



模式识别 Pattern Recognition

李泽桦，复旦大学 生物医学工程与技术创新学院

目录

1 高斯过程

2 CNN

3 医学图像分割

判别模型

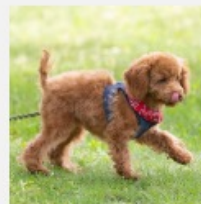
- “样本” $x \Rightarrow$ “类别” y
- 仅一个输出

生成模型

- “类别” $y \Rightarrow$ “样本” x
- 多个可能输出

discriminative

x



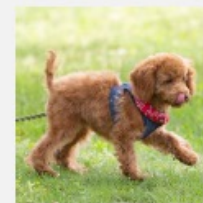
“dog”

y

generative

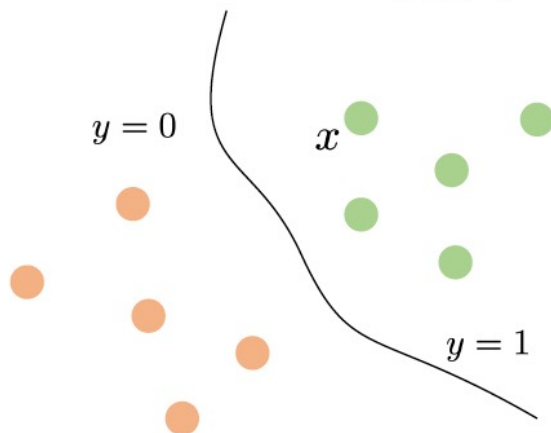
y

“dog”

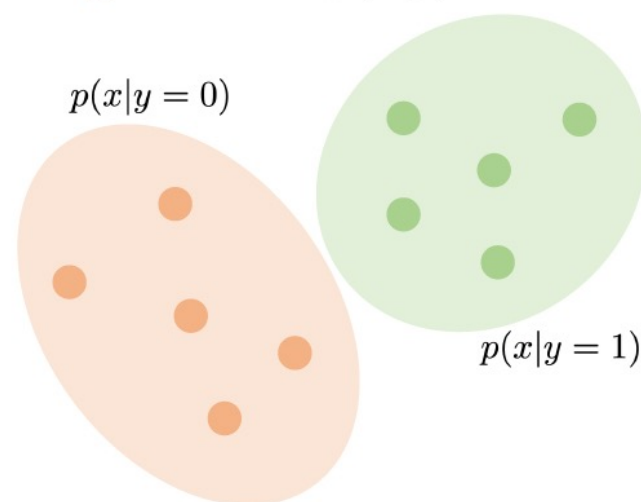


x

discriminative $p(y|x)$



generative $p(x|y)$



- 生成模型也可以具有判别性：贝叶斯法则
- 判别模型不具有生成性

- 生成模型也可以具有判别性：贝叶斯法则

$$\underset{\text{discriminative}}{p(y|x)} = \underset{\text{generative}}{p(x|y)} \frac{p(y)}{p(x)}$$

类别先验

对于特定 x ，概率固定

- 生成模型也可以具有判别性：贝叶斯法则。

$$\underset{\text{discriminative}}{p(y|x)} = \underset{\text{generative}}{p(x|y)} \frac{p(y)}{p(x)}$$

类别先验

对于特定 x ，概率固定

- 判别模型不具有生成性

$$\underset{\text{generative}}{p(x|y)} = \underset{\text{discriminative}}{p(y|x)} \frac{p(x)}{p(y)}$$

x 的分布无法获得

对于特定 y ，概率固定

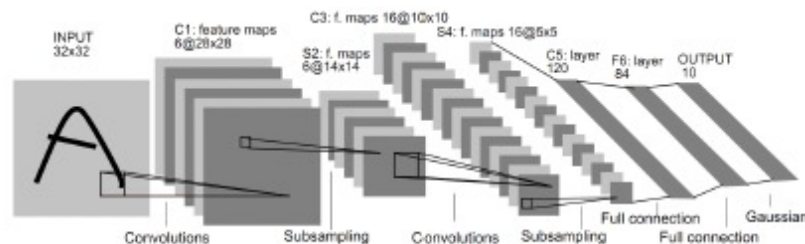
判别模型与生成模型：一体两面



- 判别模型

- 从数据映射到特征: $x \rightarrow f(x)$

- 优化损失函数: $\mathcal{L}(y, f(x))$



- 生成模型

- 从简单分布映射到复杂分布: $\pi \rightarrow g(\pi)$

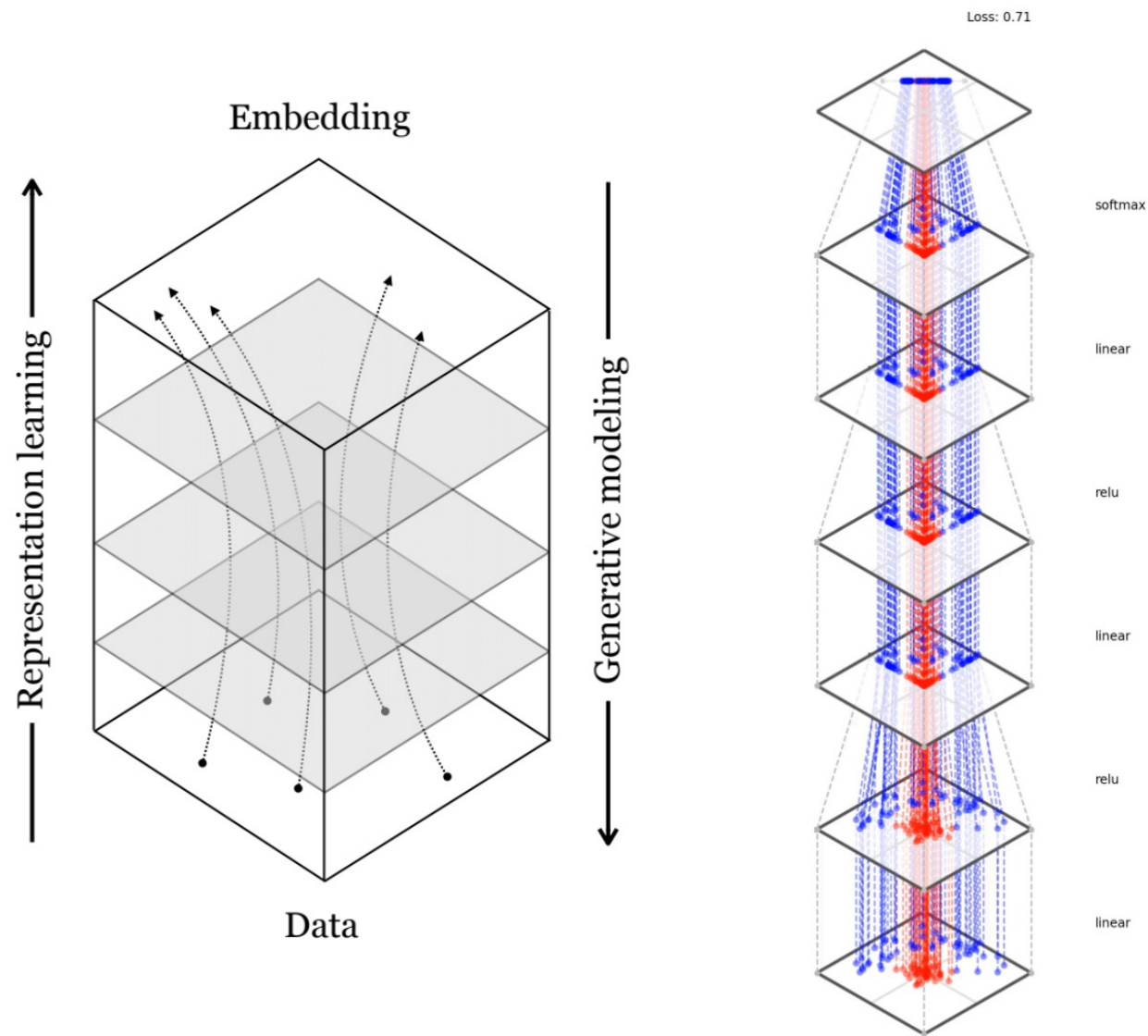
- 优化损失函数: $\mathcal{L}(p_x, g(\pi))$



判别模型与生成模型：一体两面



- 深度网络学习每一层的特征映射
- 前向传播（观测数据->潜在表征）是判别模型
- 反向传播（潜在表征->观测数据）是生成模型



- 一个常用来做判别任务的**生成模型**
- 非参数化模型 (Non-parametric functions)
- 高斯过程：观测值出现在一个连续域（例如时间或空间）的随机过程

