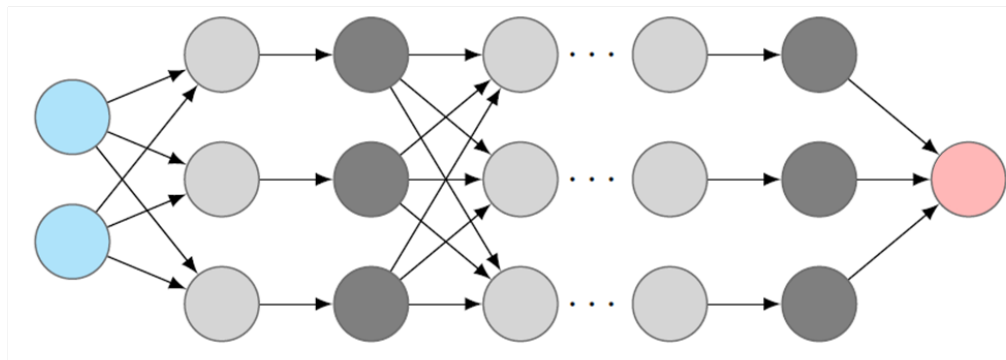
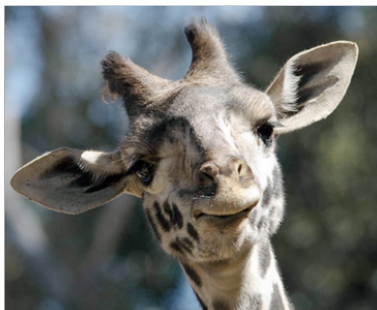


A brief look at generative models

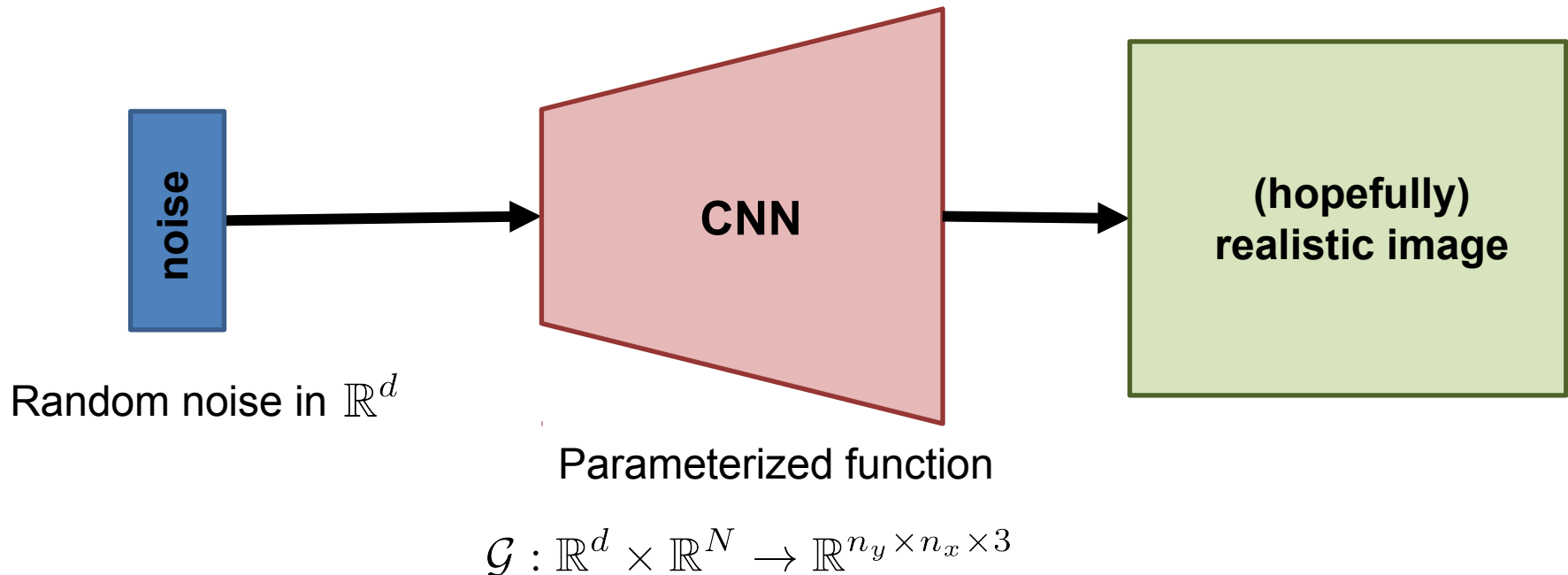
Lecturer: Michael Möller – michael.moeller@uni-siegen.de

