

ReSWD: ReSTIR'd, not shaken. Combining Reservoir Sampling and Sliced Wasserstein Distance for Variance Reduction

Mark Boss, Andreas Engelhardt, Simon Donné, and Varun Jampani

Stability AI

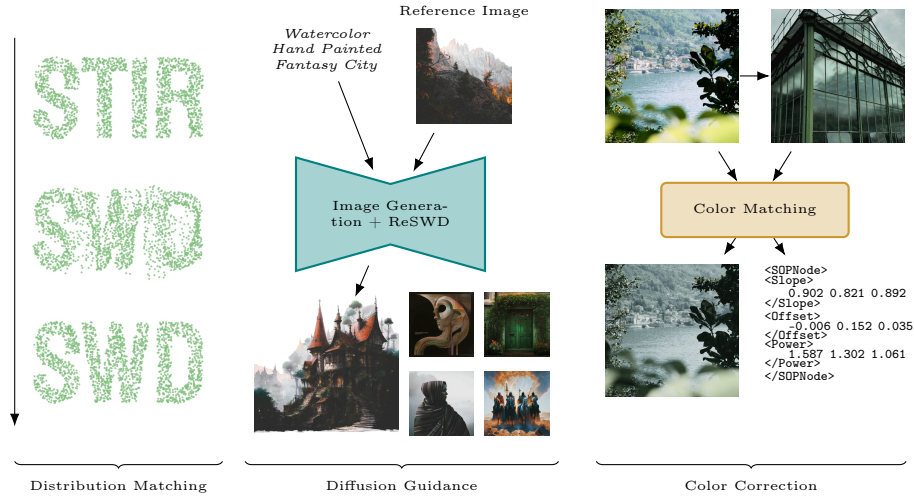


Fig. 1: Applications. SWD can be used in various applications and with ReSWD we enhance their performances. Here, we show diffusion guidance, color corrections and general distribution matching.

Abstract. Distribution matching is central to many vision and graphics tasks, where the widely used Wasserstein distance is too costly to compute for high-dimensional distributions. The *Sliced Wasserstein Distance* (SWD) offers a scalable alternative, yet its Monte Carlo estimator suffers from high variance, resulting in noisy gradients and slow convergence. We introduce *Reservoir SWD* (ReSWD), which integrates Weighted Reservoir Sampling into SWD to adaptively retain informative projection directions in optimization steps, resulting in stable gradients while remaining unbiased. Experiments on synthetic benchmarks and real-world tasks such as color correction and diffusion guidance show that ReSWD consistently outperforms standard SWD and other variance reduction baselines. Project page with code: <https://ReservoirSWD.github.io>

1 Introduction

Distribution matching is a central problem in computer vision and graphics: from classical tasks such as histogram matching to more nuanced applications such as color grading (5; 16; 39; 41), texture alignment (12; 17), and the guidance of generative models (26). Among the available metrics, the Wasserstein distance has emerged as particularly powerful due to its ability to capture shifts between distributions. However, the Wasserstein distance suffers from the curse of dimensionality and has a practical convergence rate scaled with the d -dimensional distribution $\mathcal{O}(n^{-\frac{1}{d}})$, with n data points. This limits its direct use in many iterative optimization settings.

The *Sliced Wasserstein Distance* (SWD) (5; 17; 39) offers a practical alternative by projecting distributions in random directions and averaging their 1-D Wasserstein costs. This reduces the computational burden to a sequence of 1-D sorting problems but introduces a new challenge: the expectation over projection directions is typically approximated via Monte Carlo (MC) sampling, which suffers from high variance. Increasing the directions would result in the true expected value, but the complexity of SWD scales with the number of random directions M as $\mathcal{O}(Mn \log n)$. Hence, the curse of dimensionality is introduced again indirectly with the number of projections, which should increase with the dimensionality of the data. In optimization scenarios, where such SWD metrics are used as loss functions to optimize neural networks, this variance directly results in noisy gradients and slower convergence.

A variety of variance reduction techniques have been explored for MC estimators, but most remain underutilized in distribution matching objectives. In this work, we draw inspiration from recent advances in rendering, particularly the ReSTIR resampling framework (3), and adapt *Weighted Reservoir Sampling* (WRS) (6; 11) to the SWD setting. The resulting estimator, which we term *Reservoir SWD* (ReSWD), continuously reuses and reweighs the most informative projection directions throughout the optimization. Intuitively, ReSWD preferentially retains those directions where the two distributions are most dissimilar, thus concentrating computational effort on projections that produce stronger and more stable gradients. This simple but effective modification significantly reduces stochastic variance while preserving unbiasedness due to the weighting of the WRS, leading to faster and more robust optimization.

We demonstrate the advantages of ReSWD on two real-world distribution matching problems of color correction and color diffusion guidance, as shown in Fig. 1, as well as synthetic general distribution matching problems, showing clear improvements over the standard SWD and existing variance reduction baselines.

2 Related Work

Distribution Matching and Optimal Transport. Distribution matching is the general task of aligning two distributions. A particularly effective way of achieving this is with *Optimal Transport* (OT) by minimizing the cost of

moving the probability mass, which is often performed using Wasserstein distances (37; 47). Entropic regularization accelerates OT through Sinkhorn iterations (9), leading to scalable solvers and Sinkhorn divergences that interpolate between OT and kernel *Maximum Mean Discrepancies* (MMD) (14). OT has been widely applied to textures, color transfer, and barycenters (5; 39; 41).

To address high-dimensional costs, *Sliced Wasserstein* (SW) replaces couplings with averages of 1-D transports. Early works introduced SW barycenters and Radon/Projected variants that are fast and differentiable (5; 41). Extensions include Max-SW, which selects the most contributing directions instead of averaging over them (10), and several other techniques to improve efficiency and optimization performance (22; 29; 31; 32; 33). SW has also been studied through geometry and flows: gradient flows as generative dynamics with long-time analysis (8; 46), intrinsic geometry of SW space (36), extensions to Cartan–Hadamard manifolds (4), and stereographic spherical SW for curved domains (45). Recent work further reduces variance via quasi-Monte Carlo and control variates (28; 30). In practice, SW is widely used as a loss in graphics, vision, and generative modeling *e.g.*, for textures, patch statistics, and color transfer due to stable gradients and favorable sample complexity (12; 17; 39; 49).

