



DDiT: Dynamic Patch Scheduling for Efficient Diffusion Transformers

CVPR 2026 Highlight

Dahye Kim^{1,2} Deepti Ghadiyaram² Raghudeep Gadde¹

¹Amazon ²Boston University

{dahye, dghadiya}@bu.edu raghudeep.g@gmail.com

Inference in Diffusion Models is Expensive 💰

- Generating a 5-second 720p video with Wan-2.1 takes ~**30 minutes***.

Inference in Diffusion Models is Expensive 💰

- Generating a 5-second 720p video with Wan-2.1 takes **~30 minutes***.
- Reasons:
 - Diffusion models use an iterative denoising process
 - Each step requires a full forward pass over all **spatial tokens**



Noisy image

Less noisy image

Even less noisy image

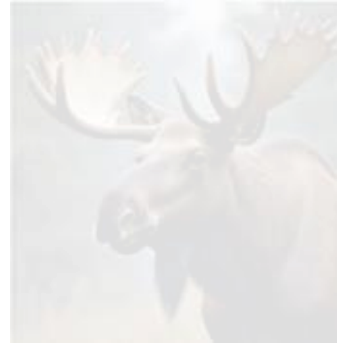
Generated image

1. Should every image be encoded at the same granularity?

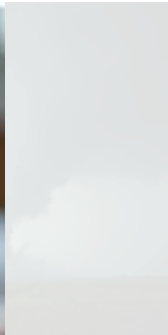
Images have very different semantic complexity



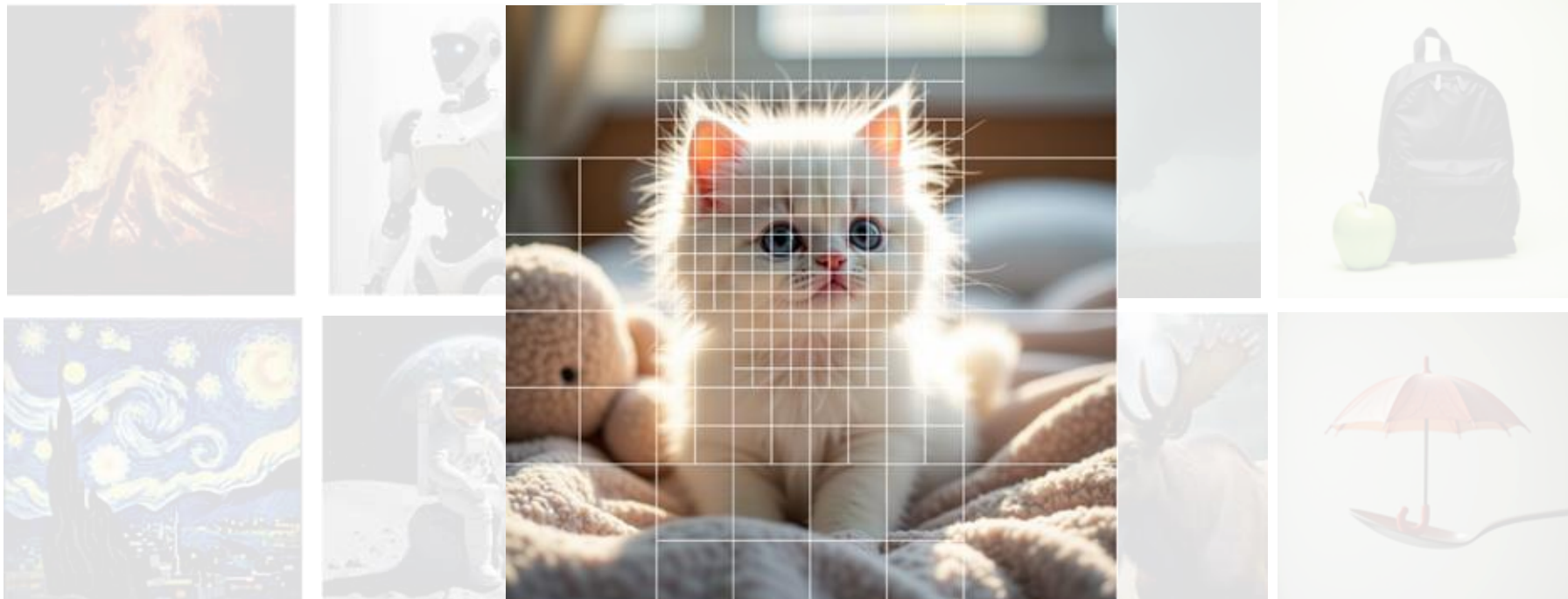
Images have very different semantic complexity



