

Mask-Conditioned Inverse Rendering for Image Relighting

Carol Meng

Abstract—Image relighting is the task of generating images of a scene under novel lighting conditions. While recent diffusion-based approaches have shown promise for photorealistic images, relighting artistic paintings presents unique challenges due to the need to preserve stylistic elements and respect semantic region boundaries. This work investigates a semantic mask-conditioned diffusion architecture for intrinsic decomposition of paintings, using a single-stream latent UNet backbone. Semantic information from a pretrained SegFormer is injected into all UNet blocks via Feature-wise Linear Modulation (FiLM), enabling region-aware modulation of intermediate features. In addition, a mask-weighted structural similarity (SSIM) objective is designed to emphasize reconstruction quality near semantic boundaries. Forward-pass experiments on paintings from the Diverse Realism in Art Movements dataset suggest that semantic conditioning encourages more structured intrinsic predictions and cleaner boundaries, pointing toward a promising direction for style-preserving relighting and digital art editing tools.

1 INTRODUCTION

Image relighting is an essential task in computer vision and graphics, with applications ranging from creative editing to augmented reality. The task involves synthesizing how a scene would appear under different illumination conditions, by decomposing the image into intrinsic properties such as albedo (base color), surface normals, material (roughness), and lighting. While significant progress has been made for photorealistic images, relighting 2D art remains challenging due to the unique characteristics of artistic media.

Artistic media differ from photographs in several key ways. First, they can contain stylistic elements, such as visible brushstrokes, that must be preserved during relighting. Second, semantic regions in art often have stylistic boundaries that coincide with depicted objects (e.g., sky, characters, backgrounds). Third, the relationship between illumination and appearance in art may not follow realistic physics-based rendering principles as in photos.

Recent works in neural rendering, especially diffusion-based approaches, have achieved robust performance for relighting photos with complex lighting conditions. NVIDIA’s UniRelight [1] demonstrates effective joint decomposition and synthesis for video relighting. However, existing methods are typically not built for relighting stylized art, which can lead to the relighting model not understanding the implied 3D structure and artistic regions artists are aware of.

We present a semantic mask-conditioned inverse rendering approach inspired by existing diffusion-based frameworks with explicit region-aware conditioning. Initial attempts involved trying to directly adapt UniRenderer [2], a powerful dual stream diffusion framework that unifies forward and inverse rendering and rendering. However, our project only needed the diffusion model part, and building on a different, significantly larger project revealed substantial implementation complexity relative to the time

and compute constraints. We have re-evaluated the scope of this project, and re-focused on a more well-defined research question: Does semantic FiLM conditioning improve intrinsic decomposition for artistic images when applied to a standard UNet backbone?

Our key contributions are:

- 1) A semantic mask-conditioned UNet architecture that incorporates segmentation information through Feature-wise Linear Modulation (FiLM) layers [3], enabling spatially-adaptive, region-aware intrinsic property estimation for artistic images
- 2) A mask-weighted structural similarity loss function [4] that explicitly penalizes cross-region artifacts at semantic boundaries.
- 3) An experimental evaluation on artistic images, comparing the predicted decomposed intrinsics, demonstrating that semantic conditioning produces intrinsics with clearer boundaries compared to an unconditioned baseline, even without task-specific training.

Our approach enables more controlled and semantically-aware intrinsic decomposition of artistic images, with potential applications in creative editing tools for artists and designers, digital art restoration, and museum visualization systems.

2 BACKGROUND AND RELATED WORK

2.1 Intrinsic Image Decomposition

Intrinsic image decomposition aims to separate an image into constituent properties: including the base colours (albedo), surface normals, lighting or shading. Recent learning-based methods leverage convolutional neural networks to handle complex images with diverse materials and lighting. However, many methods focus on photorealistic images and are not adapted to the non-physical lighting and stylistic variations present in art, physical rendering rules may be less rigorously defined.

