

Deep Equilibrium models for Poisson inverse problems via Mirror Descent

Christian Daniele, Silvia Villa, Samuel Vaiter and Luca Calatroni

SIAM Chapters Meeting for applied mathematicians
January 30, 2026

Poisson Inverse Problems

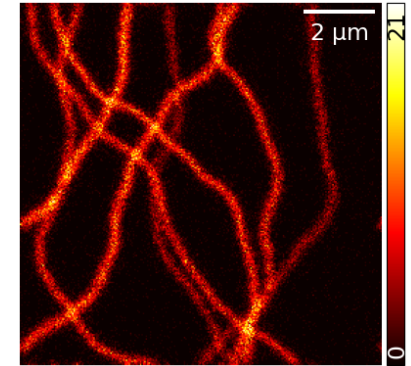
$$y = \text{Poisson}(A\mathbf{x})$$

$$A \in \mathbb{R}_{\geq 0}^{m \times n}$$

$$\mathbf{x} \in \mathbb{R}_{\geq 0}^n, \mathbf{y} \in \mathbb{R}_{\geq 0}^m$$

Motivation: Poisson noise models **Photon counts** in imaging sensors (e.g. Microscopy) $\longrightarrow$

[Bertero and Boccacci '98]



To retrieve $\mathbf{x}$:

$$\underset{\mathbf{x} \geq 0}{\operatorname{argmax}} \pi_{\mathcal{X}|\mathcal{Y}}(\mathbf{x} | \mathbf{y})$$

MAP

$\implies$

$$\underset{\mathbf{x} \geq 0}{\operatorname{argmin}} \text{KL}(\mathbf{y}, A\mathbf{x}) + R_{\theta}(\mathbf{x})$$

$$\pi_{\mathcal{X}|\mathcal{Y}}(\mathbf{x} | \mathbf{y}) \propto \pi_{\mathcal{Y}|\mathcal{X}}(\mathbf{y} | A\mathbf{x}) \pi_{\mathcal{X}}(\mathbf{x})$$



not available

$$\text{KL}(\mathbf{y}, A\mathbf{x}) = \sum_{i=1}^m d_{\text{KL}}(y_i, (A\mathbf{x})_i)$$

$$d_{\text{KL}}(y, x) = y \log \left(\frac{y}{x} \right) + x - y$$

- ≥ 0
- convex
- not Lipschitz gradient

Variational formulation and regularization

$$\operatorname{argmin}_{x \geq 0} \text{KL}(y, Ax) + R_\theta(x)$$

Model-based

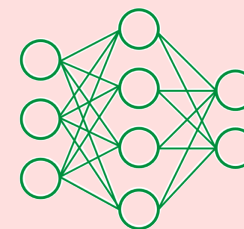
$$R_\theta(x) = \begin{cases} \theta \|x\|_1 & \text{L1} \\ \theta \text{TV}(x) & \text{Total Variation} \end{cases}$$

[Rudin, Osher and Fatemi '92]

$\theta > 0$, regularization parameter

Goal:

given $\operatorname{argmin}_{x \geq 0} \text{KL}(y, Ax) + R_\theta(x)$, learning $R_\theta =$