Images have very different semantic complexity



Images have very different semantic complexity



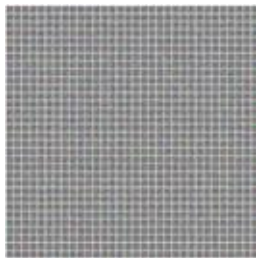
Takeaway: Different regions in the same image can be encoded at different spatial granularity.

1. *Should every image be encoded at the same granularity?*



Not necessarily...

2. Should every denoising step process the image at the same granularity?



Noisy image

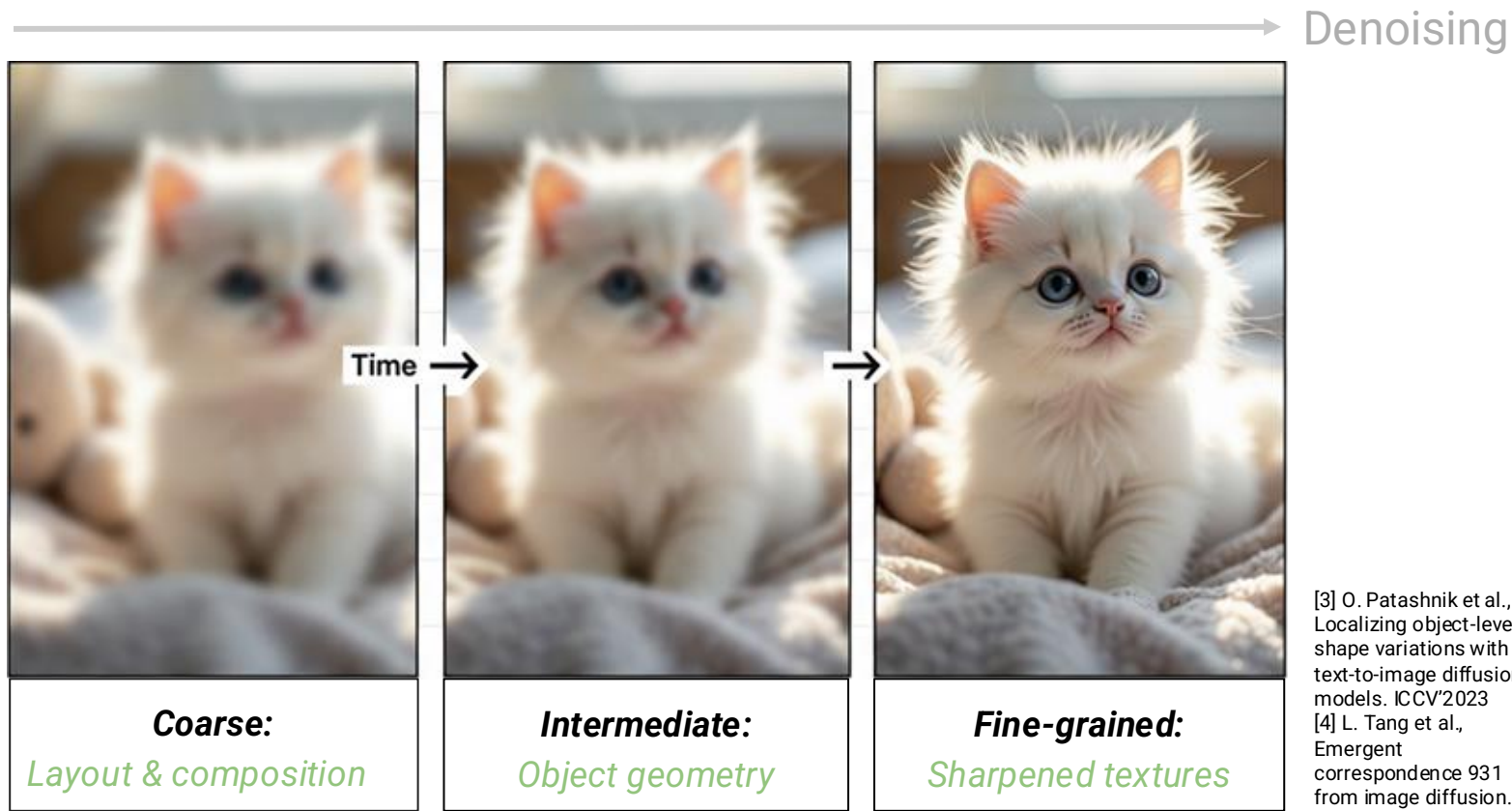


Even less noisy image



Generated image

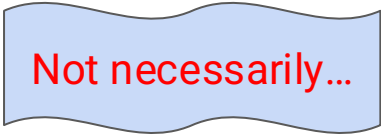
2. Should every denoising step process the image at the same granularity?



[1] D. Kim et al., Revelio: Interpreting and leveraging semantic information in diffusion models. ICCV'2025
