



(A short intro to) Deep learning for image reconstruction

Nicolas Ducros^{1, 2}

¹CREATIS, Univ Lyon, INSA-Lyon, UCB Lyon 1, CNRS, Inserm, CREATIS UMR 5220,
U1206, Lyon, France

²IUF, Institut Universitaire de France

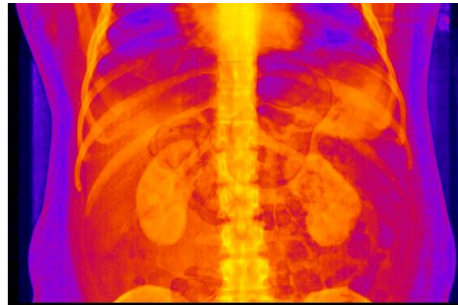
“Natural” vs Reconstructed Images

2

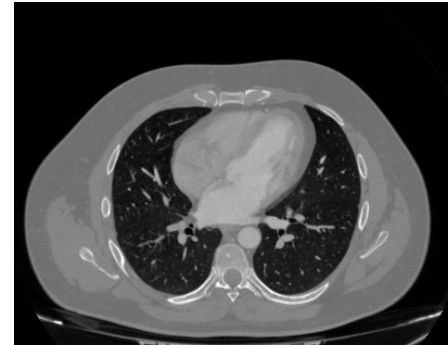
Rec



Nat



Rec



Nat



Nat



Nat



Rec



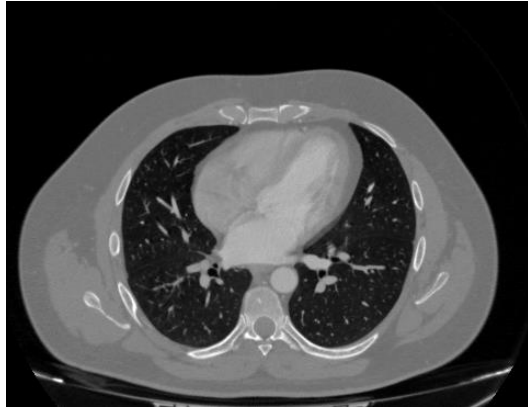
Nat



Reconstruction in Medical Imaging

3

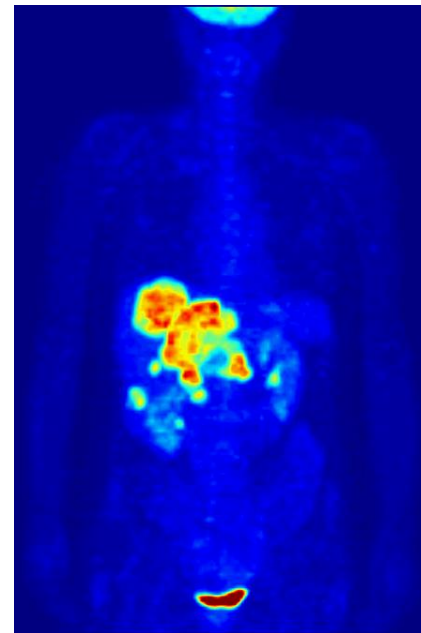
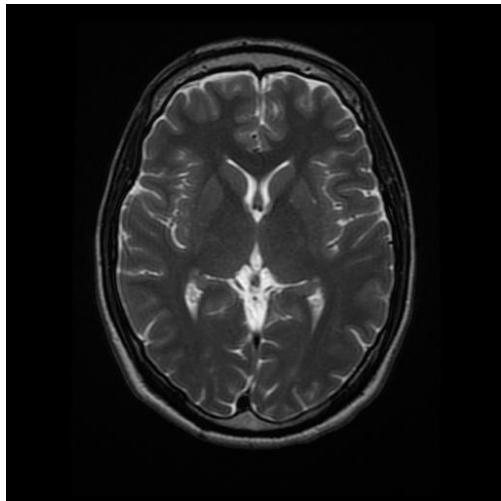
Computerized tomography (CT)



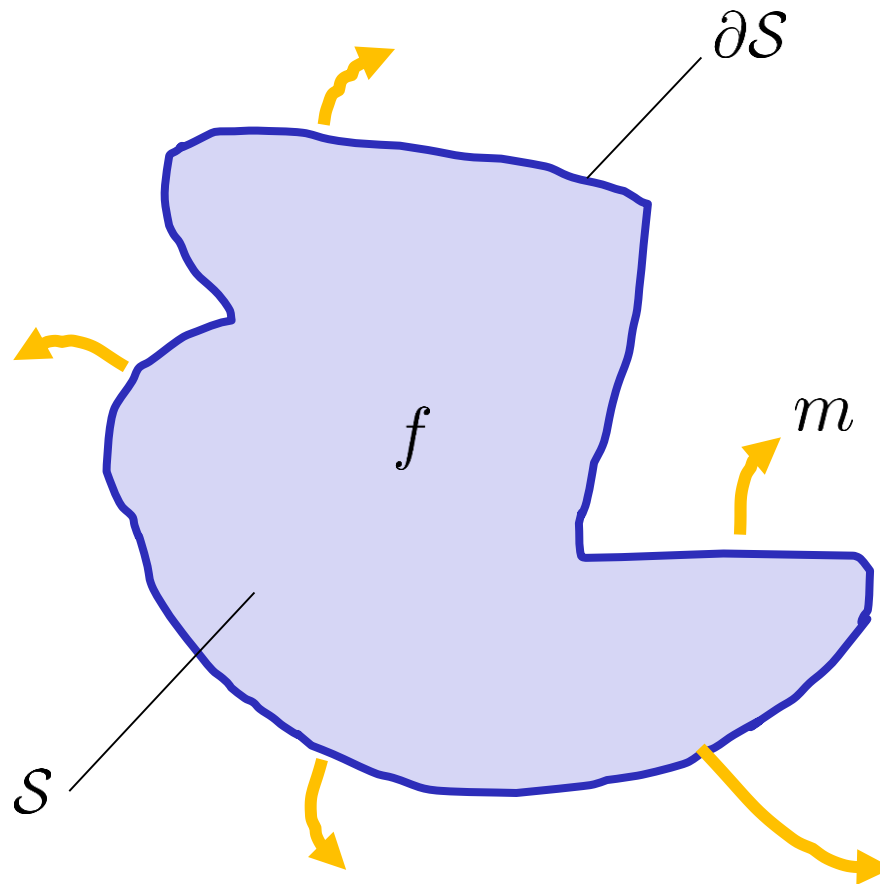
Ultrasound Imaging



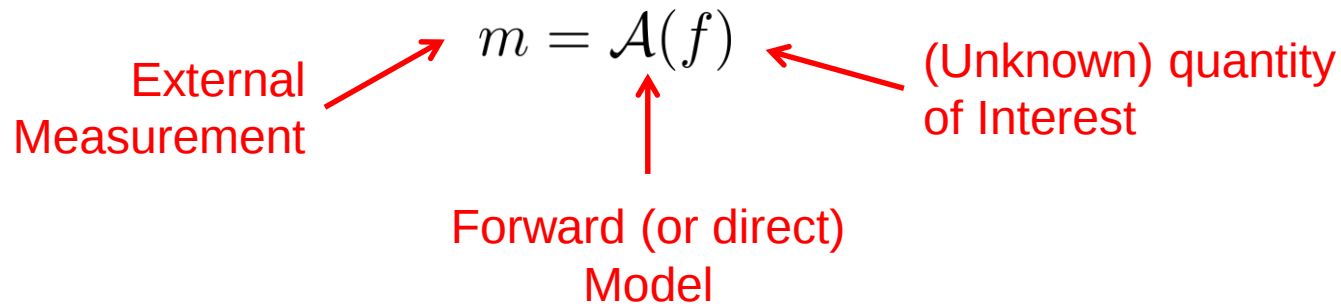
Magnetic Resonance (MRI)



Positron emission tomography (PET)



