objective function of the discriminator part is:

$$\min_G \max_D L_{cGAN}(G, D) = E_{I,J}[\log(D(I, J)) + E_I[\log(D(I, G(I)))] \quad \text{-----} \quad \text{---[2]}$$

Objective function of the generator part is:

$$\min_G L_{cGAN}(G, D) = E_I[\log(1 - D(I, G(I)))] \quad \text{-----[3]}$$

Two images are sent to the discriminator (D): (I, G(I)) and (I, J). It generates a score between 0 and 1 that indicates the probability that the input is truly a binary image when it has been trained to distinguish between the two after. True values are those greater than 0.5, and false positives are those less than 0.5. In the mini-max game, training alternates between G and D until it reaches a Nash equilibrium in which the probability of extracting D is 0.5, denoted by that simulated and actual images are indistinguishable. The resulting image G(I) is now very similar to the original image. GANs work through a two-player game theory where the discriminator and the generator compete against each other. While the classifier tries to determine whether the estimates are fake or real, the synthesizer tries to provide estimates that might mislead the classifier. Developers can build very accurate models because of this adversarial behavior, which continues until they become indistinguishable from reality

Dataset:

The SID system, which consists of two low-luminosity and ultra-thin images, [5]. Data collection was performed using different exposure times under low light conditions, resulting in 5094 images at very low light levels, and short exposures, compared with low-light, long-exposure images. The 424 indoor and outdoor scenes in the SID dataset had different illuminance levels detected by different short exposure times. The total 1467 images in the Self dataset database include 14095

images captured by 10 cameras in five different scenarios. A set of 200 images is divided into 100 images for testing and 200 images for evaluation, and training is performed on the remaining 1268 images. Unlike the jpg or png formats used in real images, Self dataset images are stored in the camera raw. When models trained on Self Dataset are applied to other images, transforming this image may lead to better results.

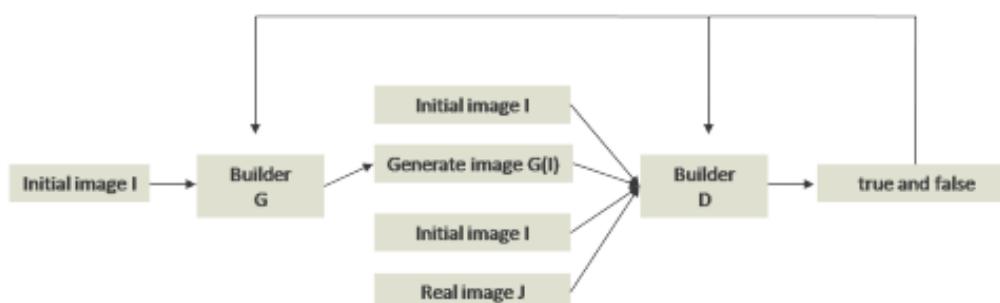


Figure:2, DEGAN Architecture

III.METHODOLOGY:

The results obtained in this study were based on the results of the two subjects for the teaching of low and high exposure scenes. This method adds a local discriminator to reduce local noise based on the traditional GAN architecture [6] with and a universal identifier. The patches of both images are passed to the local discriminator, and the generated image and the real image serve as global classifiers he. All these classifiers evaluate the accuracy of the resulting image. The optimized image is produced immediately if it is claimed to be authentic. If not, the software receives feedback, improving its output until it can confuse both differences.

$$L_{cWGAN-GP} = E_{I,J}[D(I,J)] - E_I[D(I,G(I))] + \lambda_{GP} E_I[(\|\nabla_I D(\hat{I})\|_2 - 1)^2] - [4]$$

