

pcaGAN: Improving posterior-sampling cGANs via principal component regularization

Matthew Bendel*, Rizwan Ahmad[‡], and Philip Schniter*

*Dept. ECE, The Ohio State University, {bendel.8,schniter.1}@osu.edu

[‡]Dept. BME, The Ohio State University, ahmad.46@osu.edu



THE OHIO STATE UNIVERSITY

Supported in part by the NIH under grant R01-EB029957

Asilomar Conference on Signals, Systems, and Computers
October 2024

Motivation: Imaging inverse problems

Goal: Recover image $\mathbf{x}$ from measurements $\mathbf{y} = \mathcal{M}(\mathbf{x})$:

- $\mathcal{M}(\cdot)$ masks, distorts, and/or corrupts $\mathbf{x}$ with noise.
- Applications: inpainting, super-resolution, deblurring, computed tomography, magnetic resonance imaging (MRI), etc.
- Solution typically posed as **point estimation**: find the single “best” $\hat{\mathbf{x}}$

Challenges with point-estimation:

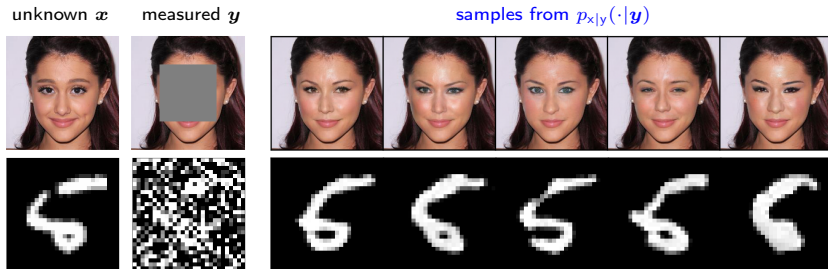
- Doesn't provide uncertainty quantification
- Doesn't navigate the perception-distortion tradeoff¹

Solution: **Sample from posterior distribution** $p_{\mathbf{x}|\mathbf{y}}(\mathbf{x}|\mathbf{y}) = \frac{p_{\mathbf{y}|\mathbf{x}}(\mathbf{y}|\mathbf{x})p_{\mathbf{x}}(\mathbf{x})}{\int p_{\mathbf{y}|\mathbf{x}}(\mathbf{y}|\mathbf{x})p_{\mathbf{x}}(\mathbf{x}) \mathrm{d}\mathbf{x}}$

¹Blau, Michaeli'18

Posterior sampling

The **posterior distribution** $p_{x|y}(\cdot|y)$ represents all knowledge about the image x given the corrupted measurements y



Typical posterior-sampling methods in imaging:

- cVAEs, cNFs, cGANs: Fast but inaccurate?
- Diffusion methods: Slow but accurate?

Our approach

We build on **rcGAN**¹, a type of Wasserstein cGAN:

- Generator G_{θ} : outputs $\hat{x}_i = G_{\theta}(z_i, \mathbf{y})$ for code realization $z_i \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$
- Discriminator D_{ϕ} : aims to distinguish true $(\mathbf{x}, \mathbf{y})$ from fake $(\hat{x}_i, \mathbf{y})$
- Training:

$$\min_{\theta} \max_{\phi} \left\{ \mathbb{E}_{\mathbf{x}, \mathbf{z}, \mathbf{y}} \{ D_{\phi}(\mathbf{x}, \mathbf{y}) - D_{\phi}(G_{\theta}(\mathbf{z}, \mathbf{y}), \mathbf{y}) \} + \mathcal{R}(\theta) - \mathcal{L}_{\text{gp}}(\phi) \right\}$$

- rcGAN's **regularization** $\mathcal{R}(\theta)$ rewards correctness in conditional mean and conditional trace-covariance

Contributions:

- A **new** $\mathcal{R}(\theta)$ that **also** enforces correctness in the K **principal components** of the conditional covariance matrix
- We call our approach **pcaGAN**