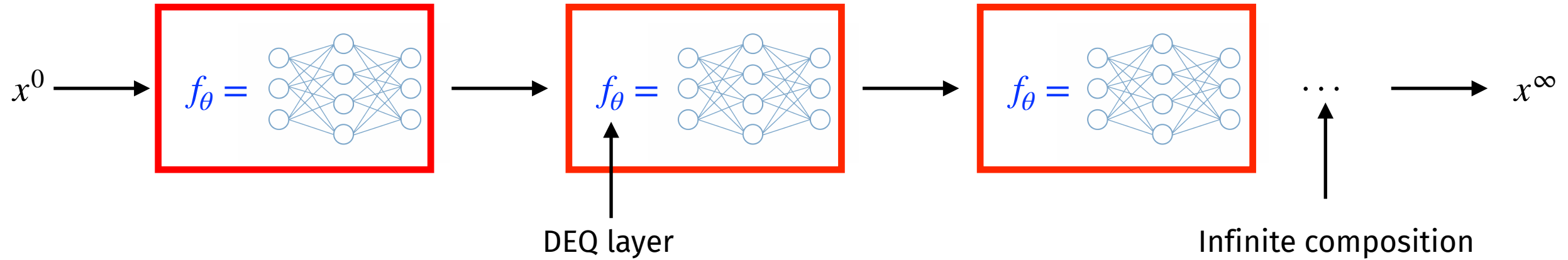


optimization

$$\text{Algo}_\theta(x^0, y, A) \longrightarrow f_\theta(\cdot; y, A)$$

Challenge: R_θ cannot be learned in an end-to-end fashion

Deep Equilibrium models (DEQs) [Bai, Kolter and Koltun '19, Gilton, Ongie and Willett '21]



Definition (Well-posedness of a DEQ)

The DEQ associated with $f_\theta: X \subset \mathbb{R}^n \rightarrow X$ is **well-posed** if:

1. $\text{Fix}(f_\theta) \neq \emptyset$;
2. For all $x^0 \in X$, $\{x^k\}_{k \in \mathbb{N}}$ given by:

$$x^k = f_\theta(x^{k-1})$$

converges to $x^\infty \in \text{Fix}(f_\theta)$.

Example: Contraction $\implies$ Well-posedness

DEQs for image reconstruction

Given:

$$y = \text{Poisson}(Ax)$$

Design $f_\theta(\cdot; y)$ s.t. $x^\infty(\theta; y) \approx x$

with $x^\infty(\theta; y) = f_\theta(x^\infty(\theta); y)$

For simplicity:

$$f_\theta(\cdot; y) = f_\theta(\cdot; y, A) \\ x^\infty(\theta; y) = x^\infty(\theta; y, x^0)$$

Ingredients:

$\{(\hat{x}_i, y_i)\}_{i=1}^N$, $y_i = \text{Poisson}(A\hat{x}_i)$, fix the structure of $f_\theta(\cdot; y_i)$

Training of a DEQ

$$\text{find } \hat{\theta} \in \underset{\theta \in \Theta}{\text{argmin}} \frac{1}{2N} \sum_{i=1}^N \|x^\infty(\theta; y_i) - \hat{x}_i\|_2^2$$

$$\text{s.t. } x^\infty(\theta; y_i) = f_\theta(x^\infty(\theta); y_i)$$

Question:

How do we choose f_θ ? $\longrightarrow$ Optimisation algorithm. Which one?

KL minimisation: Mirror Descent

$$\operatorname{argmin}_{x \geq 0} \operatorname{KL}(y, Ax) + R_\theta(x) := J_\theta(x; y)$$

Given a Legendre function h :

$$x^{k+1} = \nabla h^*(\nabla h(x^k) - \tau(\nabla \operatorname{KL}(y, Ax) + \nabla R_\theta)) =: f_\theta(x^k; y)$$

If $h = \|\cdot\|_2^2$, then

$$x^{k+1} = x^k - \tau(\nabla \operatorname{KL}(y, Ax^k) + \nabla R_\theta(x^k)) \quad \text{Gradient Descent}$$

Burg's entropy: $h(x) = -\sum_{i=1}^n \log(x_i), \quad \operatorname{dom}(h) = (0, +\infty)^n$

[Bauschke, Bolte and Teboulle '17]

Key property:

$$x^0 \in \operatorname{int}(\operatorname{dom}(h)) \implies x^k \in \operatorname{int}(\operatorname{dom}(h)) \forall k \geq 0 \longrightarrow \text{no need of projection}$$

$\operatorname{KL}(y, A\cdot)$ has a Lipschitz gradient “relative to h ”

From Mirror Descent to DEQ-MD

$$f_{\theta}(x; y) := \nabla h^*(\nabla h(x) - \tau(\nabla \text{KL}(y, Ax) + \nabla R_{\theta}(x)))$$

Proposition (Informal)

Let $R_{\theta} : \mathbb{R}^n \rightarrow \mathbb{R}$ be an **analytic network** and $h(x) = -\sum_{i=1}^n \log(x_i)$. Then, for all $x^0 \in (0, +\infty)^n$, the sequence given by $\{x^k\}_{k \in \mathbb{N}}$ **converges to** $x^{\infty}(\theta; y, x^0)$ **a critical point of** $J_{\theta}(\cdot; y) = \text{KL}(y, A\cdot) + R_{\theta}$.
In addition it holds: $\text{Fix}(f_{\theta}(\cdot; y)) \subset \text{Crit}(J_{\theta})$.