G represents generator and D represents discriminator in this structure. J means ground truth and I mean low light image. The score on the line connecting J to the surface of the gradient, G(I), is denoted by $\hat{I}$. GP acts as a selector. The loss function of WGAN-GP does not use logarithm, unlike the loss function of traditional GAN algorithms. Instead of using JS divergence, it uses Wasserstein distance to measure the difference between the simulated and true distributions. The main advantage of Wasserstein's theory is that it can accurately measure differences between distributions, even in cases where there is little or no overlap between their support groups These are the people. The effect of vanishing gradients is particularly pronounced because JS spacing becomes uniform in these cases. To solve this problem, WGAN-GP loss adds a gradient penalty term. The effect of vanishing gradients is particularly pronounced because JS spacing becomes uniform in these cases. To solve this problem, WGAN-GP loss [7] adds a gradient penalty term. Since the k-Lipschitz constraint is imposed on the discriminator, this penalty limits its scalability to some extent. However, compared to traditional convolutional tasks such as L1 or L2, the amplification of and faster synchronization, reducing graphical problems. Previous studies have shown that adding an L1 or L2 representation to a target task stimulates the motor (G) to identify biologically similar sample positions, which is useful for image translation projects. Although L2 attenuation is excellent for visual reconstruction, artifacts may still be present. Conversely, losing L1 effectively removes artifacts but requires repeated capture and reconstruction in better ears. Although L1 decomposition can reduce artifacts, it fails to capture image features and introduces reconstruction at the edges. Only L1 loss or L2 loss has been investigated in the target task in previous research. However, to better capture random and stochastic effects without introducing artifacts, the function uses a combined approach that integrates L1 and L2 losses into the target function.

$$L_1(G) = E_{I,J}[||J - G(I)||_1] \quad \text{---}[5]$$

$$L_2(G) = E_{I,J}[||J - G(I)||_2]^2 \quad \text{---}[6]$$

The calculation of previous losses is calculated as follows:

$$L^* = \min_G \max_D L_{cWGAN-GP}(G, D) + \lambda_1 L_1(G) + \lambda_2 L_2(G) \quad \text{---}[7]$$

When regional enhancements are needed, global discrimination may not be sufficient because it takes into account the entire image, and allows for all enhancements to be made. This weakness is overcome by using the global isolation scheme. Local contrast is added in addition to global contrast to enhance some regions of the image. The local discriminator identifies the origin of randomly selected positions in both the updated and original images. By reducing local disturbance, the local socialization system makes the restored areas feel more authentic and natural like. The relative discriminant provided by Alexia consists of two parts: one part estimates the likelihood of the original data being more real than the artifacts, and the other part is concerned with the possibility that artifacts may be more real than artifacts. The first part should be close to 1 and the last part to 0.

$$D_{ra}(x_r, x_f) = \sigma(C(x_r) - E_{x_f \sim P_f}[C(x_f)]) \quad \text{---}[8]$$

$$D_{ra}(x_f, x_r) = \sigma(C(x_f) - E_{x_r \sim P_r}[C(x_r)]) \quad \text{---}[9]$$

In this case, the discriminator is represented by C, the distributions of the actual and generated images are represented by x_r and x_f respectively, the probability distributions of the actual and generated images are represented by P_r and P_f , the expected value is represented by E, and the sigmoid activation function is represented by σ . The final loss functions of the global discriminator (D) and the generator (G), using the loss function proposed by Mao[23], are as follows:

$$L_G^G = E_{x_f \sim P_f}[(D_{ra}(x_f, x_r) - 1)^2] + E_{x_f \sim P_f}[D_{ra}(x_r, x_f)^2] \quad \text{----}[10]$$

$$L_D^G = E_{x_r \sim P_r}[(D_{ra}(x_r, x_f) - 1)^2] + E_{x_f \sim P_f}[(D_{ra}(x_f, x_r)^2] \quad \text{----}[11]$$