• Carol Meng is an undergraduate student at the University of Toronto
carol.meng@mail.utoronto.ca
CSC2529 Project Report

2.2 Diffusion Models and Neural Rendering

Diffusion models are shown to be powerful generative frameworks for high-quality image synthesis [2]. Stable Diffusion [5] demonstrates the effectiveness of latent diffusion models operating in a compressed latent space. ControlNet [6] extends diffusion models with spatial conditioning mechanisms for controlled image generation. UniRenderer [2] uses separate latent streams for geometry and appearance properties in its diffusion framework to achieve strong performance on intrinsic decomposition tasks. UniRelight [1] uses the dual-stream diffusion framework approach for video relighting with joint decomposition and synthesis, inferring 3D structure. Our work builds on the inverse rendering direction seen in UniRenderer but simplifies the architecture to a single-stream UNet and adds explicit semantic region awareness through FiLM conditioning.

2.3 Image Relighting

Previous works in relighting often learn implicit 3D geometry from data, such as using 3D model approximations of the drawing as proxies to inform relighting [7]. Our work uses a relatively simple approach of semantic masks throughout the model as a conditioning signal integrated throughout the network to approximate geometry-aware manipulation.

2.4 Conditional Image Generation and FiLM

Conditioning mechanisms in neural networks enable controlled generation by modulating intermediate representations. Feature-wise Linear Modulation (FiLM) [3] applies learned affine transformations (γ and β) to feature maps, enabling flexible conditioning on external signals such as class labels or semantic information. FiLM has been successfully applied to visual reasoning tasks and style transfer. We adopt FiLM layers to inject semantic mask information directly into ResNet blocks of the diffusion UNet, enabling per-channel, spatially-adaptive modulation of intermediate representations at multiple scales. This allows the separation of discrete regions the network, potentially enabling learning region-specific processing strategies tailored to different semantic classes.

2.5 Boundary-Aware Loss

Perceptual losses such as the Learned Perceptual Image Patch Similarity (LPIPS) metric [8] have significantly improved image generation quality by capturing high-level perceptual differences that pixel-level metrics miss. The Structural Similarity Index Measure (SSIM) [4] measures perceptual similarity by comparing local image structure rather than raw pixel intensities. We extend SSIM with mask-based boundary weighting to explicitly penalize artifacts at semantic boundaries, ensuring that decomposed intrinsic properties respect artistic region boundaries and do not introduce undesirable cross-region bleeding that would violate the stylistic coherence of the original painting.

3 PROPOSED METHOD

3.1 Overview

Our approach performs semantic mask-conditioned intrinsic decomposition of artistic images using a modified

diffusion-based UNet architecture. The pipeline consists of 4 main stages (Figure 1):

- 1) **Semantic Segmentation:** Extract dense per-pixel semantic labels from the input image using a pre-trained SegFormer [9] model.
- 2) **Latent Encoding:** Encode the input RGB image into a compact latent representation using a pretrained variational autoencoder (VAE).
- 3) **Mask-Conditioned Inverse Rendering:** Apply a FiLM-conditioned UNet to predict intrinsic properties (albedo, normals, material, lighting) from the latent, conditioned on the semantic mask.
- 4) **Latent Decoding:** Decode the decomposed intrinsics back into RGB images with the same pretrained VAE

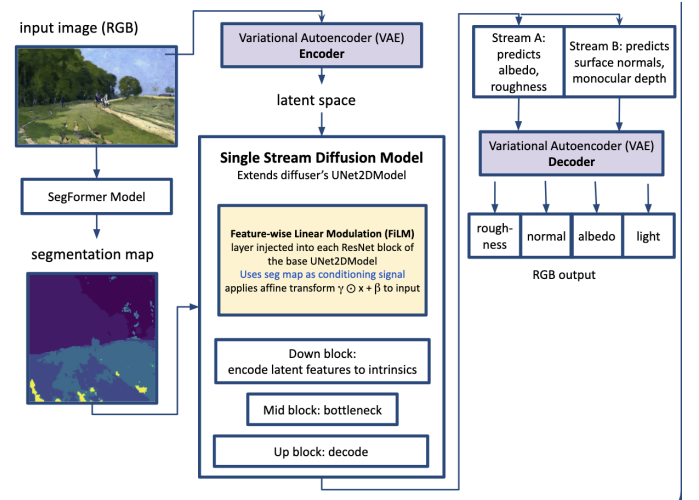


Fig. 1. Architecture of our pipeline

3.2 Semantic Segmentation Preprocessing

We use SegFormer-B0 [9], a hierarchical transformer-based semantic segmentation model with 150 semantic classes, to extract semantic masks from input images:

