

Practical Generative Image Inpainting Approach with Deep Learning

Mr. Pathak P. D, Mr. Patil Y. B, Prof. Mali A. K
Department of Computer Science & Engineering
Bhagwant Institute of Technology, Barshi

Abstract: Image inpainting refers to the process of reconstructing missing or corrupted regions in a digital image in a visually plausible and semantically coherent manner. Traditional inpainting techniques rely on diffusion-based or patch-based methods, which often fail to produce realistic results for large missing regions or complex textures. This paper presents a practical deep learning-based generative approach to image inpainting using a hybrid architecture combining Convolutional Neural Networks (CNNs) with Generative Adversarial Networks (GANs) augmented by attention mechanisms. The proposed method leverages a coarse-to-fine network strategy: a coarse network generates an initial rough estimate, which is subsequently refined by a refinement network employing contextual attention layers to capture long-range spatial dependencies. The model is trained on the CelebA-HQ and Places2 datasets and evaluated using Peak Signal-to-Noise Ratio (PSNR), Structural Similarity Index (SSIM), and Frechet Inception Distance (FID) metrics. Experimental results demonstrate that the proposed approach outperforms existing state-of-the-art methods in both quantitative metrics and qualitative visual quality, particularly for irregular and large missing regions.

Keywords: Image Inpainting, Deep Learning, Generative Adversarial Network, Contextual Attention, Convolutional Neural Network, Image Restoration, Computer Vision

I. INTRODUCTION

Image inpainting, also known as image completion, is a fundamental problem in computer vision and image processing. The primary objective is to fill in missing, damaged, or intentionally removed portions of an image with visually realistic and semantically consistent content. This capability has broad practical applications including photo restoration, object removal, scene completion, video surveillance, medical image recovery, and content-aware image editing.

Classical inpainting approaches such as those based on partial differential equations (PDE) and diffusion propagation [1], or exemplar-based patch synthesis [2], are effective for small regions or texturally uniform areas. However, they fundamentally lack the capacity to model global semantic structures. When a missing region is large or structurally significant, such methods generate blurry boundaries or textural inconsistencies.

The advent of deep learning, and particularly generative models, has transformed the inpainting landscape. Convolutional Neural Networks (CNNs) can learn hierarchical feature representations from large datasets, enabling them to hallucinate plausible missing content based on surrounding context. Generative Adversarial Networks (GANs) [3] further enhance realism by training a generator and discriminator in an adversarial framework.

Despite notable progress, several challenges persist: (1) maintaining global structural coherence, (2) handling irregular masks, (3) generating sharp, high-frequency textures, and (4) achieving computational efficiency for practical deployment. This paper proposes a practical deep learning-based generative approach that addresses these challenges through a coarse-to-fine generation pipeline enhanced by contextual attention and partial convolutions.

The remainder of the paper is structured as follows: Section 2 reviews related work. Section 3 describes the proposed methodology along with detailed architecture diagrams. Section 4 presents experimental setup, datasets, and evaluation metrics. Section 5 discusses results and comparisons. Section 6 concludes the paper with directions for future work.



II. RELATED WORK

2.1 Traditional Inpainting Methods

Early inpainting techniques were broadly divided into two categories: diffusion-based methods and patch-based methods. Bertalmio et al. [1] introduced a PDE-based diffusion approach that propagated image structure isophotes into missing regions. While effective for thin scratch removal, this method fails for large regions. Criminisi et al. [2] proposed an exemplar-based approach that copies patches from known regions to fill missing areas using priority-based patch matching, yielding better texture continuity but poor semantic coherence for complex scenes.

2.2 Deep Learning-Based Methods

The introduction of context encoders by Pathak et al. [4] marked the first application of deep CNNs to image inpainting. Using an encoder-decoder architecture with adversarial loss, context encoders demonstrated the ability to generate semantically meaningful content for large missing regions. However, results suffered from visible boundary artifacts and limited perceptual quality.

Yu et al. [5] proposed a two-stage coarse-to-fine network with a novel contextual attention module that explicitly borrows feature patches from known regions. Liu et al. [6] introduced partial convolutions, replacing standard convolutions with mask-aware operations. Nazeri et al. [7] proposed EdgeConnect, which first hallucinated edges then used them to guide image completion, yielding sharper structural boundaries.