¹Bendel, Ahmad, Schniter'22

Background on rcGAN

rcGAN¹ regularizes using an **L1 penalty** and a **standard-deviation reward**:

$$\mathcal{R}_{\text{rc}}(\theta) \triangleq \underbrace{\mathbb{E}_{\mathbf{x}, \mathbf{z}_1, \dots, \mathbf{z}_P, \mathbf{y}} \{ \|\mathbf{x} - \hat{\mathbf{x}}_{(P)}\|_1 \}}_{\triangleq \mathcal{L}_{1,P}(\theta)} - \beta_{\text{std}} \underbrace{\sum_{i=1}^P \mathbb{E}_{\mathbf{z}_1, \dots, \mathbf{z}_P, \mathbf{y}} \{ \|\hat{\mathbf{x}}_i - \hat{\mathbf{x}}_{(P)}\|_1 \}}_{\triangleq \mathcal{L}_{\text{std},P}(\theta)}$$

where $\hat{\mathbf{x}}_{(P)} \triangleq \frac{1}{P} \sum_{i=1}^P \hat{\mathbf{x}}_i$ is the average of P posterior samples.

Key points:

- β_{std} controls diversity, and is optimized during training
- Can prove¹ $\mathcal{R}_{\text{rc}}(\theta)$ yields $\mathbb{E}\{\hat{\mathbf{x}}_i | \mathbf{y}\} = \mathbb{E}\{\mathbf{x} | \mathbf{y}\}$ and $\text{tr Cov}\{\hat{\mathbf{x}}_i | \mathbf{y}\} = \text{tr Cov}\{\mathbf{x} | \mathbf{y}\}$ with Gaussian $p(\mathbf{x} | \mathbf{y})$
- Can prove¹ $\mathcal{R}(\theta) = \mathcal{L}_{2,P}(\theta) - \beta_{\text{var}} \mathcal{L}_{\text{var},P}(\theta)$ does not, for any β_{var}

¹Bendel,Ahmad,Schniter'22

The proposed pcaGAN

Goal: Ensure that $\hat{\mathbf{v}}_k = \mathbf{v}_k$ and $\hat{\lambda}_k = \lambda_k$ for $k = 1, \dots, K$

- $\{(\hat{\mathbf{v}}_k, \hat{\lambda}_k)\}_{k=1}^K$ are the principal evects/evals of $\text{Cov}\{\hat{\mathbf{x}}_i|\mathbf{y}\}$
- $\{(\mathbf{v}_k, \lambda_k)\}_{k=1}^K$ are the principal evects/evals of $\text{Cov}\{\mathbf{x}|\mathbf{y}\}$
- K is user-specified

We propose

$$\mathcal{R}_{\text{pca}}(\boldsymbol{\theta}) \triangleq \mathcal{R}_{\text{rc}}(\boldsymbol{\theta}) + \beta_{\text{pca}} \mathcal{L}_{\text{evect}}(\boldsymbol{\theta}) + \beta_{\text{pca}} \mathcal{L}_{\text{eval}}(\boldsymbol{\theta}),$$

where

$$\mathcal{L}_{\text{evect}}(\boldsymbol{\theta}) \triangleq -\mathbb{E}_{\mathbf{y}} \left\{ \mathbb{E}_{\mathbf{x}, \mathbf{z}_1, \dots, \mathbf{z}_P | \mathbf{y}} \left\{ \sum_{k=1}^K [\hat{\mathbf{v}}_k(\boldsymbol{\theta})^\top (\mathbf{x} - \boldsymbol{\mu}_{\mathbf{x}|\mathbf{y}})]^2 \middle| \mathbf{y} \right\} \right\}$$

$$\mathcal{L}_{\text{eval}}(\boldsymbol{\theta}) \triangleq \mathbb{E}_{\mathbf{y}} \left\{ \mathbb{E}_{\mathbf{x}, \mathbf{z}_1, \dots, \mathbf{z}_P | \mathbf{y}} \left\{ \sum_{k=1}^K (1 - \lambda_k / \hat{\lambda}_k(\boldsymbol{\theta}))^2 \middle| \mathbf{y} \right\} \right\}$$

and

we approximate the unknown $\{(\mathbf{v}_k, \lambda_k)\}_{k=1}^K$ and $\boldsymbol{\mu}_{\mathbf{x}|\mathbf{y}} \triangleq \mathbb{E}\{\mathbf{x}|\mathbf{y}\}$