$$(y(x_1), \dots, y(x_n)) \sim N_n(\mu, \Sigma)$$

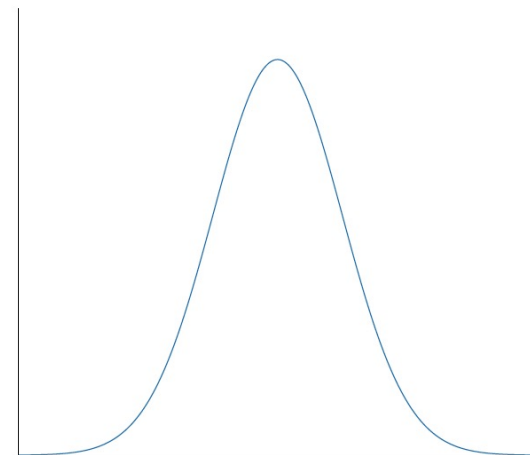
- $y(\cdot)$ 是x的function

- 高斯分布：需要计算均值 (μ) 和方差 (Σ)

$$X \sim N(\mu, \Sigma)$$

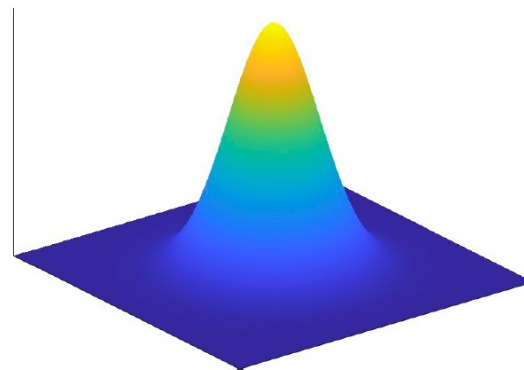
- (Univariate) Gaussians:

$$x \sim \mathcal{N}(x; \mu = 0, \sigma^2 = 1)$$



$$\mu = \begin{pmatrix} \mu_1 \\ \mu_2 \end{pmatrix} \quad \Sigma = \begin{pmatrix} \sigma_1^2 & \rho_{12}\sigma_1\sigma_2 \\ \rho_{21}\sigma_1\sigma_2 & \sigma_2^2 \end{pmatrix}$$

$$\text{Cor}(Y_i, Y_j) = \rho_{12} \text{ for } i \neq j$$



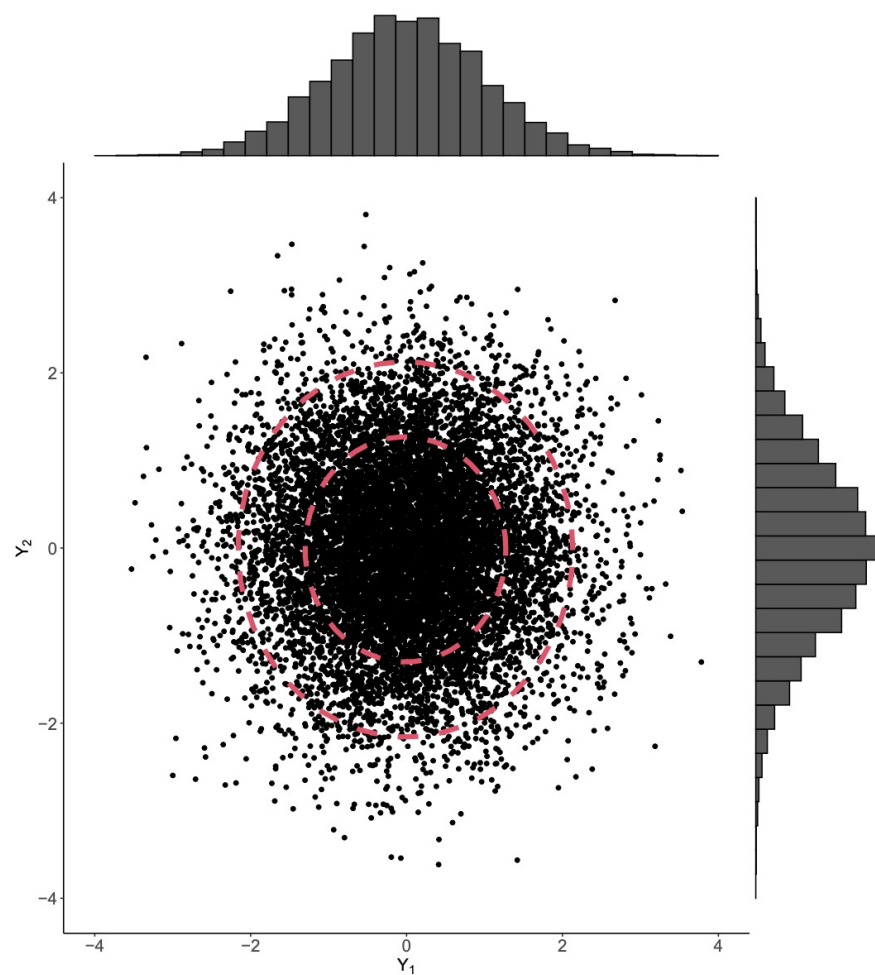
- Multivariate Gaussians:

$$\mathbf{x} = [x_1, \dots, x_D]^T$$
$$\sim \mathcal{N}(\mathbf{x}; \boldsymbol{\mu} = \mathbf{0}_D, \Sigma = I_D)$$

