

Deconvolution



15-463, 15-663, 15-862
Computational Photography
Fall 2025, Lecture 11

Course announcements

- Homework assignment 4 due next Friday.

Overview of today's lecture

- Sources of blur.
- Deconvolution.
- Blind deconvolution.

Slide credits

Most of these slides were adapted from:

- Fredo Durand (MIT).
- Gordon Wetzstein (Stanford).

Why are our images blurry?

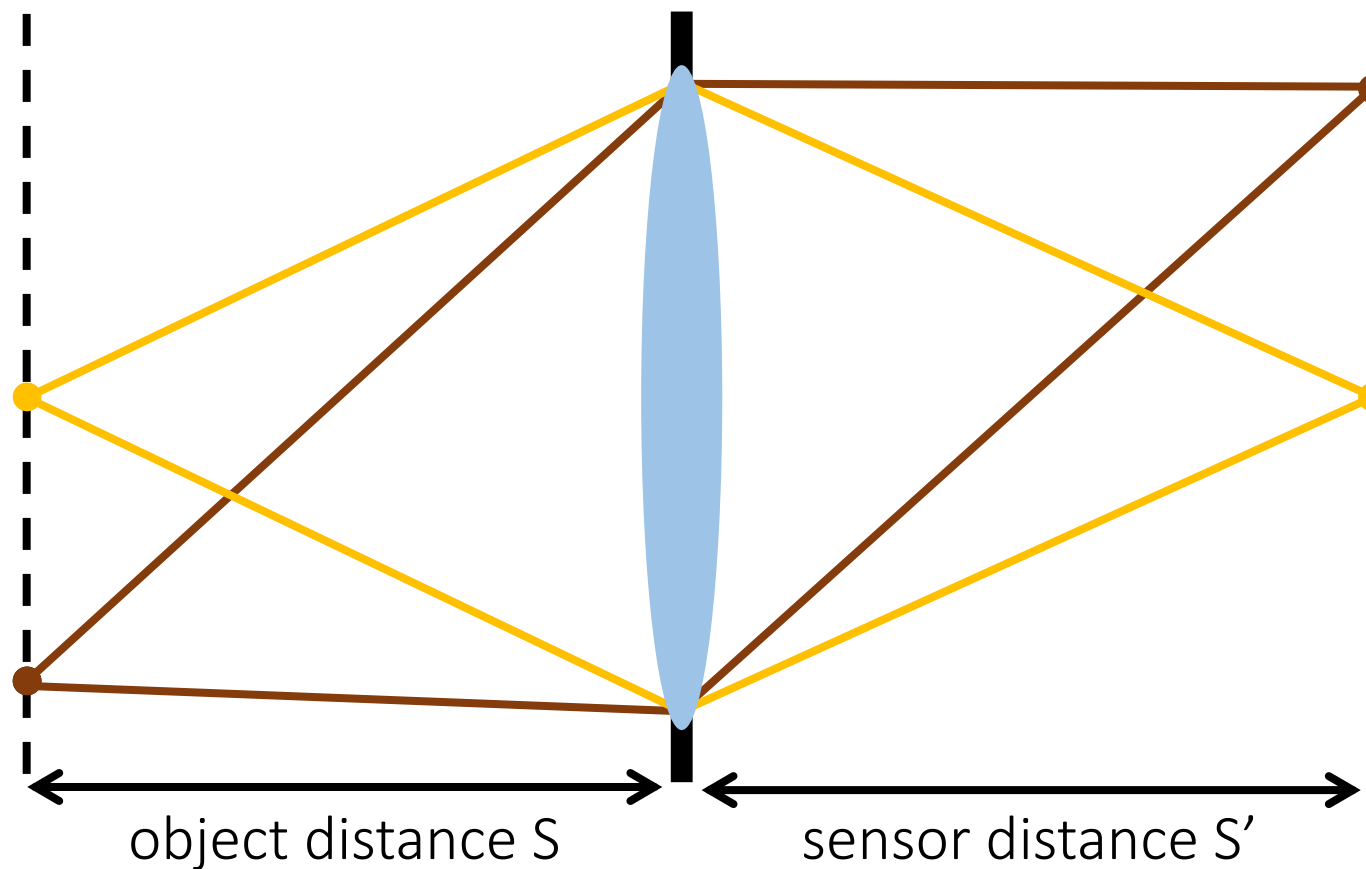
Why are our images blurry?

- Lens imperfections.
- Camera shake.
- Scene motion.
- Depth defocus.

Lens imperfections

- Ideal lens: A point maps to a point at a certain plane.

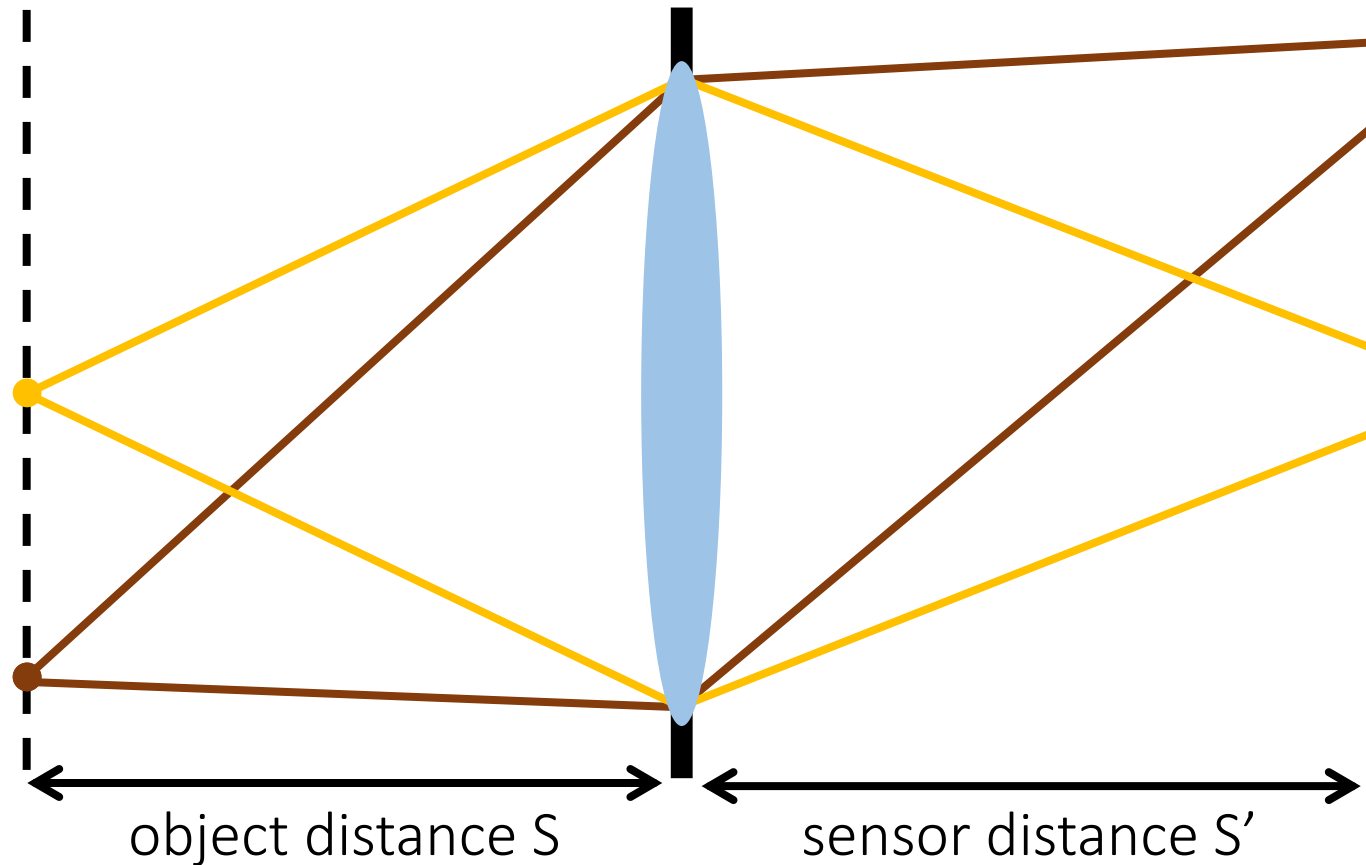
$$\frac{1}{S'} + \frac{1}{S} = \frac{1}{f}$$



Lens imperfections

- Ideal lens: A point maps to a point at a certain plane.
- Real lens: A point maps to a circle that has non-zero minimum radius among all planes.

$$\frac{1}{S'} + \frac{1}{S} = \frac{1}{f}$$

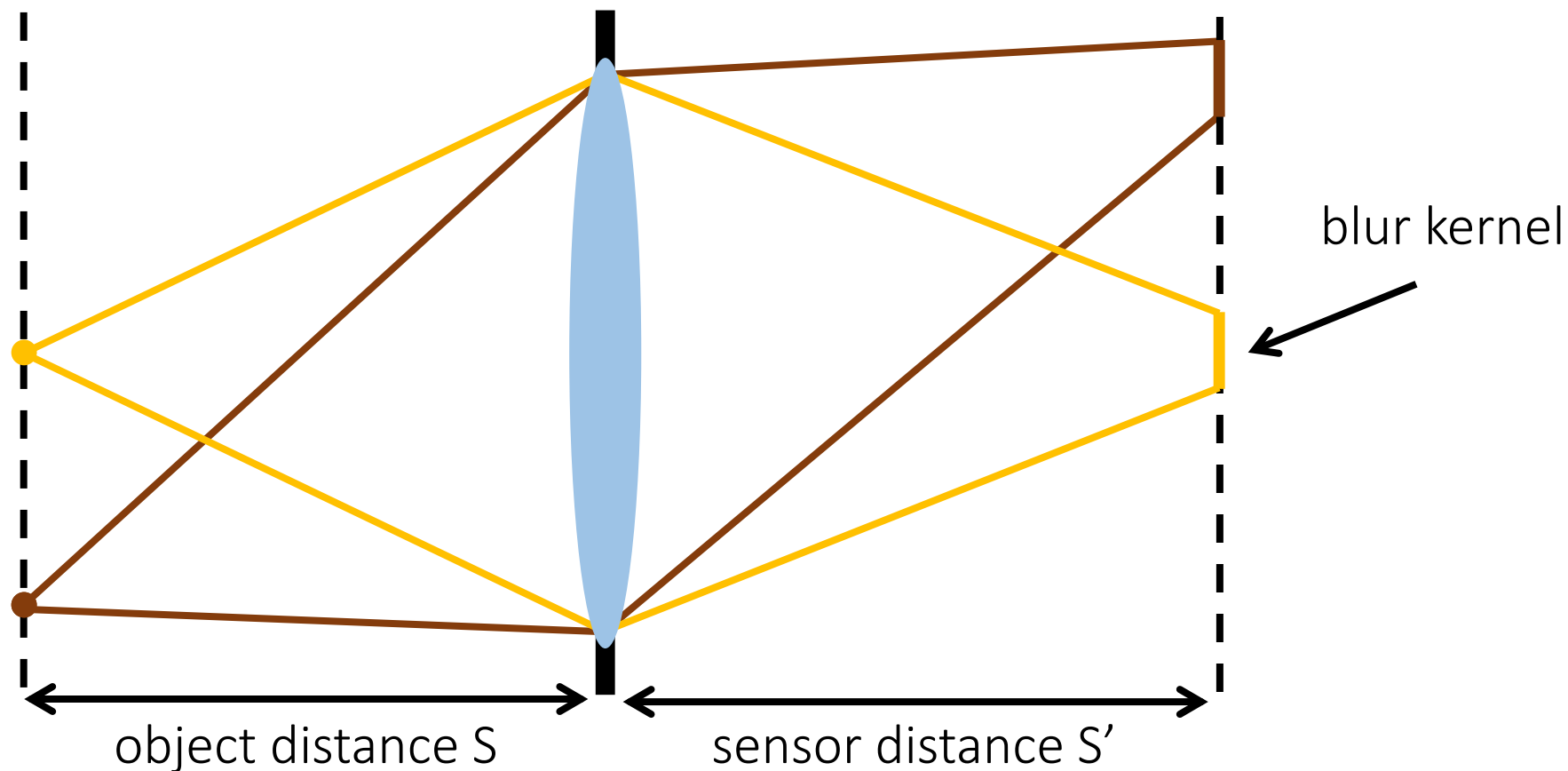


What is the effect of this on the images we capture?

Lens imperfections

- Ideal lens: A point maps to a point at a certain plane.
- Real lens: A point maps to a circle that has non-zero minimum radius among all planes.

$$\frac{1}{S'} + \frac{1}{S} = \frac{1}{f}$$



Shift-invariant blur.

Lens imperfections

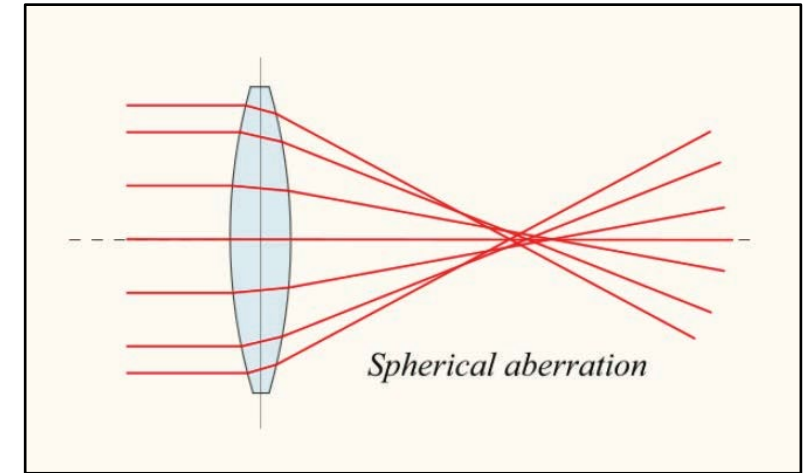
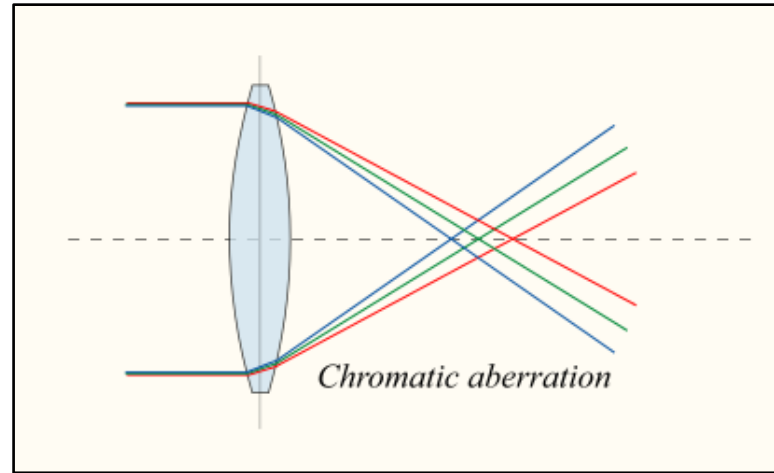
What causes lens imperfections?

Lens imperfections

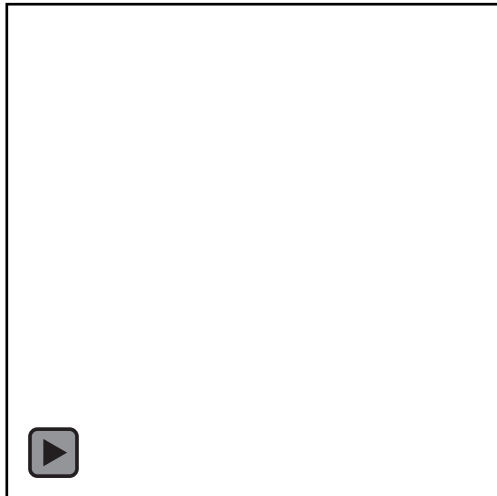
What causes lens imperfections?

- Aberrations.

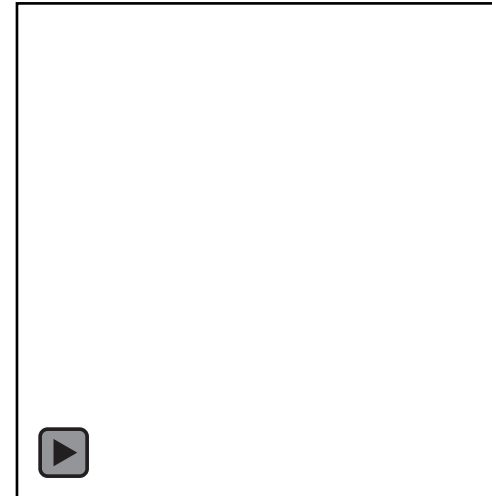
(Important note: Oblique aberrations like coma and distortion are not shift-invariant blur and we do not consider them here!)



- Diffraction.



small
aperture



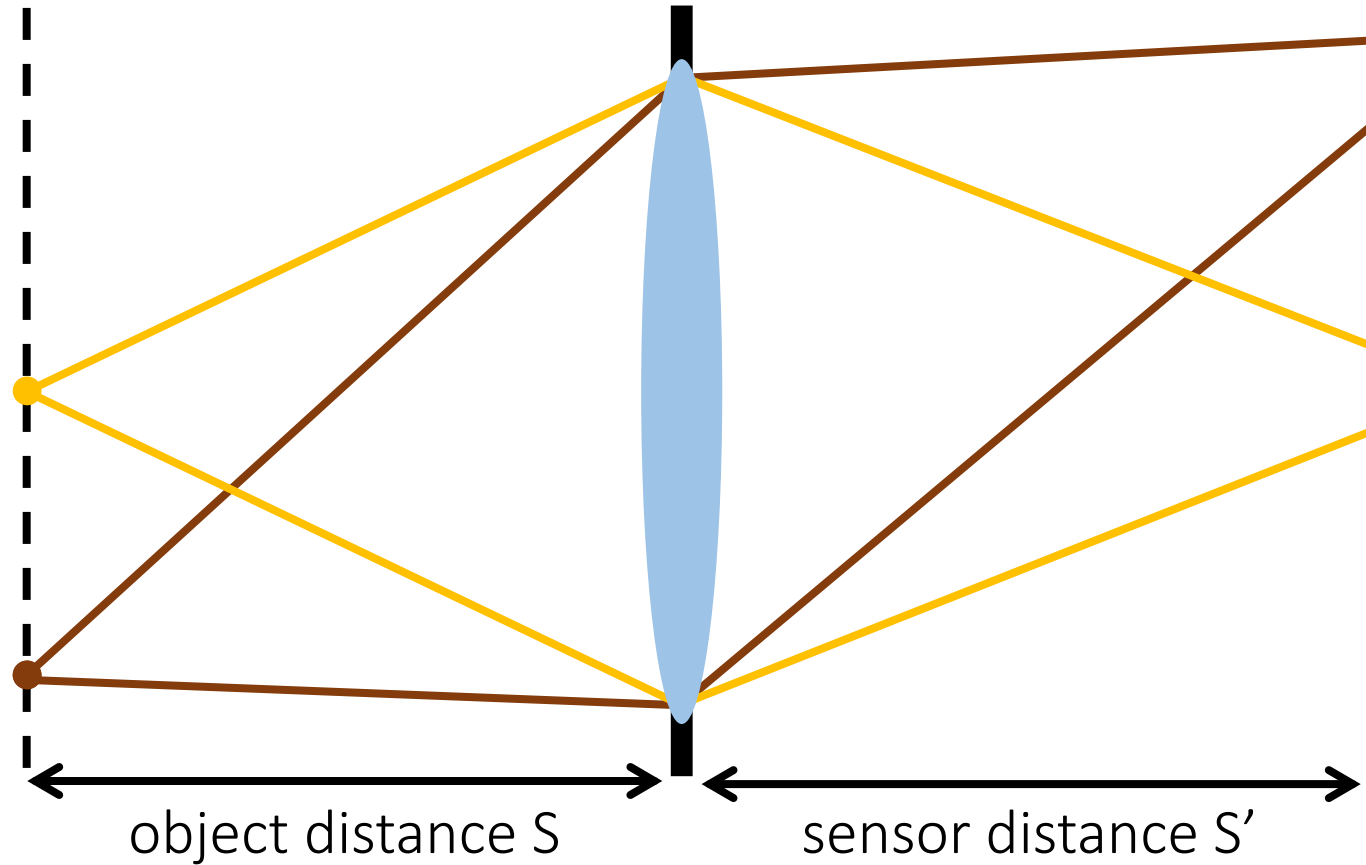
large
aperture

Lens as an optical low-pass filter

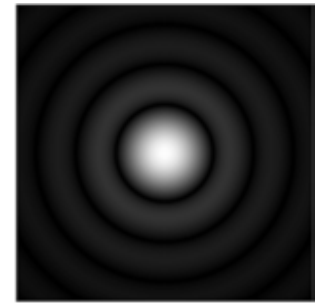
Point spread function (PSF): The blur kernel of a lens.