➤ **Internal unknowns from external measurements**



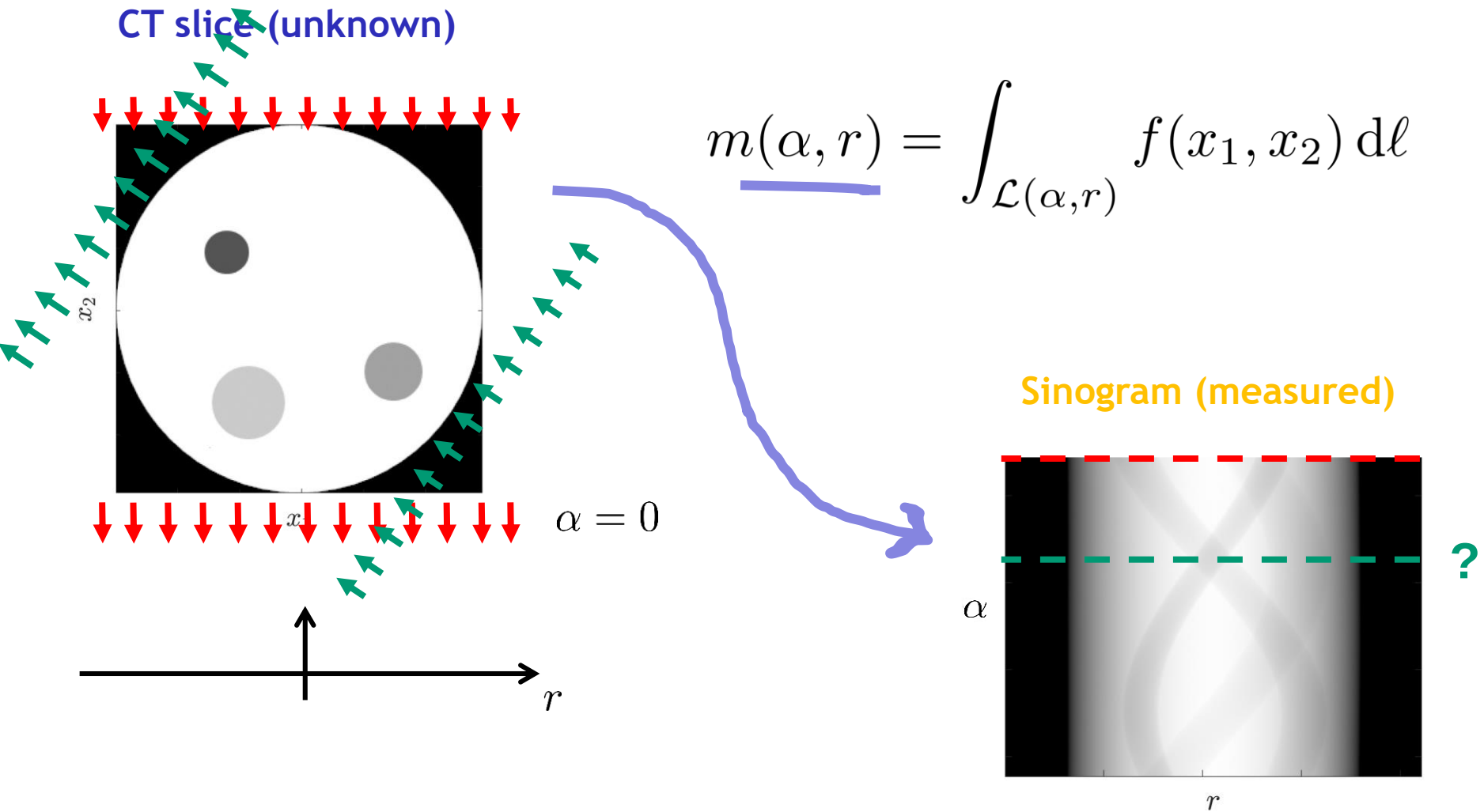
➤ **Most medical imaging problems are**

- ❖ Linear
- ❖ Corrupted by noise

➤ **In a discrete setting**

Diagram illustrating the discrete inverse problem equation $m = A f + \eta$. Red arrows point from the labels to the equation:

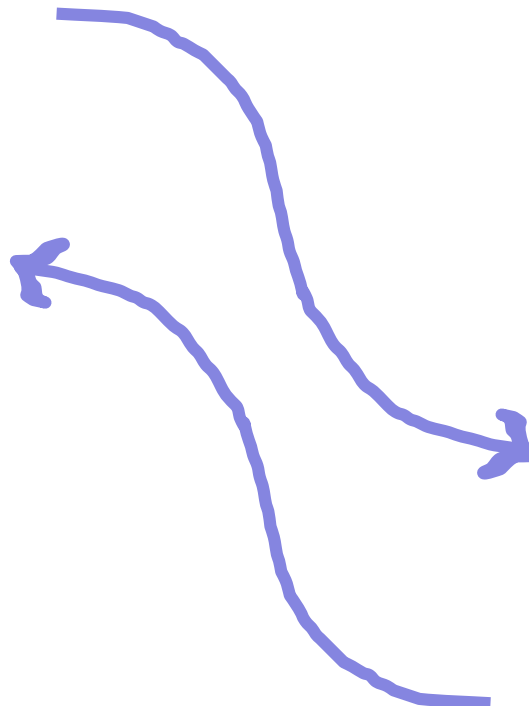
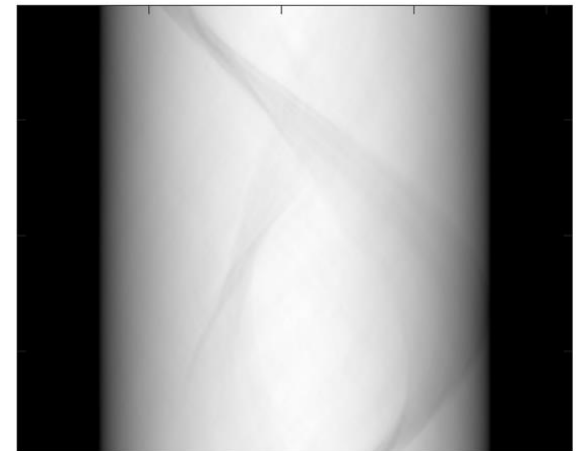
- $m \in \mathbb{R}^M$ points to m .
- $A \in \mathbb{R}^{M \times N}$ points to A .
- $f \in \mathbb{R}^N$ points to f .
- Noise points to η .



CT slice (unknown)



Sinogram (measured)



➤ 1. Inversion “by hand”

- ❖ Model the forward and invert analytically

derive $\mathcal{R}$ such that $f = \mathcal{R}(m)$

➤ 2. Optimization of handcrafted functionals

- ❖ Build cost function from prior knowledge about the solution/measurements
- ❖ Minimize the cost function

find and minimize $\mathcal{C}$ such that $\mathcal{C}(f; m)$ is small

➤ 3. “Learn” to reconstruct

- ❖ (Probably what you expect from this talk)

learn $\mathcal{R}_\theta$ such that $f = \mathcal{R}_\theta(m)$



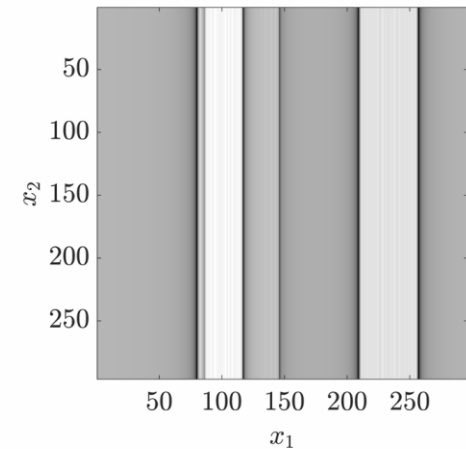
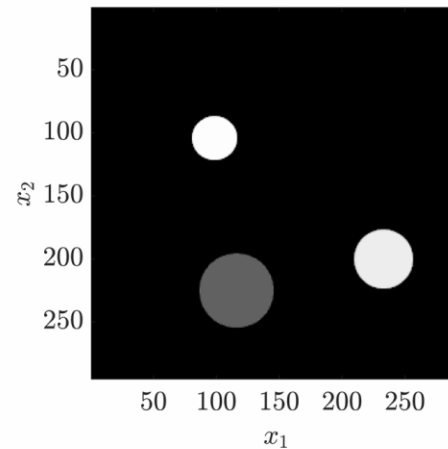
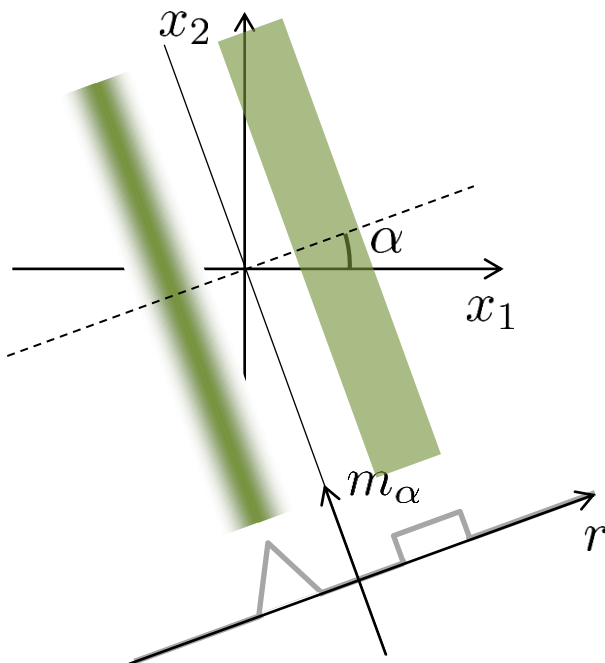
Option #1: Analytical Methods