Exercises: Hartmut Bauermeister – hartmut.bauermeister@uni-siegen.de

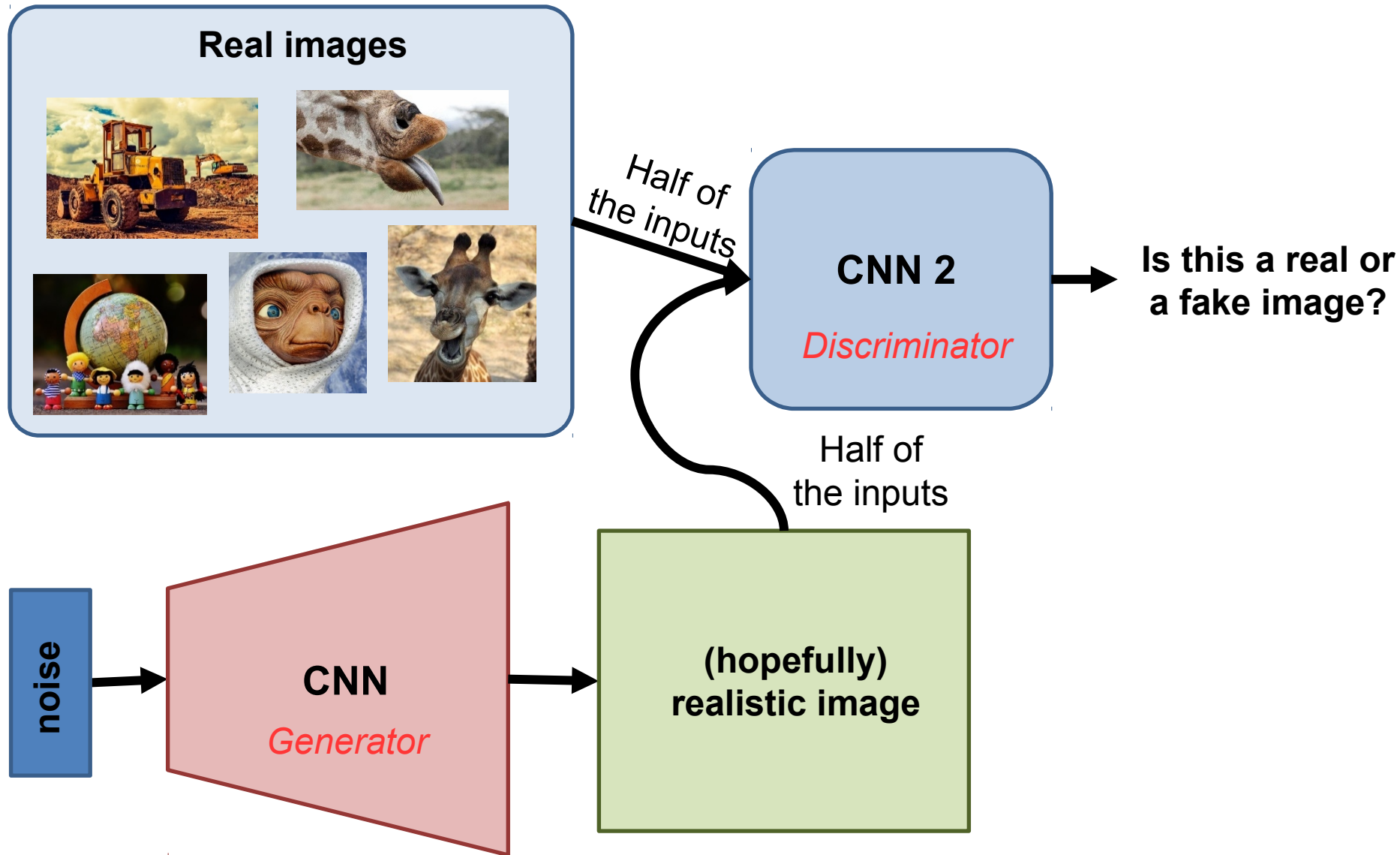


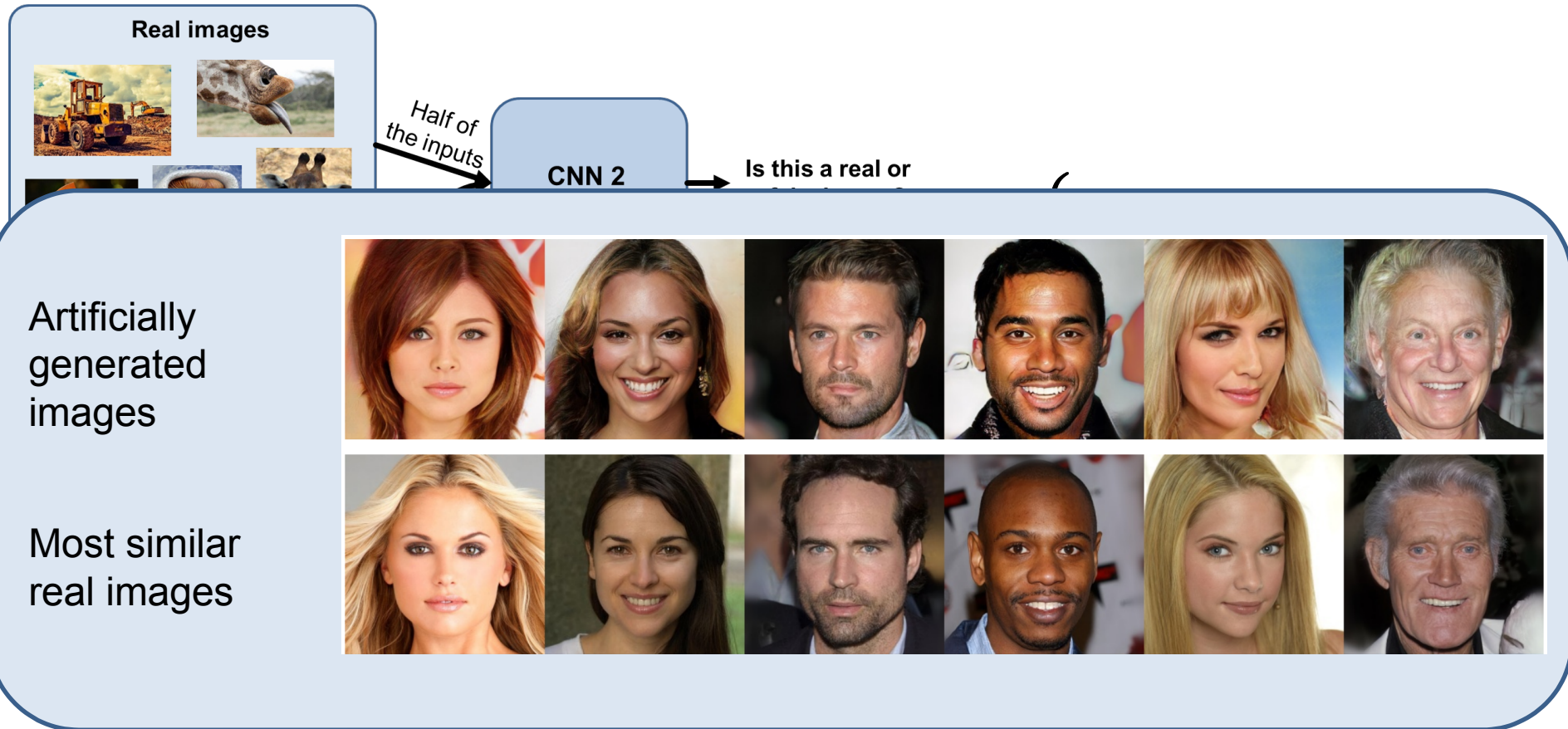
How can we generate new images that look similar to the ones we have in some training data set?

Supervised Deep Learning is a function approximation problem! Assume there is a function mapping Gaussian noise to some realistic image.



But the loss function / training data would be extremely difficult to find!





The discriminator tries to maximize its success to tell real and fake images apart. Simultaneously, the generator tries to minimize the discriminators success.