- “Diffraction-limited” PSF: No aberrations, only diffraction. Determined by aperture shape.

$$\frac{1}{S'} + \frac{1}{S} = \frac{1}{f}$$



blur kernel



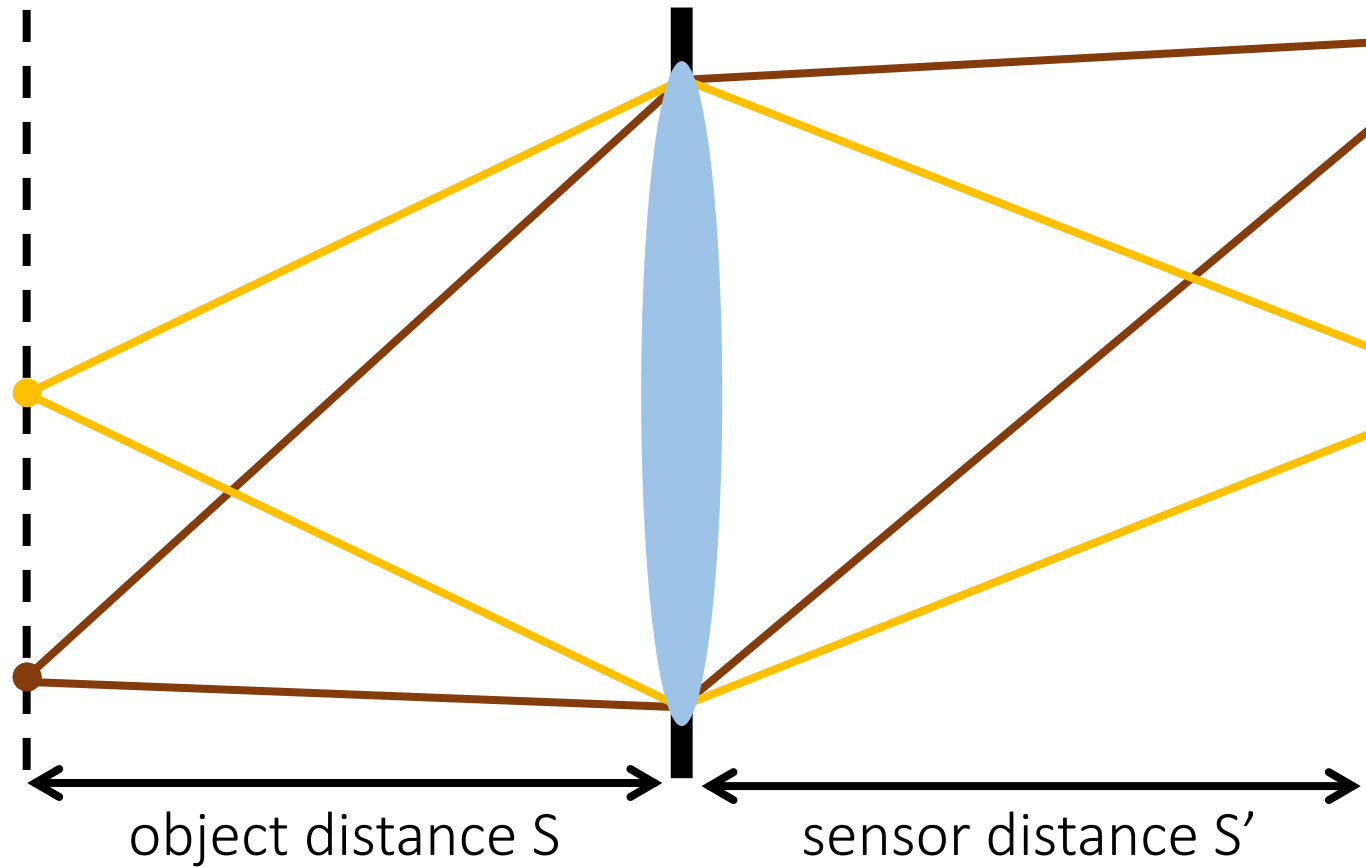
diffraction-limited
PSF of a circular
aperture
(*Airy pattern*)

Lens as an optical low-pass filter

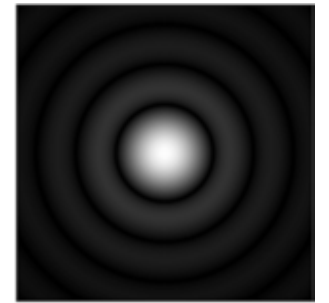
Point spread function (PSF): The blur kernel of a lens.

- “Diffraction-limited” PSF: No aberrations, only diffraction. Determined by aperture shape.

$$\frac{1}{S'} + \frac{1}{S} = \frac{1}{f}$$



blur kernel



diffraction-limited
PSF of a circular
aperture
(*Airy pattern*)

Some basics of diffraction theory

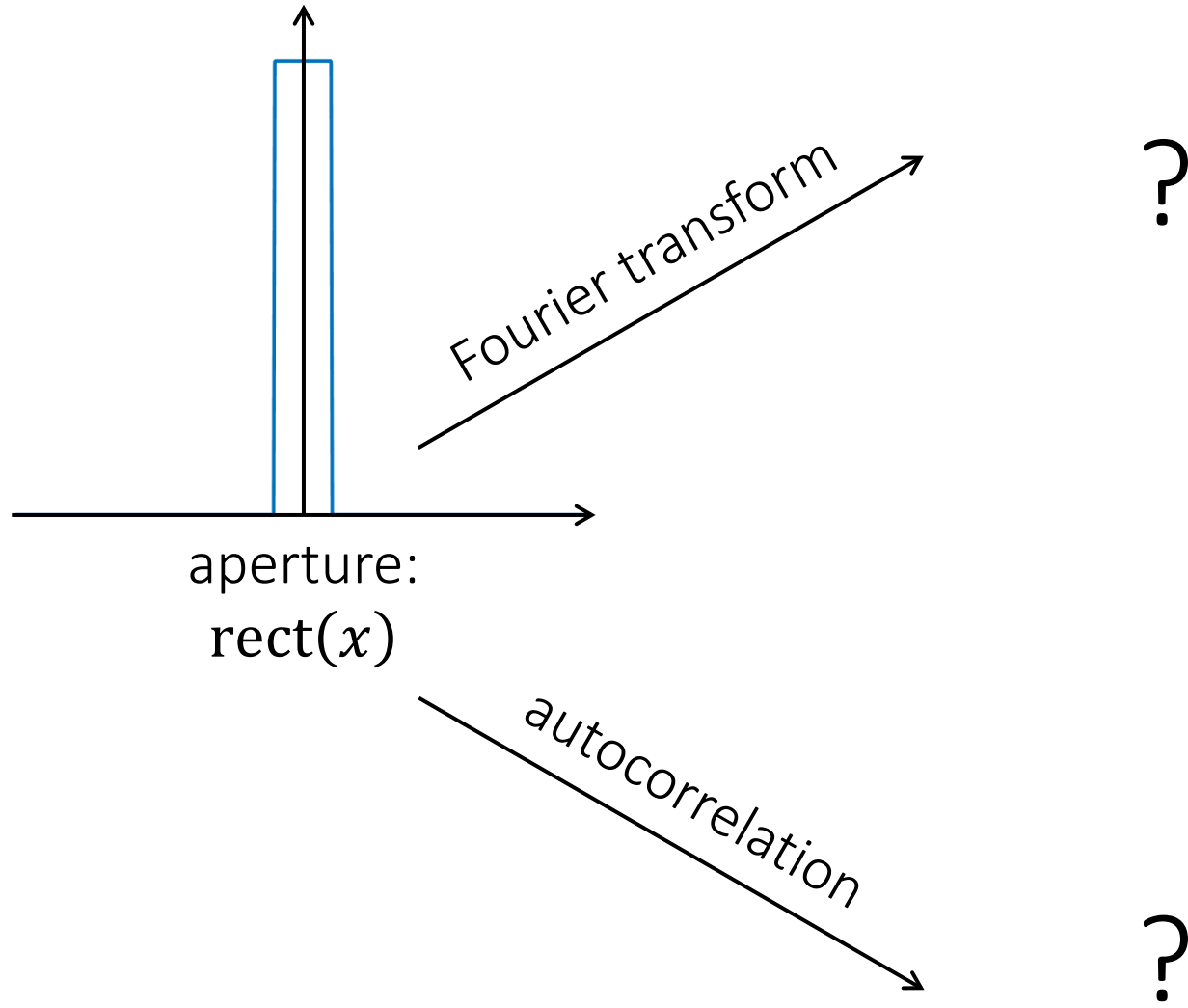
We will assume that we can use:

- *Fraunhofer diffraction* (i.e., distance of sensor and aperture is large relative to wavelength).
- *incoherent illumination* (i.e., the light we are measuring is not laser light).

We will also be ignoring various scale factors. Different functions are not drawn to scale.

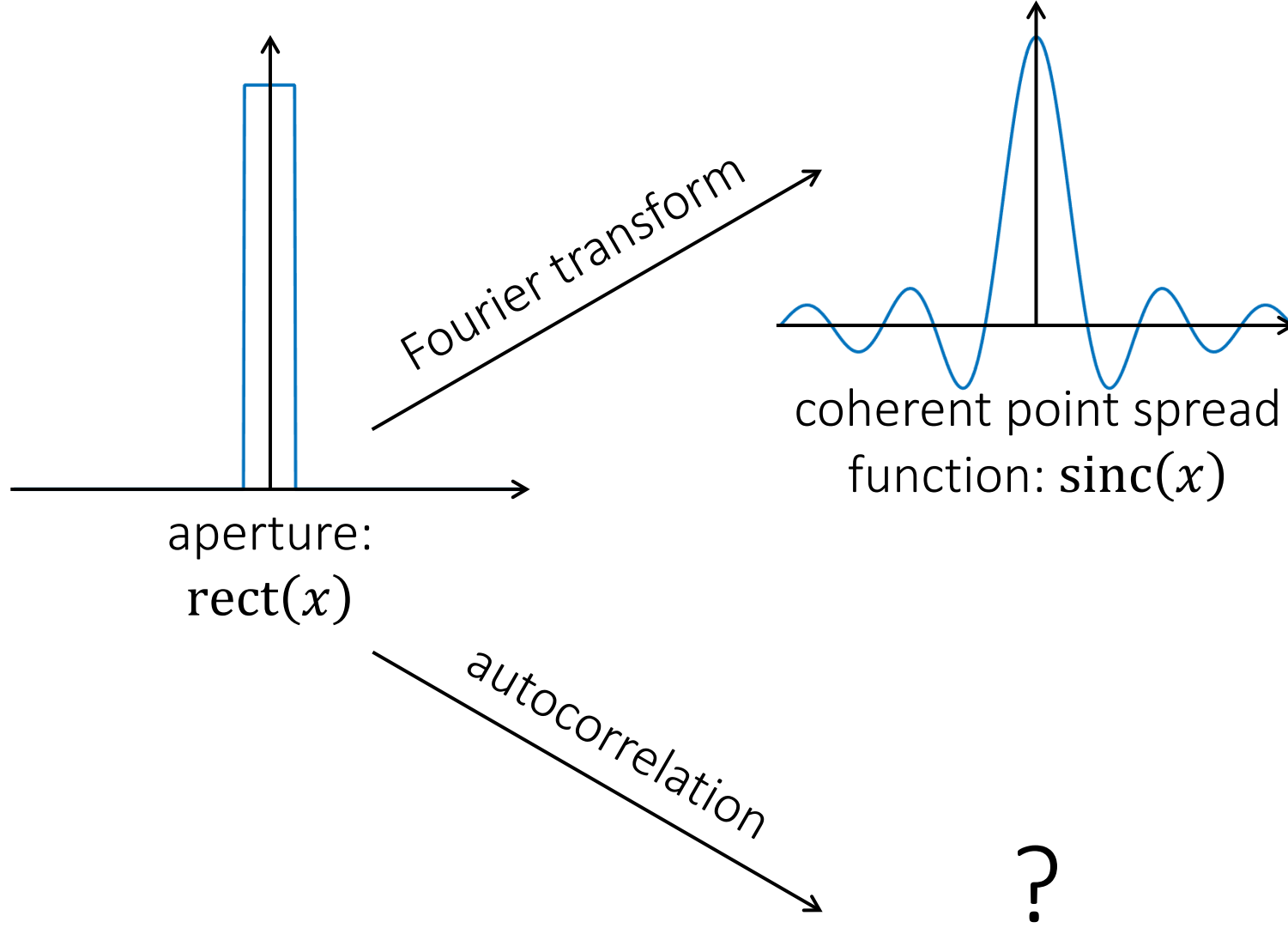
What we discuss here will make more sense when we cover Fourier optics later in this course.