Key point: Learning $f_{\theta}(\cdot; y) \implies$ learning the regularization R_{θ}

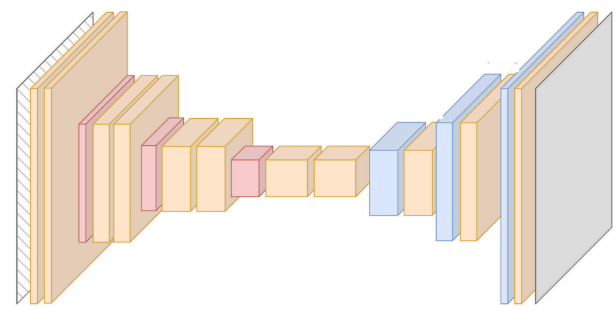
Remark: Analyticity is not restrictive!

Numerical implementation

Choice of R_θ :

$$R_\theta(x) = \frac{1}{2} \|x - N_\theta(x)\|^2, N_\theta =$$

[Romano, Elad, and Milanfar '17]
[Hurault, Leclaire and Papadakis '21]

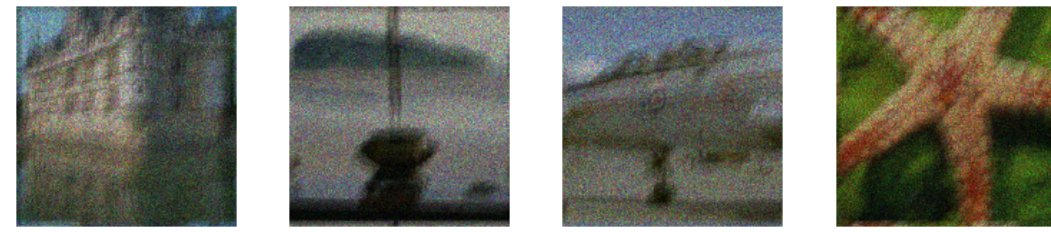


Dataset:
500 natural images of size 256 x 256

$\hat{x}_i$



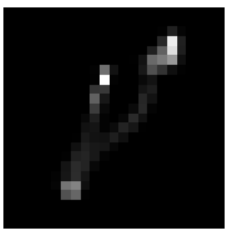
y_i



Data generation:

$$y = \text{Poiss}(\alpha Ax)$$

α = strength of Poisson noise
 $\alpha \in [20,80]$

$A =$ 

Convolution operator

Numerical tests:

