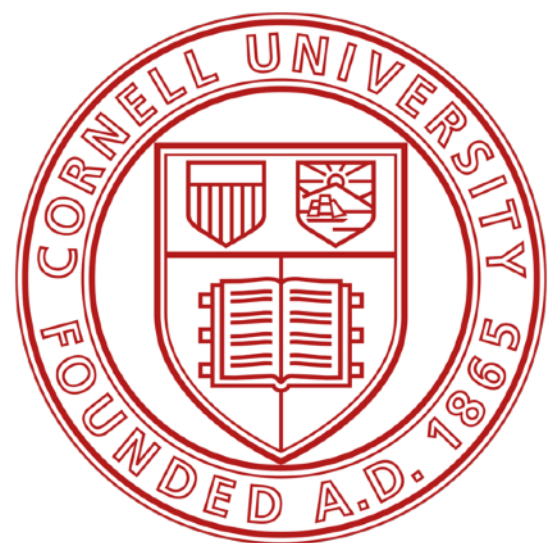


Lecture 6: Normalizing flows

CS 5788: Introduction to Generative Models

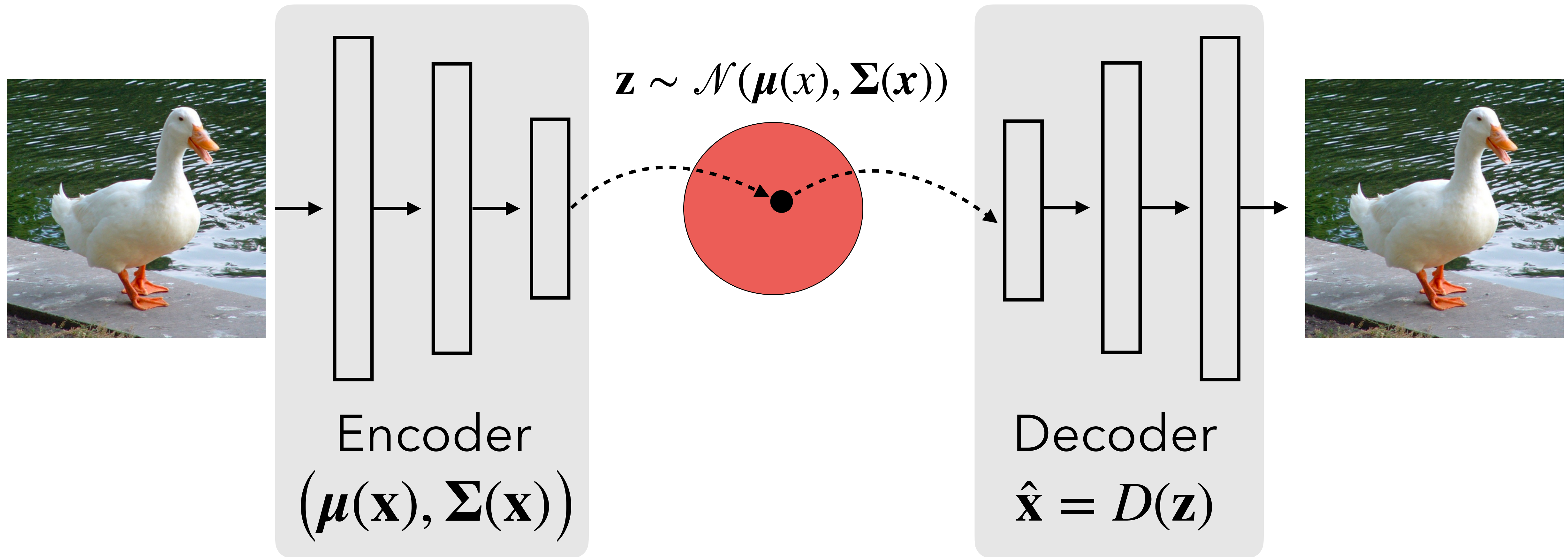


Many slides from Brubaker & Köthe and figures from Kobyzev, Prince, and Brubaker normalizing flow tutorial.

Today

- PS1 due next Tuesday
- Project info out in late February

Recall: Variational autoencoders



- Encoder net defines a distribution over $\mathbf{z}$: $q(\mathbf{z} \mid \mathbf{x}) = \mathcal{N}(\boldsymbol{\mu}(\mathbf{x}), \boldsymbol{\Sigma}(\mathbf{x}))$.
- Decoder net defines a distribution over $\mathbf{x}$: $p(\mathbf{x} \mid \mathbf{z}) = D(\mathbf{z})$.

Recall: training a VAE

$$\begin{aligned}\log p(\mathbf{x}) &= \log \int p(\mathbf{x}, \mathbf{z}) d\mathbf{z} \\ &\geq \int q(\mathbf{z} | \mathbf{x}) \log p(\mathbf{x} | \mathbf{z}) d\mathbf{z} - \int q(\mathbf{z} | \mathbf{x}) \log \frac{q(\mathbf{z} | \mathbf{x})}{p(\mathbf{z})} d\mathbf{z} \\ &= \mathbb{E}_{q(\mathbf{z}|\mathbf{x})} [\log p(\mathbf{x} | \mathbf{z})] - D_{\text{KL}}(q(\mathbf{z} | \mathbf{x}) || p(\mathbf{z}))\end{aligned}$$

Reconstruction loss:

$$\mathcal{L}_R = \|\mathbf{x} - \hat{\mathbf{x}}\|^2 \quad (\text{up to scale factor})$$

Equivalent to log likelihood of $\mathcal{N}(\hat{\mathbf{x}}, I)$: i.e.,
how likely is $\mathbf{x}$ under a Gaussian with mean $\hat{\mathbf{x}}$?

KL divergence loss:

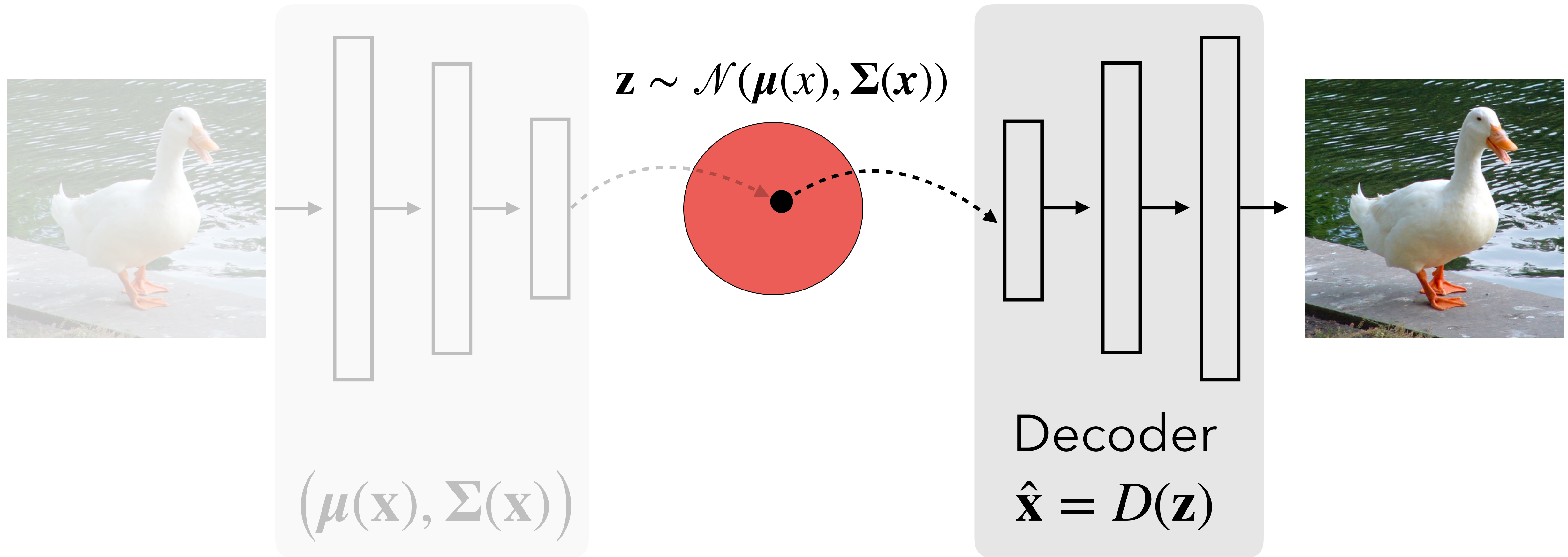
$$\mathcal{L}_P = D_{KL} (q(\mathbf{z} | \mathbf{x}) || \mathcal{N}(0, I))$$

Make $\mathbf{z}$ to match our desired source distribution.

Potential problems with VAEs

- Can't directly estimate $\log p_{\theta}(\mathbf{x})$.
 - Instead, you are optimizing a lower bound.
- Requires you to capture all variation in low-dimensional $\mathbf{z}$.
- Need separate encoder/decoder networks.