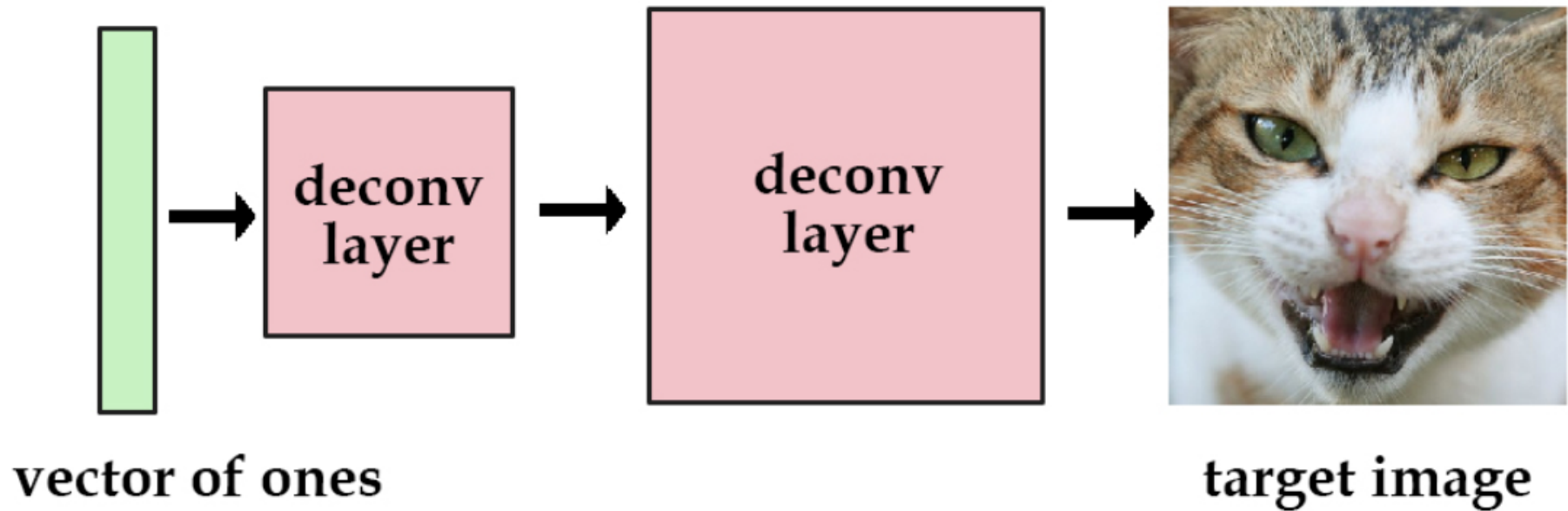
Difficult (abstract) saddle-point problem: $\min_{\theta} \max_{\xi} L(\theta; \xi)$

Drawbacks:

- No direct control over the type of image one generates – the random noise to generate from does not have a particular interpretation.
- Sometimes image-like features are sufficient to fool the discriminator, although the image is quite different from the exemplary images.