Understanding pcaGAN

Eigenvector regularization:

$$\mathcal{L}_{\text{evect}}(\theta) = -\mathbb{E}_y \left\{ \mathbb{E}_{\mathbf{x}, \mathbf{z}_1, \dots, \mathbf{z}_P | y} \left\{ \sum_{k=1}^K [\hat{\mathbf{v}}_k^T (\mathbf{x} - \boldsymbol{\mu}_{\mathbf{x}|y})]^2 \middle| \mathbf{y} \right\} \right\}$$

- If $\boldsymbol{\mu}_{\mathbf{x}|y} = \mathbb{E}\{\mathbf{x}|\mathbf{y}\}$ was known, minimizing over θ would force $\{\hat{\mathbf{v}}_k = \mathbf{v}_k\}_{k=1}^K$
- We set $\boldsymbol{\mu}_{\mathbf{x}|y} \approx \widehat{\boldsymbol{\mu}}_{\mathbf{x}|y} = \text{StopGrad}(\widehat{\mathbf{x}}_{(P_{\text{pca}})})$ after $\widehat{\mathbf{x}}_{(P_{\text{pca}})}$ stabilizes
- We compute $\{\hat{\mathbf{v}}_k\}_{k=1}^K$ using an SVD of $\{\widehat{\mathbf{x}}_i\}_{i=1}^{P_{\text{pca}}}$, where $P_{\text{pca}} \approx 10K$

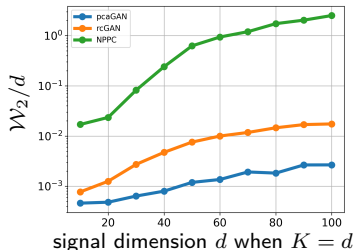
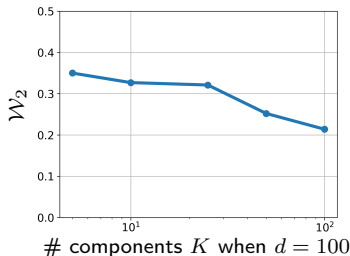
Eigenvalue regularization:

$$\mathcal{L}_{\text{eval}}(\theta) = \mathbb{E}_y \left\{ \mathbb{E}_{\mathbf{x}, \mathbf{z}_1, \dots, \mathbf{z}_P | y} \left\{ \sum_{k=1}^K (1 - \lambda_k / \widehat{\lambda}_k)^2 \middle| \mathbf{y} \right\} \right\}$$

- If λ_k was known, minimizing over θ would force $\{\widehat{\lambda}_k = \lambda_k\}_{k=1}^K$
- We set $\lambda_k \approx \text{StopGrad}(\frac{1}{P_{\text{pca}}+1} \|\widehat{\mathbf{v}}_k^T [\mathbf{x} - \widehat{\boldsymbol{\mu}}_{\mathbf{x}|y}, \widehat{\mathbf{x}}_1 - \widehat{\boldsymbol{\mu}}_{\mathbf{x}|y}, \dots, \widehat{\mathbf{x}}_{P_{\text{pca}}} - \widehat{\boldsymbol{\mu}}_{\mathbf{x}|y}]\|_2^2)$ after $\{\widehat{\mathbf{v}}_k\}$ stabilize, which is correct in expectation when $\widehat{\mathbf{v}}_k$ and $\widehat{\boldsymbol{\mu}}_{\mathbf{x}|y}$ are correct
- We compute $\{\widehat{\lambda}_k\}_{k=1}^K$ using the same SVD as above

Gaussian toy experiments

- Here we recover $x \sim \mathcal{N}(\mu_x, \Sigma_x)$ from $y = Ax + w$ with inpainting $A \in \mathbb{R}^{\frac{d}{2} \times d}$ and $w \sim \mathcal{N}(0, \sigma^2 I)$. Both μ_x and Σ_x are random
- Performance measured via Wasserstein-2 distance $\mathcal{W}_2(p_{x|y}, p_{\hat{x}|y})$



- $\mathcal{W}_2$ decreases with K , and pcaGAN beats both rcGAN and NPPC¹ in $\mathcal{W}_2$
 - NPPC trains a neural net to directly estimate $\{(\lambda_k, v_k)\}_{k=1}^K$ from y