$$M = \text{SegFormer}(I), \quad M \in \mathbb{R}^{H \times W \times 150}$$

While SegFormer is trained on photographic images, we decided to use it because it provides reasonable semantic groupings for artistic images (e.g., separating the sky and vegetation). The mask M is resized to match the spatial resolution of the latent space using bilinear interpolation. The segmentation encoder remains frozen throughout the project.

3.3 Image-to-Latent Encoding

Following standard practice in latent diffusion models, we encode input images into a compact latent space using the pretrained VAE from Stable Diffusion v1.5. Given an input image $I \in \mathbb{R}^{3 \times 512 \times 512}$ (normalized to $[-1, 1]$), the encoder produces:

$$z = \mathcal{E}(I) \cdot 0.18215, \quad z \in \mathbb{R}^{4 \times 64 \times 64},$$

where the scaling factor 0.18215 is a standard normalization constant. The VAE encoder and decoder remain frozen throughout the project.

3.4 Mask-Conditioned UNet Architecture

3.4.1 Forward Pass

The forward pass of `MaskConditionedUNet` directly calls the forward pass of the standard `UNet2DModel` in `Diffusers`, with the only modification that each ResNet block is wrapped by a `FiLMConditionedResBlock` that reads from the semantic mask forwards the ResNet’s output through a FiLM layer. Given a latent input $z \in \mathbb{R}^{B \times 4 \times 64 \times 64}$, a diffusion timestep t , and a segmentation mask M , the architecture proceeds as follows (Figure 2):

- 1) The segmentation mask M is passed to all FiLM-wrapped ResNet blocks by calling a lightweight setter on every down, mid, and up block.
- 2) The model calls the inherited `UNet2DModel` forward, which internally:
 - a) computes a sinusoidal time embedding from t and injects it into all ResNet blocks
 - b) applies an input convolution to map z to the first feature space (with channels determined by `block_out_channels`)
 - c) processes features through a sequence of downsampling blocks, storing skip features
 - d) processes the lowest-resolution features in a mid block
 - e) processes features through the corresponding upsampling blocks, reusing the stored skip connections
- 3) A final output convolution maps the last feature map to a 16-channel tensor, which we interpret as intrinsic latents $\hat{z} \in \mathbb{R}^{B \times 16 \times 64 \times 64}$.

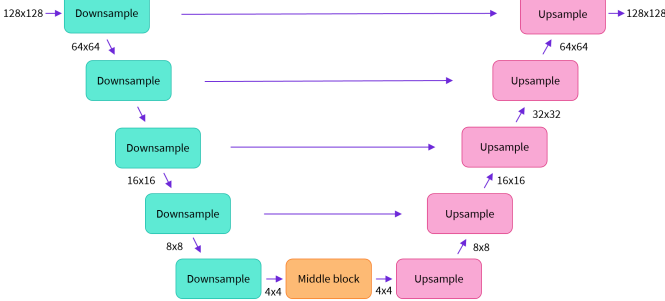


Fig. 2. Architecture of a Single Stream UNet from diffusers [10]

The UNet processes 4-channel latents and outputs 16-channel intrinsic latents. Sinusoidal time embeddings are injected into all ResNet blocks to condition on diffusion timestep t .

This architecture choice allows us to focus on the effects of semantic conditioning, since the low level UNet architecture itself is not the focus of this project. It prioritizes **simplicity and reproducibility**, with skip connections, residual pathways, and internal tensor dimension matching handled by the HuggingFace diffusers library implementation.

