

Machine Learning and Physics: A Faustian Bargain?

Jeff Byers, NRL/UniPd
Machine Learning at GGI
September 14, 2022



ML: DALL-E 2/OpenAI Text2Image Generator

INPUT_STRING="A photo of physicists discussing machine learning in Florence Italy"



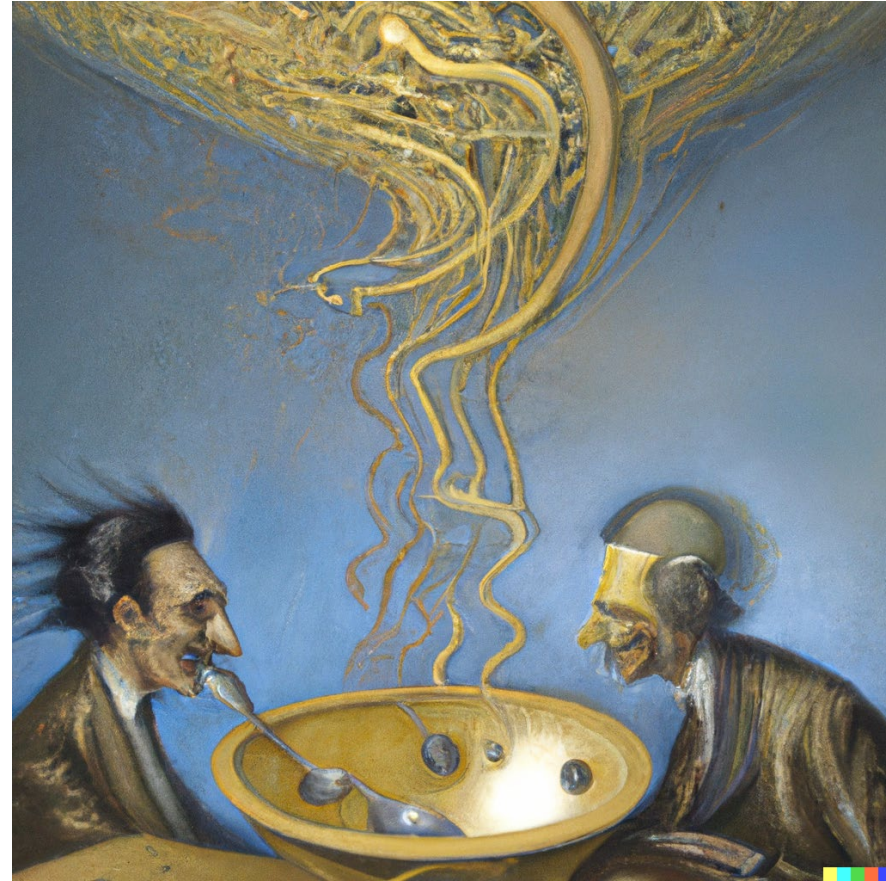
ML: DALL-E 2/OpenAI Text2Image Generator

INPUT_STRING="Impressionist painting of physicists discussing machine learning in Florence Italy"

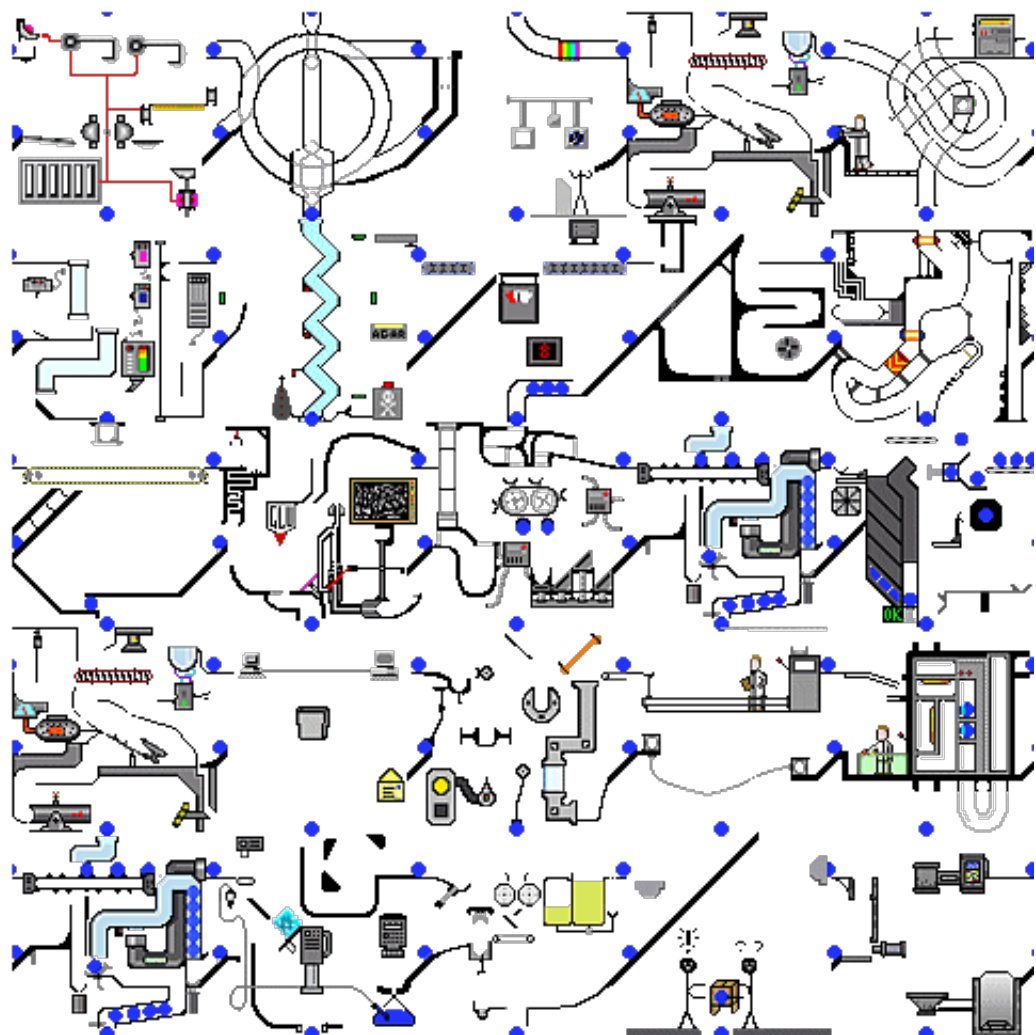


ML: DALL-E 2/OpenAI Text2Image Generator

INPUT_STRING="Dali painting of physicists discussing machine learning in a bowl of pasta"



Is Machine Learning a Faustian Bargain?



Is Machine Learning a Faustian Bargain?

<https://arxiv.org/pdf/2204.06125.pdf>

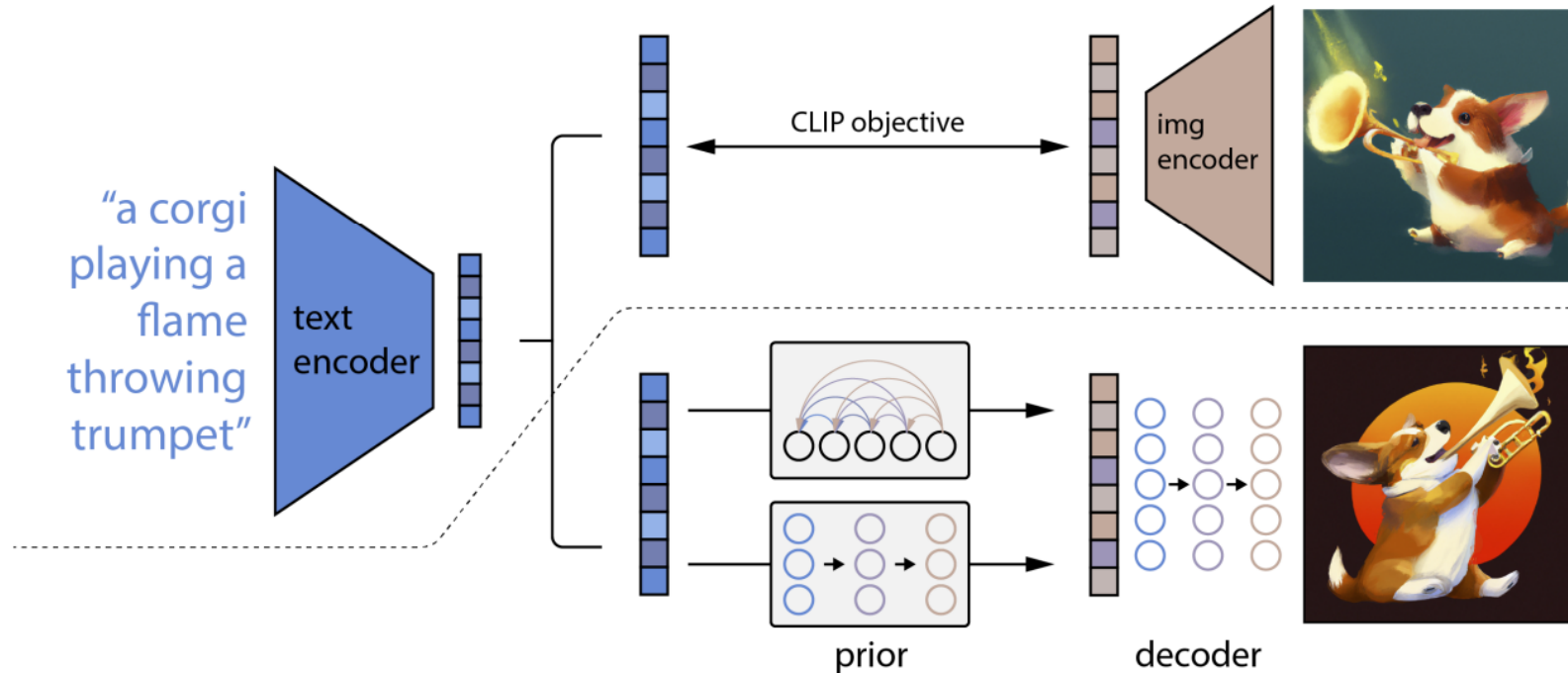
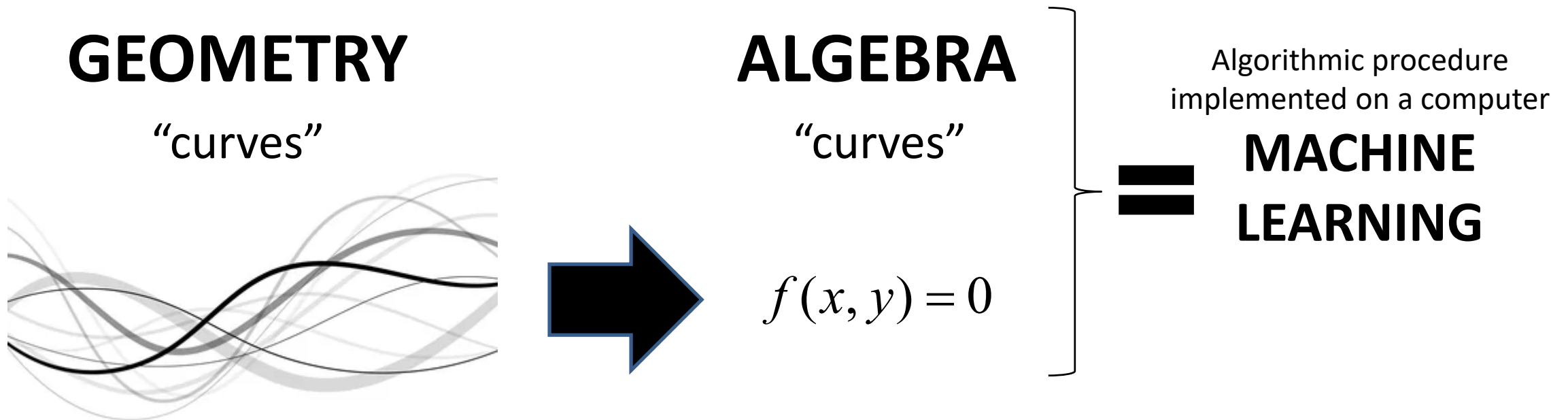


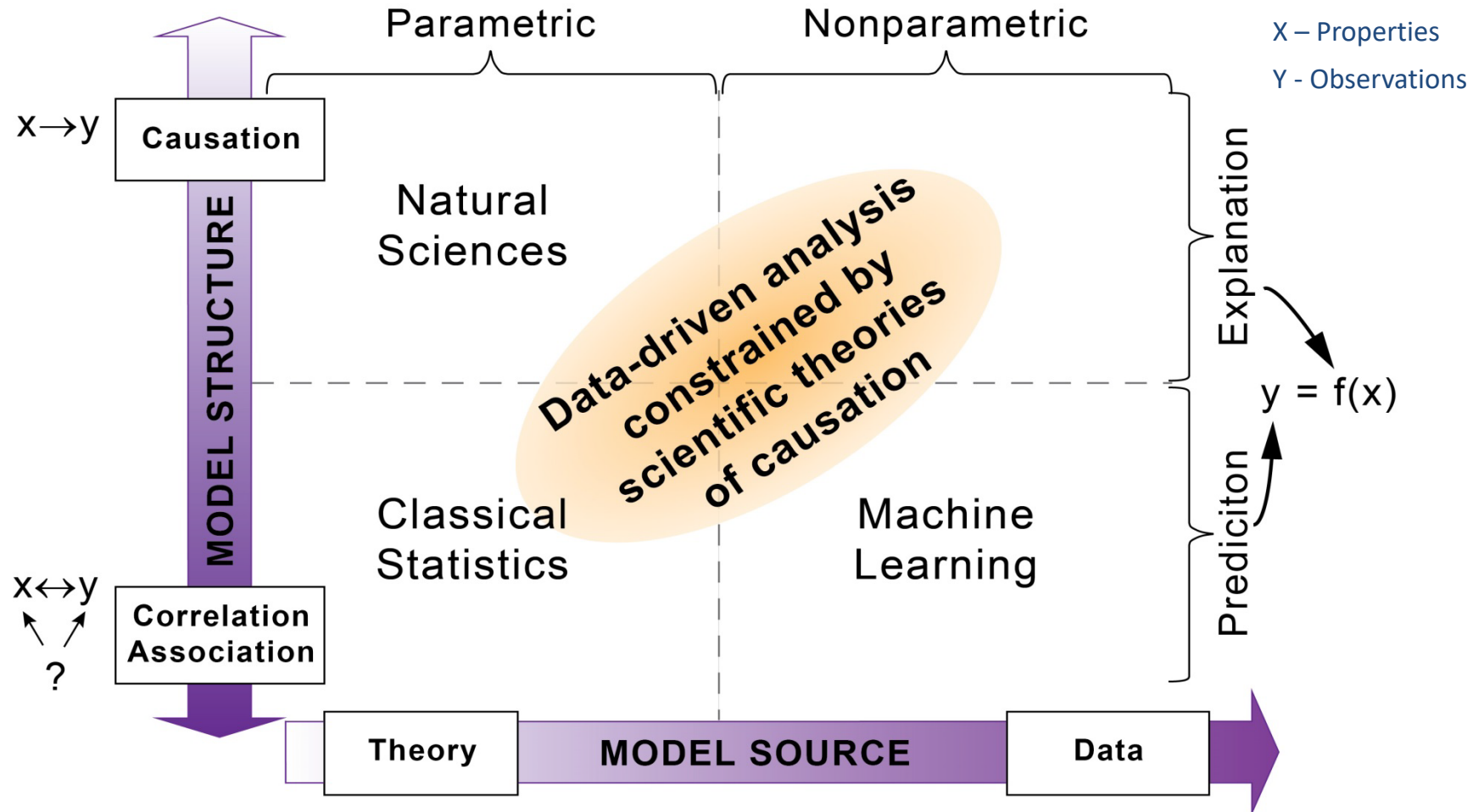
Figure 2: A high-level overview of unCLIP. Above the dotted line, we depict the CLIP training process, through which we learn a joint representation space for text and images. Below the dotted line, we depict our text-to-image generation process: a CLIP text embedding is first fed to an autoregressive or diffusion prior to produce an image embedding, and then this embedding is used to condition a diffusion decoder which produces a final image. Note that the CLIP model is frozen during training of the prior and decoder.

