

BAYESIAN IMAGE RECONSTRUCTION WITH LOCAL LINEAR REGRESSORS

Colas Schretter colas.schretter@vub,ulb,imec.be



HYBRID GENERATIVE / DISCRIMINATIVE EXPECTATION-MAXIMIZATION

- Sparse mixtures of local linear regressors are used for continuous image representation.
- Model parameters are estimated from limited data with Bayesian maximum *a posteriori* (MAP) imputation of denoised missing samples using the expected conditional variance from the model.
- progressive binary splitting on data likelihood for adaptive model selection.
- Approximate the pristine denoised and upsampled latent image, not overfitting measured data.
- Hybrid expectation-maximization (EM) combines generative (competitive) estimation for marginal Gaussian kernels and discriminative (collaborative) minimum MSE fitting for attached regressors.
- The positions and shapes of marginal kernels act as adaptive-resolution smooth pixels that are not constrained on a grid and may overlap. The model can be rasterized at any required image definition.

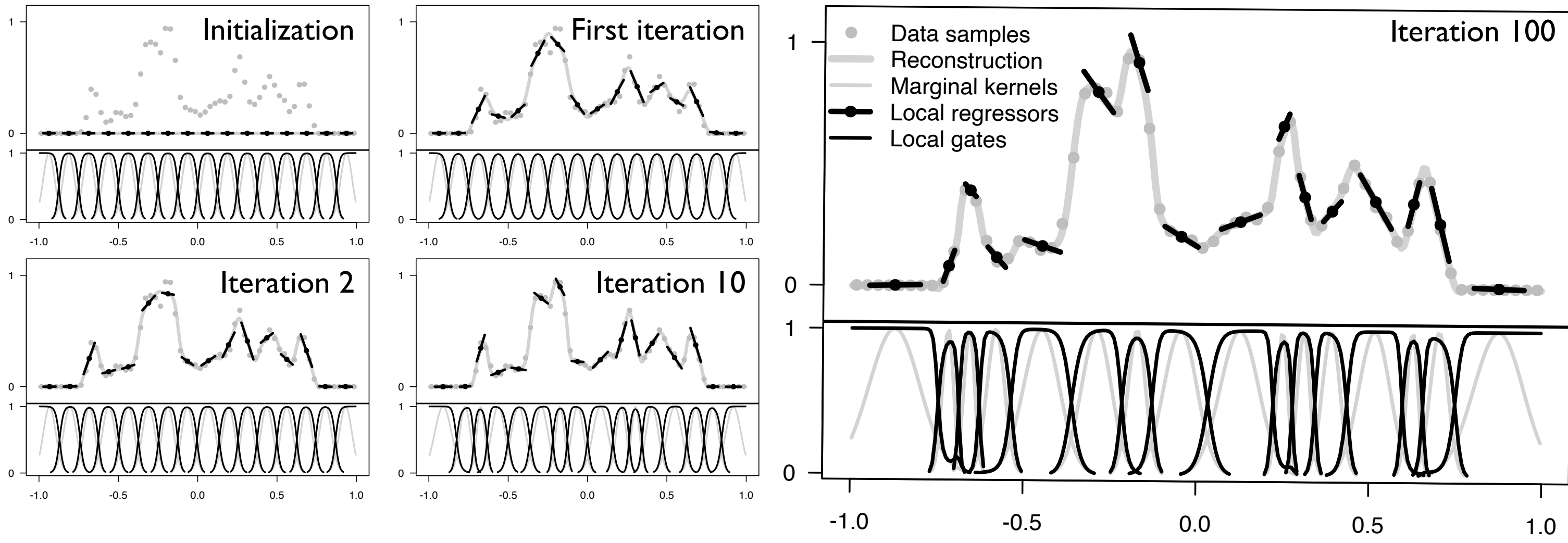
NO DEEP NETWORK
NO BLACK BOXING
NO TRAINING DATA
NO INVERSE CRIMES

GENERATIVE
Competitive — $P(X,Y)$

DISCRIMINATIVE
Collaborative — $P(Y|X)$

EXAMPLE ON A ONE-DIMENSIONAL SIGNAL

Continuous reconstruction with 16 local linear regressors estimated from 64 noise-free data samples