[2] S. Mahajan et al., Prompting hard or hardly prompting: Prompt inversion for text-to-image diffusion models. CVPR'2024

[3] O. Patashnik et al., Localizing object-level shape variations with text-to-image diffusion models. ICCV'2023
[4] L. Tang et al., Emergent correspondence 931 from image diffusion. NeurIPS'2023

2. Should every denoising step process the image at the same granularity?



Not necessarily...

Inference in Generative Models is expensive

1. *Should every image be encoded at the same granularity?*

Not necessarily...

2. *Should every denoising step process the image at the same granularity?*

Inference in Generative Models is expensive

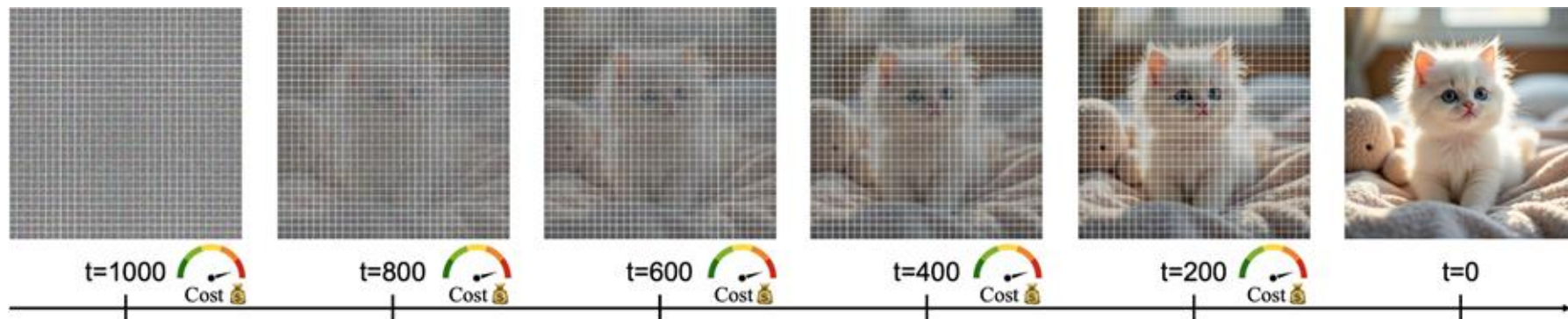
1. *Should every image be encoded at the same granularity?*