Six 32×32 images from each simulated and real image are randomly selected to use the local classifier. For salt (G) and local salt (D), the corresponding loss functions are: The deviation of the obtained images from these values is determined by P_{fp} , and the ratio of the detected images is determined by P_{rp} . The global classifiers play a role after the five steps of image reduction. The algorithm consists of a convolution layer, a batch normalization, and a Leaky ReLU activation function [8]. The sigmoid activation function after the fifth cycle of downsampling. Similar to the global discriminator, the local discriminator uses a sigmoid function to localize the image where he drew, after completing four reduction operations. The DE-GAN architecture is the main tool used in this work. However, instead of using a traditional CNN, a U-Net is used to improve the preservation of the original image features. The U-Net [9] shown in Figure 3 uses skip connections, similar to an autoencoder. When combined with map data, these connections secure the data. The generator (G) consists of 15 layers, including 8 convolutional layers and 7 non-convolutional layers. As an encoder, the first stage denoises the underwater image using convolutional techniques, extracts textual features, applies batch normalization, and activates Leaky-ReLU. The second stage acts as a decoder, reconstructing the underwater image from the original image by pooling the normal features using a deconvolutional network. In this work, the implementation of Adaptive Dense Features Fusion algorithm [10] in U-Net algorithm is presented. In contrast to the fusion approach, ADFF module. Here, $\alpha_{1xy}, \dots, \alpha_{nxy}$ are used as parameters for $\alpha_{1xy}, \dots, \alpha_{nxy}$ using the softmax function. Using $i_1 \rightarrow nxy, \dots, i_{(n-1)} \rightarrow nxy$, and i_{nxy} , the cells 1

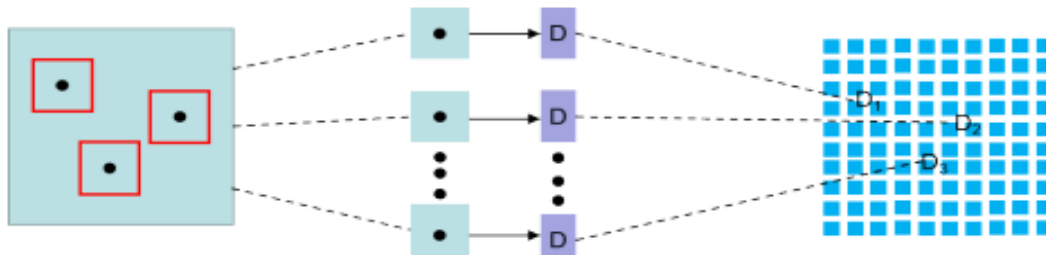


Figure:3, U-net Architecture

focuses on the features extracted from each layer. It helps in image reconstruction by combining these data in a logical way, which provides better results for further extraction or reconstruction. Specifically, it collects data from all previous layers using strong connections and feeds them into the ADF engine. To highlight the salient features during fusion, this module learns the class weights of the features of each level. Figure 4 The ADF output of the generating encoder at level $n = \{2, 3, 4, 5\}$ is shown in Figure 4.

$$L_D^G = E_{x_r \sim P_{\eta}} [(D(x_r) - 1)^2] + E_{x_f \sim P_{\theta}} [(D(x_f) - 1)^2]$$

$L_G^L = E_{x_r \sim P_{\theta}} [(D(x_r) - 1)^2]$

$$S(p) = \frac{2\mu_x\mu_y + C_1}{\mu_x^2 + \mu_y^2 + C_1} \cdot \frac{2\sigma_{xy} + C_2}{\sigma_x^2 + \sigma_y^2 + C_2}$$