MARGINAL KERNELS, LOCAL LINEAR REGRESSORS, MISSING DATA IMPUTATION

The structured image representation model allows for closed-form parametric updates

Marginal support kernels

The uniform sampling domain is covered by a mixture of marginal receptors (gating) that is represented by a mixture density of K bivariate Gaussian kernels parameterized by a mean position μ and a symmetric covariance matrix Σ :

$$d(x) = \sum_{k=1}^K \pi_k \mathcal{N}(x | \mu_k, \Sigma_k) \quad \text{with} \quad \sum_{k=1}^K \pi_k = 1, \quad (1)$$

where in the 2-dimensional case, $\pi_k \propto 2\pi \sqrt{\det(\Sigma_k)}$. These priors are proportional to the area of the first standard deviation ellipse such that kernels are peak-normalized:

$$d(x) = \sum_{k=1}^K \exp\left(-\frac{1}{2}\|x - \mu_k\|_{\Sigma_k}^2\right), \quad (2)$$

with the squared Mahalanobis distance to the kernel's position

$$\|x - \mu_k\|_{\Sigma_k}^2 = (x - \mu_k)^T \Sigma_k^{-1} (x - \mu_k). \quad (3)$$

The marginal partition of unity defines the relative spatial responsibility for each mixture component $k \in \{1, \dots, K\}$:

$$p(k|x) = \frac{1}{d(x)} \exp\left(-\frac{1}{2}\|x - \mu_k\|_{\Sigma_k}^2\right). \quad (4)$$

Local linear regressors

A local regressor is associated to each marginal kernel. Linear regressors are parameterized by a mean intensity λ and a vector of spatial gradients δ , for capturing locally the smooth intensity variations. Regression planes can be evaluated with

$$\lambda(x|k) = \lambda_k + \delta_k^T \Sigma_k^{-1} (x - \mu_k), \quad (5)$$

such that mixing regressors according to their relative spatial importance yields the average consensus from all regressors:

$$\lambda(x) = \sum_{k=1}^K p(k|x) \lambda(x|k). \quad (6)$$

Moreover, an estimate of the regression error may be computed from the discrepancy among regressor's responses at any query location x . The explained variance from the model is the expected conditional squared error from the mean:

$$\sigma^2(x) = \sum_{k=1}^K p(k|x) (\lambda(x|k) - \lambda(x))^2, \quad (7)$$

and this quantity will be useful for adaptive data-driven model selection in a Bayesian estimation by weighting the importance of regressors according to their radiometric proximity.

Missing data imputation

Let $Y(x) : [0, N_x] \times [0, N_y] \rightarrow \mathbb{R}$ be a pristine continuous intensity image to reconstruct from a grid D of $N_i \times N_j$ blurred noisy data samples $\tilde{d}_{i,j}$. One must estimate $K \times 8$ parameters:

$$\Theta = \{(\mu_1, \Sigma_1, \lambda_1, \delta_1), \dots, (\mu_K, \Sigma_K, \lambda_K, \delta_K)\} \quad (8)$$

such that the log-likelihood of observing the ground truth latent image Y is maximized, given the model's parameters. This problem is different than estimating parameters for maximizing the log-likelihood of observing input data D :

$$L(\Theta|D) = \sum_{i=1}^{N_i} \sum_{j=1}^{N_j} \log \left(\sum_{k=1}^K p(x_{i,j}, \tilde{d}_{i,j}|k) \right) \quad (9)$$