2. *Should every denoising step process the image at the same granularity?*

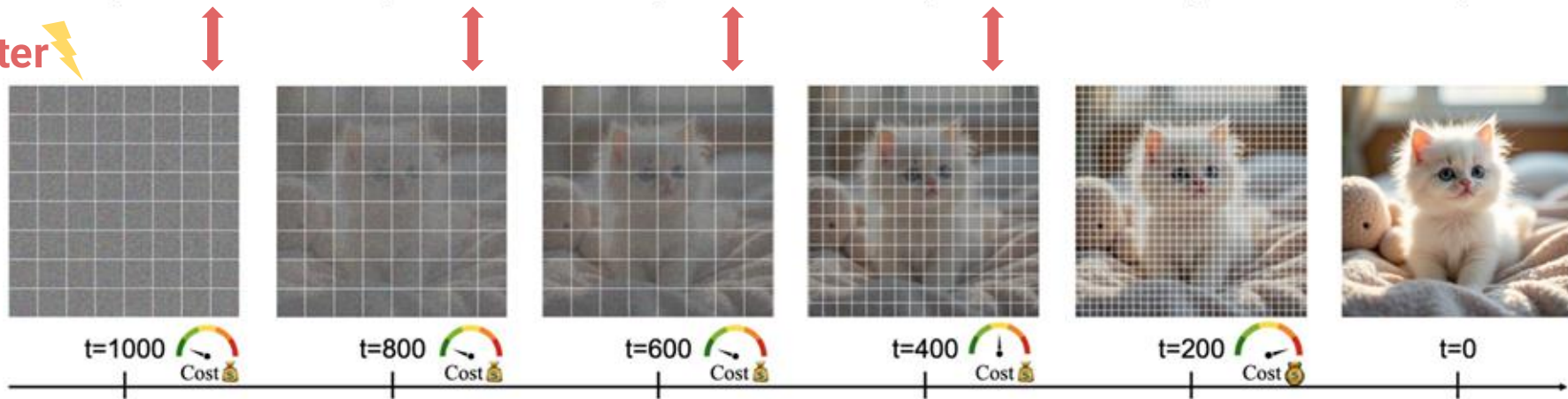
*Key idea: Dynamically vary patch sizes based on **content complexity** and the **denoising timestep**.*

*Dynamically vary patch sizes based on **content** and **timestep**.*

Baseline



After ⚡



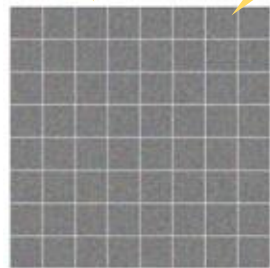
Larger patches

Responsible for coarse structure

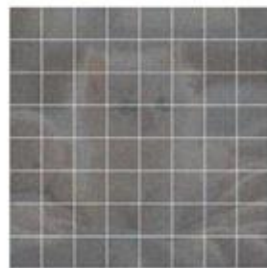
Fine-grained patches

Responsible for fine-grained structure

After



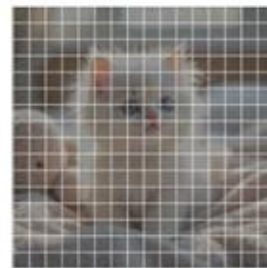
t=1000



t=800



t=600



t=400



t=200



t=0

Can we use latent's rate of change as a proxy for patch size?

How to determine the optimal patch size at each timestep?