Is Machine Learning a Faustian Bargain?

“Algebra is the offer made by the devil to the mathematicians. The devil says: I will give you this powerful machine, it will answer any question you like. All you need to do is give me your soul: give up geometry and you will have this marvelous machine.” – Michael Atiyah

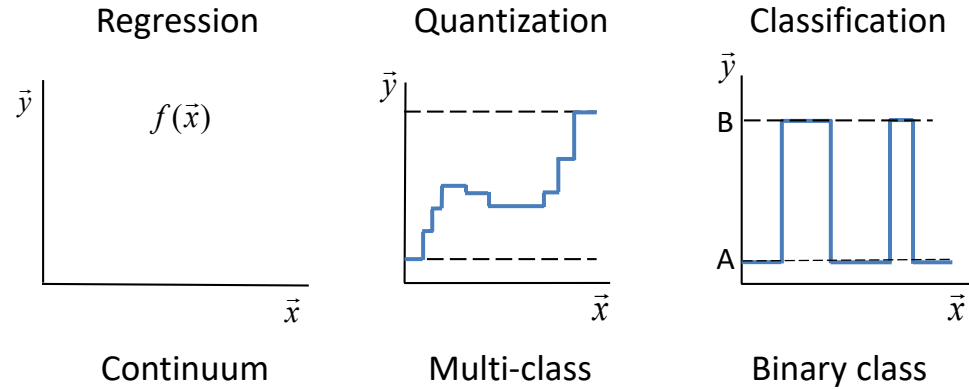


Physics and Machine Learning

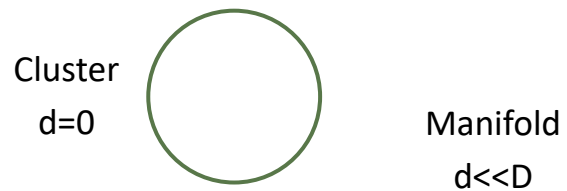


Supervised vs. Unsupervised Learning

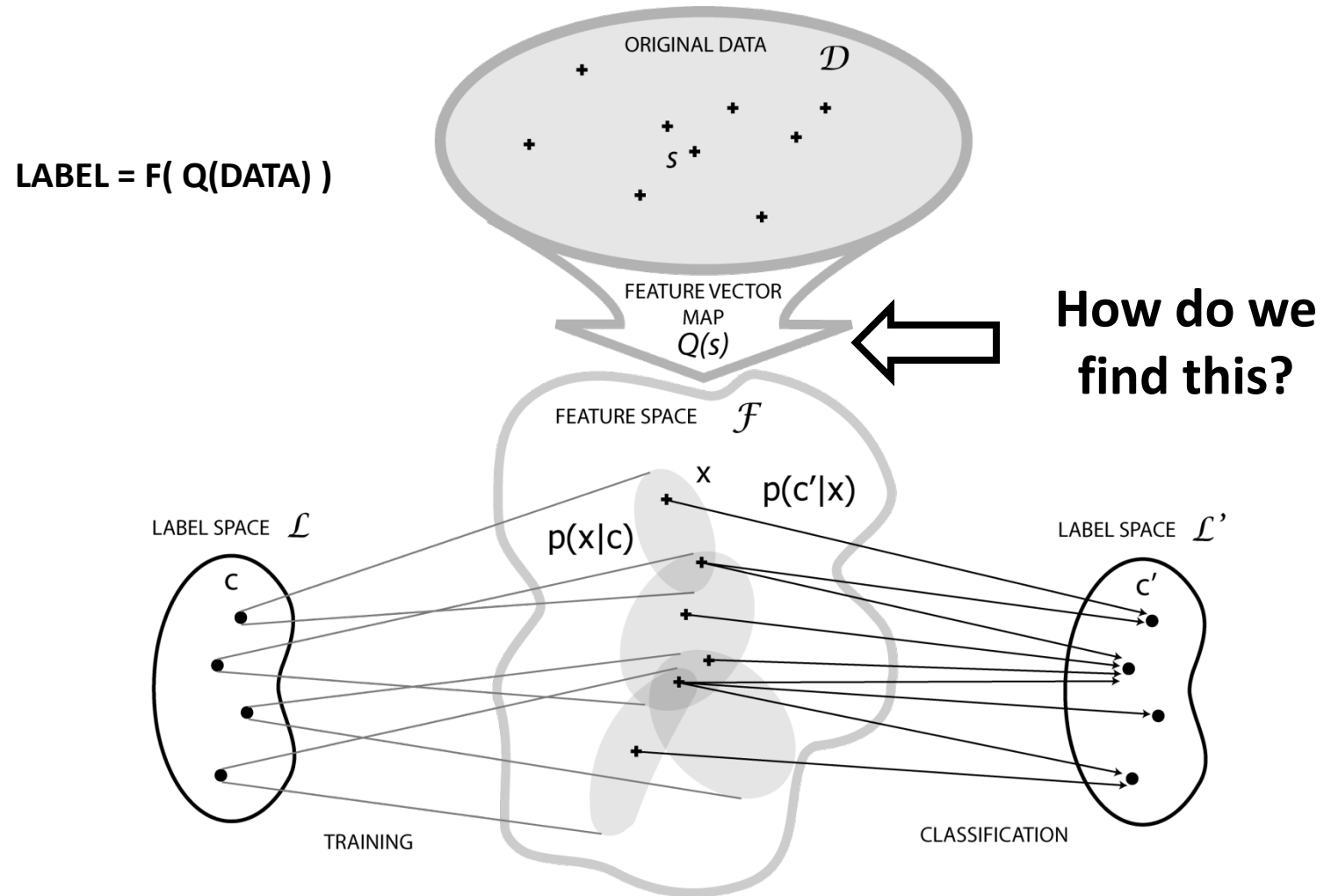
Supervised Learning



Unsupervised Learning



Supervised Learning



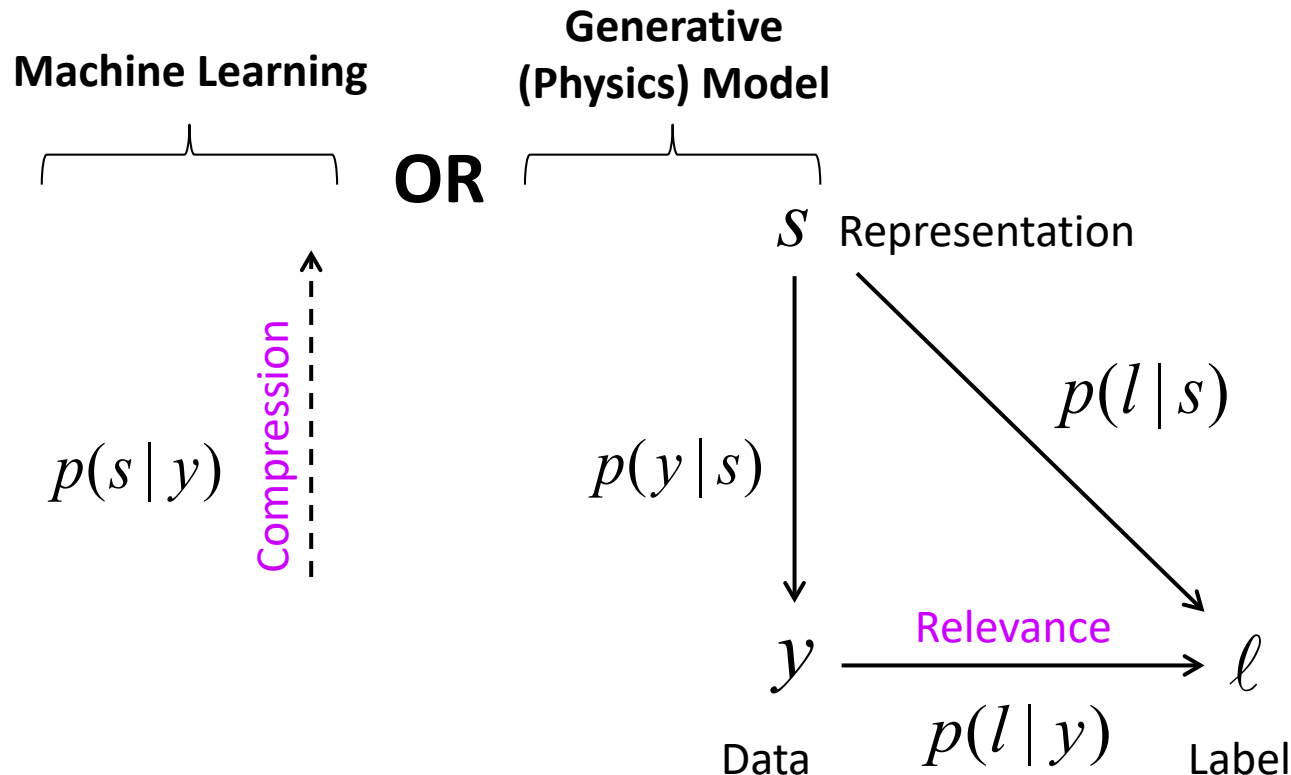
Discriminative Supervised Learning uses geometric objects as boundaries: Hyper-planes, hyper-boxes, hyper-spheres

Information Bottleneck

Goal: *Compress* the data as much as possible while retaining the *relevant* information.

Formal description:

$$p^*(s | y) = \arg \max_{p(s|y)} \underbrace{I(S, L)}_{\text{Relevance}} - \lambda \underbrace{I(Y, S)}_{\text{Compression}} = \mathcal{F}[p(s | y)]$$



We want to **maximize** $I(S, L)$ to retain **relevance**.

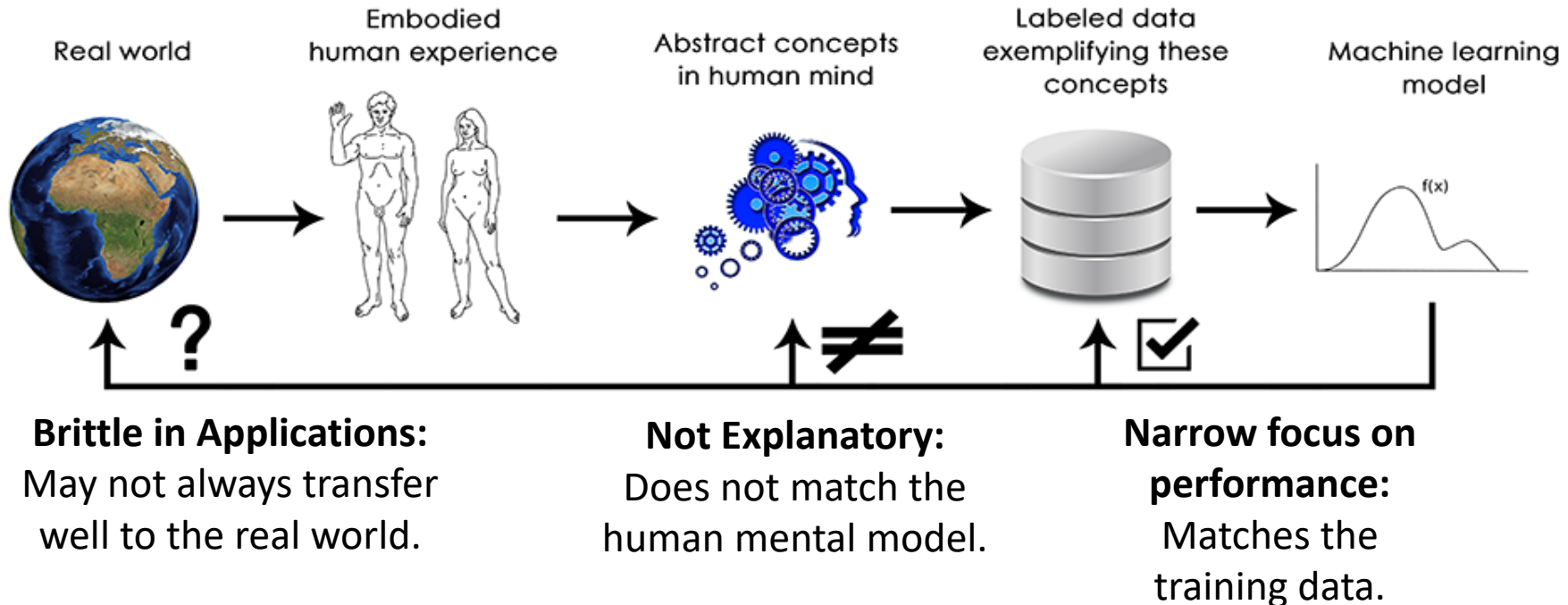
We want to **minimize** $I(Y, S)$,
to **compress** the data transferred
by the Generative model

The parameter λ is the Lagrange multiplier that
tunes between these two competing objectives.

