



1 · Teacher gradients are expensive Monte Carlo estimates

A **frozen diffusion teacher** supplies gradients to a differentiable generator with parameters θ , used in text-to-3D (SDS), one-step distillation (DMD), and data attribution. Each gradient is an expectation over a noise level t and Gaussian noise ϵ , estimated by Monte Carlo and applied through the chain rule:

$$\mathbf{u}_{\text{SDS}}(\theta) = \mathbb{E}_{\mathbf{q}, t, \epsilon} \left[\mathbf{f}(\mathbf{x}, t, \epsilon) \frac{\partial \mathbf{x}}{\partial \theta} \right] \xrightarrow{\text{Monte Carlo}} \hat{\mathbf{u}}_{\theta}^{\text{naive}} = \frac{1}{R} \sum_{r=1}^R \mathbf{f}(\mathbf{x}^{(r)}, t^{(r)}, \epsilon^{(r)}) \frac{\partial \mathbf{x}^{(r)}}{\partial \theta}$$

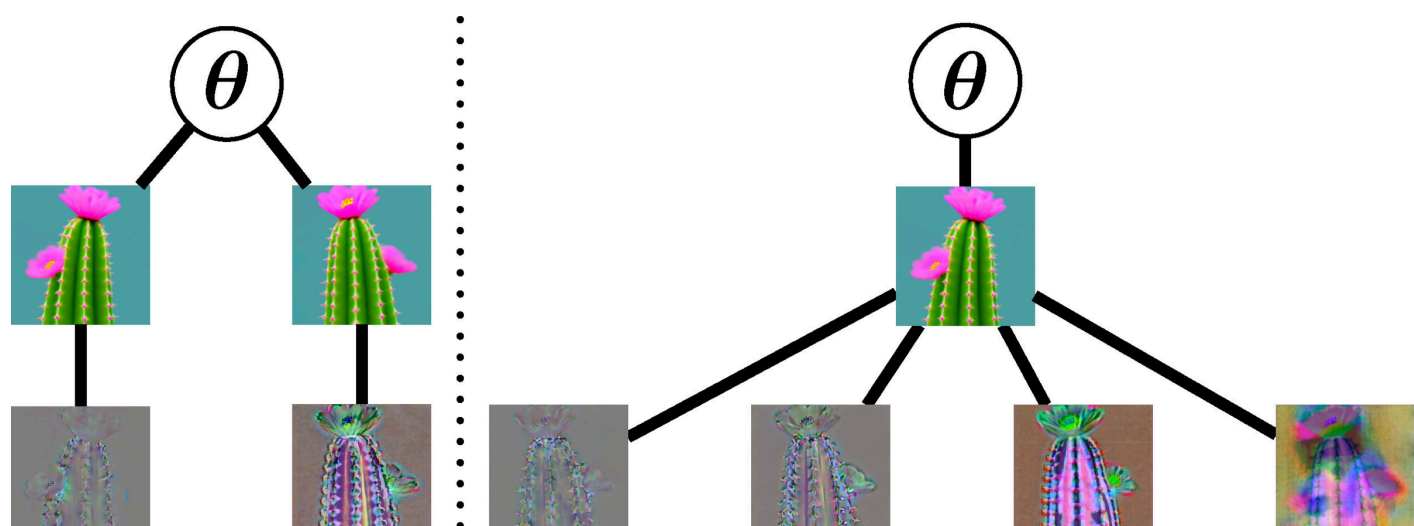
- $\mathbf{x}$ **encoded render**, $\mathbf{x} = \text{Encode}(g(\theta, \mathbf{q}))$: a differentiable render g from camera $\mathbf{q}$, encoded into the teacher's latent space.
- $\mathbf{f}(\mathbf{x}, t, \epsilon)$ **teacher direction**. For SDS, $\mathbf{f} = w_{\text{SDS}}(t) \mathbf{r}$ with residual $\mathbf{r} = \hat{\epsilon}_{\phi}(\mathbf{z}_t, t, \mathbf{c}) - \epsilon$, $\mathbf{z}_t = \alpha_t \mathbf{x} + \sigma_t \epsilon$.
- $\partial \mathbf{x} / \partial \theta$ **renderer / encoder Jacobian**: backprop through render + encode. This is the **dominant cost**.