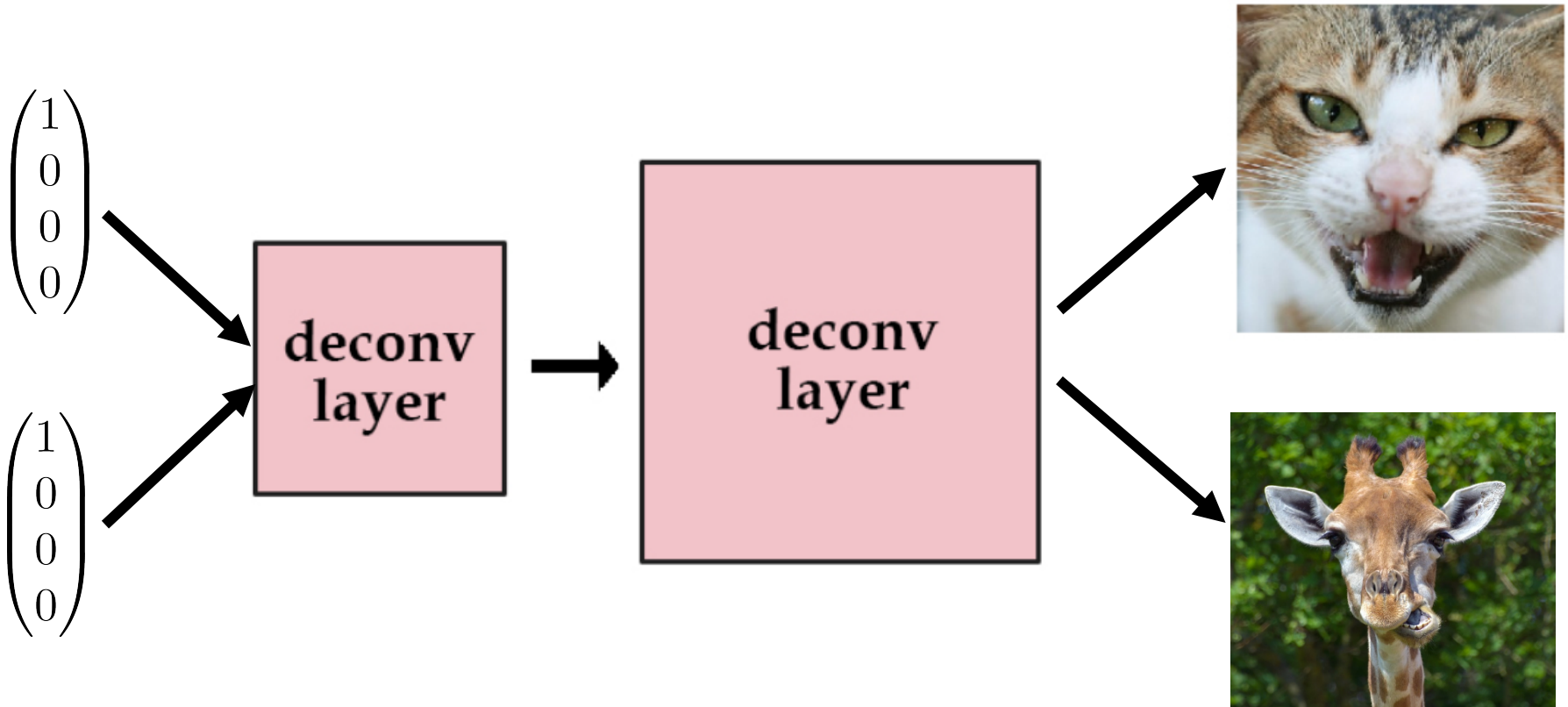


One could hard code an image in the weights of a network, e.g.



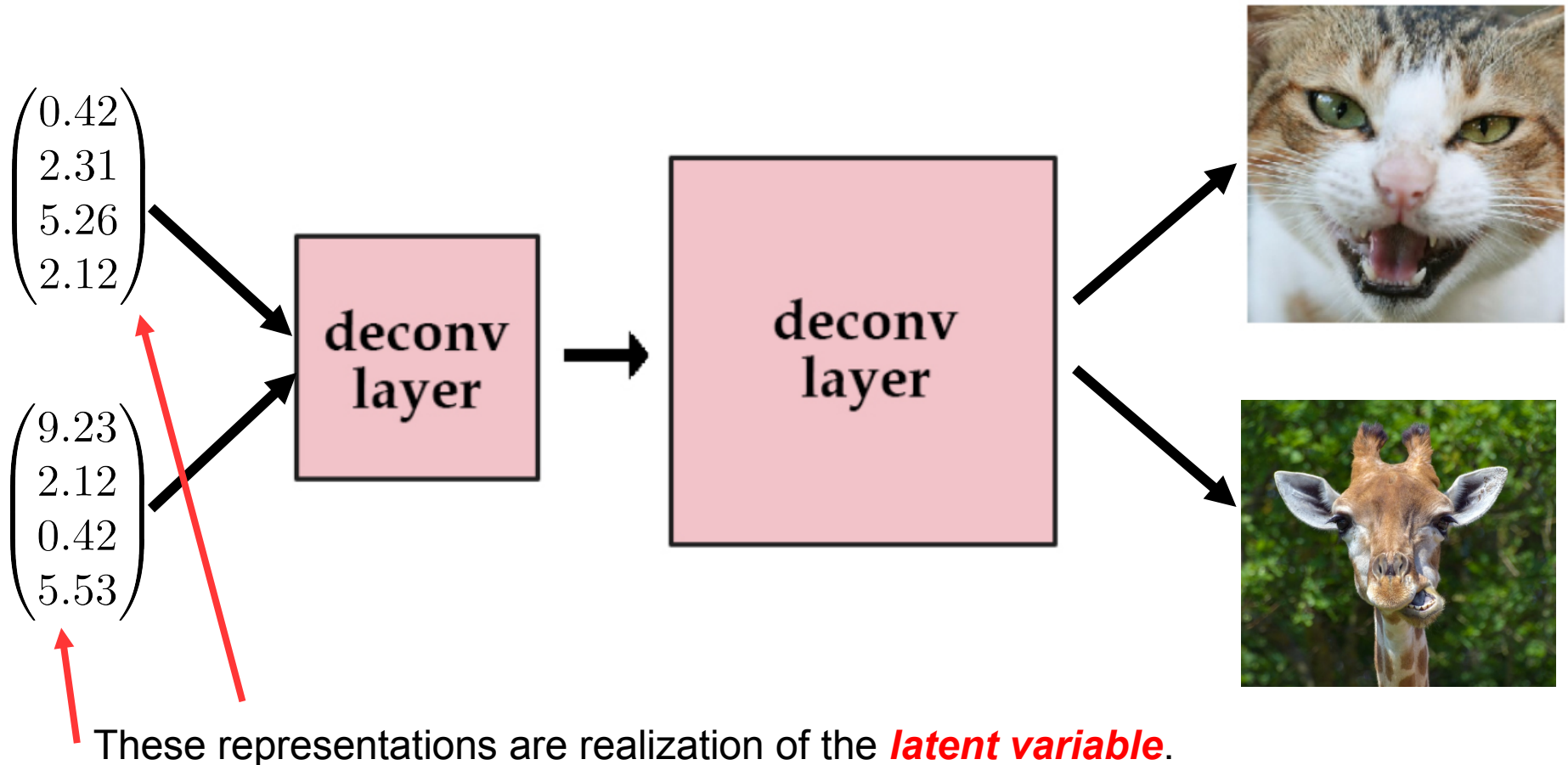
Based on <http://kvfrans.com/variational-autoencoders-explained/>

One could hard code multiple images in the weights of a network, e.g.