➤ Example with CT (filtered backprojection)

$$f(x_1, x_2) = \int_0^\pi m_\alpha^{\text{filt}}(x_1 \cos \alpha + x_2 \sin \alpha) d\alpha$$

$$\hat{m}_\alpha^{\text{filt}}(\xi) = |\xi| \hat{m}_\alpha(\xi)$$

↑
Ramp filter



➤ Pros

- ❖ Elegant
- ❖ Theoretical guarantees
- ❖ Usually fast implementation
- ❖ What if only few measurements are available?
 - For dose reduction/short scans

➤ Cons

- ❖ Not always possible to derive a solution
- ❖ Influence of noise?
- ❖ What if only few measurements are available?
 - Random or pseudo-random

Option #2: Optimization-Based Methods

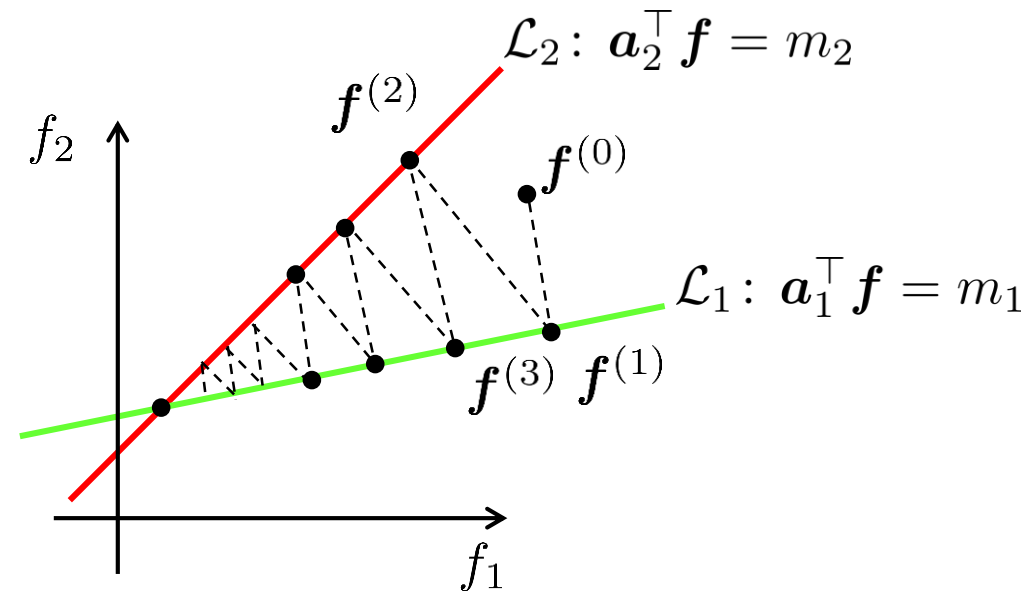
- Look for an image with small residuals

$$r = m - Af \approx 0$$

❖ A simple example:

$$\begin{bmatrix} m_1 \\ m_2 \end{bmatrix} = \begin{bmatrix} a_{1,1} & a_{1,2} \\ a_{2,1} & a_{2,2} \end{bmatrix} \begin{bmatrix} f_1 \\ f_2 \end{bmatrix}$$

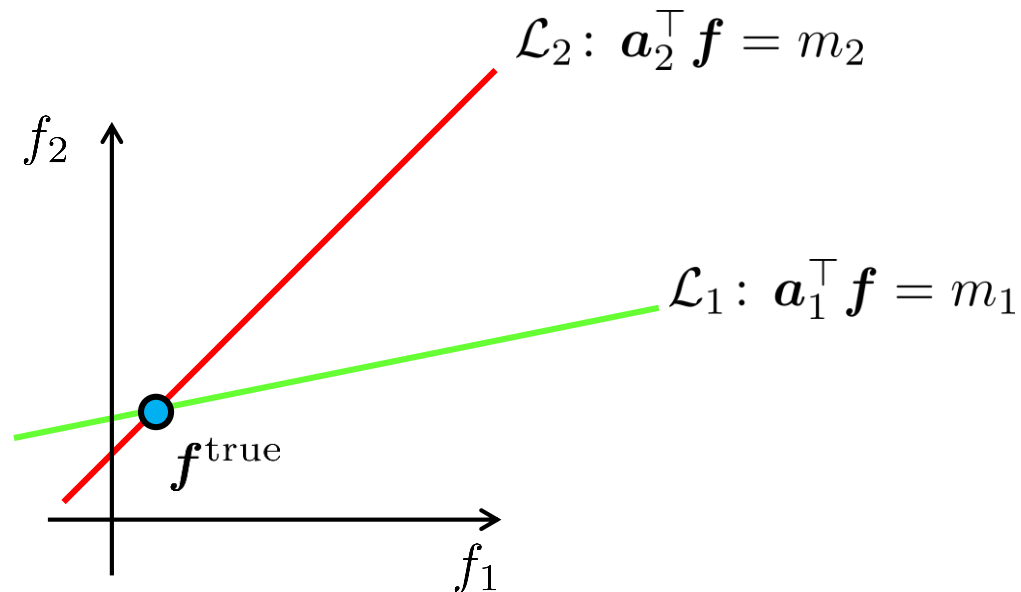
Algebraic reconstruction
technique (ART) [Gordon R., 1970]



- Look for an image with small residuals

$$r = m - Af \approx 0$$

- Influence of noise

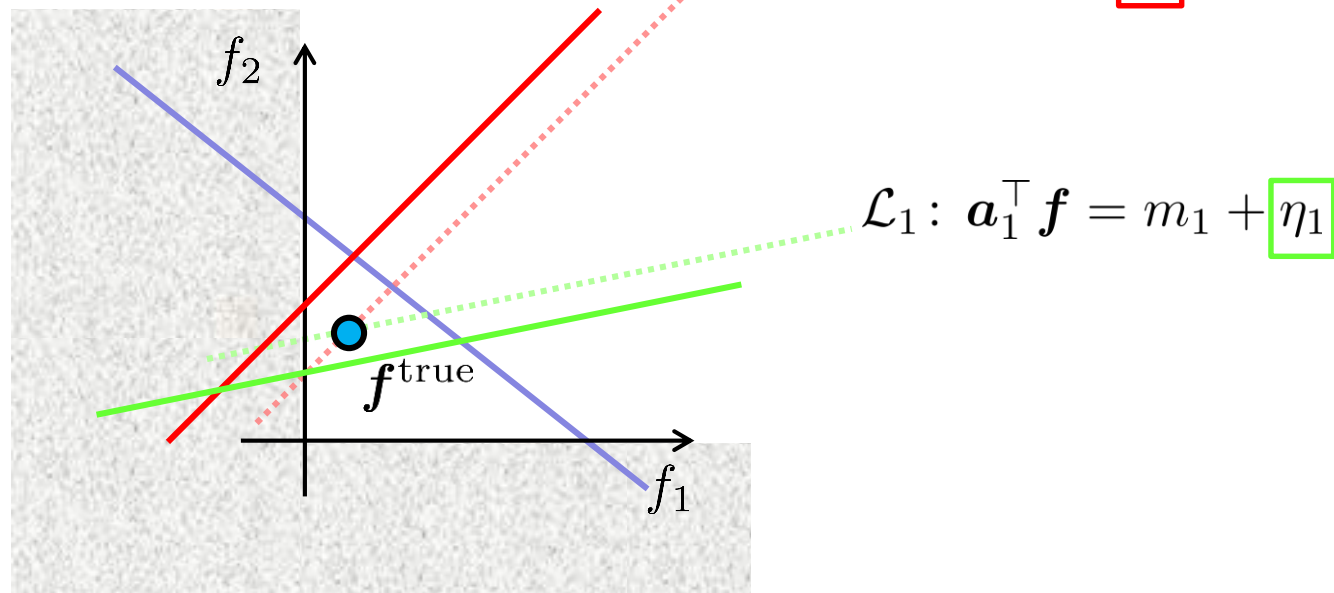


- Look for an image with small residuals

$$r = m - Af \approx 0$$

- Influence of noise

- ❖ More measurements (i.e., $M > N$)
- ❖ Prior knowledge (e.g., $f > 0$)



➤ Typical cost functions

$$\min_f \mathcal{D}(f, m) + \lambda \Psi(f)$$

Data fidelity

Regularization (prior)