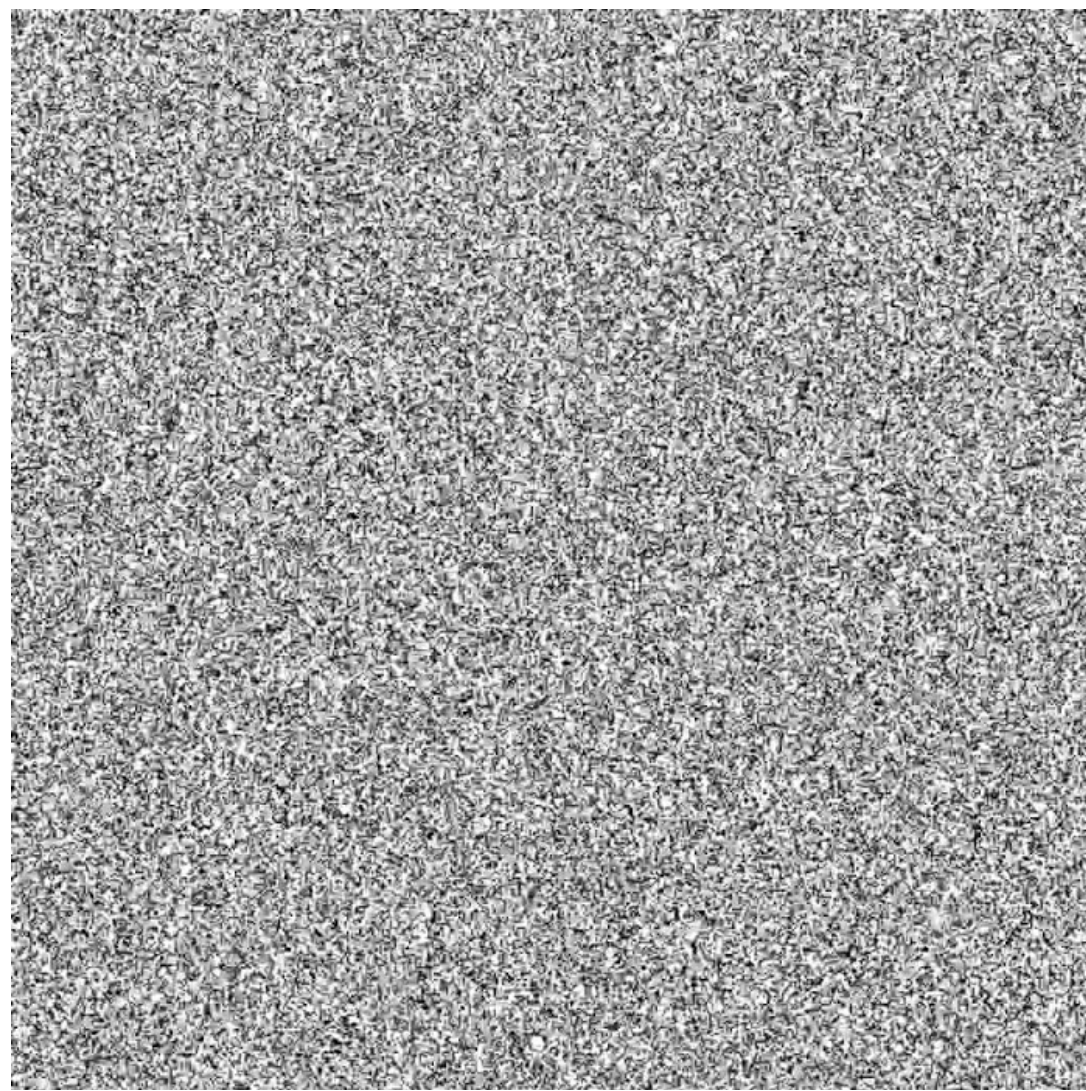
Recall: Variational autoencoders



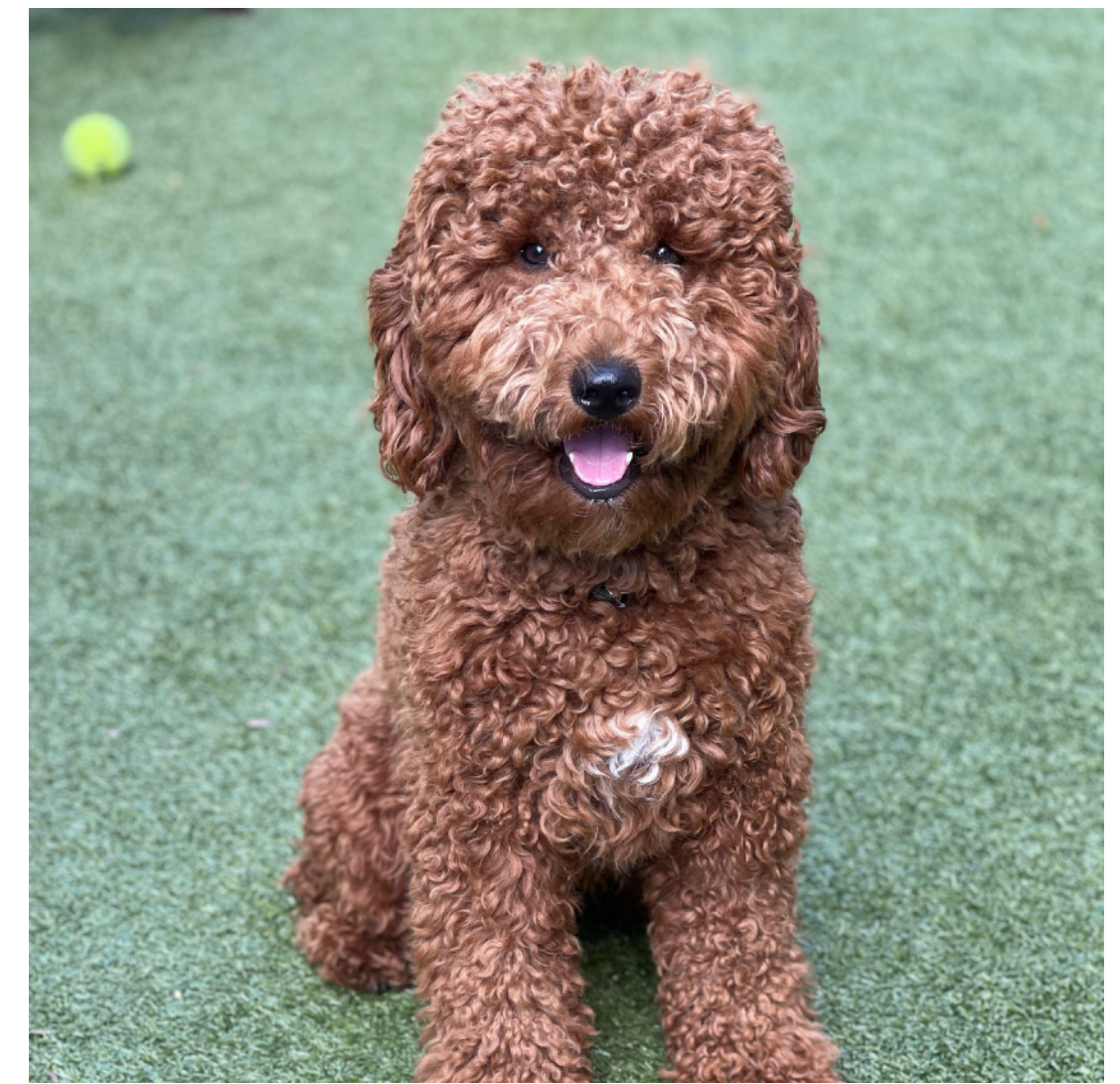
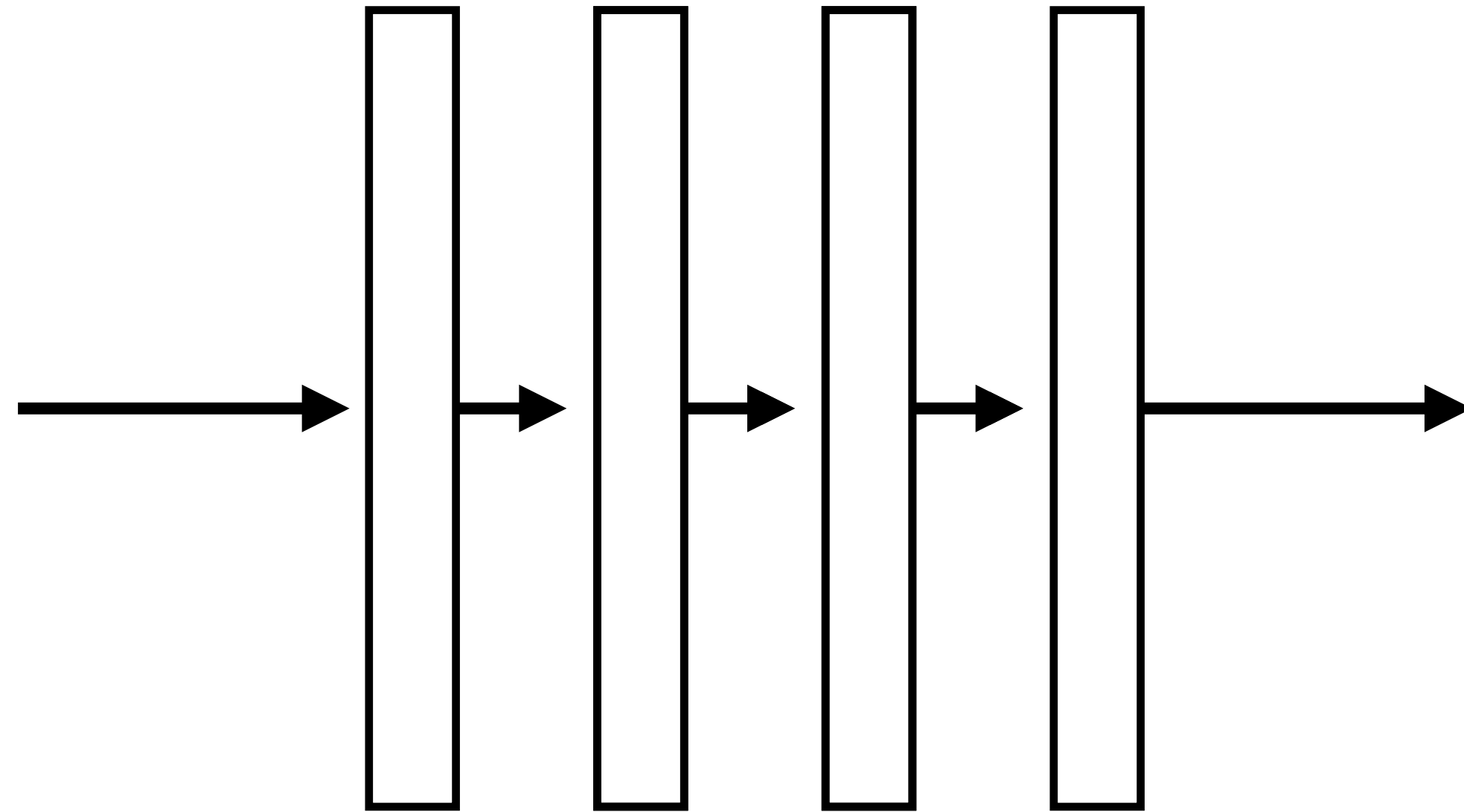
Can we get rid of the encoder?