Supervised Learning

What are the issues with labels:

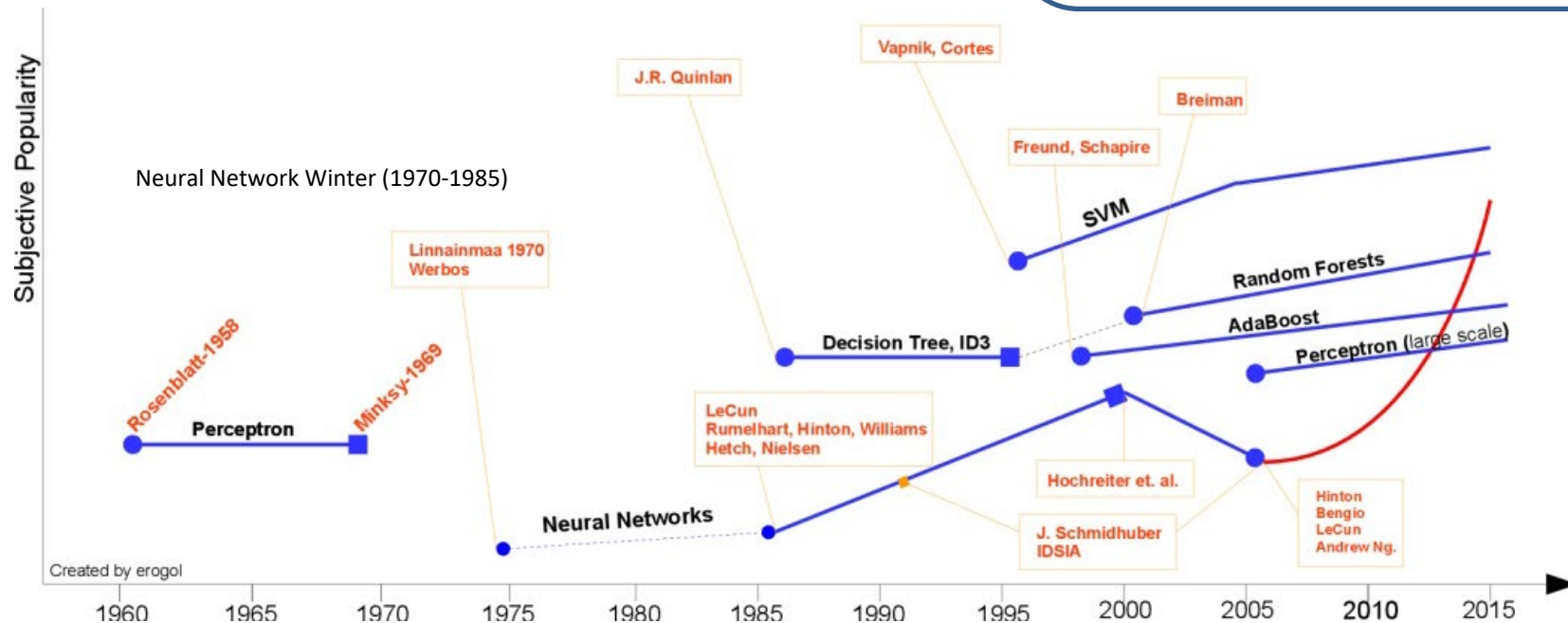
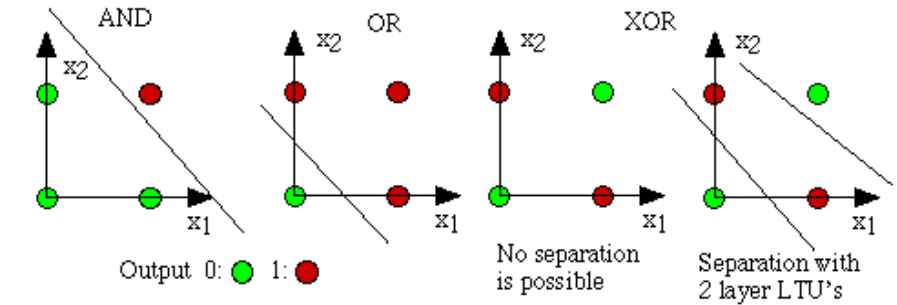
- Labels are **NOT** independent of the context that created them.
- Labels are **EXTREME** data compression preserving relevance in that context.
- Labels are a **SUPERFICIAL** transfer of intelligence into machine learning.



History of Machine Learning

Missing Bayesian techniques, PGM's, statistical techniques (e.g., Naïve Bayes, linear and logistic regression, PCA, k -NN, bootstrap, LASSO, Ridge regression), NLDR, logic and relational techniques, sampling techniques like MCMC, etc. Advancements from signal processing ... sparse dictionaries, regularization techniques, ICA, matrix factorizations (NNMF) Also, reinforcement learning, active learning,

Minsky's **XOR** problem (*Perceptrons*, 1969) and MLP solution:



Backpropagation enables training of multi-layer perceptrons (MLP), ~1985

<http://www.erogol.com/brief-history-machine-learning/>

Learning Dynamics of Neural Networks

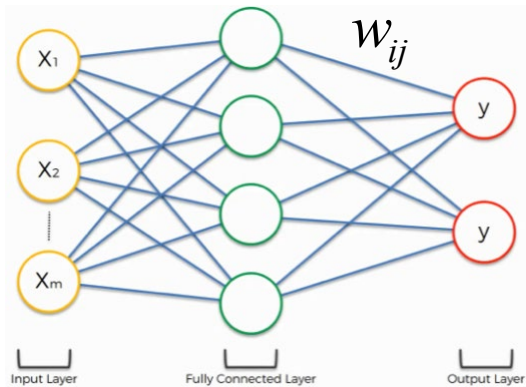
The Illusion of an Optimization Problem

Nonlinear transform of each coordinate of an affine map

$$y^i = \sigma.(w_j^i x^j + b^j)$$

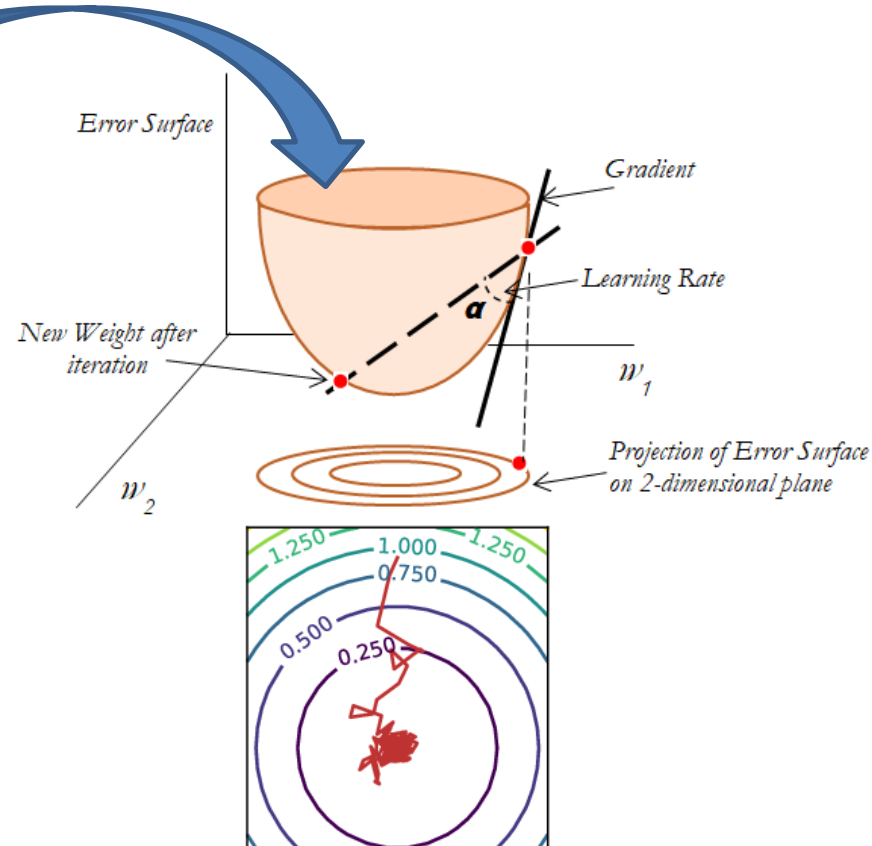
Neural Network

$$\mathbf{y} = f_{\mathbf{w}}(\mathbf{x})$$



Magic: “Inductive Bias” imposed on the network.

Cost Function:
$$C(\mathbf{W}) \approx \frac{1}{N} \sum_{n=1}^N |\mathbf{y}_n - f_{\mathbf{w}}(\mathbf{x}_n)|^2$$



Gradient descent:

Automatic
Differentiation

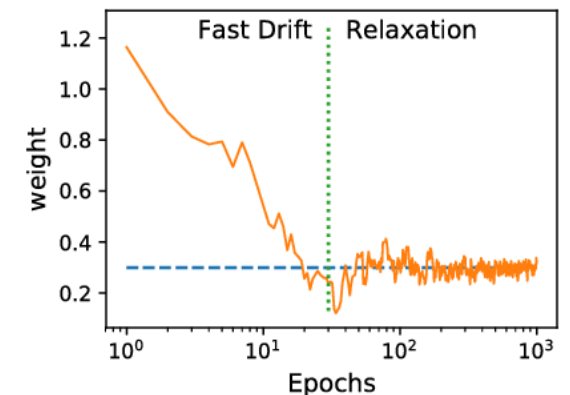
$$w_{ij}(t+1) = w_{ij}(t) - \alpha \frac{\partial C}{\partial w_{ij}}$$

Learning Rate

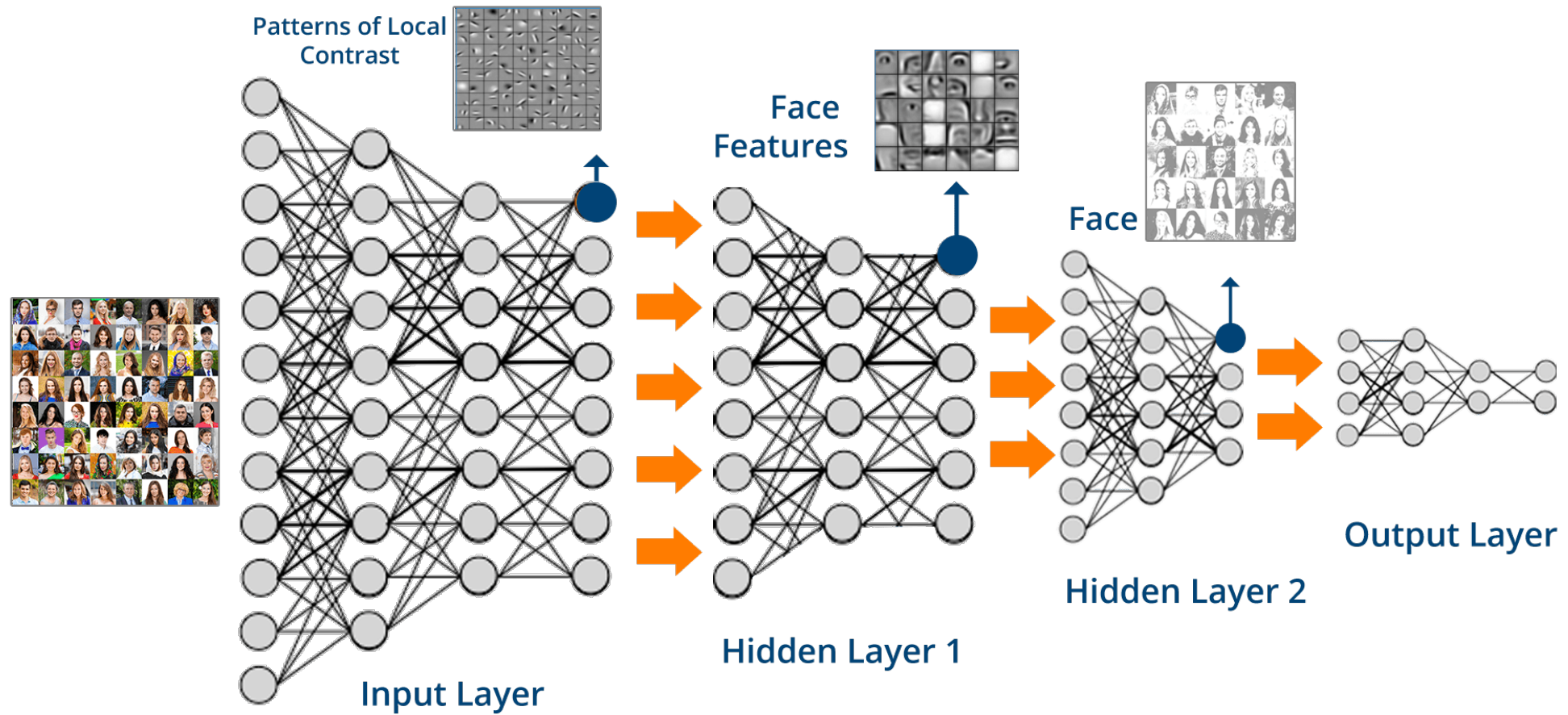
More sophisticated choices
(e.g., ADAM) are possible.

Stochastic Gradient Descent:
Backpropagation on mini-batches

Random sub-sampling of data



Deep Convolutional Neural Networks



How intelligent are neural networks?

Intriguing properties of ~~neural networks~~ ^{airliners}

2014

Christian Szegedy
Google Inc.

Wojciech Zaremba
New York University

Ilya Sutskever
Google Inc.

Joan Bruna
New York University

Dumitru Erhan
Google Inc.