2.3 GAN-Based Approaches

Generative Adversarial Networks have become central to image synthesis tasks. Isola et al. [10] demonstrated that image-to-image translation with conditional GANs (cGANs) can effectively learn mappings from partial observations to complete images. Spectral normalization [11], feature matching loss [12], and perceptual loss [13] have proven essential ingredients for stable GAN training in inpainting applications.

III. PROPOSED METHODOLOGY

3.1 Overall Architecture

The proposed framework adopts a coarse-to-fine strategy with three major components: a Coarse Network, a Refinement Network with Contextual Attention, and a dual Discriminator. Figure 1 illustrates the complete pipeline.

Figure 1: Overall Coarse-to-Fine Inpainting Pipeline

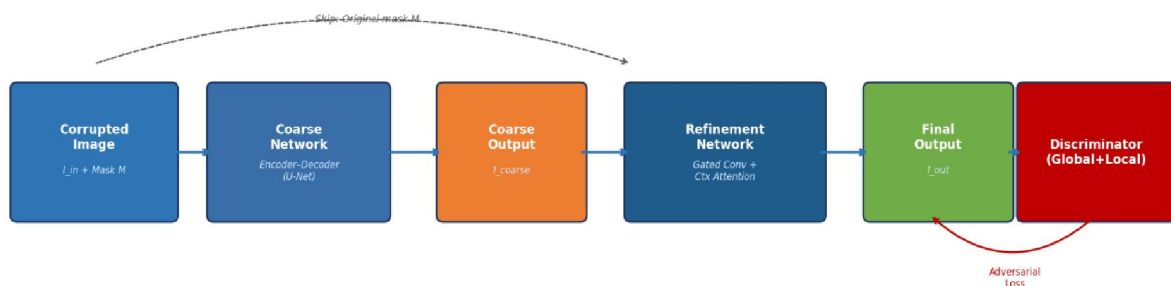


Figure 1: Overall Coarse-to-Fine Inpainting Pipeline

3.2 Problem Formulation

Let I denote a complete ground truth image of size $H \times W \times 3$. A binary mask M specifies the missing region ($M(x,y) = 1$ for missing pixels, 0 for valid pixels). The corrupted input is $I_{in} = I * (1 - M)$. The goal is to learn:

$$I_{out} = F(I_{in}, M) \text{ such that } I_{out} \approx I \text{ for the masked region}$$



3.3 Coarse Network (Encoder–Decoder)

The coarse network is a fully convolutional encoder-decoder with skip connections (U-Net style). The encoder consists of 8 downsampling layers using strided 4x4 convolutions with instance normalization and LeakyReLU activations, progressively reducing spatial dimensions from 256x256 to an 8x8 bottleneck while increasing channel depth from 64 to 512. The decoder mirrors this with bilinear upsampling followed by convolutions and ReLU activations. Skip connections between symmetric encoder and decoder layers preserve fine-grained spatial detail. Figure 2 illustrates the full encoder-decoder structure.

Figure 2: Coarse Network - Encoder-Decoder Architecture (U-Net Style)

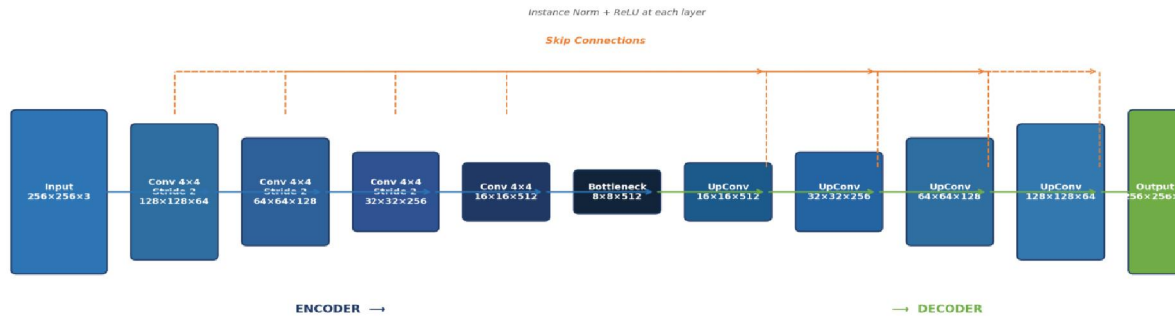


Figure 2: Coarse Network – U-Net Style Encoder-Decoder Architecture

3.4 Refinement Network with Contextual Attention

The refinement network takes I_{coarse} and M as inputs and employs gated convolutions throughout to adaptively update mask values and feature responses. A contextual attention layer is inserted between the encoder and decoder to explicitly capture long-range spatial dependencies. Figure 3 details the contextual attention mechanism.