Invertible transformations

generative direction



$\mathbf{Z}$

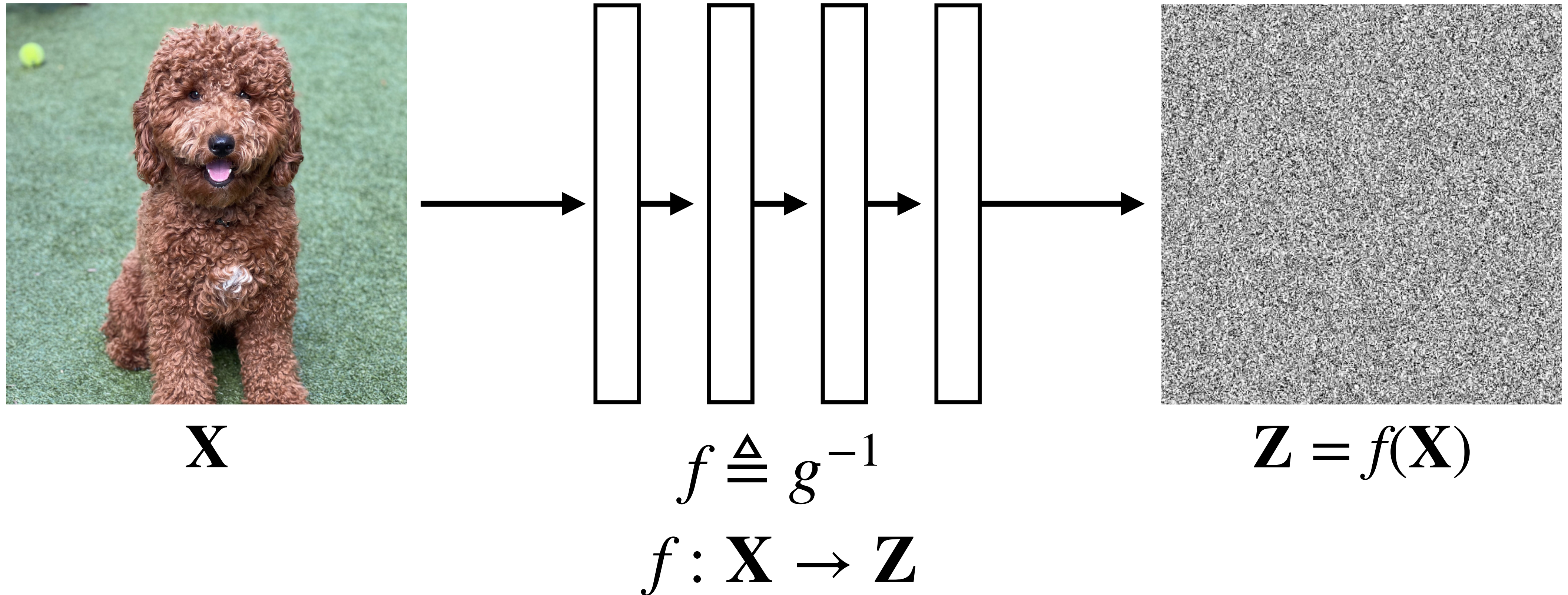


$\mathbf{X} = g(\mathbf{Z})$

$g : \mathbf{Z} \rightarrow \mathbf{X}$

Invertible transformations

normalizing direction ("flow")



Change of variables formula

Volume correction

$$\text{If } \mathbf{X} = g(\mathbf{Z}), \text{ then } p_{\mathbf{X}}(\mathbf{x}) = p_{\mathbf{Z}}(g^{-1}(\mathbf{x})) \left| \det Dg^{-1}(\mathbf{x}) \right|$$

$$p_{\mathbf{X}}(\mathbf{x}) = p_{\mathbf{Z}}(f(\mathbf{x})) \left| \det Df(\mathbf{x}) \right| \text{ for } f \triangleq g^{-1}$$

where Df is the **Jacobian**, i.e., the matrix of partial derivatives:

$$Df = \begin{bmatrix} \frac{\partial f_1}{\partial x_1} & \frac{\partial f_1}{\partial x_2} \\ \frac{\partial f_2}{\partial x_1} & \frac{\partial f_2}{\partial x_2} \end{bmatrix}$$

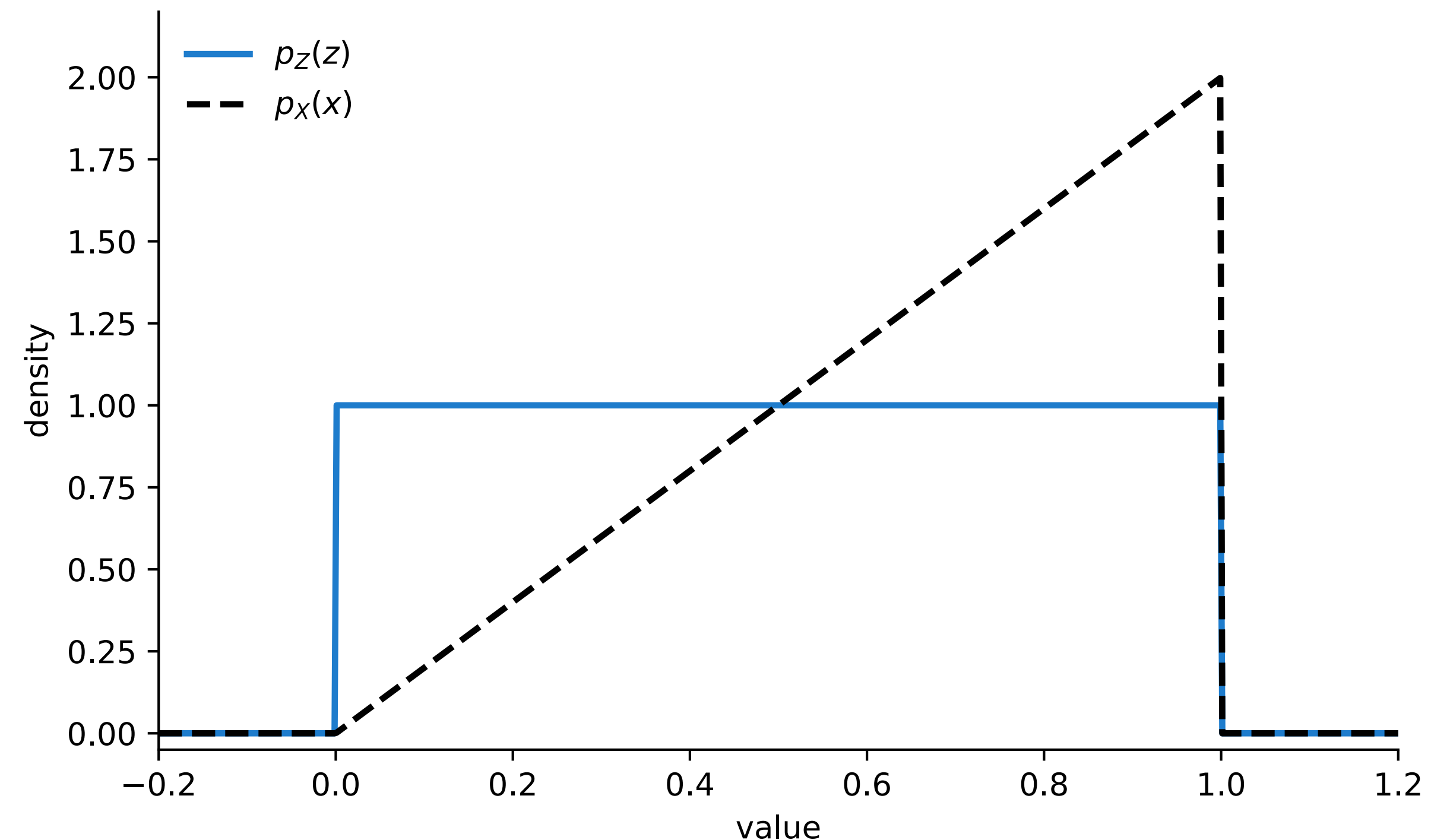
We say that $p_{\mathbf{X}}$ is the **pushforward** of $p_{\mathbf{Z}}$ by g .

Example: change of variables formula

$$\text{If } \mathbf{X} = g(\mathbf{Z}), \text{ then } p_{\mathbf{X}}(\mathbf{x}) = p_{\mathbf{Z}}(g^{-1}(\mathbf{x})) \left| \det Dg^{-1}(\mathbf{x}) \right|$$

$$\mathbf{X} = \sqrt{\mathbf{Z}} \quad \text{where } \mathbf{Z} \sim U(0, 1), \text{ i.e., } p_{\mathbf{Z}}(z) = \begin{cases} 1, & 0 < z < 1, \\ 0, & \text{otherwise.} \end{cases}$$

$$p_{\mathbf{X}}(x) = \begin{cases} 2x & 0 < x < 1 \\ 0 & \text{otherwise} \end{cases}$$

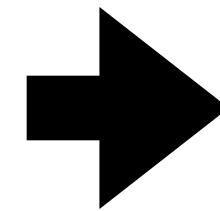


Function composition

Composition of N functions:

$$g = g_N \circ g_{N-1} \circ \dots \circ g_1$$

$$g : \mathbf{Z} \rightarrow \mathbf{X}$$



Inverse:

$$f = f_1 \circ f_2 \circ \dots \circ f_N$$

$$f : \mathbf{X} \rightarrow \mathbf{Z}$$

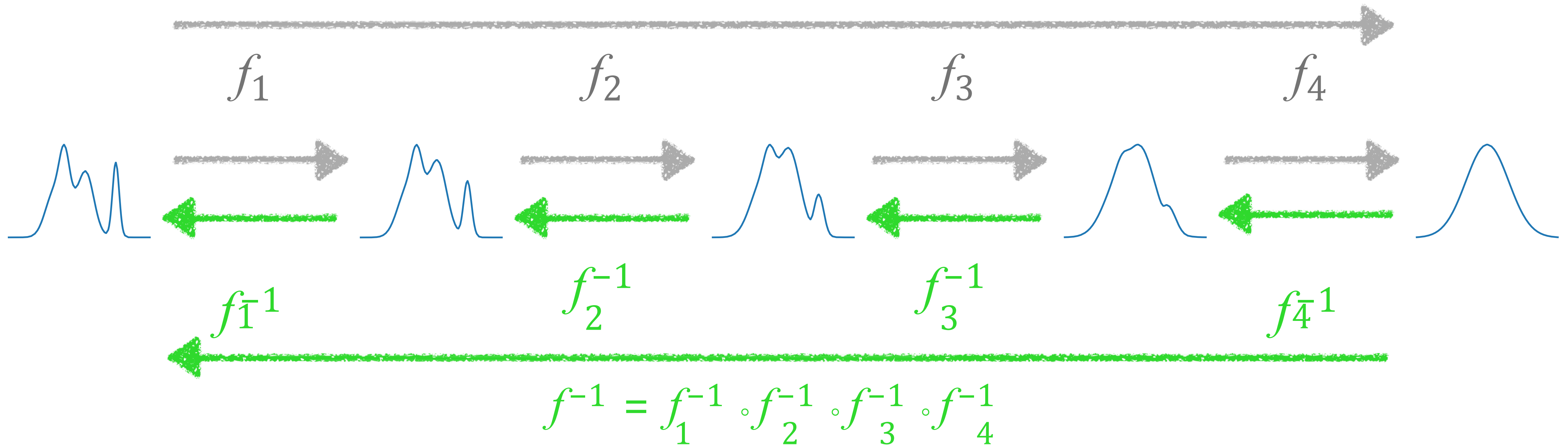
Determinant:

$$\det Df(\mathbf{x}) = \prod_{i=1}^N \det Df_i(\mathbf{y}_i)$$

$$\begin{aligned} \text{where } \mathbf{y}_i &= g_i \circ g_{i-1} \circ \dots \circ g_1(\mathbf{z}) \\ &= f_1 \circ f_2 \circ \dots \circ f_i(\mathbf{x}) \end{aligned}$$

