

\_\_\_\_\_

# DINO

**VAE**

## Diffusion Models

## 1D Tokenization

The collage consists of 12 distinct AI-generated images arranged in a 3x4 grid. The top row features a vibrant mushroom forest under a full moon, a pink unicorn sitting at a desk with a laptop, and a colorful, wavy abstract background with the text 'Go with the flow' in a cursive font. The middle row includes a wizard in a purple robe with the text 'Achievement unlocked: Diffusion models can spell now!', a profile of a human head composed of many small, colorful images, and a white truck with a sign that reads 'They left me here to rot and not away like the others'. The bottom row shows a person in a white space suit floating in blue water, a close-up of a tomato-like creature with a green stem, a ginger cat sitting next to a blue cube with a red ball on top, a black dog lying down, and a crowned avocado sitting on a throne surrounded by golden scepters.

As discussed in the paper Titok, 1D tokenization of images has great potentials for fast image generation because of its nature of small amount of encoding tokens.

Recently, multiple works have tried to add semantic information into the VAE encoding space, such methods have shown great results of improving the learning speed of diffusion models.

We summarize our contribution as:

- The diagram illustrates three different architectures for video reconstruction, each starting with an input image of a person and ending with a reconstructed image.

  - VAE (Variational Autoencoder):** The input image is processed by an **Encoder** to produce a latent representation (a 3D cube). The **Tensor Size** is [H/f: 32, W/f: 32, D: 16] (assuming f=8). This latent representation is then processed by a **Decoder** to produce the reconstructed image. The **Tensor Size** for the reconstructed image is [3, 256, 256].
  - TiTok:** The input image is processed by an **Encoder** to produce a latent representation (a 3D cube). The **Tensor Size** is [K: 64, D: 64]. This latent representation is then processed by a **Decoder** to produce the reconstructed image. The **Tensor Size** for the reconstructed image is [3, 256, 256].
  - Ours:** The input image is processed by a **Fine-Tuned DINOv3** encoder to produce a latent representation (a 3D cube). The **Tensor Size** is [K: 64, D: 64]. This latent representation is then processed by a **Decoder** to produce the reconstructed image. The **Tensor Size** for the reconstructed image is [3, 256, 256]. Additionally, the **Ours** architecture includes a **CLS Token, Patch Token...** block that receives input from the **Fine-Tuned DINOv3** encoder and outputs to the **Decoder**. A **Trash** icon indicates that the **Ours** architecture is the preferred method.

Our goal of this project is to train a 1D tokenizer with constraints of both detail reconstruction and semantic preservation. To achieve this goal, we fine tune a pretrained DINOv3 ViT as our encoder, and we design our training objective to preserve the rich semantic prior of DINO. In detail, we use a combination of the following 5 losses:

- **Pixel Level L1 loss:** Measures the mean absolute difference between reconstructed and ground-truth pixels
- **MS-SSIM loss:** Measures structural similarity across multiple scales (luminance, contrast, structure). Focuses on perceptual quality and texture consistency
- **LPIPS Loss:** Measures perceptual similarity from deep feature embeddings (e.g., VGG). Encourages semantic and texture-level similarity

- **Patch Similarity Loss:** Measures the L2 difference of the original and fine tuned DINO's image patch encodings.
- **CLS Probing Loss:** Use an additional MLP + pooling layer with random dropouts, to reconstruct the DINO CLS from the latent tokens, measures the L1/L2 loss on the reconstructed token and the original DINO's CLS token to encourage semantic information to be evenly distributed on all latent representations.



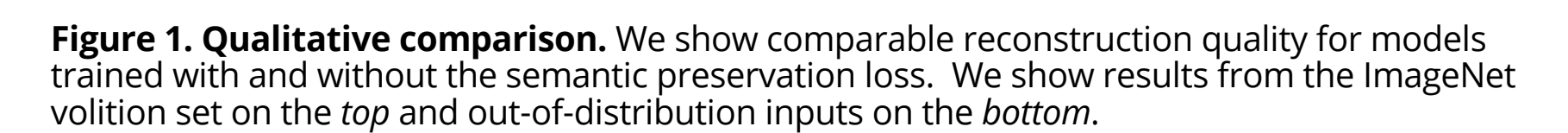
**Dataset:** We train our model on the train split of ImageNet, which contains 1.2M samples, and validate using its val split. We crop the images into square and resize them into  $256 * 256$  pixels using pillow’s implementation of BICUBIC.

**Model Architecture:** For the backbone model, we chose `dinov3-vitb16-pretrain-lvd1689m` for a balance of quality and compute. The adapter is a simple linear module.

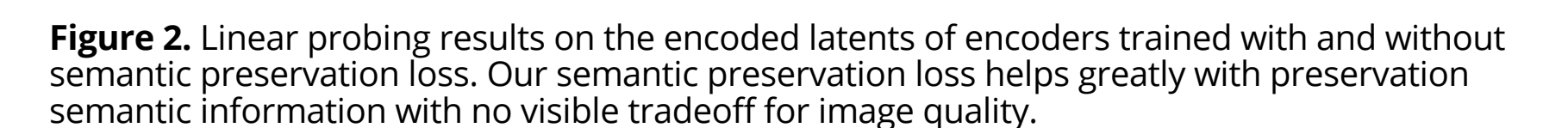
**Training Settings:** We train on an 8 A40 GPU node with DDP. Using an effective batch of 640. We train 2 copies with and without semantic preservation loss and kept every other parameters identical.

**Linear Probing Evaluation:** We employ a linear probing protocol following MAE and TiTok. The classifier consists of a single affine-free Batch Normalization layer followed by a linear projection to 1,000 ImageNet classes. This minimal architecture ensures that observed accuracy reflects the intrinsic quality of the token representations rather than any capacity in the classifier itself.

### Reconstruction Quality



## Semantic Preservation



- Add GAN discriminators for better reconstruction results
- Test the diffusibility of our proposed method by training a diffusion model based on this latent space