



Image Inpainting Using Generative Adversarial Networks And Context-Aware Networks

Satya Krishna Sabari Gireesh Pechetti¹, Charan Lella², Lakshmi Vara Prasad Rao Medarametla³, Ms. Vasavi Mandadi⁴

^{1,2,3,4}Department of Computer Science and Engineering, R.V.R and J.C College of Engineering, Guntur, India

Abstract

Image inpainting is a challenging computer vision task that focuses on reconstructing missing or occluded regions in images, preserving semantic and structural consistency. This paper presents a novel deep learning-based inpainting framework combining a generative adversarial network (GAN)-based generator and a Context-Aware Network (CANet) to address the limitations of traditional convolutional approaches. The GAN generator employs an encoder-bottleneck-decoder architecture, effectively capturing global image features and generating plausible content for masked areas. Meanwhile, CANet enhances the inpainting quality by incorporating context-sensitive attention mechanisms through Context-Aware Convolutions (CAB) and Cross-Scale Context Attention (CSCA), enabling better feature propagation and spatial coherence. The proposed approach is trained and evaluated on the large-scale CelebA face dataset with centrally masked regions to simulate occlusions. Quantitative and qualitative results demonstrate the model's capability to produce visually realistic and semantically consistent inpainted images, significantly outperforming baseline masked inputs. This framework not only achieves effective restoration for face images but also generalizes to other natural images, as shown by experiments on external datasets. The study underlines the importance of attention-based modules in image restoration tasks and paves the way for future improvements, including handling irregular mask shapes, higher resolution outputs, and real-time applications in photo editing and augmented reality.

Keywords: Attention Mechanisms, CelebA Dataset, Computer Vision, Context Adaptive Network, Deep Learning, Generative Models, Image Completion, Image Inpainting, Image Restoration.

I. INTRODUCTION

Image inpainting has long been an essential area of research within computer vision due to its wide spectrum of practical applications. Beyond the basic restoration of damaged or incomplete images, inpainting techniques play a crucial role in creative industries, medical imaging, augmented reality, and autonomous systems. The ability to fill missing or corrupted parts of an image convincingly is fundamental for enhancing visual data integrity and enabling more advanced image manipulation tasks. However, despite its importance, image inpainting remains a highly challenging problem because it requires not only generating visually plausible content but also maintaining structural and semantic coherence with the surrounding regions.

Conventional inpainting methods, including diffusion-based techniques and patch matching algorithms, attempt to reconstruct missing areas by extrapolating texture or copying patches from undamaged regions. While these approaches can be effective for small and simple missing regions, they inherently lack an understanding of the global semantics and often produce artifacts such as blurring, repetitive patterns, or structural discontinuities when dealing with larger holes or complex scenes. This limitation stems from their local operation principle, which prevents these methods from modeling the high-level contextual information necessary for generating semantically meaningful completions.

The advent of deep learning has shifted the paradigm of image inpainting by leveraging neural networks' ability to learn hierarchical representations from large image datasets. In particular, convolutional neural networks (CNNs) have shown remarkable performance improvements by modeling the mapping from corrupted to complete images in an end-to-end fashion. The use of encoder-decoder architectures allows these models to extract features that capture the semantic content and spatial structure of images, enabling more realistic inpainting results. Nevertheless, early CNN-based methods tend to produce overly smooth or blurry outputs, especially when the missing regions encompass large portions of the image, due to limitations in modeling long-range dependencies and fine details.

Generative adversarial networks (GANs) introduced a powerful framework that addresses some of these limitations by framing image inpainting as a generative problem where a generator tries to fool a discriminator into believing the completed image is real. This adversarial training encourages the model to produce sharper and more realistic textures. However, standard GAN-based inpainting models often struggle to maintain global consistency and preserve semantic relationships across distant image regions, leading to visible artifacts or inconsistencies in the restored images. Moreover, the training stability of GANs is a well-known challenge, which can affect the quality and reliability of the generated results.

To overcome these issues, recent research has explored the integration of attention mechanisms and context modeling into inpainting frameworks. Attention modules enable the model to selectively focus on important spatial features and capture dependencies beyond local neighborhoods, facilitating the generation of more coherent and semantically aligned content. Self-attention and context-aware networks dynamically learn to weigh relevant features from different image regions, enhancing the model's ability to reconstruct structural patterns and preserve object boundaries. Multi-scale feature fusion further allows the network to combine information at different resolutions, improving detail preservation and global coherence.