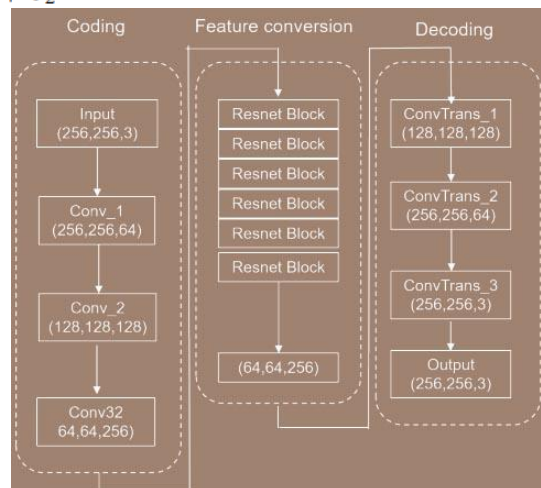


Fig:4 , Flow Chart of Generator

At encoder level n , it represents the output of the ADF unit. The feature vector at point (x, y) on this graph is denoted by in_{xy} . Before ADF integration, the feature map of encoder level n is represented by um ($il = il$), and the feature map at (x, y) on this map is um_{xy} . After updating to n , the feature maps from feature l (where $l = 1, 2, 3$) and $il \rightarrow n$ are drawn, and the feature map at (x, y) on this equilibrium map is $il \rightarrow n_{xy}$. For the simulated feature map, α_{xy} represents the spatially fixed weights learned by the network. All channels of the map share the same dimension, α_{xy} . Note that where $\alpha_{xy} \in [0, 1]$, $\alpha_{1xy} + \dots + \alpha_{nxy} = 1$. The definitions below are also important:

$$\alpha_{xy}^1 = \frac{e^{\lambda_{\alpha_{xy}}^1}}{e^{\lambda_{\alpha_{xy}}^1} + \dots + e^{\lambda_{\alpha_{xy}}^{n-1}} + e^{\lambda_{\alpha_{xy}}^n}} \quad [13]$$

$\pi_1, \dots, \pi_n$ are produced. For each stage, this approach allows the main data from all the initial stages to be combined. Two-step convolutional layers are used to downsample, and two-step convolutional layers are used to extract the samples. The downsampling is replaced by approximation in the ADF unit of the decoder, which has the same configuration as its encoder.

SSI:

Since the images are superimposed, the similarity in the images describes the morphological properties of the objects in the image. It measures brightness relative to the mean, using the mean as the variance, and correlation by variance. This similarity can be expressed as:

--[14]

where y is the generated 11×11 image, x is the 11×11 seabed image, and p is the center pixel of the image. The mean is denoted by μ , the means of x and y are denoted by σ_x and σ_y , respectively, and the variance between x and y is denoted by σ_{xy} . The values of 0.02 and 0.03 for the parameters $C1$ and $C2$, respectively. The overall deviation between the design and the submerged image due to underwater reflection is:

$$L_S^G(x, G(x)) = 1 - \frac{1}{N} \sum_{p=1}^N (S(p)) \quad \text{--}[15]$$

The correlation coefficient between the reconstructed image and the data collected from the underwater image is obtained to identify the underwater scene:

$$L_S^L(x', G(x')) = 1 - \frac{1}{M} \sum_{p=1}^N (S(p)) \quad \text{--}[16]$$

The overall loss function of the network is:

$$Loss = \lambda_1 L_S^G + \lambda_2 L_S^L + \lambda_3 L_G^G + \lambda_4 L_G^L \quad \text{--}[17]$$

The weights π_1 , α_2 , α_3 , and α_4 are determined based on experimental and theoretical results. where G is the generated image and L is the loss. The structural similarity loss, like the generator function loss, depends on the trained data as well. Therefore, the sum of the composite weights and the data loss in the dataset is the same. The structural similarity loss [11] and the generator losses are introduced to avoid excessive emphasis on local loss at the expense of global operation. This balancing method ensures a proper trade-off between local and global losses.