Our work ties into the work selecting the most contributing directions similar to Max-SW but achieves this by building a reservoir of high contributing directions during the training. This improves efficiency, as we keep more directions in optimization. It also ties in with the work on variance reduction while remaining unbiased, as our technique is inspired by variance reduction techniques in real-time path tracing (3).

Color Transfer and Correction. Classical color transfer methods framed recoloring as histogram or distribution matching across images, often using statistical metrics or optimal transport (5; 39; 41). These techniques established the foundation for perceptual color differences in imaging. More recently, multiscale OT and Wasserstein-based metrics have been proposed as perceptually faithful color difference measures (16).

With deep learning, style transfer methods enabled more expressive color and appearance manipulation. CNN-based approaches introduced artistic and photorealistic stylization (2; 7; 15; 19; 21; 24; 25; 27; 50), with adaptive instance normalization and feature transforms widely adopted to control color and tone. These methods evolved from hand-crafted global mappings to neural architectures capable of spatially aware and semantically coherent color control.

With the advance of text-guided diffusion models (13; 44), additional mechanisms to control the style such as LoRA (20), IP-Adapter (52), ControlNets (51) are introduced. A recent work by Lobashev *et al.* (26) also demonstrated that traditional constraints based on OT can be used during generation to enforce a color distribution.

Our work can be used to enhance traditional color matching works with improved efficiency and to enable more novel color-guided diffusion tasks.

3 Method

In our setting, we only ever compare empirical *discrete* distributions, that is, two sets of samples $X = \{x_1, \dots, x_N\} \subset \mathbb{R}^d$ and $Y = \{y_1, \dots, y_N\} \subset \mathbb{R}^d$, drawn from the underlying distributions $\mathcal{X}$ and $\mathcal{Y}$.

3.1 Preliminaries

Wasserstein Distance. The Wasserstein p -distance measures differences between two empirical distributions. In the 1D case relevant to our setting, given two sets of samples X and Y , the Wasserstein distance can be computed simply by sorting both sets and comparing the corresponding order statistics:

$$W_p(X, Y) = \left(\frac{1}{n} \sum_{i=1}^n |x_i - y_i|^p \right)^{1/p}, \quad (1)$$

This efficient formulation has complexity $\mathcal{O}(n \log n)$ due to sorting, in contrast to the cubic complexity of the general d -dimensional case (5; 39; 41).

Sliced Wasserstein Distance (SWD). The Sliced Wasserstein Distance proposes to stochastically approximate the true Wasserstein distance in multi-dimensional distributions. It is defined as (12; 17; 39):

$$S_p(X, Y) = \mathbb{E}_{\theta \sim U(S^{d-1})} \left[W_p(\pi_\theta X, \pi_\theta Y) \right], \quad (2)$$

where π_θ indicates the projection onto the unit-normal vector θ (uniformly sampled from S^{d-1} , the d -dimensional unit sphere $\mathbb{R}^d$).

In practice, the expectation is often approximated via Monte Carlo (MC) integration:

$$S_p(X, Y) \approx w_i \sum_{i=1}^L W_p(\pi_{\theta_i} X, \pi_{\theta_i} Y). \quad (3)$$

with random uniformly sampled directions θ_i and w_i being the sampling weights, *i.e.* $w_i = \frac{1}{L}$ with uniform direction sampling. The resulting complexity is therefore $\mathcal{O}(Ln \log n)$. The SWD is an unbiased estimate of the true Wasserstein distance (17; 39). With modern frameworks such as PyTorch or Tensorflow, the entire metric can be trivially implemented fully differentiable, which enables usage in optimization settings. However, due to the MC integration, SWD can lead to noisy gradients, which we aim to solve with ReSWD.

3.2 Reservoir Sliced Wasserstein Distance (ReSWD)

Although SWD is highly efficient in the calculation at high dimensions, it still suffers from high variance due to the MC integration from random directions. Especially, when SWD is used as a loss during an optimization, this leads to noisy gradients. In computer graphics, variance is often also an issue, especially

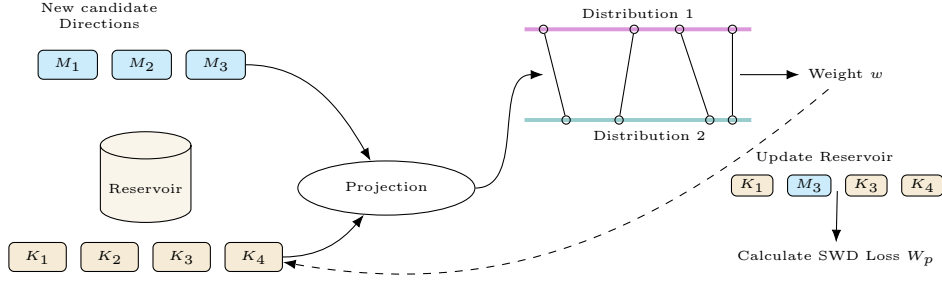


Fig. 2: Overview. During the optimization we keep a reservoir of highly influential directions. We use the Sliced Wasserstein Distance (SWD) as a proxy metric for the reservoir update and the final optimization loss. Note that only directions in the reservoir influence the optimization to remain unbiased.

in real-time path tracing where MC integration is also usually employed. Inspired by ReSTIR’s *resampled-importance-sampling* (3), we incorporate the Weighted Reservoir Sampling (WRS) mechanism into SWD. This allows the reuse of information between optimization steps, and hence faster convergence overall with minimal performance penalties.