Naive variance reduction is wasteful: each Monte Carlo estimate of the gradient averages R samples, and every sample pays the full render/encode cost ($c_{\text{denoise}} \ll c_{\text{render+encode}}$), so reducing variance by naively growing the batch size R scales this dominant cost. Downstream pipelines inherit the teacher's schedule and average ad hoc, often biasing the objective, without accounting for which randomness dominates or how it trades against compute.

CARV is a compute-aware variance-accounting framework: it measures gradient variance against wall-clock to expose which randomness actually dominates, then reduces it with **three unbiased drop-ins** that leave the objective unchanged.

2 · Method: three unbiased drop-ins

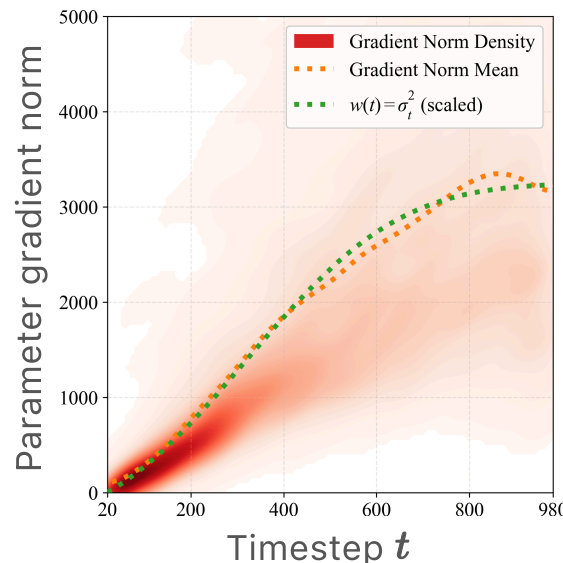
1 Amortized compute reuse



$$\hat{\mathbf{u}}_{\theta}^{\text{reuse}} = \frac{1}{R} \sum_{r=1}^R \left(\frac{1}{K} \sum_{k=1}^K \mathbf{f}(\mathbf{x}^{(r)}, t^{(r,k)}, \epsilon^{(r,k)}) \right) \frac{\partial \mathbf{x}^{(r)}}{\partial \theta}$$

The naive estimator renders once per teacher evaluation ($K=1$). Instead, reuse render $\mathbf{x}$ and average K teacher evaluations with fresh (t, ϵ) : the shared vector-Jacobian product through $\partial \mathbf{x} / \partial \theta$ is computed once, so within-render variance falls at little added cost.

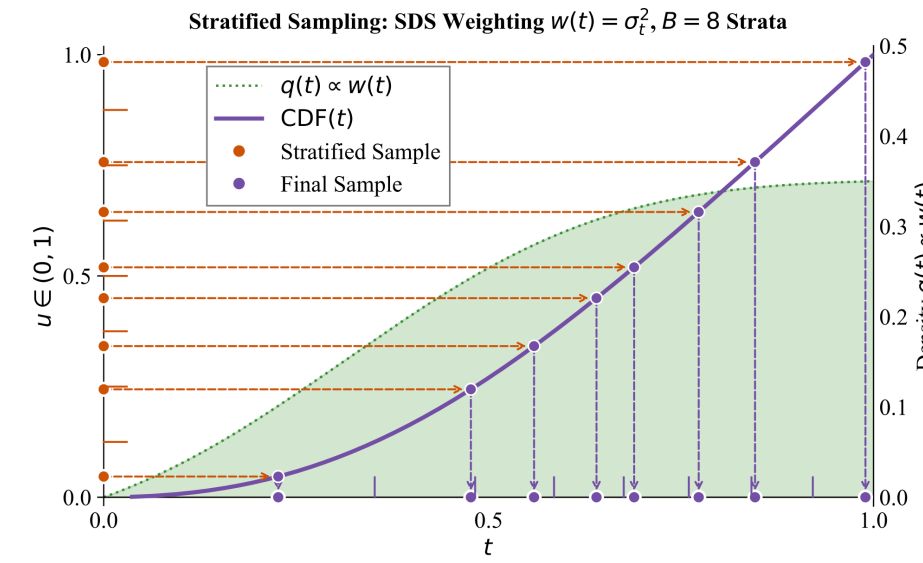
2 Timestep importance weighting (IW)