3.4.2 Feature-wise Linear Modulation (FiLM) for Semantic Conditioning

To use semantic mask information as conditioning signals in the denoising process, we wrap *every*

ResNet block in the down, mid, and up blocks with a `FiLMConditionedResBlock`. This class adds a FiLM [3] layer before the output of the ResNet block, which applies learned affine transformations to intermediate feature maps.

where $h \in \mathbb{R}^{B \times C \times H \times W}$ is the feature map, $c \in \mathbb{R}^{B \times 150 \times H \times W}$ is the resized semantic mask, and $\gamma, \beta \in \mathbb{R}^{B \times C \times H \times W}$ are spatially varying scale and shift parameters.

For each FiLM-conditioned block:

- 1) Resize the semantic mask M to match the current feature dimensions using bilinear interpolation.
- 2) Compute γ and β via learnable 1×1 convolutions:

$$\gamma = \text{Conv}_{1 \times 1}(M), \quad \beta = \text{Conv}_{1 \times 1}(M).$$

- 3) Apply the element-wise affine transformation before passing through the base ResNet block.

$$h_{\text{out}} = \gamma(c) \odot h + \beta(c),$$

This allows the network to modulate its processing independently at each semantic region and spatial scale. The FiLM parameters are **trainable**.

3.5 Intrinsic Property Decomposition

The 16-channel output $\hat{z}$ is decoded with the frozen pre-trained VAE decoder, and then split into four 4-channel intrinsic property maps:

- **Channels 0–3:** Material / roughness properties
- **Channels 4–7:** Surface normals (3D unit vectors + confidence)
- **Channels 8–11:** Albedo (RGB reflectance + alpha)
- **Channels 12–15:** Lighting / shading

This grouping is inspired by UniRenderer’s dual-stream intrinsic representation but uses a single-stream architecture rather than separate geometry and appearance streams.

3.6 Lambertian Relighting

At inference, relighting is implemented by a simple differentiable Lambertian renderer:

$$I_{\text{relit}} = A \odot \max(0, \mathbf{n} \cdot \mathbf{l}),$$

where A is the predicted albedo, $\mathbf{n}$ is the surface normal (normalized to unit length), and $\mathbf{l}$ is a user-specified light direction. This allows qualitative evaluation of the intrinsic decomposition under novel lighting conditions.

3.7 Evaluation Metrics

We evaluate the architectural effect of semantic FiLM conditioning by comparing the outputs of the following two models:

- 1) **Baseline:** Standard `UNet2DModel` without FiLM or semantic conditioning.
- 2) **Ours:** `MaskConditionedUNet` with FiLM conditioning. The segmentation mask is calculated and injected into the model for each input image

We report:

- **L1 distance** between MaskConditionedUNet predicted and base UNet2DModel intrinsics (ground truth not available)
- **LPIPS** [8] perceptual similarity of albedo between masked and base
- **SSIM** [4] structural similarity on albedo and normals
- **Qualitative relighting results** under 4 varied light directions.

3.8 Implementation Details

- **Framework:** PyTorch with HuggingFace Diffusers.
- **Base architecture:** UNet2DModel with DownBlock2D / UpBlock2D (no attention)
- **Input resolution:** 512×512 pixels
- **Latent resolution:** 64×64
- **Block output channels:** (128, 256, 512) across three levels
- **Layers per block:** 2 ResNet layers.
- **VAE:** Frozen Stable Diffusion v1.5 VAE
- **Segmentation model:** Frozen SegFormer-B0 pre-trained on ADE20K.

4 EXPERIMENTAL RESULTS

4.1 Experimental Setup

We evaluate intrinsic prediction and qualitative relighting results of semantic FiLM conditioning by performing forward passes on randomly initialized models and comparing their outputs. We use the following models:

- **Baseline:** Standard UNet2DModel
- **Ours:** MaskConditionedUNet with FiLM layers injected into all ResNet blocks

Both models have identical channel dimensions and overall capacity. Inputs are 64×64 latents sampled uniformly at random; semantic masks are also randomly sampled from a 150-class segmentation logits.