Some basics of diffraction theory



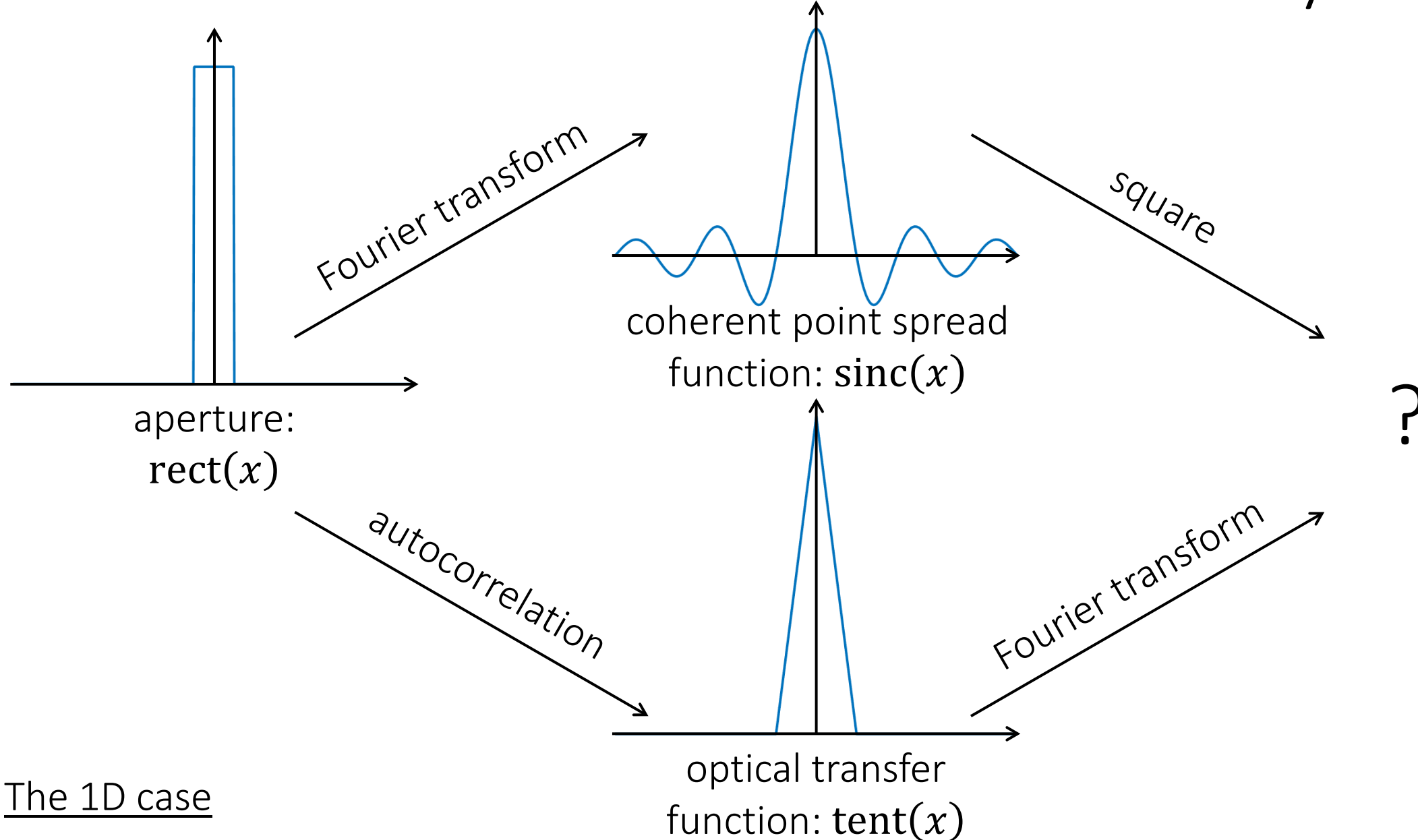
The 1D case

Some basics of diffraction theory



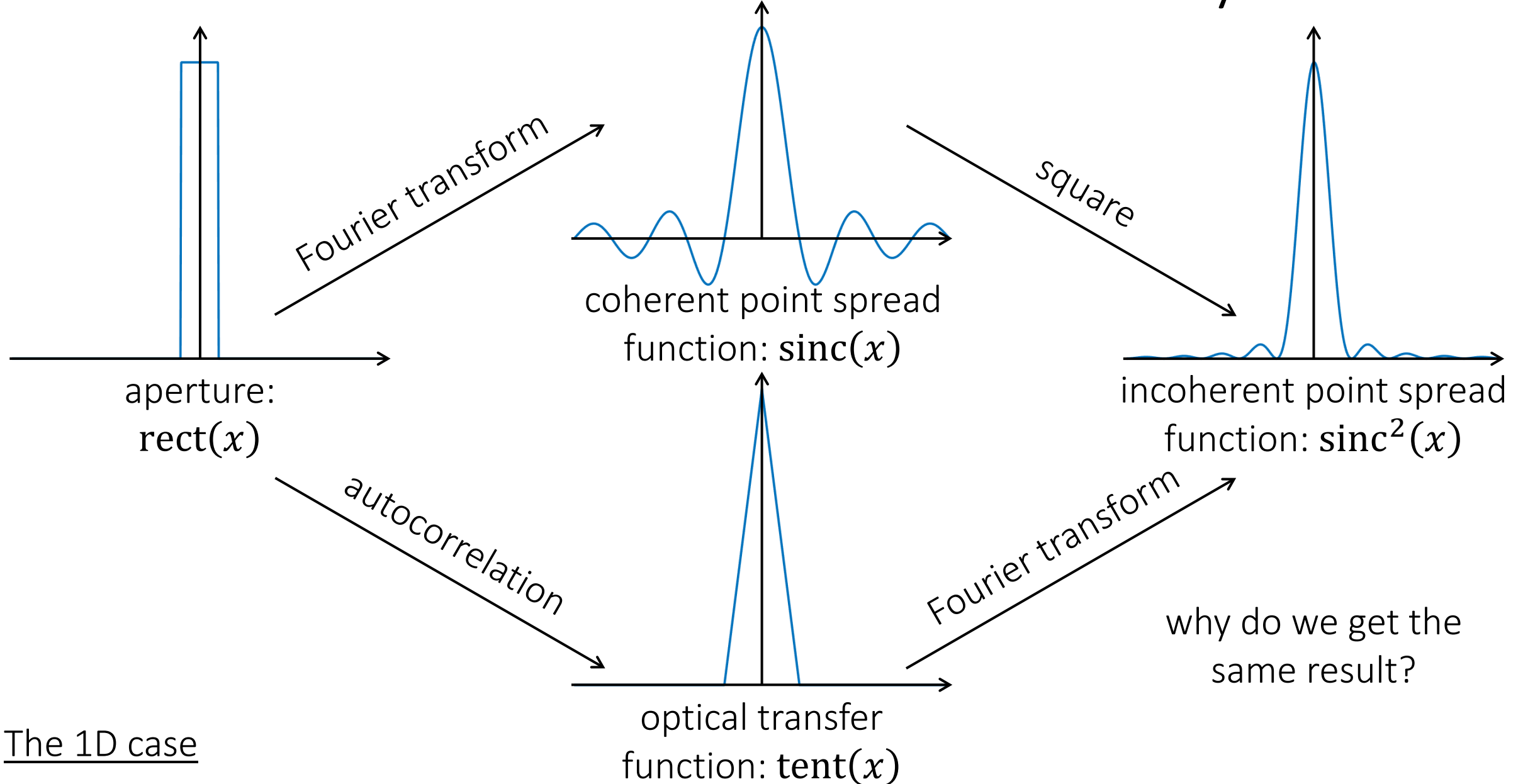
The 1D case

Some basics of diffraction theory

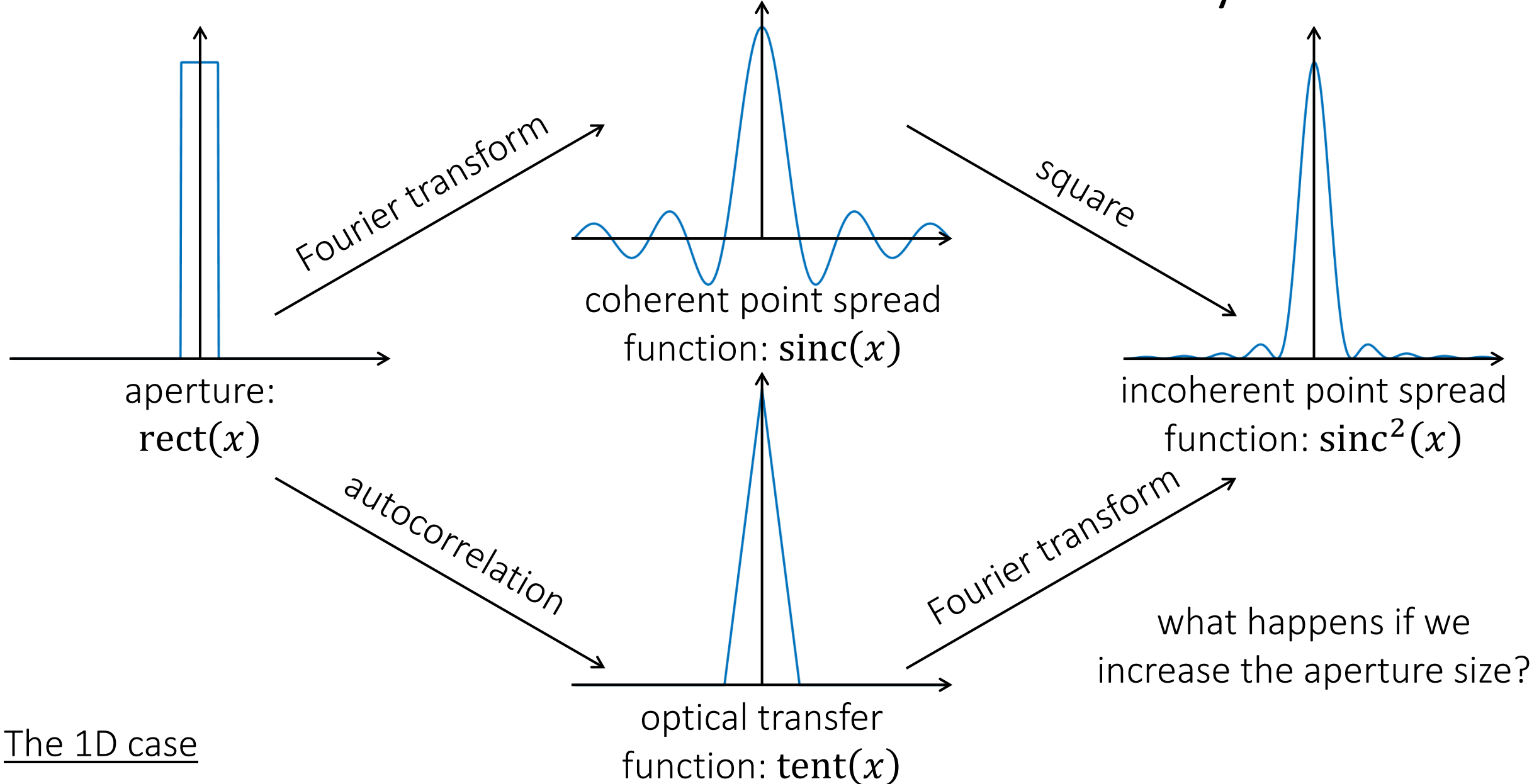


The 1D case

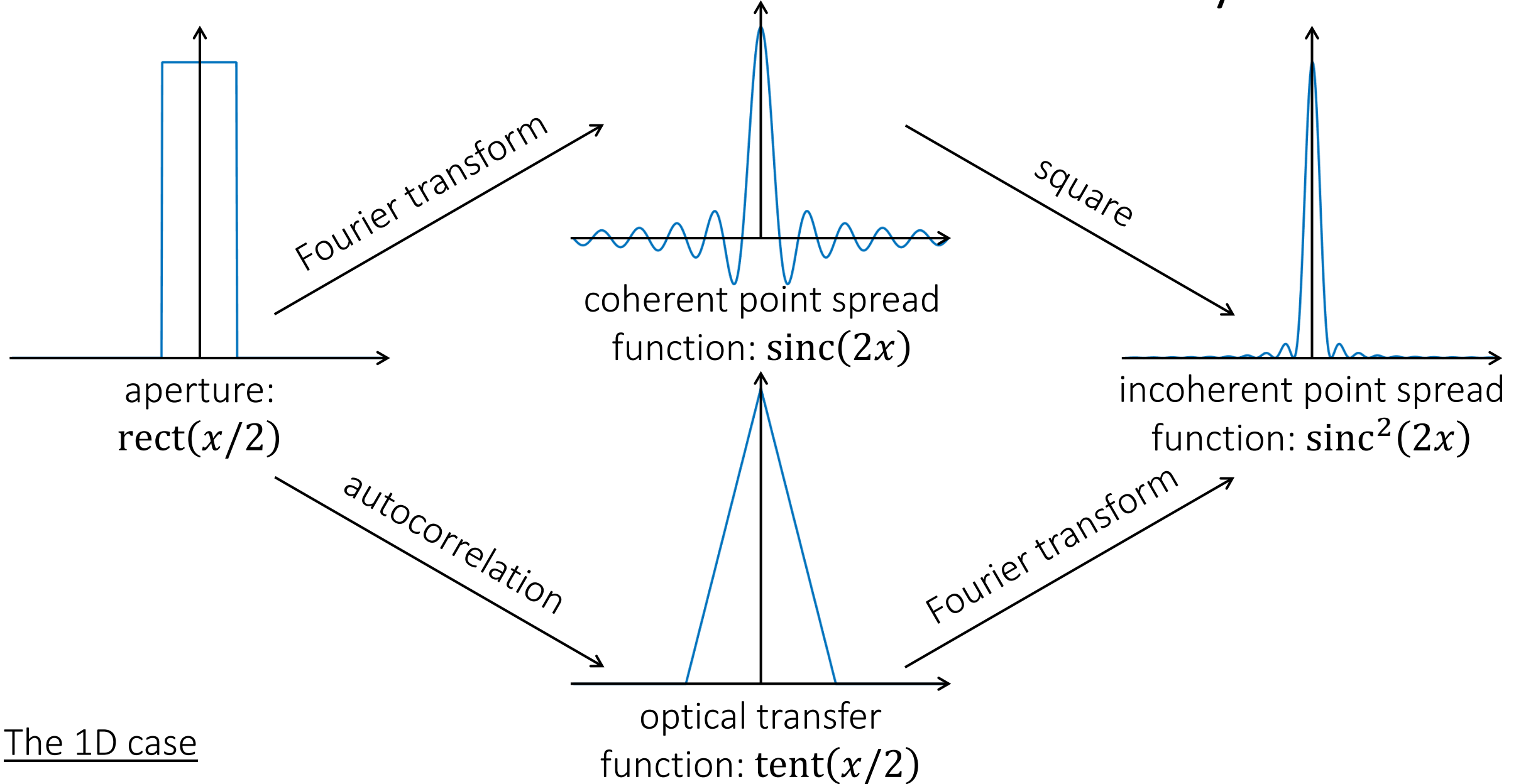
Some basics of diffraction theory



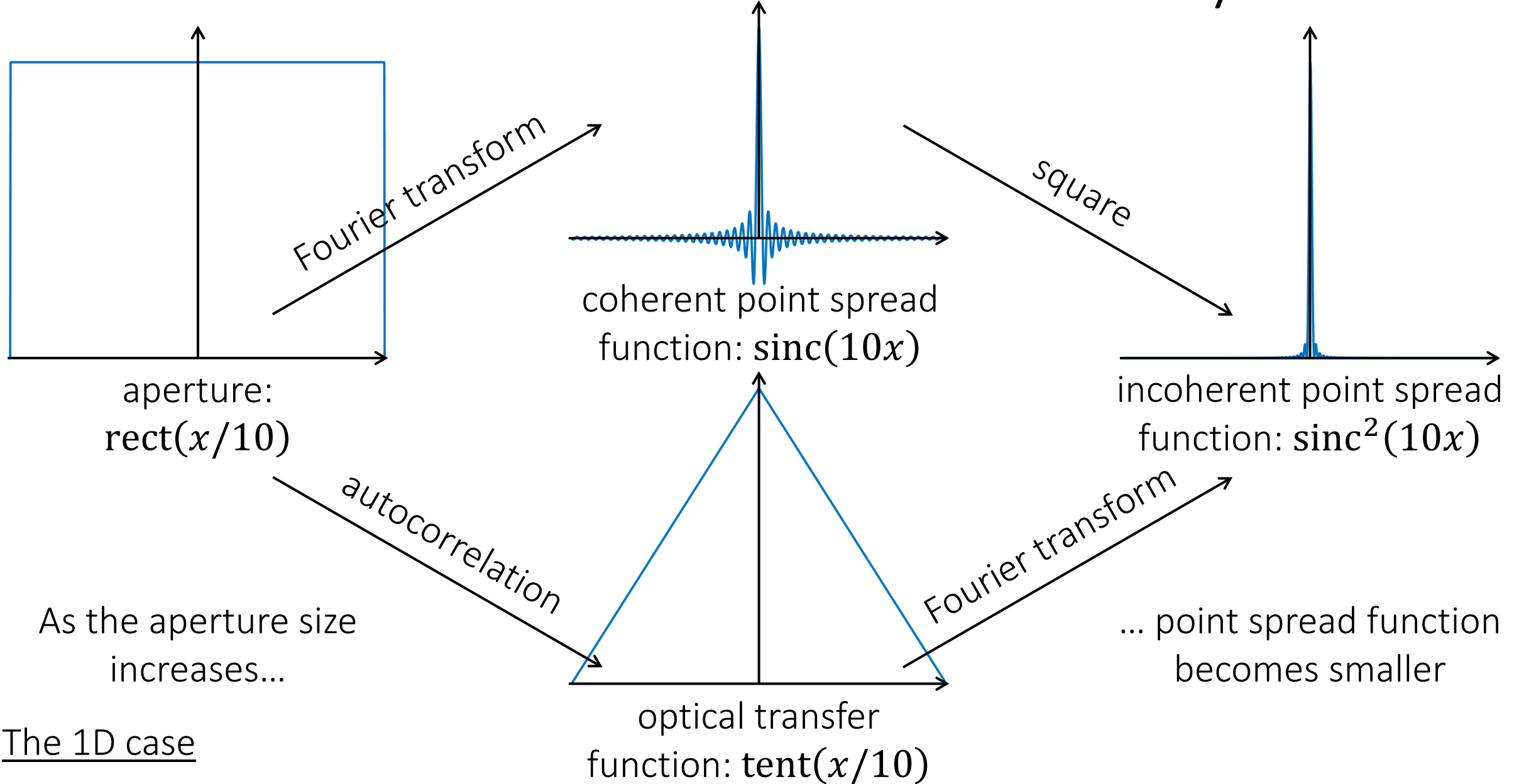
Some basics of diffraction theory



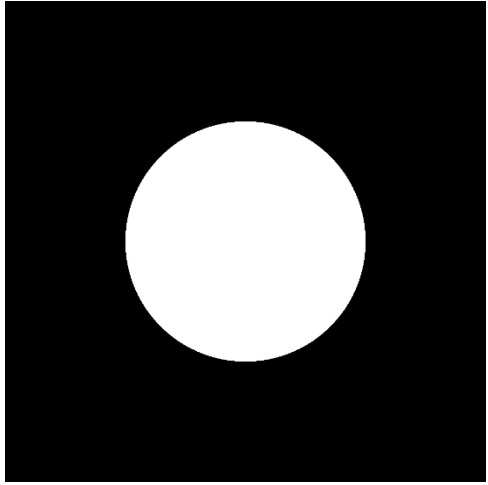
Some basics of diffraction theory



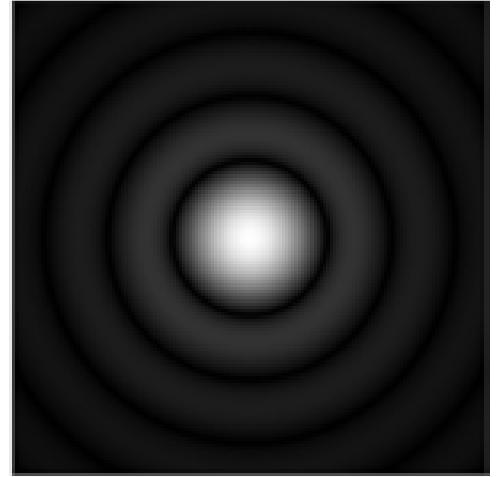
Some basics of diffraction theory



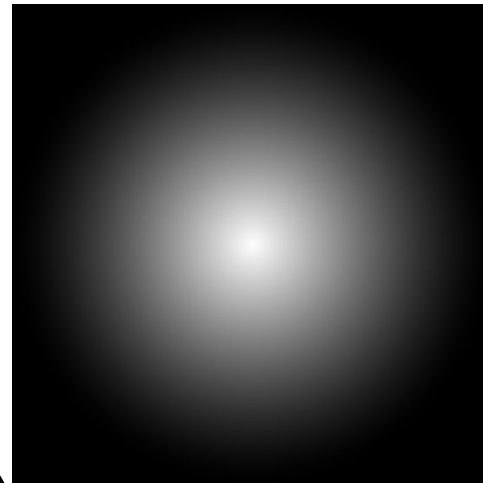
Some basics of diffraction theory



aperture



incoherent point spread
function



optical transfer
function

autocorrelation

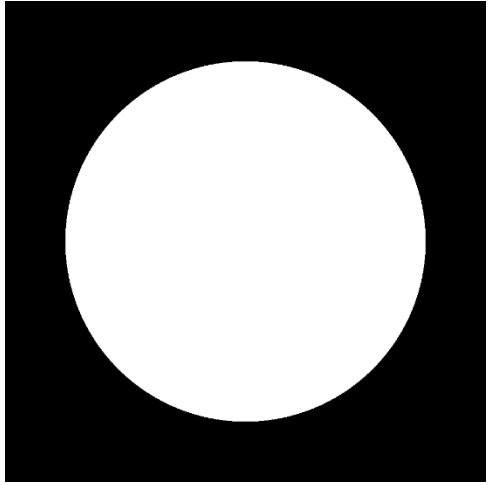
Fourier transform

As the aperture size
increases...

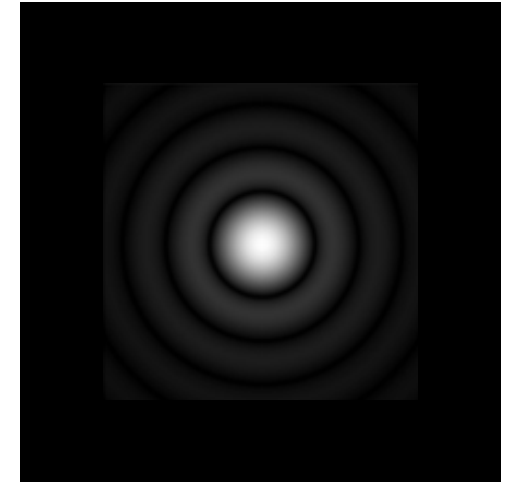
... point spread function
becomes smaller

The 2D case

Some basics of diffraction theory



aperture



incoherent point spread
function



optical transfer
function

autocorrelation

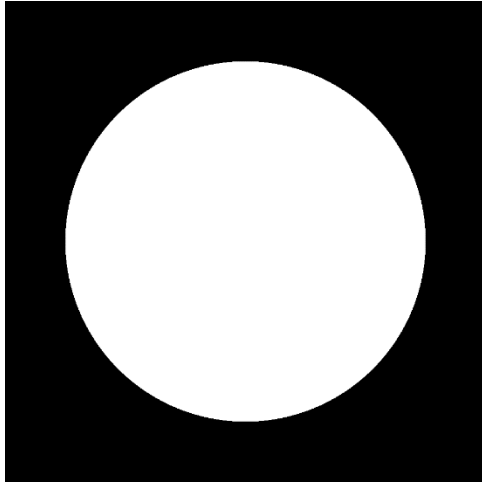
Fourier transform

As the aperture size
increases...

... point spread function
becomes smaller

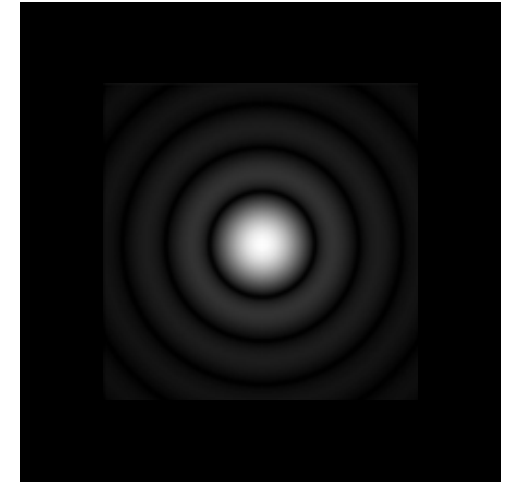
The 2D case

Some basics of diffraction theory



aperture

Why do we prefer circular apertures?

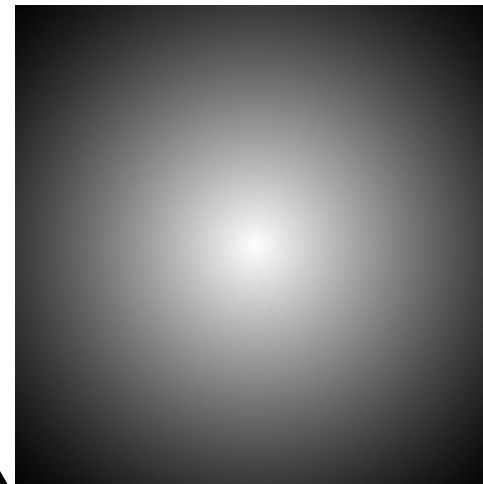


incoherent point spread
function

autocorrelation

As the aperture size
increases...

The 2D case

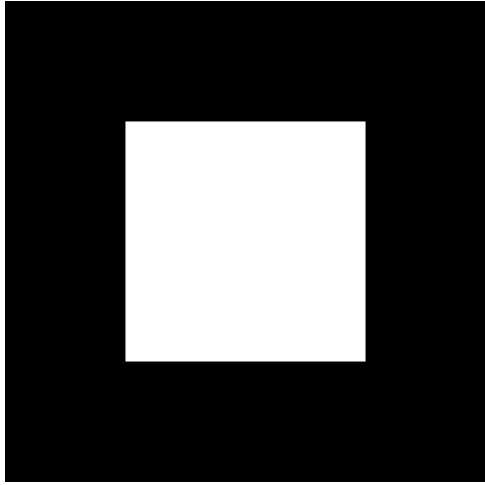


optical transfer
function

Fourier transform

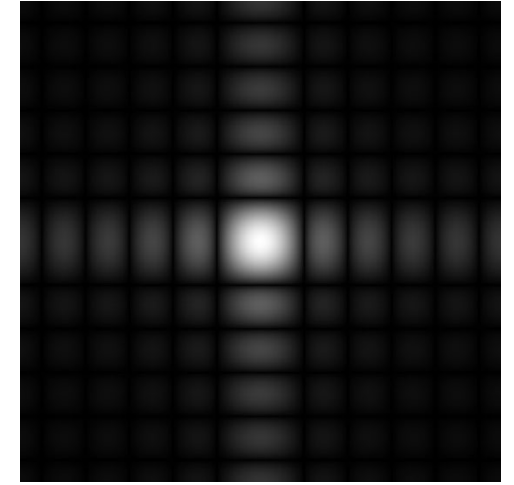
... point spread function
becomes smaller

Some basics of diffraction theory



aperture

Other shapes produce very anisotropic blur.

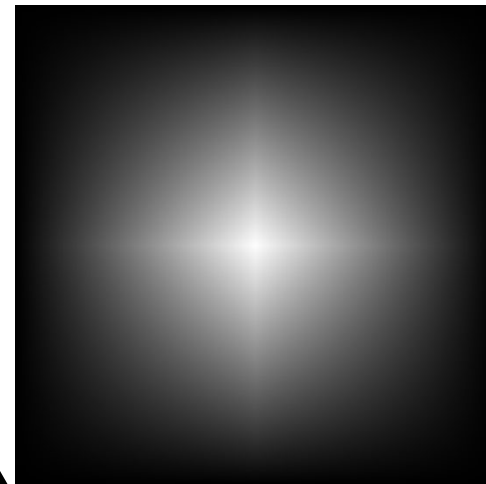


incoherent point spread function

autocorrelation

As the aperture size increases...

The 2D case



optical transfer function

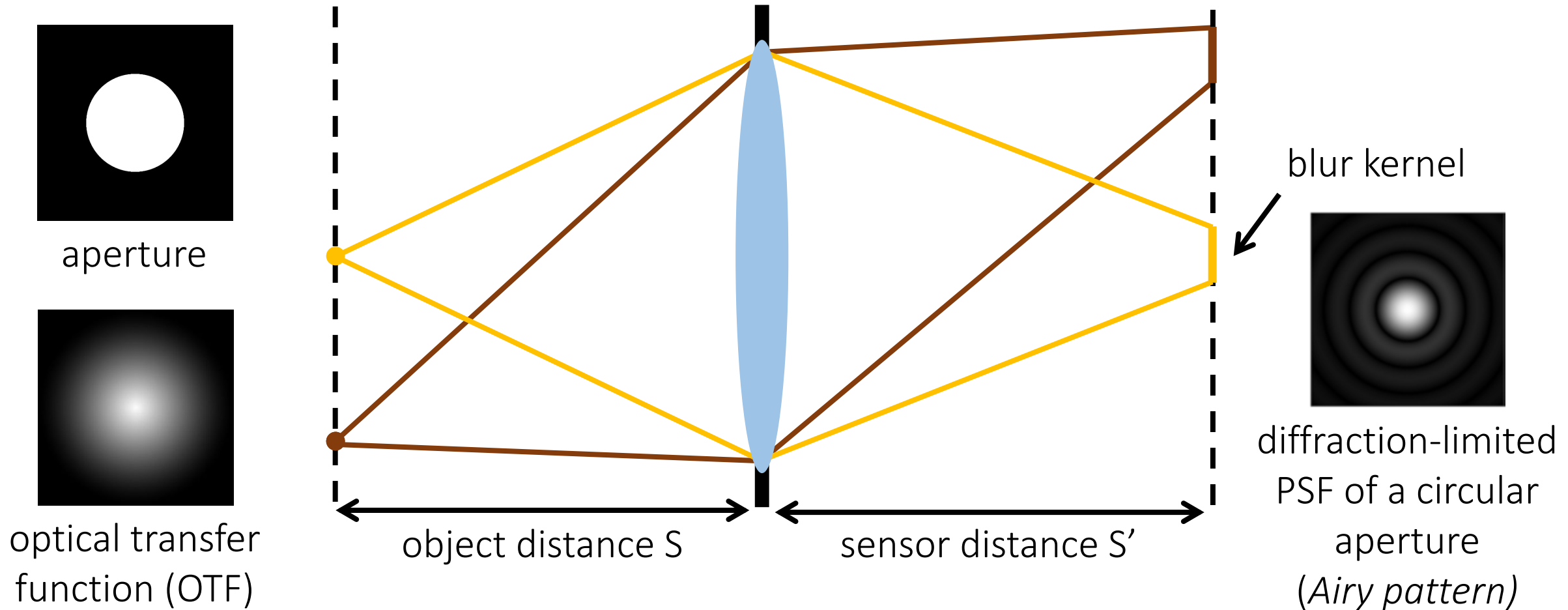
Fourier transform

... point spread function becomes smaller

Lens as an optical low-pass filter

Point spread function (PSF): The blur kernel of a lens.

- “Diffraction-limited” PSF: No aberrations, only diffraction. Determined by aperture shape.



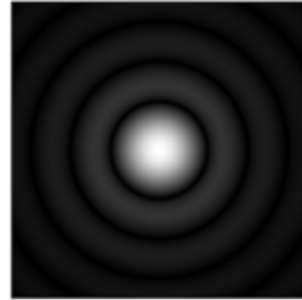
Lens as an optical low-pass filter



image from a perfect lens

i

$*$



$=$



image from imperfect lens

$*$

k

$=$

b

Lens as an optical low-pass filter

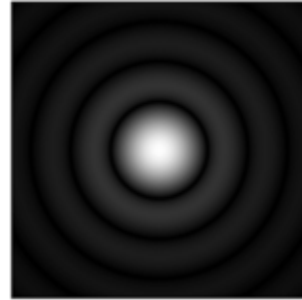
If we know b and k , can we recover i ?



image from a perfect lens

i

$*$



$=$



image from imperfect lens

$*$

k

$=$

b

Deconvolution

$$i * k = b$$

If we know k and b , can we recover i ?