¹Nehme, Yair, Michaeli'23

MNIST denoising experiments

We recovered MNIST digits x from measurements $y = x + w$ with $w \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$

- We measured by rMSE, residual error magnitude (REM) at $K = 5$:

$$\text{REM}_5 \triangleq \mathbb{E}_{x,y} \left\{ \left\| (\mathbf{I} - \hat{\mathbf{V}}_5 \hat{\mathbf{V}}_5^T)(x - \widehat{\mu}_{x|y}) \right\|_2 \right\} \text{ where } \hat{\mathbf{V}}_5 \triangleq [\hat{v}_1, \dots, \hat{v}_5]$$

- We also measured Conditional FID (CFID)¹, which is similar to FID but measures the discrepancy between $p_{\hat{x}|y}$ and $p_{x|y}$ (not between $p_{\hat{x}}$ and p_x)

- Results (128 test images):

Model	rMSE ↓	REM ₅ ↓	CFID ↓	Time (128 samples) ↓
NPPC ($K = 5$)	3.94	3.63	–	112 ms
rcGAN	4.04	<u>3.41</u>	63.44	<u>118</u> ms
pcaGAN (ours, $K = 5$)	<u>4.02</u>	3.31	61.48	<u>118</u> ms

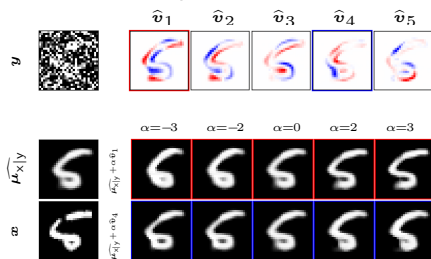
- **pcaGAN wins in both REM and CFID** (note NPPC is not generative)

¹Soloveitchik'21

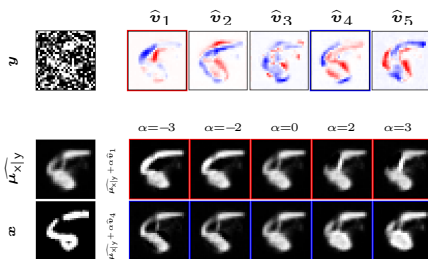
MNIST: visualizing the principal uncertainty components

- Principal eigenvectors $\{\mathbf{v}_k\}_{k=1}^K$ are shown below for $K = 5$
- Also shown are $\widehat{\mu_{\mathbf{x}|y}} \pm \alpha \widehat{\mathbf{v}}_k$ for $k \in \{1, 4\}$ and $\alpha \in \{-3, -2, 0, 2, 3\}$

pcaGAN:



NPPC:



- pcaGAN's eigenvectors show much more meaningful structure

Large-scale image completion/inpainting

We **inpainted large random masks** on 256x256 FFHQ face images

■ Results (20k test images):

Model	CFID↓	FID↓	LPIPS↓	Time (40 samples)↓
DPS ¹ (1000 NFEs)	<u>7.26</u>	<u>2.00</u>	<u>0.1245</u>	14 min
DDNM ² (100 NFEs)	11.30	3.63	0.1409	30 s
DDRM ³ (20 NFEs)	13.17	5.36	0.1587	5 s
pscGAN ⁴	18.44	8.40	0.1716	325 ms
CoModGAN ⁵	7.85	2.23	0.1290	325 ms
rcGAN	7.51	2.12	0.1262	325 ms
pcaGAN (ours, $K = 2$)	7.08	1.98	0.1230	325 ms

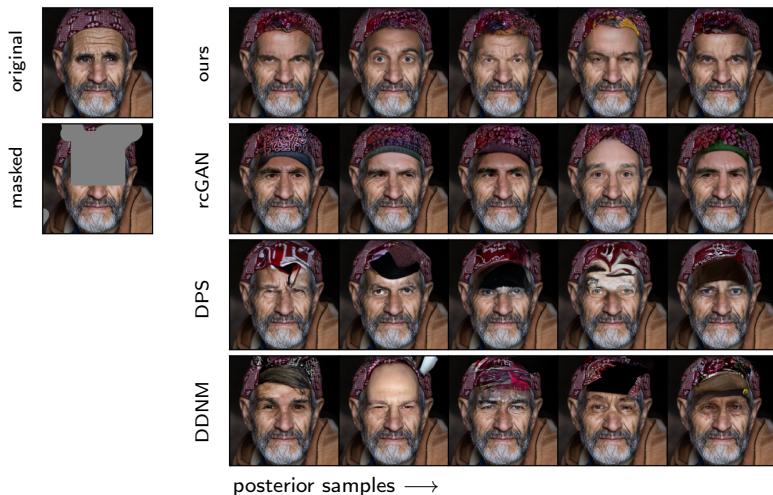
■ pcaGAN **outperformed all diffusion and cGAN competitors!**