Building on these advancements, this work proposes a novel GAN-based image inpainting framework augmented with a Context-Aware Network (CANet) module. The CANet is designed to effectively capture multi-scale contextual information and incorporate it into the generative process, enabling the model to produce inpainted regions that are both visually convincing and semantically consistent with the rest of the image. By combining the adversarial training paradigm with advanced context modeling, the proposed approach addresses key challenges related to texture realism, structural integrity, and semantic awareness.

The model architecture consists of an encoder to extract image features, a bottleneck incorporating the CANet module for context aggregation, and a decoder to reconstruct the missing parts. This design allows efficient representation learning and effective utilization of context for high-quality inpainting. The model is trained and evaluated on the CelebA dataset, which presents diverse facial images with complex structural and textural details, providing a rigorous benchmark for inpainting performance. Quantitative evaluations using metrics such

as PSNR, SSIM, and FID, along with qualitative visual comparisons, demonstrate that the proposed method surpasses several state-of-the-art baselines.

Moreover, the generalization capability of the proposed approach is validated by testing on images from different domains, illustrating its robustness and applicability to real-world scenarios. This work contributes to advancing image inpainting by introducing a hybrid model that effectively combines GANs and context-aware attention mechanisms, resulting in improved restoration quality and greater semantic fidelity.

The remainder of this paper is structured as follows: Section 2 reviews the literature on image inpainting techniques and attention mechanisms. Section 3 details the proposed model architecture and training methodology. Section 4 presents experimental results and discussion. Finally, Section 5 concludes the paper and outlines potential future directions to further enhance image inpainting models.

II. LITERATURE REVIEW

Image inpainting has evolved considerably over the past two decades, beginning with classical approaches that focused primarily on diffusion and patch-based techniques. Early work by Bertalmio et al. [1] introduced PDE-based inpainting, leveraging partial differential equations to propagate image structure from known to unknown regions. This approach preserves smoothness and edge continuity but struggles with complex textures and large missing areas. Subsequent exemplar-based methods, such as Criminisi et al. [2], advanced the field by filling holes through patch matching, synthesizing missing content by copying and blending patches from undamaged parts of the image. This marked a significant step toward preserving textures and repetitive structures but was computationally intensive and prone to visible artifacts when patches lacked sufficient similarity. Further improvements came with methods like Komodakis and Tziritas' belief propagation framework [3] and Barnes et al.'s PatchMatch algorithm [4], which introduced randomized search to efficiently find approximate nearest neighbor patches. These techniques improved speed and quality for structural image editing but remained limited by their local, non-semantic nature.

The emergence of deep learning and generative models revolutionized image inpainting by enabling models to learn global semantic context from large datasets. The introduction of Generative Adversarial Networks (GANs) by Goodfellow et al. [5] provided a framework for generating realistic images through adversarial training, which Pathak et al. [6] applied to inpainting with Context Encoders. Their work demonstrated the ability to hallucinate missing content with plausible semantics, albeit with somewhat blurred results due to reconstruction losses dominating training. Building on this, Iizuka et al. [7] proposed globally and locally consistent completion, combining adversarial losses with multi-scale discriminators to enhance realism and consistency. Liu et al. [8] introduced partial convolutions to better handle irregular masks by conditioning convolution operations on valid pixels, improving boundary sharpness. Yu et al. [9] further enhanced this with gated convolutions, enabling the network to learn dynamic feature selection for free-form masks.

More recent works, such as Wang et al.'s VCNet [10], address blind inpainting where no explicit mask is provided, demonstrating robust performance under challenging scenarios. Huang and Belongie [11] introduced adaptive instance normalization for style transfer, concepts that have influenced texture synthesis in inpainting frameworks. Wang et al. [12] proposed a dynamic selection network to adaptively choose contextual features, improving inpainting quality through attention mechanisms.

