

Cross-Space Distillation

Teaching One-Step Students with Modern Diffusion Teachers

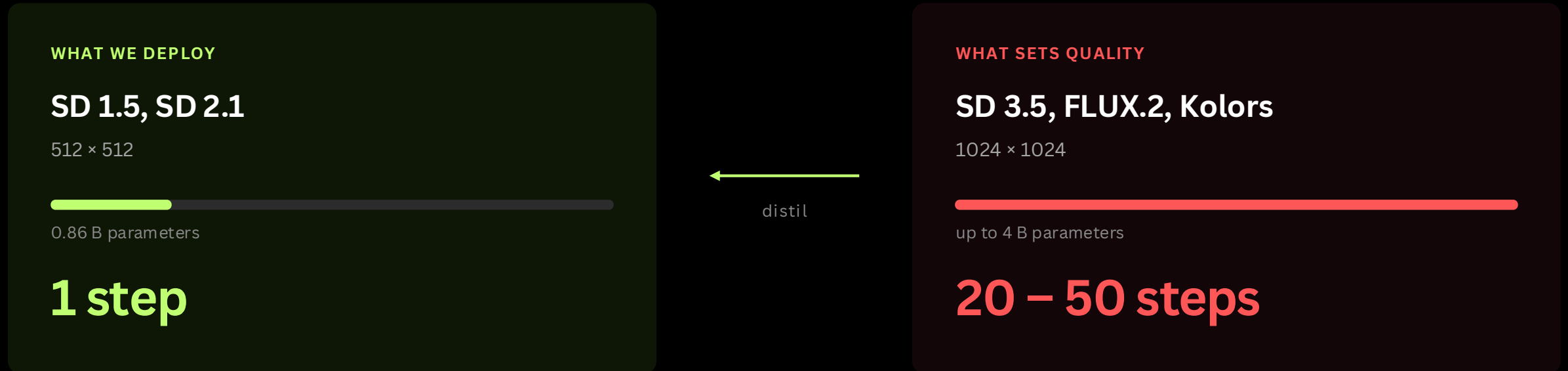
Anh Nguyen^{1*} Ngan Nguyen^{1*} Duc Vu^{1*} Trung Dao² Viet Nguyen³ Quan Dao⁴ Kien Nguyen¹
Chi Tran¹ Phong Nguyen¹ Khoi Nguyen¹ Cuong Pham¹ Dimitris N. Metaxas⁴ Vishal M. Patel³ Anh Tran¹

¹ Qualcomm AI Research ² University of Wisconsin–Madison ³ Johns Hopkins University ⁴ Rutgers University

ECCV 2026

Motivation: one-step models are what we deploy

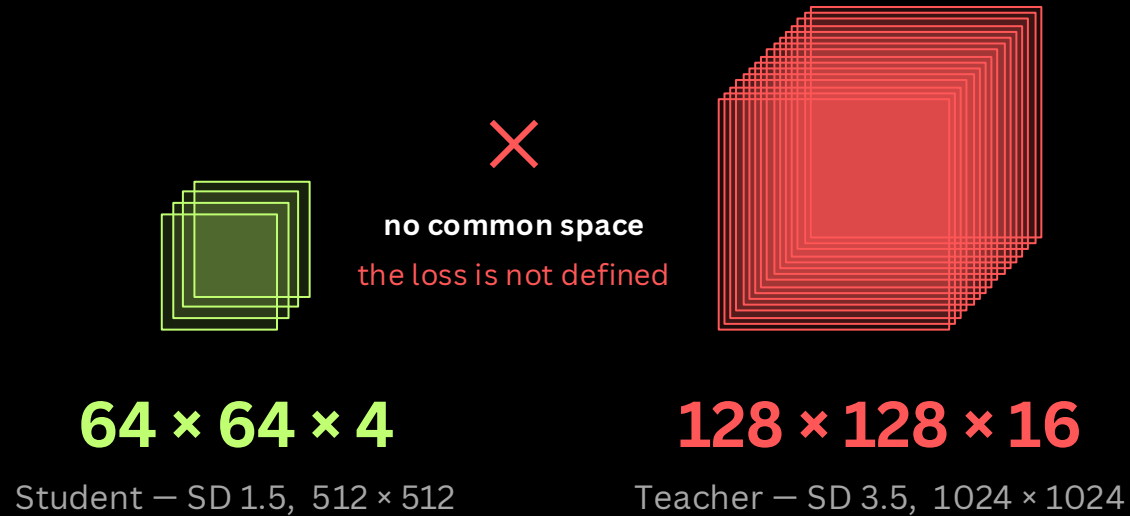
The backbones small enough to ship are generations behind the models that set quality.