$$q^*(t) \propto p(t) \sqrt{\mathbb{E}[\|\mathbf{f}(t, \xi)\|_2^2 | t]} \rightarrow q(t) \propto p(t) w_{\text{SDS}}(t)$$

The variance-optimal proposal is intractable. After backprop, the parameter-gradient norm tracks the **teacher weight** $w_{\text{SDS}}(t) = \sigma_t^2$ (green), a free, near-oracle surrogate.

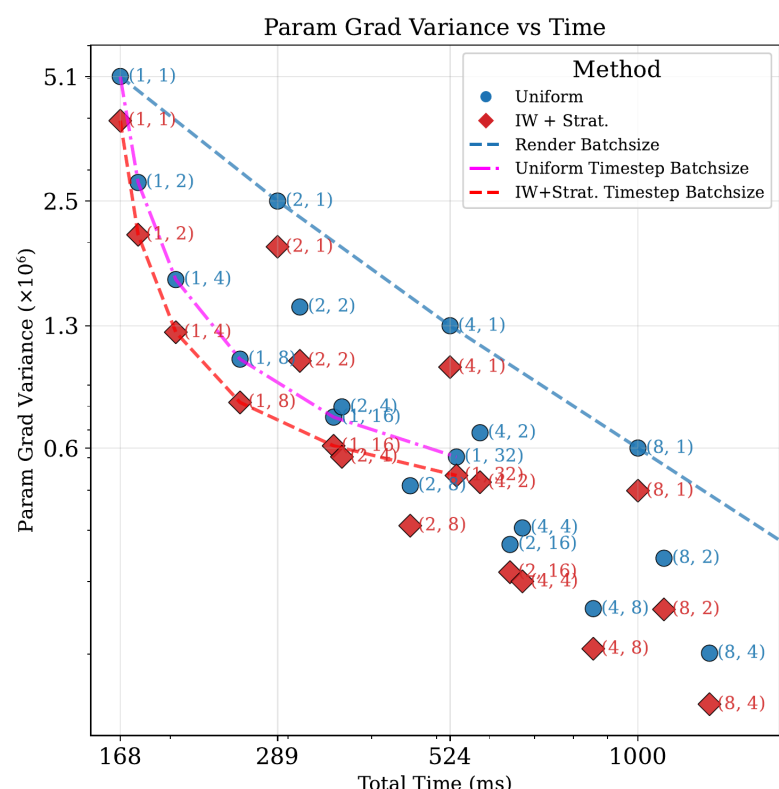
3 Stratified inverse-CDF



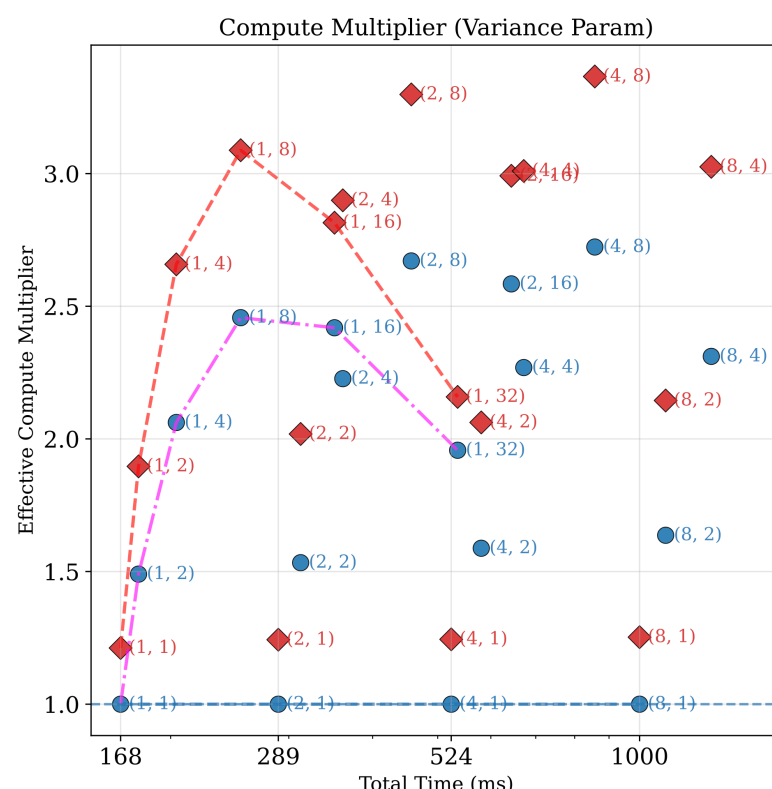
$$u_b^{(r)} = \frac{b-1 + \xi_b^{(r)}}{B}, \quad t_b^{(r)} = \text{CDF}_q^{-1}(u_b^{(r)})$$

IID draws of t clump and leave gaps. Forcing **one draw per equal-mass bin** and mapping through the proposal's inverse-CDF combines stratification with importance weighting at **no extra compute**.

3 · Results: 2–3× effective compute on text-to-3D



Gradient variance vs wall-clock. Each point is an (R, K) pair; at every budget, IW + stratification (red) sits below the uniform baseline (blue).