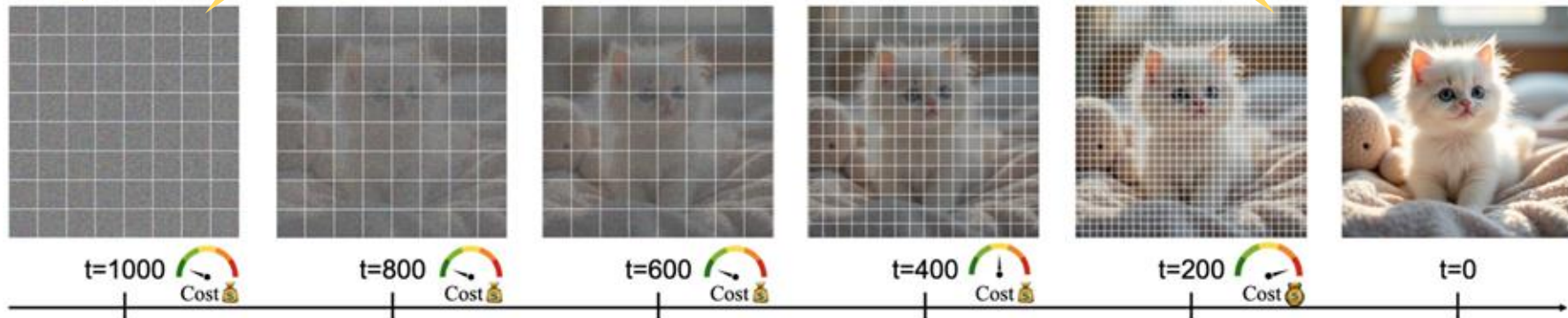
Larger patches

Responsible for coarse structure

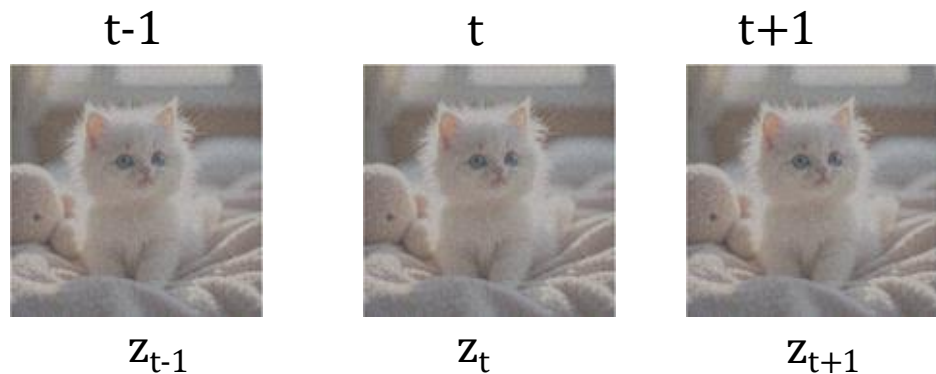
Fine-grained patches

Responsible for fine-grained structure

After



How to measure a latent's evolution?



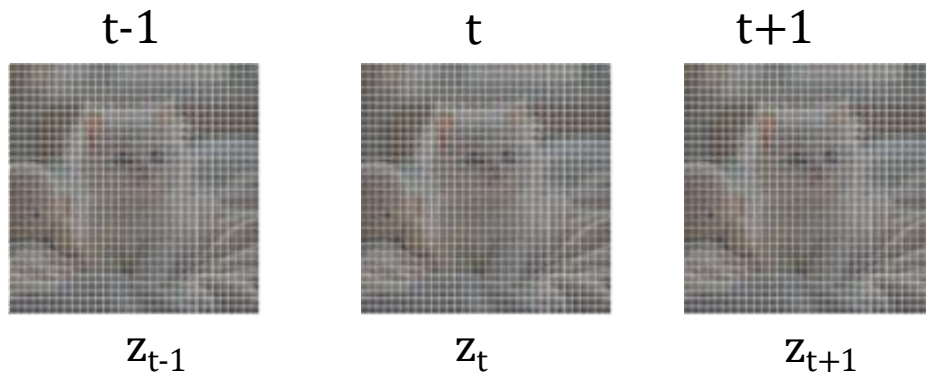
Displacement $\Delta \mathbf{z}_t = \mathbf{z}_t - \mathbf{z}_{t+1}.$

Velocity $\Delta^{(2)} \mathbf{z}_{t-1} = \Delta \mathbf{z}_{t-1} - \Delta \mathbf{z}_t.$

Acceleration $\Delta^{(3)} \mathbf{z}_{t-1} = \Delta^{(2)} \mathbf{z}_{t-1} - \Delta^{(2)} \mathbf{z}_t$ Rate of change of latent.

How to measure a latent's evolution?

Patch size = 2



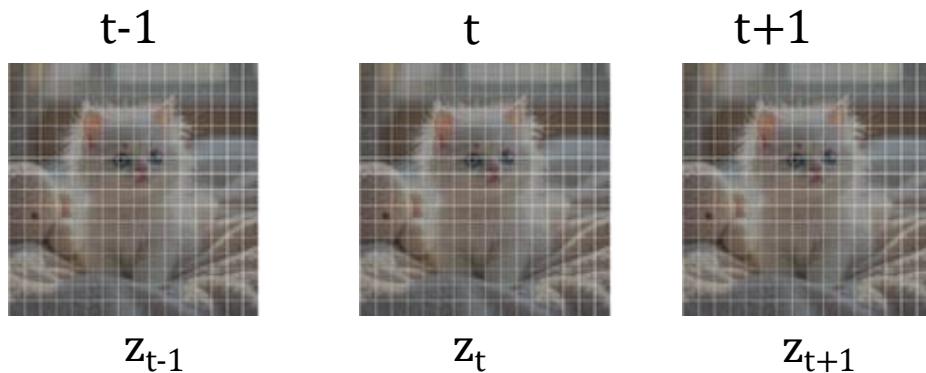
Acceleration $\Delta^{(3)} z_{t-1} = \Delta^{(2)} z_{t-1} - \Delta^{(2)} z_t$

Compute standard deviation of acceleration at different patch sizes

Low standard deviation? Latent evolves slowly -> go more coarse grained.