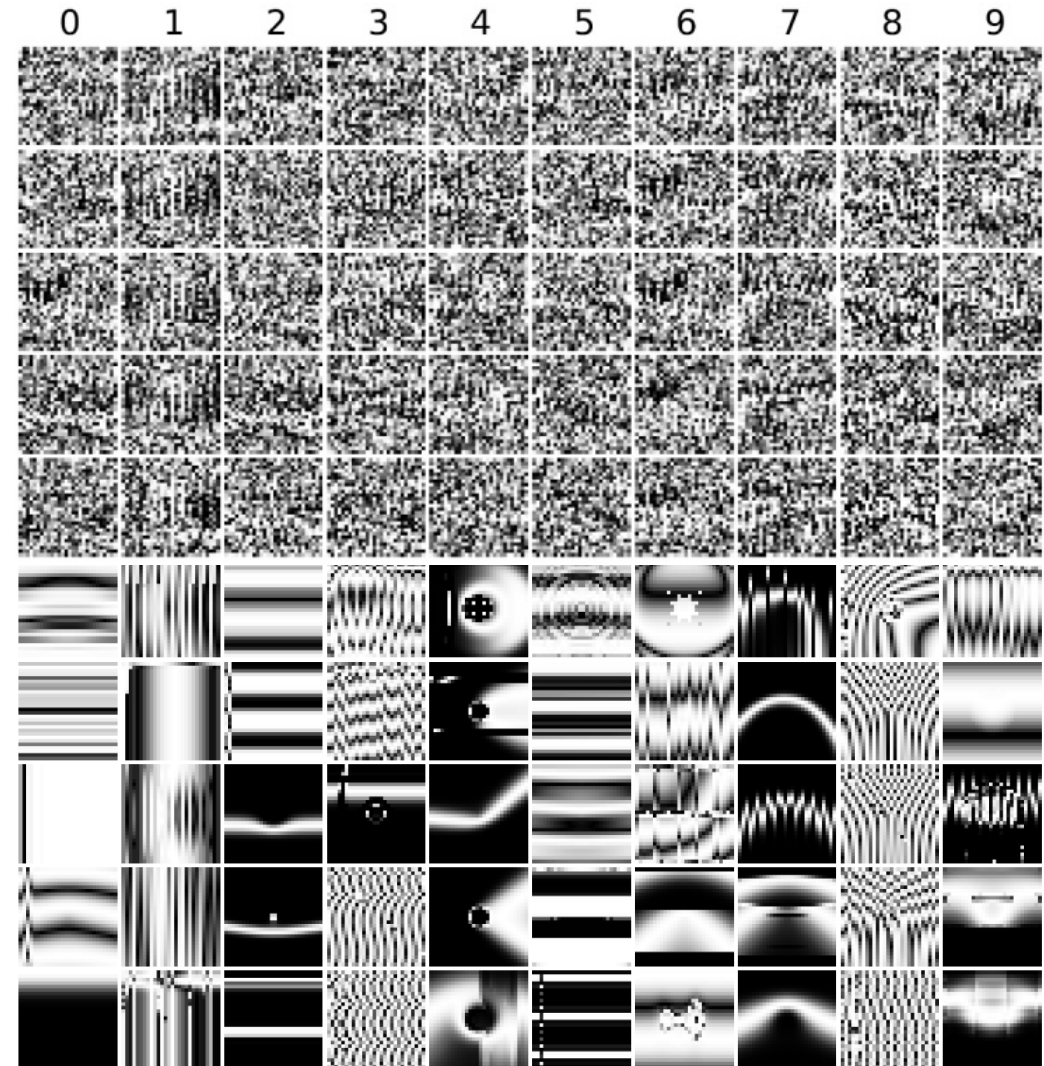
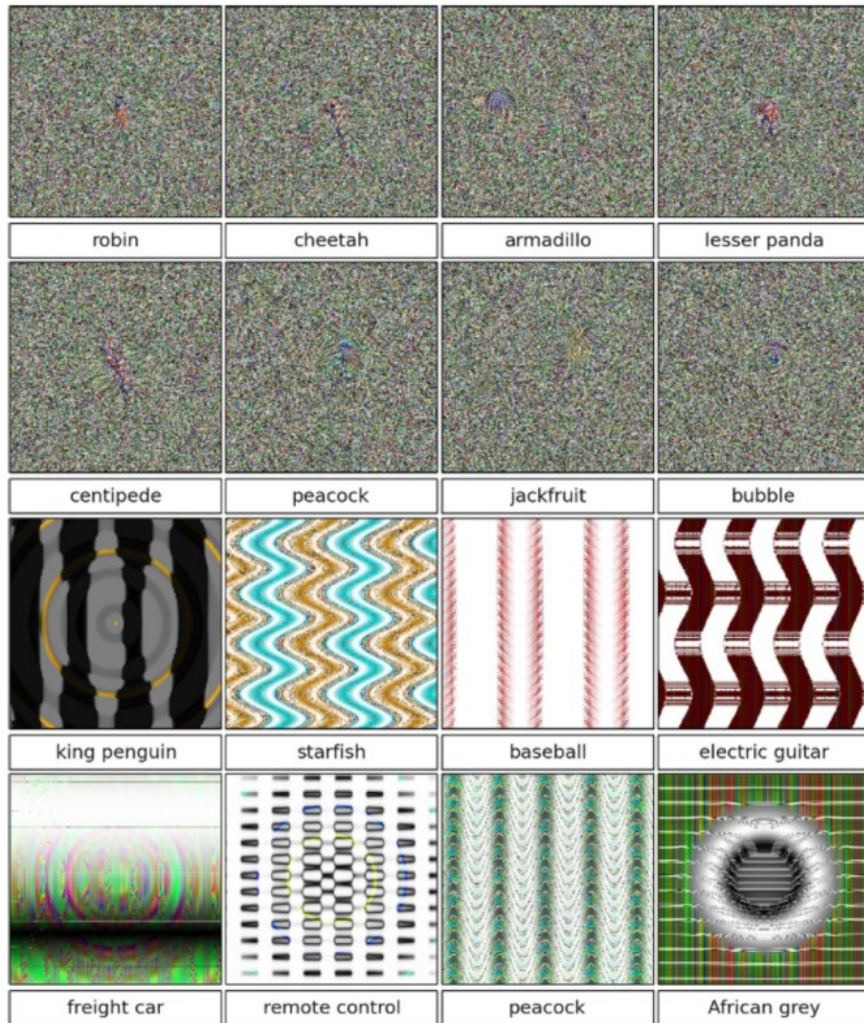
Ian Goodfellow
University of Montreal

Rob Fergus
New York University
Facebook Inc.



Deep Learning: Adversarial Examples

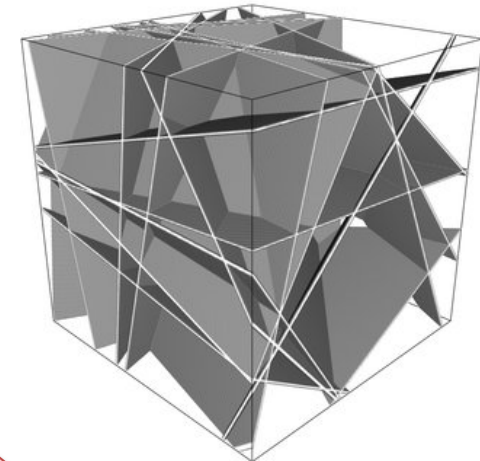
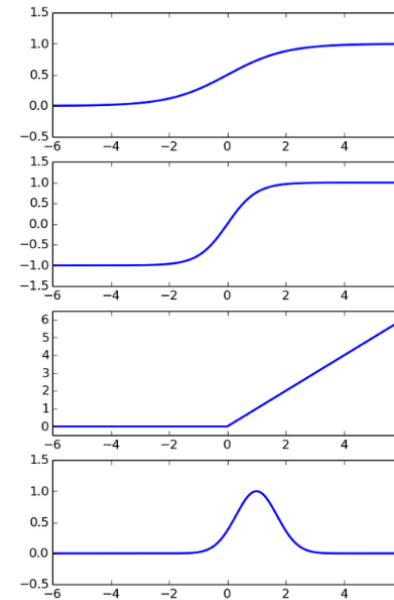
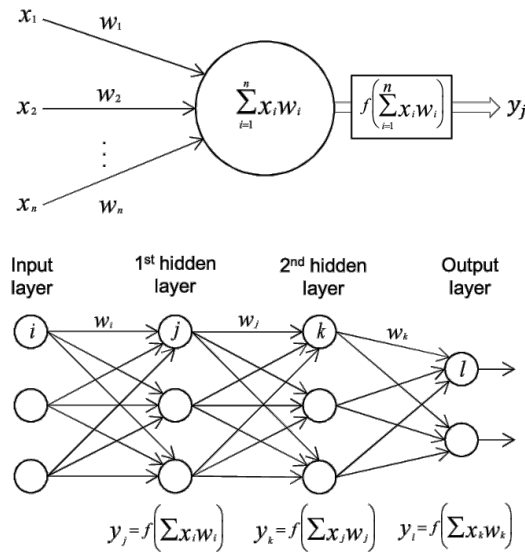
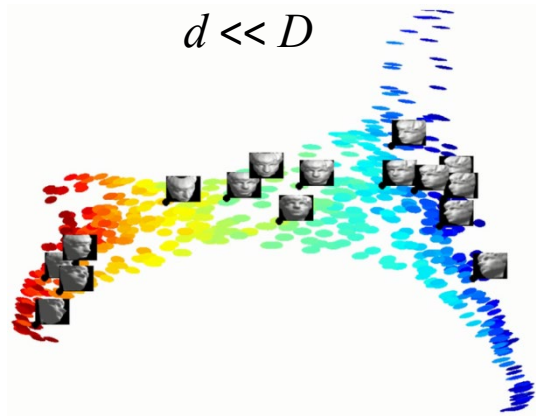
The algorithm is >99.6% confident of these labels



Geometry of Neural Networks

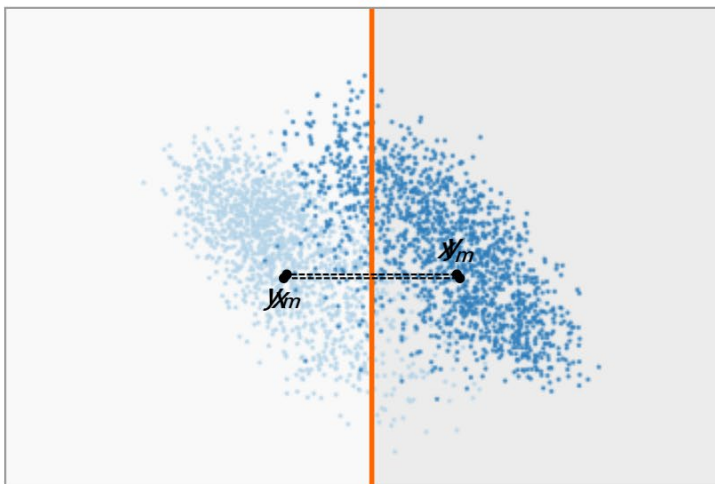
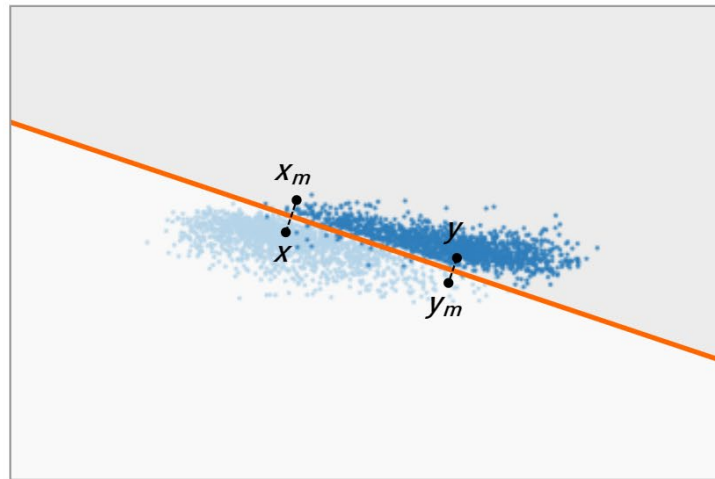
The data lies along a manifold constrained to low-dimensions by the generative mechanism.

However, the Neural Network typically chops up the space with hyperplanes to form non-local decision boundaries for classes.



Exception!

Deep Learning: Adversarial Examples



Real data



Image x



Image y



Image x



Image y

Adversarial data



Mirror image x_m



Mirror image y_m

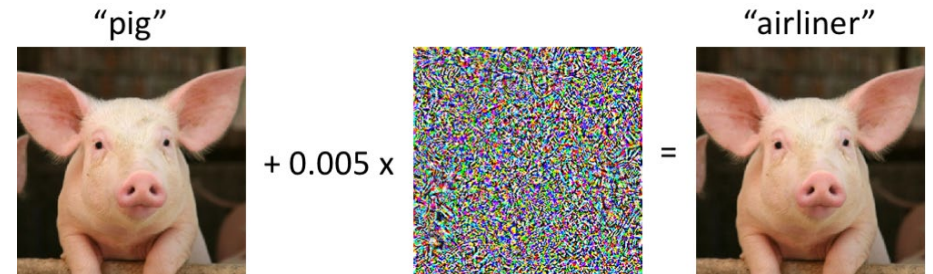


Mirror image x_m



Mirror image y_m

Nudging images in high-dimensional spaces



MIT CSAIL

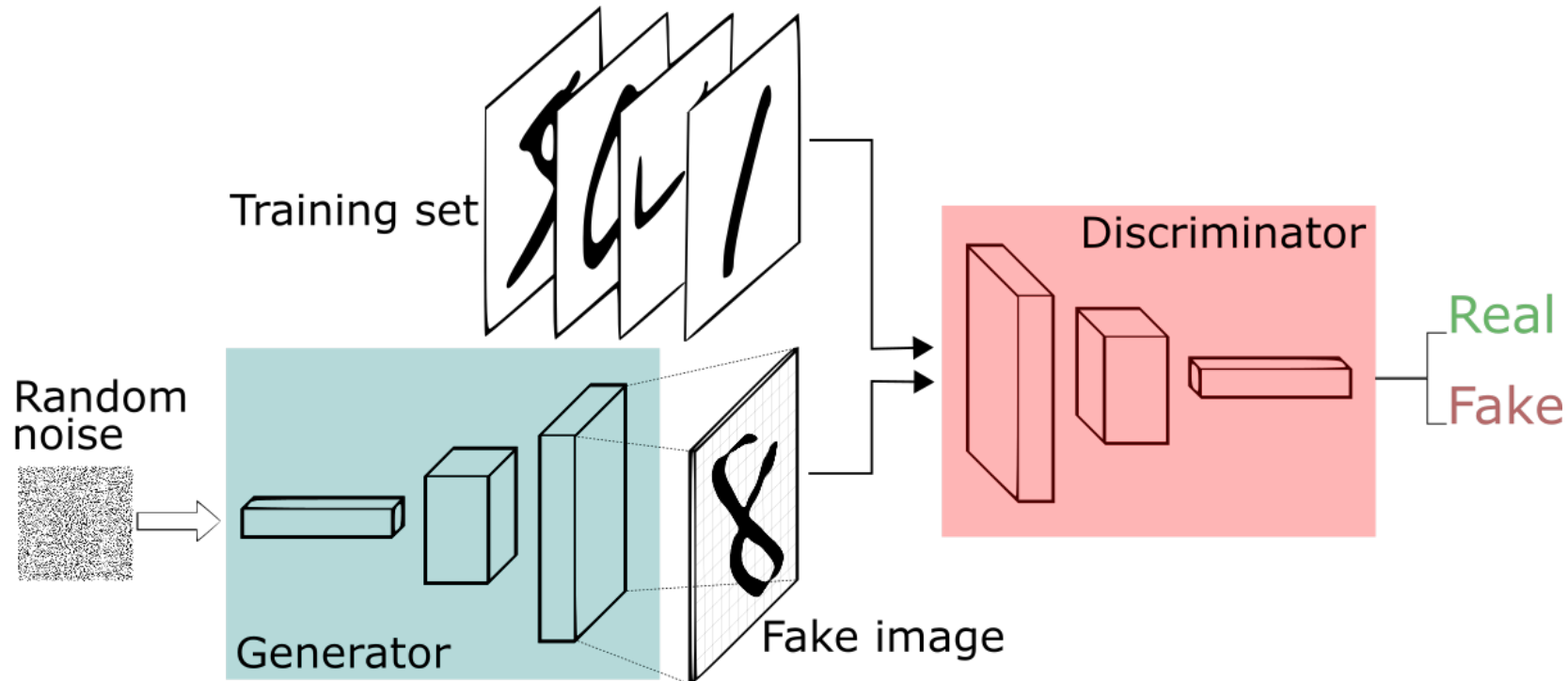


Use of ShapeShifter by Shang-Tse Chen, Ga.Tech

Generative Adversarial Network (GAN): Bug as Feature

“This, and the variations that are now being proposed is the most interesting idea in the last 10 years in ML, in my opinion.”
-Yann LeCun (2016)