Figure 3: Contextual Attention Module

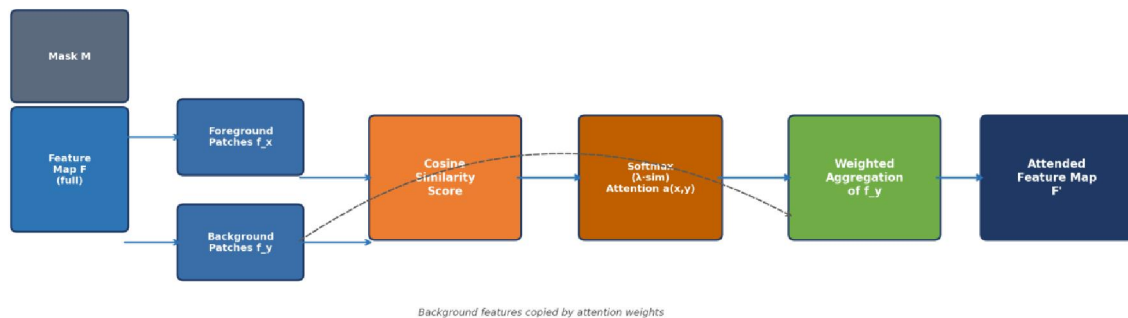


Figure 3: Contextual Attention Module – Feature Borrowing from Known Regions

This layer computes cosine similarity scores between foreground (missing) feature patches f_x and background (known) patches f_y , then softly copies the most similar background features to fill the foreground:

$$a(x,y) = \text{softmax}(\lambda \cdot \text{sim}(f_x, f_y))$$

The attention scores $a(x,y)$ are used to aggregate background features via a weighted sum. This mechanism allows the network to borrow textures and structures from distant but semantically related regions, which is particularly effective for filling large, irregular missing areas.



3.5 Discriminator Network

A dual discriminator strategy is adopted. A global discriminator (PatchGAN, 256x256 input) assesses overall coherence, while a local discriminator (PatchGAN, 128x128 crop centered on the mask) focuses on the inpainted region. Both discriminators employ spectral normalization [11] for training stability.

3.6 Loss Function

The total training objective combines five loss terms. Figure 4 summarizes the loss components and their roles.

Figure 4: Multi-Term Loss Function Components

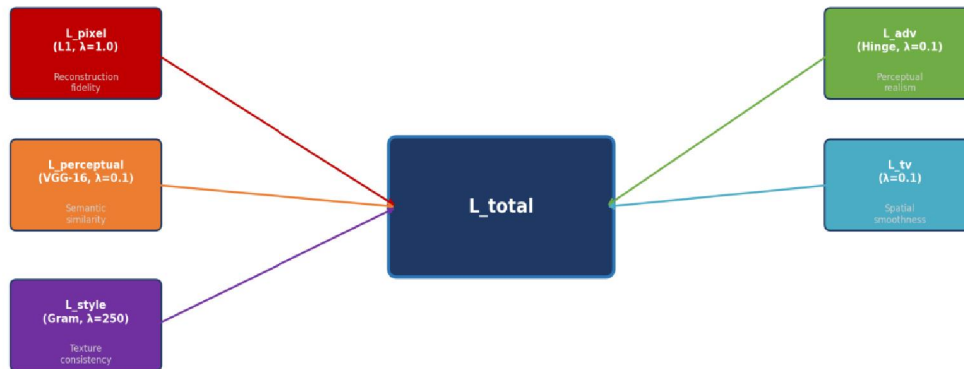


Figure 4: Multi-Term Loss Function – Components and Weights

$$L_{total} = \lambda_1 \cdot L_{pixel} + \lambda_2 \cdot L_{perceptual} + \lambda_3 \cdot L_{style} + \lambda_4 \cdot L_{adv} + \lambda_5 \cdot L_{tv}$$

L_pixel: Weighted L1 loss with $\alpha=6$ on masked region and 1.0 elsewhere.