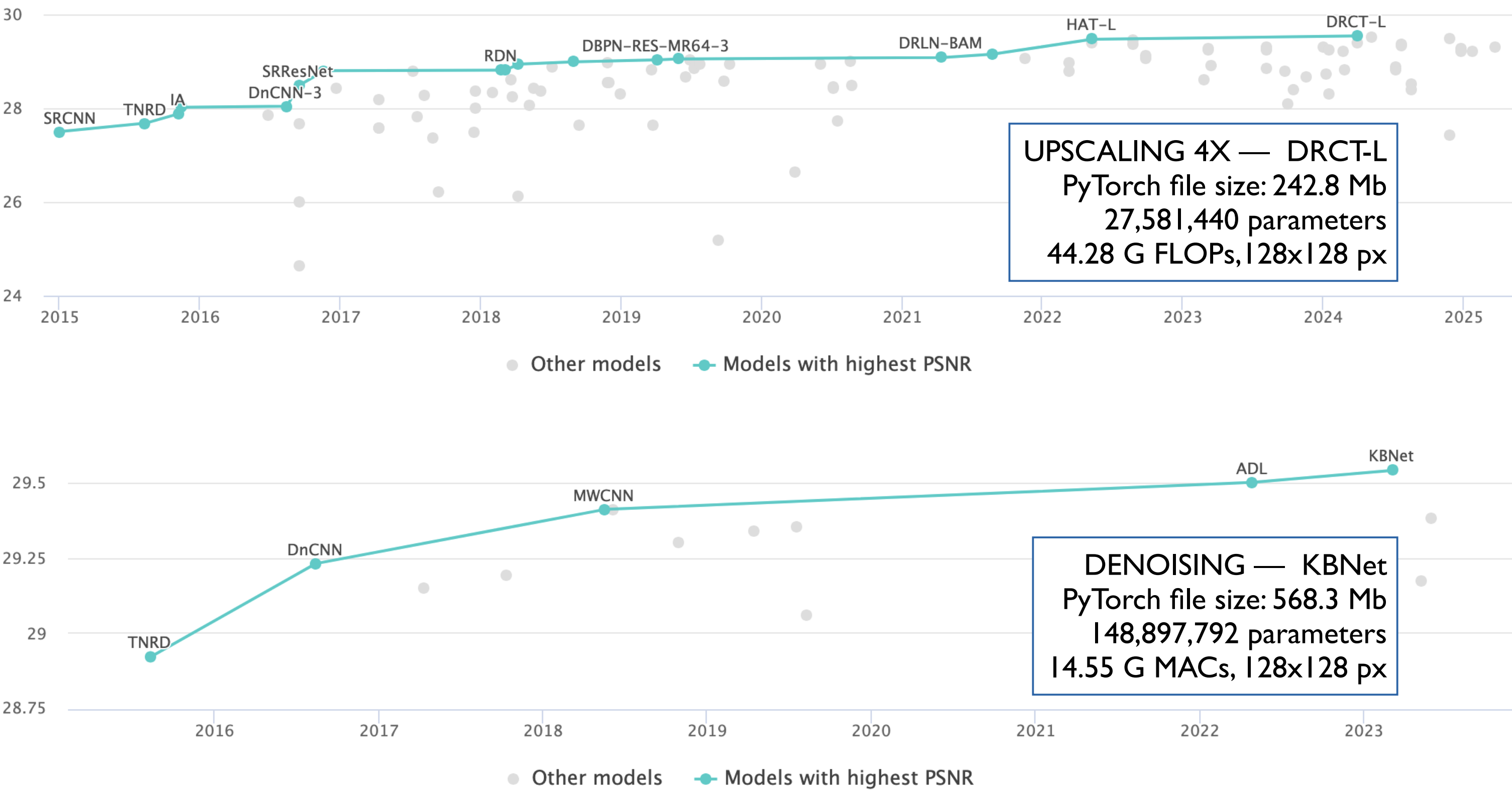
while the log-likelihood from the continuous latent image is

$$L(\Theta|Y) = \int_0^{N_i} \int_0^{N_j} \log \left(\sum_{k=1}^K p([u,v]^T, Y([u,v]^T)|k) \right) du dv \quad (10)$$

In the above formulations, $L(\Theta|D) \approx L(\Theta|Y)$ only if a very large number of noise-free samples is captured. Imputation is achieved by computing the current upsampled regression at the subpixel locations and add the blurred data residual.

