

Table 1: Comparison of performance: w/o DoSG vs. different accelerated samplers on the *RealSR* and *DRealSR* datasets with same model(RealESRNet preprocessing + DiffBIR). In all setups, inference is carried out over 5 steps.

Method	Corr. Sampler	RealSR Dataset				DrealSR Dataset			
		CLIPQA $\uparrow$	MUSIQ $\uparrow$	MANIQA $\uparrow$	TOPIQ $\uparrow$	CLIPQA $\uparrow$	MUSIQ $\uparrow$	MANIQA $\uparrow$	TOPIQ $\uparrow$
w/o DoSG (DDPM)	DDIM($\eta = 1$)	0.5176	57.95	0.4293	0.5286	0.4732	48.22	0.3518	0.4316
	EDM	0.5351	62.08	0.4445	0.5789	0.5341	53.13	0.3917	0.5082
	DPM-Solver++ -3	0.5323	62.67	0.4384	0.5807	0.5379	54.09	0.3932	0.5180
w/ DoSG (DoSSR)	DoS SDE-Solver -1	0.6874	66.55	0.5574	0.6588	0.5907	59.12	0.4686	0.5907
	DoS SDE-Solver -2	0.7025	69.27	0.5794	0.6966	0.6749	64.09	0.5196	0.6571
	DoS SDE-Solver -3	0.7025	69.42	0.5781	0.6985	0.6776	64.40	0.5214	0.6618

Table 2: Comparison of performance with other methods on the *RealSRSet* and *RealSR* datasets. NFE represents the number of function evaluations in the inference. * involves retraining using the same training data and identical network architecture as our model.

Method	Training Dataset	RealSRSet		RealSR		NFE $\downarrow$
		MUSIQ $\uparrow$	MANIQA $\uparrow$	MUSIQ $\uparrow$	MANIQA $\uparrow$	
ColdDiff*	DIV2K+DIV8K+Flickr2K+OST($\sim 15k$)	58.19	0.3194	47.42	0.2783	5
ResShift*	DIV2K+DIV8K+Flickr2K+OST($\sim 15k$)	63.90	0.4505	56.01	0.4001	5
DoSSR	DIV2K+DIV8K+Flickr2K+OST($\sim 15k$)	73.35	0.6169	69.42	0.5781	5
FlowIE	ImageNet($\sim 1280k$)	61.63	0.3611	56.51	0.3284	1
FlowIE*	DIV2K+DIV8K+Flickr2K+OST($\sim 15k$)	60.48	0.3644	50.82	0.3228	1
DoSSR	DIV2K+DIV8K+Flickr2K+OST($\sim 15k$)	69.42	0.5554	62.69	0.5115	1

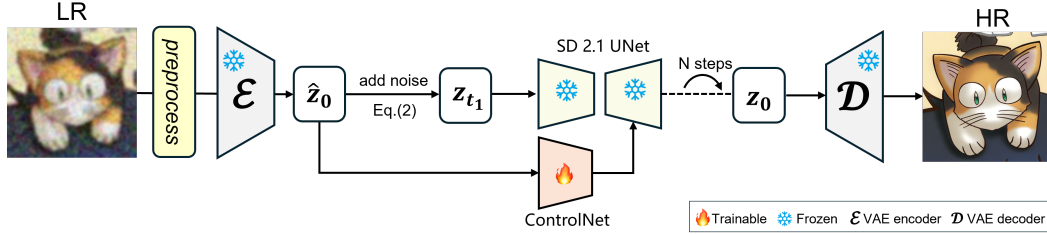


Figure 1: The overall framework of DoSSR. During training, we introduce noise to facilitate the gradual transition from the HR to LR domain, integrating it with the standard diffusion process, and incorporate preprocessed LR as a conditioning input for the denoising process, following the ControlNet approach. During inference, we add noise to LR latent according to Eq.(2) and perform inference starting from t_1 .