Can a deployable Student learn from these Teachers?

Mismatched latent spaces cause existing distillation methods to fail

Distribution matching (DMD2, SiD) needs both models in one latent space.

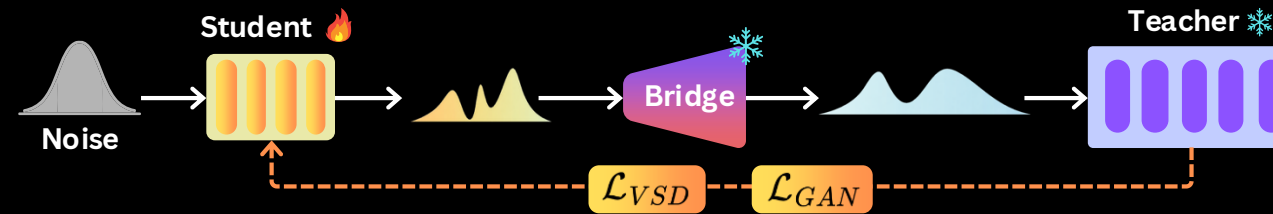


We call this setting **Cross-Space Distillation**

Latent resolution and VAE parameterization differ at the same time.

Idea: a **Bridge** between the two spaces

One stage per axis.



Where the Bridge sits in the distillation loop.

the Bridge



Training the Bridge

Two terms, both computed in Teacher space.

Match the Teacher's latent

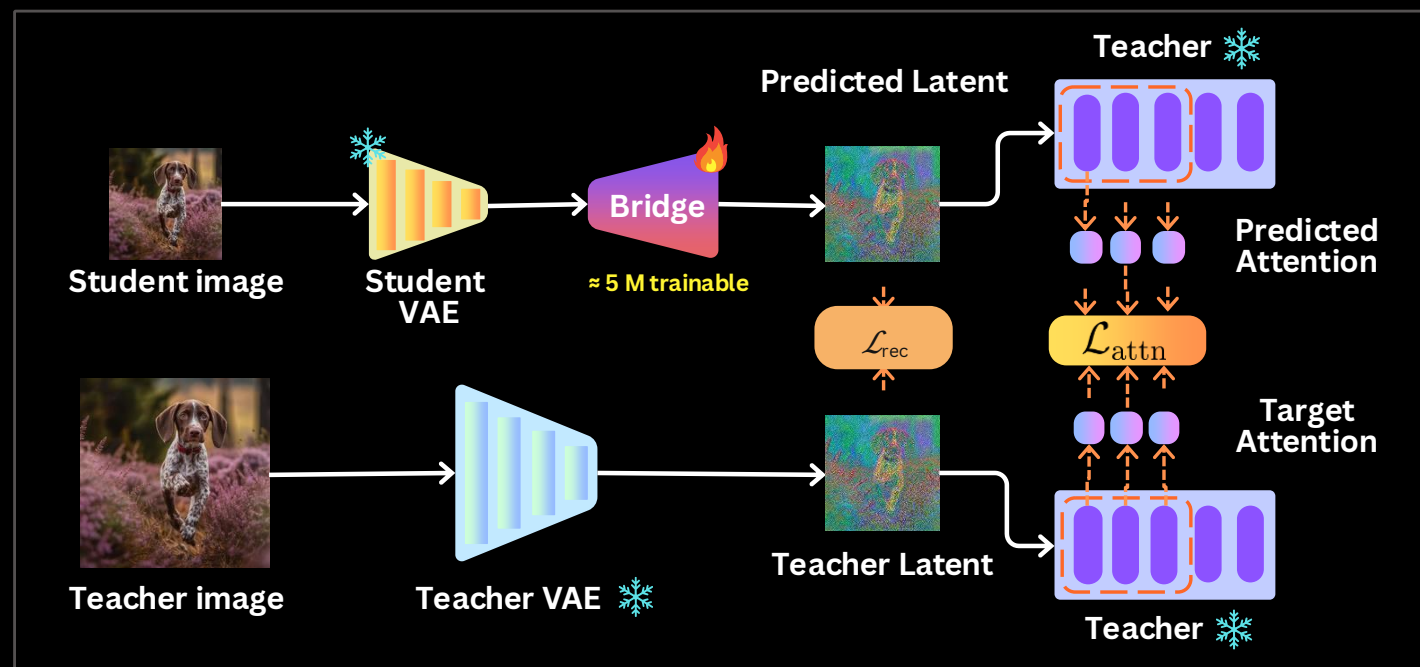
L1 distance between the mapped Student latent and the true one

Match the Teacher's attention

reverse KL between the Teacher's self-attention on each latent

HOW THE BRIDGE IS TRAINED

≈ 2 M images $\cdot 8 \times H100, 20$ h



Both terms are measured in the Teacher's space.