IV. EXPERIMENTS:

The analysis includes 3800 evening and 3800-day samples, and [12] Training uses Adam optimizer with a learning rate of 0.0001 and a batch size of 16. The proposed scheme is compared with four low lighting image enhancement schemes using both independent measurements and statistical measurements. The test dataset consists of 3,600 evening images. The proposed method is compared with four other image enhancement algorithms in Figure 4. The image reconstruction results in the hyper parathyroid technique revealed hypersensitivity at distant locations.[8] Having too many parameters affects the strength of the other method. The third method produces images with poor local contrast, and objects appear blurry at the beginning. Another method gives poor local definition, complete darkness, and obvious image noise. Finally, one method has poor local contrast and has trouble removing local color artifacts. These four algorithms do not provide complete visualization results because they only focus on updating the global map, ignoring local changes. On the other hand, the proposed technology enhances the overall light reflectance and provides a better, more realistic viewing experience by improving and achieving targets in source areas. The results of the proposed method are compared with the results of four low illumination image enhancement algorithms in the SID [13] on it, and enlarged sections are shown in Figure 3 to better illustrate the particle retention. Compared with day images and images optimized by four other algorithms, the proposed method produces images with different resolution, as seen in the significantly by reconstructing the particles in the human-edited image, enhancing the lighting and creating clearer and more realistic images. In terms of computation and calculation, it performs better than other algorithms. normalized sections. This shows how the proposed method can restore excellent features, resulting in better image quality

Day image metrics, day imaging difference metrics, low light image quality, and distribution entropy are the four metrics used to analyze day images. Two well-known image quality metrics are IQM and CIQE. IQM uses three methods to evaluate reconstructed images: IConM, color calibration, and different time imaging. IConM and IQM are the most widely used algorithms. Higher scores indicate better overall image quality when CIQE evaluates enhanced images based on color, saturation, and contrast. The complexity of an image is measured by its entropy; higher entropy indicates higher content and higher image quality. The results of the proposed approach compared with traditional methods on the self-dataset database are shown in Figure 4

Experimental Verification and Analysis:

The number of trials was 3600 to reduce unnecessary errors. Comparisons with other datasets Each algorithm is shown in Table 1. The method used in this study measures variability and IQM in a broad sense. Conclusions with IConM Compared with previous reports, the improved model in this study has better overall accuracy, and high-contrast, while enriching the media. The algorithm used in this paper achieved inconsistent results for the IQM index and the CIQE index. The result shows an improvement in low light imaging. Too much saturation can make an image look better, but too much saturation results in artifacts with unnatural textures. Enhanced images perform better on metrics such as content, contrast, and overall impressions. According to experimental results, the proposed method removes artifacts and improves contrast .

Table 1: comparison of the numerous objective indices used in each method.

Method	IQM	IConM	CIQE	Entropy
UPND [12]	3.9895	0.797 1	0.638 5	7.472 1
ODM [9]	4.1396	0.855 5	0.579 9	7.043 4
FU-GAN [10]	4.3071	0.916 5	0.541 0	7.473 6
UGAN [11]	4.335	0.907 9	0.564 5	7.557 7
Our method	4.3977	0.926 9	0.583 9	7.608 9

ABLATION RESEARCH AND ANALYSIS

To illustrate the performance of each component of the model in this research, a series of method evaluation graphs were constructed and compared with The network configuration is as follows: 1)