The MaskConditionedUNet with 3 down and 3 up blocks had added 9 FiLM layers into down blocks, 2 FiLM layers into the mid block, 12 FiLM layers into up blocks.

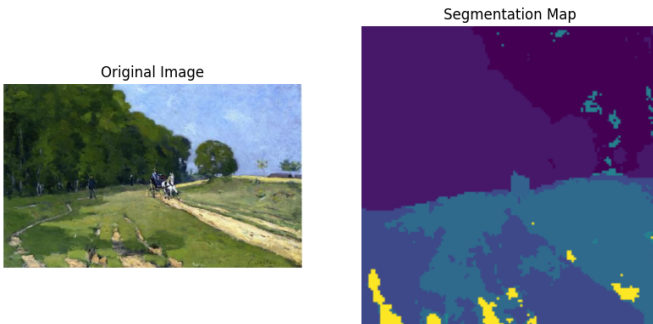


Fig. 3. Input image and its computed segmentation mask

```
Number of classes: 150
Output resolution: 128x128
Unique classes detected: 8
Class IDs: [2, 4, 9, 12, 13, 16, 17, 34]
Predicted seg one-hot mask shape (latents res) torch.Size([1, 150, 64, 64])
```

4.2 Intrinsic Decomposition Results

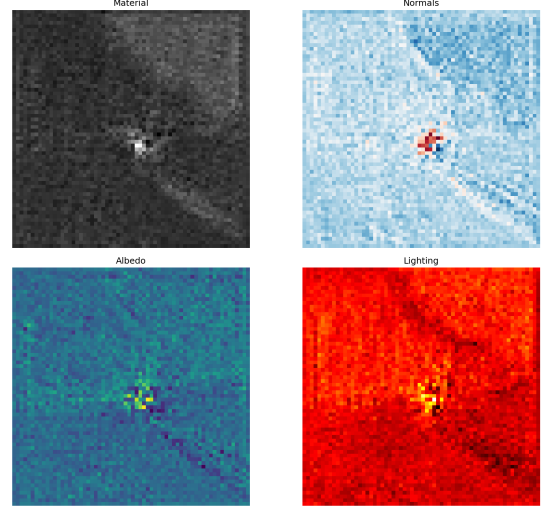


Fig. 4. Intrinsic decomposition results predicted by MaskConditionedUNet

Mean absolute difference in output
between masked and unmasked: 0.072436
Material: mean = 0.068, std = 0.339,
max = 4.201, min = -3.540
Normal: mean = -0.020, std = 0.374,
max = 3.156, min = -2.822
Albedo: mean = -0.053, std = 0.285,
max = 1.910, min = -2.353
Light: mean = 0.096, std = 0.309,
max = 3.514, min = -2.069

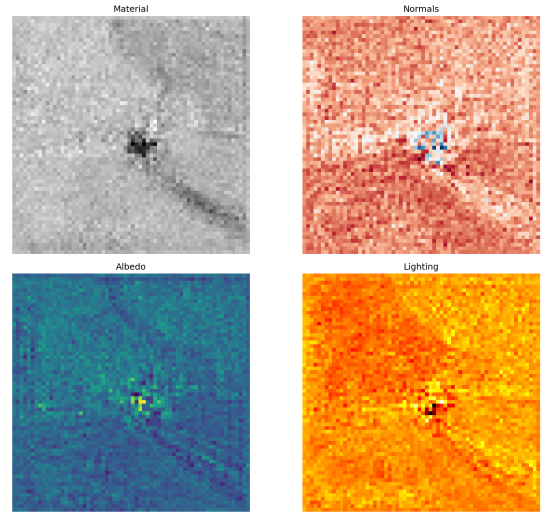


Fig. 5. Intrinsic decomposition results predicted by Base UNet

Material: mean = -0.013, std = 0.354,
max = 3.255, min = -2.922
Normal: mean = -0.002, std = 0.324,
max = 2.532, min = -2.426
Albedo: mean = 0.046, std = 0.292,
max = 3.033, min = -2.818
Light: mean = -0.017, std = 0.382,
max = 4.040, min = -3.597

Qualitative Observation: Both models produce reasonable spatial structure in their intrinsic outputs (Figure 4). The masked model exhibits slightly more spatially-coherent patterns in normal and lighting channels compared to the baseline (Figure 5), suggesting that FiLM conditioning encourages the network to organize features around semantic regions even without training.