Ablation: the choice of training objective

Same Student, same Teacher, same Bridge. Only the objective changes.

	L1 ↓	PSNR ↑	SSIM ↑
Reconstruction only per-position signal only	0.041	23.58	0.70
+ E-LatentLPIPS no measurable gain	0.042	23.62	0.71
+ attention fidelity (ours) adds the long-range signal	0.035	24.97	0.75

WHY BOTH TERMS

+1.39 dB

PSNR, over reconstruction alone

Reconstruction is a per-position signal. Attention fidelity is what supplies the long-range one.

E-LatentLPIPS: + 0.04 dB.

Measured after the Teacher's decoder. L1 reconstruction in every row.

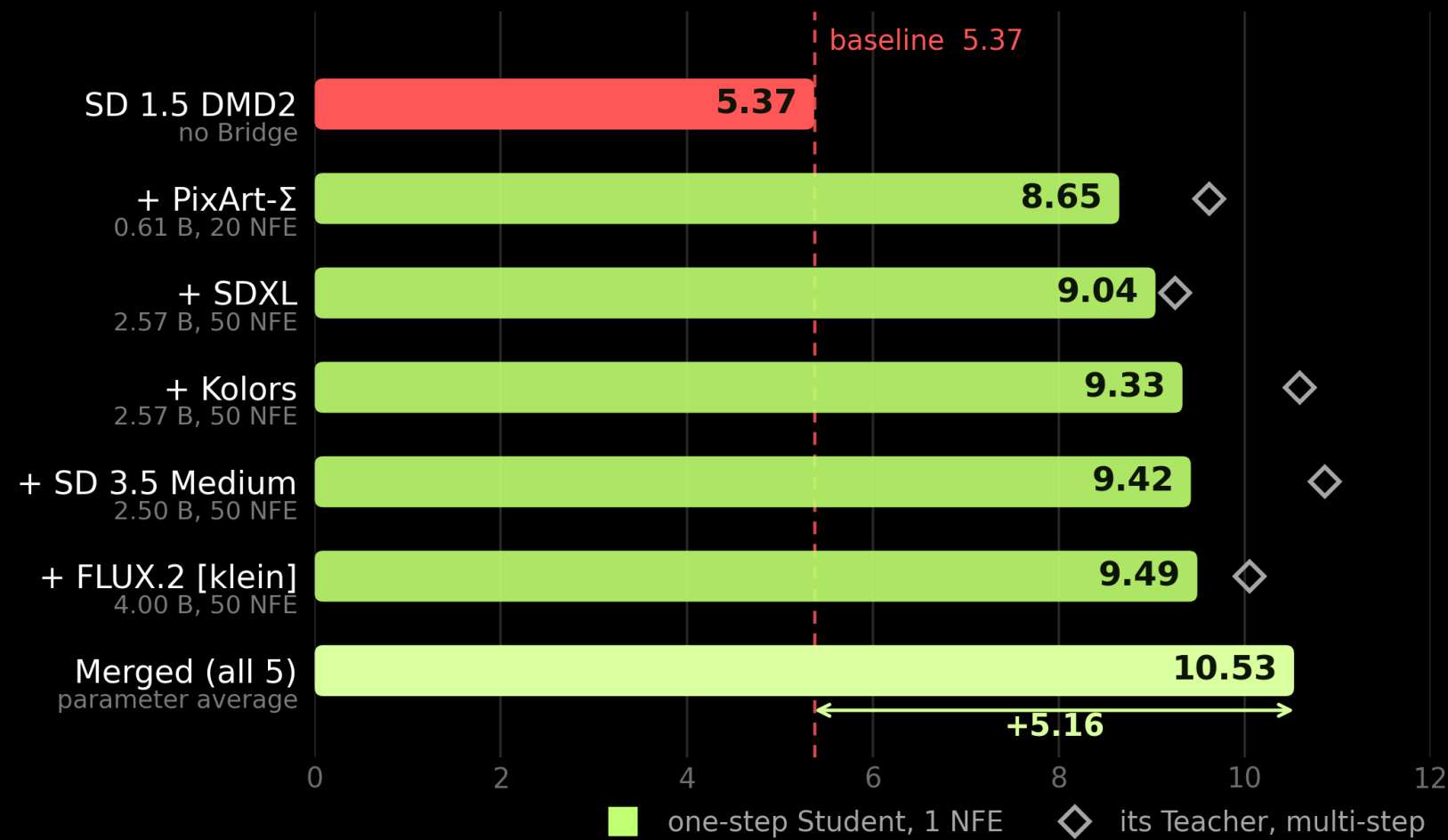
Result: the same Student, five modern Teachers