Weighted Reservoir Sampling is a family of algorithms that draw a fixed subset from a weighted stream of candidates, *i.e.* directions, in a single pass (6; 11). Suppose that we wish to maintain a reservoir of K candidates while processing a candidate pool of $K + M$ items, where M denotes the number of newly drawn candidates. Each candidate θ_j is associated with a non-negative weight w_j , and the goal is to select exactly K survivors such that the marginal inclusion probability of each element is proportional to w_j . This is achieved by assigning to every candidate a random key (11)

$$k_j = u_j^{1/w_j}, \quad u_j \sim \mathcal{U}(0, 1), \quad (4)$$

retaining the K elements with the smallest keys. The resulting selection is unbiased: the expected inclusion probability of each element is proportional to its weight, while the reservoir size remains fixed (11). In our context, WRS allows efficient reallocation of computational effort toward projections with higher contribution to the optimization loss, while preserving Monte Carlo unbiasedness.

In Fig. 2, we present an overview of our proposed method. Let $\mathcal{P}$ be a pool of $K + M$ 1-D projection directions $\theta \in S^{d-1}$. For every $\theta \in \mathcal{P}$ we compute the 1-D p -power Wasserstein cost

$$D(\theta) = W_p(\pi_\theta \mu, \pi_\theta \nu) \quad (5)$$

Then each distance calculation performs the following steps.

Step 0: Time-decay reweighting. As SWD is most often used in optimization scenarios, we introduce τ as the time decay constant to place more emphasis on

newer random directions to account for the shifting optimization field. Before drawing new candidates, we *age* the stored reservoir weights:

$$\tilde{w}_i \leftarrow w_i \exp[-(t - t_i)/\tau], \quad \tilde{k}_i \leftarrow k_i \exp[-(t - t_i)/\tau], \quad (6)$$

where t_i is the step when θ_i entered the reservoir. This exponential decay (enabled when $\tau > 0$) steadily forgets stale projections so that the sampler adapts to *nonstationary* optimization trajectories.

Step 1: Reservoir construction. At optimization step t we maintain a persistent reservoir $\mathcal{R}_{t-1}$ of K directions from the previous optimization step $t-1$ and draw M new directions $\mathcal{N}_t$. The candidate set is $\mathcal{P}_t = \mathcal{R}_{t-1} \cup \mathcal{N}_t$.

Step 2: Weighted reservoir sampling. Following Efraimidis & Spirakis (11) we assign each $\theta \in \mathcal{P}_t$ a key $k(\theta) = u^{1/D(\theta)}$, $u \sim \mathcal{U}(0, 1)$ and keep the K largest keys. This selects each θ with probability $q(\theta) = \frac{D(\theta)}{\sum_{\theta' \in \mathcal{P}_t} D(\theta')}$, *i.e.* proportionally to its contribution to the loss.

Step 3: Self-normalized weights. For the survivors $\{\theta_i\}_{i=1}^K$ we calculate the loss with weights as:

$$\hat{S}_p(\mu, \nu) = \sum_{i=1}^K \frac{1/q(\theta_i)}{\underbrace{\sum_j 1/q(\theta_j)}_{w_i}} D(\theta_i), \quad (7)$$

This loss can then be used in optimizations with a detached gradient calculation for importance weights w_i . This remains an *unbiased* MC estimate of $\mathbb{E}_\theta[W_p]$, while focusing on the computation on projections that highlight larger differences.

Step 4: ESS-based reservoir reset. We monitor the effective sample size (3) $\text{ESS} = (\sum_i w_i)^2 / \sum_i w_i^2$ and flush the reservoir whenever $\text{ESS} < \alpha K$ ($\alpha = 0.5$ in all experiments), preventing weight collapse.

Handling unequal sample counts. If the sample sizes of X and Y are not equal, we repeat the shorter projection so that both 1-D arrays have the same length before sorting. Repetitions are picked uniformly with replacement, following Elnekave *et al.* (12).

Complexity. The dominant cost is sorting n scalars for each of $(K+M)$ directions: $\mathcal{O}((K+M)n \log n)$. Key generation, ESS evaluation, and gradient accumulation are $\mathcal{O}(K+M+d)$.

Algorithmic summary. The overview of the algorithm is shown in Algorithm 1.

3.3 Applications

We showcase our SWD modifications on two real-world applications.

Color Correction. Often in movie production two shots have to be matched in terms of their color appearance. We tackle this approach by matching a source image using a Color Decision List (CDL) (38) with a reference image. The source

Algorithm 1 ReSWD estimator (per optimisation step)

Require: batches $\{x_n\}_{n=1}^{N_\mu}$, $\{y_m\}_{m=1}^{N_\nu}$, reservoir $\mathcal{R}_{t-1}$, hyper-params K, M, p, α, τ

- 1: **for** $(\theta_i, w_i, k_i, t_i) \in \mathcal{R}_{t-1}$ **do** ▷ Step 0: Time-decay
- 2: $w_i \leftarrow w_i \exp[-(t-t_i)/\tau]$
- 3: $k_i \leftarrow k_i \exp[-(t-t_i)/\tau]$
- 4: **end for**
- 5: $\mathcal{N}_t \leftarrow \text{DRAW DIRECTIONS}(M, d)$ ▷ Step 1: Reservoir Construction
- 6: $\hat{x}, \hat{y} \leftarrow \text{DUPLICATE SAMPLES}(x_n, y_m)$ ▷ See “Handling unequal sample counts”
- 7: **for** $\theta \in \mathcal{N}_t$ **do** ▷ Step 2: Reservoir Sampling
- 8: $D(\theta) \leftarrow W_p(\pi_\theta \hat{x}, \pi_\theta \hat{y})$
- 9: $k_t(\theta) \leftarrow u^{1/D(\theta)}, u \sim \mathcal{U}(0, 1)$
- 10: **end for**
- 11: $k(\theta) \leftarrow R_{t-1}[k] \cup k_t(\theta)$ ▷ Join reservoir with new keys
- 12: $\mathcal{R}_t \leftarrow \mathcal{R}_{t-1}[\theta] \cup \mathcal{N}_t$ ▷ Join reservoir with new directions
- 13: $\mathcal{R}_t \leftarrow K$ directions with largest $k(\theta)$
- 14: compute $q(\theta)$, weights w_i , and $\widehat{W}_p$ via Eq. (7) ▷ Step 3: Self normalizing weights
- 15: **if** $\text{ESS} < \alpha K$ **then return** $\widehat{W}_p, \emptyset$ ▷ Step 4: ESS-Reset
- 16: **else return** $\widehat{W}_p, \mathcal{R}_t$
- 17: **end if**