Deconvolution

$$i * k = b$$

Reminder: convolution is multiplication in Fourier domain:

$$F(i) \cdot F(k) = F(b)$$

If we know k and b , can we recover i ?

Deconvolution

$$i * k = b$$

Reminder: convolution is multiplication in Fourier domain:

$$F(i) \cdot F(k) = F(b)$$

Deconvolution is division in Fourier domain:

$$F(i_{\text{est}}) = F(b) \setminus F(k)$$

After division, just do inverse Fourier transform:

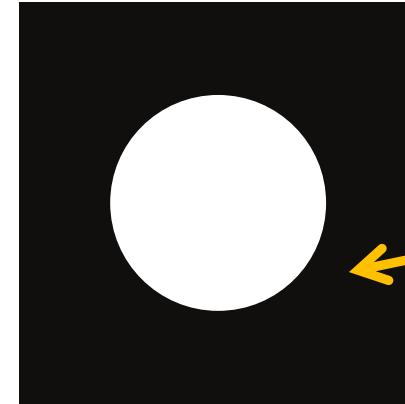
$$i_{\text{est}} = F^{-1} (F(b) \setminus F(k))$$

Deconvolution

Any problems with this approach?

Deconvolution

- The OTF (Fourier of PSF) is a low-pass filter
- The measured signal includes noise



zeros at high
frequencies

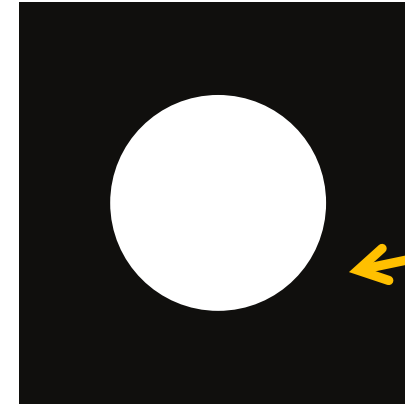
$$b = k * i + n$$



noise term

Deconvolution

- The OTF (Fourier of PSF) is a low-pass filter



zeros at high
frequencies

- The measured signal includes noise

$$b = k * i + n$$



noise term

- When we divide by zero, we amplify the high frequency noise

Naive deconvolution

Even tiny noise can make the results awful.

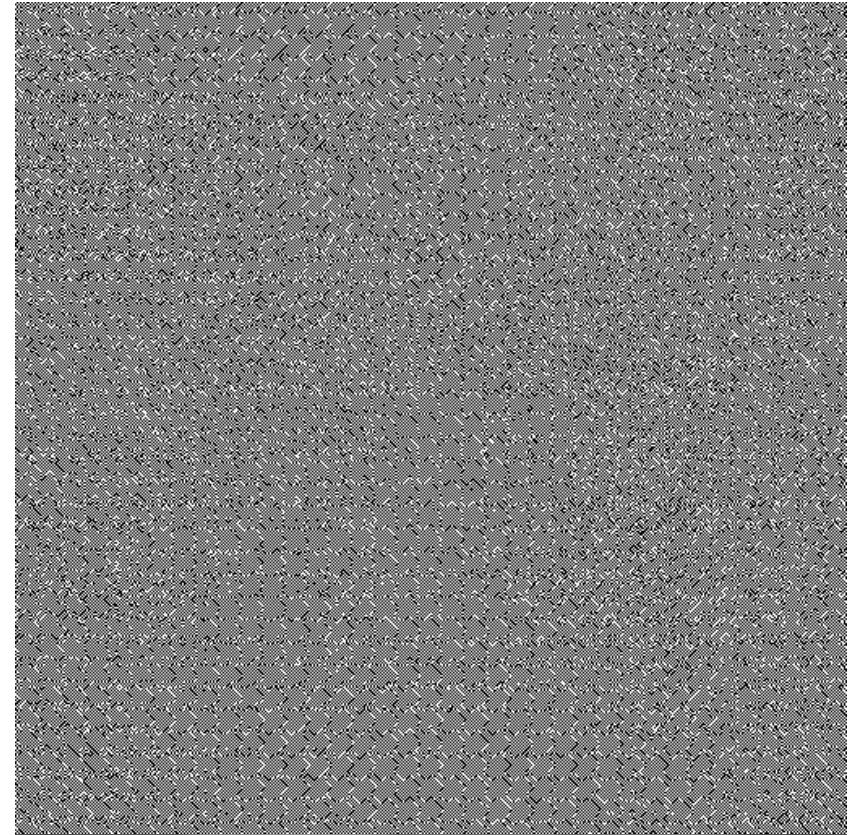
- Example for Gaussian of $\sigma = 0.05$



b

$$* \left[\text{Gaussian kernel} \right]^{-1} =$$

$$* k^{-1} =$$



i_{est}

Wiener Deconvolution

Apply inverse kernel and do not divide by zero:

$$i_{\text{est}} = F^{-1} \left(\frac{|F(k)|^2}{|F(k)|^2 + 1/\text{SNR}(\omega)} \cdot \frac{F(b)}{F(k)} \right)$$

noise-dependent damping factor



- Derived as solution to maximum-likelihood problem under Gaussian noise assumption
- Requires noise of signal-to-noise ratio at each frequency

$$\text{SNR}(\omega) = \frac{\text{signal variance at } \omega}{\text{noise variance at } \omega}$$

Wiener Deconvolution

Apply inverse kernel and do not divide by zero:

$$i_{\text{est}} = F^{-1} \left(\frac{|F(k)|^2}{|F(k)|^2 + 1/\text{SNR}(\omega)} \cdot \frac{F(b)}{F(k)} \right)$$

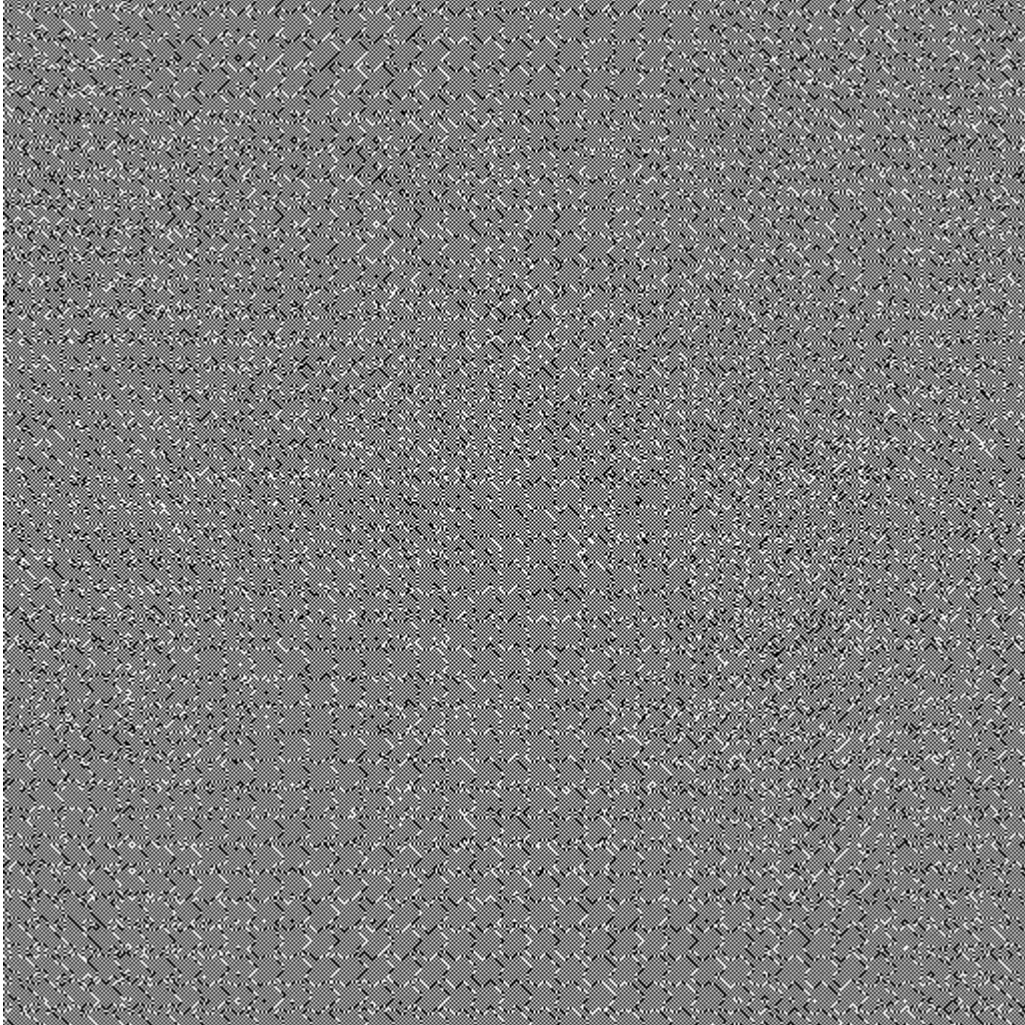
noise-dependent damping factor



Intuitively:

- When SNR is high (low or no noise), just divide by kernel.
- When SNR is low (high noise), just set to zero.

Deconvolution comparisons



naive deconvolution



Wiener deconvolution

Deconvolution comparisons



$\sigma = 0.01$



$\sigma = 0.05$



$\sigma = 0.1$

Derivation

Sensing model:

$$b = k * i + n$$

Noise n is assumed to be zero-mean and independent of signal i .

Derivation

Sensing model:

$$b = k * i + n$$

Noise n is assumed to be zero-mean and independent of signal i .

Fourier transform:

$$B = K \cdot I + N$$



Why multiplication?

Derivation

Sensing model:

$$b = k * i + n$$

Noise n is assumed to be zero-mean and independent of signal i .

Fourier transform:

$$B = K \cdot I + N$$

Convolution becomes multiplication.

Problem statement: Find function $H(\omega)$ that minimizes *expected error in Fourier domain*.

$$\min_H E[\|I - HB\|^2]$$

Derivation

Replace B and re-arrange loss:

$$\min_H E[\|(1 + HK)I - HN\|^2]$$

Expand the squares:

$$\min_H \|1 - HK\|^2 E[\|I\|^2] - 2(1 - HK)E[IN] + \|H\|^2 E[\|N\|^2]$$

Derivation

When handling the cross terms:

- Can I write the following?

$$E[IN] = E[I]E[N]$$

Derivation

When handling the cross terms:

- Can I write the following?

$$E[IN] = E[I]E[N]$$

Yes, because I and N are assumed independent.

- What is this expectation product equal to?

Derivation

When handling the cross terms:

- Can I write the following?

$$E[IN] = E[I]E[N]$$

Yes, because I and N are assumed independent.

- What is this expectation product equal to?

Zero, because N has zero mean.

Derivation

Replace B and re-arrange loss:

$$\min_H E[\|(1 + HK)I - HN\|^2]$$

Expand the squares:

$$\min_H \|1 - HK\|^2 E[\|I\|^2] - 2(1 - HK)E[IN] + \|H\|^2 E[\|N\|^2]$$

 cross-term is zero

Simplify:

$$\min_H \|1 - HK\|^2 E[\|I\|^2] + \|H\|^2 E[\|N\|^2]$$

How do we solve this optimization problem?

Derivation

Differentiate loss with respect to H , set to zero, and solve for H :

$$\frac{\partial \text{loss}}{\partial H} = 0$$

$$\Rightarrow -2(1 - HK)E[\|I\|^2] + 2HE[\|N\|^2] = 0$$

$$\Rightarrow H = \frac{KE[\|I\|^2]}{K^2E[\|I\|^2] + E[\|N\|^2]}$$

Divide both numerator and denominator with $E[\|I\|^2]$, extract factor $1/K$, and done!

Wiener Deconvolution

Apply inverse kernel and do not divide by zero:

$$i_{est} = F^{-1} \left(\frac{|F(k)|^2}{|F(k)|^2 + 1/SNR(\omega)} \cdot \frac{F(b)}{F(k)} \right)$$

noise-dependent damping factor 

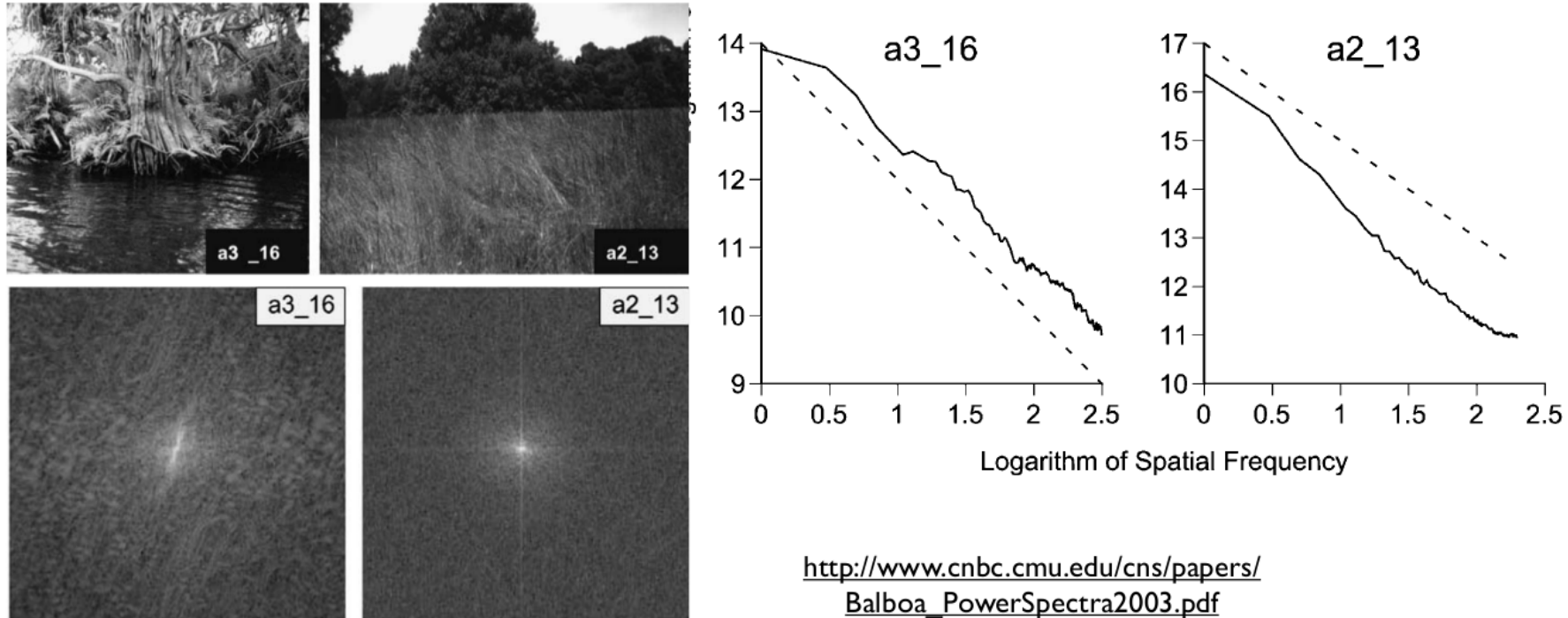
- Derived as solution to maximum-likelihood problem under Gaussian noise assumption
- Requires estimate of signal-to-noise ratio at each frequency

$$SNR(\omega) = \frac{\text{signal variance at } \omega}{\text{noise variance at } \omega}$$

Natural image and noise spectra

Natural images *tend* to have spectrum that scales as $1 / \omega^2$

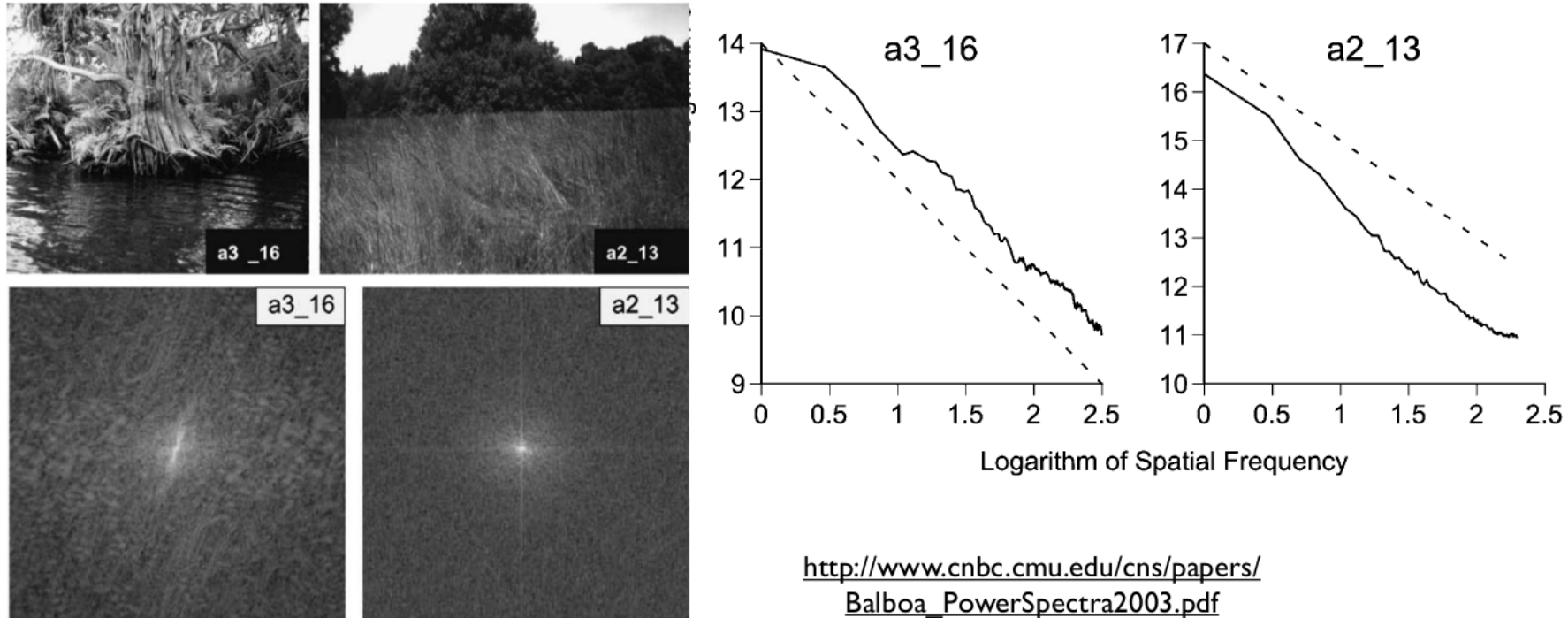
- This is a *natural image statistic*



Natural image and noise spectra

Natural images *tend* to have spectrum that scales as $1 / \omega^2$

- This is a *natural image statistic*



Noise tends to have flat spectrum, $\sigma(\omega) = \text{constant}$

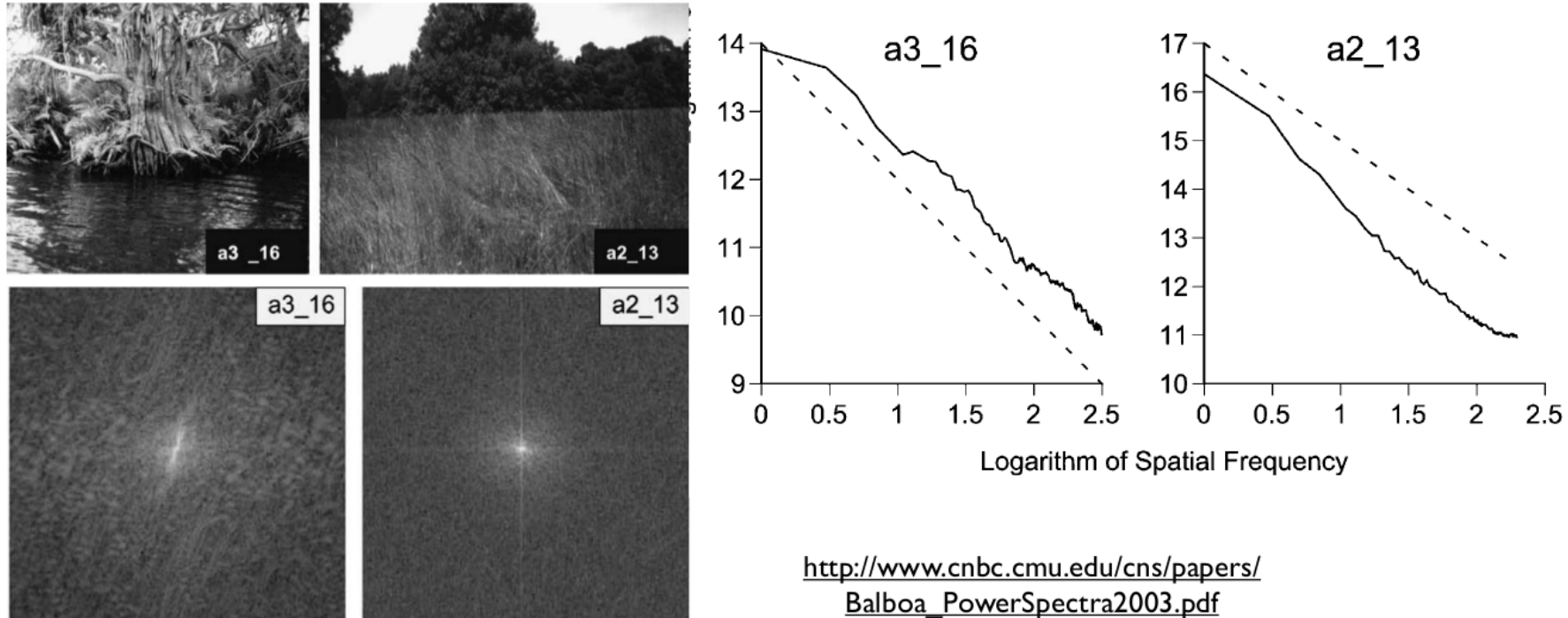
- We call this white noise

What is the SNR?