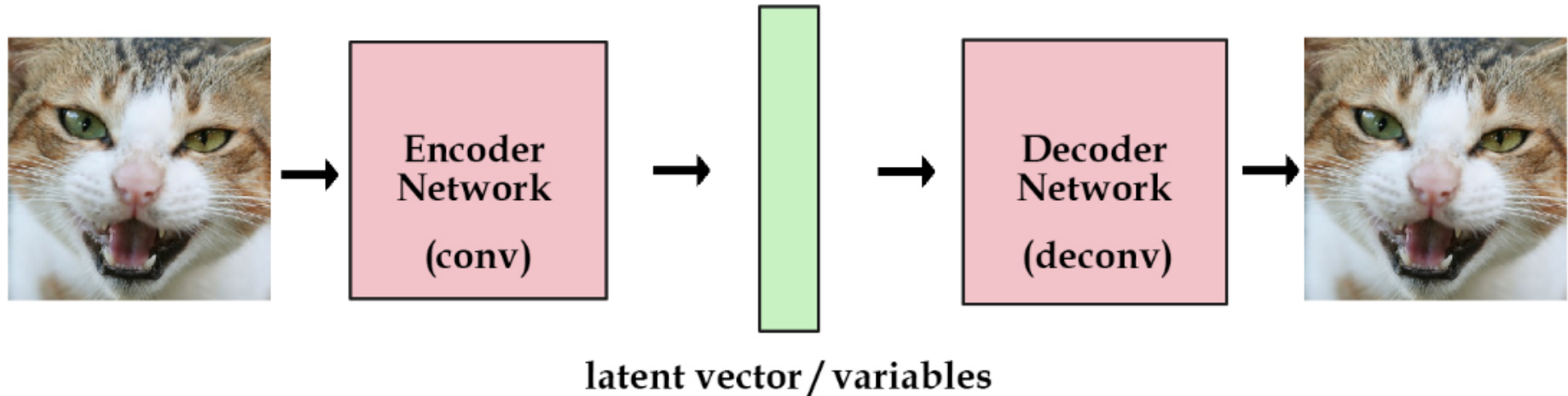
Based on <http://kvfrans.com/variational-autoencoders-explained/>

One could hard code multiple images in the weights of a network, e.g.



Based on <http://kvfrans.com/variational-autoencoders-explained/>

Better idea: Learn the latent representation!