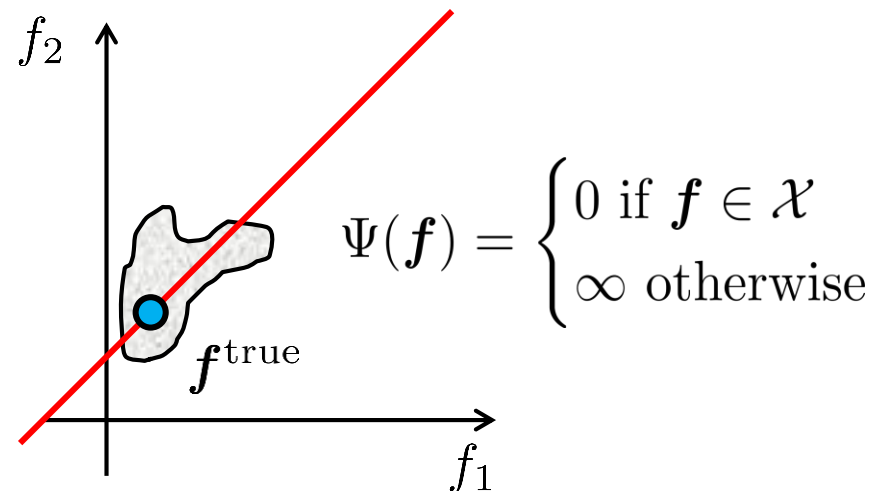
- ❖ Data fidelity is related the noise model/measurements confidence.
E.g.

$$\mathcal{D}(f; m) = \|m - Af\|_2^2$$

$$\mathcal{D}(f; m) = \text{KL}(m, Af)$$

- ❖ Regularization convey prior knowledge about the solution.

$$M \ll N$$



➤ Typical regularizers

❖ Quadratic/Tikhonov regularization

$$\Psi(\mathbf{f}) = \|\mathbf{f}\|_2^2$$

... leads to

$$\mathbf{f}^* = \mathbf{A}^\top (\mathbf{A}\mathbf{A}^\top + \lambda \mathbf{I}_M)^{-1} \mathbf{m}$$

❖ Sparsity-promoting

$$\Psi(\mathbf{f}) = \|\Phi \mathbf{f}\|_1$$

... requires iterative algorithms

$$\mathbf{z}^{(k)} = \mathbf{f}^{(k-1)} - \eta \mathbf{A}^\top (\mathbf{A} \mathbf{f}^{(k-1)} - \mathbf{m})$$

$$\mathbf{f}^{(k)} = \text{prox}_{\lambda \Psi}(\mathbf{z}^{(k)})$$

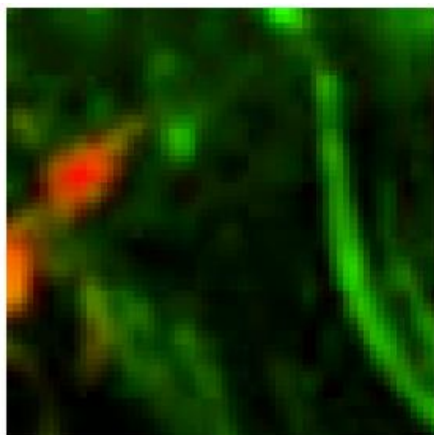
gradient of
data fidelity

proximal
operator of
regularizer

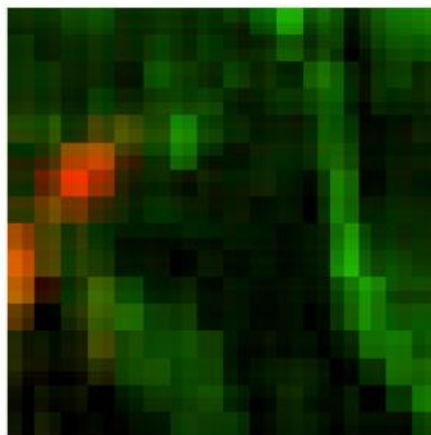
➤ Illustrative results

$N = 64 \times 64$ image
 $M = 333$ measurements
 $N / M \approx 8\%$

Ground-Truth



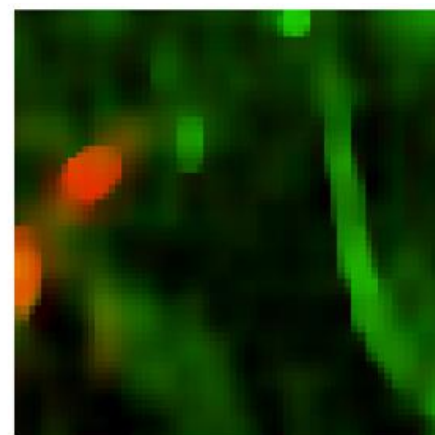
$$\Psi(\mathbf{f}) = \|\mathbf{f}\|_2^2$$



++ Analytical solution
(fast computation)

-- Image quality

$$\Psi(\mathbf{f}) = \|\nabla \mathbf{f}\|_1$$



-- Iterative algorithms
(time consuming)

++ Image quality



Option #3: Learning-Based Methods

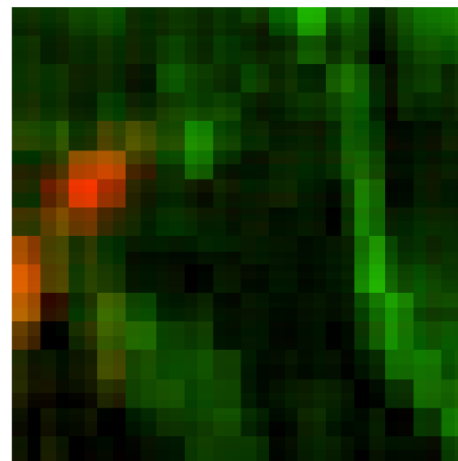
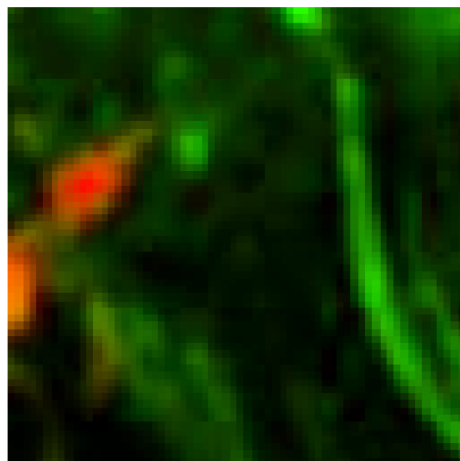
3. Learning-Based Methods

20

➤ Optimization- vs learning-based methods

$N = 64 \times 64$ image
 $M = 333$ measurements
 $N / M \approx 8\%$

Ground-Truth

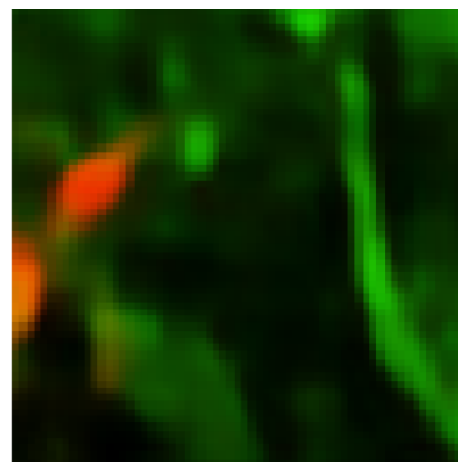
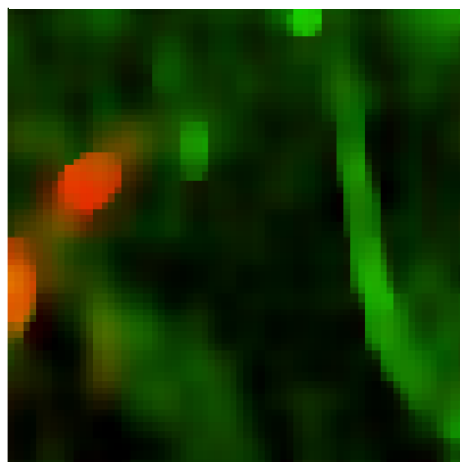


$$\Psi(f) = \|f\|_2^2$$

(pinv)

$$\Psi(f) = \|\nabla f\|_1$$

+2.5 dB
(wrt pinv)



CNN
+3.5 dB
(wrt pinv)

[N. Ducros et al., IEEE
ISBI, 2019]

➤ **Our dream is to find**

$$\mathcal{R}^* : \mathbb{R}^M \mapsto \mathbb{R}^N \text{ such that } \mathcal{R}^*(\mathbf{m}) = \mathbf{f}^{\text{true}}$$