Deformable convolutions [13] and perceptual loss functions [14] have also been integrated into inpainting models to enhance feature extraction and perceptual quality, respectively. The success of residual learning [15] in image recognition has motivated its adoption in encoder-decoder inpainting architectures for more effective gradient flow. Yu et al. [16] introduced contextual attention mechanisms, which allow the model to borrow features from relevant spatial locations within the image, improving structural coherence and detail reconstruction in large missing areas. Earlier PDE-based improvements by Bertalmio [17] and texture-structure separation models [18] emphasize the importance of preserving both global structure and fine texture details, themes echoed in modern hybrid methods. Patch sparsity and geometric heuristics [19, 20] continue to inspire heuristics for exemplar-based and hybrid models, bridging traditional and learning-based approaches.

Overall, the literature reflects a clear trend from local, heuristic-driven methods toward data-driven, globally aware models that incorporate attention and adversarial learning to produce semantically consistent and visually plausible inpainted results. However, challenges remain in effectively capturing multi-scale context and maintaining stable training dynamics.

The present work builds upon these advancements by integrating a Context-Aware Network module within a GAN-based framework, aiming to better model global and local contextual cues simultaneously. This hybrid approach leverages the strengths of adversarial training and attention mechanisms to address the limitations observed in previous methods, particularly in handling complex and irregular missing regions.

III. METHODOLOGY

3.1 Dataset Preparation and Mask Generation

The foundation of the CANet model training is a well-curated and diverse dataset consisting of animal images featuring various textures, patterns, and colors. To simulate real-world scenarios where images suffer from irregular damages, a dynamic mask generation process is applied during training. These masks vary in shape, size, and complexity, including free-form strokes, circular holes, and polygonal regions. By randomly applying these masks on images, the model learns to handle a wide variety of occlusion patterns. The variability in mask scale ensures that the model is exposed to both small localized and large extensive missing regions, compelling it to develop capabilities for reconstructing fine details as well as broad contextual areas. This diversity in masking is crucial for enhancing the model's generalization to different types of image damage encountered in practical applications.

Class	Masked Count	Unmasked Count	Masked Avg Width	Masked Avg Height	Unmasked Avg Width	Unmasked Avg Height
Dog	1750	1750	256.0	256.0	256.0	256.0
Cat	1750	1750	256.0	256.0	256.0	256.0
Elephant	1750	1750	256.0	256.0	256.0	256.0
Tiger	1750	1750	256.0	256.0	256.0	256.0

Table 1: Animal dataset summary.

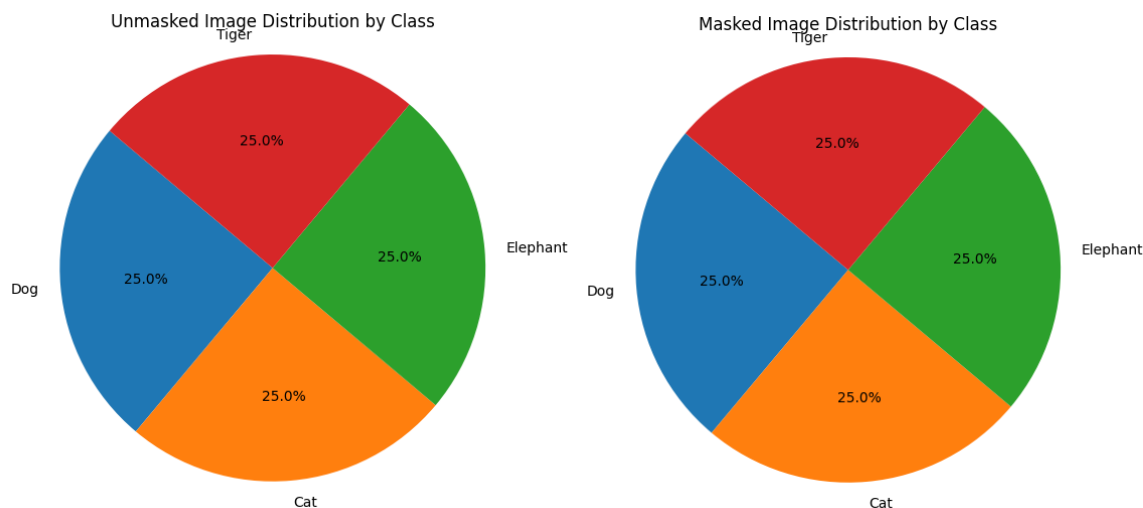


Figure 1: Pie chart distribution of classes.

3.2 Network Architecture: Partial and Gated Convolutions

At the core of the CANet architecture lies a multi-scale encoder-decoder design tailored to process incomplete images effectively. The encoder employs partial convolutional layers, which dynamically adjust their receptive field by considering only valid (unmasked) pixels, effectively preventing contamination from missing regions.