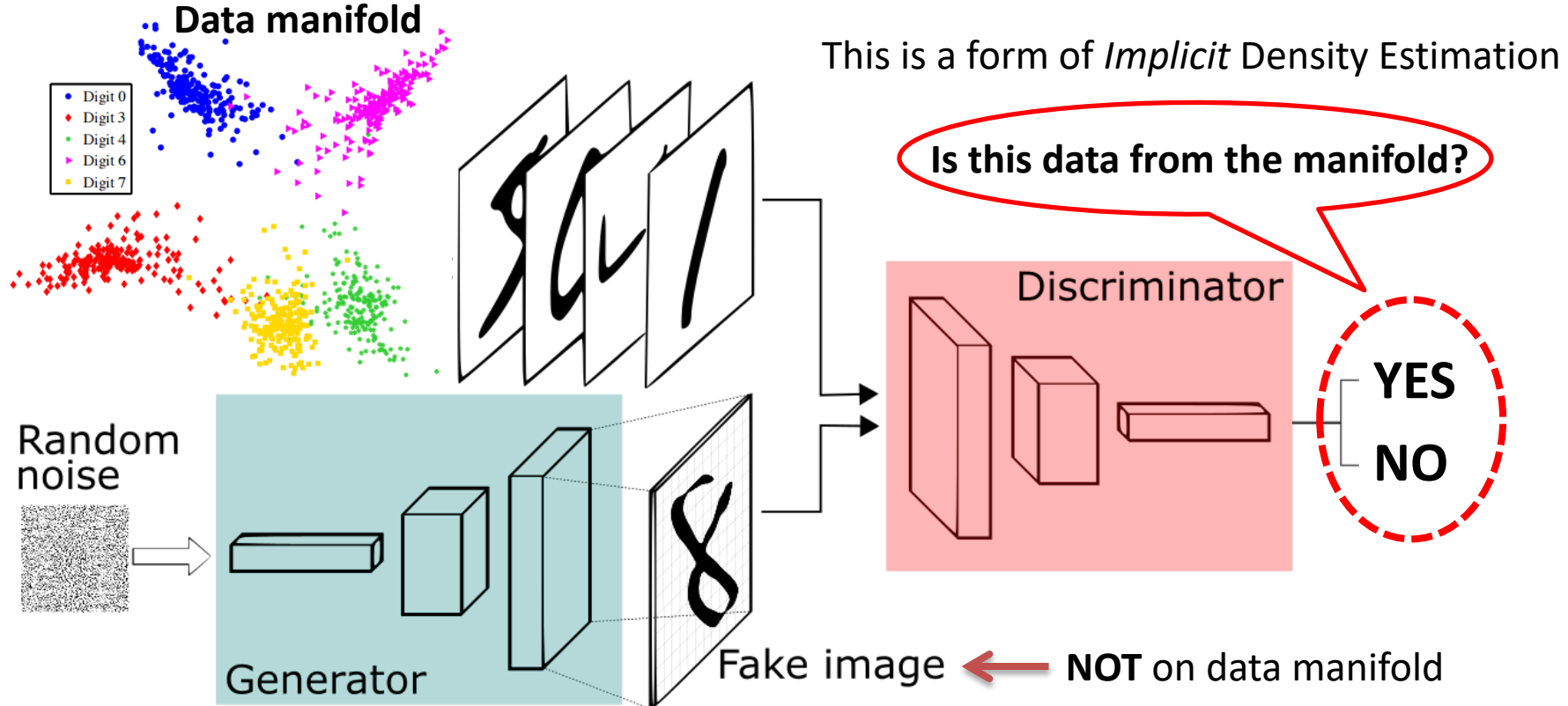
More Supervised Learning to the rescue ... **LABEL** = { *Real*, *Fake* }



Generative Adversarial Network (GAN): Bug as Feature

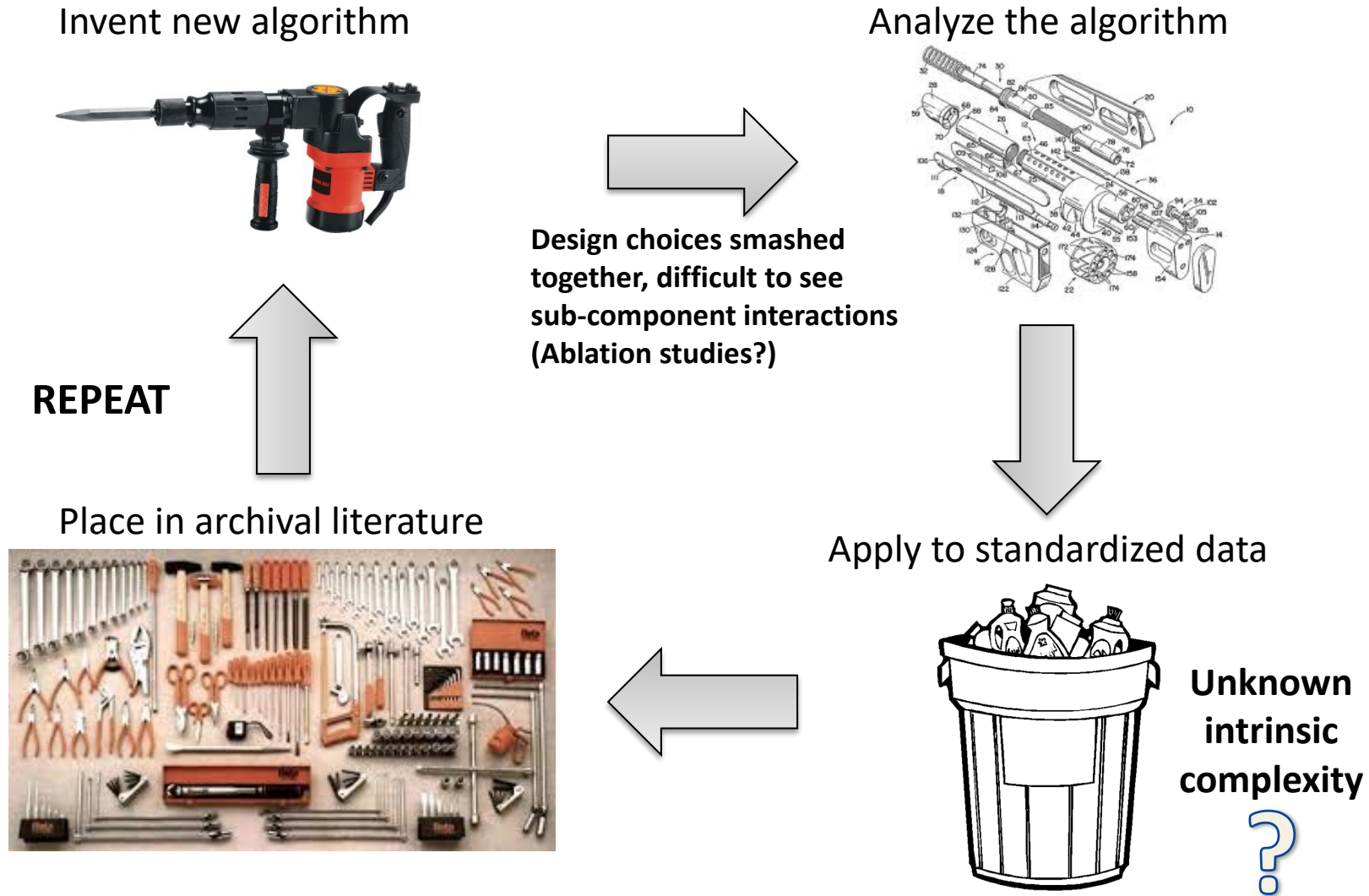
“This, and the variations that are now being proposed is the most interesting idea in the last 10 years in ML, in my opinion.”
-Yann LeCun (2016)

More Supervised Learning to the rescue ... **LABEL** = { *Real*, *Fake* }



To be fair, a bit of a caricature of ...

ML Work Flow leads to Gamification



My algorithm is like your brain ...

Bio-inspiration!



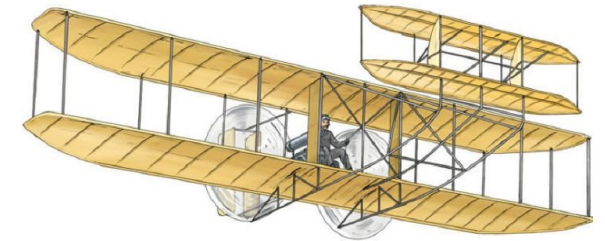
Biomimetic solution?



Neuromorphic computing
Neural networks

“Flapping wings” = “Firing neurons”

Or not?


$$\frac{\text{Propeller}}{\text{Flapping wings}} = \frac{?}{\text{Firing neurons}}$$

- Birds are a solution to an engineering problem with physical constraints. The Wright brothers understood the physical problem and then found an appropriate engineering solution different from birds.
- Brains are an evolutionary solution to the statistical constraints of inference from experience. What are those statistical constraints?



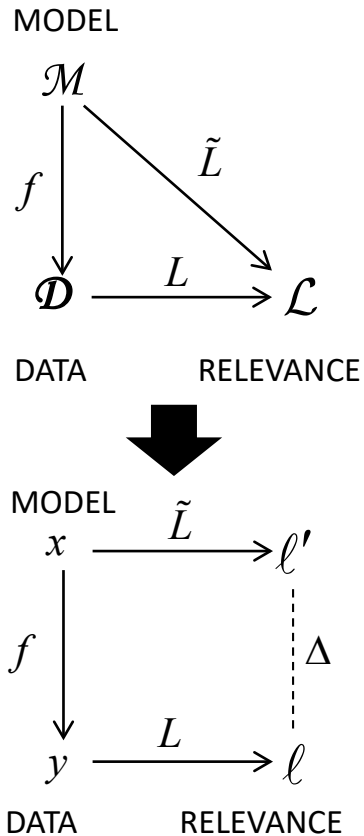
"It sort of makes you stop and think, doesn't it."

Ti fa solo fermare e pensare, vero?

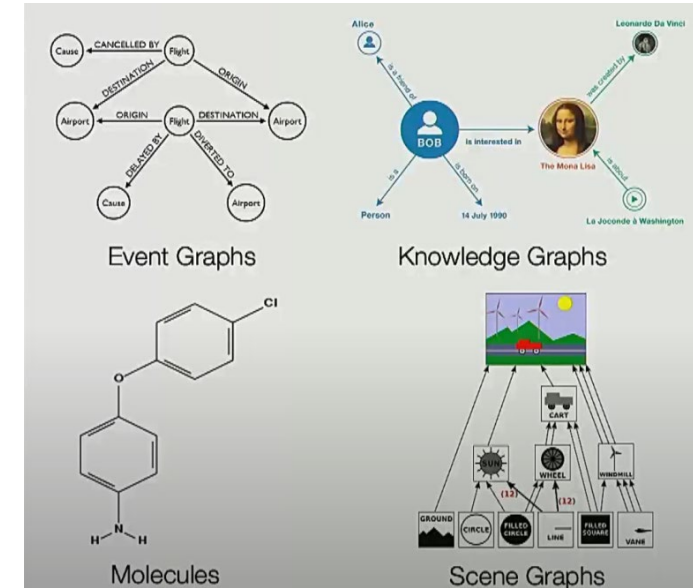
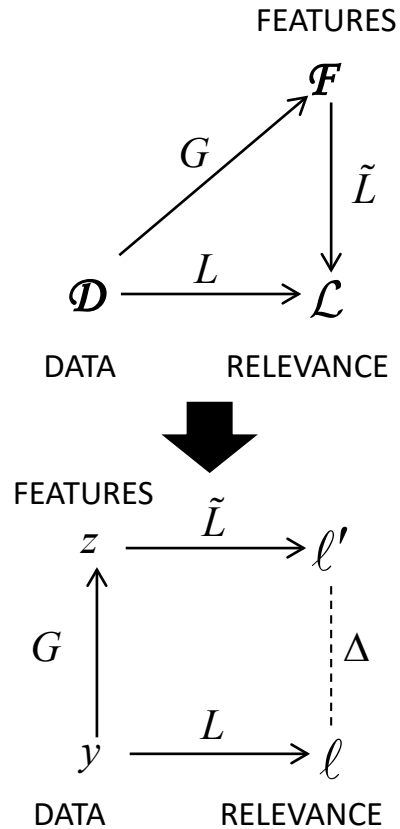
Geometry of Machine Learning

Diverse and complex data structures

Physics



Machine Learning



But we want to use Linear Algebra and Multivariable Calculus!

So just work in a huge vector space! $\mathbb{R}^D$

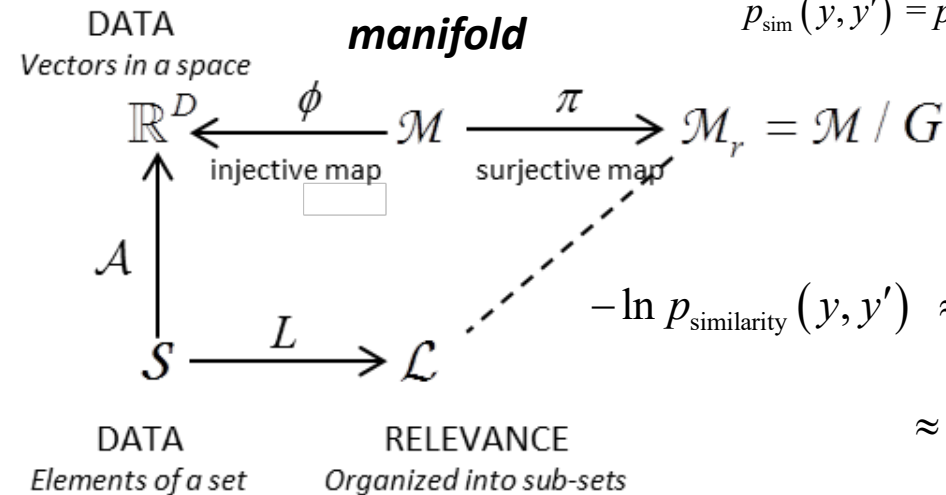
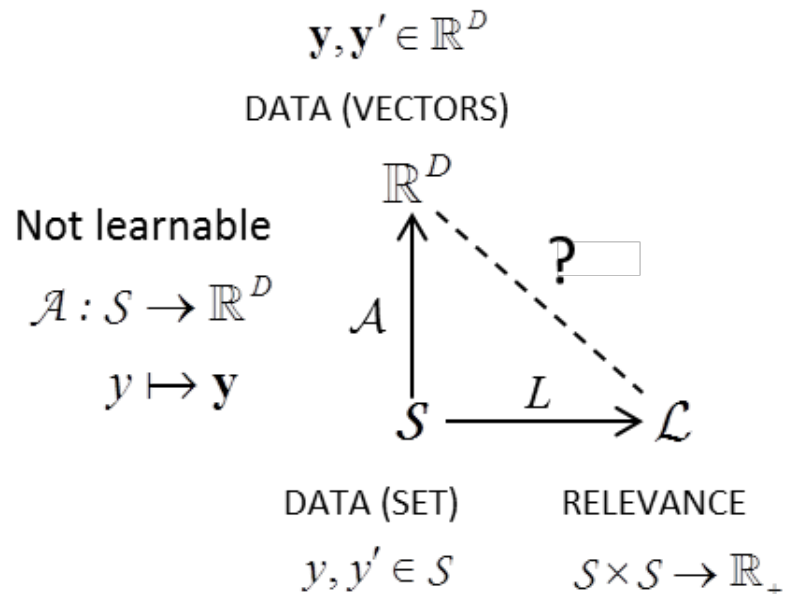
Functions acting on these vectors should have: Continuity $\rightarrow$ Smoothness $\rightarrow$ Differentiability } **manifold**

Model regularization (prior)

Learning via gradient descent

Geometry of Machine Learning

- STEPS:**
1. Form a DATA space, $\mathcal{D}$, as a high-dimensional vector space, $\mathbb{R}^D$.
 2. Identify transformations, S , on the DATA space that leave the similarity measure invariant.
 3. Learn the underlying MODEL space, $\mathcal{M}$, or “embedding” that preserves the invariances of the similarity measure.



