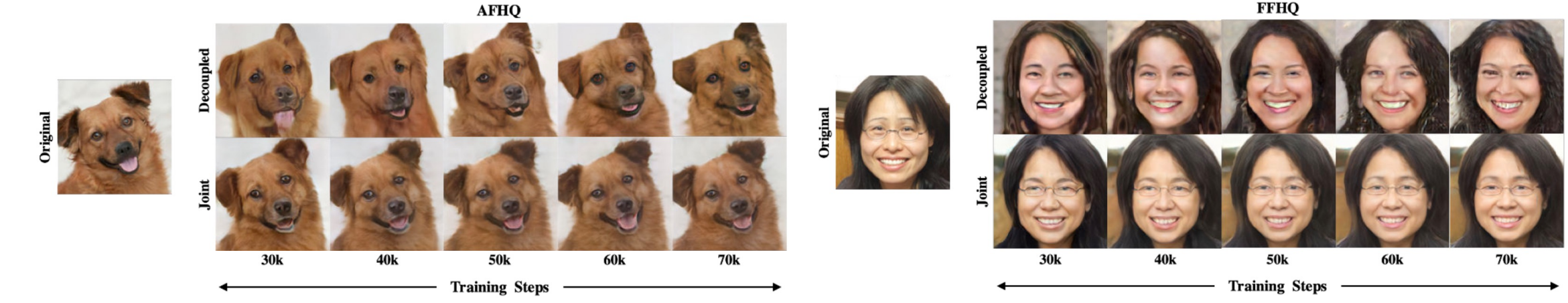


AE-StyleGAN: Improved Training of Style-Based Auto-Encoders

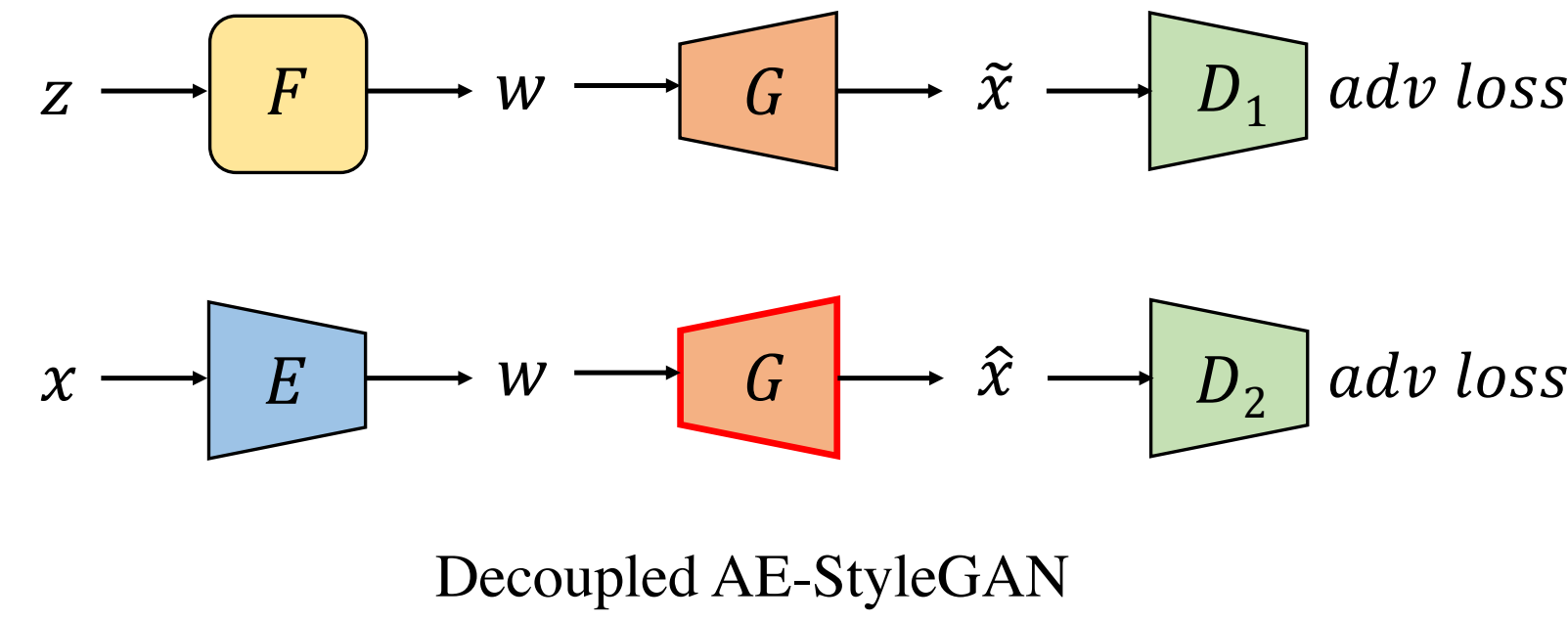
Ligong Han^{*1}, Sri Harsha Musunuri^{*1}, Martin Renqiang Min², Ruijiang Gao³, Yu Tian¹, Dimitris Metaxas¹
¹Rutgers University ²NEC Labs America ³The University of Texas at Austin