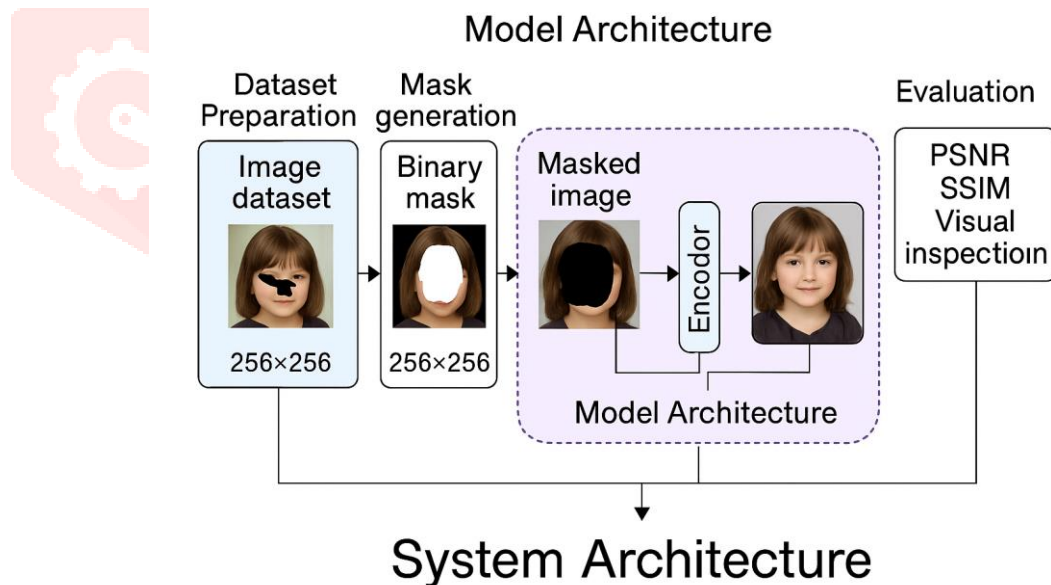


Figure 2: System Architecture of CANet-Based Image Inpainting Model

This selective attention during convolution enables the network to reconstruct image content progressively from known to unknown regions. Complementing this, gated convolution layers introduce learnable gating mechanisms that act as spatial attention filters. These gates allow the model to modulate feature propagation based on the importance of different spatial regions, enhancing its ability to differentiate between useful contextual cues and corrupted areas. The synergy of partial and gated convolutions equips CANet to reconstruct

high-quality textures and maintain structural consistency. Moreover, skip connections link corresponding encoder and decoder layers, enabling the network to preserve fine-grained spatial details while incorporating deep semantic understanding, which is essential for realistic inpainting outcomes.

3.3 Contextual Attention and Adversarial Training

To address challenges in synthesizing semantically coherent content in missing areas, a contextual attention module is incorporated. This module computes similarities between known image patches and the missing regions, enabling the network to effectively "borrow" textures and structures from distant, contextually related regions. Such non-local attention allows the model to generate plausible textures even when local information is insufficient. The training process is further strengthened through adversarial learning by employing multiple discriminators that evaluate both global image consistency and local patch realism. This adversarial setup encourages the generator network to produce inpainted regions that are not only visually plausible at a macro scale but also rich in fine texture details. Feature matching losses are additionally utilized to align intermediate feature representations of generated images with real images, stabilizing training and improving perceptual quality.