... able to reconstruct well any image, i.e., something like

$$\mathcal{R}^* \in \arg \min_{\mathcal{R}} \frac{1}{L} \sum_{\ell} \|\mathcal{R}(\mathbf{m}^{\ell}) - \mathbf{f}^{\ell}\|_2^2$$

Minimum mean
square error
(MMSE) estimator

... Often intractable

➤ **We have to reduce the dimension of the solution space**

❖ E.g.,

$$\mathcal{R}(\mathbf{m}) = \mathbf{W}\mathbf{m} + \mathbf{b},$$

Linear MMSE
estimator

3. Learning-Based Methods

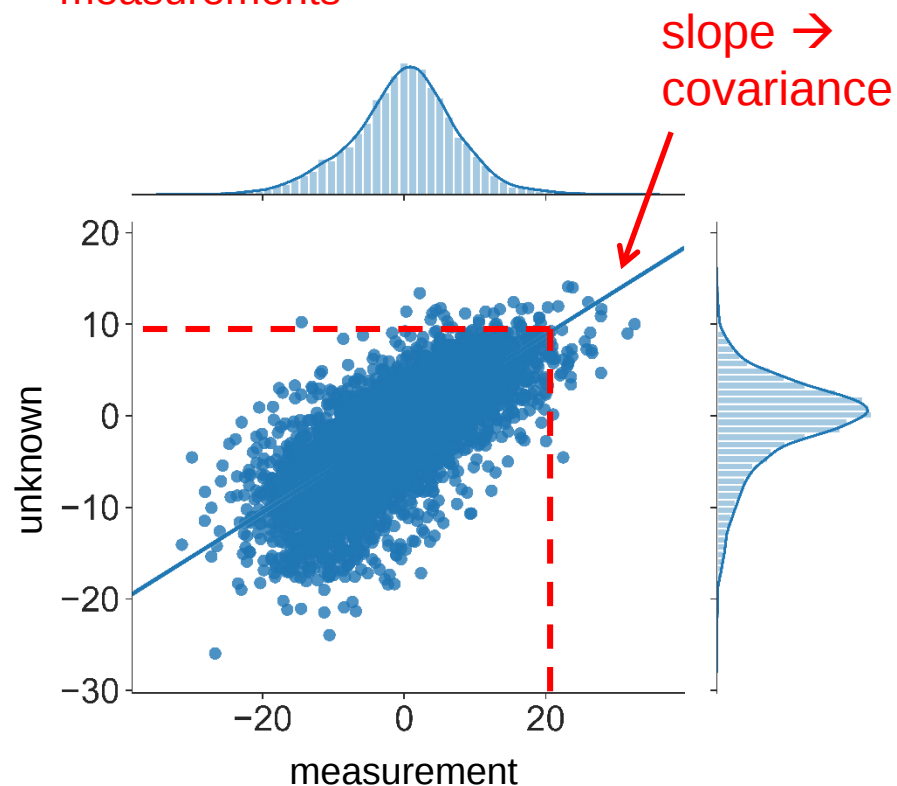
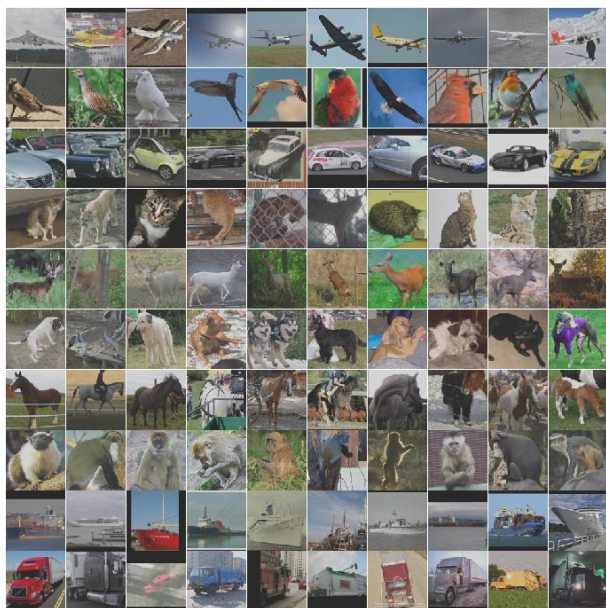
22

➤ Linear MMSE

$$f^* = \bar{f} + \Sigma_{f,m} \Sigma_m^{-1} (m - \bar{m})$$

Covariance between
measurements and
unknowns

Covariance of
measurements



- Learning approaches only reduce the dimension of the solution space to a family of non linear mappings

$$\theta^* \in \arg \min_{\theta} \frac{1}{L} \sum_{\ell} \|\mathcal{R}(\theta; m^{\ell}) - f^{\ell}\|_2^2$$

❖ Training phase

- Image-measurement pairs $\{f^{(\ell)}; m^{(\ell)}\}_{1 \leq \ell \leq L}$
- Loss (e.g., MSE)
- Optimization machinery (i.e., PyTorch or TensorFlow)

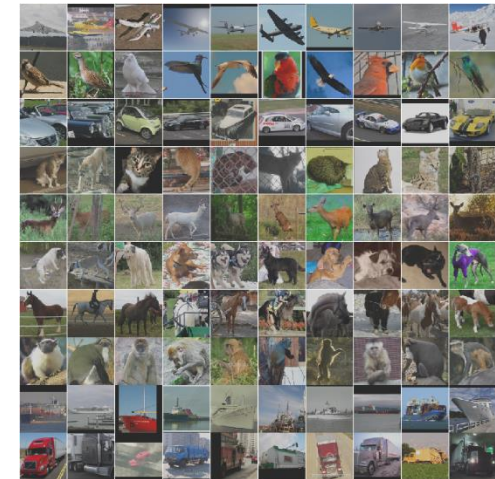
D.P. Kingma and J.L Ba,
ICRL, 2015 (> 215k citations)

A. Paszke *et al.*, NEURIPS,
2019 (> 22k citations)

❖ Reconstruction phase

$$f^* = \mathcal{R}_{\theta^*}(m)$$

STL-10 dataset



➤ Pros

❖ Reconstruction performance

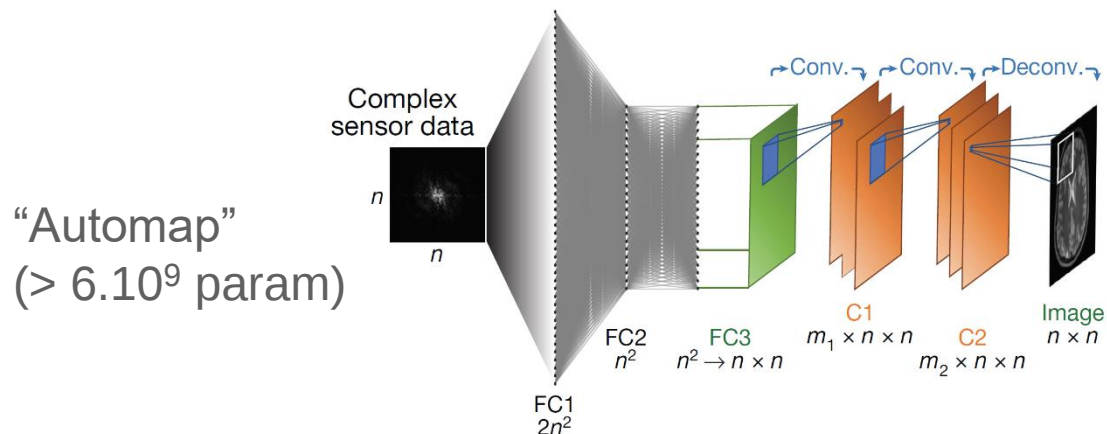
- Empirically excellent (i.e., almost always outperform optimization-based approaches)

❖ Computation times

- Training phase is slow, i.e., several hours or days
- Inference is fast, i.e., tens or hundreds of milliseconds

➤ Cons

- ❖ No reconstruction guarantees (*mathematicians don't like it*)
- ❖ Black box (*radiologists don't like it*)
- ❖ Practical issue
 - How to choose the model?