二元高斯分布



- Y_1 和 Y_2 无关

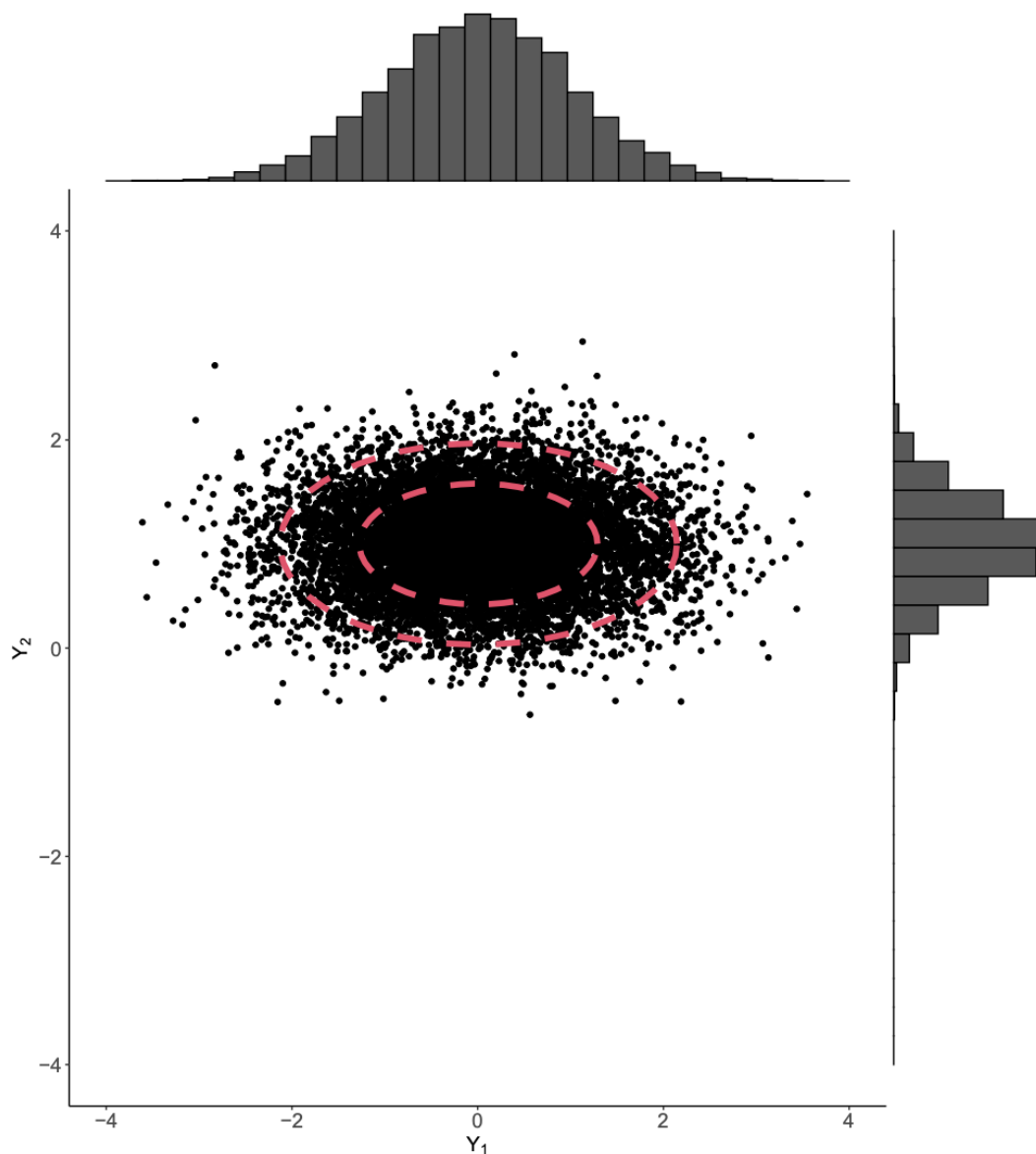


$$\mu = \begin{pmatrix} 0 \\ 0 \end{pmatrix}$$

$$\Sigma = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix}$$



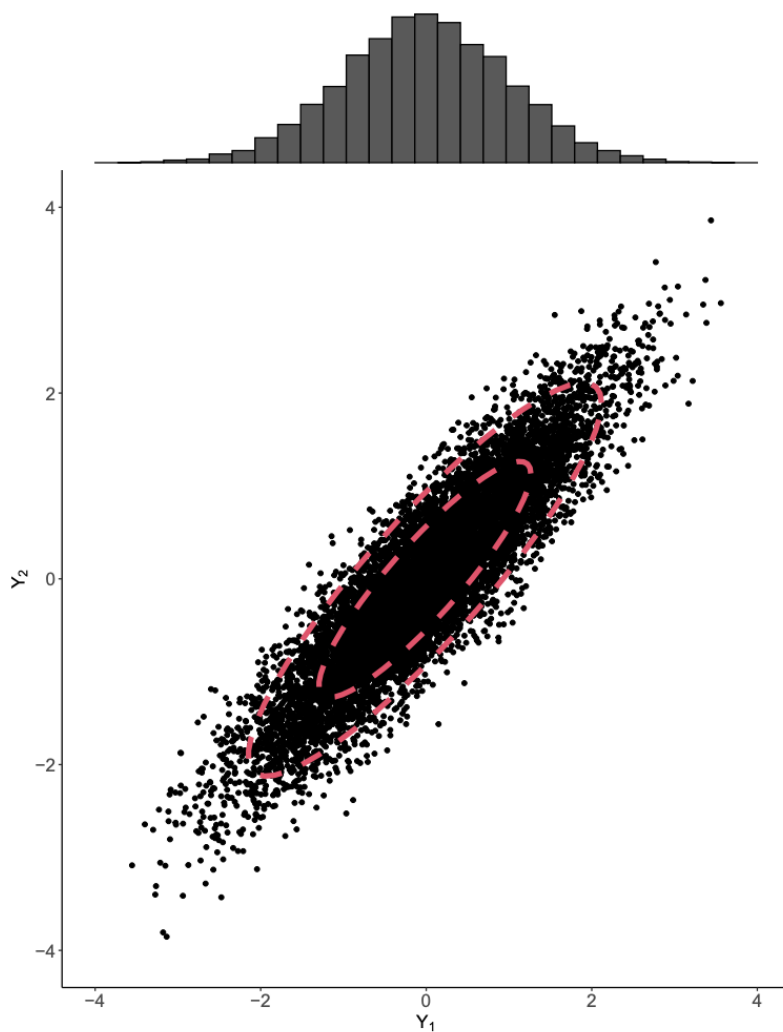
- Y_1 和 Y_2 无关



$$\mu = \begin{pmatrix} 0 \\ 1 \end{pmatrix}$$