image x is first passed through a differentiable CDL implementation, applying the *slope* s , *offset* o , *power* p and λ *saturation* adjustments according to $x' = S((s \times x + o)^p; \lambda)$. The saturation is defined as $S(x; \lambda) = L(x) + \lambda(x - L(x))$ and $L(x) = 0.2126x_r + 0.7152x_g + 0.0722x_b$, where r, g, b define the color channels. The source and reference are then converted to the CIELAB space. We found that this space improves the color accuracy, which is consistent with a recent paper proposing a perceptual color metric (16). As we perform the matching on a pixel level, we find that lower resolutions result in faster matching speeds with similar performance. Thus, we leverage a resolution of 128 in the maximum dimension. Within 150 steps, the CDL parameters are optimized and can be applied to the full-res image or even a video clip.

Diffusion Guidance. Lobashev *et al.* (26) proposed to incorporate SWD guidance in diffusion models to shift the final generation towards a certain color distribution. The main concept is to optimize a small offset on top of the current latents with i SWD steps in each diffusion step. Here, we use the predicted x_0 in each step and decode it using the VAE decoder. This prediction is then matched with the reference image. We apply this approach to the more recent flow matching model SD3.5 (1). This required several modifications. Lobashev *et al.* (26) calculated the full gradient from the input, through the VAE and the U-net. With the increased complexity of larger transformer models, we opted to employ a gradient stop after the backbone similar to SDS (40). Hence, we only backpropagate through the VAE decoder. Similarly to our color-correction approach, we also opted to perform the matching in the CIELAB space. In addition, we replace the simple gradient descent of Lobashev *et al.* (26) with an Adam optimizer. With these changes, we use a learning rate of $3e-3$ and a total

of 6 steps for 95% of the total denoising steps. We do not reset the reservoir after each denoising step, but twice during the generation, as we only perform 6 SWD steps, which would not suffice to build a reservoir. This general method can then be applied to medium, large, and large-turbo SD3.5 models using the recommended CFG and step counts.

4 Results

We perform synthetic as well as real-world tests with ReSWD and study the influence of our main hyperparameter, the number of fresh candidates in each optimization step.

Table 1: Comparison on 1D distribution matching. Mean- W_1 over 1000 distribution matches for various methods alongside the respective running time. Here, we can see that ours provides the best performance with a comparatively low run-time cost.

Method	Mean- W_1 [10^3] ↓	σ [10^3] ↓	95% CI [10^3] ↓	Grad SNR ↑	Time [ms] ↓
SWD	0.733	0.050	± 0.003	0.215	1.03
(30) (LCV)	0.735	0.044	± 0.003	0.218	1.81
(30) (UCV)	0.726	0.045	± 0.003	0.200	2.19
(10) (Max-SW)	29.152	38.278	± 2.373	0.0497	1.13
(29) (IS-EBSW-e)	0.698	0.075	± 0.005	0.2291	1.42
(28) (QMC)	0.670	0.043	± 0.003	0.208	1.38
ReSWD	0.622	0.075	± 0.005	0.278	1.92
ReSWD + QMC	0.610	0.076	± 0.005	0.274	2.10

Table 2: Influence of fresh candidates. With a fixed budget of 64 projections, we analyze the influence of the fresh candidates.

# Candidates	Mean- W_1 [10^3] ↓	σ [10^3] ↓	95% CI [10^3] ↓	Time [ms] ↓
2	0.721	0.257	± 0.016	1.99
4	0.673	0.151	± 0.009	1.96
8	0.622	0.075	± 0.005	1.92
16	0.746	0.122	± 0.008	1.91
32	1.192	1.283	± 0.080	1.98
48	2.122	3.280	± 0.203	1.93
56	3.811	3.317	± 0.206	1.85

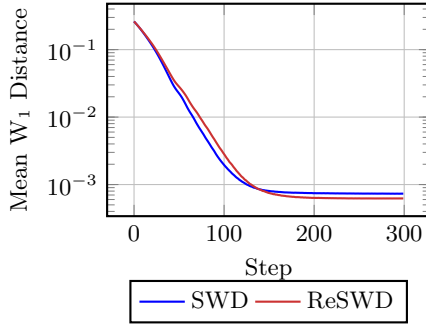


Fig. 3: True Wasserstein metric over steps. The effect of reservoir warmup is clear when comparing ReSWD with the true Wasserstein distance during optimization. Initially, ReSWD performs slightly worse due to lower projections in the loss, but can outperform SWD in the end.

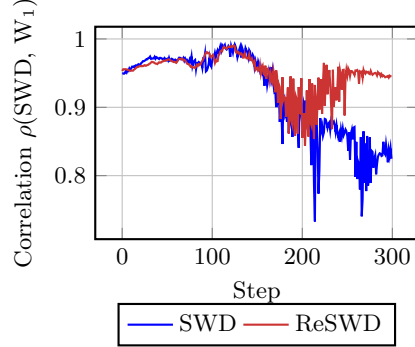


Fig. 4: Pearson correlation with the true Wasserstein Distance. Our method achieves a high correlation with the true Wasserstein loss, while improving upon pure SWD (See Fig. 3). This indicates the unbiased nature of our proposed method.

General Distribution Matching. For a general test of our proposed method we create 1000 $d = 3$ distribution pairs (normal, uniform, bimodal normal) and align them based on 1024 samples in 300 steps. All methods leverage 64 projections. In Table 1, our method clearly outperforms previous SWD techniques with a slight additional performance cost. To evaluate the optimization behavior of our method, we measure the gradient signal-to-noise ratio (SNR) by taking the mean gradient as the signal and the standard deviation as the noise. A higher value here indicates that the method provides more stable gradients. Additionally, we calculate the mean W_1 true Wasserstein score for each dimension in each step. This allows us to plot the convergence behavior of our method in Fig. 3. Here, it is evident that pure SWD and our ReSTIR-based modification have a similar trajectory, but building the reservoir initially results in slightly slower trajectory, which produces better results after roughly 140 steps. This also allows us to investigate the correlation between our proposed loss and the true Wasserstein distance. We also investigate the combination of ReSWD with QMC in Table 1 and observe an additional performance improvement with minimal additional overhead. In Fig. 4, our method achieves a high correlation with the true loss, indicating the unbiased nature of our method.