How to measure a latent's evolution?

Patch size = 4



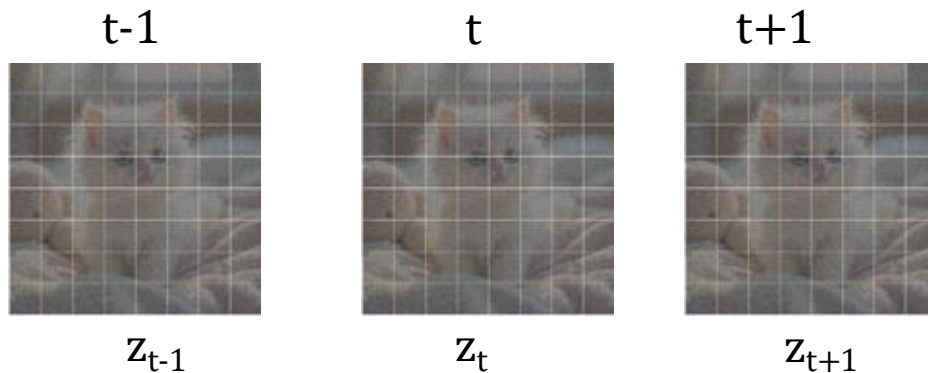
Acceleration $\Delta^{(3)} z_{t-1} = \Delta^{(2)} z_{t-1} - \Delta^{(2)} z_t$

Compute standard deviation of acceleration at different patch sizes

Low standard deviation? Latent evolves slowly -> go more coarse grained.

How to measure a latent's evolution?

Patch size = 8



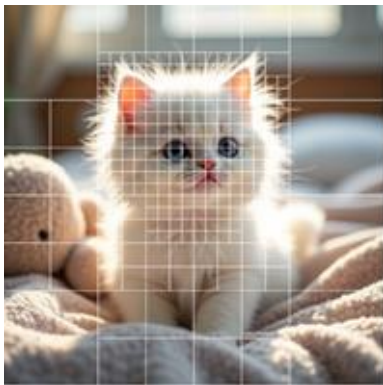
Acceleration

$$\Delta^{(3)} z_{t-1} = \Delta^{(2)} z_{t-1} - \Delta^{(2)} z_t$$

Compute standard deviation of acceleration at different patch sizes

How to measure a latent's evolution?

- Hypothesis: Fine-grained image region leads to rapid latent evolution



Solution: **Dynamic patch scheduler**

1. Use standard deviation of acceleration at different patch sizes.
2. Determine the right patch size.



Prompt 1:
Several Zebra
are together
through the
fence.

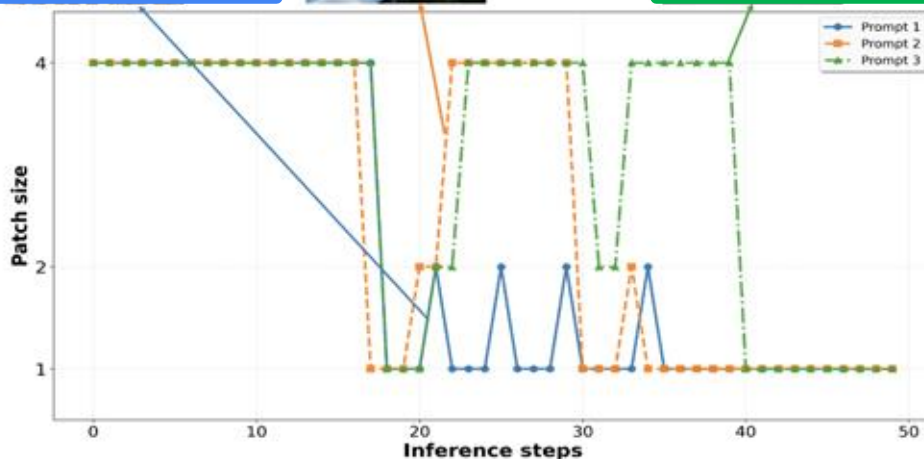
adapts to each



Prompt2: A strange
landscape: the
ground on the left
is covered with
snow but the ground
on the right is
covered with green
grass.