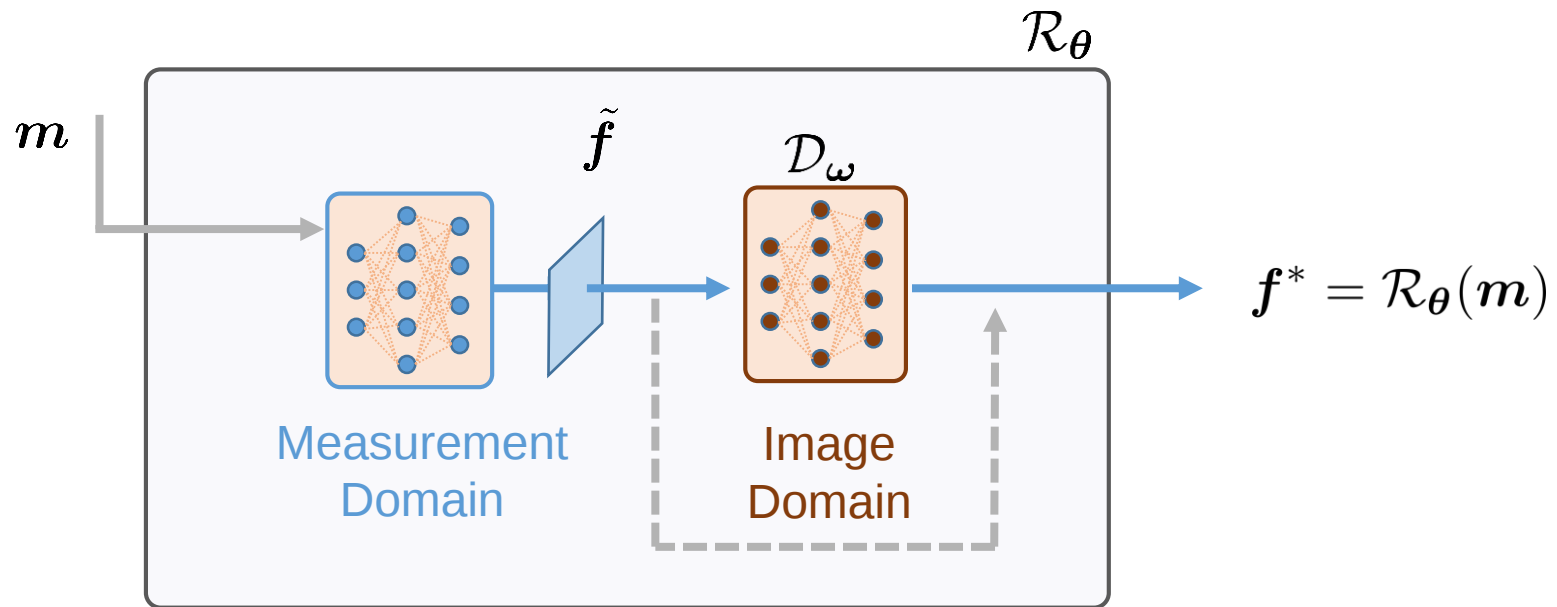


[B. Zhu et al., Nature Letters, 2018] ($> 1.5k$ citations)

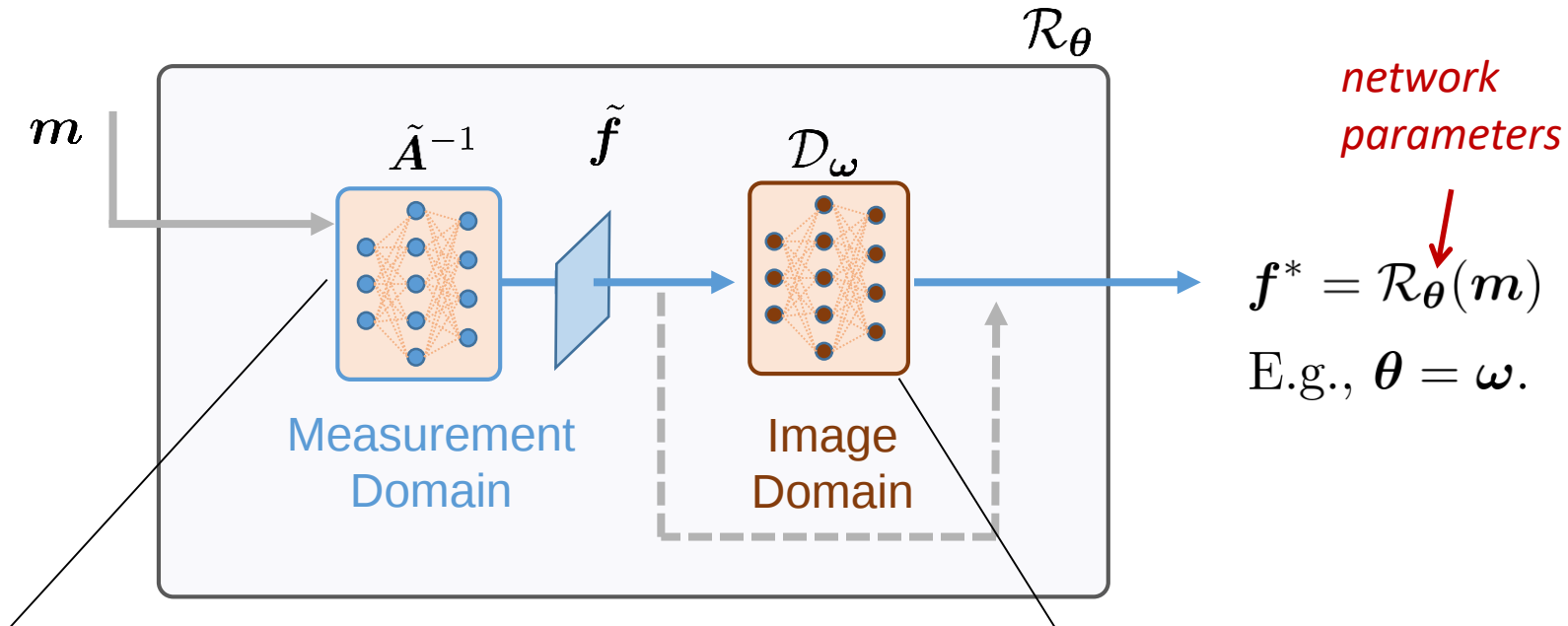
➤ Two-step methods

$$\begin{aligned}\tilde{f} &= \tilde{A}^{-1}m \\ f^* &= \mathcal{D}_\omega(\tilde{f}) + \tilde{f}\end{aligned}$$

where $\tilde{A}^{-1}$ is an approximate inverse of the forward, i.e., $\tilde{A}^{-1}Af \approx f$



- Equivalent to a neural network with a frozen layer



```
Atilde = nn.Linear(..., bias=False, ...)
Atilde.weight.requires_grad = False
```

$$\tilde{\mathbf{A}}^{-1} = \mathbf{A}^\top$$

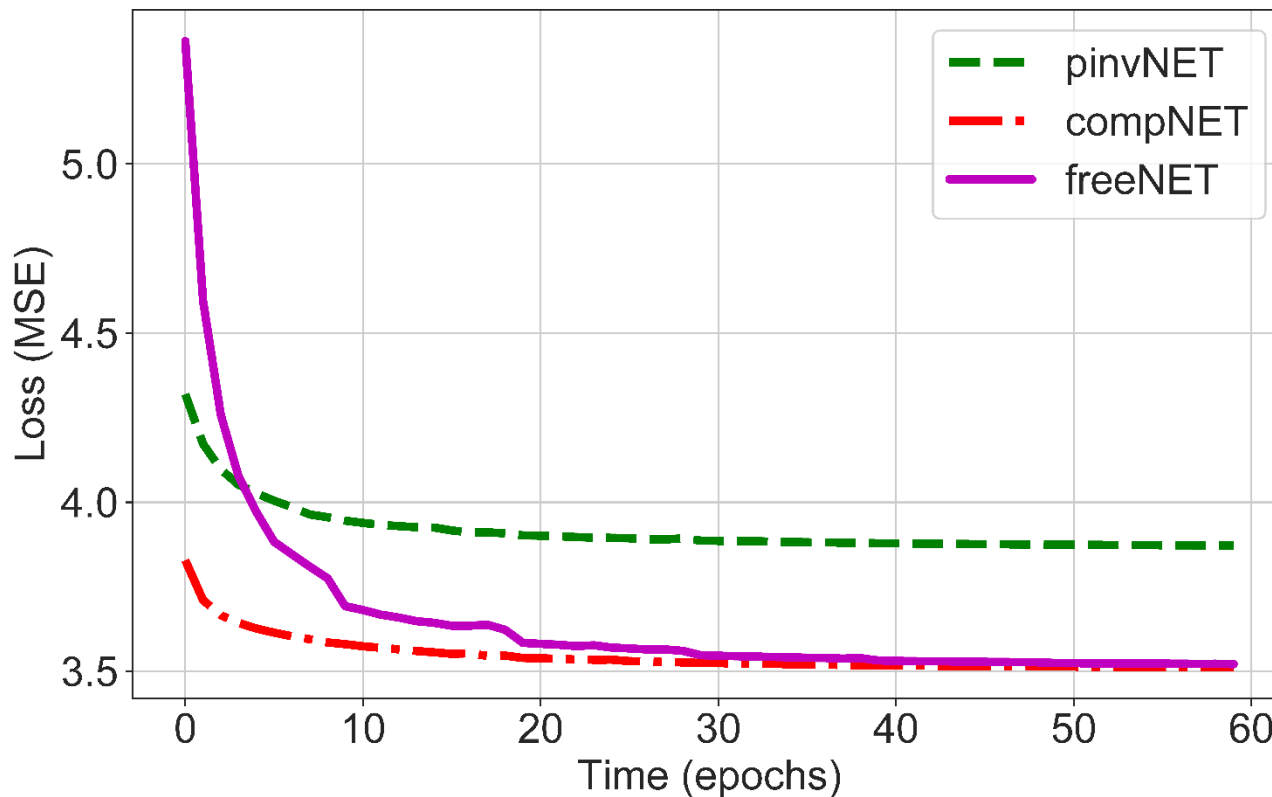
$$\tilde{\mathbf{A}}^{-1} = \mathbf{A}^\dagger \text{ (E.g., } \mathbf{A}^\top (\mathbf{A}\mathbf{A}^\top)^{-1} \text{ for full row rank } \mathbf{A})$$