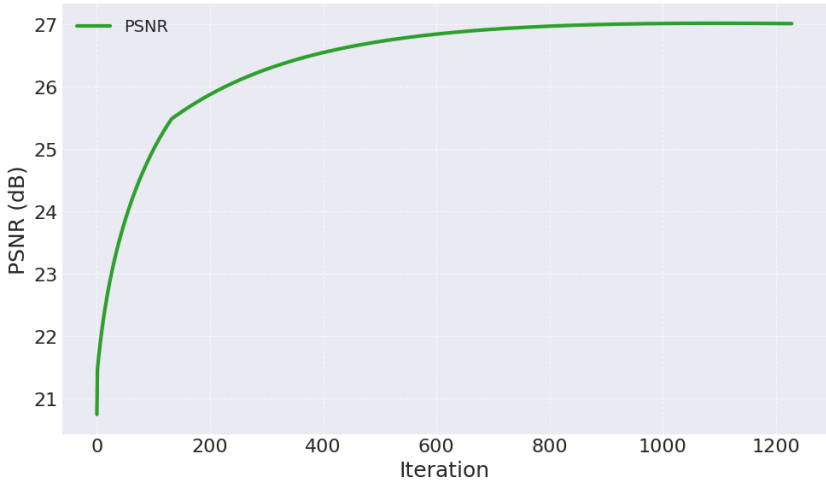
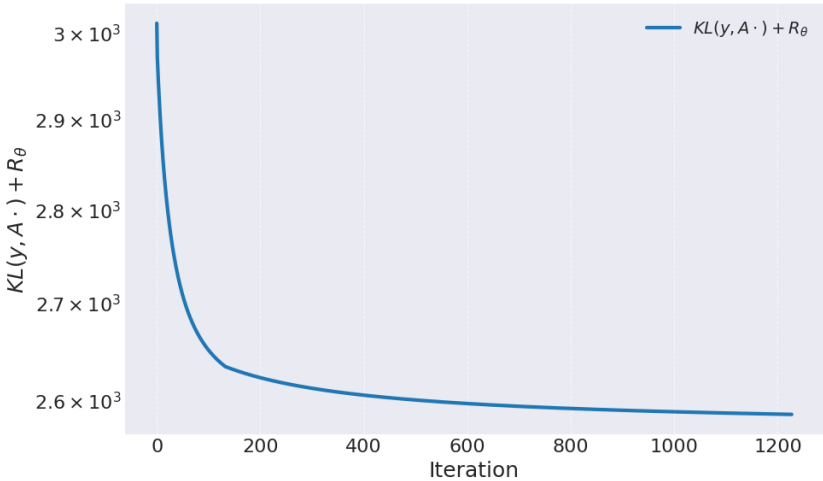
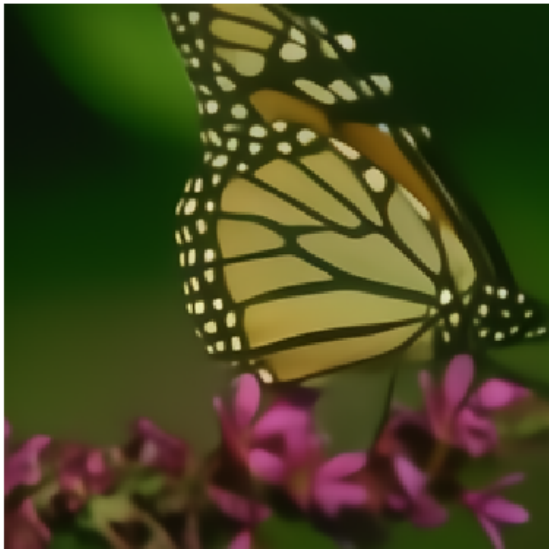
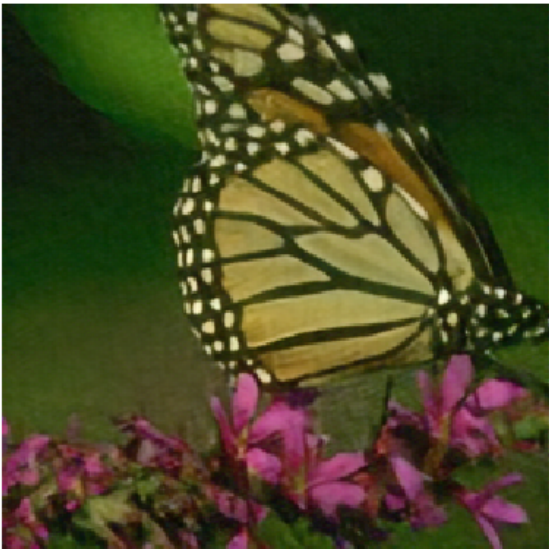
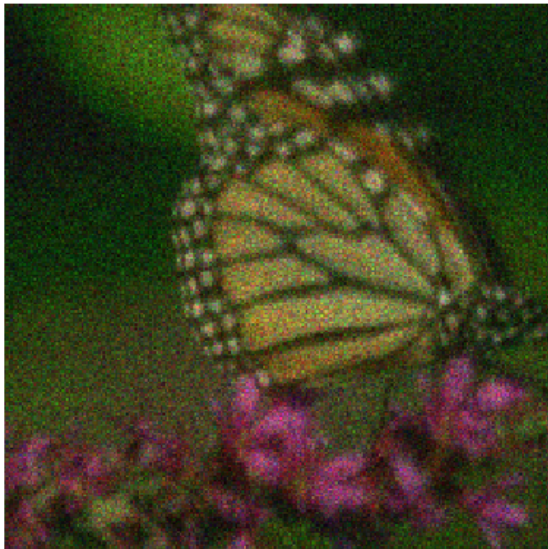
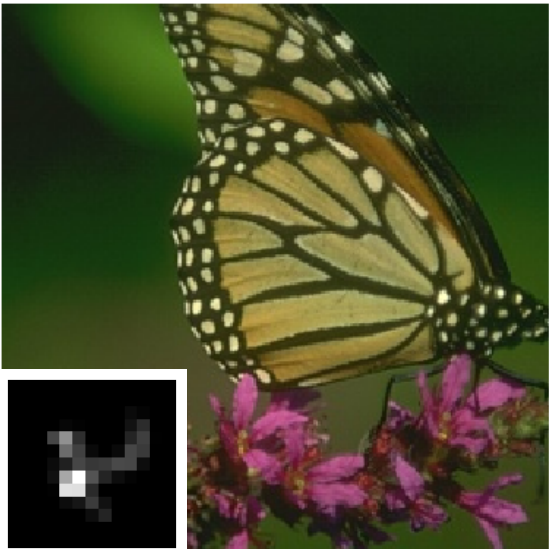
$\alpha = 40$