One SD 1.5 backbone, one step. Only the Teacher changes.

baseline	taught by a modern Teacher				
					
SD 1.5 DMD2	+ SDXL	+ Kolos	+ PixArt-Σ	+ SD 3.5 Medium	+ FLUX.2 [klein]
5.37	9.04	9.33	8.65	9.42	9.49

One prompt. Scores are HPSv3 over the benchmark's 12,000 prompts.

Result: HPSv3 across all five Teachers



Every Teacher gives at least **+3.28** over the baseline.

10.53

merged Student, one step

Above SDXL (2.57 B) and FLUX.2 [klein] (4.00 B) at fifty steps.

Model merging

Average the parameters of the five distilled Students.



One-step samples from the merged SD 1.5 Student.

10.53

SD 1.5, from 9.49

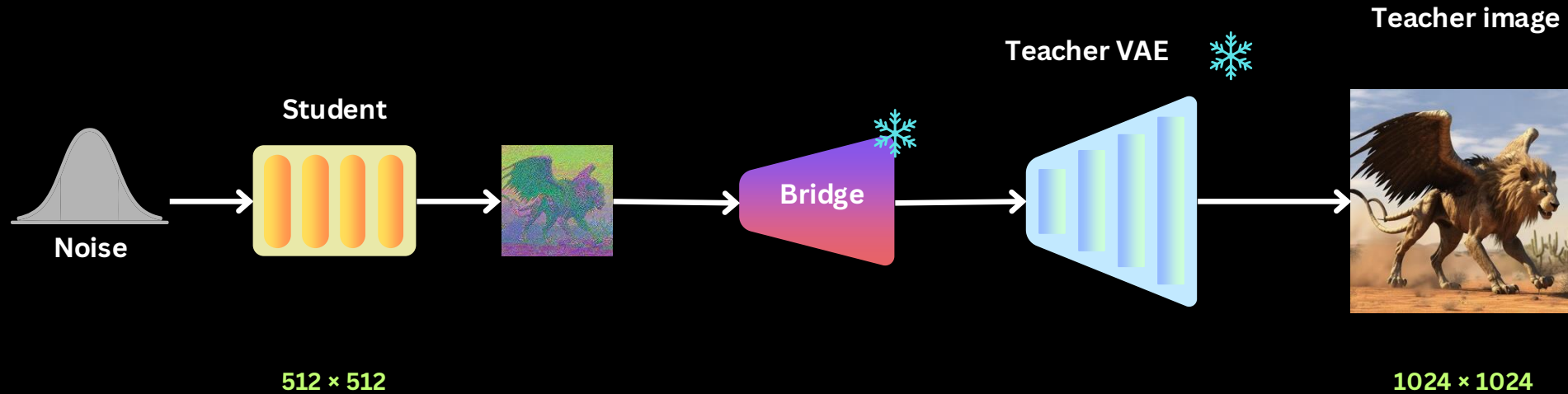
9.75

SD 2.1, from 8.74

Three Teacher architecture families — U-Net, DiT, MM-DiT — in one 0.86 B Student.

Bonus: 1024×1024 from a Student trained at 512

The Bridge outputs a Teacher latent, so the Teacher's decoder can read it.



No additional training, still one function evaluation.

Conclusion

The teacher set is no longer fixed by the Student's latent space.

SAMPLE EXPERIMENT

5

Teachers, three architecture families

≈ 5 M

Bridge, outside the backbone

1

function evaluation

5.37 → 10.53

HPSv3 on SD 1.5

Teacher and Student can be chosen independently.

A deployed Student can be taught by a Teacher that arrives later.

arXiv:2606.32020 Qualcomm AI Research · UW–Madison · Johns Hopkins · Rutgers