Natural image and noise spectra

Natural images *tend* to have spectrum that scales as $1 / \omega^2$

- This is a *natural image statistic*



Noise tends to have flat spectrum, $\sigma(\omega) = \text{constant}$

- We call this white noise

Therefore, we have that: $\text{SNR}(\omega) = 1 / \omega^2$

Wiener Deconvolution

Apply inverse kernel and do not divide by zero:

$$i_{\text{est}} = F^{-1} \left(\frac{|F(k)|^2}{|F(k)|^2 + \omega^2} \cdot \frac{F(b)}{F(k)} \right)$$

amplitude-dependent damping factor 

- Derived as solution to maximum-likelihood problem under Gaussian noise assumption
- Requires noise of signal-to-noise ratio at each frequency

$$\text{SNR}(\omega) = \frac{1}{\omega^2}$$

Wiener Deconvolution

For natural images and white noise, equivalent to the minimization problem:

$$\min_i ||b - k * i||^2 + ||\nabla i||^2$$

gradient *regularization*



How can you prove this equivalence?

Wiener Deconvolution

For natural images and white noise, it can be re-written as the minimization problem

$$\min_i ||b - k * i||^2 + ||\nabla i||^2$$

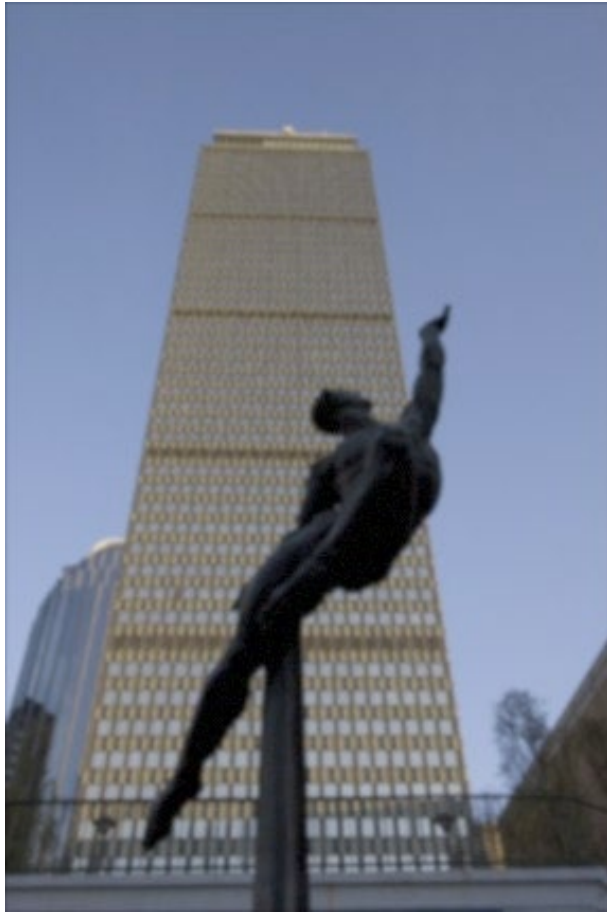
gradient *regularization*



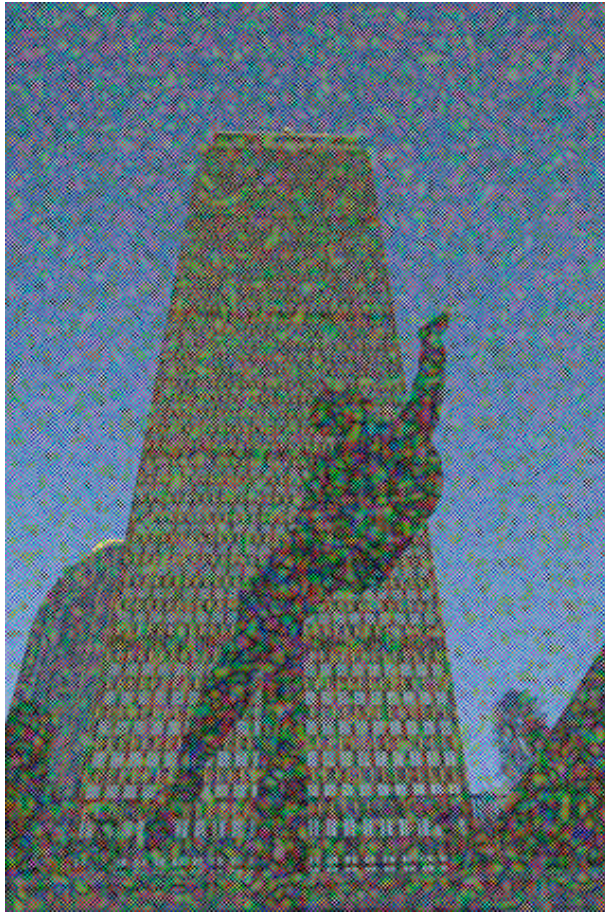
How can you prove this equivalence?

- Convert to Fourier domain and repeat the proof for Wiener deconvolution.
- Intuitively: The ω^2 term in the denominator of the special Wiener filter is the square of the Fourier transform of ∇i , which is $j \cdot \omega$.

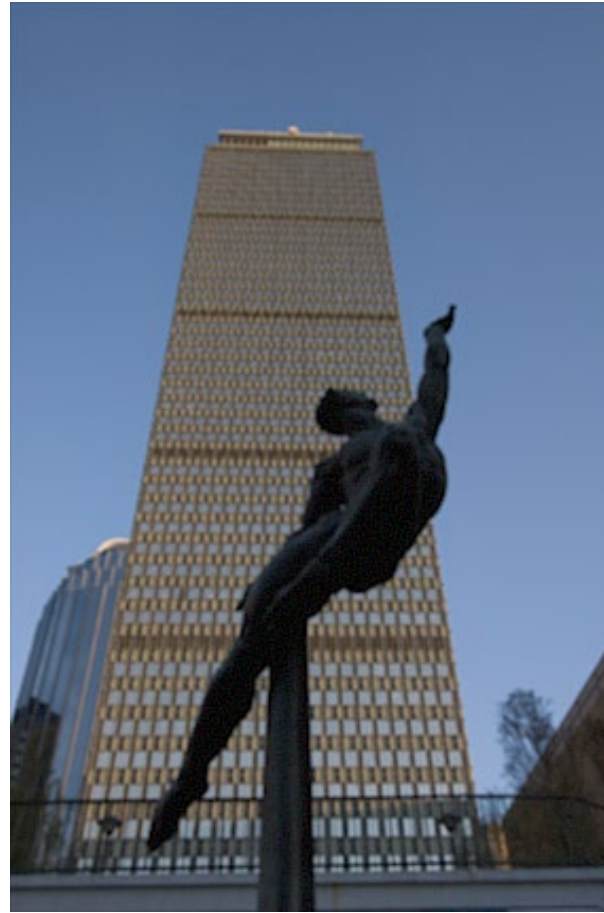
Deconvolution comparisons



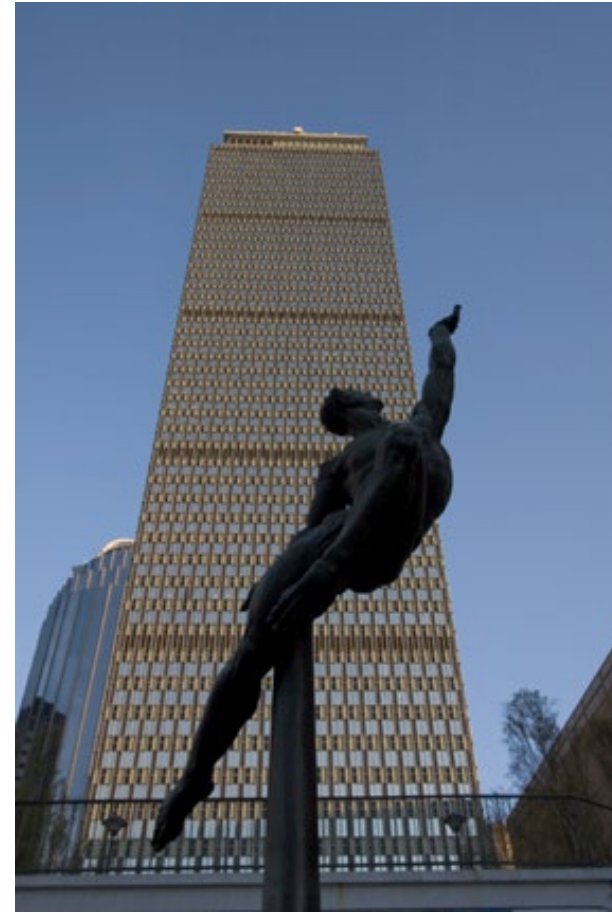
blurry input



naive deconvolution

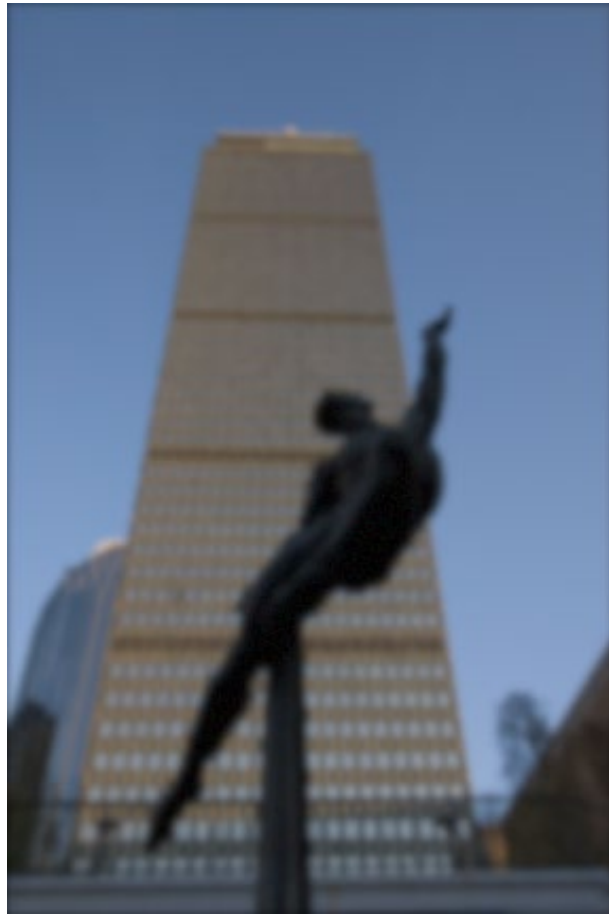


gradient regularization



original

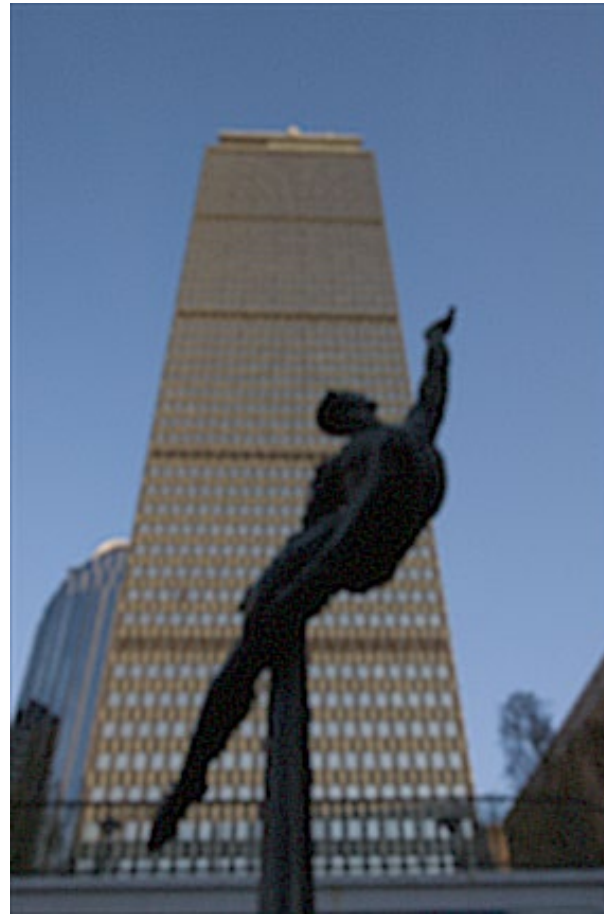
Deconvolution comparisons



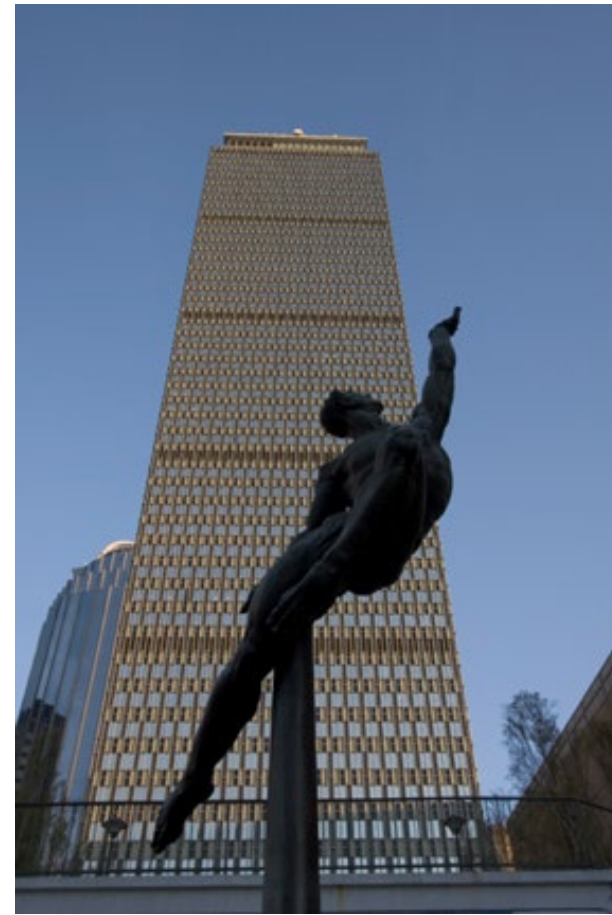
blurry input



naive deconvolution

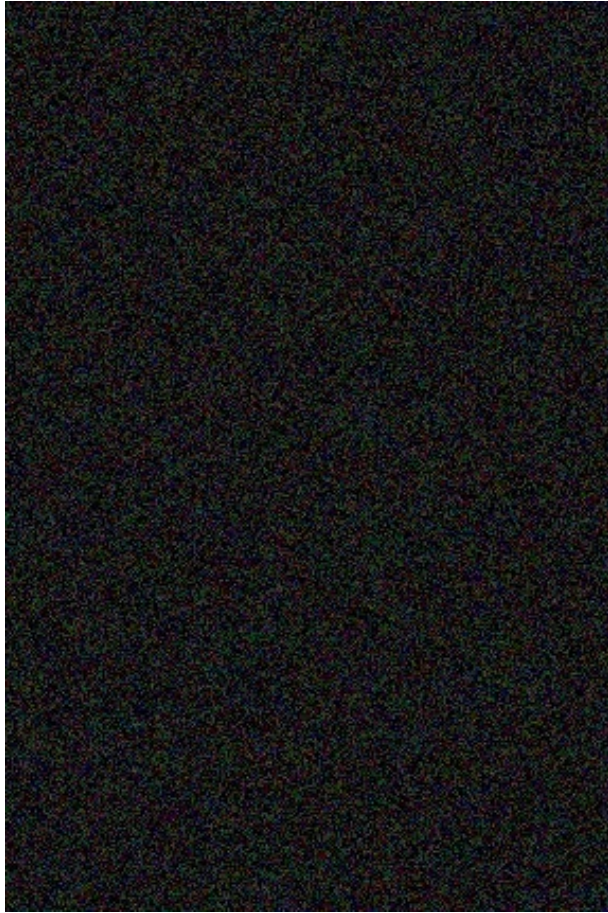


gradient regularization



original

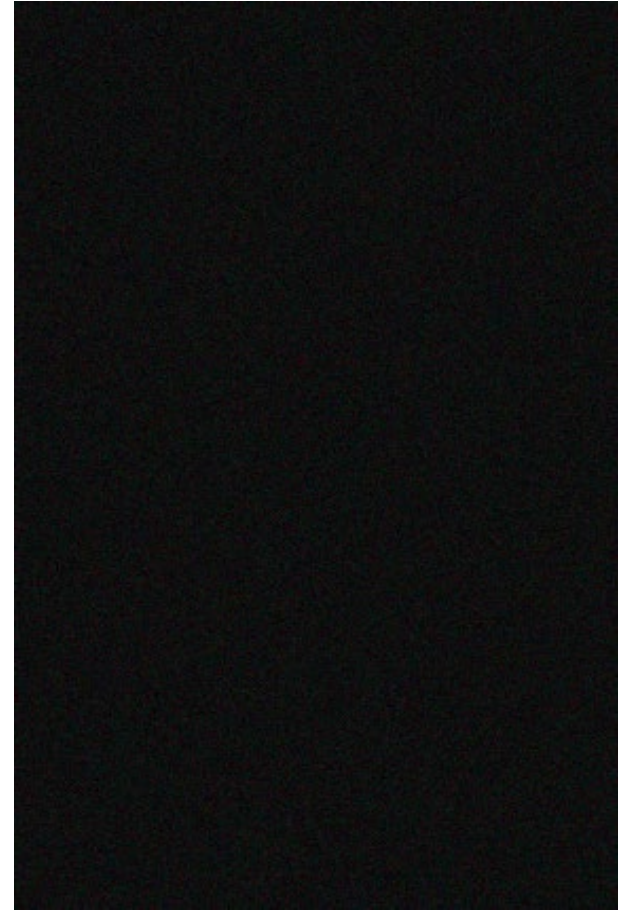
... and a proof-of-concept demonstration



noisy input



naive deconvolution



gradient regularization

Question

Can we undo lens blur by deconvolving a PNG or JPEG image without any preprocessing?

Question

Can we undo lens blur by deconvolving a PNG or JPEG image without any preprocessing?

- All the blur processes we discuss today happen *optically* (before capture by the sensor).
- Blur model is accurate only if our images are *linear*.

Are PNG or JPEG images linear?

Question

Can we undo lens blur by deconvolving a PNG or JPEG image without any preprocessing?

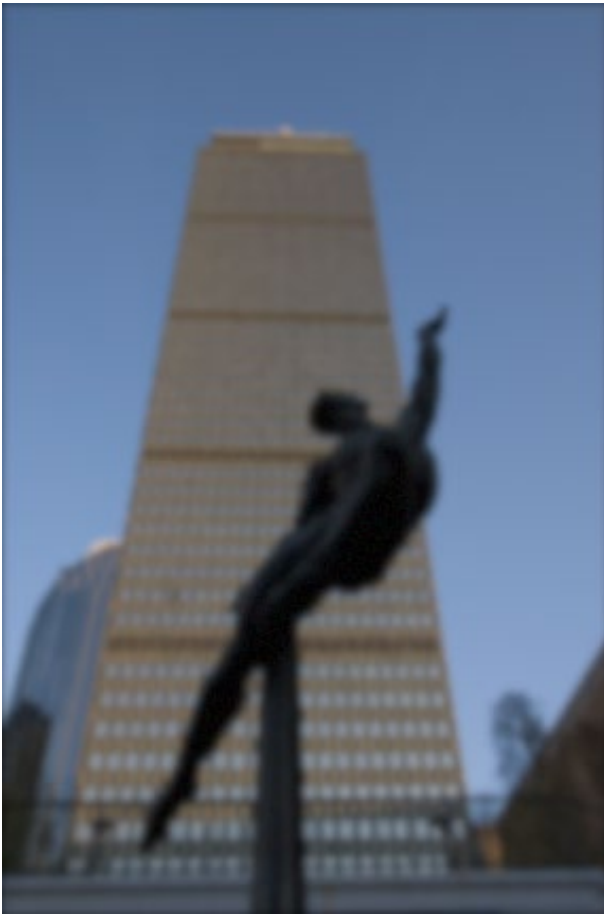
- All the blur processes we discuss today happen *optically* (before capture by the sensor).
- Blur model is accurate only if our images are *linear*.

Are PNG or JPEG images linear?

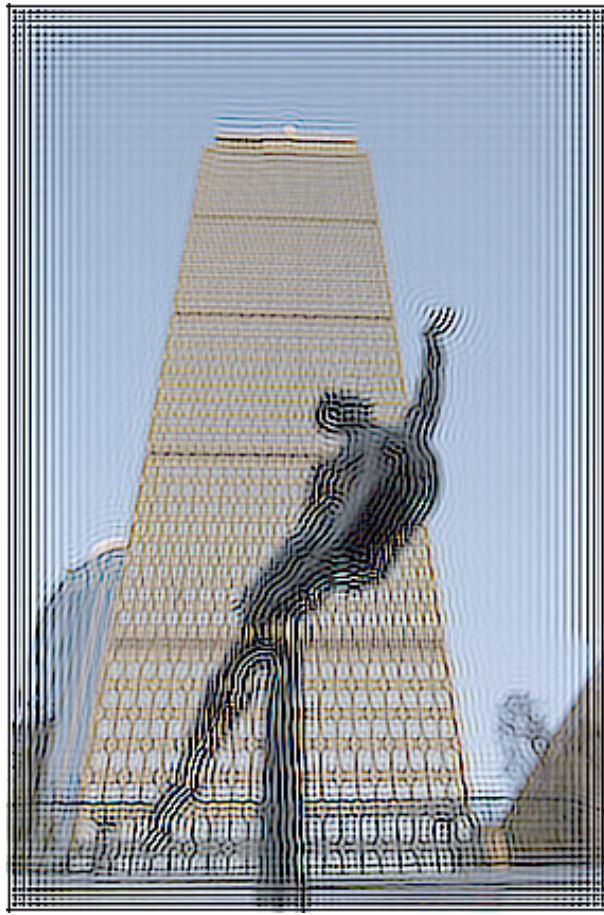
- No, because of gamma encoding.
- Before deblurring, you must linearize your images.

How do we linearize PNG or JPEG images?