STALLED IMPROVEMENT IN NEURAL NETWORKS FOR IMAGE RESTORATION

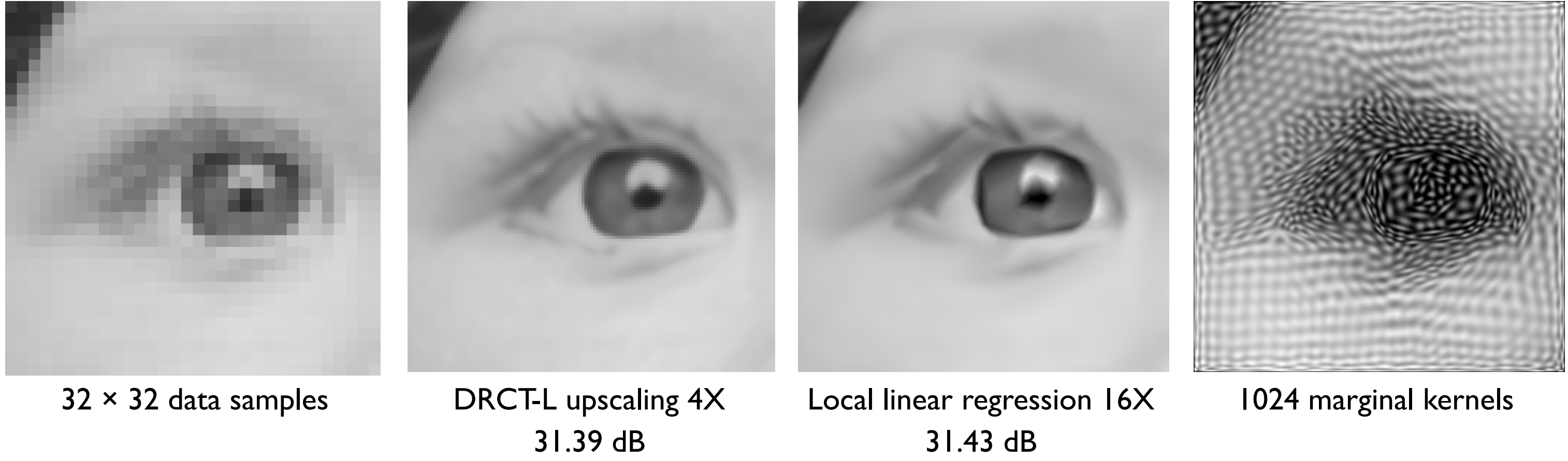
SGD training with ADAM overfits training sets, hence models are very large and not robust



State of the Art from www.paperswithcode.com (Acquired by Meta in 2019, shutdown in July 2025)