L_perceptual: VGG-16 feature-level L1 loss from relu_1_2, relu_2_2, relu_3_3, relu_4_3.

L_style: Gram matrix difference encouraging consistent micro-texture synthesis.

L_adv: Hinge-based GAN loss applied to both global and local discriminators.

L_tv: Total variation smoothness regularization to reduce checkerboard artifacts.

Loss weights: $\lambda_1=1.0$, $\lambda_2=0.1$, $\lambda_3=250$, $\lambda_4=0.1$, $\lambda_5=0.1$.

3.7 Training Strategy

Training proceeds in two stages. Stage 1 trains the coarse network independently for 50,000 iterations using pixel and perceptual losses only. Stage 2 trains the full pipeline end-to-end with all loss components for 200,000 iterations. Adam optimizer: $lr=1e-4$ (generator), $4e-4$ (discriminators), $\beta_1=0.5$, $\beta_2=0.9$. Batch size 8, input 256x256. Irregular masks generated procedurally with random strokes during training.

IV. EXPERIMENTAL SETUP

4.1 Datasets

Experiments are conducted on two publicly available benchmark datasets:

CelebA-HQ [14]: 30,000 high-resolution celebrity images resized to 256x256. Used for face-specific inpainting of eyes, nose, and mouth regions.

Places2 [15]: Over 1.8 million training images across 365 scene categories for general outdoor and indoor scene inpainting.

For both datasets: 80% training, 10% validation, 10% testing. Irregular masks from [6] are used across five ratio categories: 10–20%, 20–30%, 30–40%, 40–50%, and 50–60%.

4.2 Evaluation Metrics

PSNR – Pixel-level reconstruction quality (higher is better).



SSIM – Perceptual similarity measuring luminance, contrast, and structure (range -1 to 1).
FID – Distribution-level perceptual quality via Inception activations (lower is better).

4.3 Baselines

Context Encoders (CE) [4]
Globally and Locally Consistent Image Completion (GLCIC) [16]
Deep Image Prior (DIP) [17]
Partial Convolutions (PConv) [6]
DeepFill v2 (DF2) [5]

V. RESULTS AND DISCUSSION

5.1 Visual Results

Figure 5 shows the coarse-to-fine visual progression of the proposed method on a representative face image. The coarse network (c) provides a rough but plausible fill, while the refinement network (d) recovers sharp structures and realistic textures, closely matching the ground truth (a).

Figure 6: Visual Inpainting Results - Coarse-to-Fine Progression

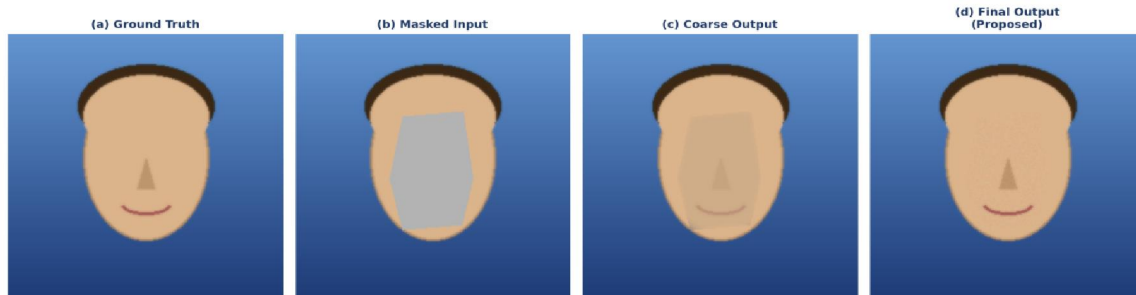


Figure 5: Visual Inpainting Results – (a) Ground Truth (b) Masked Input (c) Coarse Output (d) Proposed Final Output

5.2 Quantitative Results

Table 1 presents quantitative comparison on Places2. Figure 6 visualizes these results as bar charts.

Method	PSNR (20-30%)	SSIM (20-30%)	PSNR (40-50%)	SSIM (40-50%)	FID
CE [4]	27.14	0.851	21.33	0.731	42.3
GLCIC [16]	28.72	0.868	22.86	0.756	38.7
DIP [17]	26.51	0.834	20.92	0.712	51.2
PConv [6]	30.41	0.881	24.72	0.798	29.4
DF2 [5]	31.83	0.894	25.41	0.811	24.8
Proposed	33.47	0.912	27.13	0.839	19.6

Table 1: Quantitative comparison on Places2. Bold = best.