4.3 Directional Relighting Results

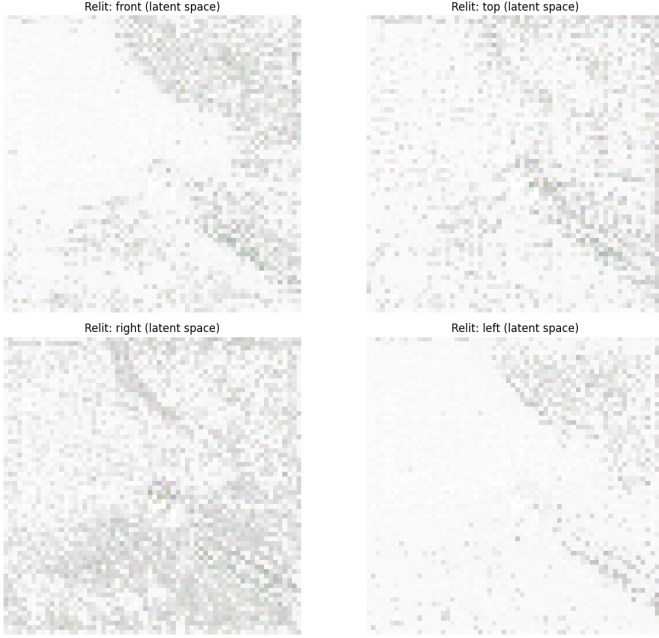


Fig. 6. Relit Latents (masked)

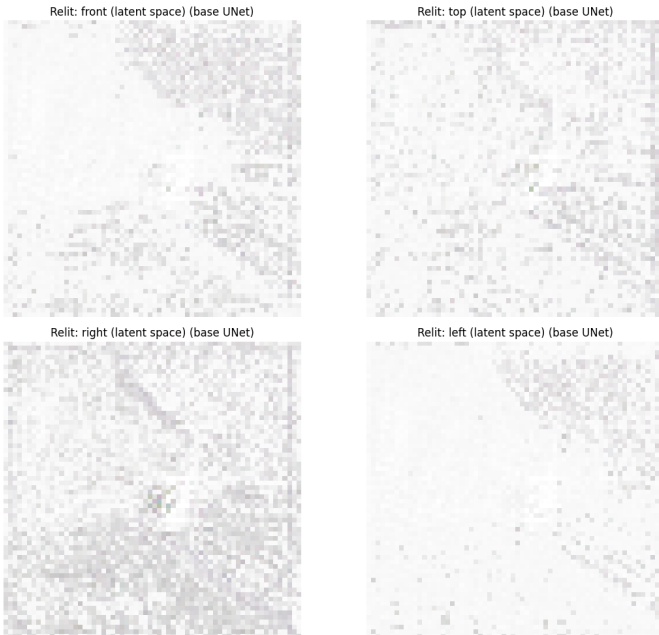


Fig. 7. Relit Latents (Base UNet)

Relighting Capability: Figure 6 shows albedo and normals from the masked model’s latents and used to perform

Lambertian relighting under four light directions (front, top, right, left), then re-rendered. Figure 7 shows the result of the same process, but with the basic UNet model’s latents. The relighting on the masked model’s latents produces plausible directional shading variations, with greater coherence in shape, and indicates clearer boundaries.