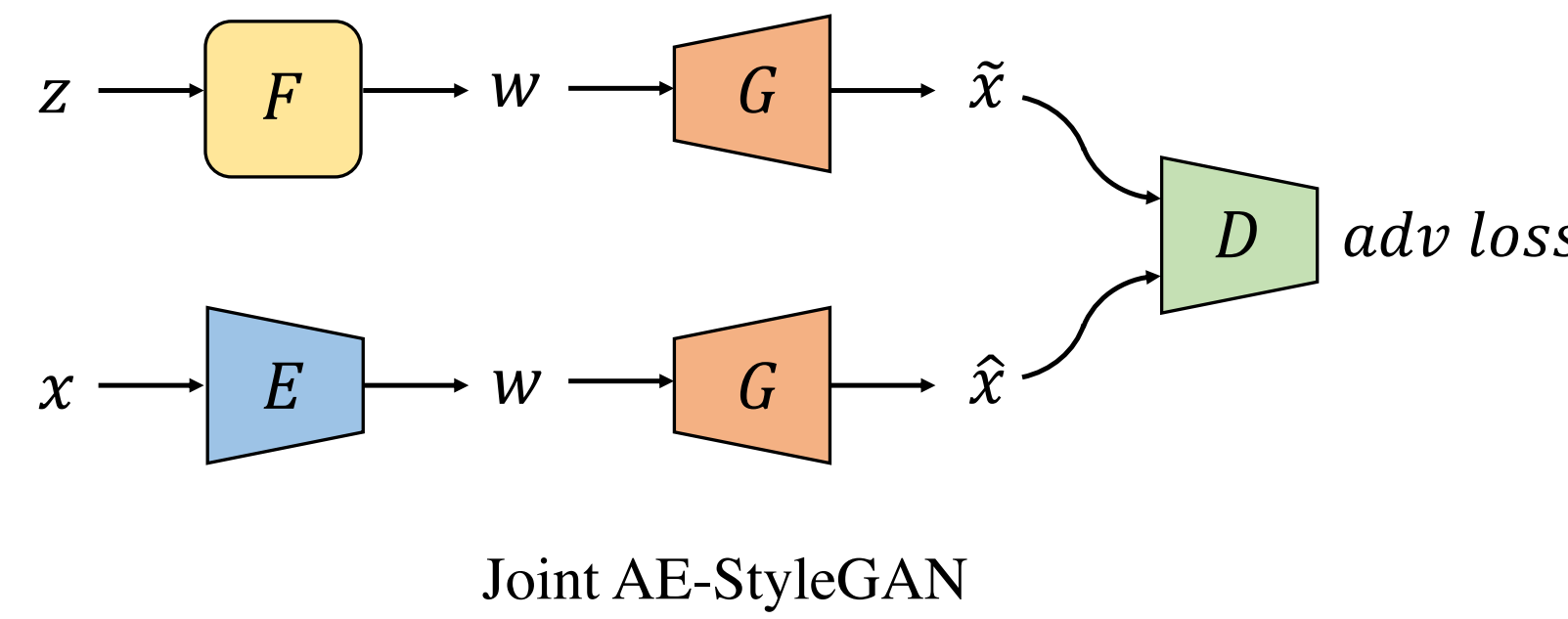
Introduction



Decoupled AE-StyleGAN

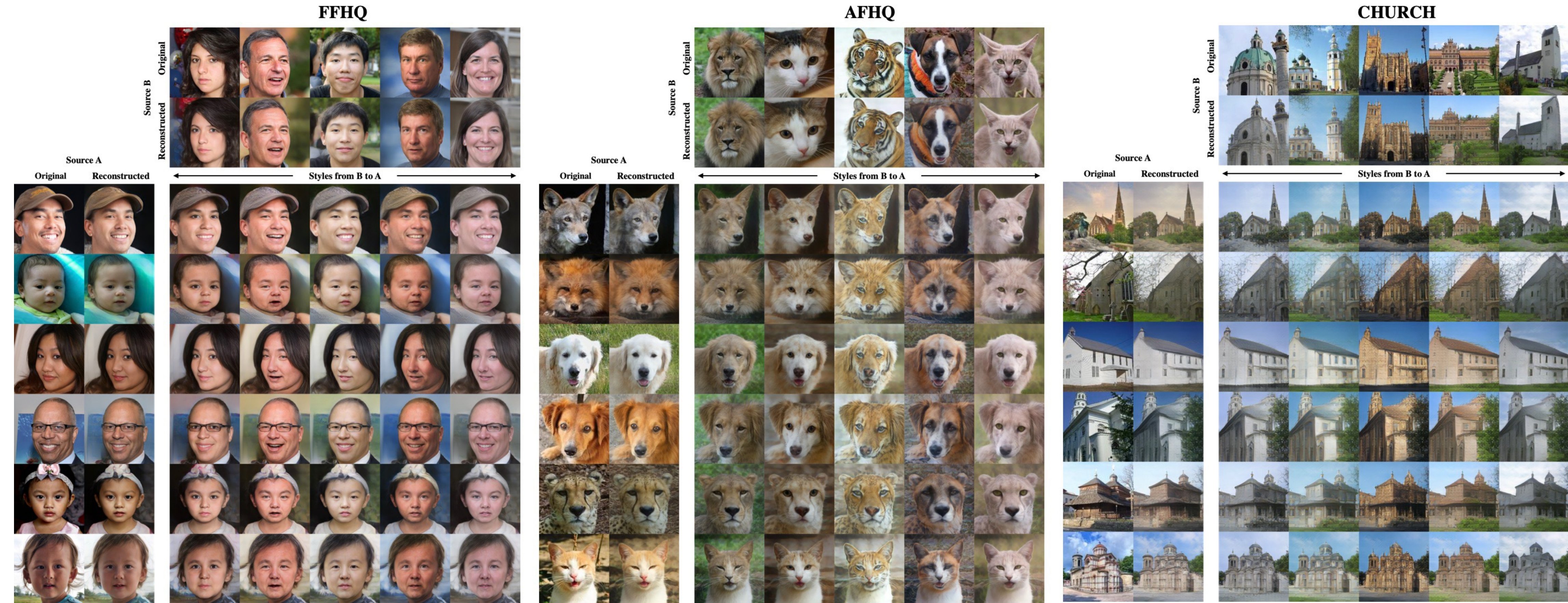


Joint AE-StyleGAN



Experiments

- FFHQ, AFHQ, CHURCH, METFACES – Reconstructions & Style Transfers



Method

- Loss Functions

$$V_{\text{GAN}}(G \circ F, D) = -\mathbb{E}_{x \sim P_X} \mathcal{A}(-\tilde{D}(x)) - \mathbb{E}_{z \sim P_Z} \mathcal{A}(\tilde{D}(G(F(z)))) \quad (1)$$

Decoupled AE-StyleGAN. One straightforward way to train encoder and generator end-to-end is to simultaneously train an encoder with GAN inversion algorithms along with the generator. Here we choose in-domain inversion. To keep the generator’s generating ability intact, one can decouple GAN training and GAN inversion training by introducing separate discriminator models D_1 and D_2 , and freezing G in inversion step. Specifically, D_1 is involved in $V_{\text{GAN}}(G \circ F, D_1)$ in Equation 1 for GAN training steps, and D_2 is involved in $L_{\text{idinv}}(E, D_2, G)$ in Equation 2. Training of D_2 follows:

$$V_{\text{GAN}}(G \circ F, D_2) = -\mathbb{E}_{x \sim P_X} \mathcal{A}(-\tilde{D}_2(x)) - \mathbb{E}_{x \sim P_X} \mathcal{A}(\tilde{D}_2(G(E(x)))) \quad (4)$$