The importance of linearity



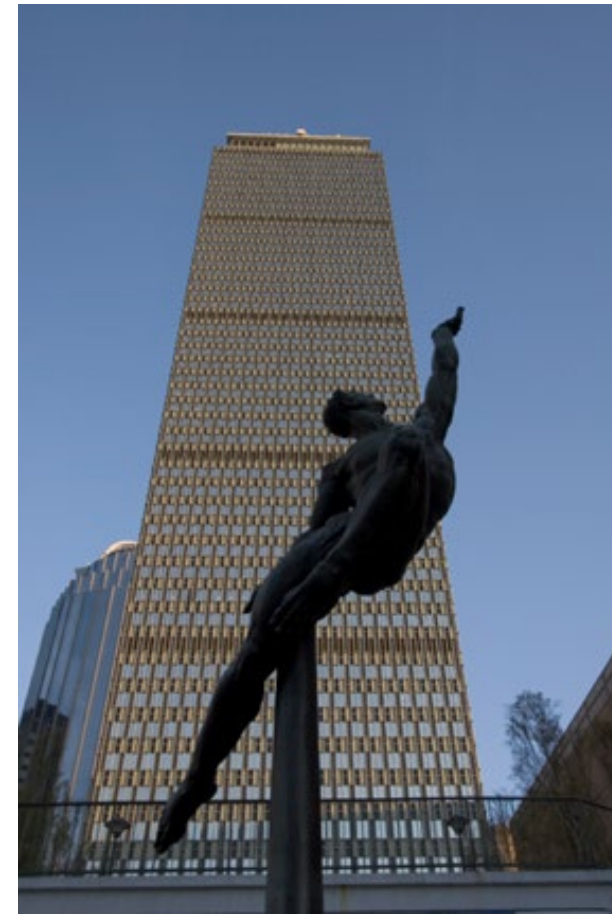
blurry input



deconvolution without
linearization



deconvolution after
linearization



original

Can we do better than that?

Can we do better than that?

Use different gradient regularizations:

- L_2 gradient regularization (Tikhonov regularization, same as Wiener deconvolution)

$$\min_i ||b - k * i||^2 + ||\nabla i||_2^2$$

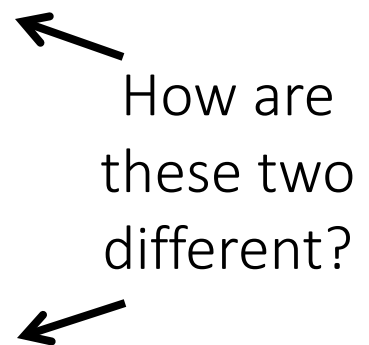
- L_1 gradient regularization (sparsity regularization, *isotropic total variation*)

$$\min_i ||b - k * i||^2 + ||\nabla i||_1^1$$

- *Anisotropic total variation*

$$\min_i ||b - k * i||^2 + ||\nabla i||_2$$

How are
these two
different?



All of these are motivated by natural image statistics. Active research area.

Total Variation

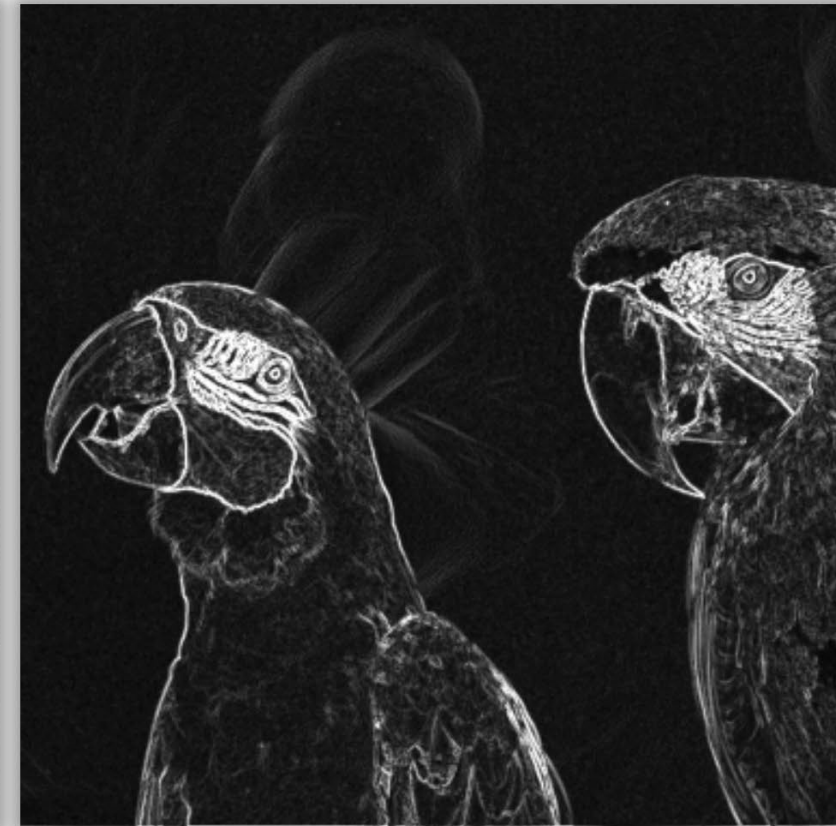
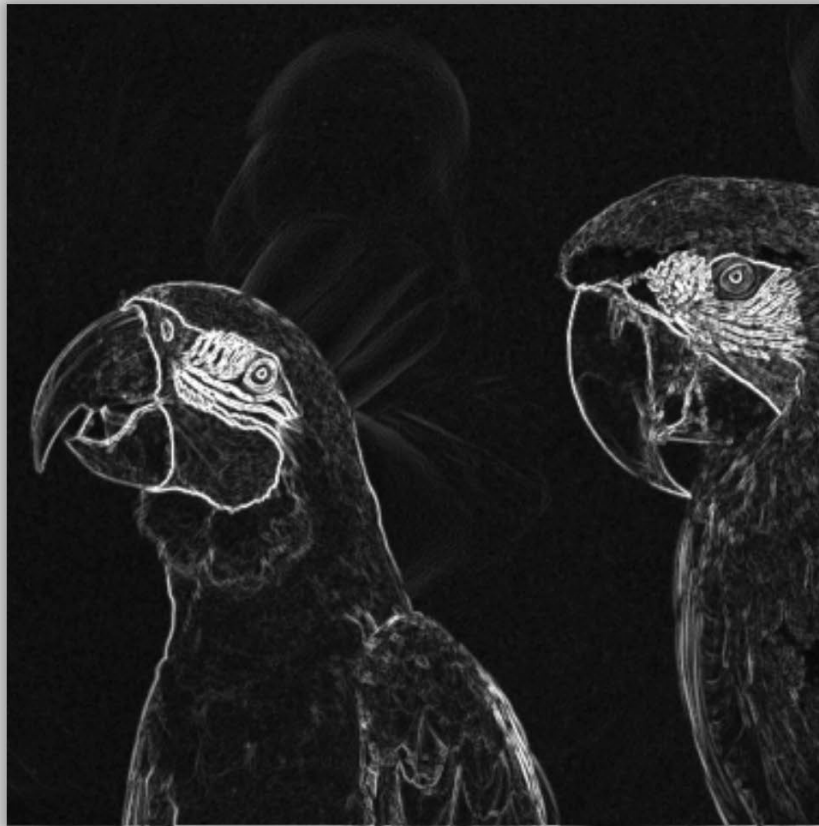
better: isotropic

$$\sqrt{(\nabla_x x)^2 + (\nabla_y x)^2}$$

easier: anisotropic

$$\sqrt{(\nabla_x x)^2} + \sqrt{(\nabla_y x)^2}$$

x



Total Variation

$$\underset{x}{\text{minimize}} \|Cx - b\|_2^2 + \lambda TV(x) = \underset{x}{\text{minimize}} \|Cx - b\|_2^2 + \lambda \|\nabla x\|_1$$

$$\|x\|_1 = \sum_i |x_i|$$

- idea: promote sparse gradients (edges)

- ∇ is finite differences operator, i.e. matrix
$$\begin{bmatrix} -1 & 1 & & & \\ & -1 & 1 & & \\ & & \ddots & \ddots & \\ & & & -1 & 1 \\ & & & & -1 \end{bmatrix}$$

Total Variation

- for simplicity, this lecture only discusses anisotropic TV:

$$TV(x) = \|\nabla_x x\|_1 + \|\nabla_y x\|_1 = \left\| \begin{bmatrix} \nabla_x \\ \nabla_y \end{bmatrix} x \right\|_1$$

- problem: l1-norm is not differentiable, can't use inverse filtering
- however: simple solution for data fitting along and simple solution for TV alone $\rightarrow$ split problem!

Deconvolution with ADMM

- split deconvolution with TV prior:

$$\begin{aligned} & \text{minimize} && \|Cx - b\|_2^2 + \lambda \|z\|_1 \\ & \text{subject to} && \nabla x = z \end{aligned}$$

- general form of ADMM (alternating direction method of multipliers):

$$\begin{aligned} & \text{minimize} && f(x) + g(z) \\ & \text{subject to} && Ax + Bz = c \end{aligned}$$

$$f(x) = \|Cx - b\|_2^2$$

$$g(z) = \lambda \|z\|_1$$

$$A = \nabla, B = -I, c = 0$$

ADMM

minimize $f(x) + g(z)$

subject to $Ax + Bz = c$

- Lagrangian (bring constraints into objective = penalty method):

$$L(x, y, z) = f(x) + g(z) + y^T (Ax + Bz - c)$$



dual variable or Lagrange multiplier

ADMM

minimize $f(x) + g(z)$

subject to $Ax + Bz = c$

- augmented Lagrangian is differentiable under mild conditions (usually better convergence etc.)

$$L_{\rho}(x, y, z) = f(x) + g(z) + y^T (Ax + Bz - c) + (\rho / 2) \|Ax + Bz - c\|_2^2$$

$$\begin{array}{ll}\text{minimize} & f(x) + g(z) \\ \text{subject to} & Ax + Bz = c\end{array} \quad \text{ADMM}$$

- ADMM consists of 3 steps per iteration k :

$$x^{k+1} := \arg \min_x L_\rho(x, z^k, y^k)$$

$$z^{k+1} := \arg \min_z L_\rho(x^{k+1}, z, y^k)$$

$$y^{k+1} := y^k + \rho(Ax^{k+1} + Bz^{k+1} - c)$$

ADMM

minimize $f(x) + g(z)$

subject to $Ax + Bz = c$

- ADMM consists of 3 steps per iteration k:

constant



$$x^{k+1} := \arg \min_x \left(f(x) + (\rho / 2) \left\| Ax + \boxed{Bz^k - c + u^k} \right\|^2 \right)$$

$$z^{k+1} := \arg \min_z \left(g(z) + (\rho / 2) \left\| \boxed{Ax^{k+1}} + Bz - c + u^k \right\|^2 \right)$$

$$u^{k+1} := u^k + Ax^{k+1} + Bz^{k+1} - c$$

scaled dual variable: $u = (1 / \rho)y$

ADMM

minimize $f(x) + g(z)$

subject to $Ax + Bz = c$

- ADMM consists of 3 steps per iteration k:

split $f(x)$ and $g(x)$ into independent problems!

$$x^{k+1} := \arg \min_x \left(f(x) + (\rho / 2) \left\| Ax + Bz^k - c + u^k \right\|_2^2 \right) \quad \downarrow \text{(u connects them)}$$

$$z^{k+1} := \arg \min_z \left(g(z) + (\rho / 2) \left\| Ax^{k+1} + Bz - c + u^k \right\|_2^2 \right)$$

$$u^{k+1} := u^k + Ax^{k+1} + Bz^{k+1} - c$$

scaled dual variable: $u = (1 / \rho)y$

Deconvolution with ADMM

minimize $\frac{1}{2} \|Cx - b\|_2^2 + \lambda \|z\|_1$

subject to $\nabla x - z = 0$

- ADMM consists of 3 steps per iteration k:

$$x^{k+1} \quad \coloneqq \quad \arg \min_x \left(\frac{1}{2} \|Cx - b\|_2^2 + (\rho / 2) \|\nabla x - z^k + u^k\|_2^2 \right)$$

$$z^{k+1} \quad \coloneqq \quad \arg \min_z \left(\lambda \|z\|_1 + (\rho / 2) \|\nabla x^{k+1} - z + u^k\|_2^2 \right)$$

$$u^{k+1} \quad \coloneqq \quad u^k + \nabla x^{k+1} - z^{k+1}$$

Deconvolution with ADMM

minimize $\frac{1}{2} \|Cx - b\|_2^2 + \lambda \|z\|_1$

subject to $\nabla x - z = 0$

constant, say $v = z^k - u^k$

1. x-update:

$$x^{k+1} := \arg \min_x \left(\frac{1}{2} \|Cx - b\|_2^2 + (\rho / 2) \|\nabla x - z^k + u^k\|_2^2 \right)$$

solve normal equations $(C^T C + \rho \nabla^T \nabla) x = (C^T b + \rho \nabla^T v)$

$$\nabla^T v = \begin{bmatrix} \nabla_x \\ \nabla_y \end{bmatrix}^T \quad v = \nabla_x^T v_1 + \nabla_y^T v_2$$

minimize $\frac{1}{2} \|Cx - b\|_2^2 + \lambda \|z\|_1$ Deconvolution with ADMM

subject to $\nabla x - z = 0$ constant, say $v = z^k - u^k$

1. x-update: $x^{k+1} := \arg \min_x \left(\frac{1}{2} \|Cx - b\|_2^2 + (\rho / 2) \|\nabla x - \boxed{z^k + u^k}\|_2^2 \right)$

$$x = (C^T C + \rho \nabla^T \nabla)^{-1} (C^T b + \rho \nabla^T v)$$

• inverse filtering: $x^{k+1} = F^{-1} \left\{ \frac{\boxed{F\{c\}^* \cdot F\{b\}} + \rho \left(\boxed{F\{\nabla_x\}^* \cdot F\{v_1\}} + \boxed{F\{\nabla_y\}^* \cdot F\{v_2\}} \right)}{\boxed{F\{c\}^* \cdot F\{c\} + \rho \left(F\{\nabla_x\}^* \cdot F\{\nabla_x\} + F\{\nabla_y\}^* \cdot F\{\nabla_y\} \right)}} \right\}$

precompute!

minimize $\frac{1}{2} \|Cx - b\|_2^2 + \lambda \|z\|_1$ Deconvolution with ADMM

subject to $\nabla x - z = 0$ constant, say $a = \nabla x^{k+1} + u^k$

2. z-update: $z^{k+1} := \arg \min_z \left(\lambda \|z\|_1 + (\rho / 2) \|\nabla x^{k+1} - z + u^k\|_2^2 \right)$

Deconvolution with ADMM

minimize $\frac{1}{2} \|Cx - b\|_2^2 + \lambda \|z\|_1$

subject to $\nabla x - z = 0$

for k=1:max_iters

$$x^{k+1} := \arg \min_x \left(\frac{1}{2} \left\| \begin{bmatrix} C \\ \rho \nabla \end{bmatrix} x - \begin{bmatrix} b \\ \rho v \end{bmatrix} \right\|_2^2 \right) \quad \text{inverse filtering}$$

$$z^{k+1} := S_{\lambda/\rho}(\nabla x^{k+1} + u^k) \quad \text{element-wise threshold}$$

$$u^{k+1} := u^k + \nabla x^{k+1} - z^{k+1} \quad \text{trivial}$$

Deconvolution comparisons



Wiener deconvolution



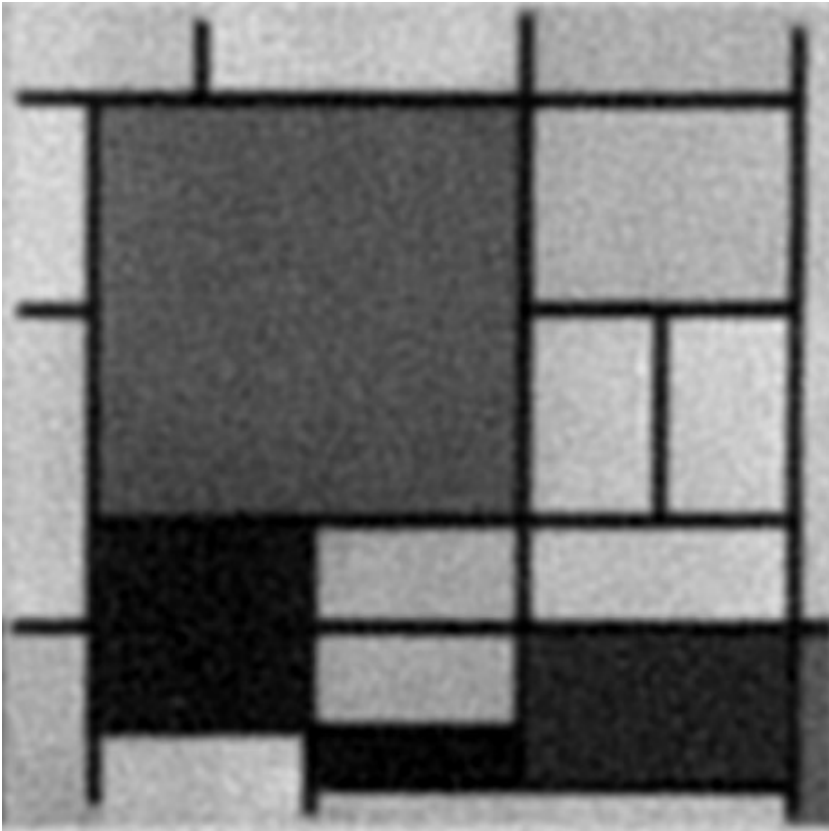
ADMM + TV, $\lambda = 0.01$



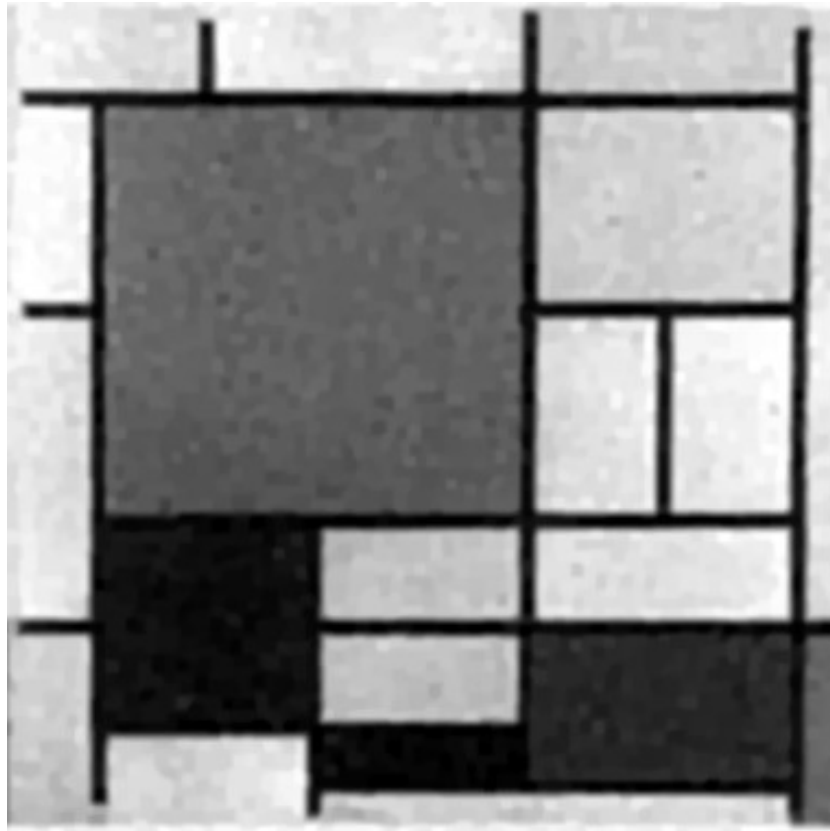
ADMM + TV, $\lambda = 0.1$

- image becomes too flat as we increase weight of TV prior
- Image becomes too noisy as we decrease weight of TV prior

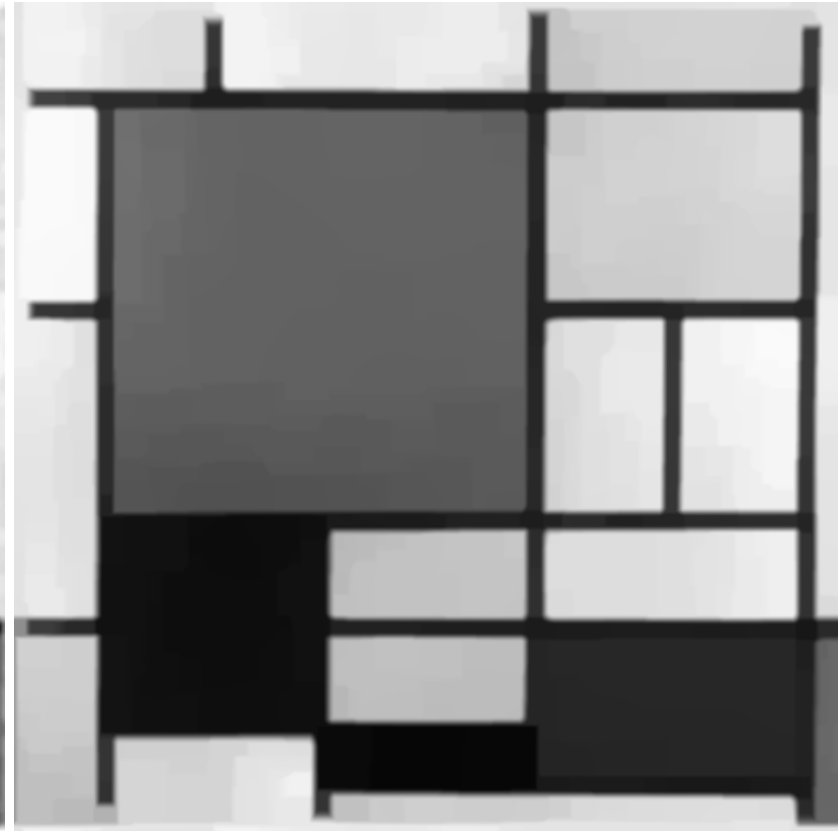
Deconvolution comparisons



Wiener deconvolution



ADMM + TV, $\lambda = 0.01$



ADMM + TV, $\lambda = 0.1$

- image becomes too flat as we increase weight of TV prior
- Image becomes too noisy as we decrease weight of TV prior

Outlook ADMM

- powerful tool for many computational imaging problems
- include generic prior in $g(z)$, just need to derive proximal operator

$$\underset{x}{\text{minimize}} \underbrace{\frac{1}{2} \|Ax - b\|_2^2}_{\text{data fidelity}} + \underbrace{\Gamma(x)}_{\text{regularization}} \quad \longrightarrow \quad \underset{\{x,z\}}{\text{minimize}} \quad f(x) + g(z)$$

subject to $Ax = z$

- example priors: noise statistics, sparse gradient, smoothness, ...
- weighted sum of different priors also possible
- anisotropic TV is one of the easiest priors

Can we do better than that?