$$\Sigma = \begin{pmatrix} 1 & 0 \\ 0 & 0.2 \end{pmatrix}$$

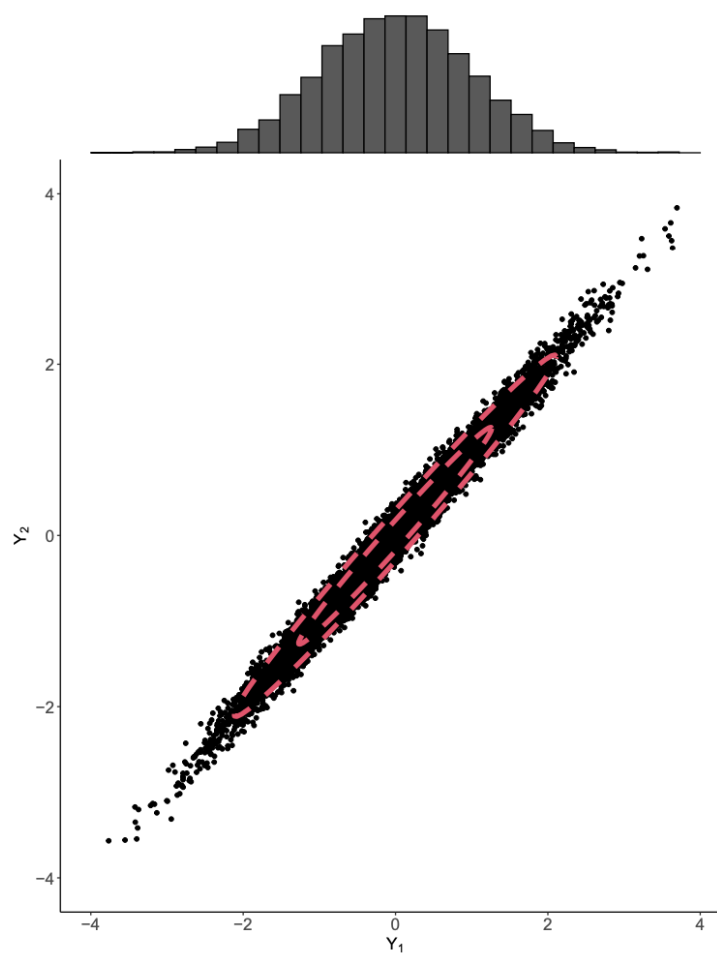
- Y_1 和 Y_2 相关



$$\mu = \begin{pmatrix} 0 \\ 0 \end{pmatrix}$$

$$\Sigma = \begin{pmatrix} 1 & 0.9 \\ 0.9 & 1 \end{pmatrix}$$

- Y_1 和 Y_2 特别相关



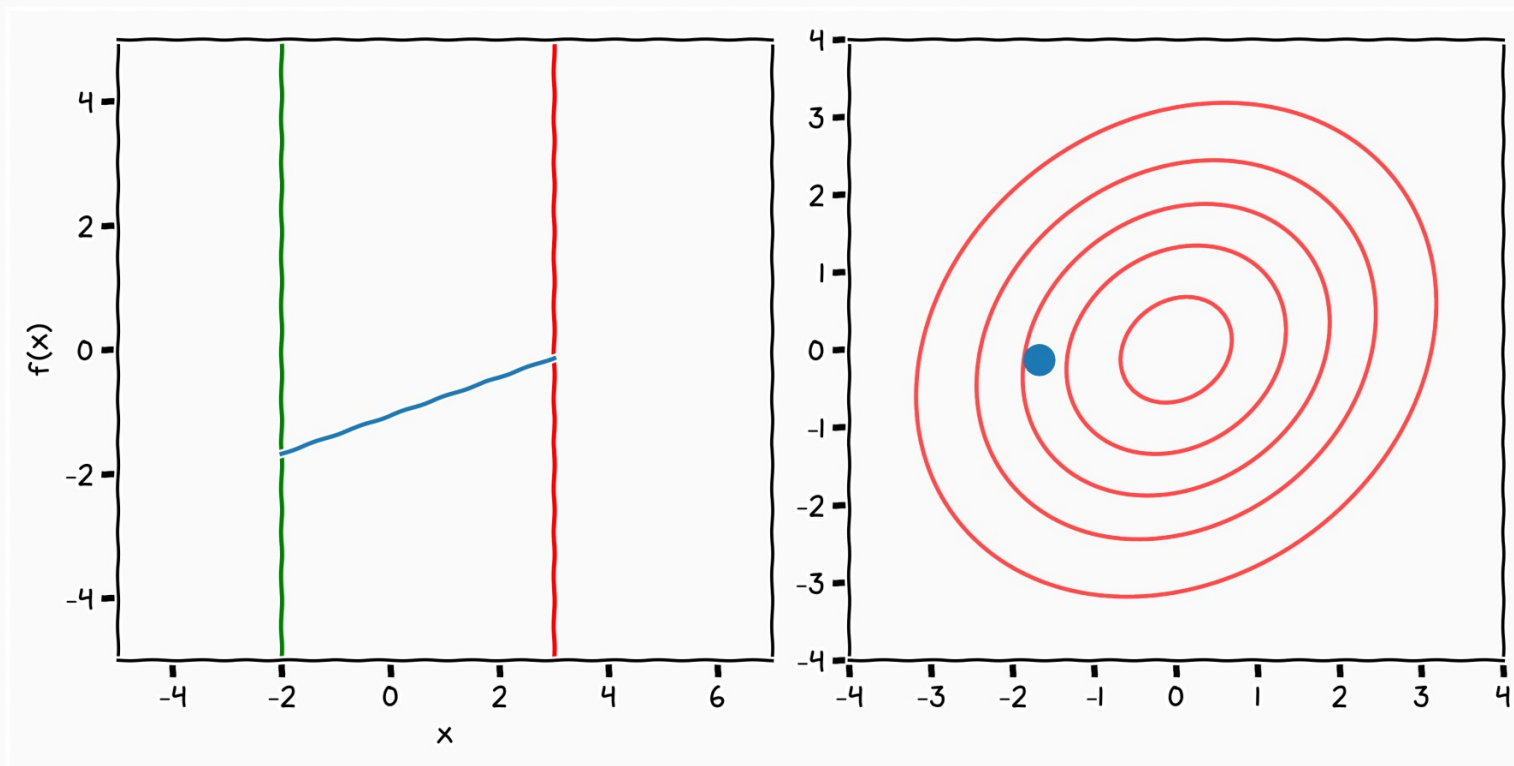
$$\mu = \begin{pmatrix} 0 \\ 0 \end{pmatrix}$$