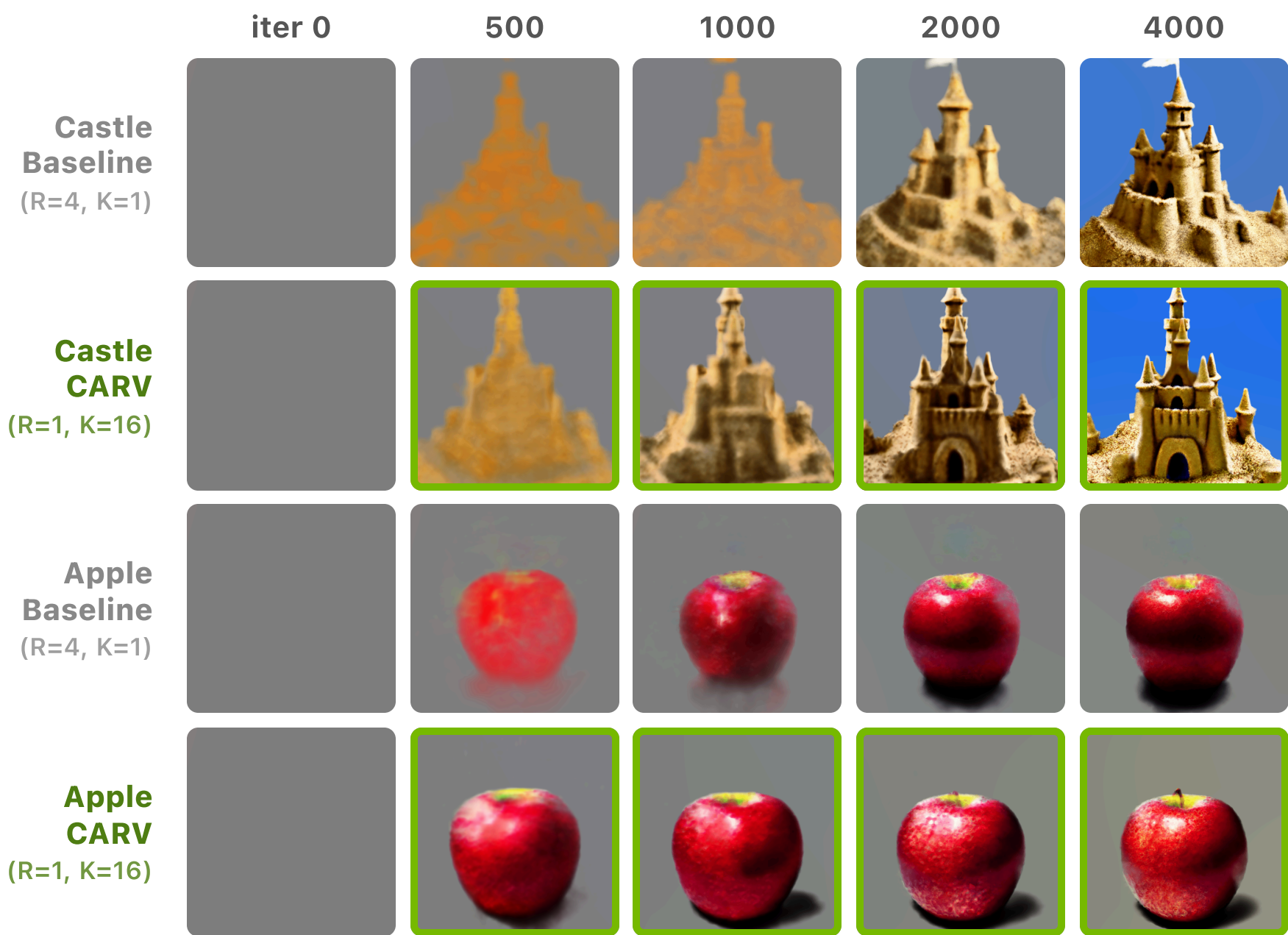


Effective compute multiplier ($\text{cost}_{\text{baseline}} / \text{cost}_{\text{method}} |_{\text{equal Var}}$): baseline compute needed to match CARV's variance. Reuse (blue) drives most of the lift; IW + stratification (red) peaks at **3.3x** at (1, 8).

K	Uniform	IW	Strat.	IW+Strat.
1	1.00	1.24	1.00	1.24
2	1.57	1.93	1.66	2.05
4	2.24	2.68	2.47	2.94
8	2.63	3.01	2.96	3.29
16	2.52	2.75	2.75	2.94
32	1.98	2.09	2.08	2.18