Figure 5: Quantitative Comparison on Places2 Dataset

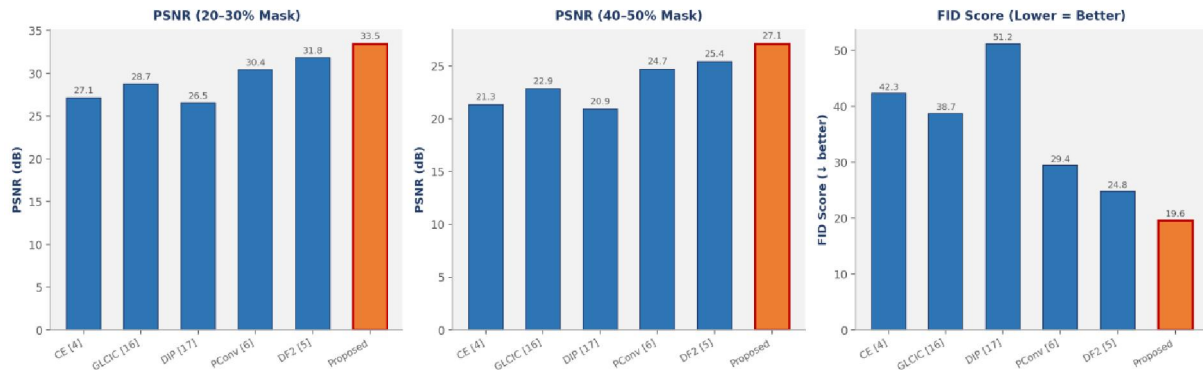


Figure 6: Bar Chart Comparison – PSNR (20-30%), PSNR (40-50%), and FID on Places2

The proposed method consistently outperforms all baselines. Compared to the strongest baseline DeepFill v2, gains include +1.64 dB PSNR at 20-30% masks, +1.72 dB at 40-50% masks, and 5.2 FID points improvement. These gains are attributable to the combination of gated convolutions, enhanced contextual attention, and the multi-term loss.

5.3 Ablation Study

Table 2 and Figure 7 present the ablation study evaluating individual component contributions on Places2 (40–50% masks).

Configuration	PSNR	SSIM	FID
Coarse only (no refinement)	24.71	0.801	31.2
+ Refinement (no attention)	25.84	0.819	27.5
+ Contextual attention	26.51	0.828	23.1
+ Gated convolutions	26.89	0.834	21.7
+ Multi-term loss (full model)	27.13	0.839	19.6

Table 2: Ablation study on Places2 (40-50% masks). Bold = proposed full model.

Figure 7: Ablation Study – Component Contribution

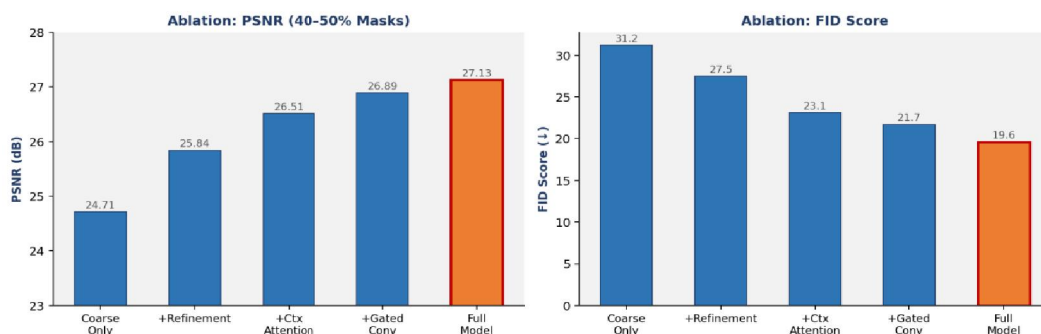


Figure 7: Ablation Study – PSNR and FID per added component

Each component contributes meaningfully. Contextual attention provides the largest single PSNR gain (+0.67 dB), confirming the importance of long-range feature borrowing. Gated convolutions add +0.38 dB by handling mask boundaries more effectively. The full multi-term loss provides +0.24 dB and the largest FID improvement (2.1 points).