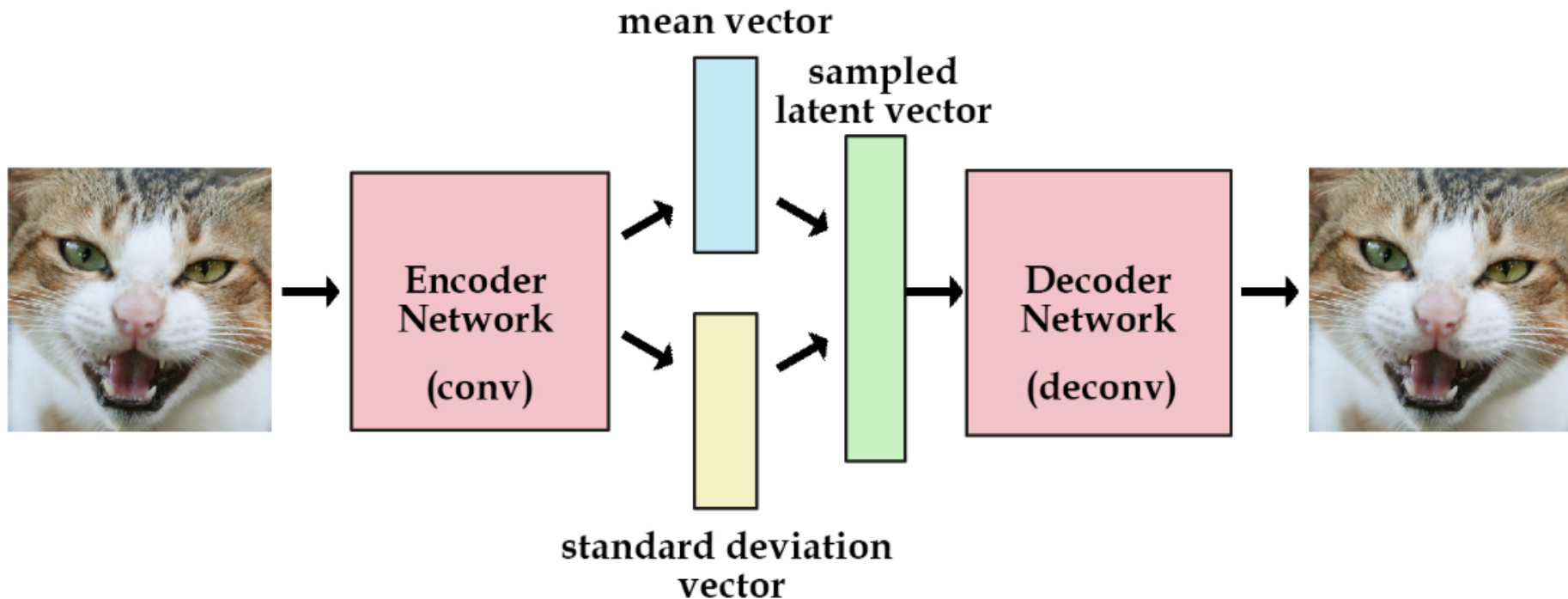


This architecture is called a (standard) autoencoder!

Problem: The latent variables will likely not have a meaning! They could be isolated points!

Based on <http://kvfrans.com/variational-autoencoders-explained/>

Solution: Force the latent variables to roughly follow a unit Gaussian distribution.
Then sample from this distribution to generate new images.