Color Matching. To evaluate the color matching based on our differentiable color correction pipeline, we created a dataset of 10 scenes with two different illumination settings each. For each setting, we also take a photo using a calibration color chart. We match the illumination pairs and compute the following metrics from the color checker data extracted from the additional image pairs: Color data PSNR and RMSE (between adjusted and ground truth color checker images), transform error, i.e. deviation from an identity transform between target and adjusted colors as RMSE, as well as an adapted CTQM metric (35) describing the overall color transfer quality. We select Reinhard *et al.* (43) and Nguyen *et al.* (34) as conventional baselines, Yoo *et al.* (50), Ho *et al.* (18) and Larchenko *et al.* (23) as neural network based comparisons, as well as our proposed method but without the addition of ReSWD. Table 3 shows that our approach achieves the lowest transform error in a competitive runtime. In Fig. 5, we present the visual comparison between the methods. It is also evident that our method outperforms the existing baseline in qualitative results as well. Additionally, it is worth pointing out that our results using a parametric model also offer real-world advantages because of better pipeline integration and further manual edits being possible.

Diffusion Guidance. We select Lobashev *et al.* (26) as our main comparison. Furthermore, we also implemented our modifications to the matching in the SDXL model to compare our alterations with the baseline, which is based on SDXL. This delineates the influence of the improved SD3.5 model from our matching and alterations. For evaluations, we follow Lobashev *et al.* and select 1000 prompts from the ContraStyles dataset¹ and 1000 images from Unsplash Lite². We provide the W_2 Wasserstein distance, CLIP-IQA (48) and CLIP-T (42)

¹ <https://huggingface.co/datasets/tomg-group-umd/ContraStyles>

² <https://unsplash.com/data>

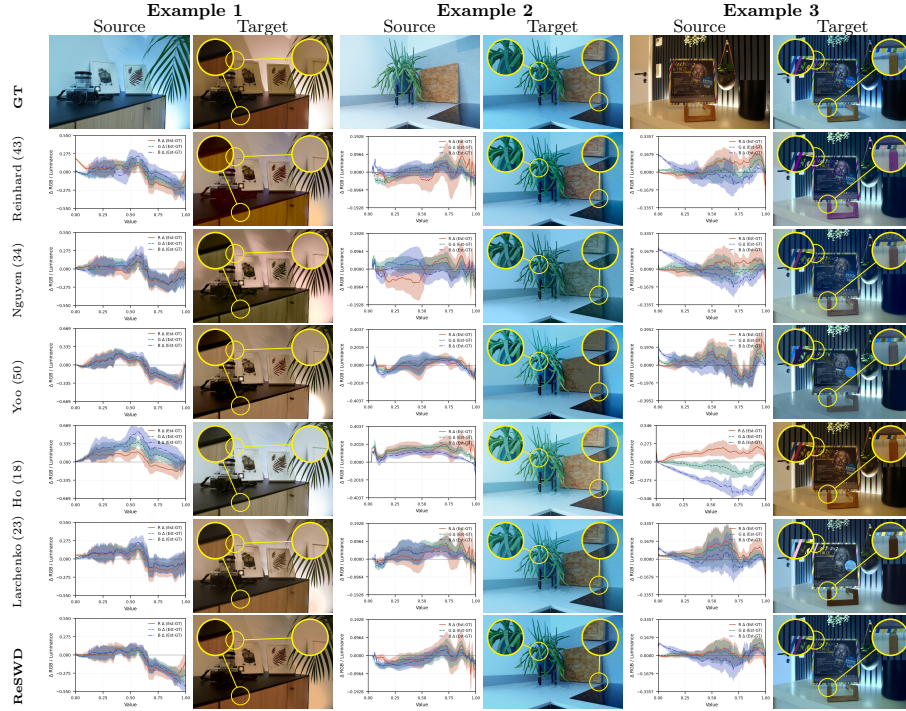


Fig. 5: Color Matching. Notice the consistent images our method produces with accurate shot matching. ReSWD consistently produces good results which match the final color distribution and without introducing any artifacts. We highlighted challenging areas which our method handled well. This is also evident by comparing the delta with the ground truth for each color channel. Ours produces fewer deviations with less variance. Note that the reference implementations of Larchenko *et al.* (23), Yoo *et al.* (50) and Nguyen *et al.* (34) do not support high resolution images.

to present color matching performance, highlight the quality of the generation and prompt adherence, respectively.

In Table 4 our modifications clearly outperform the prior state-of-the-art technique of Lobashev *et al.* Even with the same base model (SDXL), our method is faster and provides better results, and this holds true with additionally ablating the optimizer improvements and without the gradient stopping. As seen here, our gradient stopping implementation does not influence on the guidance negatively but creates a large speed-up as well as a significant memory reduction which allows us to use larger base models. With the upgraded base model, ReSWD clearly outperforms the previous method. Here, it is interesting that the distilled Turbo variants provides better quality, which we attribute to the straighter trajectories.

We show example generations of this dataset in Fig. 6 along with the reference and the prompts. As evident, the improved generation capacity of SD3.5 along

Table 3: Comparison on color matching. Here, we compare ReSWD with a baseline SWD method and prior color matching works by Reinhard *et al.* (43), Nguyen *et al.* (34) as well as neural network based methods by Yoo *et al.* (50), Ho *et al.* (18) and Larchenko *et al.* (23). We report the errors between the color charts of adjusted and ground truth images as PSNR, transform and CDL error and the image quality as CTQM. Additionally, we compare the runtimes of the methods.