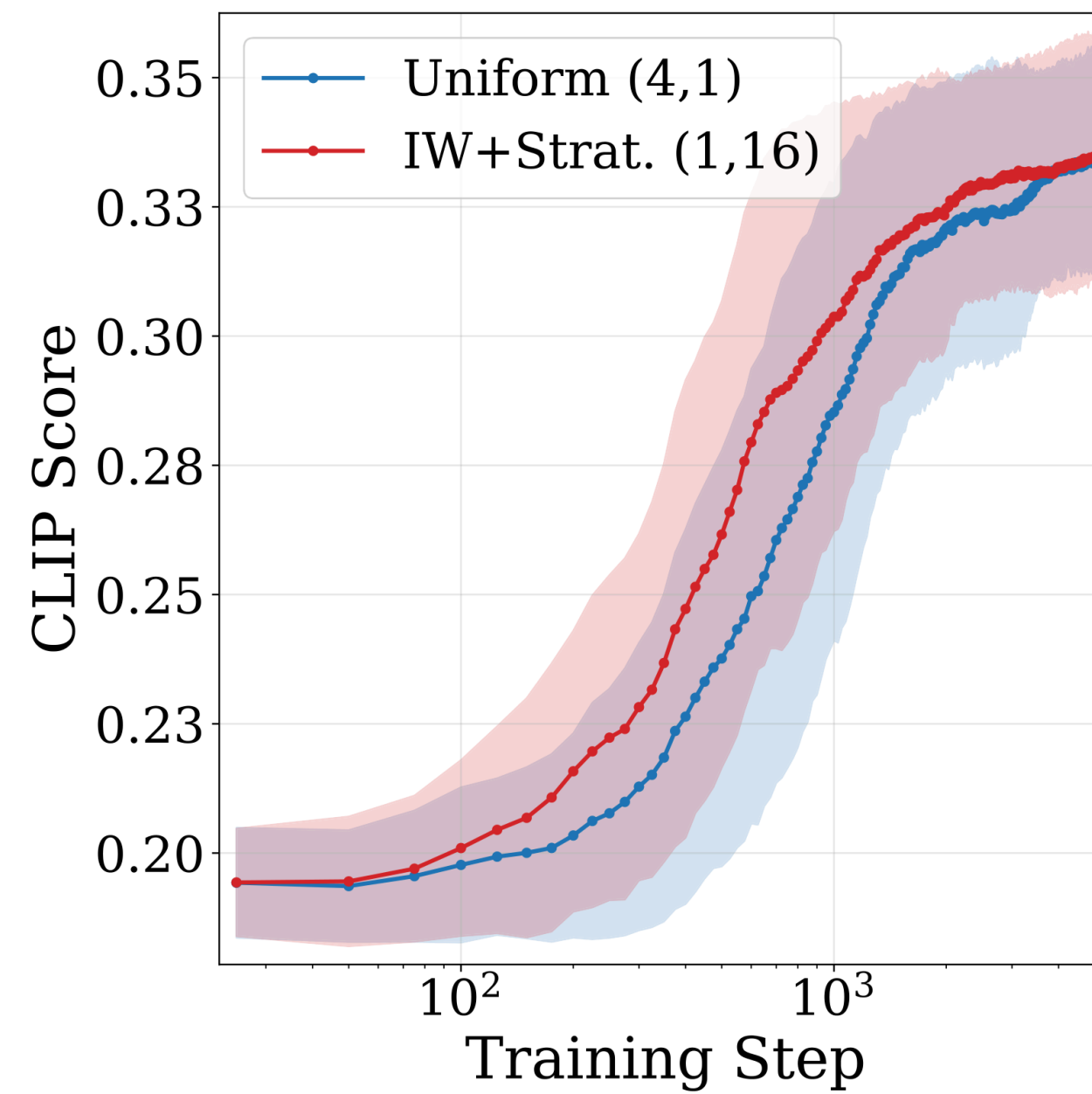
ECM by re-noisings K (avg. over renders R , 5 runs). Ranking **IW+Strat > IW $\approx$ Strat > Uniform** is stable; the three drop-ins are complementary.

Equal-cost training trajectory · “a castle-shaped sandcastle” and “a red apple”