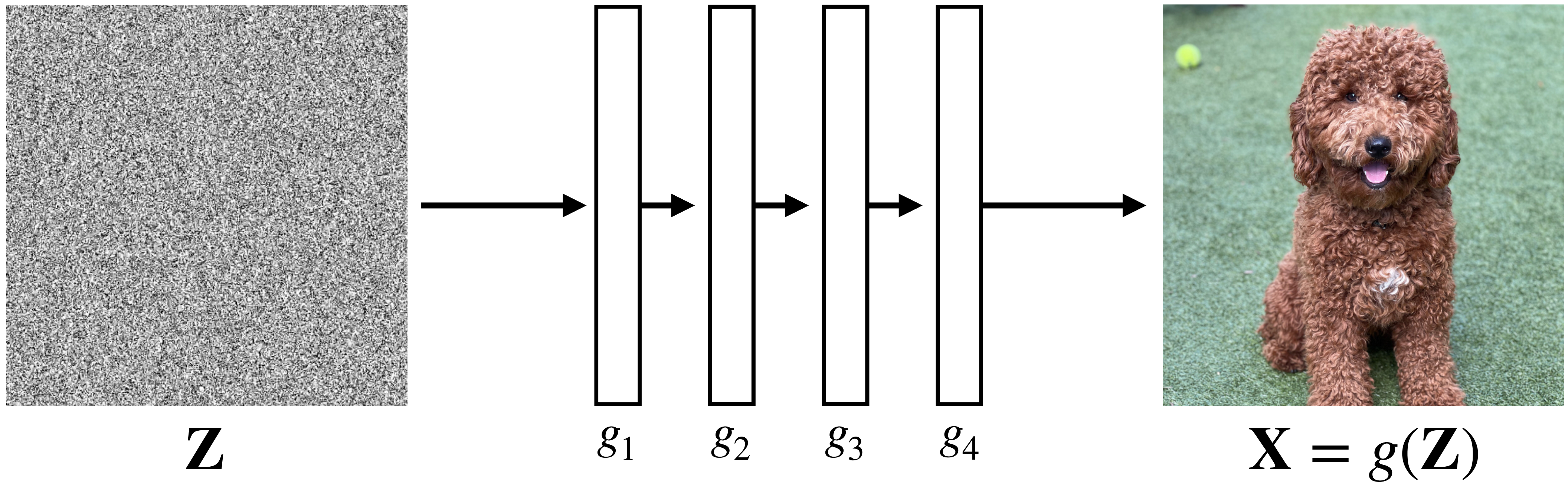
Composition of Flows

$$f = f_4 \circ f_3 \circ f_2 \circ f_1$$



Building a deep generative model

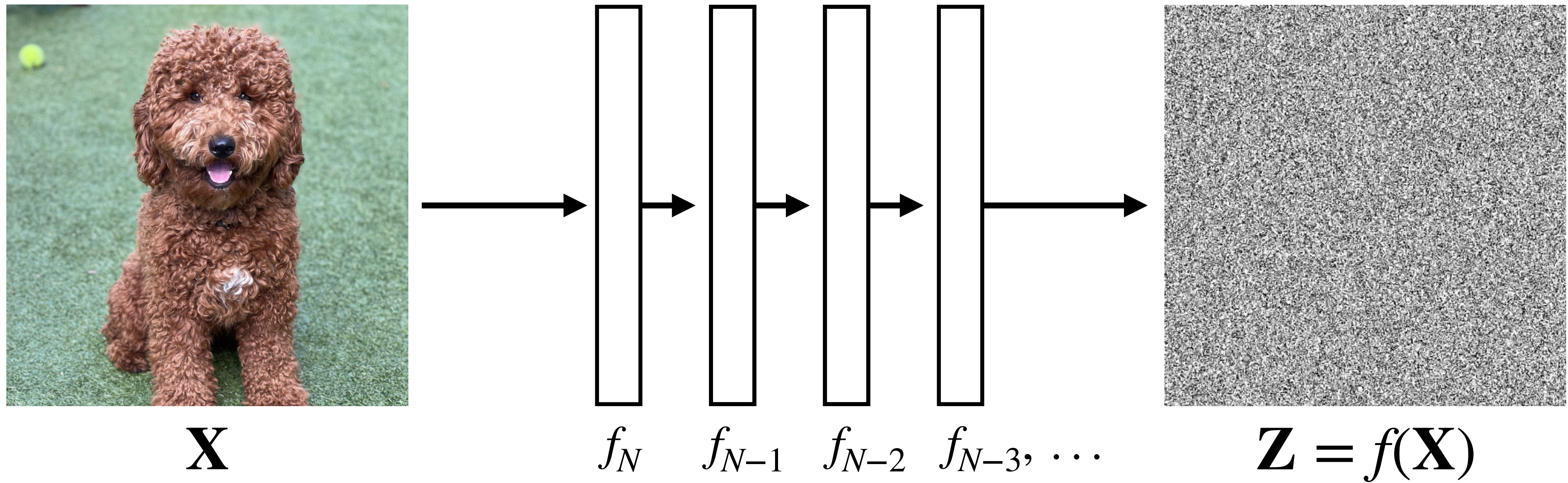
Normalizing flow model



$$g_N \circ g_{N-1} \circ \dots \circ g_1(\mathbf{z})$$

Normalizing flow model

For estimating the density and learning:



$$\log p_{\theta}(\mathbf{X}) = \sum_{i=1}^M \log p_{\mathbf{Z}}(f(\mathbf{x}_i)) + \log \left| \det(\mathbf{D}f(\mathbf{x}_i)) \right| \text{ where } \mathbf{x}_i \text{ is a training example.}$$

Exact likelihood estimation (unlike VAE, which uses lower bound).

What's the catch?

Let's look at the change-of-variables formula again:

$$p_{\mathbf{X}}(\mathbf{x}) = p_{\mathbf{Z}}\left(g^{-1}(\mathbf{x})\right) \left| \det Dg^{-1}(\mathbf{x}) \right|$$

Network architecture?

$$\log p_{\theta}(\mathbf{X}) = \sum_{i=1}^M \log p_{\mathbf{Z}}(f(\mathbf{x}_i)) + \log \left| \det(Df(\mathbf{x}_i)) \right|$$

The change-of-variables formula places constraints on network:

- Each layer g_i must be invertible
- Must be able to compute $\log \det Df_i$
- Latent $\mathbf{z}$ needs to be same dimension as $\mathbf{x}$
- All of these must be efficient to compute!

Pointwise nonlinearities

$$f(\mathbf{y}) = \begin{bmatrix} g(y_1) \\ g(y_2) \\ \dots \\ g(y_d) \end{bmatrix} \quad \text{if } g \text{ is invertible, so is } f$$

Pointwise nonlinearities

Jacobian:

$$f(\mathbf{y}) = \begin{bmatrix} g(y_1) \\ g(y_2) \\ \dots \\ g(y_d) \end{bmatrix}$$

$$D = \begin{bmatrix} g'(y_1) & 0 & 0 & 0 \\ 0 & g'(y_2) & 0 & 0 \\ 0 & 0 & g'(y_3) & 0 \\ 0 & 0 & 0 & \dots \end{bmatrix}$$

$$\det D = \prod_{i=1}^N g'(y_i)$$

(Jacobian of f^{-1} is similar, but using f for simplicity)

Linear layers

Inverse? $f^{-1}(\mathbf{y}) = A^{-1}(\mathbf{y} - \mathbf{b})$

$$f(\mathbf{y}) = A\mathbf{y} + \mathbf{b}$$

but A needs to be full rank.

$$f: \mathbb{R}^n \rightarrow \mathbb{R}^n$$

Determinant? $\det Df(\mathbf{x}) = \det A$

...but expensive! $\mathcal{O}(N^3)$

Linear layers: special cases

Diagonal matrix:

$$g(\mathbf{y}) = D\mathbf{y} + \mathbf{b}$$

$$g : \mathbb{R}^n \rightarrow \mathbb{R}^n$$

$$D = \text{diag}(a, b, c, \dots) = \begin{bmatrix} a & 0 & 0 & 0 \\ 0 & b & 0 & 0 \\ 0 & 0 & c & 0 \\ 0 & 0 & 0 & \dots \end{bmatrix}$$

$$\det(D) = a \times b \times c \times \dots$$

computable in $O(n)$ time.

Linear layers: special cases

Triangular matrix:

$$g(\mathbf{y}) = L\mathbf{y} + \mathbf{b}$$

$$g : \mathbb{R}^n \rightarrow \mathbb{R}^n$$

$$L = \begin{bmatrix} a & 0 & 0 & 0 \\ b & c & 0 & 0 \\ d & e & f & 0 \\ g & h & i & j \end{bmatrix}$$

- Determinant is product of diagonal.
- Invertible in $O(n^2)$ time.

Still, limited expressiveness...