Method	Color PSNR $\uparrow \pm \sigma$	Transform err. (RMSE) $\downarrow \pm \sigma$	CDL err. (RMSE) $\downarrow \pm \sigma$	CTQM $\uparrow \pm \sigma$	Time per match [s] $\downarrow$
Reinhard (43)	21.94 $\pm$ 5.10	0.31 $\pm$ 0.28	0.14 $\pm$ 0.09	5.12 $\pm$ 1.82	1
Nguyen (34)	18.76 $\pm$ 3.69	1.27 $\pm$ 1.32	0.33 $\pm$ 0.23	5.16 $\pm$ 1.78	3
Yoo (50)	20.97 $\pm$ 5.44	0.44 $\pm$ 0.38	0.21 $\pm$ 0.17	5.09 $\pm$ 1.80	20
Ho (18)	10.43 $\pm$ 2.53	0.54 $\pm$ 0.29	0.45 $\pm$ 0.25	5.11 $\pm$ 1.77	2
Larchenko (23)	14.80 $\pm$ 3.47	0.47 $\pm$ 0.34	0.20 $\pm$ 0.14	5.05 $\pm$ 1.82	24
Ours with SWD	24.30 $\pm$ 6.12	0.34 $\pm$ 0.30	0.11 $\pm$ 0.08	5.15 $\pm$ 1.81	5
ReSWD	24.64 $\pm$ 5.70	0.31 $\pm$ 0.31	0.10 $\pm$ 0.07	5.17 $\pm$ 1.80	5

Table 4: Comparison on diffusion guidance. Here, we compare ReSWD with Lobashev *et al.* (26) on color guidance in image generation. As our main proposed method uses the more recent SD3.5, we also compare with our modifications on SDXL. Even on the same base model our modifications improve upon the base technique.

Method	CLIP-IQA $\uparrow \pm \sigma$	CLIP-T $\uparrow \pm \sigma$	Mean-W ₂ [10 ²] $\downarrow \pm \sigma$	Time per generation [s] $\downarrow$
Lobashev (26)	0.671 $\pm$ 0.224	14.861 $\pm$ 2.75	1.941 $\pm$ 1.924	124
ReSWD w/o Adam & GradStop @ SDXL	0.691 $\pm$ 0.225	14.848 $\pm$ 2.795	1.832 $\pm$ 1.142	126
ReSWD w/o Adam @ SDXL	0.681 $\pm$ 0.217	14.771 $\pm$ 2.732	1.846 $\pm$ 1.129	34
ReSWD @ SDXL	0.696 $\pm$ 0.203	14.870 $\pm$ 2.710	1.213 $\pm$ 0.561	34
ReSWD @ SD3.5-medium	0.786 $\pm$ 0.167	14.783 $\pm$ 2.768	0.815 $\pm$ 0.650	32
ReSWD @ SD3.5-turbo	0.800 $\pm$ 0.137	14.882 $\pm$ 2.774	0.675 $\pm$ 0.397	4
ReSWD @ SD3.5-large	0.793 $\pm$ 0.179	14.817 $\pm$ 2.745	0.55 $\pm$ 0.574	69

with our necessary changes resulted in drastically better generation which are faithful to the color distribution defined by the input image.

Influence of Dimensionality. Our method is designed to maintain good directions in a reservoir that will be reused during optimization in an unbiased way. However, with higher dimensionalities the chance of drawing a good kernel is drastically lowered. Therefore, the reservoir will not be updated and the directions become stale. This effect can be seen in Table 5, where our solution drastically degrades in a 12 dimension problem. This could be circumvented with more aggressive ESS clearings or time-based decay, but at this point our method would revert back to plain SWD. We can also combine our method with QMC (28) to alleviate some of the issues, but even with that the chance of sampling good filters is drastically decreased.

Ablation. In Table 2, we assess the influence of the number of new candidates in our general distribution matching task. Too few new candidates starve the reservoir of new directions, resulting in static directions during optimization. Too many new candidates reduce the reservoir size, which drives the optimization solely, resulting in too few directions supporting the optimization. With 8 candidates, we achieved the best overall results in all applications given 64 total projections.

Limitations. Our current technique is limited to simple matrix projections, and an extension to learned convolution kernels similar to Elnekave *et al.* (12) did not provide satisfactory results, as similar to the dimensionality problem

Table 5: Influence of dimensions. Increasing the dimensionality of the problem introduces a decrease in quality. This comes from the fact that the reservoir mechanism becomes less likely to receive good candidates as randomly sampling them becomes unlikely; thus, the reservoir is not updated as often, and the direction becomes stale.

Method	Mean- W_1 [10^3] ↓						
	3	6	12	24	32	48	64
SWD	0.733	0.945	1.190	3.386	4.863	6.503	7.922
LCV (30)	0.735	0.963	2.203	4.186	5.271	7.152	8.826
UCV (30)	0.726	0.941	1.935	3.815	4.919	6.573	8.009
Max-SW (10)	29.152	72.572	127.321	172.492	186.748	201.986	210.266
IS-EBSW-e (29)	0.698	0.887	1.536	2.812	3.518	4.774	6.044
QMC (28)	0.670	0.736	1.166	3.904	4.899	6.509	7.854
ReSWD	0.622	0.939	9.768	15.350	17.114	19.578	21.446
ReSWD + QMC	0.615	0.823	4.532	8.953	14.389	18.375	20.862

for 1D distributions, the search space for kernels is too large to randomly obtain drastically better kernels. This degraded the results similar to limiting the projection directions and reducing the redrawing of new directions to every few optimization steps.

5 Conclusion

Our novel combination of real-time rendering-influenced variance reduction techniques in SWD optimization offers a more efficient and unbiased solution compared to other recent variance reduction techniques. This results in reduced overhead while maintaining optimal performance. Our technique has demonstrated superior performance in various real-world and synthetic applications, achieving state-of-the-art results.