At matched per-iteration cost, CARV resolves coherent geometry and texture several thousand iterations earlier than the baseline.

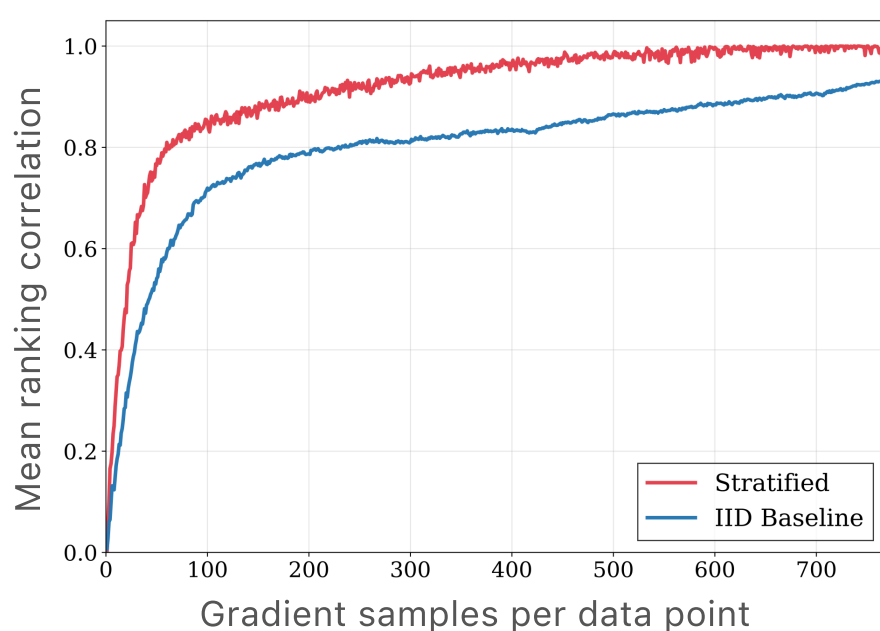
Lower variance → faster convergence



At equal cost, CARV (red, (1, 16)) sits above the uniform baseline (blue, (4, 1)), reaching the converged score in **~half the iterations** (30 prompts × 3 seeds).