Linear layers: permutation matrix

Permutation matrix:

What does this matrix do?

$$g(\mathbf{y}) = P\mathbf{y} + \mathbf{b}$$

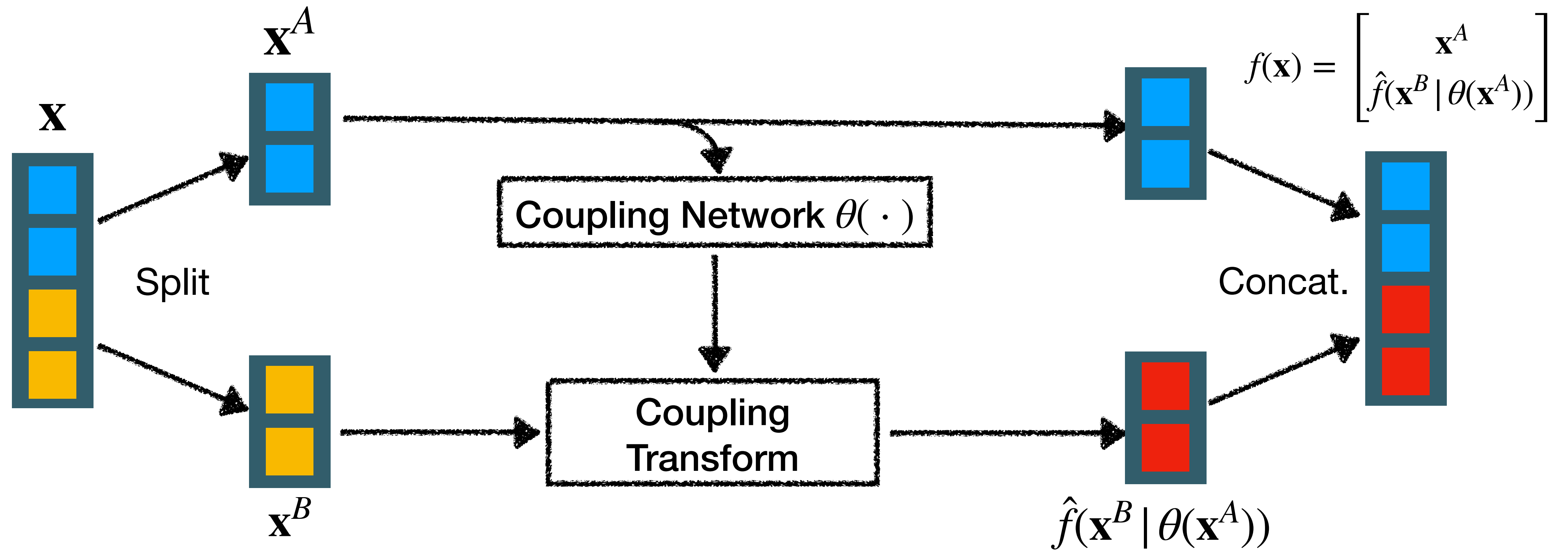
$$g : \mathbb{R}^n \rightarrow \mathbb{R}^n$$

$$P = \begin{bmatrix} 0 & 1 & 0 & 0 \\ 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 \\ 0 & 0 & 1 & 0 \end{bmatrix} \quad \det(P) = 1$$

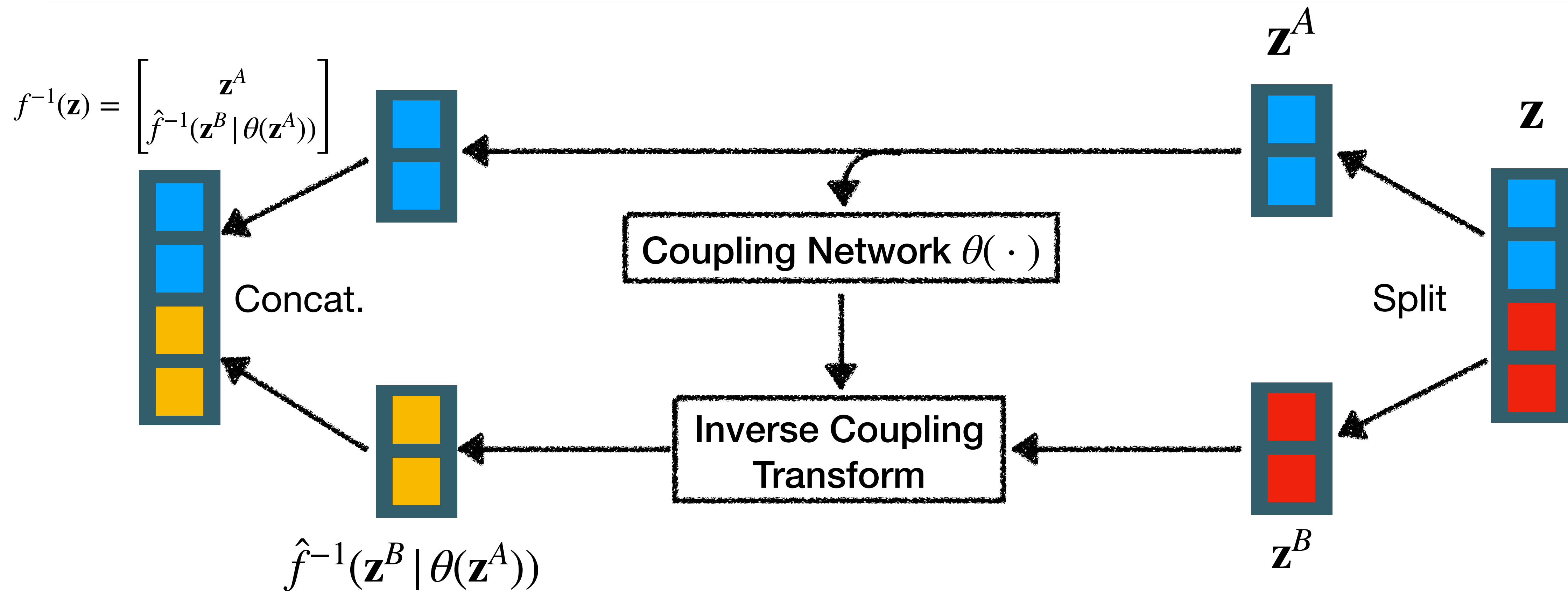
Can chain together with lower triangular matrix to improve expressiveness.

Other useful special cases

Coupling flows



Coupling flows: inverse



Coupling flows

$$f(\mathbf{x}) = \left[\begin{array}{c} \mathbf{x}^A \\ \hat{f}(\mathbf{x}^B \mid \theta(\mathbf{x}^A)) \end{array} \right] \quad \mathrm{D}f(\mathbf{x}) = \left[\begin{array}{c} \phantom{\hat{f}(\mathbf{x}^B \mid \theta(\mathbf{x}^A))} \\ \phantom{\hat{f}(\mathbf{x}^B \mid \theta(\mathbf{x}^A))} \end{array} \right]$$

Coupling flows

$$f(\mathbf{x}) = \left[\hat{f}(\mathbf{x}^B \mid \theta(\mathbf{x}^A)) \right] \quad \mathbf{D}f(\mathbf{x}) = \begin{bmatrix} I & \mathbf{0} \\ \frac{\partial}{\partial \mathbf{x}^A} \hat{f}(\mathbf{x}^B \mid \theta(\mathbf{x}^A)) & \mathbf{D} \hat{f}(\mathbf{x}^B \mid \theta(\mathbf{x}^A)) \end{bmatrix}$$

Coupling flows

$$Df(\mathbf{x}) = \begin{bmatrix} I & \mathbf{0} \\ \frac{\partial}{\partial \mathbf{x}^A} \hat{f}(\mathbf{x}^B \mid \theta(\mathbf{x}^A)) & D\hat{f}(\mathbf{x}^B \mid \theta(\mathbf{x}^A)) \end{bmatrix}$$

Useful property:

$$\det Df(\mathbf{x}) = ?$$

$$\det \begin{bmatrix} A & \mathbf{0} \\ C & D \end{bmatrix} = \det(A) \det(D)$$

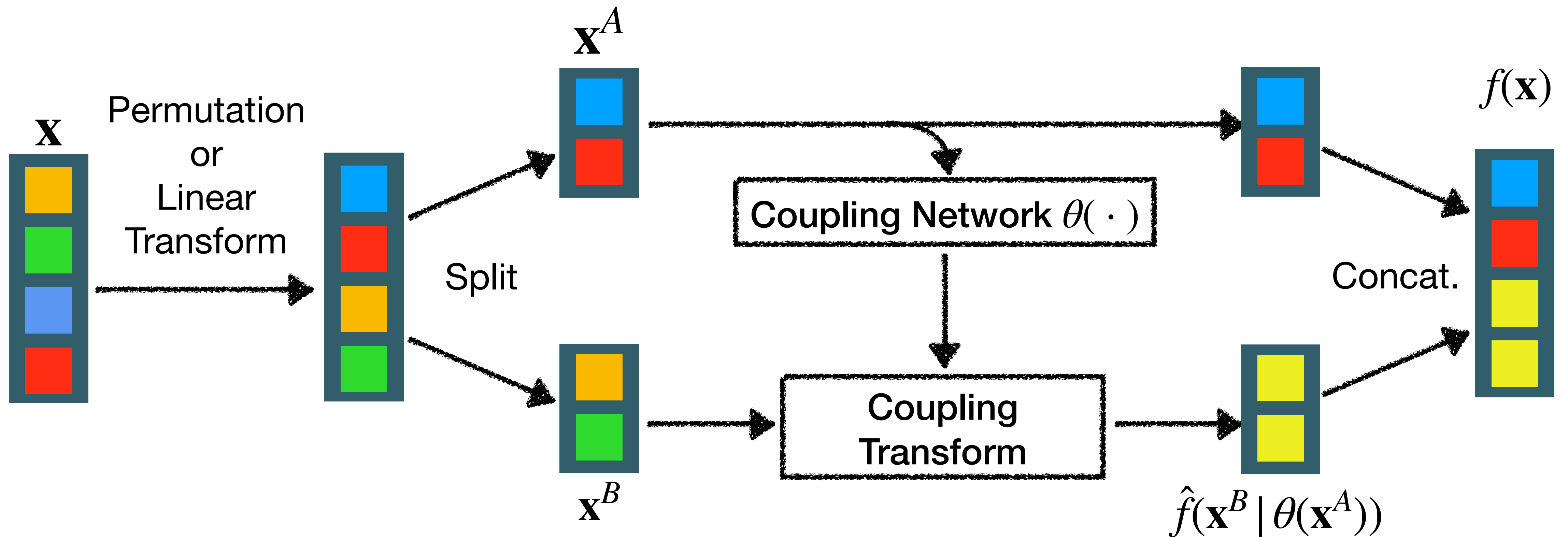
Coupling flows

$$Df(\mathbf{x}) = \begin{bmatrix} I & \mathbf{0} \\ \frac{\partial}{\partial \mathbf{x}^A} \hat{f}(\mathbf{x}^B \mid \theta(\mathbf{x}^A)) & D\hat{f}(\mathbf{x}^B \mid \theta(\mathbf{x}^A)) \end{bmatrix}$$

$$\det Df(\mathbf{x}) = \det \left(D\hat{f}(\mathbf{x}^B \mid \theta(\mathbf{x}^A)) \right)$$

$\theta(\mathbf{x}^A)$ is arbitrary!
(e.g., MLP, CNN, etc.)