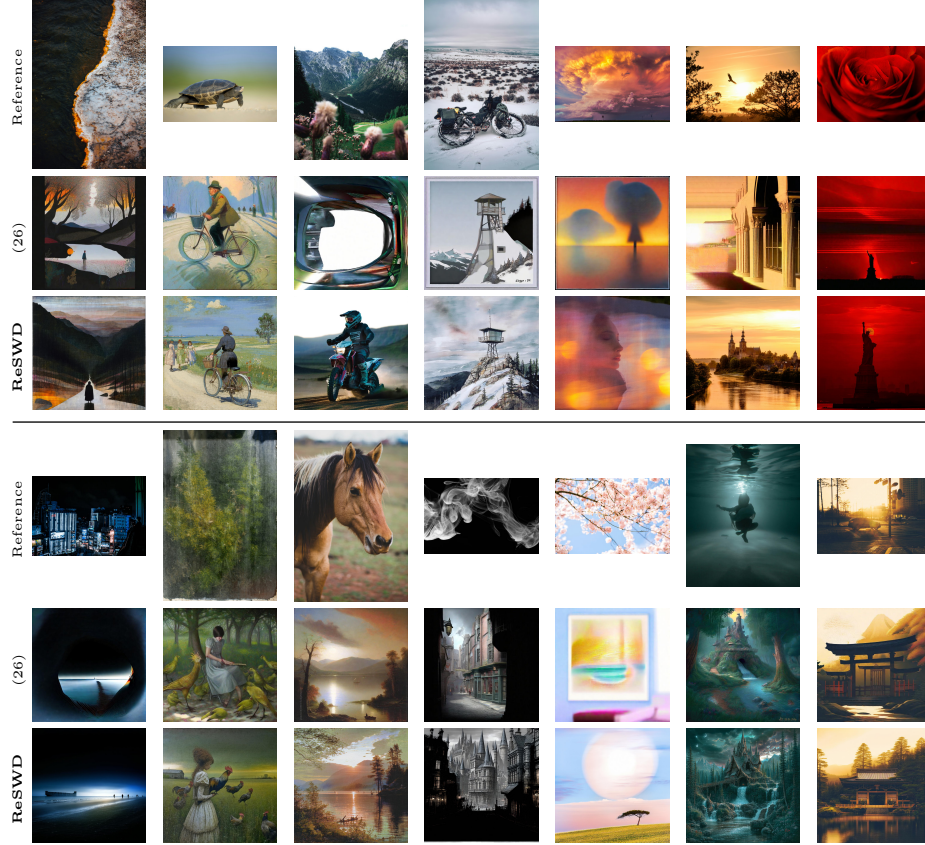


Fig. 6: Diffusion Guidance. Visualization of our color guidance with SD3.5-Large. Notice, how closely the colors match the reference and the detailed generation quality of our method. Also we noticed that Lobashev *et al.* (26) sometimes degenerates (Third example top, First example bottom) whereas ReSWD handles these challenging guidance signal well.

Bibliography

- [1] AI, S.: Stable diffusion 3.5. <https://github.com/Stability-AI/stablediffusion> (2025), version 3.5, accessed June 2025
- [2] An, J., Xiong, H., Huan, J., Luo, J.: Ultrafast photorealistic style transfer via neural architecture search. In: AAAI (2020)
- [3] Bitterli, B., Wyman, C., Pharr, M., Shirley, P., Lefohn, A., Jarosz, W.: Spatiotemporal Reservoir Resampling for Real-Time Ray Tracing with Dynamic Direct Lighting. ACM SIGGRAPH (2020)
- [4] Bonet, C., Drumetz, L., Courty, N.: Sliced-wasserstein distances and flows on cartan-hadamard manifolds. *Journal of Machine Learning Research* **26**, 1–76 (2025)
- [5] Bonneel, N., Rabin, J., Peyré, G., Pfister, H.: Sliced and radon wasserstein barycenters of measures. *Journal of Mathematical Imaging and Vision* (2015)
- [6] Chao, M.T.: A general purpose unequal probability sampling plan. *Biometrika* (1982)
- [7] Chiu, T.Y., Gurari, D.: Photowct2: Compact autoencoder for photorealistic style transfer resulting from blockwise training and skip connections of high-frequency residuals. In: WACV (2022)
- [8] Cozzi, G., Santambrogio, F.: Long-time asymptotics of the sliced-wasserstein flow. *SIAM Journal on Imaging Sciences* (2024), also available as arXiv:2405.06313
- [9] Cuturi, M.: Sinkhorn distances: Lightspeed computation of optimal transport. In: NeurIPS (2013)
- [10] Deshpande, I., Hu, Y.T., Sun, R., Pyrros, A., Siddiqui, N., Koyejo, S., Zhao, Z., Forsyth, D., Schwing, A.G.: Max-sliced wasserstein distance and its use for gans. In: CVPR (2019)
- [11] Efraimidis, P.S., Spirakis, P.G.: Weighted random sampling with a reservoir. *Information Processing Letters* (2006)
- [12] Elnekave, A., Weiss, Y.: Generating natural images with direct patch distributions matching. In: ECCV (2022)
- [13] Esser, P., Rombach, R., Ommer, B.: Taming transformers for high-resolution image synthesis. In: CVPR (2021)
- [14] Feydy, J., Séjourné, T., Vialard, F., Amari, S., Trounev, A., Peyré, G.: Interpolating between optimal transport and mmd using sinkhorn divergences. In: AISTATS (2019)
- [15] Gatys, L.A., Ecker, A.S., Bethge, M.: Image style transfer using convolutional neural networks. In: CVPR (2016)
- [16] He, J., Wang, Z., Wang, L., Liu, T.I., Fang, Y., Sun, Q., Ma, K.: Multi-scale sliced Wasserstein distances as perceptual color difference measures. In: ECCV (2024)
- [17] Heitz, E., Vanhoey, K., Chambon, T., Belcour, L.: A sliced wasserstein loss for neural texture synthesis. In: CVPR (2021)