A group, G , acting on the set, S , that leaves the similarity measure *invariant*:

$$p_{\text{sim}}(y, y') = p_{\text{sim}}(g_\alpha y, g_\beta y') \text{ where } g_\alpha, g_\beta \in G.$$

$$-\ln p_{\text{similarity}}(y, y') \approx \frac{1}{2} \text{dist}_{\mathcal{M}}^2(\mathcal{A}(y), \mathcal{A}(y'))$$

$$\approx \frac{1}{2} \|\mathbf{y} - \mathbf{y}'\|_2^2$$

Need sufficient sampling so the data is “close” enough to use this approximation.

Geometry of ML: How much data is enough?

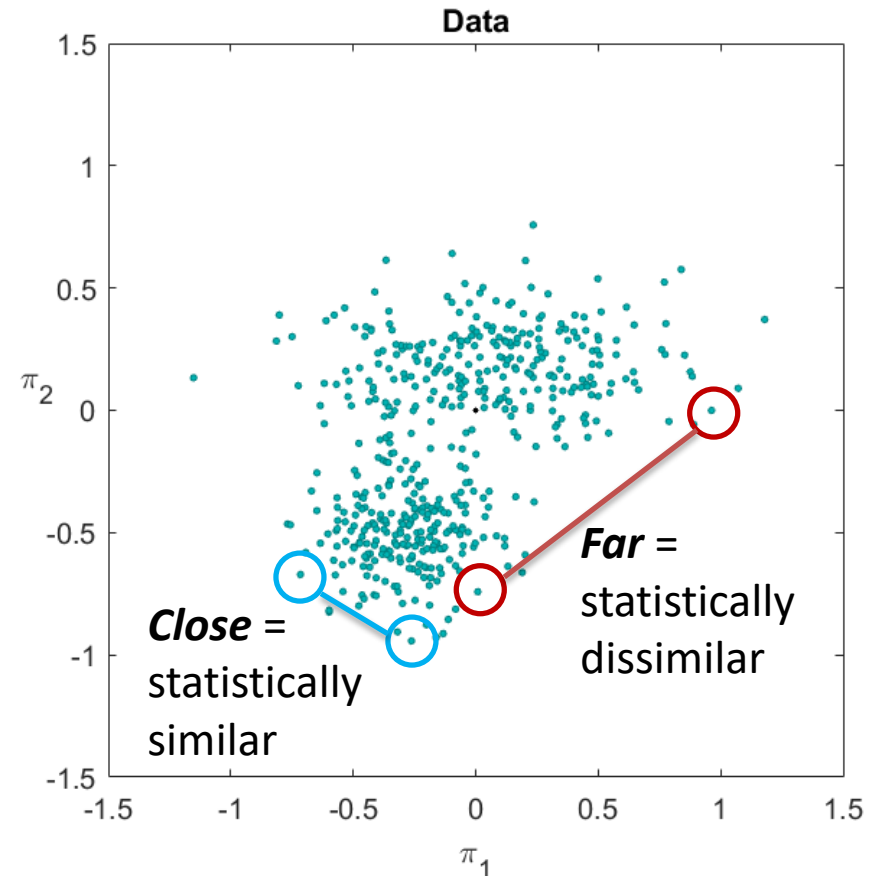
Answer: It's not about having "a lot of data", it's about having *enough* data in the right places to answer a particular question.

Assumption: The *distance* encodes information about *statistical similarity*.

$$-\ln p_{\text{similarity}}(y, y') \approx \frac{1}{2} \|\mathbf{y} - \mathbf{y}'\|_2^2$$

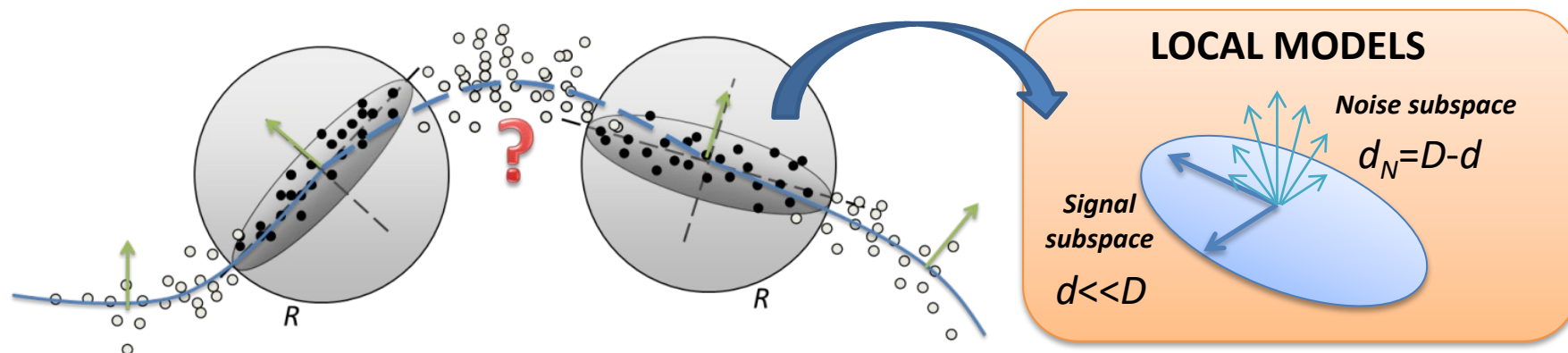


If distant data points are mapped nearby in the model space, then there has to be significant inferential evidence (more data), and likewise if local data points are mapped to large separations by the model.



Geometry of ML: How much data is enough?

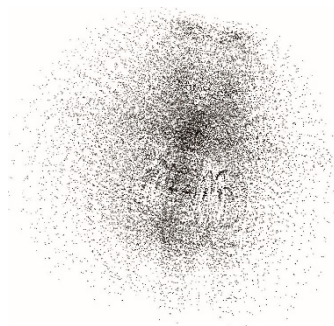
To find the best local linear model: We need enough data *locally* to distinguish curvature from finite sampling noise



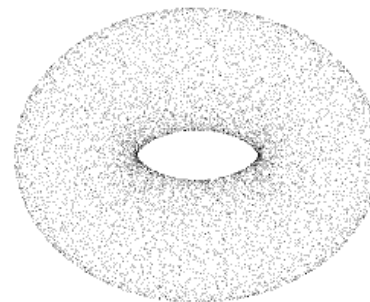
Local structure
High- D Noise Balls



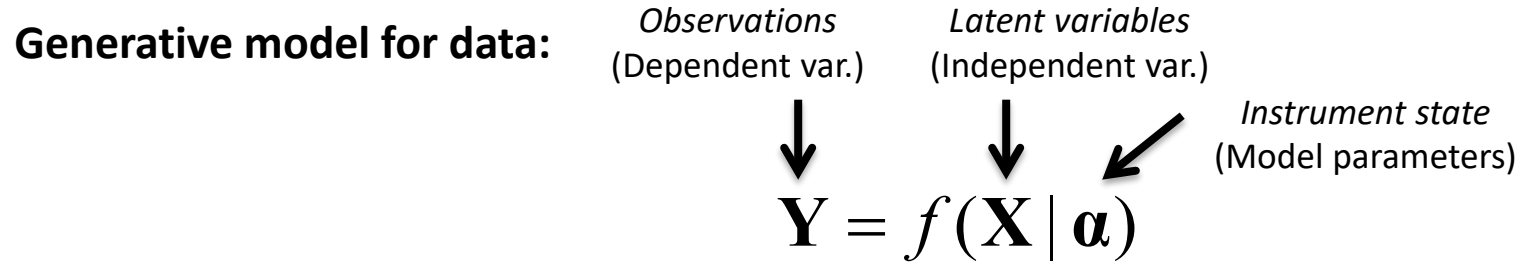
Crossover Regime
Noisy Tangent Spaces
(with an envelope)



Large-scale structure
Low- d Manifold

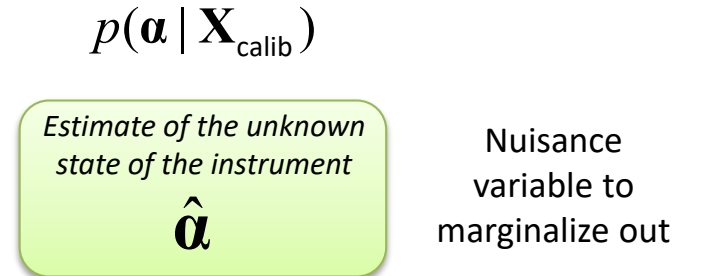
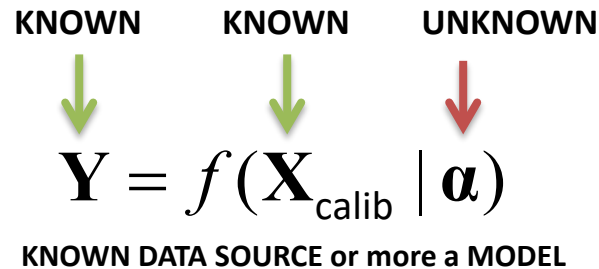


Geometry of ML: Experimental Calibration



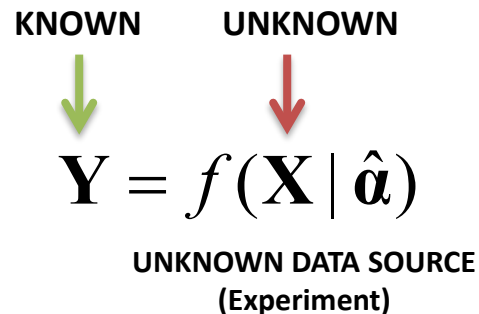
1. CALIBRATION

Regression with
training data

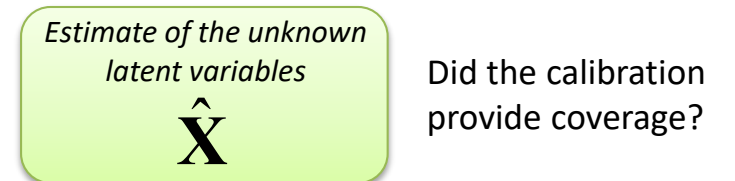


2. MEASUREMENT

Inferring latent
variables



$$p(\mathbf{X} | \mathbf{X}_{\text{calib}}) = \int d\boldsymbol{\alpha} p(\mathbf{X} | \boldsymbol{\alpha}) \cdot p(\boldsymbol{\alpha} | \mathbf{X}_{\text{calib}})$$

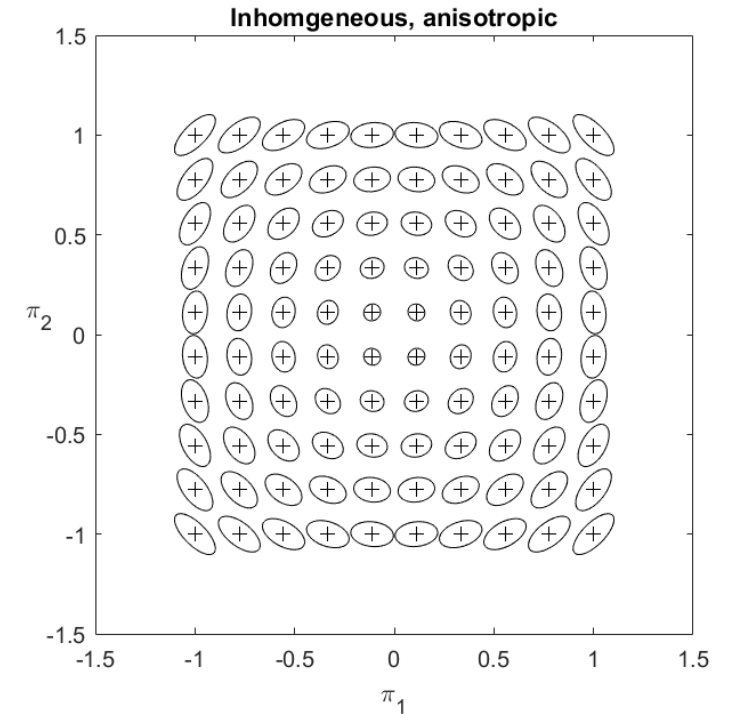
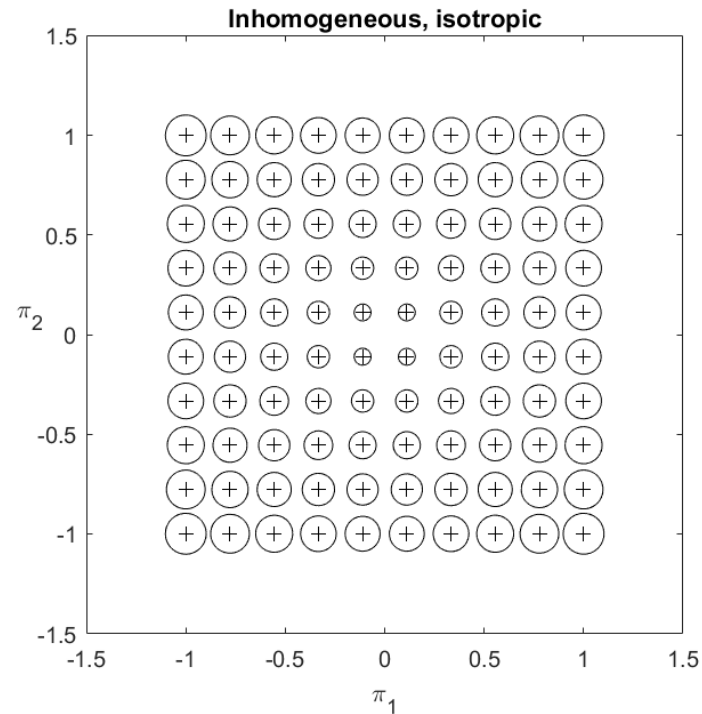
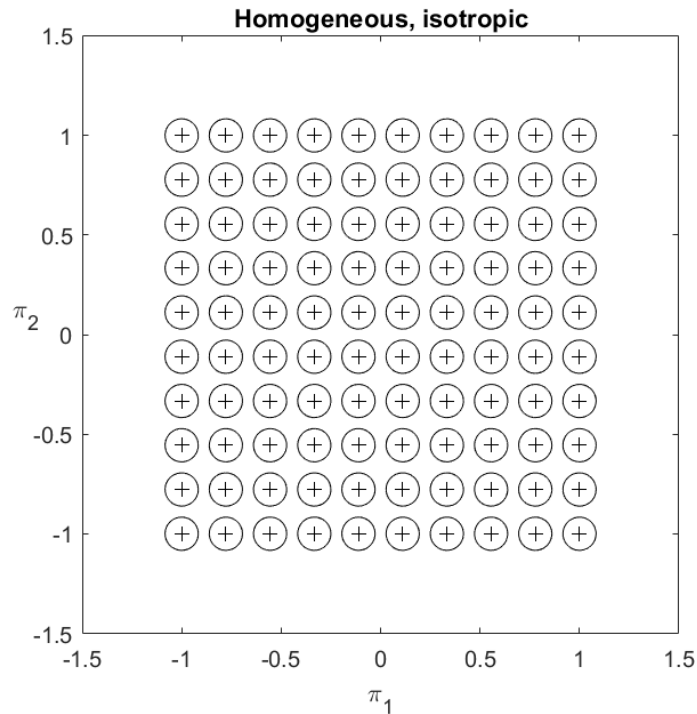


Geometry of ML: Experimental Calibration

Examples of different possible calibrations across the **DATA** space ...