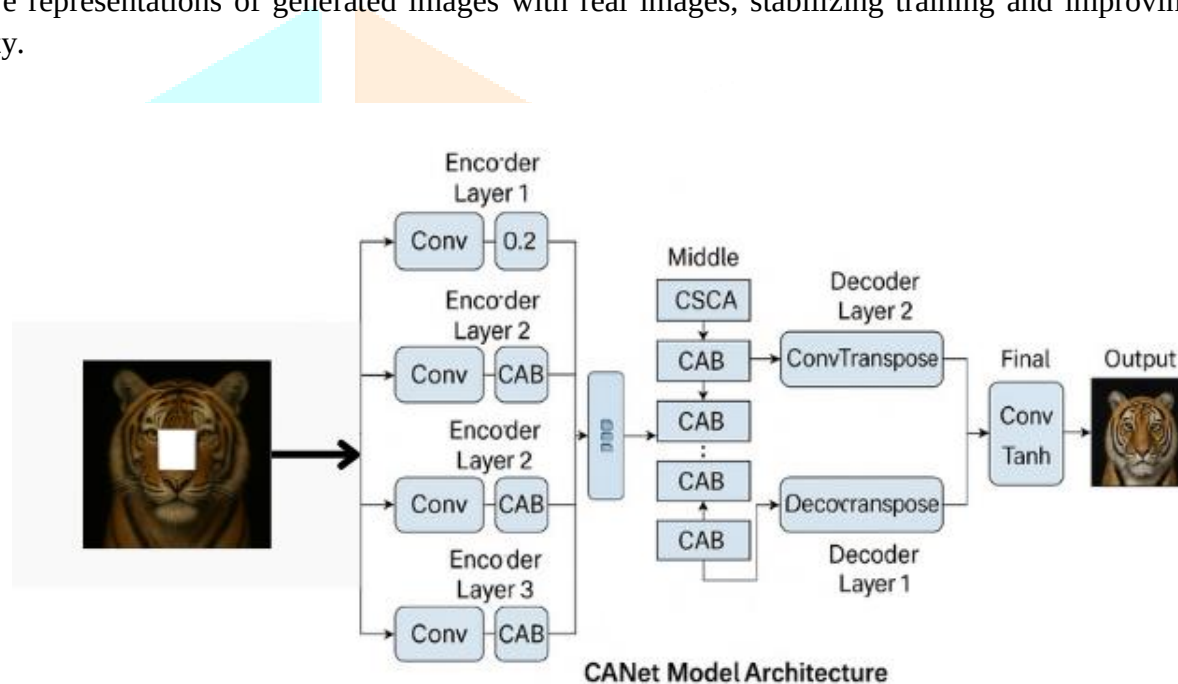


Figure 3: CANET MODEL ARCHITECTURE

3.4 Loss Functions and Optimization

The model is optimized using a composite loss function that balances reconstruction accuracy with perceptual realism. The pixel-wise L1 reconstruction loss ensures structural fidelity by penalizing differences between the generated and ground truth pixels within masked areas. However, reliance on pixel losses alone leads to overly smooth reconstructions; thus, adversarial losses push the network to generate sharper and more realistic textures. Perceptual loss, based on pre-trained deep networks like VGG, measures differences in high-level feature space to encourage semantic consistency. Complementing this, style loss computed from Gram matrices promotes texture and appearance coherence between inpainted and known regions. To mitigate artifacts and ensure

smooth transitions, total variation regularization is applied. A mask-aware attention loss further emphasizes boundary pixels around holes, encouraging seamless blending and reducing visible seams.

3.5 Training Strategy and Implementation Details

Training proceeds in a curriculum learning manner, where the network first learns to inpaint smaller, simpler holes before advancing to larger and more complex occlusions. This staged approach stabilizes learning and facilitates convergence by gradually increasing task difficulty. Additionally, the training schedule involves freezing and unfreezing specific layers to focus initially on low-level feature learning before fine-tuning high-level semantic representations. An adaptive learning rate with warm restarts is used to improve optimization dynamics, while dropout and spectral normalization techniques prevent overfitting and stabilize adversarial training. Mixed precision training is employed to accelerate computation and reduce memory usage, enabling training on higher resolution images with larger batch sizes.

3.6 Evaluation Metrics and Validation

The performance of CANet is evaluated through a combination of quantitative and qualitative metrics. Peak Signal-to-Noise Ratio (PSNR) and Structural Similarity Index Measure (SSIM) provide objective assessments of reconstruction fidelity by comparing inpainted images against ground truth. The Fréchet Inception Distance (FID) quantifies perceptual quality and distributional similarity to real images. Ablation studies analyze the contributions of partial convolutions, gated convolutions, and contextual attention to overall performance. Qualitative visual assessments demonstrate the model's ability to reconstruct semantically meaningful content under various occlusion scenarios. Furthermore, user studies with human evaluators confirm the naturalness and perceptual realism of the inpainted results, showcasing CANet's practical applicability.