- [18] Ho, M.M., Zhou, J.: Deep preset: Blending and retouching photos with color style transfer. In: Proceedings of the IEEE/CVF Winter Conference on Applications of Computer Vision (WACV). pp. 2113–2121 (January 2021)
- [19] Hong, K., Jeon, S., Yang, H., Fu, J., Byun, H.: Domain-aware universal style transfer. In: CVPR (2021)
- [20] Hu, E.J., Shen, Y., Wallis, P., Allen-Zhu, Z., Li, Y., Wang, L., Wang, W.: Lora: Low-rank adaptation of large language models. In: ICLR (2021)
- [21] Huang, X., Belongie, S.: Arbitrary style transfer in real-time with adaptive instance normalization. In: CVPR (2017)
- [22] Kolouri, S., Nadjahi, K., Simsekli, U., Badeau, R., Rohde, G.: Generalized sliced wasserstein distances. NeurIPS (2019)
- [23] Larchenko, M., Lobashev, A., Guskov, D., Palyulin, V.V.: Color transfer with modulated flows. In: AAAI. arxiv (2025). <https://doi.org/10.1609/aaai.v39i4.32470>, <http://arxiv.org/abs/2503.19062>
- [24] Li, Y., Fang, C., Yang, J., Wang, Z., Lu, X., Yang, M.H.: Universal style transfer via feature transforms. NeurIPS (2017)
- [25] Li, Y., Liu, M.Y., Li, X., Yang, M.H., Kautz, J.: A closed-form solution to photorealistic image stylization. In: ECCV (2018)
- [26] Lobashev, A., Larchenko, M., Guskov, D.: Color conditional generation with sliced wasserstein guidance. arXiv.org (2025). <https://doi.org/10.48550/arXiv.2503.19034>
- [27] Luan, F., Paris, S., Shechtman, E., Bala, K.: Deep photo style transfer. In: CVPR (2017)
- [28] Nguyen, K., Bariletto, N., Ho, N.: Quasi-monte carlo for 3d sliced wasserstein. ICLR (2024)
- [29] Nguyen, K., Ho, N.: Energy-based sliced wasserstein distance. NeurIPS **36** (2023)
- [30] Nguyen, K., Ho, N.: Sliced Wasserstein Estimation with Control Variates. In: ICLR (2024)
- [31] Nguyen, K., Ho, N., Pham, T., Bui, H.: Distributional sliced-wasserstein and applications to generative modeling. In: ICLR (2021)
- [32] Nguyen, K., Ren, T., Nguyen, H., Rout, L., Nguyen, T., Ho, N.: Hierarchical sliced wasserstein distance. In: ICLR (2023)
- [33] Nguyen, K., Zhang, S., Le, T., Ho, N.: Sliced wasserstein with random-path projecting directions. ICML (2024)
- [34] Nguyen, R.M.H., Kim, S.J., Brown, M.S.: Illuminant Aware Gamut-Based Color Transfer. Comput. Graph. Forum **33**(7), 319–328 (oct 2014). <https://doi.org/10.1111/cgf.12500>
- [35] Panetta, K., Bao, L., Agaian, S.: Novel multi-color transfer algorithms and quality measure. IEEE Transactions on Consumer Electronics **62**(3), 292–300 (aug 2016). <https://doi.org/10.1109/TCE.2016.7613196>
- [36] Park, S., Slepčev, D.: Geometry and analytic properties of the sliced wasserstein space. Journal of Functional Analysis (2025)
- [37] Peyré, G., Cuturi, M.: Computational optimal transport. Foundations and Trends in Machine Learning **11**(5-6), 355–607 (2019)

- [38] Pines, J., Reisner, D.: ASC Color Decision List (ASC CDL) transfer functions and interchange syntax, version 1.2. Tech. rep., American Society of Cinematographers Technology Committee, DI Subcommittee (2009)
- [39] Pitie, F., Kokaram, A.C., Dahyot, R.: N-dimensional probability density function transfer and its application to color transfer. In: ICCV (2005)
- [40] Poole, B., Jain, A., Barron, J.T., Mildenhall, B.: Dreamfusion: Text-to-3d using 2d diffusion. In: ICLR (2022). <https://doi.org/10.48550/arXiv.2209.14988>
- [41] Rabin, J., Peyré, G., Delon, J., Bernot, M.: Wasserstein barycenter and its application to texture mixing. In: Scale Space and Variational Methods in Computer Vision (2012)
- [42] Radford, A., Kim, J.W., Hallacy, C., Ramesh, A., Goh, G., Agarwal, S., Sastry, G., Askell, A., Mishkin, P., Clark, J., et al.: Learning transferable visual models from natural language supervision. In: ICML (2021)
- [43] Reinhard, E., Ashikhmin, M., Gooch, B., Shirley, P.: Color transfer between images. *IEEE Computer Graphics and Applications* **21**(5), 34–41 (2001)
- [44] Rombach, R., Blattmann, A., Lorenz, D., Esser, P., Ommer, B.: High-resolution image synthesis with latent diffusion models. In: CVPR (2022)
- [45] Tran, H., Bai, Y., Kothapalli, A., Shahbazi, A., Liu, X., Martin, R.D., Kolouri, S.: Stereographic spherical sliced wasserstein distances. *ICML* (2024)
- [46] Vauthier, C., Korba, A., Mérigot, Q.: Properties of wasserstein gradient flows for the sliced-wasserstein distance. *arXiv* (2025)
- [47] Villani, C., et al.: Optimal transport: old and new, vol. 338. Springer (2008)
- [48] Wang, J., Chan, K.C., Loy, C.C.: Exploring clip for assessing the look and feel of images. In: AAAI (2023)
- [49] Wu, R., Liu, R., Vondrick, C., Zheng, C.: Sin3dm: Learning a diffusion model from a single 3d textured shape. In: ICLR (2024)
- [50] Yoo, J., Uh, Y., Chun, S., Kang, B., Ha, J.W.: Photorealistic style transfer via wavelet transforms. In: International Conference on Computer Vision (ICCV) (2019)
- [51] Zhang, L., Rao, A., Agrawala, M.: Adding conditional control to text-to-image diffusion models. In: CVPR (2023)
- [52] Zhang, M., Chen, H., Zeng, Y., Cheng, M., Yang, L., Xu, Z., Gu, J.: IP-Adapter: Text compatible image prompt adapter for text-to-image diffusion models. In: ICCV (2023)