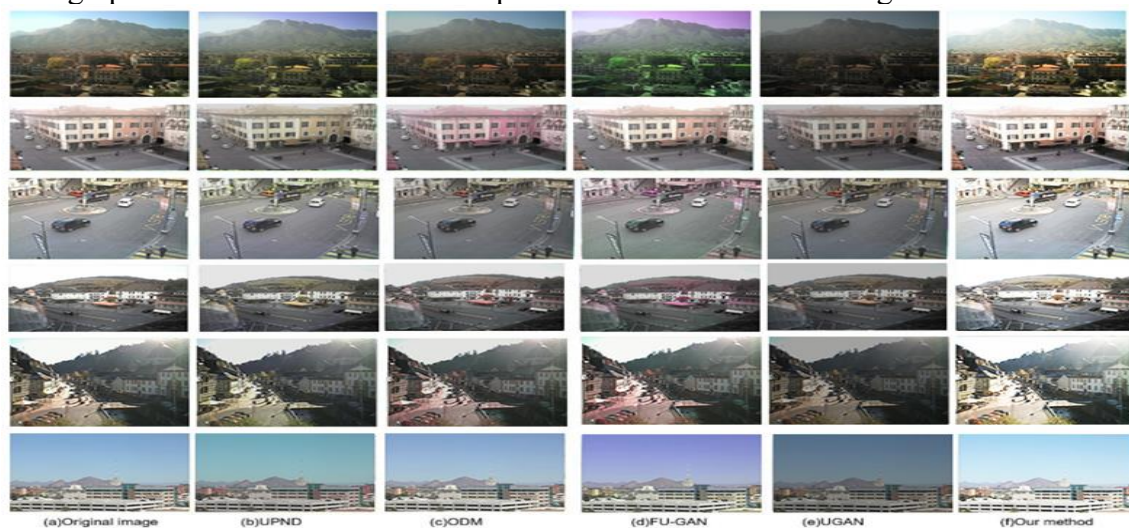


Figure:4, Experimental results of proposed method compared with classical algorithms on SID dataset

respectively. Table 3 shows that by increasing the learning rate of the network, using residual blocks in the vehicle improves the performance compared to U-Net due to the traditional structure. Moreover, the performance of the adaptive densefeature fusion is demonstrated by comparing the proposed structure (with ADFF) with that without it. The key features of all the previous fields are collected and presented by the ADFF program. The results in Figure 2 are consistent with the results in Figure 6 that residual blocks improve the smoothness of detail extraction from low light images. The ADFF module[14] facilitates feature fusion, enabling the network to leverage key information for generating clearer images Simple U grid; 2) residual error only, no ADFF block U-grid convolution (Ours-woADFF); 3) On top of Ours-woADFF, an ADFF module is added (algorithm in this paper). Table 3 and Figure 5 present qualitative and quantitative results,



Figure:5, Experimental results of proposed method compared with classical algorithms on self-dataset Overall, the proposed algorithm achieves the best performance in underwater image enhancement, as demonstrated by SSIM and PSNR [15] metrics and visual results, validating the effectiveness of the network model. Table 3 presents objective indicators for the enhancement effects of the FGF, Patch-net, and the proposed algorithms. The proposed method effectively improves the Average Gradient, Entropy, and Extended Maximum Entropy of the images.

Data set	Evaluation index	Original	FGF	Patch-net	Our method
Self	AG	3.57	4.57	5.67	5.67
	EN	34.82	56.46	60.46	60.46
	EME	6.73	7.88	8.21	8.21

Table 2: Comparison of different image enhancement method

	U-Net	Ours-woADFF	Our method
SSIM	0.752	0.653	0.878
PSNR	21.641	35.1734	33.452

Table 3: quantitative evaluation of DE-GAN on artificial low light photos using various generator configuration

MODEL SIMPLIFICATION TEST:

This work implements related proposed components, such as k-path feed-forward,[16] and easy, for testing and evaluation purposes. Simplified structures are discussed in this paper. Figure 6 shows the process from training to convergence using a 22-layer network and SID dataset for testing. The

multipath connection is represented by the blue line, and the VDSR system [17] is represented by the black line. It can also be seen that during the aggregation, the PSNR value varies significantly. The overall speed is improved and the bandwidth is enhanced by the pixel encoding algorithm