Use different gradient regularizations:

- L_2 gradient regularization (Tikhonov regularization, same as Wiener deconvolution)

$$\min_i ||b - k * i||^2 + ||\nabla i||_2^2$$

- L_1 gradient regularization (sparsity regularization, same as *total variation*)

$$\min_i ||b - k * i||^2 + ||\nabla i||_1^1$$

- $L_{n<1}$ gradient regularization (fractional regularization)

$$\min_i ||b - k * i||^2 + ||\nabla i||_{0.8}^{0.8}$$

All of these are motivated by natural image statistics. Active research area.

Comparison of gradient regularizations



input



squared gradient
regularization



fractional gradient
regularization

Derivation

Sensing model:

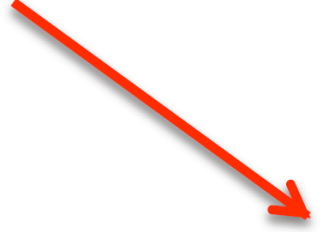
$$b = k * i + n$$

Noise n is assumed to be zero-mean and independent of signal i .



Is this a reasonable noise model?

Richardson-Lucy Algorithm + TV



- log-likelihood function:

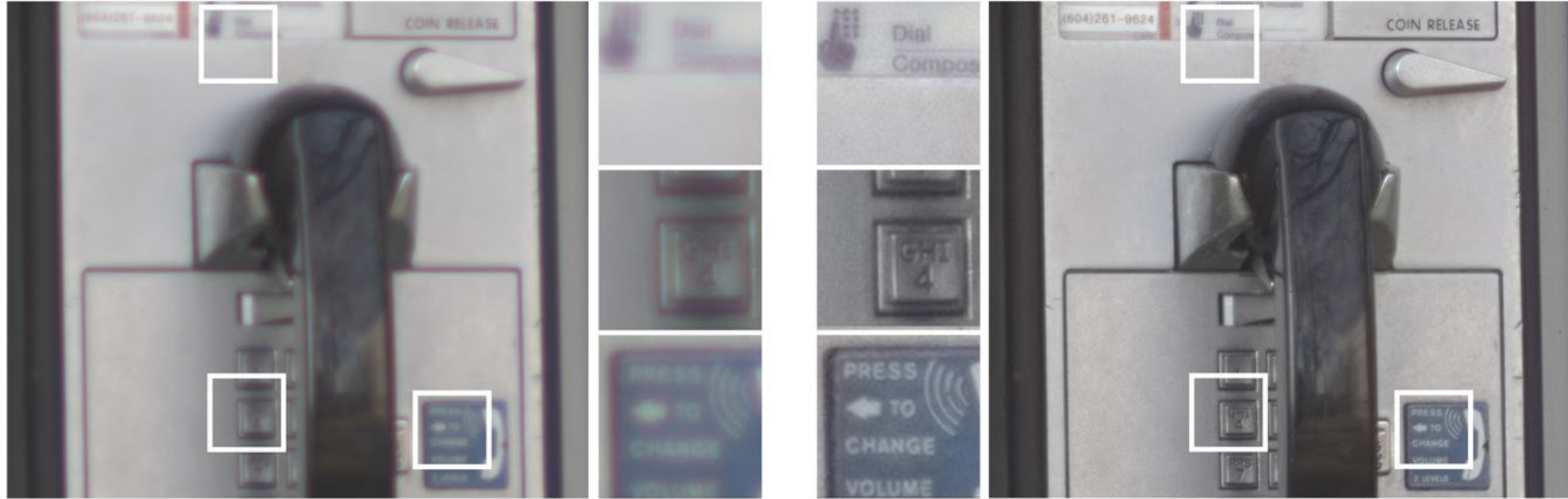
$$\log(L_{TV}(\mathbf{x})) = \log(p(\mathbf{b}|\mathbf{x})) + \log(p(\mathbf{x})) = \log(\mathbf{Ax})^T \mathbf{b} - (\mathbf{Ax})^T \mathbf{1} - \sum_{i=1}^M \log(\mathbf{b}_i!) - \lambda \|\mathbf{Dx}\|_1$$

- gradient:

$$\nabla \log(L_{TV}(\mathbf{x})) = \mathbf{A}^T \text{diag}(\mathbf{Ax})^{-1} \mathbf{b} - \mathbf{A}^T \mathbf{1} + \nabla \lambda \|\nabla \mathbf{x}\|_1 = \mathbf{A}^T \left(\frac{\mathbf{b}}{\mathbf{Ax}} \right) - \mathbf{A}^T \mathbf{1} - \nabla \lambda \|\mathbf{Dx}\|_1$$

- recover signal by setting gradient to zero
- generally challenging

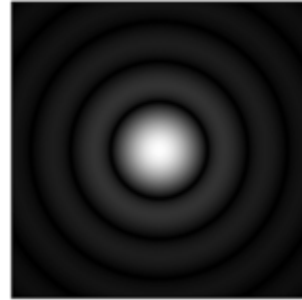
High quality images using cheap lenses



[Heide et al., "High-Quality Computational Imaging Through Simple Lenses," TOG 2013]

Deconvolution

If we know b and k , can we recover i ?

 i $*$  $*$ k $=$ $=$ b

How do we
measure this?



PSF calibration

Take a photo of a point source



Image of PSF

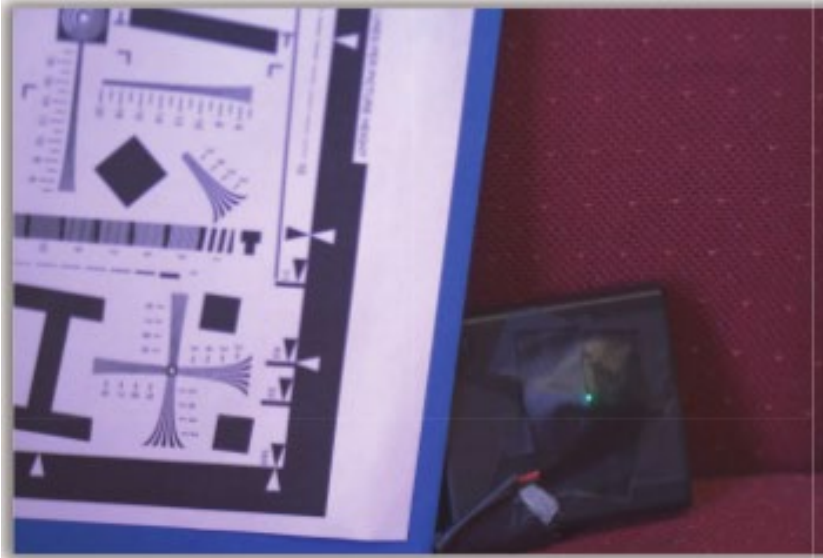


Image with sharp lens



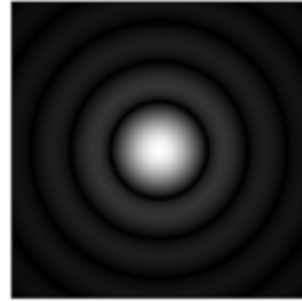
Image with cheap lens

Deconvolution

If we know b and k , can we recover i ?



*



=



i

*

k

=

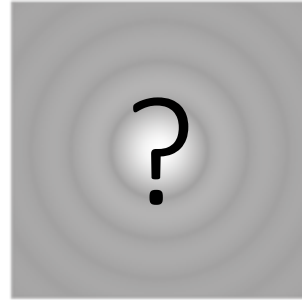
b

Blind deconvolution

If we know b , can we recover i and k ?



*



=



i

*

k

=

b

Camera shake

Removing Camera Shake from a Single Photograph

Rob Fergus¹ Barun Singh¹ Aaron Hertzmann² Sam T. Roweis² William T. Freeman¹

¹MIT CSAIL ²University of Toronto



Figure 1: *Left*: An image spoiled by camera shake. *Middle*: result from Photoshop “unsharp mask”. *Right*: result from our algorithm.

Camera shake as a filter

If we know b , can we recover i and k ?



image from static camera

*



=



image from shaky camera

i

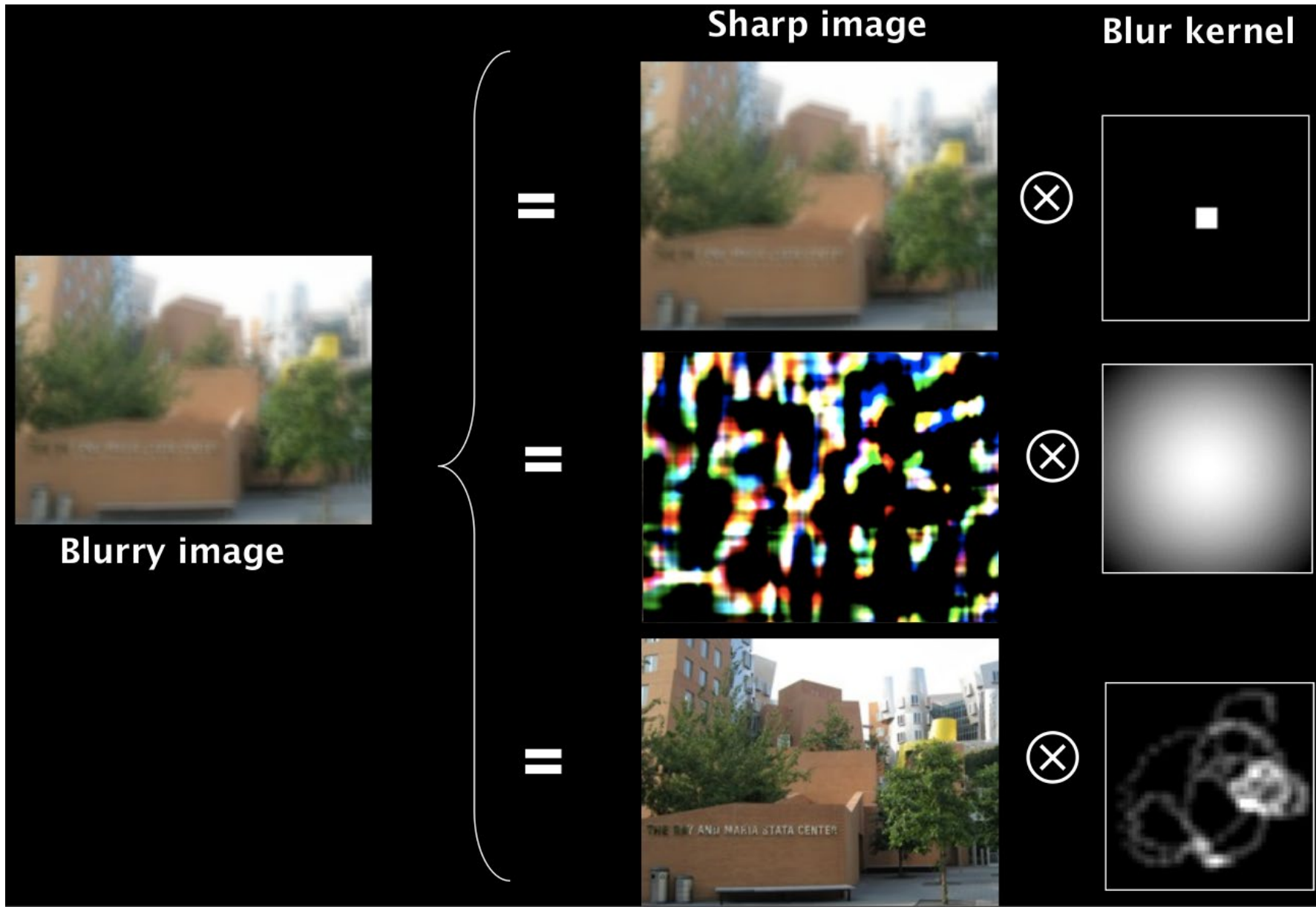
*

k

=

b

Multiple possible solutions



How do we
detect this
one?

Use prior information

Among all the possible pairs of images and blur kernels, select the ones where:

- The image “looks like” a natural image.
- The kernel “looks like” a motion PSF.

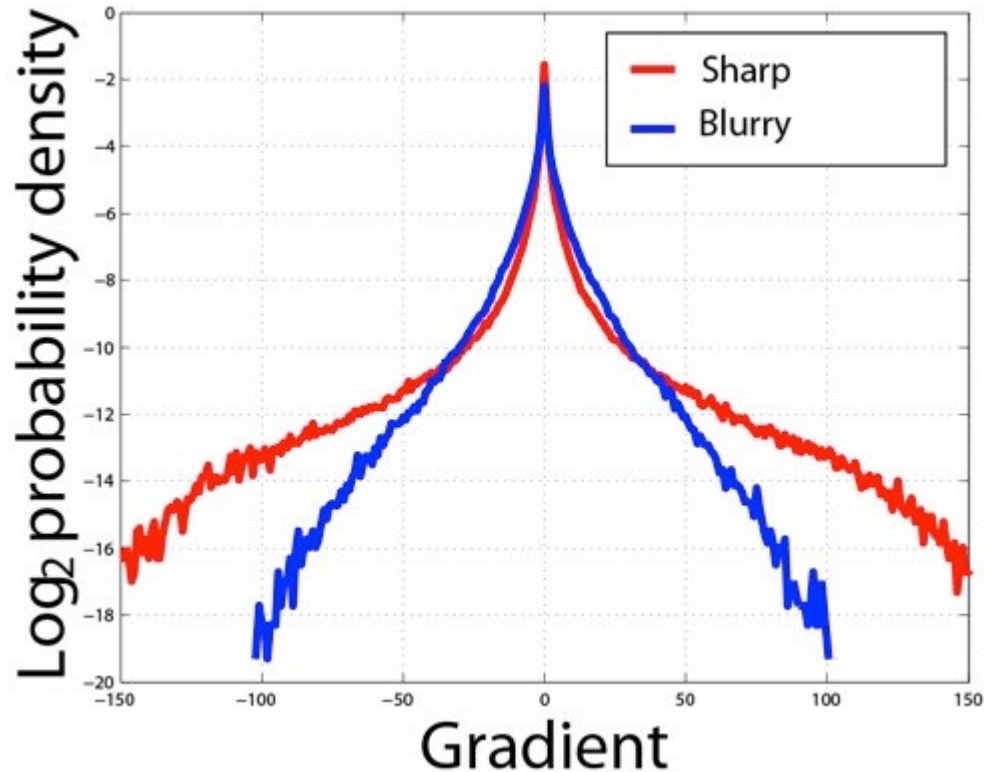
Use prior information

Among all the possible pairs of images and blur kernels, select the ones where:

- The image “looks like” a natural image.
- The kernel “looks like” a motion PSF.

Natural image statistics

Gradients in natural images follow a characteristic “heavy-tail” distribution.



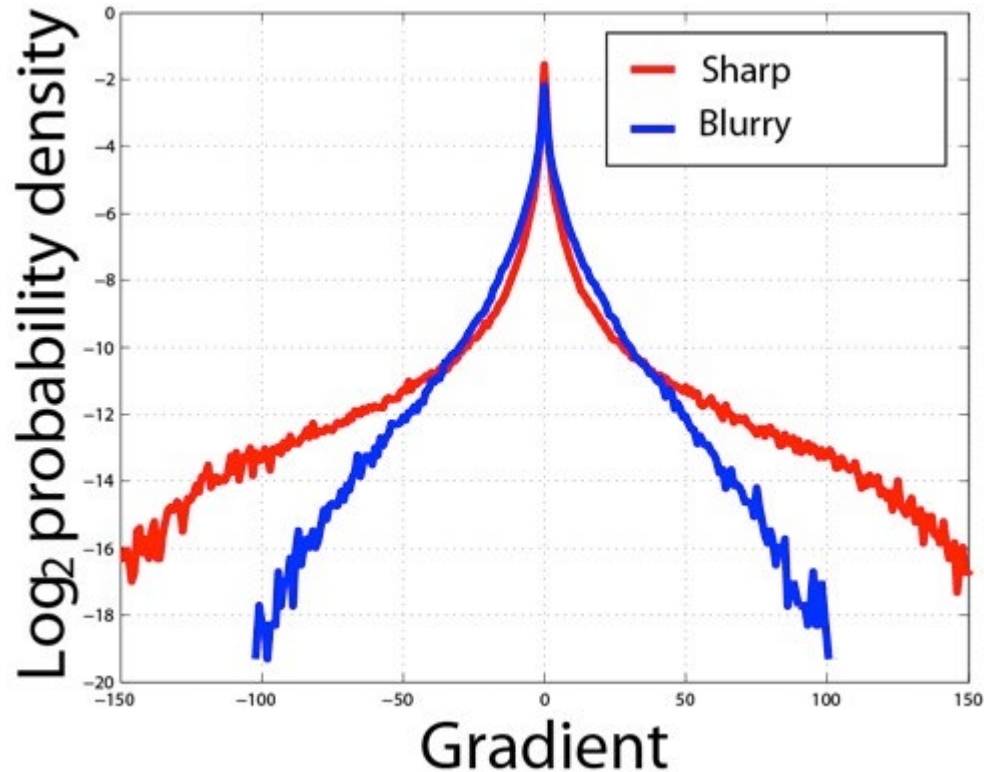
sharp
natural
image



blurry
natural
image

Natural image statistics

Gradients in natural images follow a characteristic “heavy-tail” distribution.



Can be approximated by $\|\nabla i\|^{0.8}$



sharp
natural
image



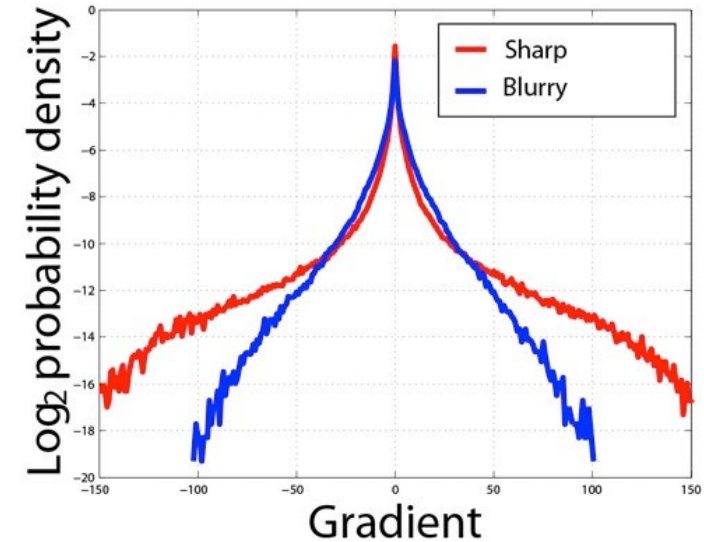
blurry
natural
image

Use prior information

Among all the possible pairs of images and blur kernels, select the ones where:

- The image “looks like” a natural image.

Gradients in natural images follow a characteristic “heavy-tail” distribution.



- The kernel “looks like” a motion PSF.

Shake kernels are very sparse, have continuous contours, and are always positive



How do we use this information for blind deconvolution?

Regularized blind deconvolution

Solve regularized least-squares optimization

$$\min_{i,k} ||b - k * i||^2 + ||\nabla i||^{0.8} + ||k||_1$$

What does each term in this summation correspond to?

Regularized blind deconvolution

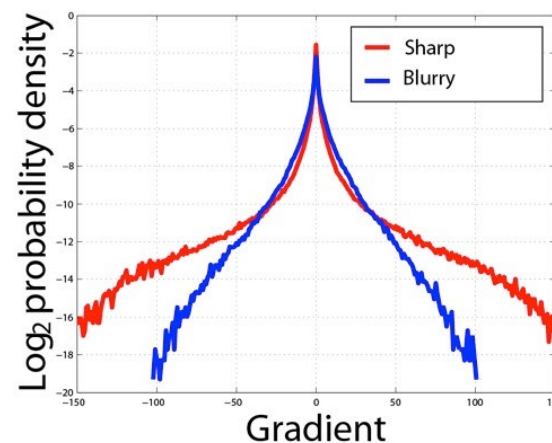
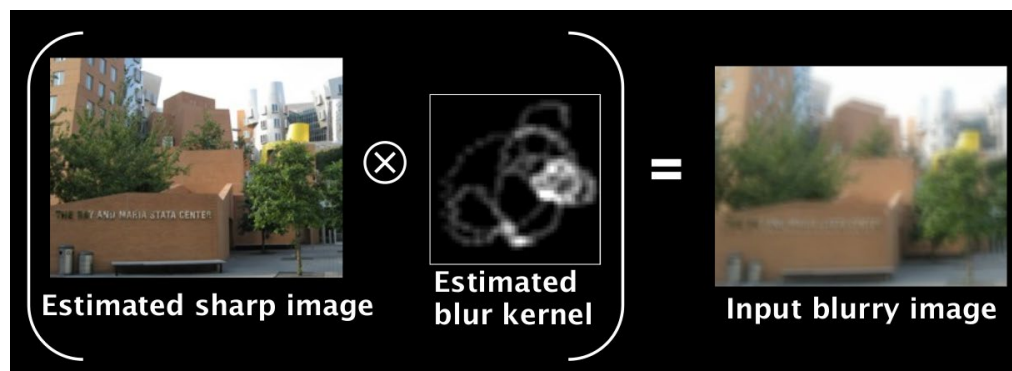
Solve regularized least-squares optimization

$$\min_{i,k} ||b - k * i||^2 + ||\nabla i||^{0.8} + ||k||_1$$

data term

natural image prior

shake kernel prior



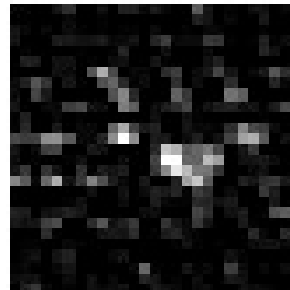
Note: Solving such optimization problems is complicated (no longer *linear* least squares).