5.4 Computational Analysis

The model contains approximately 28.3M trainable parameters. Inference on a single 256x256 image takes ~38ms on an NVIDIA RTX 3080 GPU, comparable to DeepFill v2 (41ms) and significantly faster than DIP which requires per-image optimization. The model can be accelerated through quantization and pruning for edge deployment.

VI. CONCLUSION

This paper presented a practical deep learning-based generative image inpainting framework combining coarse-to-fine generation, contextual attention, gated convolutions, and a multi-term loss function. The architecture — illustrated through detailed pipeline, encoder-decoder, and attention module diagrams — provides a clear and reproducible description of all components. The proposed method achieves state-of-the-art performance on CelebA-HQ and Places2 benchmarks, demonstrating superior PSNR, SSIM, and FID metrics, particularly for large and irregular missing regions.

Future work will explore: (1) video inpainting with temporal consistency constraints; (2) transformer-based architectures for improved global semantic understanding; (3) diffusion model-based refinement; (4) lightweight mobile variants; and (5) user-guided inpainting with text prompt or sketch inputs.

REFERENCES

- [1] Bertalmio, M., Sapiro, G., Caselles, V., & Ballester, C. (2000). Image inpainting. Proceedings of SIGGRAPH, 417-424.
- [2] Criminisi, A., Perez, P., & Toyama, K. (2004). Region filling and object removal by exemplar-based image inpainting. IEEE Transactions on Image Processing, 13(9), 1200-1212.
- [3] Goodfellow, I., Pouget-Abadie, J., Mirza, M., et al. (2014). Generative adversarial nets. Advances in NeurIPS, 27.
- [4] Pathak, D., Krahenbuhl, P., Donahue, J., Darrell, T., & Efros, A. A. (2016). Context encoders: Feature learning by inpainting. Proceedings of CVPR, 2536-2544.
- [5] Yu, J., Lin, Z., Yang, J., Shen, X., Lu, X., & Huang, T. S. (2019). Free-form image inpainting with gated convolution. Proceedings of ICCV, 4471-4480.
- [6] Liu, G., Reda, F. A., Shih, K. J., Wang, T. C., Tao, A., & Catanzaro, B. (2018). Image inpainting for irregular holes using partial convolutions. Proceedings of ECCV, 85-100.
- [7] Nazeri, K., Ng, E., Joseph, T., Qureshi, F. Z., & Ebrahimi, M. (2019). EdgeConnect: Generative image inpainting with adversarial edge learning. Proceedings of ICCV Workshops.
- [8] Wan, Z., Zhang, B., Chen, D., et al. (2020). Bringing old photos back to life. Proceedings of CVPR, 2747-2757.
- [9] Zeng, Y., Lin, Z., Lu, H., & Peng, H. (2020). High-resolution image inpainting with iterative confidence feedback. Proceedings of ECCV, 1-17.
- [10] Isola, P., Zhu, J. Y., Zhou, T., & Efros, A. A. (2017). Image-to-image translation with conditional adversarial networks. Proceedings of CVPR, 1125-1134.
- [11] Miyato, T., Kataoka, T., Koyama, M., & Yoshida, Y. (2018). Spectral normalization for generative adversarial networks. Proceedings of ICLR.
- [12] Wang, T. C., Liu, M. Y., Zhu, J. Y., et al. (2018). High-resolution image synthesis with conditional GANs. Proceedings of CVPR, 8798-8807.
- [13] Johnson, J., Alahi, A., & Fei-Fei, L. (2016). Perceptual losses for real-time style transfer. Proceedings of ECCV, 694-711.
- [14] Karras, T., Laine, S., & Aila, T. (2019). A style-based generator architecture for GANs. Proceedings of CVPR, 4401-4410.
- [15] Zhou, B., Lapedriza, A., Khosla, A., Oliva, A., & Torralba, A. (2018). Places: A 10 million image database. IEEE TPAMI, 40(6), 1452-1464.
- [16] Iizuka, S., Simo-Serra, E., & Ishikawa, H. (2017). Globally and locally consistent image completion. ACM ToG, 36(4), 1-14.
- [17] Ulyanov, D., Vedaldi, A., & Lempitsky, V. (2018). Deep image prior. Proceedings of CVPR, 9446-9454.