$$\Sigma = \begin{pmatrix} 1 & 0.99 \\ 0.99 & 1 \end{pmatrix}$$

二元高斯分布



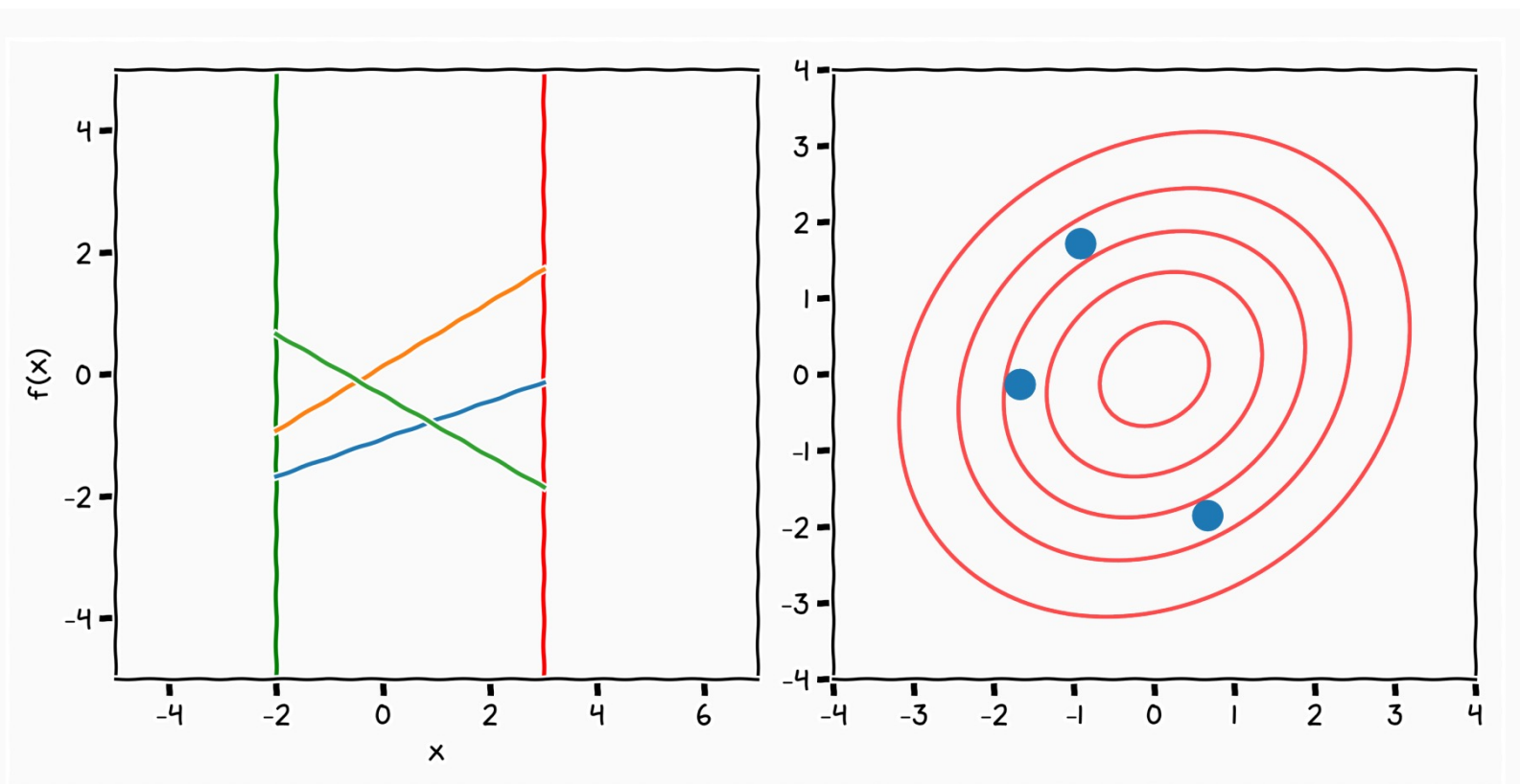
- 可以看作是从 Y_1 到 Y_2 的推断
- 右边的点=左边的线=单次采样



二元高斯分布



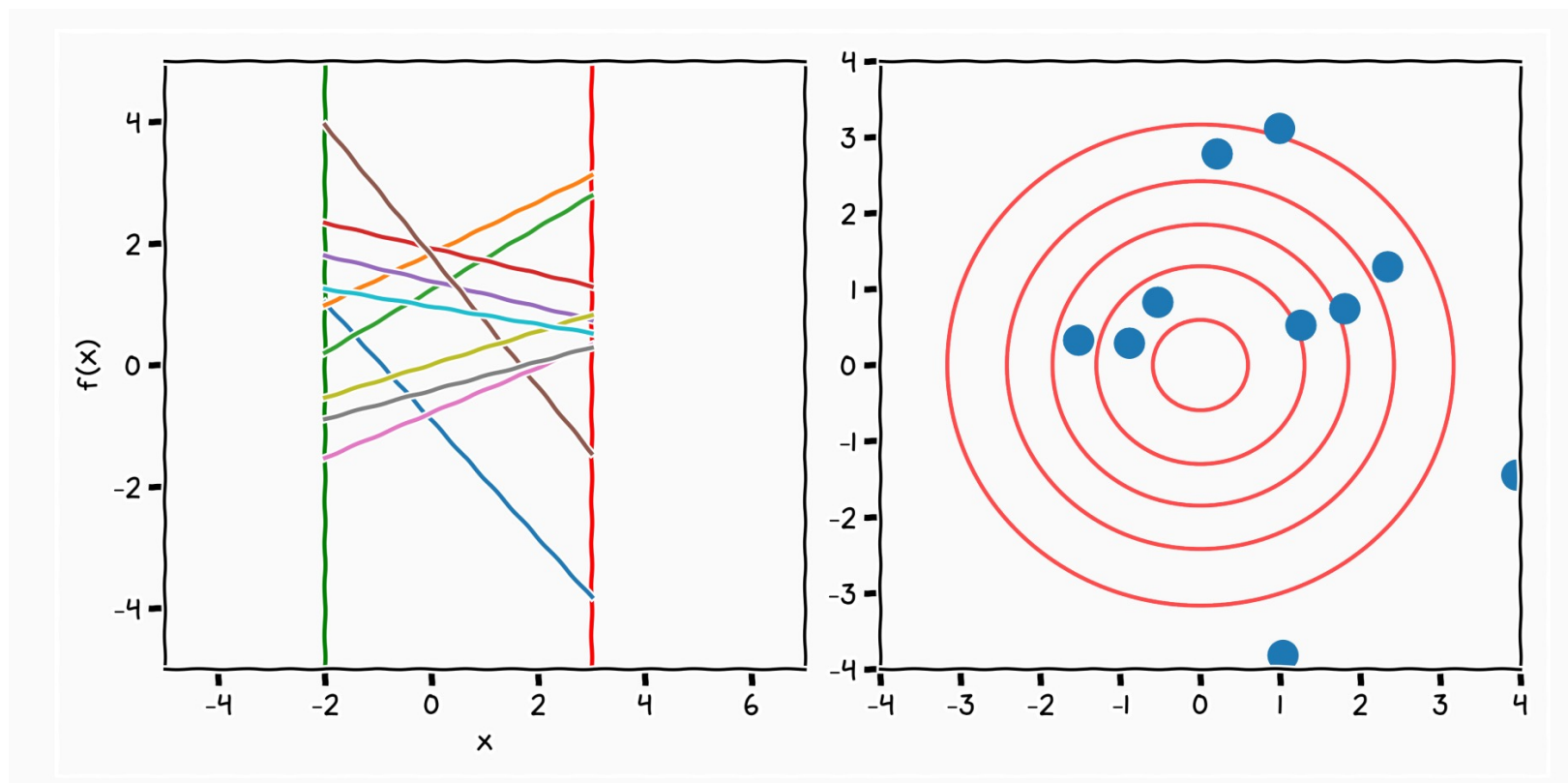
- 可以看作是从 Y_1 到 Y_2 的推断
- 右边的点=左边的线=多次采样



二元高斯分布



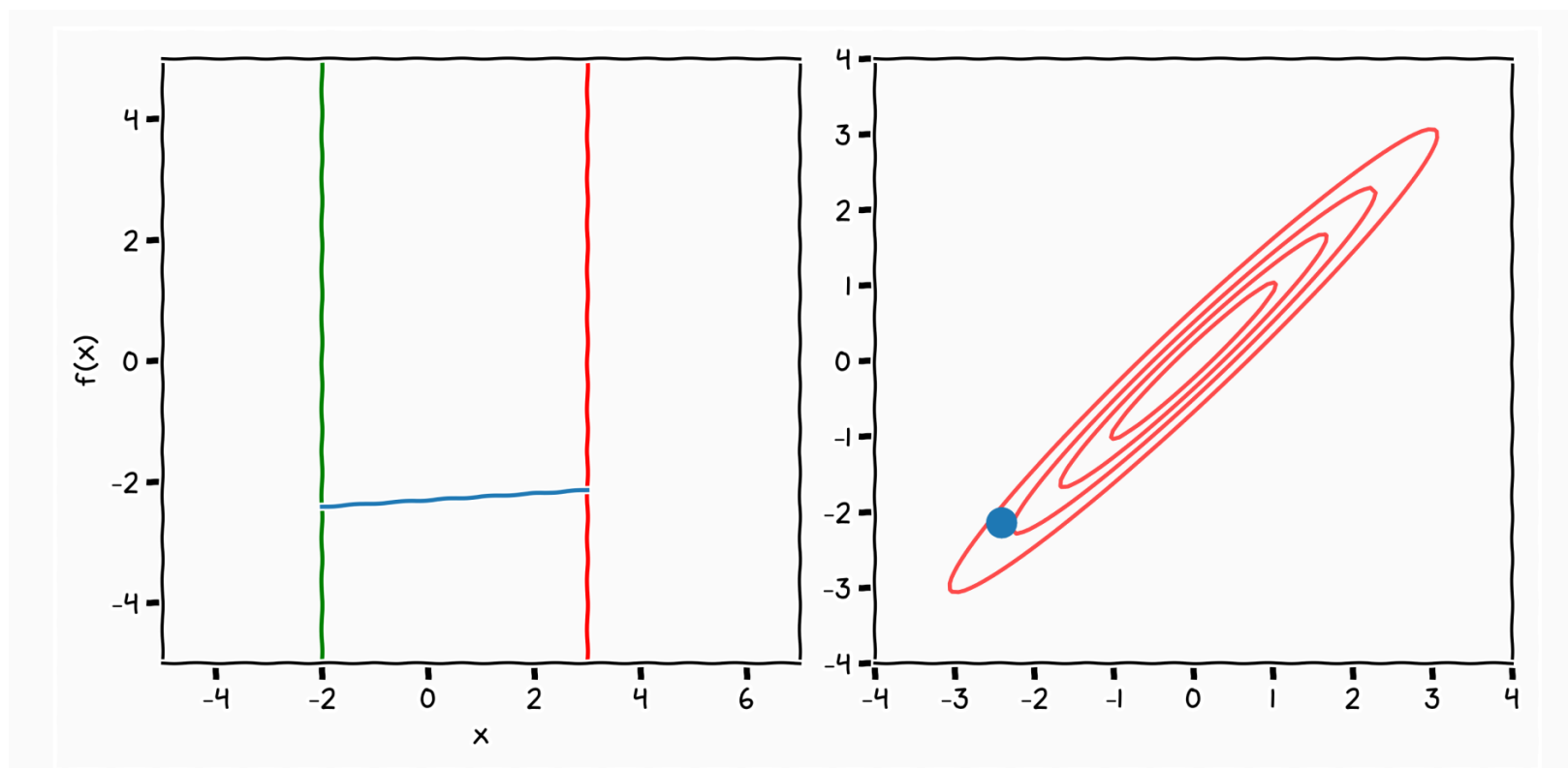
- 可以看作是从 Y_1 到 Y_2 的推断
- 右边的点=左边的线=多次采样，不相关