■ cGANs are 15x to 2500x faster than the diffusion methods

¹Chung et al'23, ²Wang et al'23, ³Kawar et al'22, ⁴Ohayon et al'21, ⁵Zhao et al'21

Example FFHQ inpainting

pcaGAN generates samples that are both **high quality** and **diverse**



Results on accelerated MR image recovery

We reconstructed multicoil **fastMRI**¹ T2 brain images at acceleration $R = 8$

■ Results (74 test images):

Model	CFID↓	FID↓	PSNR↑	SSIM↑	LPIPS↓	DISTS↓	Time (4 samples)↓
E2E-VarNet ²	36.86	44.04	36.49	0.9220	0.0575	0.1253	316ms
Langevin (Jalal ³)	48.59	52.62	33.90	0.9137	0.0579	0.1086	14 min
cGAN (Adler ⁴)	59.94	31.81	33.51	0.9111	0.0614	0.1252	217 ms
pscGAN	39.67	43.39	34.92	0.9222	0.0532	0.1128	217 ms
rcGAN	<u>24.04</u>	<u>28.43</u>	35.42	<u>0.9257</u>	<u>0.0379</u>	<u>0.0877</u>	217 ms
pcaGAN (ours, $K = 1$)	21.65	28.35	<u>35.94</u>	0.9283	0.0344	0.0799	217 ms

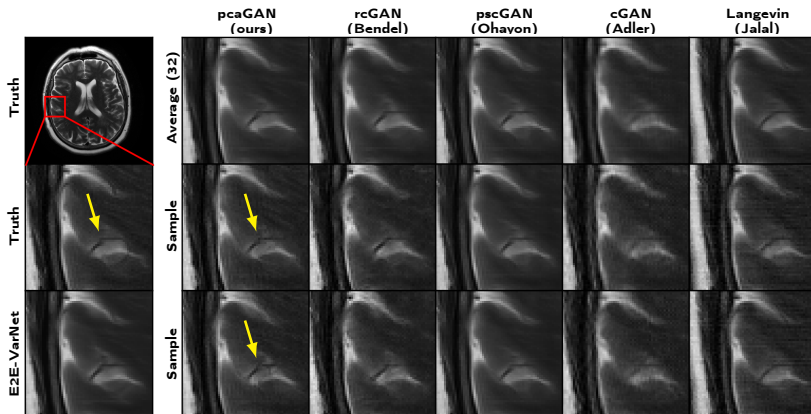
■ pcaGAN won in all metrics but PSNR!

■ The cGANs generated samples 3800x faster than the Langevin method

■ Note: PSNR, SSIM, LPIPS, DISTS computed using $\hat{x}_{(P)}$ for the optimal P

¹Zbontar et al'18, ²Sriram et al'19, ³Jalal et al'21, ⁴Adler,Öktem'18

Example MRI recoveries



- posterior samples from pcaGAN show meaningful variations (see arrows)
- posterior samples from pscGAN show no variation
- posterior samples from cGAN (Adler) and Langevin (Jalal) show unwanted artifacts

Summary

Goal:

- Fast and accurate posterior sampling for imaging inverse problems

Contribution:

- A cGAN with a regularization that encourages correctness in the y -conditional mean, trace-covariance, and K principal covariance components

Experimental results

- Considered MNIST denoising, FFHQ inpainting, and multicoil MRI
- Outperforms existing cGANs and diffusion methods (DPS, DDRM, DDNM, Langevin) in PSNR, SSIM, LPIPS, DISTS, FID, and CFID
- Runs 10x to 4000x faster than diffusion methods
- Outperforms NNPC (direct estimation of principal covariance components)