1. CALIBRATION

(1,1)-tensor field of covariance



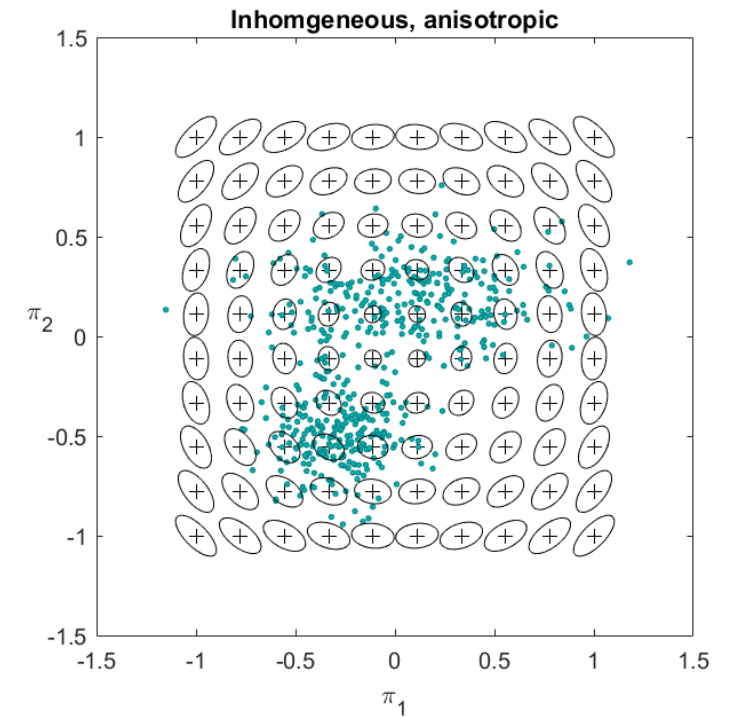
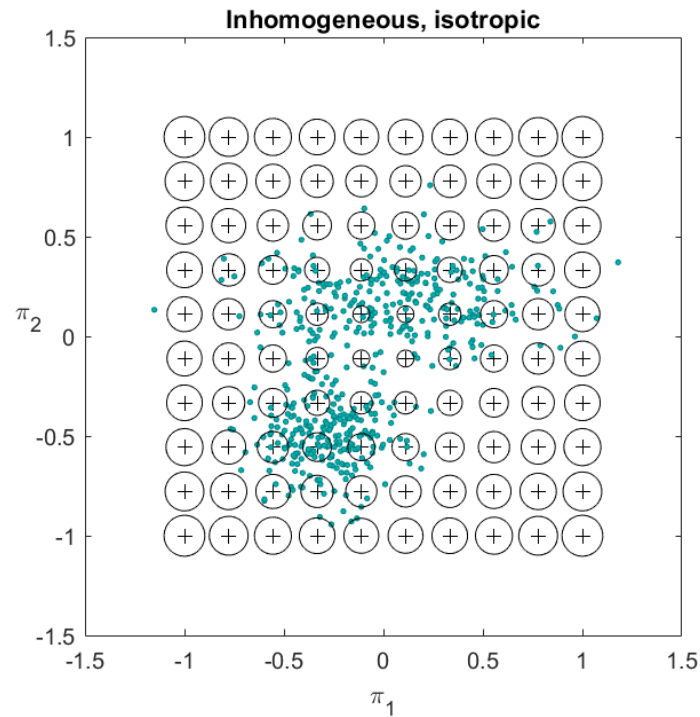
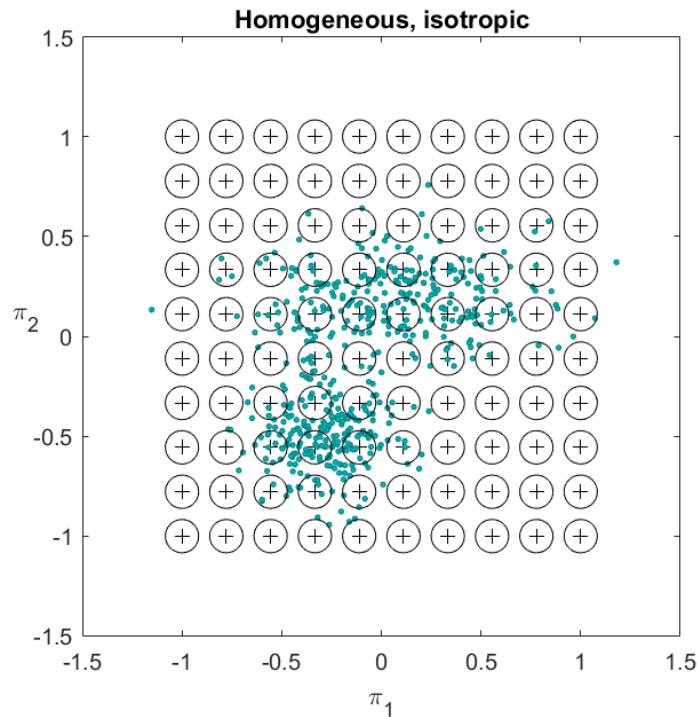
How do we use this information in Machine Learning?

Geometry of ML: Experimental Calibration

.... with the “same” data measurements.

2. MEASUREMENT

(1,1)-tensor field of covariance

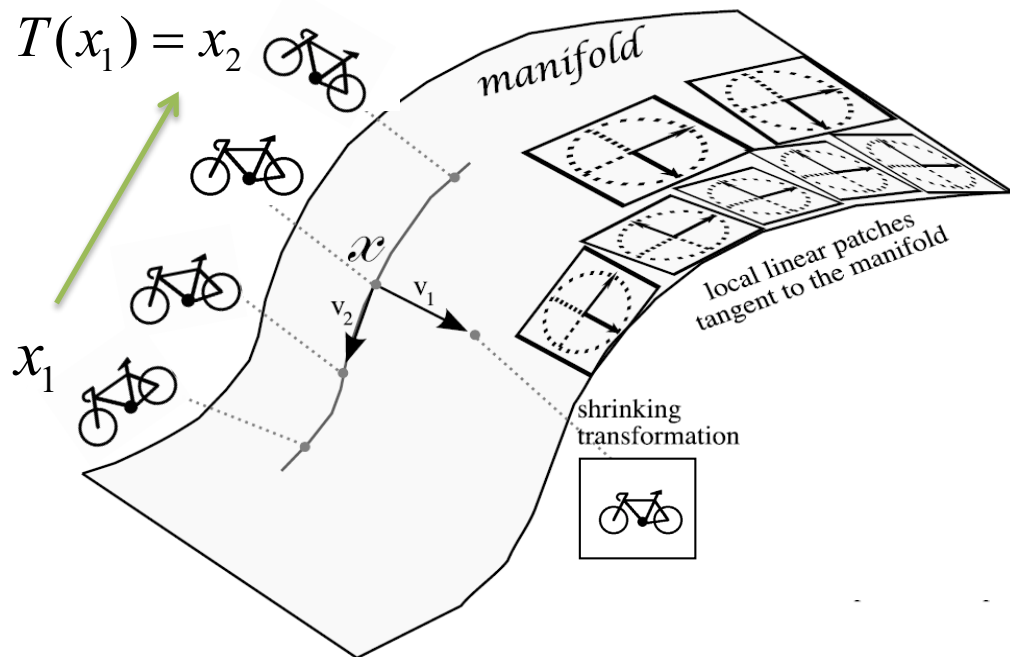
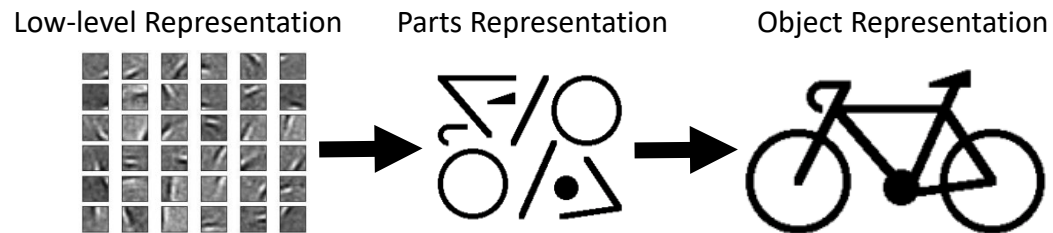


Use the covariance to estimate the metric tensor of the measurement space

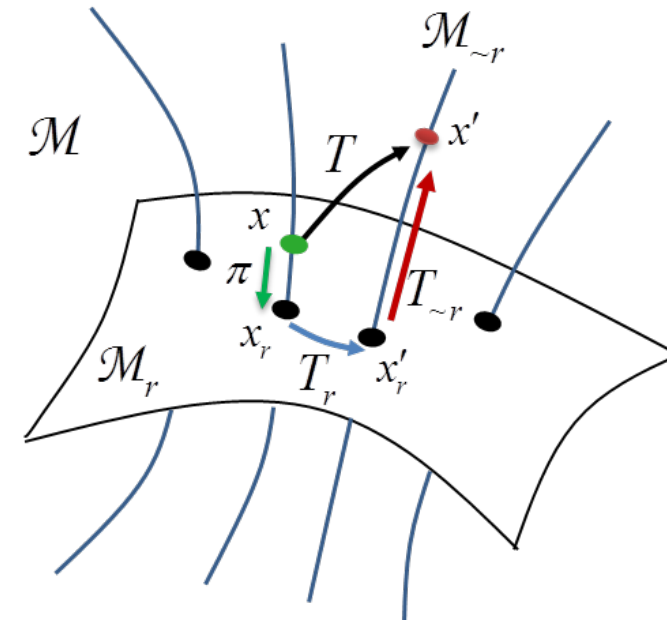
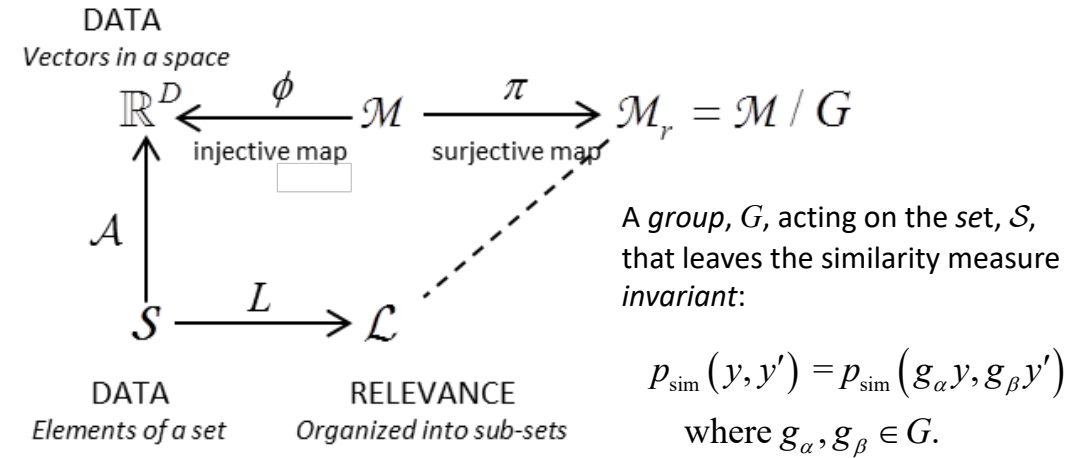
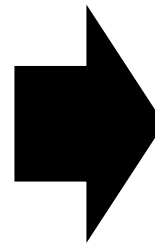
$$\mathbf{g}(y) = \Sigma^{-1}(y) = \mathbf{J}_y^T \mathbf{J}_y \rightarrow \tilde{\mathbf{g}} = \mathbb{I}_D$$

Now, the ML Euclidean space ansatz is true!

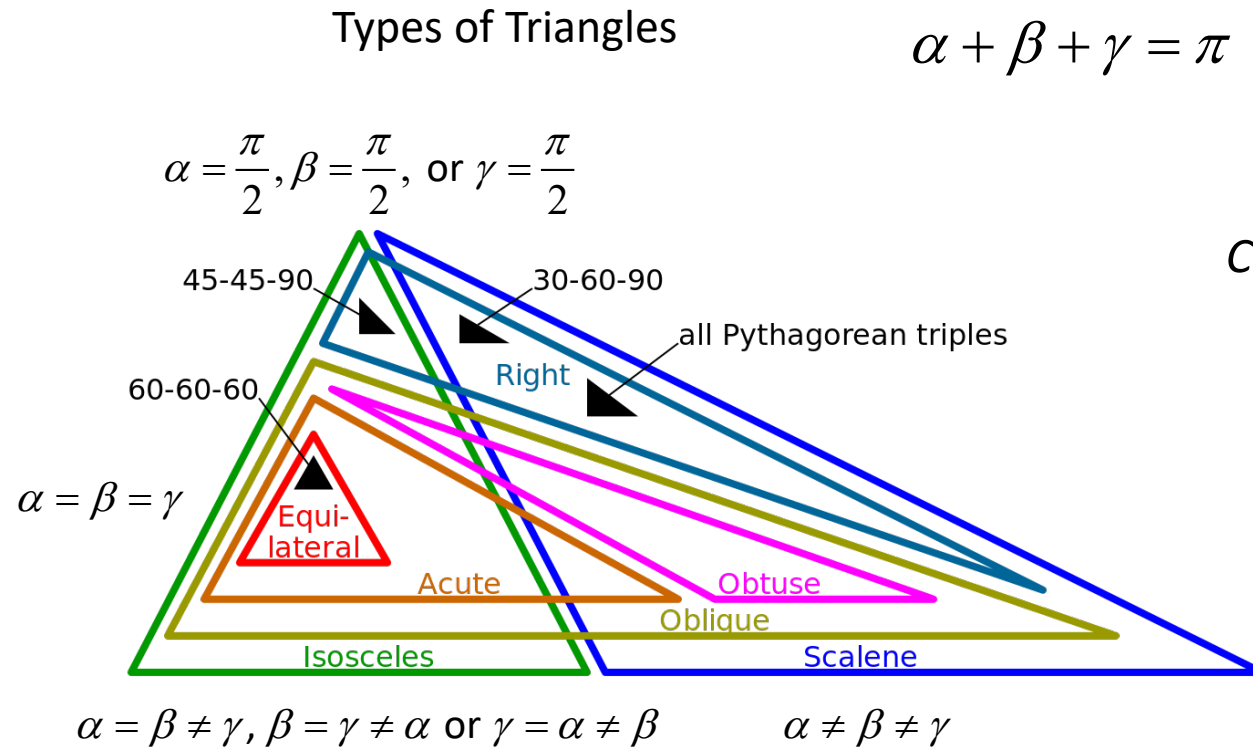
Geometry of ML: Manifolds and Relevance



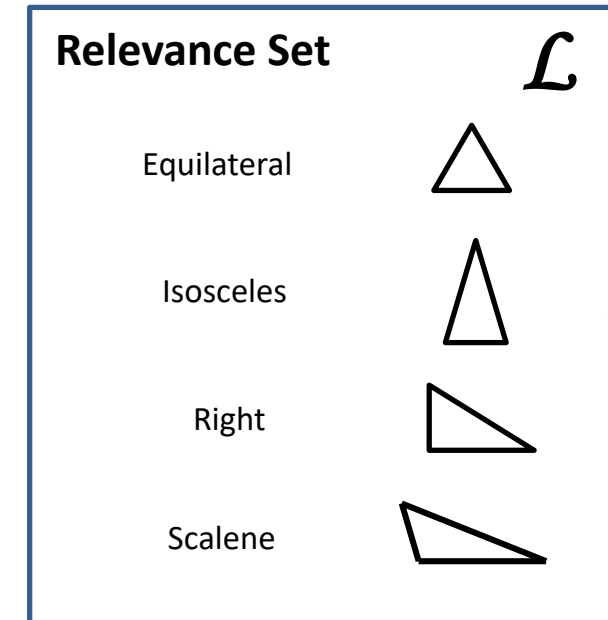
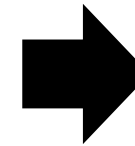
Re-organize
the space as a
fiber bundle