二元高斯分布



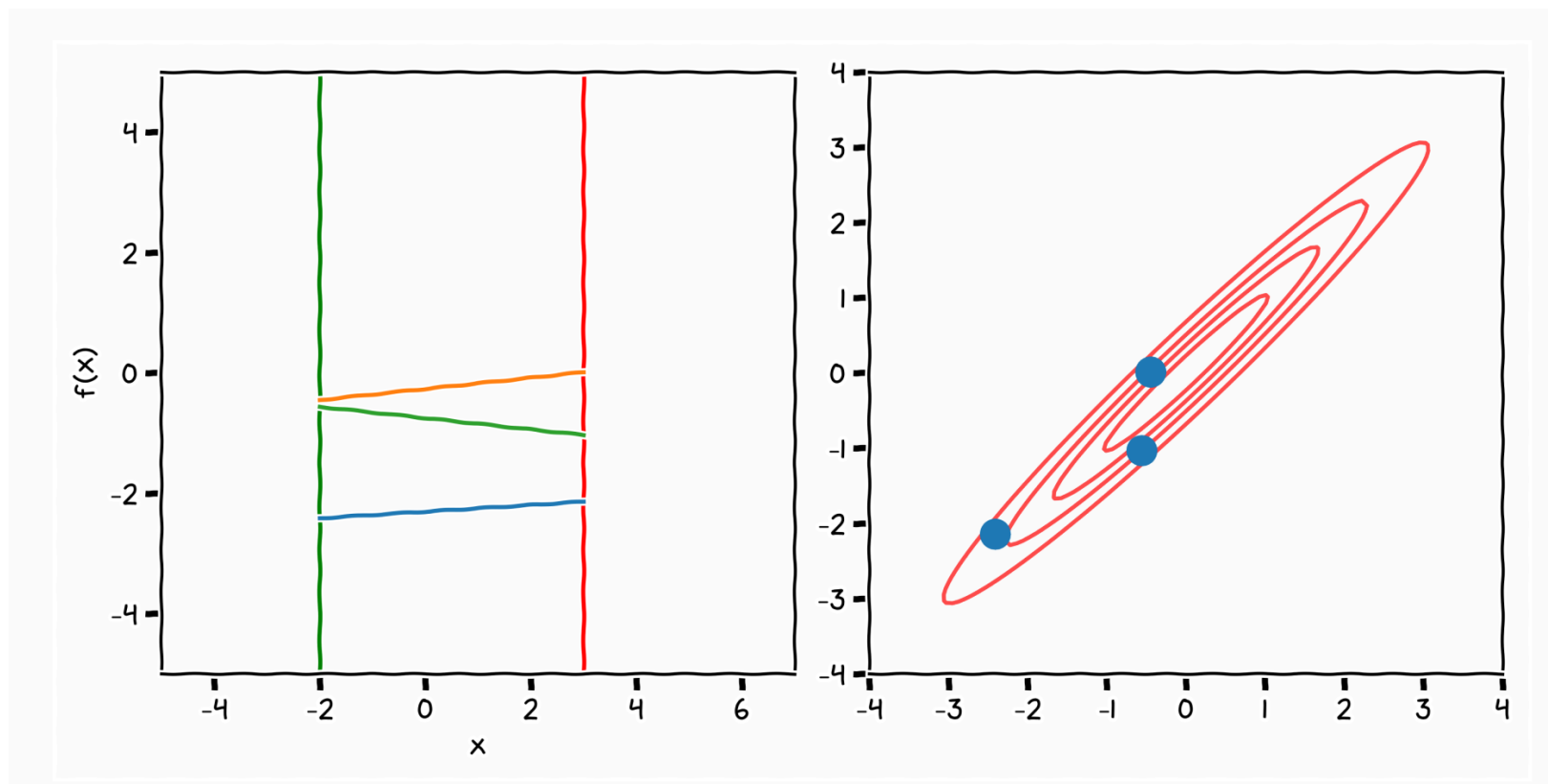
- 可以看作是从 Y_1 到 Y_2 的推断
- 右边的点=左边的线=单次采样



二元高斯分布



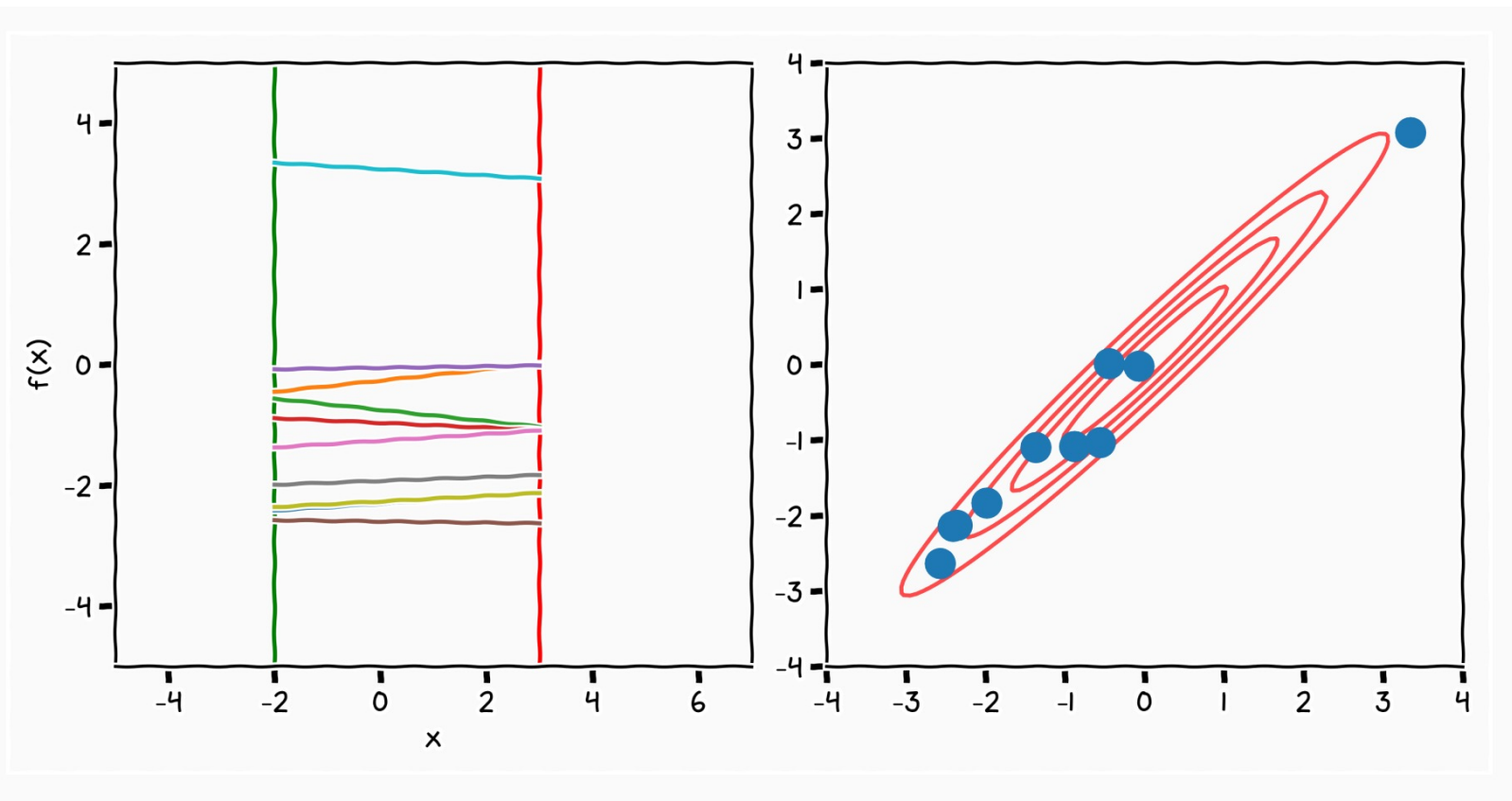
- 可以看作是从 Y_1 到 Y_2 的推断
- 右边的点=左边的线=多次采样



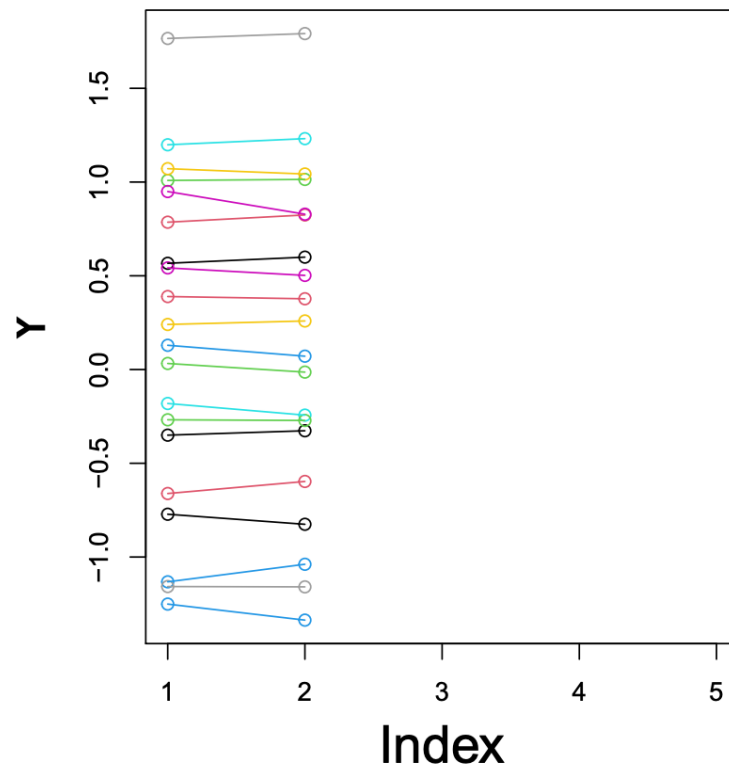
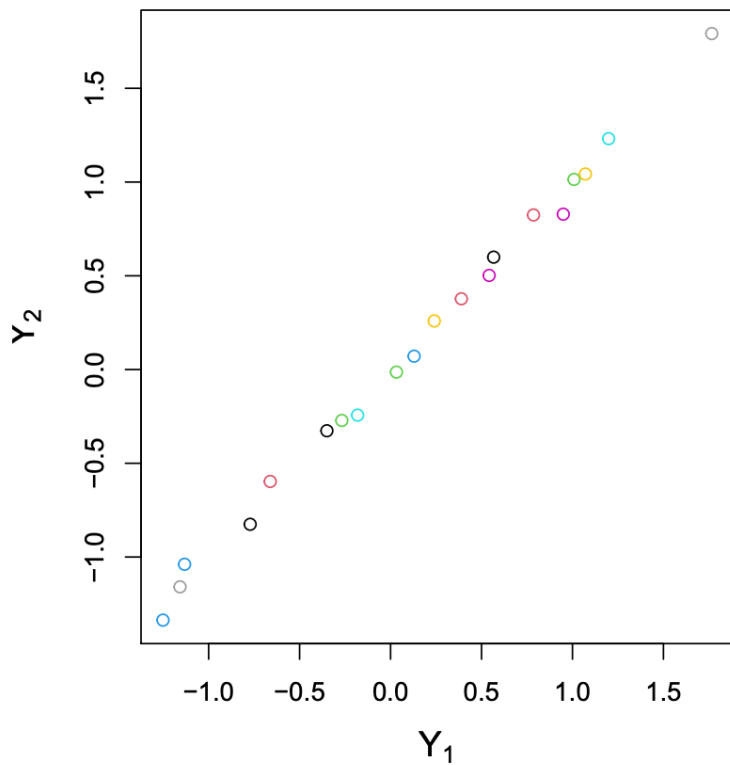
二元高斯分布



- 可以看作是从 Y_1 到 Y_2 的推断
- 右边的点=左边的线=多次采样，相关



- 可以看作是从 Y_1 到 Y_2 的推断

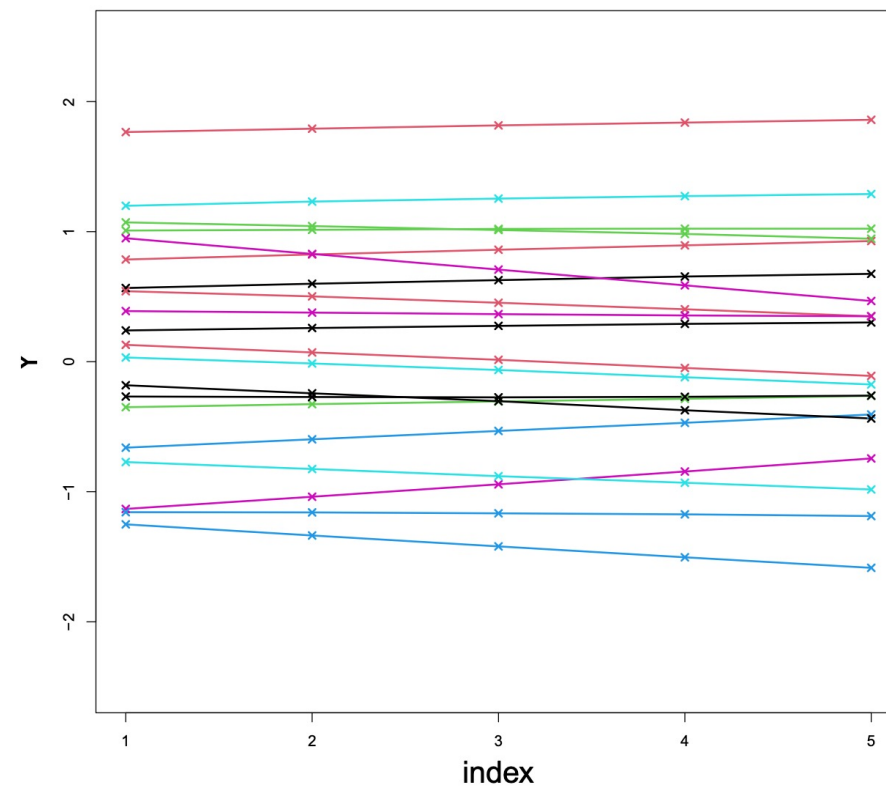


多元高斯分布



- $d=5$

$$\mu = \begin{pmatrix} 0 \\ 0 \\ \vdots \\ 0 \end{pmatrix} \quad \Sigma = \begin{pmatrix} 1 & 0.99 & 0.98 & 0.97 & 0.96 \\ 0.99 & 1 & 0.99 & 