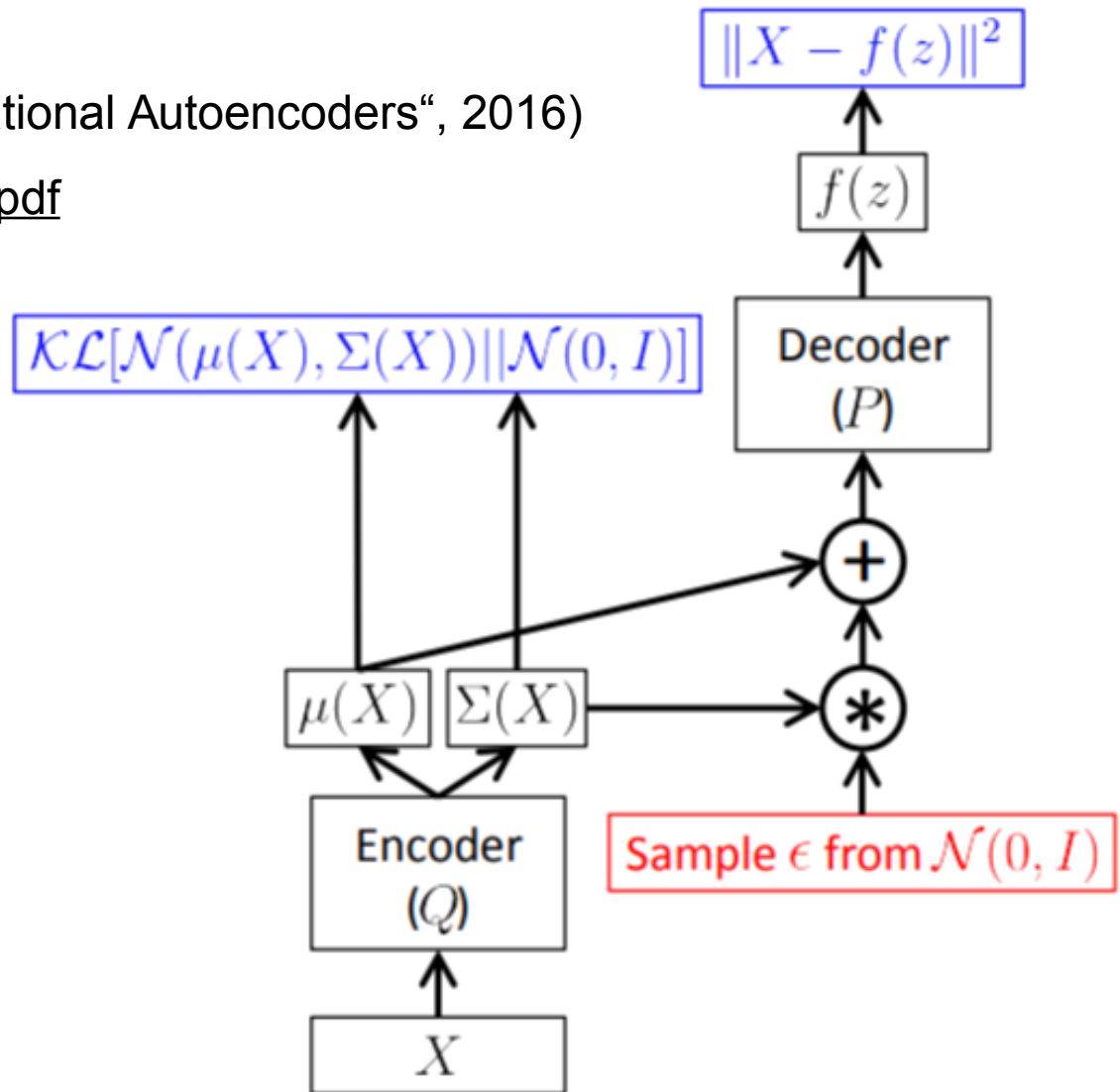


Based on <http://kvfrans.com/variational-autoencoders-explained/>

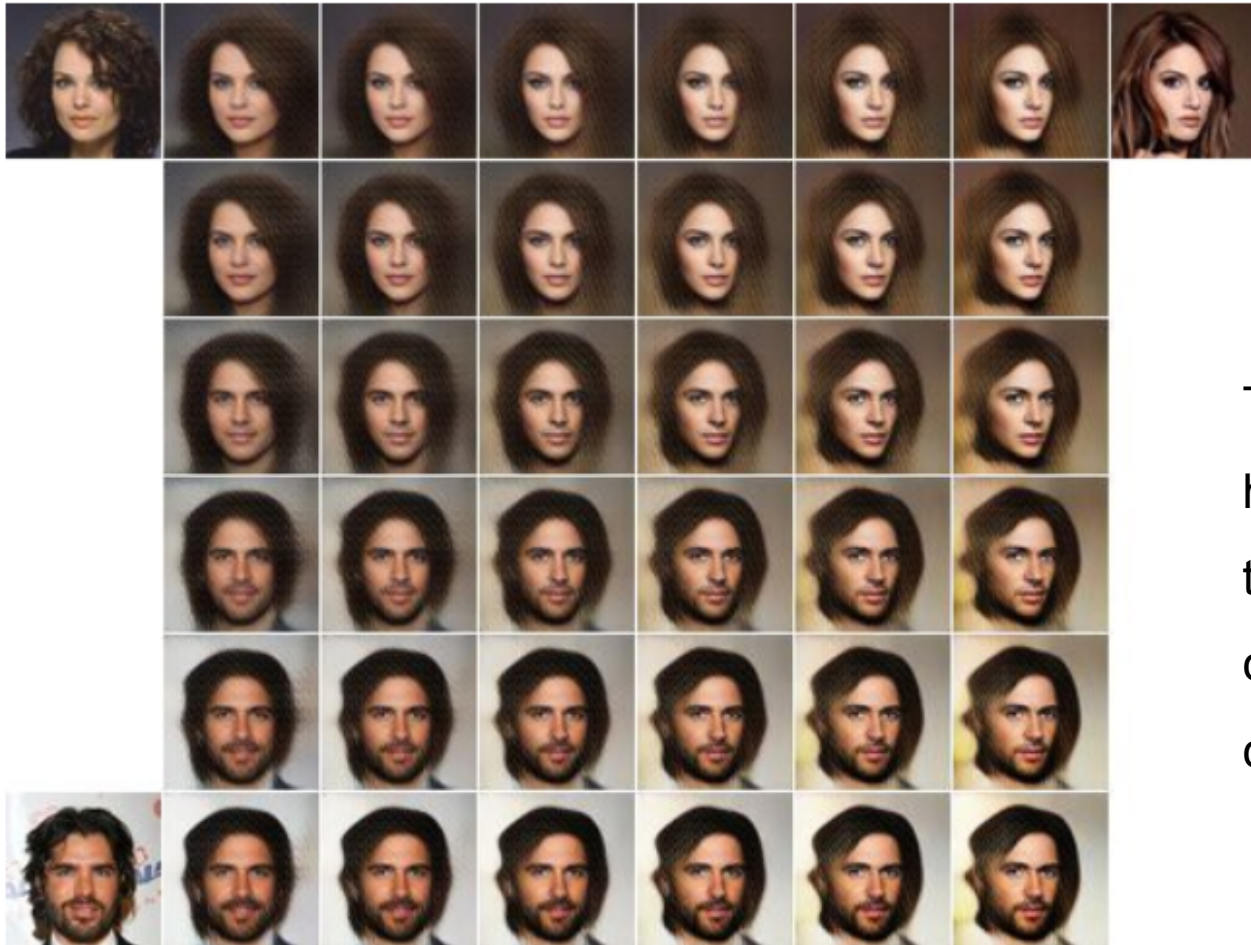
A little more details

(from Doersch, „Tutorial on Variational Autoencoders“, 2016)

<https://arxiv.org/pdf/1606.05908.pdf>



Linear interpolation and arithmetics in latent space are meaningful!



The results of VAEs are, however, often more blurry than GANs. Various combinations exist in the current literature.

<https://hackernoon.com/latent-space-visualization-deep-learning-bits-2-bd09a46920df>