Results

	StyleGAN		ALAE		AE-StyleGAN ($\mathcal{W}$)		AE-StyleGAN ($\mathcal{W}^+$)	
	FID ↓	LPIPS ↑	FID ↓	LPIPS ↑	FID ↓	LPIPS ↑	FID ↓	LPIPS ↑
FFHQ	7.359	0.432	12.574	0.438	8.176	0.448	7.941	0.451
AFHQ	7.992	0.496	21.557	0.508	15.655	0.522	10.282	0.518
MetFaces	29.318	0.465	41.693	0.462	29.710	0.469	29.041	0.471
LSUN Church	27.780	0.520	29.999	0.552	29.387	0.603	29.358	0.592

	StyleGAN		ALAE		AE-StyleGAN ($\mathcal{W}$)		AE-StyleGAN ($\mathcal{W}^+$)	
	Full	End	Full	End	Full	End	Full	End
FFHQ	173.09	173.68	192.60	193.94	181.03	180.23	166.70	165.85
AFHQ	244.83	240.68	229.54	232.53	247.75	248.75	233.86	231.91
MetFaces	231.40	232.77	235.39	237.63	238.84	238.60	240.01	235.41
LSUN Church	245.22	239.62	298.06	295.01	241.46	231.73	240.01	231.49

Algorithm

Algorithm 1 Decoupled AE-StyleGAN Training

- 1: $\theta_E, \theta_{D_1}, \theta_{D_2}, \theta_F, \theta_G \leftarrow$ Initialize network parameters
- 2: **while** not converged **do**
- 3: $x \leftarrow$ Random mini-batch from dataset
- 4: $z \leftarrow$ Samples from $\mathcal{N}(0, I)$
- 5: Step I. Update D_1, D_2
- 6: $L_D \leftarrow -V_{\text{GAN}}(G \circ F, D_1) - V_{\text{GAN}}(G \circ E, D_2)$ in [1]
- 7: $\theta_{D_1}, \theta_{D_2} \leftarrow \text{ADAM}(\nabla_{\theta_{D_1}, \theta_{D_2}} L_D, \theta_{D_1}, \theta_{D_2})$
- 8: Step II. Update E
- 9: $L_E \leftarrow L_{\text{idinv}}(E, D_2, G)$ in [2] or [6]
- 10: $\theta_E \leftarrow \text{ADAM}(\nabla_{\theta_E} L_E, \theta_E)$
- 11: Step III. Update F, G
- 12: $L_G \leftarrow V_{\text{GAN}}(G \circ F, D_1)$ in [1]
- 13: $\theta_F, \theta_G \leftarrow \text{ADAM}(\nabla_{\theta_F, \theta_G} L_G, \theta_F, \theta_G)$
- 14: **end while**

Algorithm 2 Joint AE-StyleGAN Training

- 1: $\theta_E, \theta_D, \theta_F, \theta_G \leftarrow$ Initialize network parameters
- 2: **while** not converged **do**
- 3: $x \leftarrow$ Random mini-batch from dataset
- 4: $z \leftarrow$ Samples from $\mathcal{N}(0, I)$
- 5: Step I. Update D
- 6: $L_D \leftarrow -V_{\text{AEGAN}}(G \circ F, D)$ in [5]
- 7: $\theta_D \leftarrow \text{ADAM}(\nabla_{\theta_D} L_D, \theta_D)$
- 8: Step II. Update E, G
- 9: $L_E \leftarrow L_{\text{idinv}}(E, D, G)$ in [2] or [6]
- 10: $\theta_E, \theta_G \leftarrow \text{ADAM}(\nabla_{\theta_E, \theta_G} L_E, \theta_E, \theta_G)$
- 11: Step III. Update F, G
- 12: $L_G \leftarrow V_{\text{AEGAN}}(G \circ F, D)$ in [5]
- 13: $\theta_F, \theta_G \leftarrow \text{ADAM}(\nabla_{\theta_F, \theta_G} L_G, \theta_F, \theta_G)$
- 14: **end while**

Summary

In this paper, we proposed AE-StyleGAN, a novel algorithm that jointly trains an encoder with a style-based generator. With empirical analysis, we confirmed that this methodology provides an easy-to-invert encoder for real image editing. Extensive results showed that our model has superior image generation and reconstruction capability than baselines. We have explored the problem of training an end-to-end autoencoder. With improved generation fidelity and reconstruction quality, the proposed AE-StyleGAN model can serve as a building-block for further development and applications.