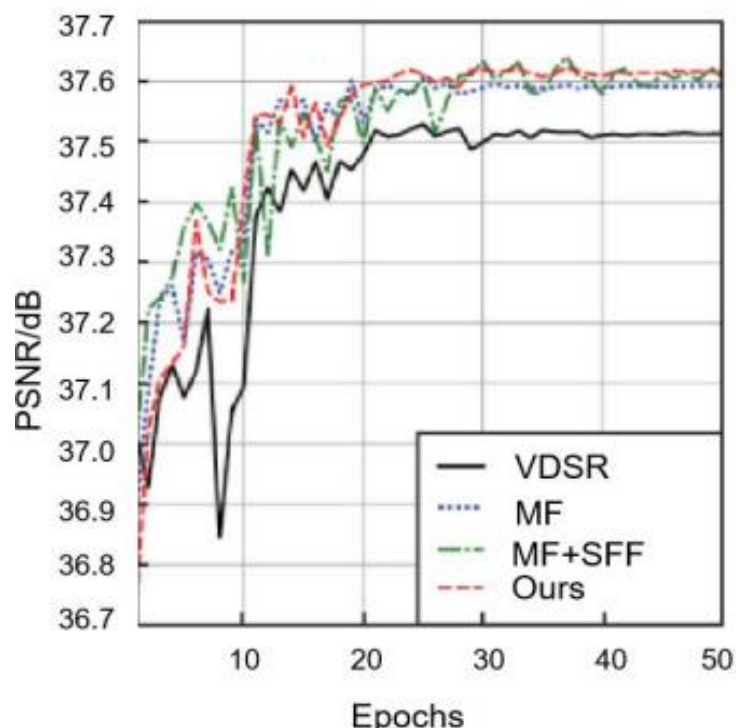


Figure 6: Convergence analysis of different magnifications

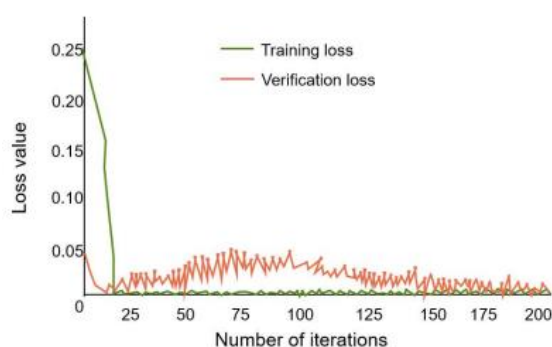


Figure 7: Convergence analysis of loss value

MODEL LOSS CURVE:

Variance = 0.1, mean = 0, dropout = 0.5, bias = 0.1, and learning rate = 0.01 were the initial values of the metrics.[18] The training and validation spectra are shown in Figure 13. After about 20 cycles, both the training and validation curves[19] are at the same level, indicating that the method converges quickly. This indicates that the target detection [20] task has been successfully trained in the model.

CONCLUSION

The paper presents a complete technique to adaptive photo enhancement underneath low illumination situations, and different times in a day . Acknowledging the challenges posed by means of varying brightness and assessment because of environmental elements together with climate and time of day. By spotting that a one- length-fits-all enhancement set of rules is insufficient, the study proposes a

tailor-made technique that differentiates between day and night pictures and employs awesome processing techniques for every.