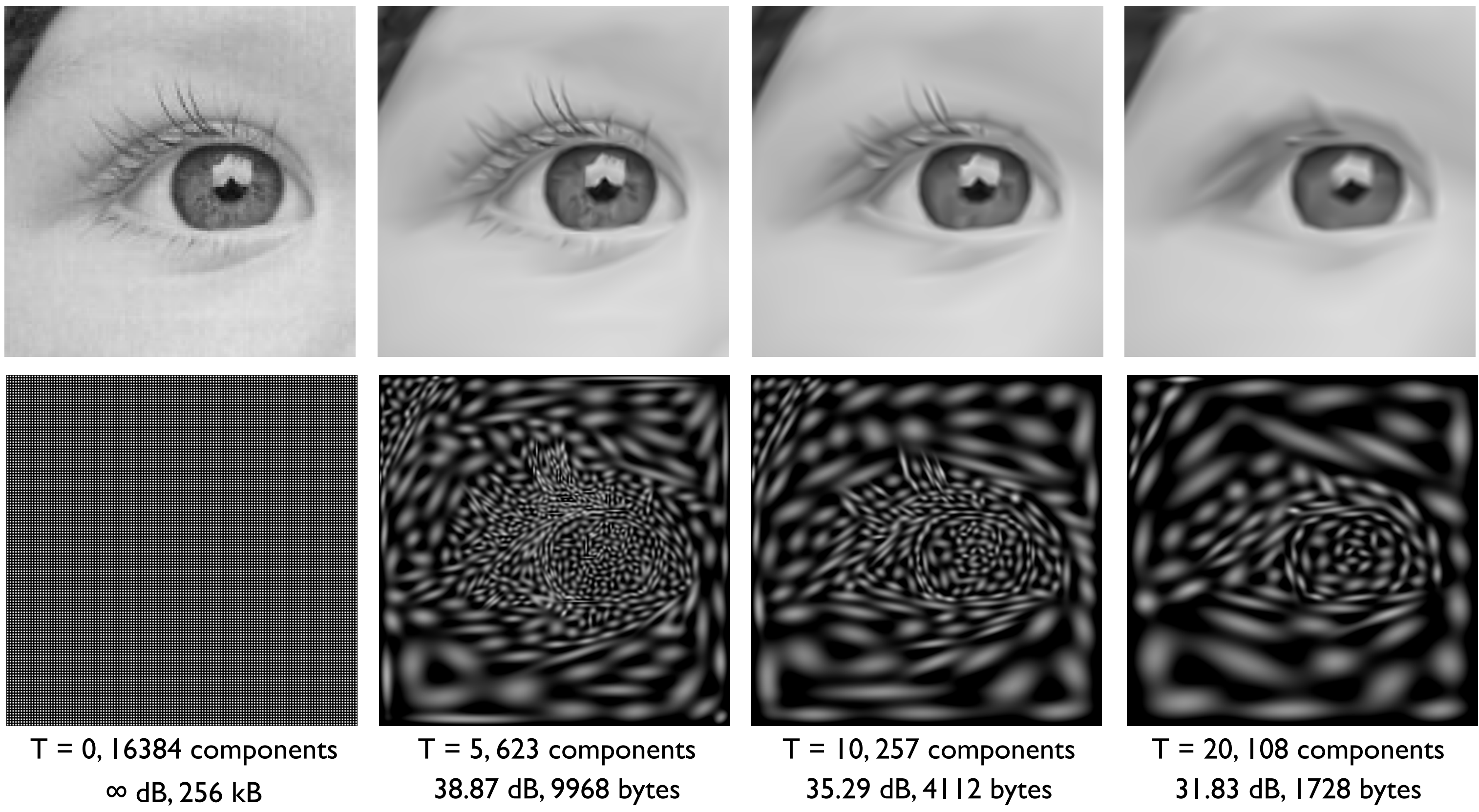
NAIVE IMAGE UPSCALING WITH DENSE MODELS — ONE KERNEL PER DATA SAMPLE

Comparison between local linear regressors and a supervised learning model for image upscaling