$$\tilde{\mathbf{A}}^{-1} = \mathbf{A}^\top (\mathbf{A}\mathbf{A}^\top + \mathbf{I})^{-1}$$

```
D = nn.Module(...)
requires_grad = True
```

- **STL-10 (training: ~100k images; test: 8k images)**

$$\sum_{\ell \in \mathcal{I}_{\text{test}}} \|f^{(\ell)} - \mathcal{R}_{\theta}(m^{(\ell)})\|^2$$



$$\begin{aligned}\tilde{\mathbf{A}}^{-1} &= \mathbf{A}^{\dagger} \\ \tilde{\mathbf{A}}^{-1} &= \mathbf{A}^{\top}(\mathbf{A}\mathbf{A}^{\top} + \mathbf{\Sigma})^{-1} \\ \tilde{\mathbf{A}}^{-1} &= \text{Linear}\end{aligned}$$

1.3M vs 8k trainable parameters

[N. Ducros et al., IEEE ISBI, 2019]

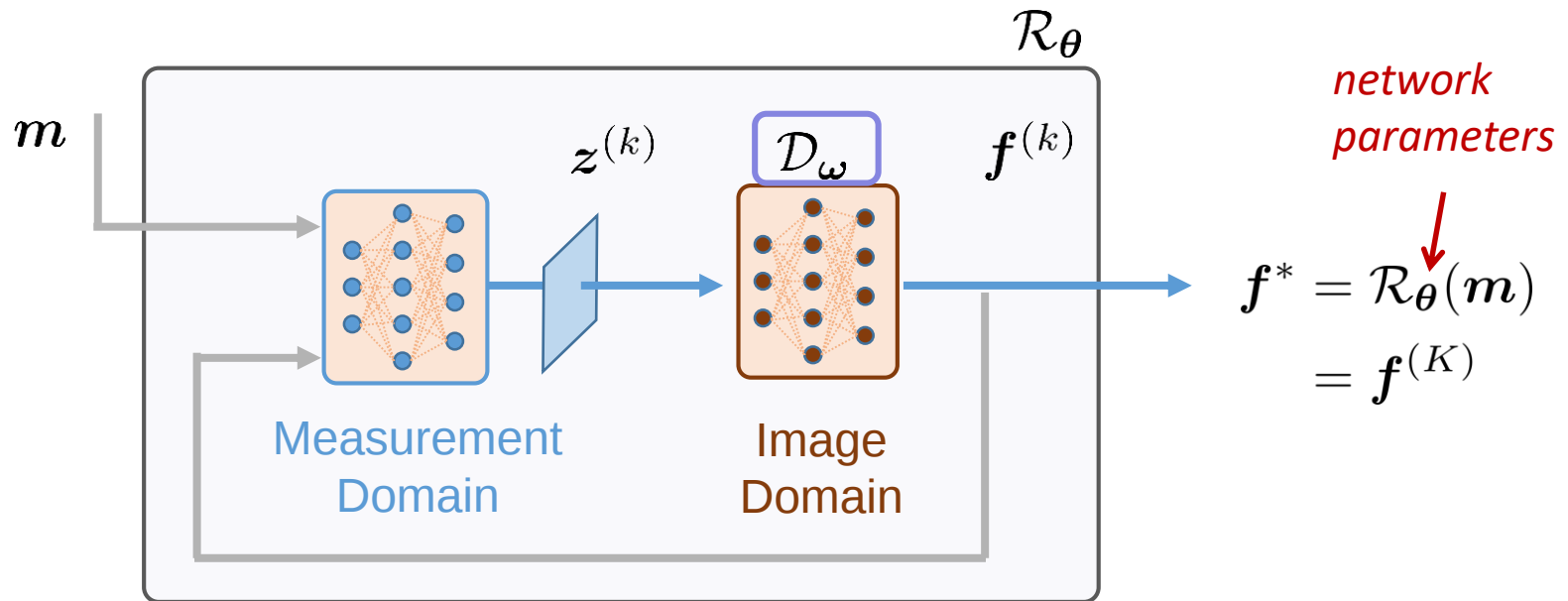
➤ Iterative methods

$$z^{(k)} = f^{(k-1)} - \eta \mathbf{A}^\top (\mathbf{A} f^{(k-1)} - m)$$

$$f^{(k)} = \text{prox}_{\lambda \Psi}(z^{(k)})$$

proximal operator

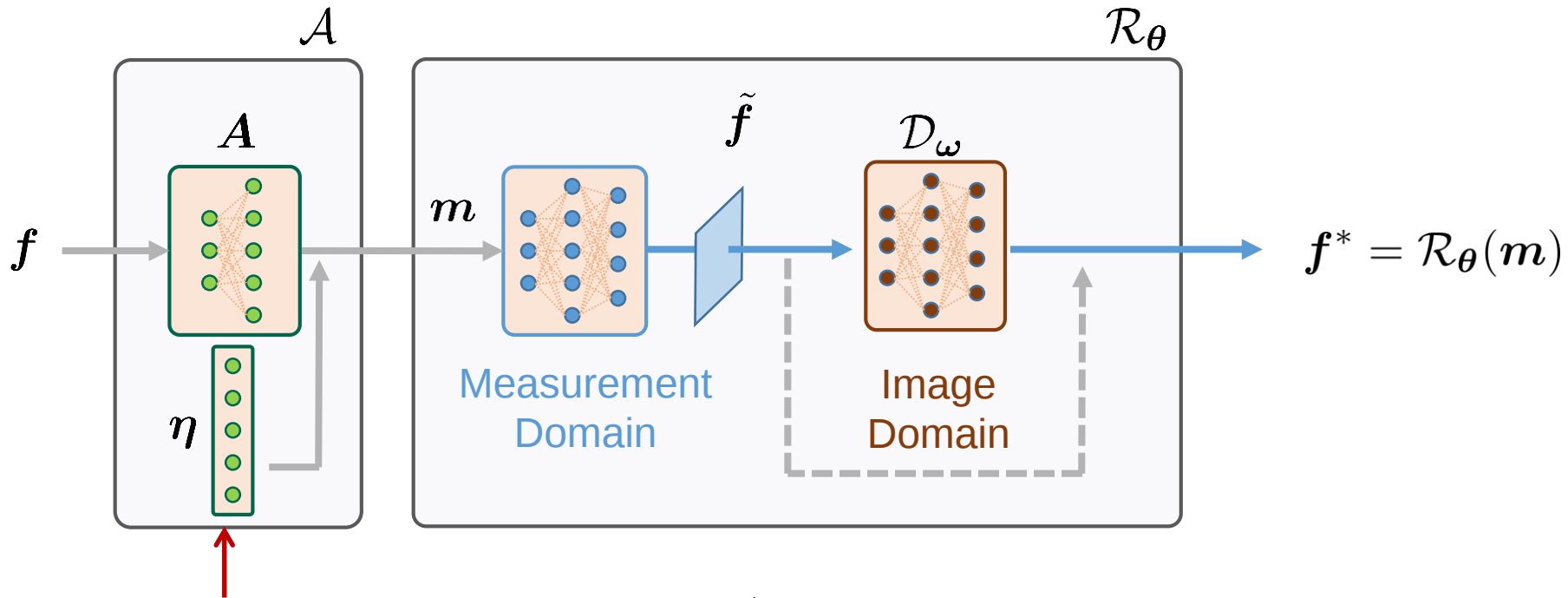
data fidelity



Parameters can be shared across iterations or not

E.g., $\theta = \omega$,
or $\theta = [\omega^{(1)}, \dots, \omega^{(K)}]$.