A demonstration

input



deconvolved image and kernel



A demonstration

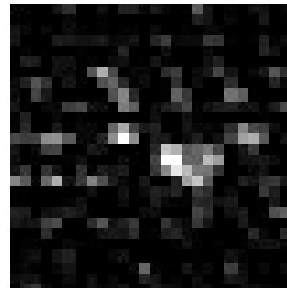
input



deconvolved image and kernel



This image looks worse than the original...



This doesn't look like a plausible shake kernel...

Regularized blind deconvolution

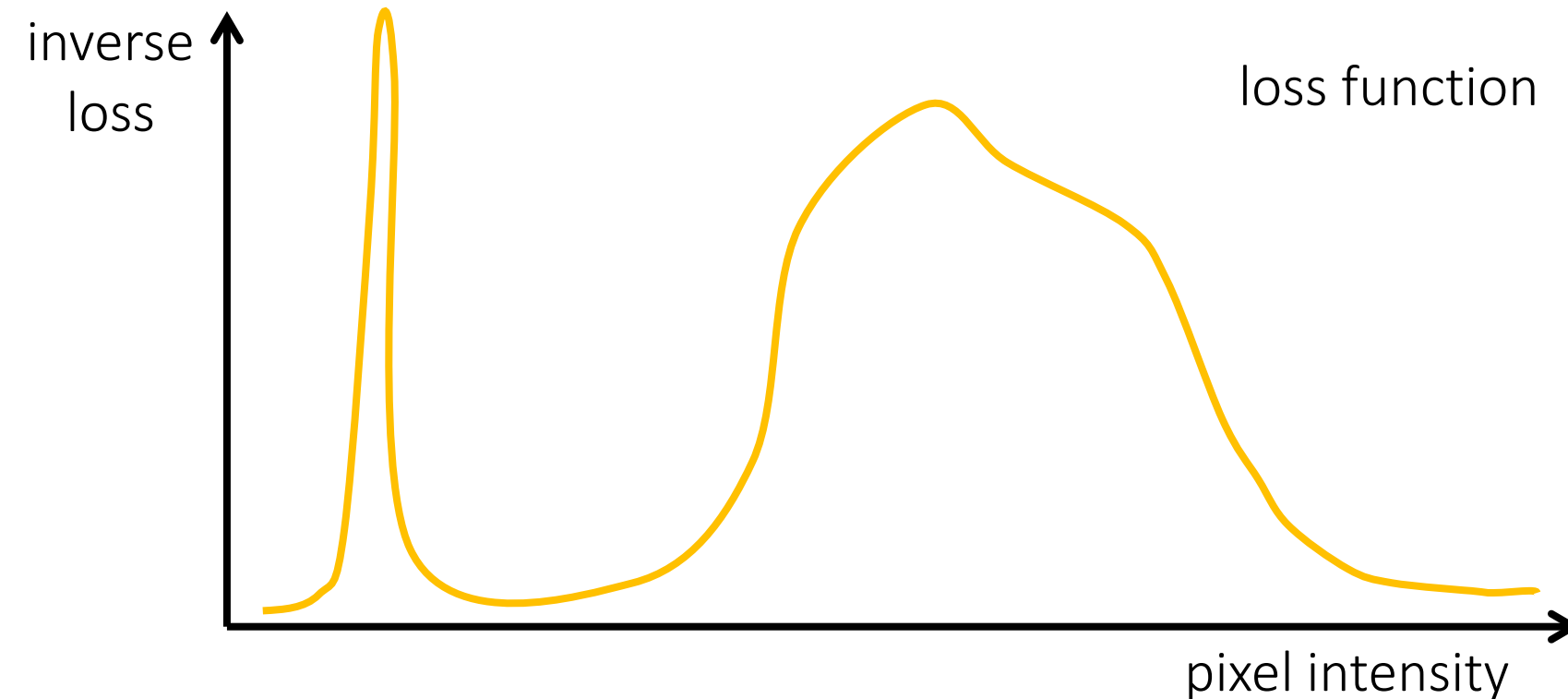
Solve regularized least-squares optimization

$$\min_{i,k} \underbrace{\|b - k * i\|^2 + \|\nabla i\|^{0.8} + \|k\|_1}_{\text{loss function}}$$

Regularized blind deconvolution

Solve regularized least-squares optimization

$$\min_{i,k} \underbrace{\|b - k * i\|^2 + \|\nabla i\|^{0.8} + \|k\|_1}_{\text{loss function}}$$

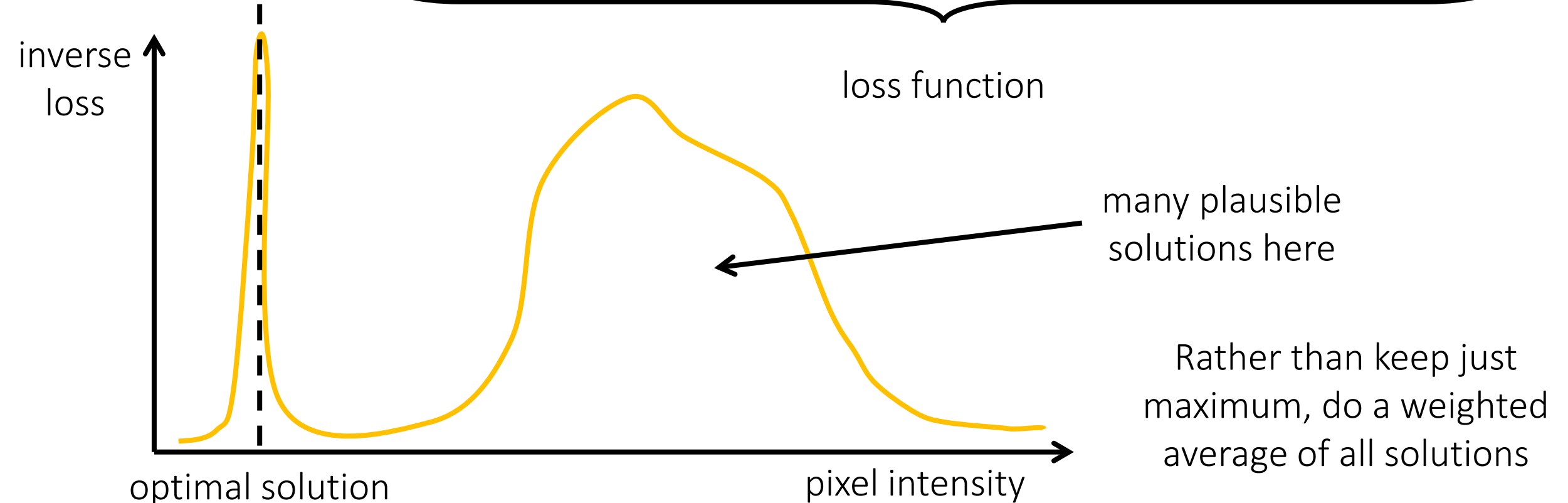


Where in this graph is the solution we find?

Regularized blind deconvolution

Solve regularized least-squares optimization

$$\min_{i,k} \underbrace{\|b - k * i\|^2 + \|\nabla i\|^{0.8} + \|k\|_1}_{\text{loss function}}$$



A demonstration

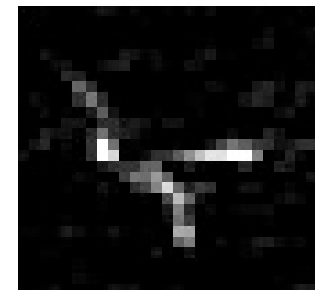
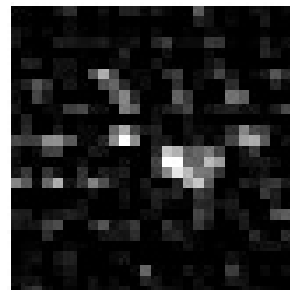
input



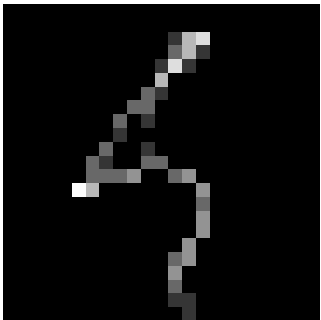
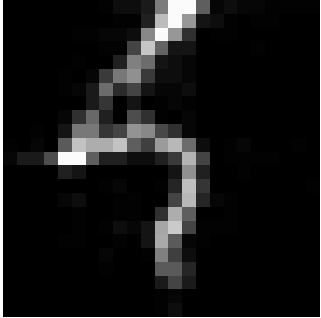
maximum-only



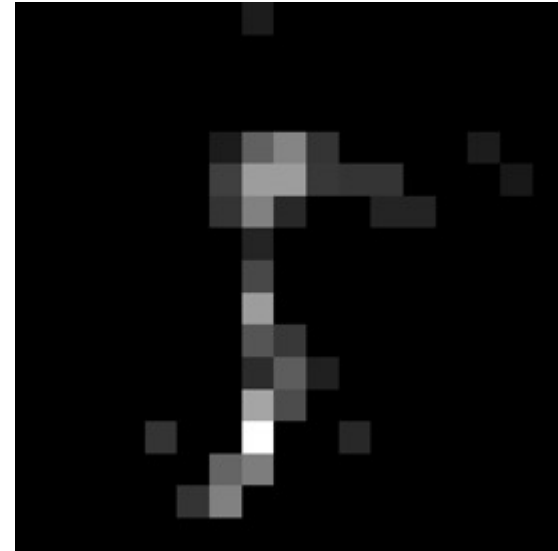
average



More examples



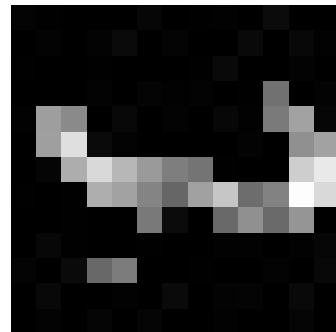
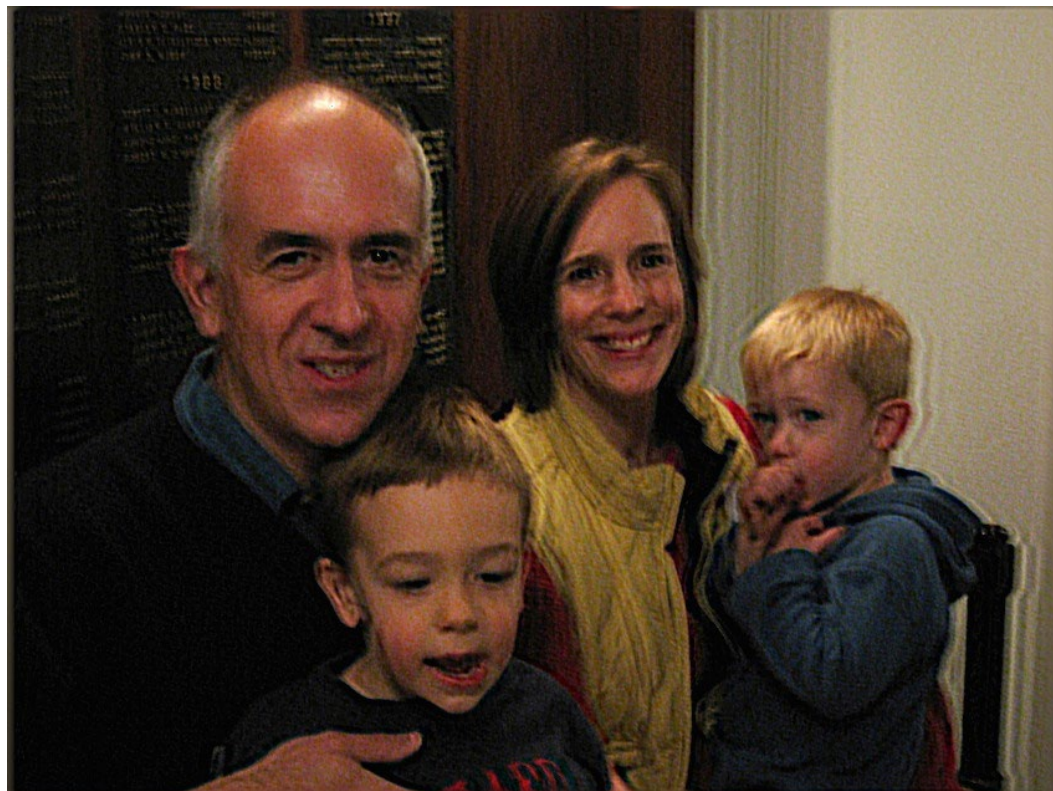
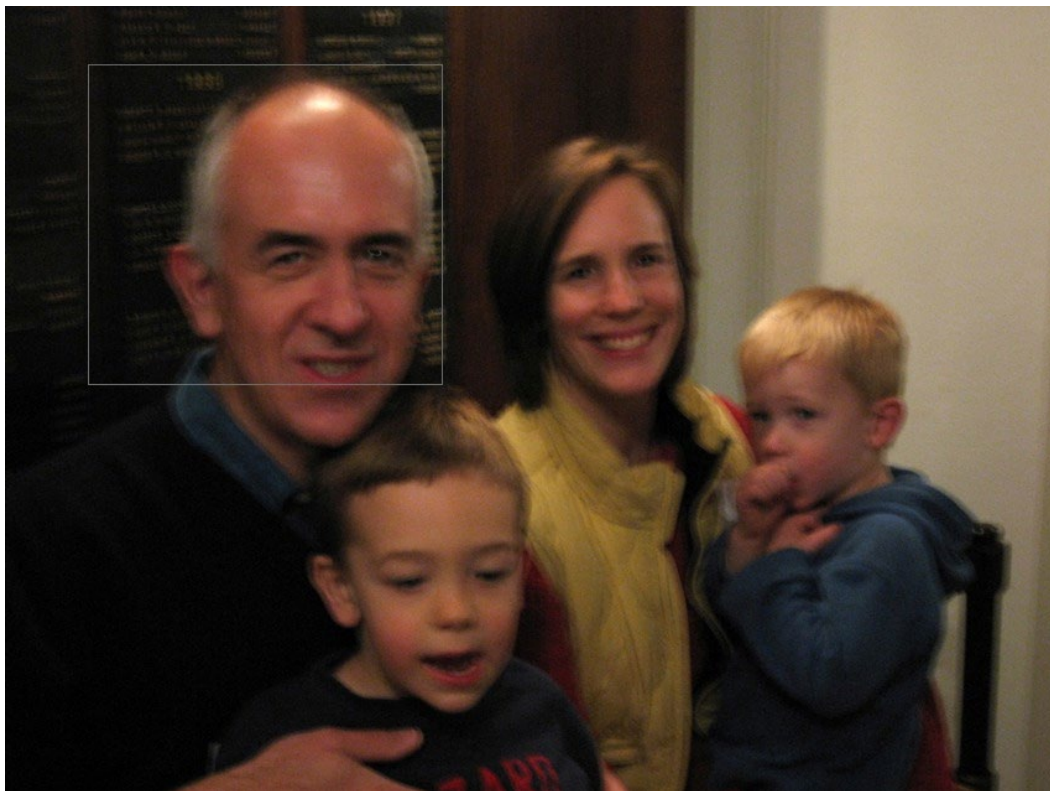
Results on real shaky images



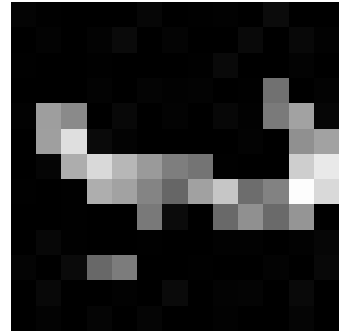
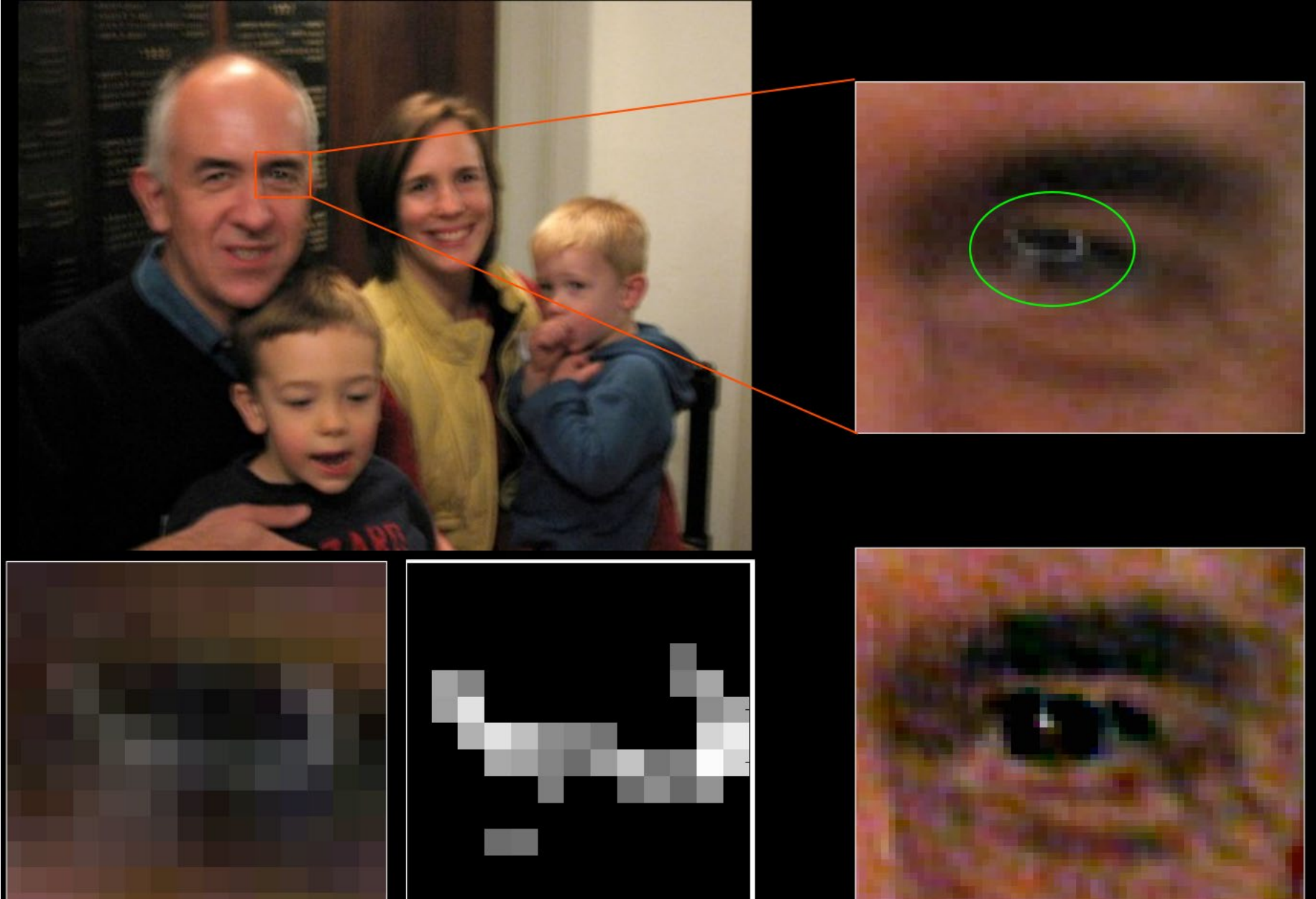
Results on real shaky images



Results on real shaky images



Results on real shaky images



More advanced motion deblurring



[Shah et al., High-quality Motion Deblurring from a Single Image, SIGGRAPH 2008]

Why are our images blurry?

- Lens imperfections.
- Camera shake.
- Scene motion.
- Depth defocus.

Can we solve all of these problems using (blind) deconvolution?

Why are our images blurry?

- Lens imperfections.
- Camera shake.
- Scene motion.
- Depth defocus.

Can we solve all of these problems using (blind) deconvolution?

- We can deal with (some) lens imperfections and camera shake, because their blur is shift invariant.
- We cannot deal with scene motion and depth defocus, because their blur is not shift invariant.
- See coded photography lecture.

References

Basic reading:

- Szeliski textbook, Sections 3.4.3, 3.4.4, 10.1.4, 10.3.
- Fergus et al., “Removing camera shake from a single image,” SIGGRAPH 2006.
the main motion deblurring and blind deconvolution paper we covered in this lecture.

Additional reading:

- Heide et al., “High-Quality Computational Imaging Through Simple Lenses,” TOG 2013.
the paper on high-quality imaging using cheap lenses, which also has a great discussion of all matters relating to blurring from lens aberrations and modern deconvolution algorithms.
- Levin, “Blind Motion Deblurring Using Image Statistics,” NIPS 2006.
- Levin et al., “Image and depth from a conventional camera with a coded aperture,” SIGGRAPH 2007.
- Levin et al., “Understanding and evaluating blind deconvolution algorithms,” CVPR 2009 and PAMI 2011.
- Krishnan and Fergus, “Fast Image Deconvolution using Hyper-Laplacian Priors,” NIPS 2009.
- Levin et al., “Efficient Marginal Likelihood Optimization in Blind Deconvolution,” CVPR 2011.
a sequence of papers developing the state of the art in blind deconvolution of natural images, including the use Laplacian (sparsity) and hyper-Laplacian priors on gradients, analysis of different loss functions and maximum a-posteriori versus Bayesian estimates, the use of variational inference, and efficient optimization algorithms.
- Minskin and MacKay, “Ensemble Learning for Blind Image Separation and Deconvolution,” AICA 2000.
the paper explaining the mathematics of how to compute Bayesian estimators using variational inference.
- Shah et al., “High-quality Motion Deblurring from a Single Image,” SIGGRAPH 2008.
a more recent paper on motion deblurring.