REFERENCES

1. Dong, Z. and Tian, X., 2015. Multi-level photo quality assessment with multi-view features. *Neurocomputing*, 168, pp.308-319.
2. Shen, D., Wu, G. and Suk, H.I., 2017. Deep learning in medical image analysis. *Annual review of biomedical engineering*, 19(1), pp.221-248.
3. Lu, Y., Chen, D., Olaniyi, E. and Huang, Y., 2022. Generative adversarial networks (GANs) for image augmentation in agriculture: A systematic review. *Computers and Electronics in Agriculture*, 200, p.107208.
5. Chen, C., Chen, Q., Xu, J. and Koltun, V., 2018. Learning to see in the dark. In *Proceedings of the IEEE conference on computer vision and pattern recognition* (pp. 3291-3300).
6. Uzun, C., Çolakoğlu, M.B. and İnceoğlu, A., 2020. GAN as a generative architectural plan layout tool: A case study for training DCGAN with Palladian Plans and evaluation of DCGAN outputs. *A|Z ITU Journal of Faculty of Architecture*, 17, pp.185-198.
7. Moghadam, A.Z., Azarnoush, H. and Seyyedsalehi, S.A., 2021, May. Multi WGAN-GP loss for pathological stain transformation using GAN. In 2021 29th Iranian Conference on Electrical Engineering (ICEE) (pp. 927-933). IEEE.
8. Maniatopoulos, A. and Mitianoudis, N., 2021. Learnable leaky relu (LeLeLU): An alternative accuracy-optimized activation function. *Information*, 12(12), p.513.
9. Ronneberger, O., Fischer, P. and Brox, T., 2015. U-net: Convolutional networks for biomedical image segmentation. In *Medical image computing and computer-assisted intervention—MICCAI 2015: 18th international conference, Munich, Germany, October 5-9, 2015, proceedings, part III 18* (pp. 234-241). Springer international publishing.
10. Sahoo, S. and Nanda, P.K., 2021. Adaptive feature fusion and spatio-temporal background modeling in KDE framework for object detection and shadow removal. *IEEE Transactions on Circuits and Systems for Video Technology*, 32(3), pp.1103-1118.
11. Huang, Y., Song, R., Xu, K., Ye, X., Li, C. and Chen, X., 2020. Deep learning-based inverse scattering with structural similarity loss functions. *IEEE Sensors Journal*, 21(4), pp.4900-4907.
12. Al-Shemarry, M.S.; Li, Y.; Abdulla, S. An Efficient Texture Descriptor for the Detection of License Plates From Vehicle Images in Difficult Conditions. *IEEE Trans. Intell. Transp. Syst.* 2019, 21, 553–564
13. El-Shair, Z.A., Abu-raddaha, A., Cofield, A., Alawneh, H., Aladem, M., Hamzeh, Y. and Rawashdeh, S.A., 2024, July. SID: Stereo Image Dataset for Autonomous Driving in Adverse
14. Long, M., Zhou, J., Zhang, L.B., Peng, F. and Zhang, D., 2024. ADFF: Adaptive de-morphing factor framework for restoring accomplice's facial image. *IET Image Processing*, 18(2), pp.470-480.
15. Hore, A. and Ziou, D., 2010, August. Image quality metrics: PSNR vs. SSIM. In *2010 20th international conference on pattern recognition* (pp. 2366-2369). IEEE.
16. Gong, W., Sun, T., Bai, H., Tsai, J.Y., Ling, H. and Yan, Q., 2024. Graph Transformer Networks for Accurate Band Structure Prediction: An End-to-End Approach. *arXiv preprint arXiv:2411.16483*.
17. Vint, D., Di Caterina, G., Soraghan, J.J., Lamb, R.A. and Humphreys, D., 2019, May. Evaluation of performance of VDSR super resolution on real and synthetic images. In *2019 Sensor Signal Processing for Defence Conference (SSPD)* (pp. 1-5). IEEE.
18. Yang, J. and Yang, G., 2018. Modified convolutional neural network based on dropout and the stochastic gradient descent optimizer. *Algorithms*, 11(3), p.28.

19. Viering, T. and Loog, M., 2022. The shape of learning curves: a review. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 45(6), pp.7799-7819.
20. Vink, J.P. and de Haan, G., 2015. Comparison of machine learning techniques for target detection. *Artificial Intelligence Review*, 43, pp.125-139.
21. Fan, Z., Bi, D., He, L., Shiping, M., Gao, S. and Li, C., 2017. Low-level structure feature extraction for image processing via stacked sparse denoising autoencoder. *Neurocomputing*, 243, pp.12-20.
22. Goodfellow, I.J., Pouget-Abadie, J., Mirza, M., Xu, B., Warde-Farley, D., Ozair, S., Courville, A. and Bengio, Y., 2014. Generative adversarial nets. *Advances in neural information processing systems*, 27.
23. Mao, A., Mohri, M. and Zhong, Y., 2023, July. Cross-entropy loss functions: Theoretical analysis and applications. In International conference on Machine learning (pp. 23803-23828). pmlr.