GT

Measurement

DEQ, PSNR: **28.23**, SSIM: **0.83**

RAM, PSNR: 27.66, SSIM: 0.82



Numerical tests:

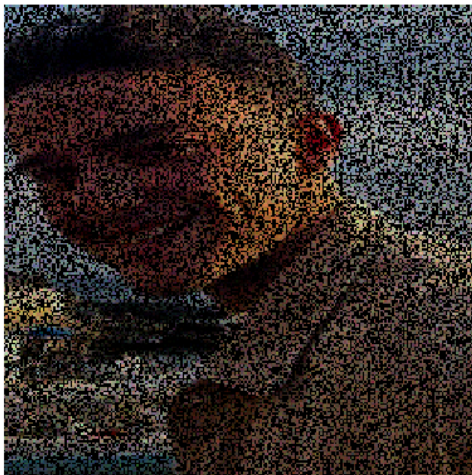
$\alpha = 60$

President

GT



Measurement



DEQ, PSNR: **29.19**, SSIM: **0.83**



RAM, PSNR: 28.21, SSIM: 0.80



Y-SIMAI
representative

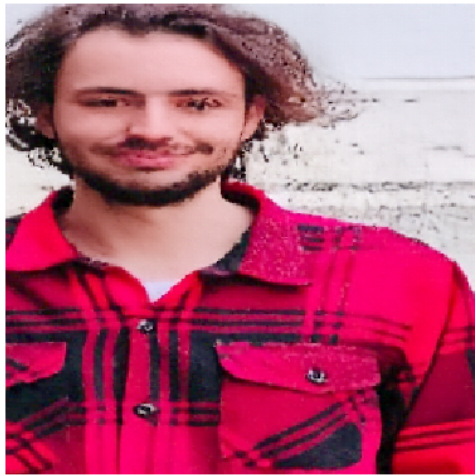
GT



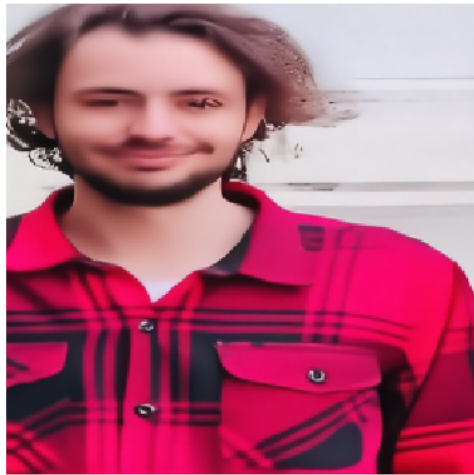
Measurement



DEQ, PSNR: **25.39**, SSIM: **0.77**



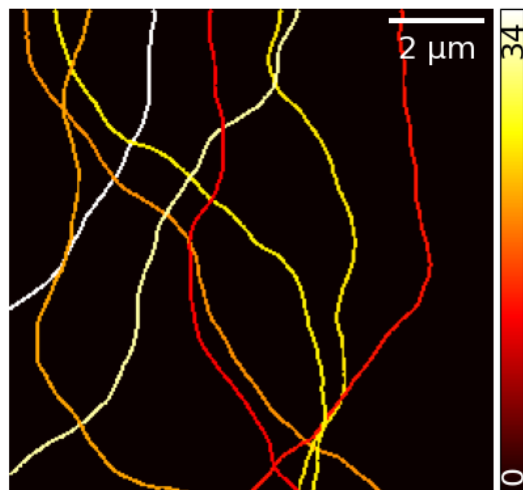
RAM, PSNR: 25.10, SSIM: 0.71



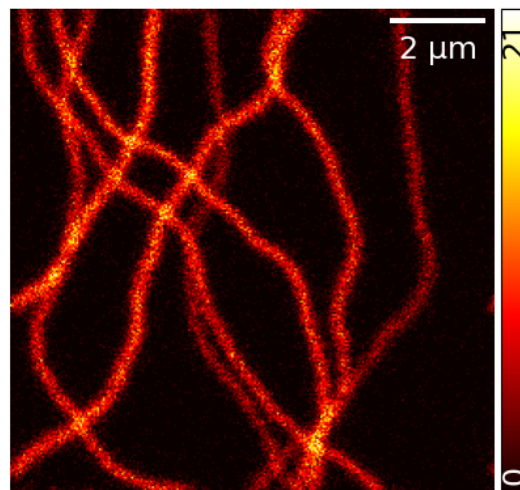
Work in Progress: Testing on microscopy

Supervised training on simulated data

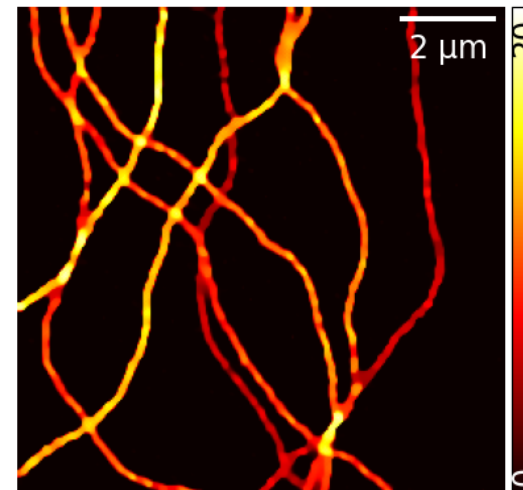
GT



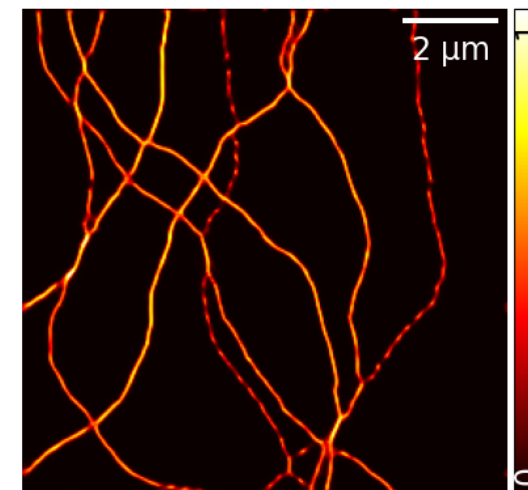
Measurement



TV, PSNR: 17.83, SSIM: 0.78



DEQ, PSNR: **21.05**, SSIM: **0.91**



Conclusion & Outlook

- Learned regularisation for Poisson IPs through MD-type deep equilibrium models
- Interpretable inference and reliable results

Future work:

- Extend training of DEQ-MD in the self-supervised setting
- Extend to CT, MRI

Thanks for your attention!

UniGe



Preprint on Arxiv:



Github:



Weights the action of the regularizer

$$f_{\theta_1, \theta_2}(x; y) = \nabla h^*(\nabla h(x) - \tau(\nabla \text{KL}(y, Ax) + g_{\theta_2}(y) \nabla R_{\theta_1}(x)))$$

More precisely:

Burg's entropy: $h(x) = - \sum_{i=1}^n \log(x_i), \quad \text{dom}(h) = (0, +\infty)^n$

[Bauschke, Bolte and Teboulle '17]

Key points:

$$x^0 \in \text{int}(\text{dom}(h)) \implies x^k \in \text{int}(\text{dom}(h)) \forall k \geq 0$$

$$\tau \propto \frac{1}{L}, L > 0 \text{ is s.t. } Lh - (\text{KL} + R_\theta) \text{ is convex on } \text{int}(\text{dom}(h))$$