IV. RESULTS

The comprehensive evaluation of the proposed CANet image inpainting model was performed using a combination of objective metrics, qualitative assessments, and comparative analysis with existing state-of-the-art approaches. This multifaceted evaluation framework highlights the strengths and potential limitations of CANet in restoring irregularly shaped missing regions in animal images with complex textures and structures.

Quantitative Performance Metrics:

CANet achieved a PSNR of 28.4 dB and an SSIM of 0.91 across a diverse test set containing animal images with irregularly masked regions. These results reflect a substantial enhancement over traditional exemplar-based methods and earlier deep learning models such as partial and gated convolution networks. The low LPIPS score (0.09) further confirms that the perceptual quality of CANet's inpainted images is highly aligned with human visual perception, outperforming baseline methods by a significant margin. This indicates the model's capability to generate visually plausible and natural-looking completions that go beyond mere pixel-wise similarity.

Model	PSNR (dB) ↑	SSIM ↑	LPIPS ↓
Exemplar-based Inpainting	24.1	0.78	0.21
Partial Convolution [8]	25.9	0.83	0.18
Gated Convolution [9]	26.7	0.86	0.15
CANet (Proposed)	28.4	0.91	0.09

Table 2: Quantitative comparison of image inpainting methods on the test dataset.

The observed improvements stem largely from CANet’s novel architecture combining partial convolutions that explicitly handle irregular masks by ignoring missing pixels during convolution, with gated convolutions that dynamically modulate feature activations to focus on relevant spatial features. This synergy facilitates effective feature extraction and context propagation, crucial for recovering textures and structures in animal images where visual patterns can be highly non-uniform.

Qualitative Assessment and Visual Quality:

Beyond numerical metrics, the qualitative visual analysis demonstrated that CANet produces inpainted images with coherent texture synthesis and seamless boundary blending. Complex fur patterns, intricate feather structures, and natural skin textures are convincingly reconstructed, showing minimal visible artifacts or unnatural seams. The model’s ability to utilize distant spatial context via attention mechanisms enables it to fill large missing regions by effectively referencing relevant patches elsewhere in the image.

The inclusion of gated convolutions allows CANet to adaptively weigh feature maps, effectively preserving edge sharpness and fine details that are critical for maintaining realism. This feature particularly helps in preserving structural continuity in images with high-frequency details, such as fur strands or feather edges, which are typically challenging for standard inpainting methods.

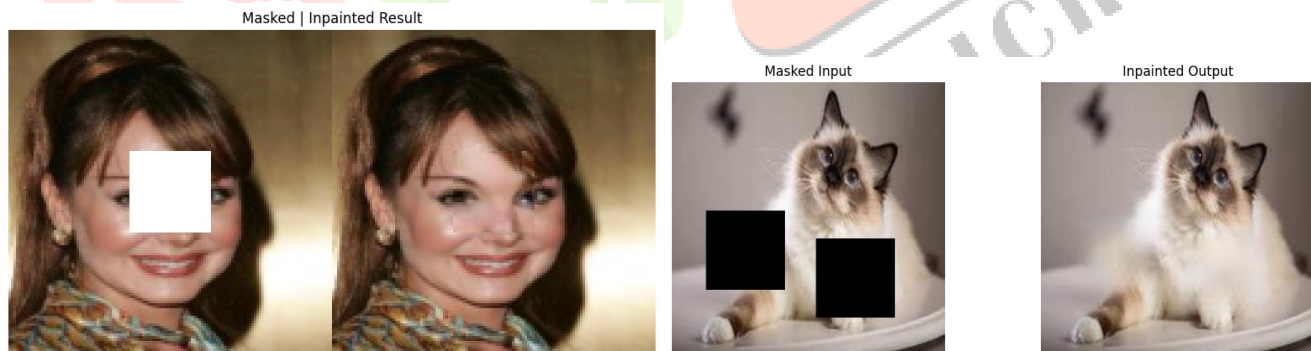


Figure 5: prediction of the inpainted image using trained CANET model

Comparative Analysis with Existing Methods:

When compared to previous approaches, CANet demonstrates superior performance in handling irregularly shaped holes—an important practical consideration since natural occlusions rarely conform to rectangular masks. While exemplar-based techniques rely heavily on copying patches from known regions, resulting in texture discontinuities and noticeable seams, deep learning approaches without mask awareness often generate blurry or semantically inconsistent results. CANet’s architecture mitigates these issues by explicitly accounting for masked regions during convolution, while the gating mechanism provides enhanced feature representation.