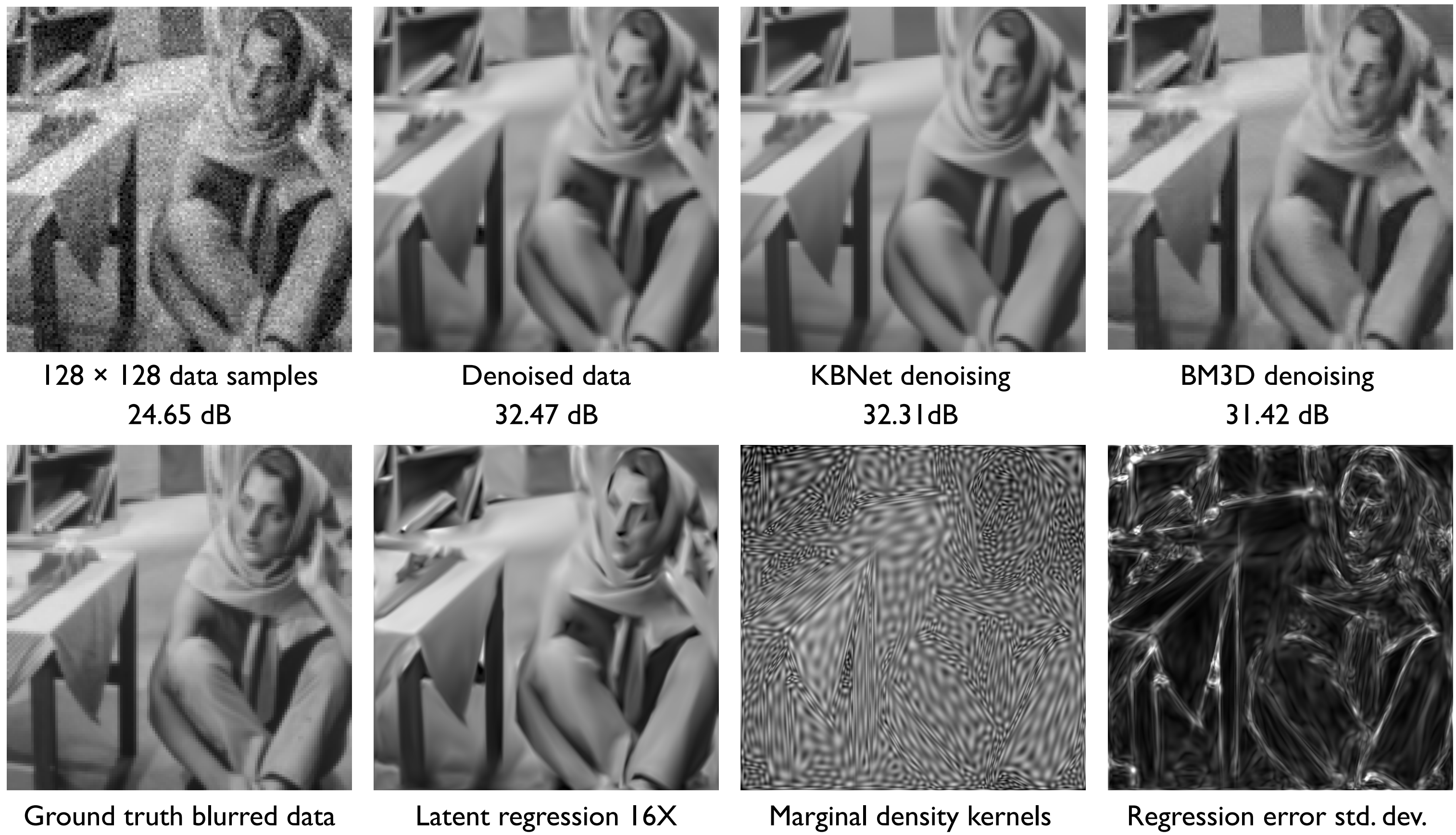
PROGRESSIVE MODEL SELECTION FOR IN-CORE IMAGE COMPRESSION

Approximations with decreasing model sizes in function of the splitting threshold parameter T



JOINT DATA DENOISING, IMAGE UPSCALING AND SPARSE MODEL SELECTION

Restoration with 1619 local regressors, estimated from 16384 blurred and noisy data samples



ADAPTIVE MODELS WITH BINARY SPLITTING BASED ON LOCAL LIKELIHOOD

Components are split in two where local linear regression poorly explains the imputed data