0.98 & 0.97 \\ 0.98 & 0.99 & 1 & 0.99 & 0.98 \\ 0.97 & 0.98 & 0.99 & 1 & 0.99 \\ 0.96 & 0.97 & 0.98 & 0.99 & 1 \end{pmatrix}$$

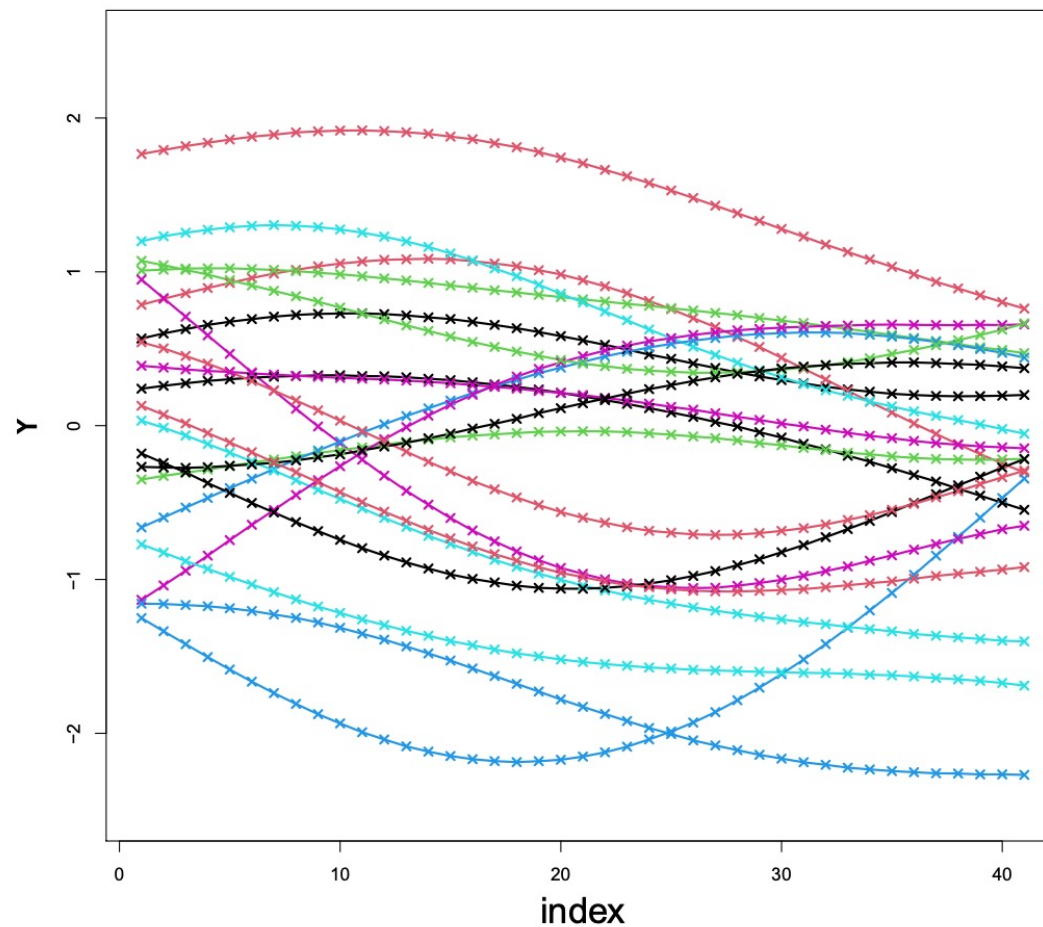


多元高斯分布



- $d=50$

$$\mu = \begin{pmatrix} 0 \\ 0 \\ \vdots \\ 0 \end{pmatrix} \quad \Sigma = \begin{pmatrix} 1 & 0.99 & 0.98 & 0.97 & 0.96 & \dots \\ 0.99 & 1 & 0.99 & 0.98 & 0.97 & \dots \\ 0.98 & 0.99 & 1 & 0.99 & 0.98 & \dots \\ 0.97 & 0.98 & 0.99 & 1 & 0.99 & \dots \\ 0.96 & 0.97 & 0.98 & 0.99 & 1 & \dots \\ \vdots & \vdots & \vdots & \vdots & \vdots & \ddots \end{pmatrix}$$