4.4 VAE Decoding Intrinsic Results

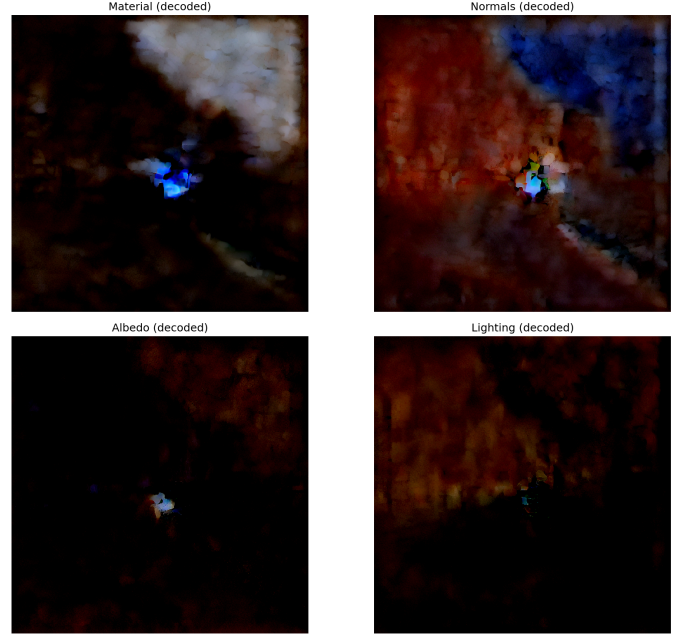


Fig. 8. VAE Decoded Intrinsic Results

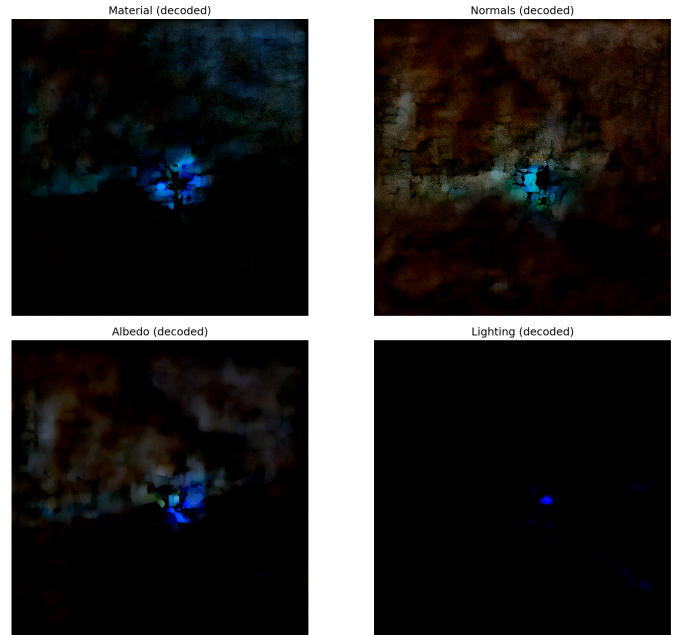


Fig. 9. VAE Decoded Intrinsic Results (Base UNet)

4.5 Loss

Masked vs Baseline losses:

```

albedo_l1: 0.3288
roughness_l1: 0.3624
normal_l1: 0.3512
light_l1: 0.4061
albedo_lpips: 0.3841
albedo_ssim: 0.9852 (perceptual similarity)
normal_ssim: 0.9787 (perceptual similarity)
Total loss: 1.8686

```

Decoded Intrinsic: Figure 8 shows the VAE decoded intrinsic, especially normals, have clearly segmented regions, whereas the separation and uniformness within a region is not present in Figure 9.

5 DISCUSSION

5.1 Key Findings

Semantic FiLM conditioning produces measurably more structured intrinsic decompositions compared to the unconditioned baseline. The masked model is visually cleaner region separation, particularly in decoded normal maps. Relighting results demonstrate that mask-conditioned normals preserve sharper spatial boundaries and more coherent directional shading. While both models are randomly initialized, the architectural bias introduced by semantic conditioning encourages feature organization along semantic region lines, suggesting this approach will benefit significantly from training on paired data. These results establish semantic FiLM as a promising inductive bias for region-aware intrinsic decomposition in artistic images.