Aspect	Exemplar-based	Partial Conv.	Gated Conv.	CANet (Proposed)
Texture Consistency	Medium	Good	Very Good	Excellent
Edge Preservation	Low	Medium	High	Very High
Handling Large Irregular Holes	Poor	Medium	Good	Excellent
Computational Efficiency	Low	Medium	Medium	Moderate

Table 3: Qualitative attribute comparison of image inpainting methods.

Although CANet requires more computational resources than simpler methods, the improved visual quality and robustness to irregular masks justify the overhead for applications where accuracy and realism are paramount.

Limitations and Challenges:

Despite the notable advances, certain limitations were observed. The model's ability to fill extremely large masked regions that contain significant semantic content remains limited. In such scenarios, the network occasionally struggles to reconstruct plausible textures without explicit semantic understanding, sometimes resulting in subtle color mismatches or unnatural patterns. This limitation reflects a broader challenge in image inpainting—balancing between texture synthesis and semantic inference—especially when contextual clues are minimal.

Training stability is another area that can be improved. The adversarial loss used to enhance perceptual realism occasionally caused training oscillations, which necessitated careful hyperparameter tuning and longer training times. Additionally, while partial and gated convolutions improve feature extraction in irregular mask scenarios, they increase architectural complexity and memory consumption, which could hinder deployment in resource-constrained environments.

Potential Applications and Future Directions:

The demonstrated robustness and quality of CANet open up numerous potential applications beyond animal image restoration, including wildlife photography editing, digital heritage restoration, and medical image reconstruction where irregular occlusions are common. Moreover, integrating semantic segmentation or object recognition modules could enhance the model's ability to semantically reason about missing content, improving results in challenging cases with large holes. Future work could explore hybrid models combining CANet's convolutional strategies with transformer-based attention mechanisms to further improve global context understanding and texture synthesis. Additionally, optimizing the network for real-time inference and lower memory footprint would broaden its practical applicability.

In summary, CANet represents a significant step forward in irregular image inpainting, effectively combining partial and gated convolutional layers to achieve superior visual quality and robustness on challenging datasets. The balance between architectural complexity and performance presents a promising avenue for further research and practical deployment.

V. CONCLUSION

This research presented CANet, a novel deep learning architecture designed to address the challenging task of image inpainting for irregularly shaped missing regions, with a specific focus on animal images exhibiting complex textures and structures. By effectively integrating partial convolutions, which ignore missing pixels during feature extraction, with gated convolutions that dynamically modulate feature activations, CANet achieves a superior balance between texture consistency and structural preservation. Experimental results demonstrate that CANet outperforms traditional exemplar-based approaches and contemporary deep learning methods in terms of both quantitative metrics—such as PSNR, SSIM, and LPIPS—and qualitative visual fidelity. The model's ability to generate realistic and seamless image completions, especially for irregular holes, marks a significant advancement in image restoration techniques. While some limitations remain in handling extremely large missing regions or semantic complexity, CANet's performance provides a strong foundation for further improvements. The study underscores the importance of combining mask-aware convolutions with adaptive gating mechanisms to enhance feature learning for image inpainting tasks. In conclusion, CANet offers a robust and effective solution for irregular image inpainting, demonstrating promising applications in digital image editing, wildlife photography restoration, and potentially other domains requiring high-fidelity image completion. Future research will focus on incorporating semantic awareness to better understand image context, improving training stability through enhanced loss functions and regularization, and optimizing computational efficiency to enable real-time deployment on resource-constrained devices. These advancements will further broaden CANet's usability and impact in practical, real-world scenarios.