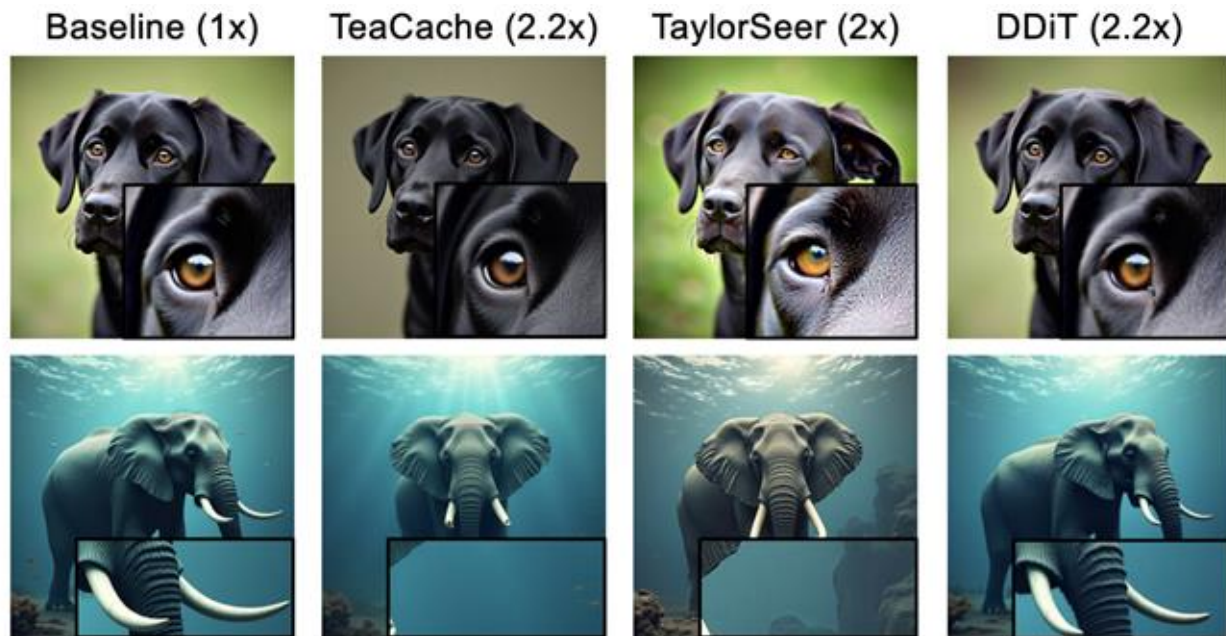
Prompt 3: A
simple red
apple in
black
background.



Observations:

1. High textured regions take smaller patches $\Rightarrow$ same compute as baseline.
2. Smooth regions take larger patches $\Rightarrow$ less compute

Results: Text to Image Generation



DDiT preserves

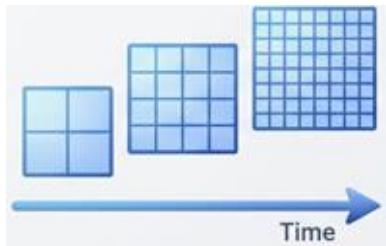
- ✓ fine-grained details
- ✓ pose
- ✓ spatial layout, and
- ✓ overall color distribution

Results: Text to Video Generation



- ✓ 3.2X speedup
- ✓ Maintains perceptual quality

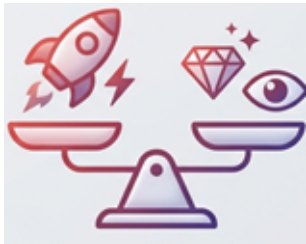
Summary: DDiT



Adaptive Granularity



Dynamic Patching



Max Efficiency & Quality

Text-to-Image

FLUX-1.Dev
(Baseline)
ImgR: 1.0291
CLIP: 0.3156



+ DDiT
1.88x speedup 🚀
ImgR: 1.0271
CLIP: 0.3148



+ DDiT
2.2x speedup 🚀
ImgR: 1.0284
CLIP: 0.3136



+ DDiT
3.5x speedup 🚀
ImgR: 1.0182
CLIP: 0.3117



Text-to-Video

Wan 2.1
(Baseline)
V-Bench: 81.24



+ DDiT
1.6x speedup 🚀
V-Bench: 81.17



+ DDiT
2.1x speedup 🚀
V-Bench: 80.97



+ DDiT
3.2x speedup 🚀
V-Bench: 80.53