4 · Applications beyond text-to-3D

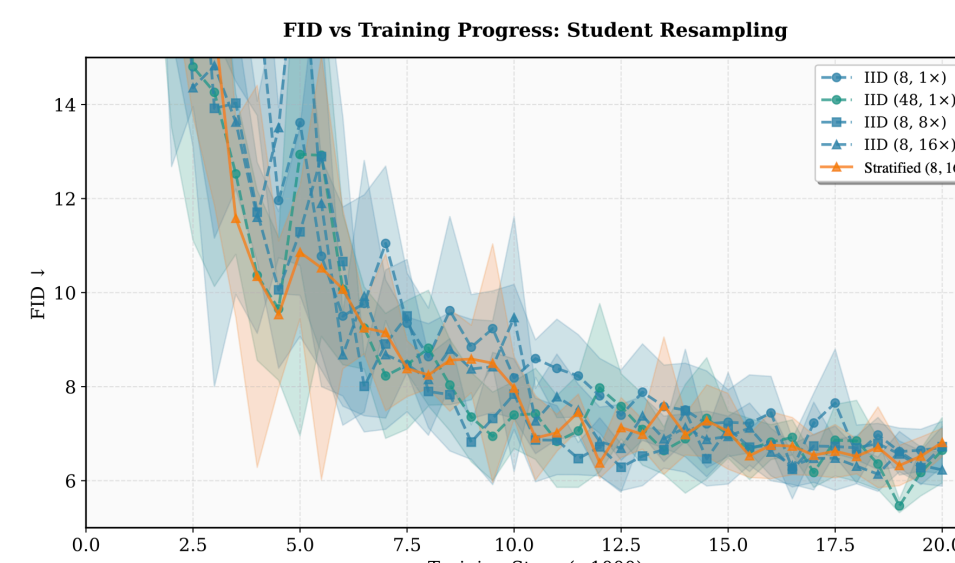
Video data attribution: the same idea drops in



$$I(\text{query}, n) = \frac{1}{|\mathcal{T}|} \sum_{(t, \epsilon) \in \mathcal{T}} \frac{\mathbf{g}_{\text{query}}(t, \epsilon)^{\top} \mathbf{g}_n(t, \epsilon)}{\|\mathbf{g}_{\text{query}}(t, \epsilon)\|_2 \|\mathbf{g}_n(t, \epsilon)\|_2}$$

MOTIVE-style influence on **Wan2.1** video generation (Outstanding Paper Honorable Mention). Stratified sampling (red) matches ground-truth influence rankings with far fewer per-clip gradients than IID (blue): a **2–3x** multiplier at practical budgets, with global stratification most effective.

Single-step distillation: when variance isn't the bottleneck



$$\nabla_{\theta} D_{\text{KL}} \simeq \mathbb{E}_{\epsilon, t, \epsilon'} \left[w(t) \alpha_t (\mathbf{s}_{\text{fake}}(\mathbf{z}_t, t) - \mathbf{s}_{\text{real}}(\mathbf{z}_t, t)) \frac{\partial G_{\theta}(\epsilon)}{\partial \theta} \right]$$

The same drop-ins cut DMD gradient variance by an **order of magnitude**, yet FID does not move. DMD marks where MC variance **stops being the bottleneck**; auxiliary stabilizers and input diversity bind instead. CARV stays unbiased, and its variance-per-compute accounting exposes the regime where MC variance is the bottleneck, and, as in DMD, **where it isn't**.

When CARV helps: the MC teacher gradient dominates the update · rendering / encoding costs more than denoising ($c_{\text{render+encode}} > c_{\text{denoise}}$) · variance limits convergence.