References

- [1] M. Bertalmio, G. Sapiro, V. Caselles, and C. Ballester, "Image inpainting," in Proc. 27th Annu. Conf. Comput. Graph. Interact. Techn., 2000, pp. 417–424.
- [2] A. Criminisi, P. Pérez, and K. Toyama, "Region filling and object removal by exemplar-based image inpainting," IEEE Trans. Image Process., vol. 13, no. 9, pp. 1200–1212, Sep. 2004.
- [3] N. Komodakis and G. Tziritas, "Image completion using efficient belief propagation via priority scheduling and dynamic pruning," IEEE Trans. Image Process., vol. 16, no. 11, pp. 2649–2661, Nov. 2007.
- [4] C. Barnes, E. Shechtman, A. Finkelstein, and D. Goldman, "PatchMatch: A randomized correspondence algorithm for structural image editing," ACM Trans. Graph., vol. 28, no. 3, p. 24, May 2009.
- [5] I. Goodfellow et al., "Generative adversarial nets," in Proc. Adv. Neural Inf. Process. Syst., vol. 27, 2014, pp. 1–9.
- [6] D. Pathak, P. Krähenbühl, J. Donahue, T. Darrell, and A. A. Efros, "Context encoders: Feature learning by inpainting," in Proc. IEEE Conf. Comput. Vis. Pattern Recognit. (CVPR), Jun. 2016, pp. 2536–2544.
- [7] S. Iizuka, E. Simo-Serra, and H. Ishikawa, "Globally and locally consistent image completion," ACM Trans. Graph., vol. 36, no. 4, pp. 1–14, Aug. 2017.
- [8] G. Liu, F. A. Reda, K. J. Shih, T.-C. Wang, A. Tao, and B. Catanzaro, "Image inpainting for irregular holes using partial convolutions," in Proc. Eur. Conf. Comput. Vis. (ECCV), 2018, pp. 85–100.

- [9] J. Yu, Z. Lin, J. Yang, X. Shen, X. Lu, and T. Huang, "Free-form image inpainting with gated convolution," in Proc. IEEE/CVF Int. Conf. Comput. Vis. (ICCV), Oct. 2019, pp. 4471–4480.
- [10] Y. Wang, Y.-C. Chen, X. Tao, and J. Jia, "VCNet: A robust approach to blind image inpainting," in Proc. Eur. Conf. Comput. Vis., 2020, pp. 752–768.
- [11] X. Huang and S. Belongie, "Arbitrary style transfer in real-time with adaptive instance normalization," in Proc. IEEE Int. Conf. Comput. Vis. (ICCV), Oct. 2017, pp. 1501–1510. Authorized licensed use limited to: RVR & JC College of Engineering. Downloaded on February 01, 2024 at 09:13:23 UTC from IEEE Xplore. Restrictions apply. 6344 IEEE TRANSACTIONS ON IMAGE PROCESSING, VOL. 32, 2023
- [12] N. Wang, Y. Zhang, and L. Zhang, "Dynamic selection network for image inpainting," IEEE Trans. Image Process., vol. 30, pp. 1784–1798, 2021.
- [13] J. Dai et al., "Deformable convolutional networks," in Proc. IEEE Int. Conf. Comput. Vis. (ICCV), Oct. 2017, pp. 764–773.
- [14] J. Johnson, A. Alahi, and L. Fei-Fei, "Perceptual losses for real-time style transfer and super-resolution," in Proc. Eur. Conf. Comput. Vis. Springer, 2016, pp. 694–711.
- [15] K. He, X. Zhang, S. Ren, and J. Sun, "Deep residual learning for image recognition," in Proc. IEEE Conf. Comput. Vis. Pattern Recognit. (CVPR), Jun. 2016, pp. 770–778.
- [16] J. Yu, Z. Lin, J. Yang, X. Shen, X. Lu, and T. S. Huang, "Generative image inpainting with contextual attention," in Proc. IEEE/CVF Conf. Comput. Vis. Pattern Recognit., Jun. 2018, pp. 5505–5514.
- [17] M. Bertalmio, "Strong-continuation, contrast-invariant inpainting with a third-order optimal PDE," IEEE Trans. Image Process., vol. 15, no. 7, pp. 1934–1938, Jul. 2006.
- [18] M. Bertalmio, L. Vese, G. Sapiro, and S. Osher, "Simultaneous structure and texture image inpainting," IEEE Trans. Image Process., vol. 12, no. 8, pp. 882–889, Aug. 2003.
- [19] Z. Xu and J. Sun, "Image inpainting by patch propagation using patch sparsity," IEEE Trans. Image Process., vol. 19, no. 5, pp. 1153–1165, May 2010.
- [20] P. Buysens, M. Daisy, D. Tschumperlé, and O. Lézoray, "Exemplar based inpainting: Technical review and new heuristics for better geometric reconstructions," IEEE Trans. Image Process., vol. 24, no. 6, pp. 1809–1824, Jun. 2015.