Coupling flows



Can increase expressivity with other operations.

Coupling transforms

Lots of possible choices.

- Additive [“NICE”, Dinh et al 2014]

$$\hat{f}(\mathbf{x} \mid \mathbf{t}) = \mathbf{x} + \mathbf{t}$$

(where $\mathbf{t}$ is a function that maps $\mathbf{x}^A$ to a vector the same size as $\mathbf{x}^B$).

- Affine [“RealNVP”, Dinh et al 2016], where $\odot$ is elementwise product

$$\hat{f}(\mathbf{x} \mid \mathbf{s}, \mathbf{t}) = \mathbf{x} \odot \exp(\mathbf{s}) + \mathbf{t}$$

Coupling flows (affine model)

$$\hat{f}(\mathbf{x} \mid \mathbf{s}, \mathbf{t}) = \mathbf{x} \odot \exp(\mathbf{s}) + \mathbf{t}$$

Jacobian:

$$Df(\mathbf{x}) = \begin{bmatrix} I & \mathbf{0} \\ \frac{\partial}{\partial \mathbf{x}^A} \hat{f}(\mathbf{x}^B \mid \theta(\mathbf{x}^A)) & D\hat{f}(\mathbf{x}^B \mid \theta(\mathbf{x}^A)) \end{bmatrix}$$

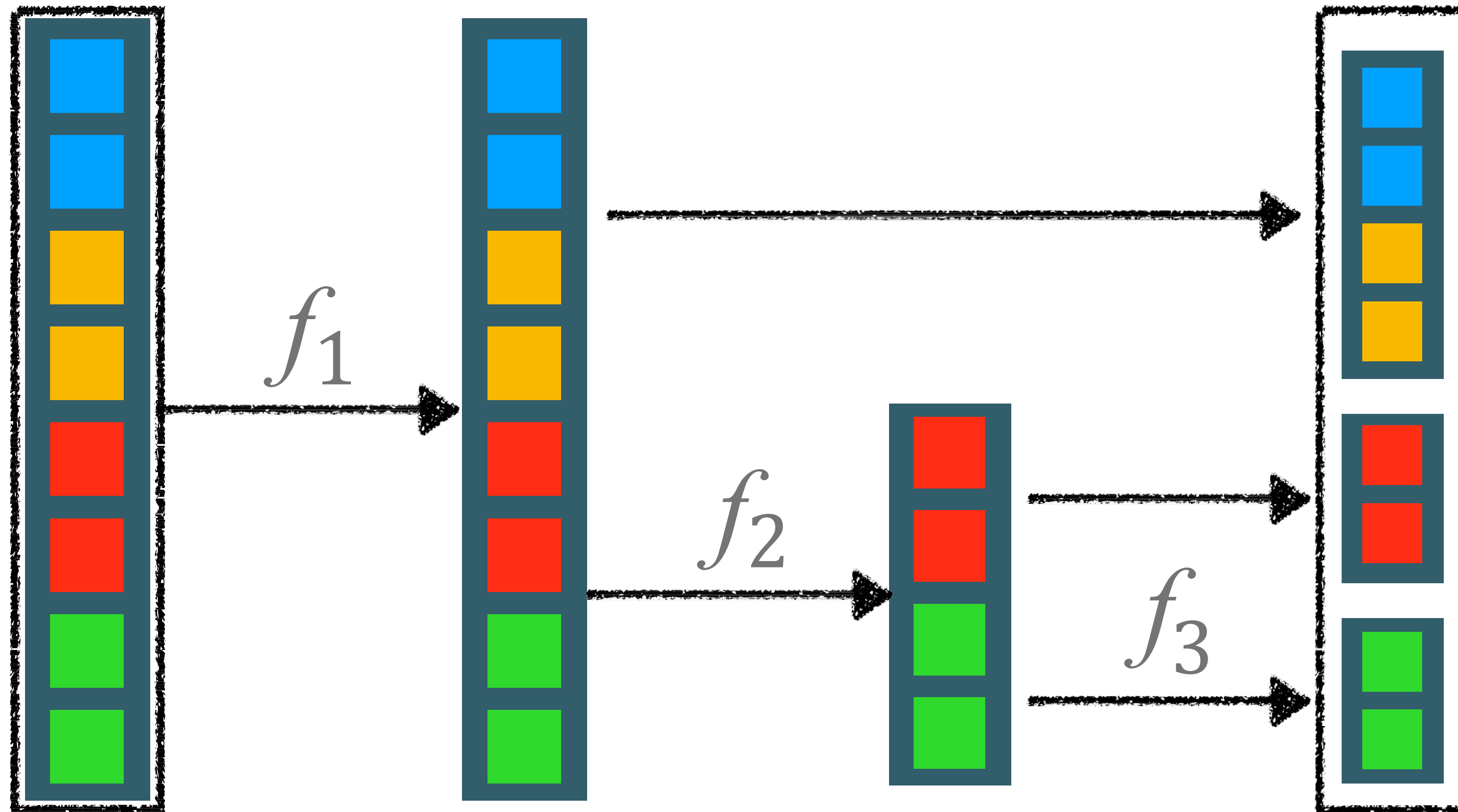
$$Df(\mathbf{x}) = \begin{bmatrix} I & \mathbf{0} \\ \frac{\partial}{\partial \mathbf{x}^A} \hat{f}(\mathbf{x}^B \mid \theta(\mathbf{x}^A)) & \text{diag}(\exp(\mathbf{s}_1), \dots, \exp(\mathbf{s}_K)) \end{bmatrix}$$

Log determinant:

$$\det(Df(\mathbf{x})) = \exp\left(\sum_{i=1}^K \mathbf{s}_i\right)$$

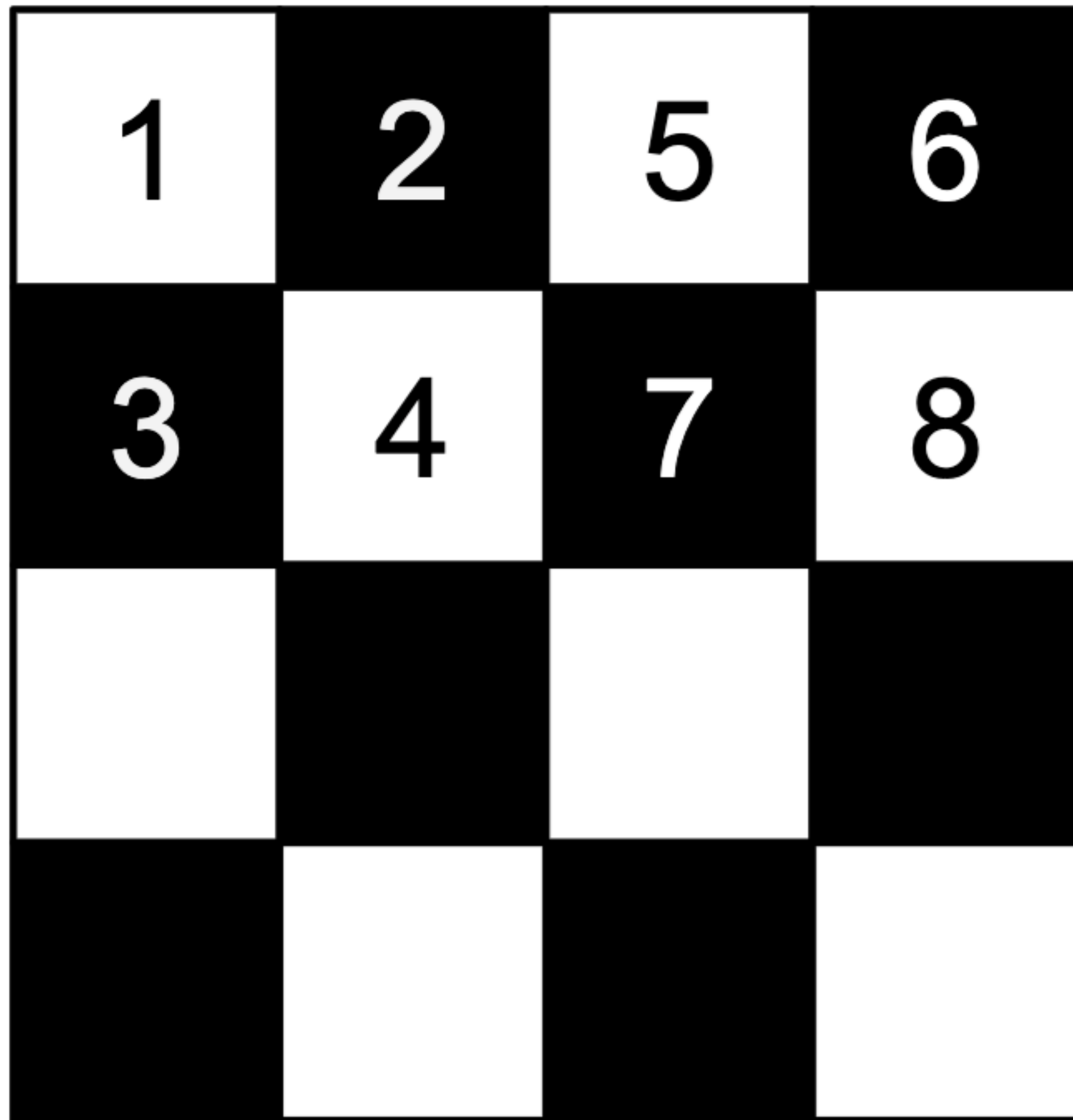
Handling images

Multiscale flows

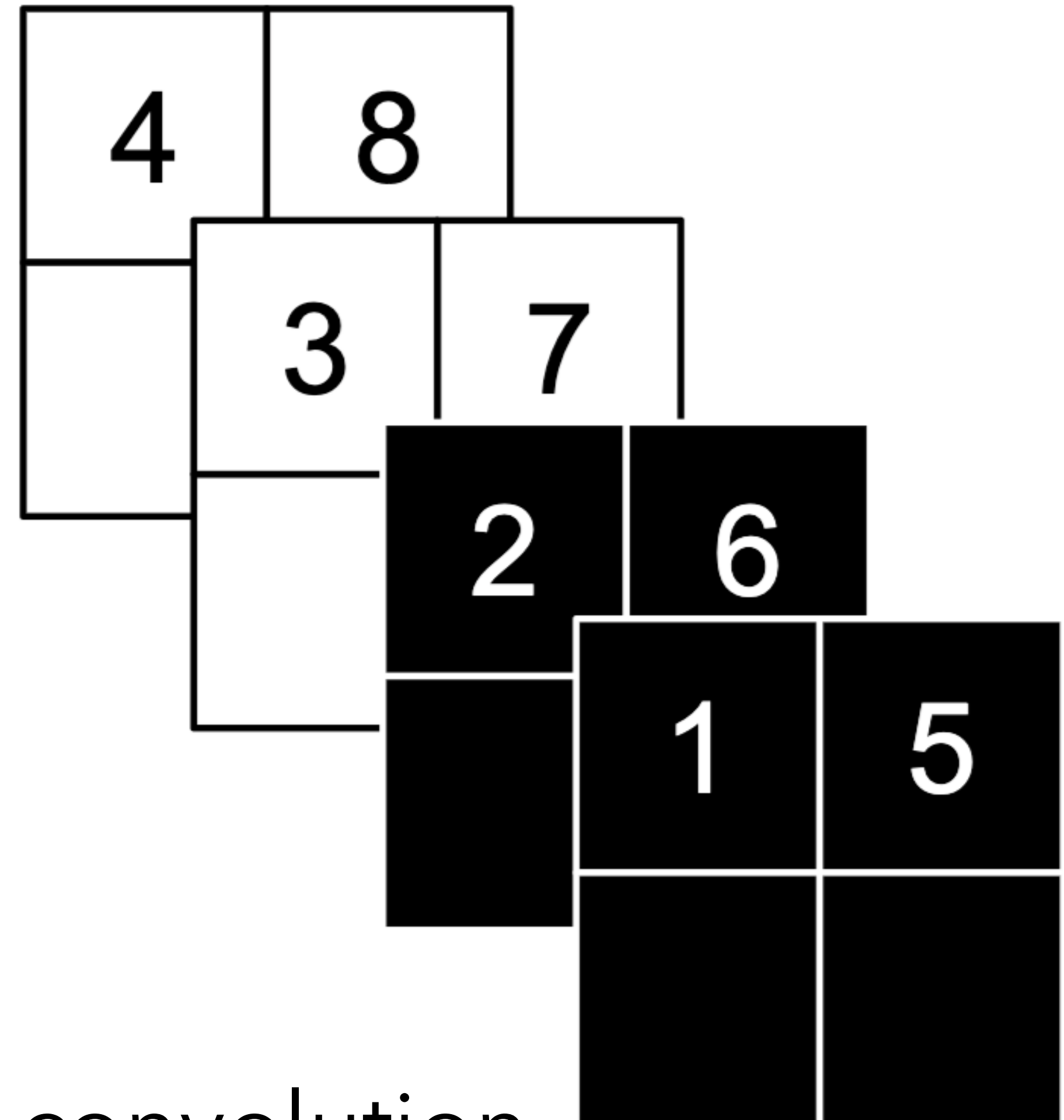


Normalizing flows force same dimension. But we can just stop using subsets of dimensions as we go!

Masked convolution



Mask features before convolution



Masked convolution

1	2	5	6
3	4	7	8

$$\mathbf{y} = \mathbf{b} \odot \mathbf{x} + (\mathbf{1} - \mathbf{b}) \odot \left(\mathbf{x} \odot \exp(s(\mathbf{b} \odot \mathbf{x})) + t(\mathbf{b} \odot \mathbf{x}) \right)$$

Masked convolution

1	2	5	6
3	4	7	8

binary mask

$$y = \mathbf{b} \odot \mathbf{x} + (1 - \mathbf{b}) \odot \left(\mathbf{x} \odot \exp(s(\mathbf{b} \odot \mathbf{x})) + t(\mathbf{b} \odot \mathbf{x}) \right)$$

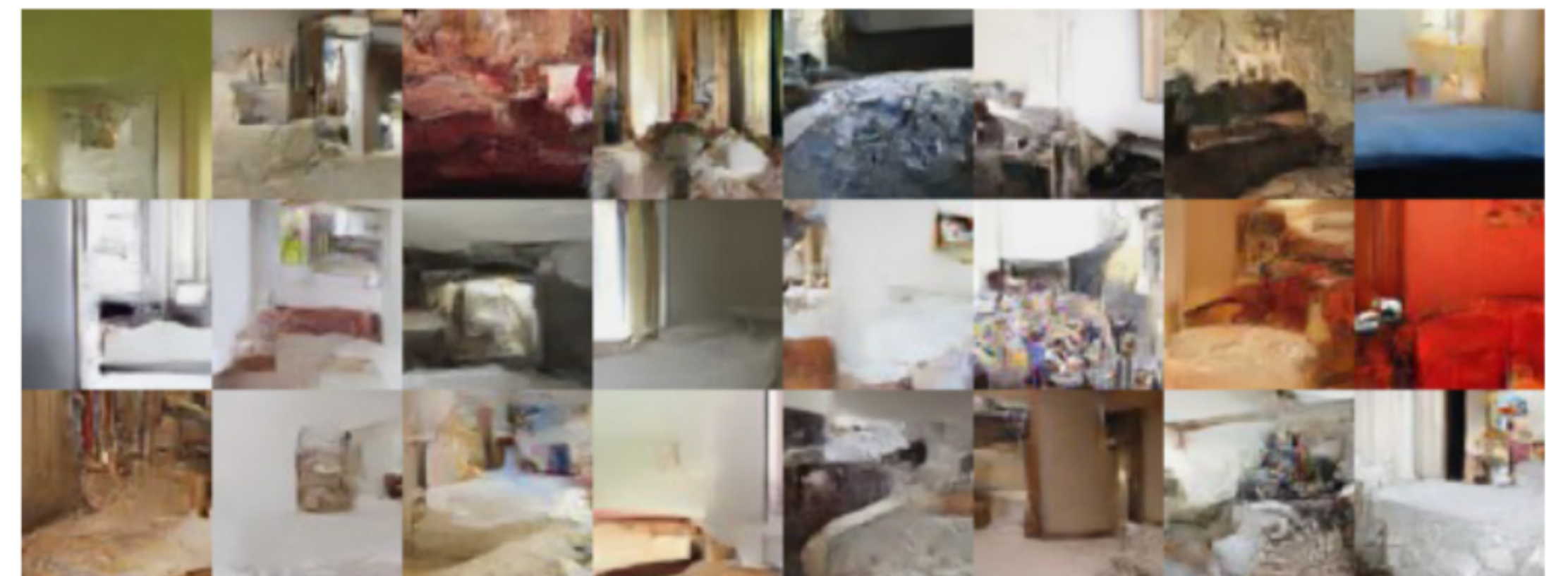
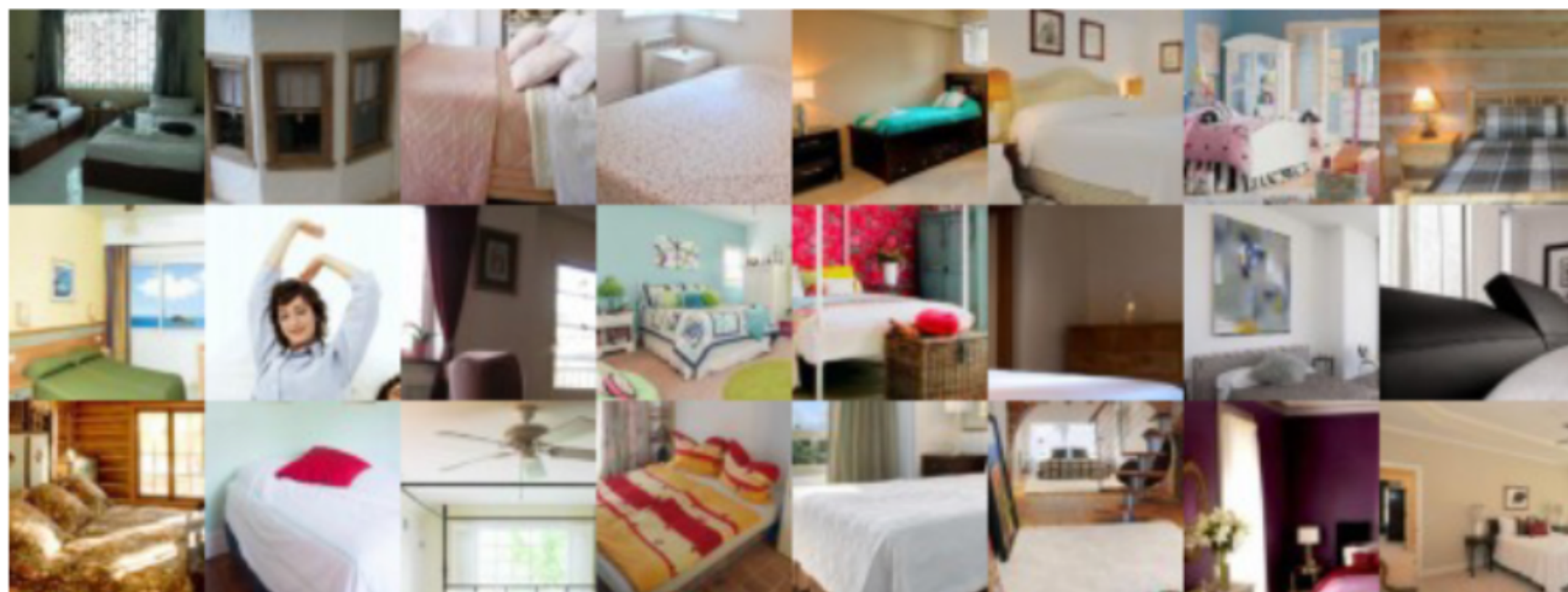
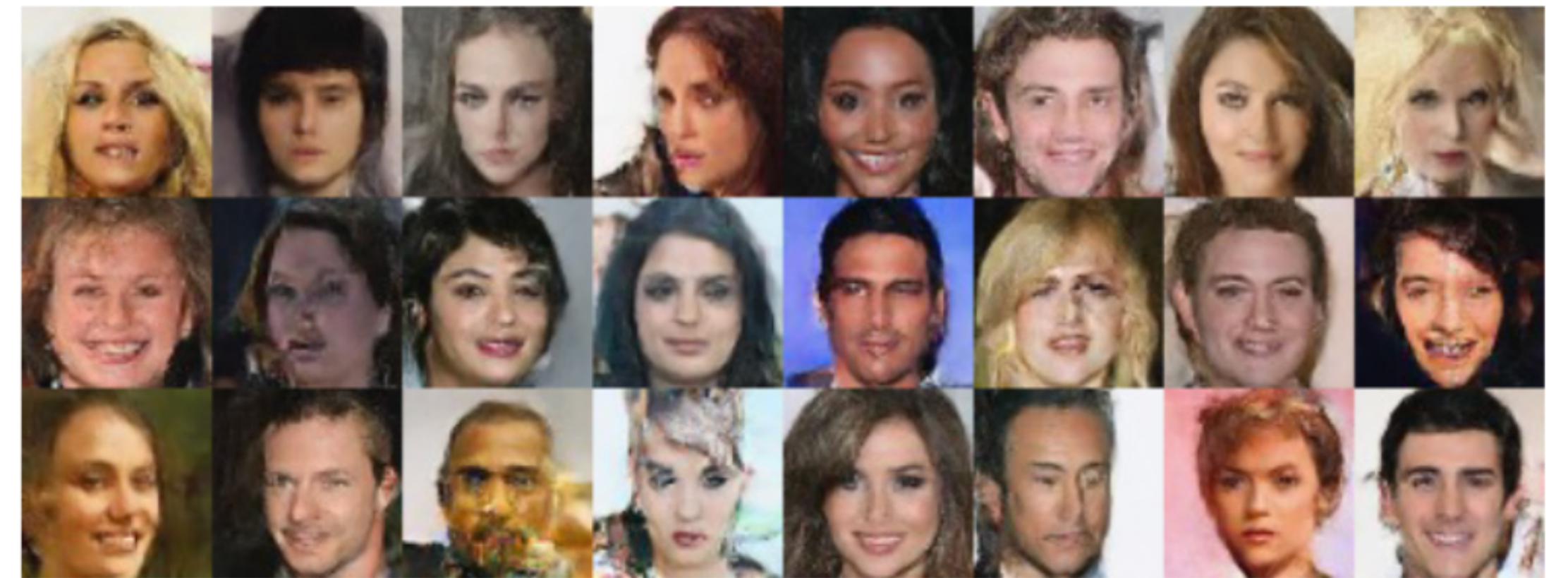
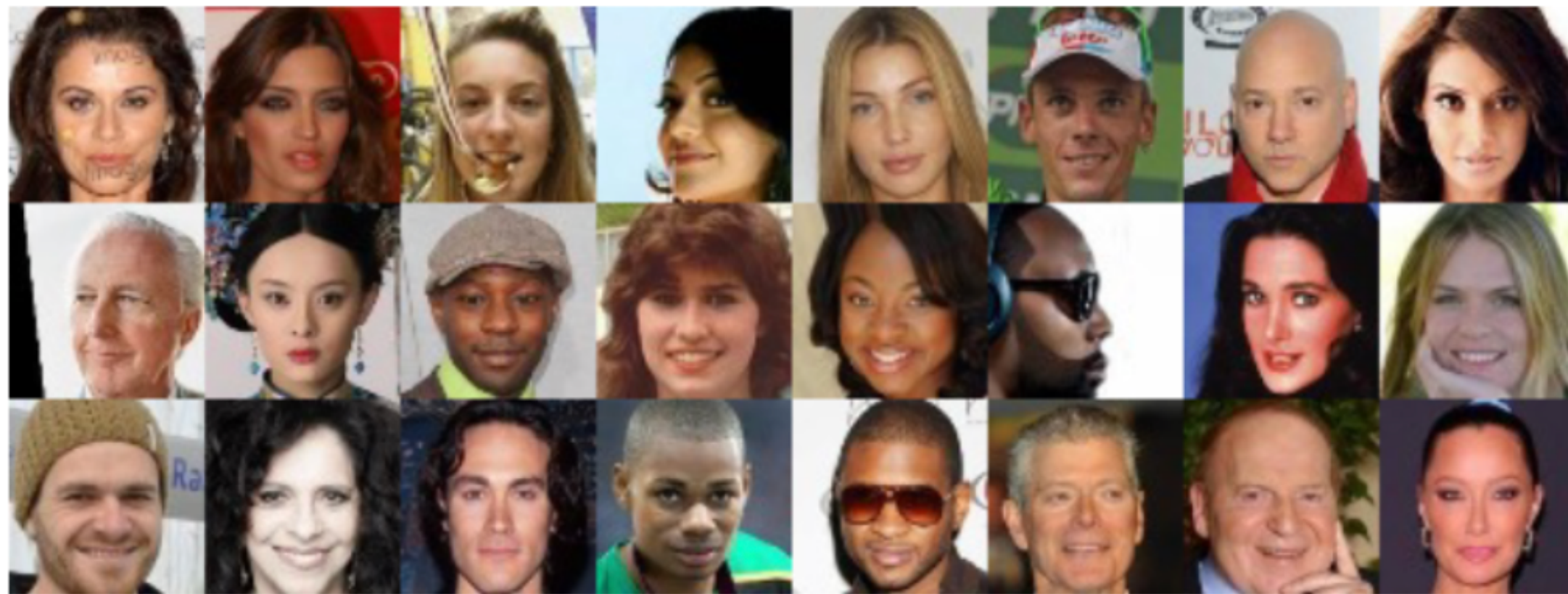
Masked convolution