NoLip property

[Bauschke, Bolte and Teboulle '17]

L-SMAD

[Bolte, Sabach, Teboulle and Vaisbourd '18]

Relative smoothness

[Lu, Freund and Nesterov '18]

Gradient computation:

$$\{(x_i, y_i)\}_{i=1}^N, y_i = \text{Poisson}(Ax_i)$$

$$\text{find } \hat{\theta} \in \operatorname{argmin}_{\theta \in \Theta} \frac{1}{N} \sum_{i=1}^N \|x^\infty(\theta; y_i, x_i^0), x_i\|_2^2$$

$$\bar{\theta} \in \Theta, \nabla_{\theta} \|x^\infty(\bar{\theta}; y, x^0), x\|_2^2 = \frac{\partial x^\infty(\bar{\theta}; y, x^0)}{\partial \theta}^\top (x^\infty(\bar{\theta}; y, x^0) - x) \quad \frac{\partial x^\infty(\bar{\theta}; y, x^0)}{\partial \theta} = \frac{f_{\bar{\theta}}(x^\infty(\bar{\theta}; y, x^0))}{\partial \theta}$$

$$\frac{\partial x^\infty(\bar{\theta}; y, x^0)}{\partial \theta} = \left(I - \frac{\partial f_{\bar{\theta}}(x)}{\partial x} \bigg|_{x=x^\infty} \right)^{-1} \frac{\partial f_{\bar{\theta}}(x^\infty(\bar{\theta}; y, x^0))}{\partial \theta}$$

Constant memory

Gradient computation:

$$\{(x_i, y_i)\}_{i=1}^N, y_i = \text{Poisson}(Ax_i)$$

$$\text{find } \hat{\theta} \in \operatorname{argmin}_{\theta \in \Theta} \frac{1}{N} \sum_{i=1}^N \|x^\infty(\theta; y_i, x_i^0), x_i\|_2^2$$

$$\bar{\theta} \in \Theta, \nabla_{\theta} \|x^\infty(\bar{\theta}; y, x^0), x\|_2^2 = \frac{\partial x^\infty(\bar{\theta}; y, x^0)^\top}{\partial \theta} (x^\infty(\bar{\theta}; y, x^0) - x) \quad \frac{\partial x^\infty(\bar{\theta}; y, x^0)}{\partial \theta} = \frac{f_{\bar{\theta}}(x^\infty(\bar{\theta}; y, x^0))}{\partial \theta}$$

$$\frac{\partial x^\infty(\bar{\theta}; y, x^0)}{\partial \theta} \approx \left(I - \frac{\partial f_{\bar{\theta}}(x^\infty(\bar{\theta}; y, x^0))}{\partial x} \right)^{-1} \frac{\partial f_{\bar{\theta}}(x^\infty(\bar{\theta}; y, x^0))}{\partial \theta}$$

Jacobian Free Backpropagation/
One step differentiation

[Fung, Heaton, Li, McKenzie, Osher and Yin '22]
[Bolte, Pauwels and Vaiter '23]

$$\left\| \frac{\partial f_{\bar{\theta}}}{\partial x} \right\| < 1$$

Our case: $\left\| \frac{\partial f_{\bar{\theta}}}{\partial x} \right\| \leq 1$

Forward pass

Challenges:

Fixed point iteration $x^{k+1} = f_{\theta}(x^k; y)$ slow (sublinear rate)

$\tau \propto \frac{1}{L}$ with $L > 0$ s.t. $Lh - (\text{KL} + R_{\theta})$ is convex, hard to compute

Backtracking for DEQ-MD [Hurault, Kamilov, Leclaire and Papadakis '23]

Input: $\gamma \in (0,1), \eta \in (0,1)$

Define: $T_{\tau}(x) = \nabla h^*(\nabla h(x) - \tau(\nabla \text{KL}(y, Ax) + \nabla R_{\theta}(x)))$

while $\Psi_{\theta}(x^k) - \Psi_{\theta}(T_{\tau}(x^k)) < \frac{\gamma}{\tau} D_h(T_{\tau}(x^k), x^k)$ **do**

$\tau \leftarrow \eta \tau$

Update: $x^{k+1} = T_{\tau}(x^k)$