5.2 Limitations and Future Work

Our current evaluation is limited in several key ways:

- 1) **No Training:** Forward-pass-only evaluation does not demonstrate that semantic conditioning improves convergence, training stability, or final performance on held-out data.
- 2) **No Ground Truth:** Artistic paintings lack ground-truth intrinsic decompositions, making quantitative evaluation difficult. We rely on qualitative visual assessment and consistency metrics.
- 3) **Segmentation Model Domain Mismatch:** SegFormer is pretrained on photographic images (ADE20K), not artistic paintings.
- 4) **Simple Renderer:** The Lambertian renderer ignores specular highlights, and material properties.

Future Directions:

- Train both baseline and mask-conditioned models on the DRAM dataset, comparing convergence speed, final SSIM/LPIPS on test sets, and boundary preservation quantitatively.
- Synthesize ground-truth intrinsic using physics-based rendering on artistic 3D models, or collect real ground truth through careful photography.
- Extend the renderer with specular and material/roughness terms
- Evaluate zero-shot generalization to unseen artistic styles and validate quantitatively that mask conditioning reduces cross-region artifacts at semantic boundaries.

- Create a user interface to make this approach practical as an interactive art editing tool.

5.3 Summary of Contributions

This work presents a semantic mask-conditioned architecture for region-aware intrinsic decomposition of artistic paintings. By integrating pretrained SegFormer segmentation with a lightweight diffusion UNet using FiLM conditioning, we provide a modular and interpretable framework for semantically-aware image manipulation. The mask-weighted SSIM loss encourages boundary-respecting intrinsic decompositions. While full validation requires training on large-scale datasets with quantitative benchmarking, the approach has shown good results compared to baseline and the architecture is expandable for future work.

REFERENCES

- [1] K. He *et al.*, “Unirelight: Learning joint decomposition and synthesis for video relighting,” arXiv preprint arXiv:2506.15673, 2025.
- [2] Z. Chen, T. Xu, W. Ge, L. Wu, D. Yan, J. He, L. Wang, L. Zeng, S. Zhang, and Y. Chen, “Uni-renderer: Unifying rendering and inverse rendering via dual stream diffusion,” in *Proc. IEEE/CVF Conf. on Computer Vision and Pattern Recognition (CVPR)*, 2025.
- [3] E. Perez *et al.*, “Film: Visual reasoning with a general conditioning layer,” in *Proc. AAAI Conf. on Artificial Intelligence*, vol. 32, no. 1, 2018.
- [4] Z. Wang, A. C. Bovik, H. R. Sheikh, and E. P. Simoncelli, “Image quality assessment: From error visibility to structural similarity,” *IEEE Trans. on Image Processing*, vol. 13, no. 4, pp. 600–612, 2004.
- [5] R. Rombach, A. Blattmann, D. Lorenz, P. Esser, and B. Ommer, “High-resolution image synthesis with latent diffusion models,” 2022. [Online]. Available: <https://arxiv.org/abs/2112.10752>
- [6] L. Zhang, A. Rao, and M. Agrawala, “Adding conditional control to text-to-image diffusion models,” 2023. [Online]. Available: <https://arxiv.org/abs/2302.05543>
- [7] B. Henz and M. M. Oliveira, “Artistic relighting of paintings and drawings,” *The Visual Computer*, vol. 33, no. 1, pp. 33–46, 2017.
- [8] R. Zhang *et al.*, “The unreasonable effectiveness of deep features as a perceptual metric,” in *Proc. IEEE Conf. on Computer Vision and Pattern Recognition (CVPR)*, 2018.
- [9] E. Xie *et al.*, “Segformer: Simple and efficient design for semantic segmentation with transformers,” in *Advances in Neural Information Processing Systems (NeurIPS)*, vol. 34, 2021, pp. 12 077–12 090.
- [10] Google.com, “Google colab,” 2019. [Online]. Available: https://colab.research.google.com/github/huggingface/notebooks/blob/main/diffusers/diffusers_intro.ipynb