1	2	5	6
3	4	7	8

convolutional nets

$$y = \mathbf{b} \odot \mathbf{x} + (\mathbf{1} - \mathbf{b}) \odot \left(\mathbf{x} \odot \exp(\mathbf{s}(\mathbf{b} \odot \mathbf{x})) + \mathbf{t}(\mathbf{b} \odot \mathbf{x}) \right)$$

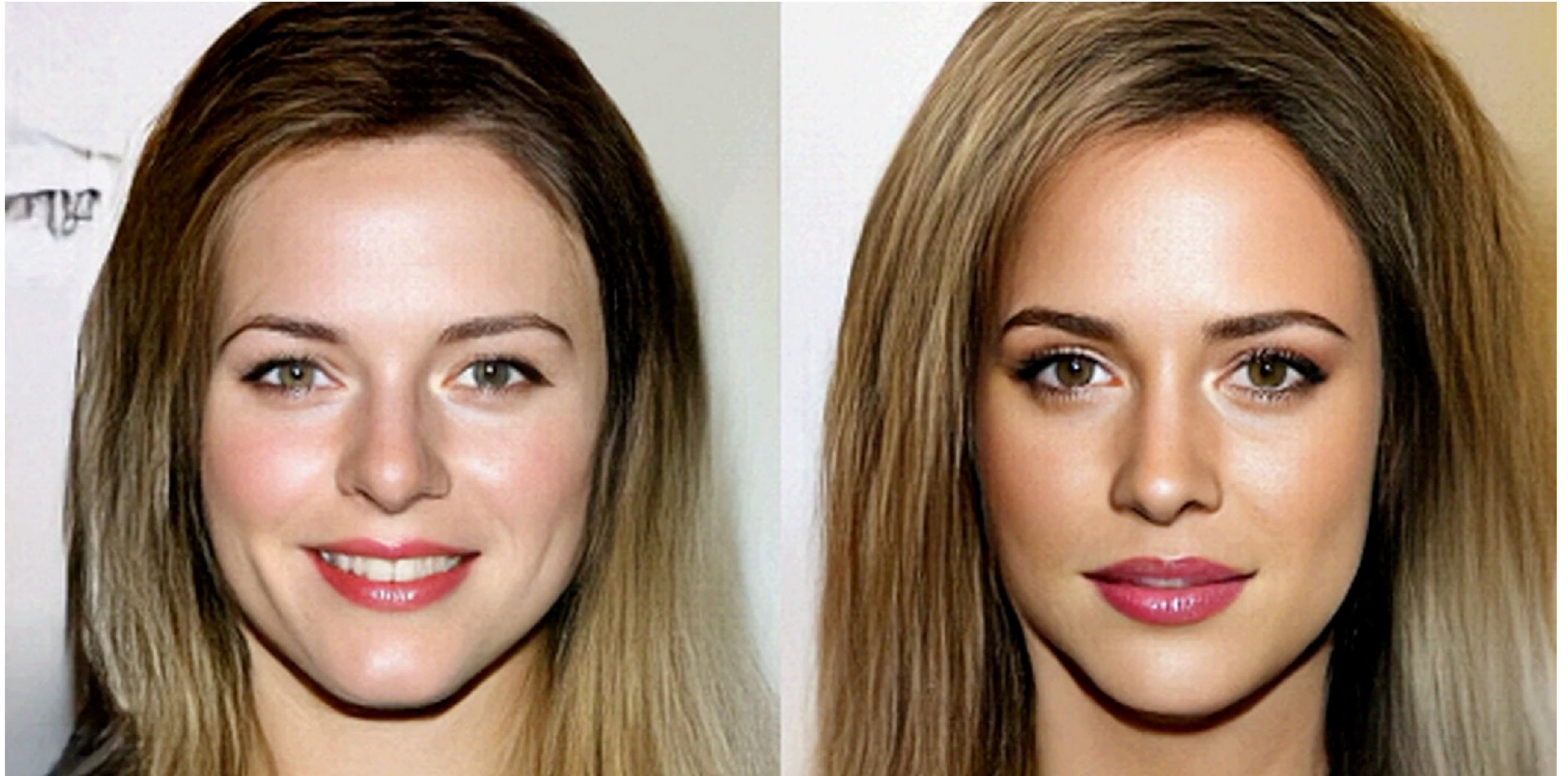
Generated images



Real

Generated

Generated images



Improved model [Kingma and Dhariwal, "GLOW", NeurIPS 2018]

Source: M. Brubaker & Köthe

Normalizing flows vs. VAEs

- Normalizing flows give exact log likelihood, while VAEs allow you to approximately estimate a lower bound.
- Architecture for VAE is much less constrained.
 - Normalizing flow layers must be invertible
 - Must also be efficient to compute determinant of Jacobian
- VAEs typically reduce dimensionality in encoder, whereas normalizing flow must preserve input size throughout.

Next class: Generative adversarial networks