- With a “physical” module: no need for meas/image pairs



physical
module

$$\theta^* \in \arg \min_{\theta} \frac{1}{L} \sum_{\ell} \|\mathcal{R}_{\theta}(\underline{m}^{\ell}) - \underline{f}^{\ell}\|_2^2$$

$$\mathcal{A}(f) = Af + \eta$$

$$\theta^* \in \arg \min_{\theta} \frac{1}{L} \sum_{\ell} \|(\mathcal{R}_{\theta} \circ \mathcal{A})(\underline{f}^{\ell}) - \underline{f}^{\ell}\|_2^2$$

- **Idea: use denoisers (e.g., BM3D) in place of proximal operators**

$$\begin{aligned} z^{(k)} &= f^{(k-1)} - \eta A^\top (A f^{(k-1)} - m) \\ f^{(k)} &= \boxed{\text{prox}_{\lambda\Psi}}(z^{(k)}) \quad \longrightarrow \quad f^{(k)} = \boxed{\text{Denoise}}(z^{(k)}; \sigma) \end{aligned}$$

proximal operator Denoiser

Noise level 

- **The denoiser can be data-driven. E.g.** $\text{Denoise} = \text{CNN}_{\theta^*}$ **with**

$$\theta^* \in \arg \min_{\theta} \frac{1}{L} \sum_{\ell} \|\text{CNN}_{\theta}(f^{\ell} + \underline{\sigma\epsilon}; \sigma) - f^{\ell}\|_2^2$$

Gaussian noise with variance σ^2  $\epsilon \sim \mathcal{N}(0, 1)$

- **Idea: use denoisers in place of proximal operators**

$$\mathbf{z}^{(k)} = \mathbf{f}^{(k-1)} - \eta \mathbf{A}^\top (\mathbf{A} \mathbf{f}^{(k-1)} - \mathbf{m})$$

$$\mathbf{f}^{(k)} = \text{Denoise}(\mathbf{z}^{(k)}; \sigma)$$

- **Pros**

- ❖ Training is independent of the forward model
 - Flexibility
 - Applies to any inverse problem
- ❖ Adapt to varying noise levels via hyperparameter
- ❖ Convergence can be guaranteed

- **Cons**

- ❖ Manual tuning of hyperparameter
- ❖ Many iterations required compared to supervised methods (E.g., $K = 100\text{—}1,000$ vs $K = 1\text{—}10$)
 - Longer reconstruction times
 - Higher memory requirement
- ❖ Convergence is not always guaranteed (e.g., Lipschitz constant of denoiser)
- ❖ Underlying prior not always known

➤ Deep Generative Models

$$f = \mathcal{R}_\theta(z)$$

Random vector

$$z^* \leftarrow \min_z \|A\mathcal{R}_\theta(z) - m\|_2^2$$

$$f^* = \mathcal{R}_\theta(z^*)$$

➤ Pros

- ❖ Only requires measurements from a single acquisition
- ❖ Theoretical guarantees (based on compressed sensing, e.g., considering Gaussian random matrices)

➤ Cons

- ❖ Long and challenging reconstruction
- ❖ Training of DGM is challenging (lots of data/long times)
- ❖ Stability issues (arbitrary forward models, out-of-distribution images, etc.)

➤ Deep Image Priors

$$f = \mathcal{R}_{\theta}(z)$$

Fixed random vector

$$\theta^* \leftarrow \min_{\theta} \|A\mathcal{R}_{\theta}(z) - m\|_2^2$$

$$f^* = \mathcal{R}_{\theta^*}(z)$$

Note: Minimization must be stopped before convergence (tends to noise otherwise)

➤ Pros

- ❖ Only requires measurements from a single acquisition
- ❖ The reconstruction quality is surprisingly good

➤ Cons

- ❖ Long reconstruction times
- ❖ No guarantees

Conclusions

	Memory Requirement	Recon- struction Time (inference)	Training	Hyperparam/ Comment
Supervised	Low to intermediate	1—10	$\mathbf{A} \quad \{f^\ell\}$	No adaptation (forward, noise)
PnP	Intermediate to high	100—1,000	$\{f^\ell\}$	Noise level
Untrained	Usually low	> 1,000	—	Number of iterations

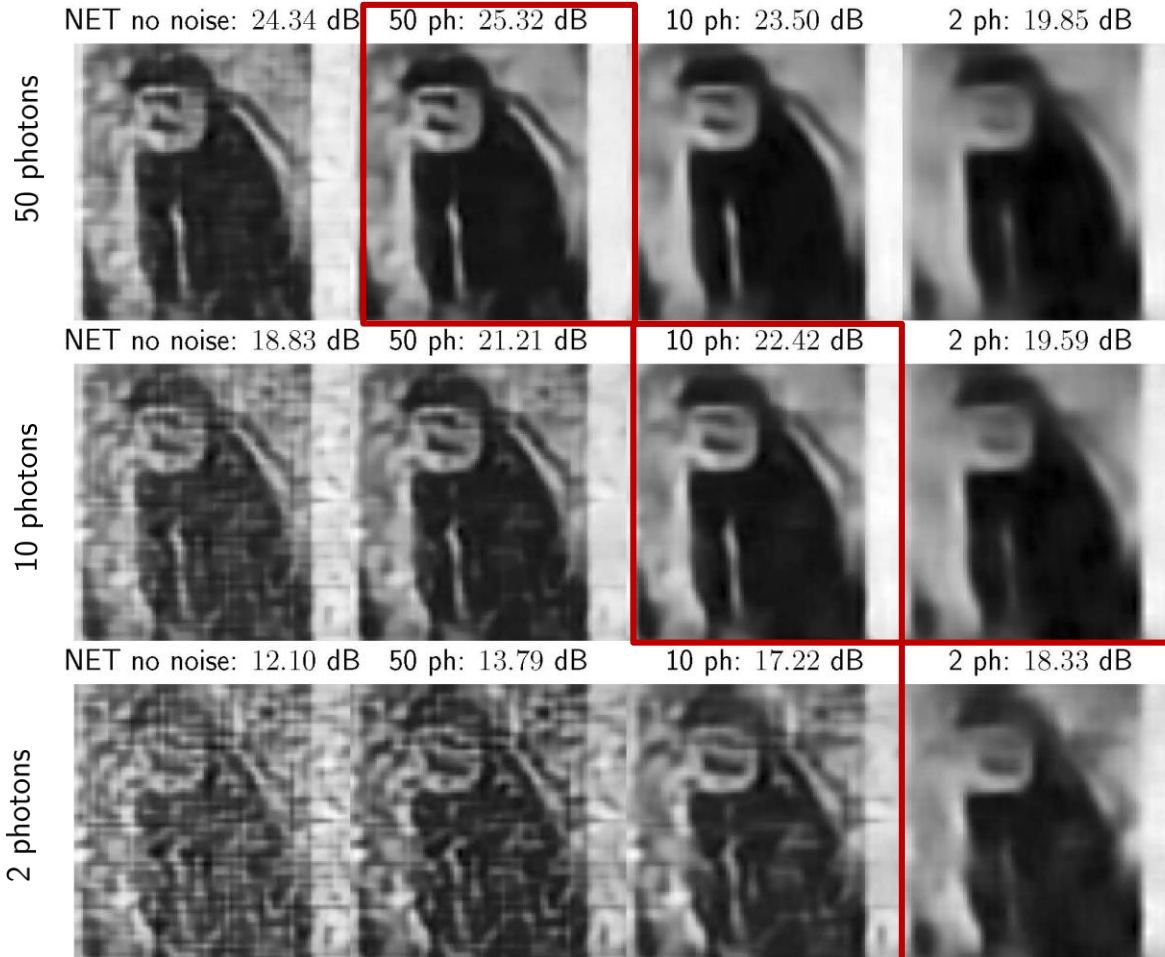
Increasing training noise →

Increasing test noise ↓

Ground-truth



C-Net



[N. Ducros, ISTE Book chapter, 2022])

➤ **Data-driven DL-based approaches for image reconstruction are**

- ❖ Powerful!
- ❖ No longer black boxes
- ❖ Supervised, PnP, based on generative models, untrained, etc.

➤ **Supervised vs PnP methods**

- ❖ Supervised methods usually require fewer parameters
- ❖ Supervised methods performs very well
- ❖ PnP methods adapts to different
 - Imaging modalities (i.e., forward models)
 - Noise levels

➤ **Warning**

- ❖ Handling noise is still an issue.
 - Evaluate the robustness to noise level deviations
 - Train with noise (supervised)
 - Tune hyperparameters (PnP)

→ **Hands-on session on Friday at 2 pm! ←**