Example: Choosing the Relevance of Triangles



CHOOSE



$y = \text{IMAGE} \rightarrow "f^{-1}" = \text{Vertex Detector} \rightarrow x, \dim(\mathcal{M})=6 \rightarrow x_r, \dim(\mathcal{M}_r)=2 \rightarrow$

$\ell, |\mathcal{L}| = 4$

$$\ell = L(y)$$

Example: Shapes and Geometric Invariants

Data Space

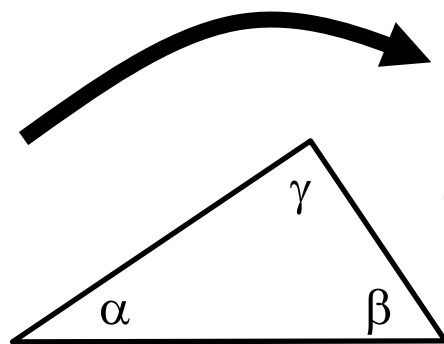
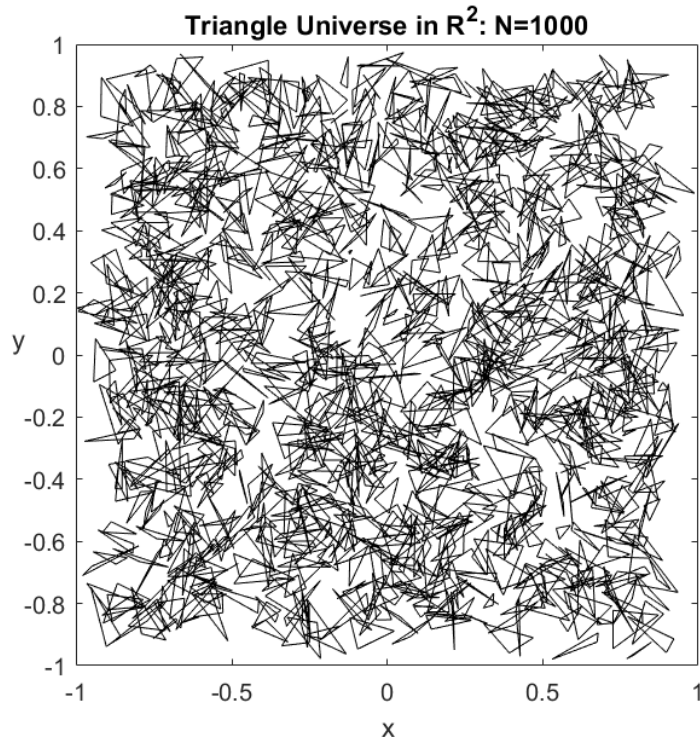
$\mathcal{D}$

$$\mathcal{M} = \mathbb{R}^6$$

$$\dim \mathcal{M} = 6$$

$$\dim \mathcal{M}_r = 2$$

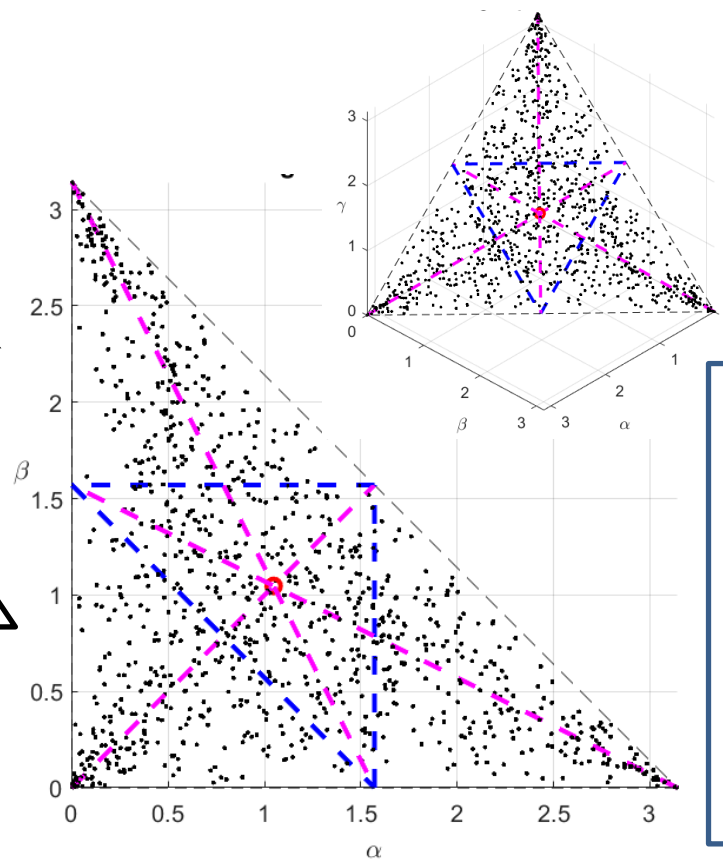
$$x_r = \pi \circ f^{-1}(y)$$



$$\alpha + \beta + \gamma = \pi$$

Model Relevance Space

$\mathcal{M}_r$



Relevance Set



Equilateral



Isosceles



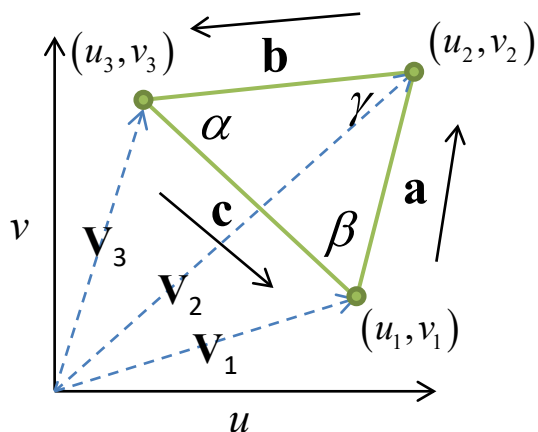
Right



$\mathcal{L}$

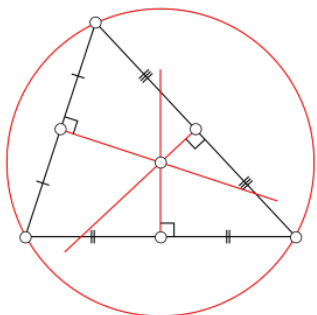
Transformations: From Model to Model Relevance Space

MODEL space



Diameter of the circumscribed circle:

$$s = \frac{a}{\sin \alpha} = \frac{b}{\sin \beta} = \frac{c}{\sin \gamma}$$



$$\mathbf{x}_r = \pi(\mathbf{x})$$

$$\mathbf{x} = \begin{pmatrix} \mathbf{v}_1 \\ \mathbf{v}_2 \\ \mathbf{v}_3 \end{pmatrix} = \begin{pmatrix} u_1 \\ v_1 \\ u_2 \\ v_2 \\ u_3 \\ v_3 \end{pmatrix}$$



Translation invariance

$$\mathbf{s} = \begin{pmatrix} \mathbf{a} \\ \mathbf{b} \\ \mathbf{c} \end{pmatrix} = \begin{pmatrix} u_2 - u_1 \\ v_2 - v_1 \\ u_3 - u_2 \\ v_3 - v_2 \\ u_1 - u_3 \\ v_1 - v_3 \end{pmatrix}$$



Scale invariance

$$\hat{\mathbf{s}} = \begin{pmatrix} \hat{\mathbf{a}} \\ \hat{\mathbf{b}} \\ \hat{\mathbf{c}} \end{pmatrix} = \begin{pmatrix} \frac{u_2 - u_1}{a} \\ \frac{v_2 - v_1}{a} \\ \frac{u_3 - u_2}{b} \\ \frac{v_3 - v_2}{b} \\ \frac{u_1 - u_3}{c} \\ \frac{v_1 - v_3}{c} \end{pmatrix}$$



$$\begin{pmatrix} \alpha \\ \beta \\ \gamma \end{pmatrix} = \begin{pmatrix} \cos^{-1}(-\hat{\mathbf{b}}^T \hat{\mathbf{c}}) \\ \cos^{-1}(-\hat{\mathbf{c}}^T \hat{\mathbf{a}}) \\ \cos^{-1}(-\hat{\mathbf{a}}^T \hat{\mathbf{b}}) \end{pmatrix}$$



Rotation invariance

$$\begin{aligned} a &= \sqrt{\mathbf{a}^T \mathbf{a}} = \sqrt{(u_2 - u_1)^2 + (v_2 - v_1)^2} \\ b &= \sqrt{\mathbf{b}^T \mathbf{b}} = \sqrt{(u_3 - u_2)^2 + (v_3 - v_2)^2} \\ c &= \sqrt{\mathbf{c}^T \mathbf{c}} = \sqrt{(u_1 - u_3)^2 + (v_1 - v_3)^2} \end{aligned}$$

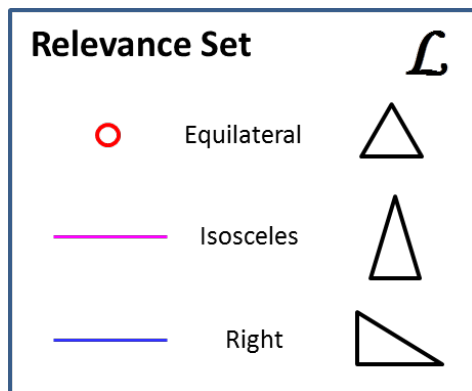


$$\mathbf{x}_r = \begin{pmatrix} \xi \\ \eta \\ s \end{pmatrix} = \begin{pmatrix} \frac{1}{\sqrt{2}}(\beta - \alpha) \\ \frac{1}{\sqrt{6}}(2\gamma - \alpha - \beta) \\ \frac{a+b+c}{\sin \alpha + \sin \beta + \sin \gamma} \end{pmatrix}$$

s is **NOT** scale invariant

$$\begin{aligned} \sin \gamma &= \sin(\pi - \alpha - \beta) \\ &= \sin \pi \cdot \cos(\alpha + \beta) - \cos \pi \cdot \sin(\alpha + \beta) \\ &= \sin(\alpha + \beta) \end{aligned}$$

Model compression: Triangle Example

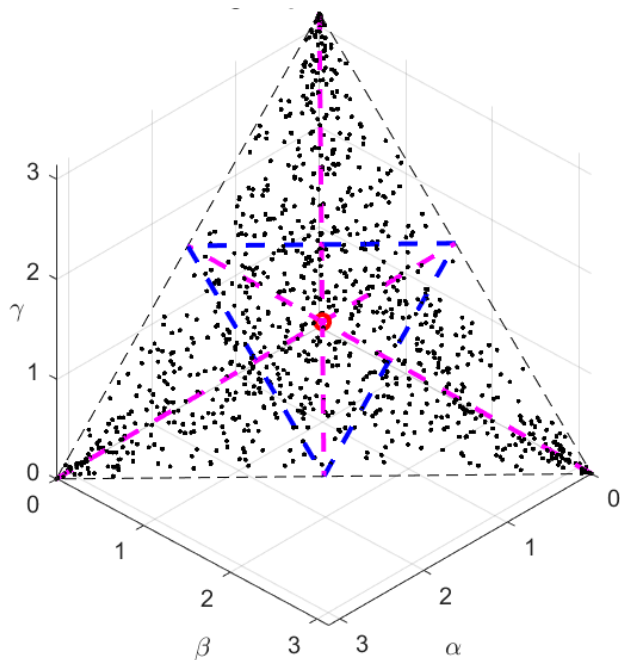


$$\mathcal{M} = \mathbb{R}^6$$

$$\dim \mathcal{M} = 6$$

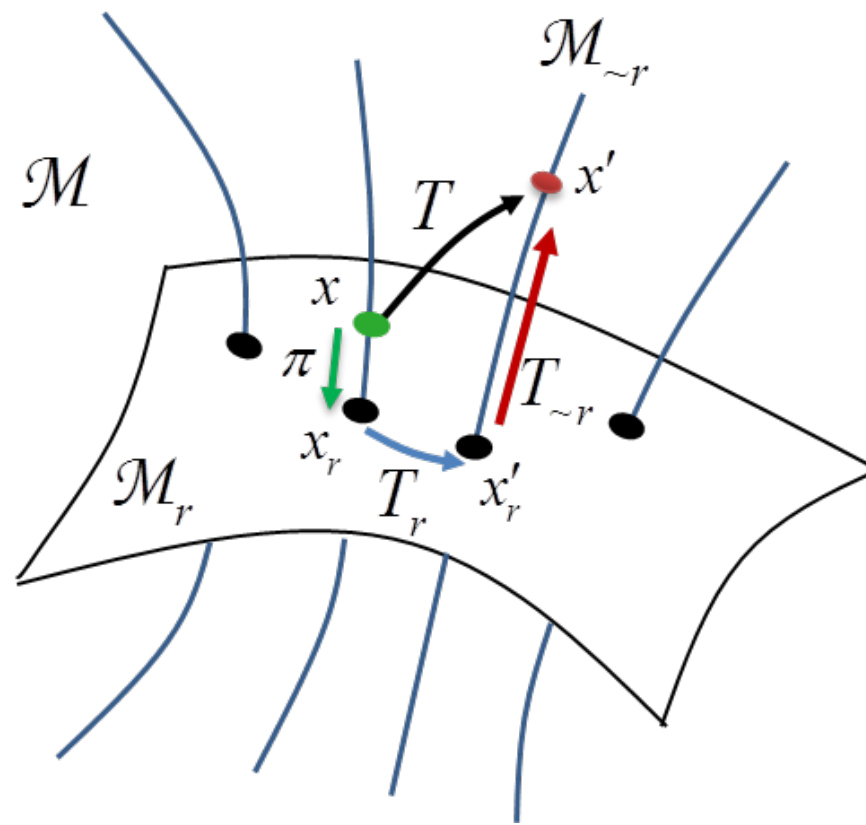
$$\dim \mathcal{M}_r = 2$$

$$x_r = \pi \circ f^{-1}(y)$$



sim(2)

- Translation invariance
- Scale invariance
- Rotation invariance



Physics and Machine Learning