- 高斯过程可视为函数的无限维分布

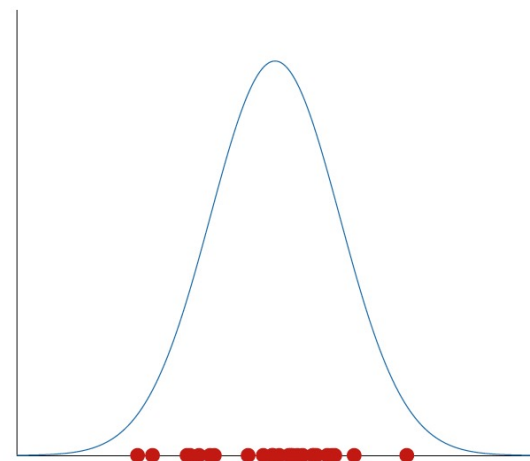
$$p(\mathbf{f}) = \mathcal{N} \left(\begin{bmatrix} f_1 \\ f_2 \\ \vdots \\ f_N \\ \vdots \end{bmatrix} \middle| \begin{bmatrix} \mu(x_1) \\ \mu(x_2) \\ \vdots \\ \mu(x_N) \\ \vdots \end{bmatrix}, \begin{bmatrix} k(x_1, x_1) & k(x_1, x_2) & \dots & k(x_1, x_N) & \dots \\ k(x_2, x_1) & k(x_2, x_2) & \dots & k(x_2, x_N) & \dots \\ \vdots & \vdots & \ddots & \vdots & \vdots \\ k(x_N, x_1) & k(x_N, x_2) & \dots & k(x_N, x_N) & \dots \\ \vdots & \vdots & \dots & \vdots & \ddots \end{bmatrix} \right)$$

- 高斯分布：需要计算均值 (μ) 和方差 (Σ)

$$X \sim N(\mu, \Sigma)$$

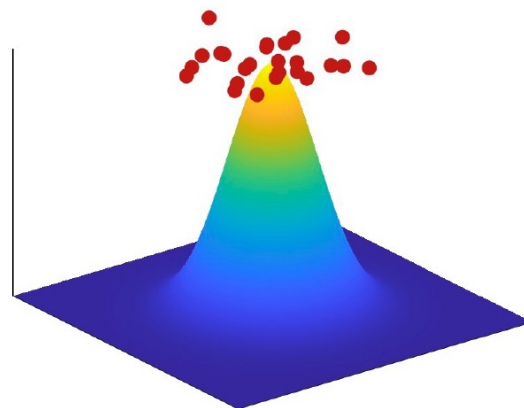
- (Univariate) Gaussians:

$$x \sim \mathcal{N}(x; \mu = 0, \sigma^2 = 1)$$



$$\mu = \begin{pmatrix} \mu_1 \\ \mu_2 \end{pmatrix} \quad \Sigma = \begin{pmatrix} \sigma_1^2 & \rho_{12}\sigma_1\sigma_2 \\ \rho_{21}\sigma_1\sigma_2 & \sigma_2^2 \end{pmatrix}$$

$$\text{Cor}(Y_i, Y_j) = \rho_{12} \text{ for } i \neq j$$



- Multivariate Gaussians:

$$\mathbf{x} = [x_1, \dots, x_D]^T \\ \sim \mathcal{N}(\mathbf{x}; \boldsymbol{\mu} = \mathbf{0}_D, \Sigma = I_D)$$

高斯过程：函数的高斯分布

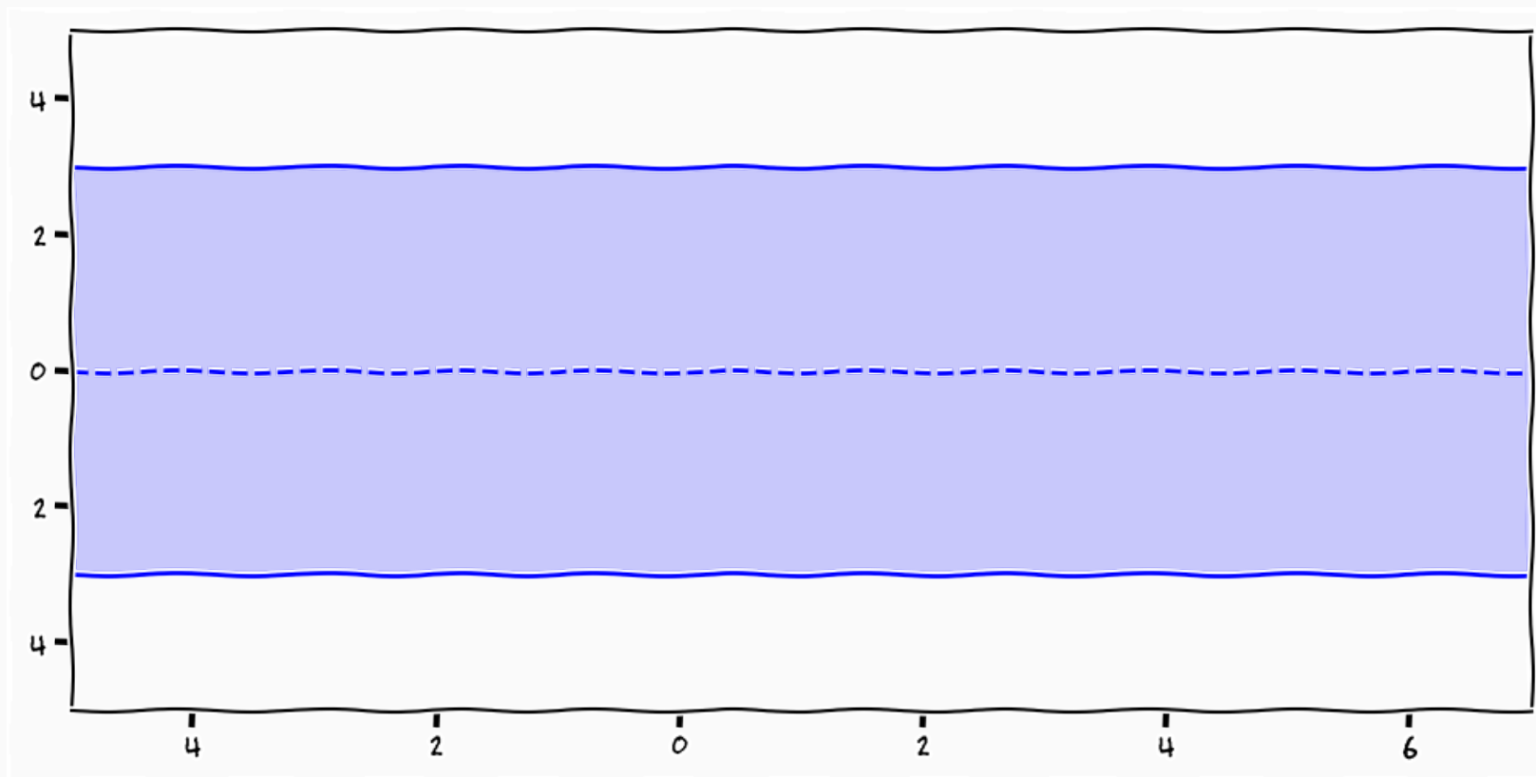


- 高斯过程：需要计算函数均值 ($m(\cdot)$) 和函数方差 ($k(\cdot, \cdot)$)

$$y(\cdot) \sim GP(m(\cdot), k(\cdot, \cdot))$$

$$\mathbb{E}(y(x)) = m(x)$$

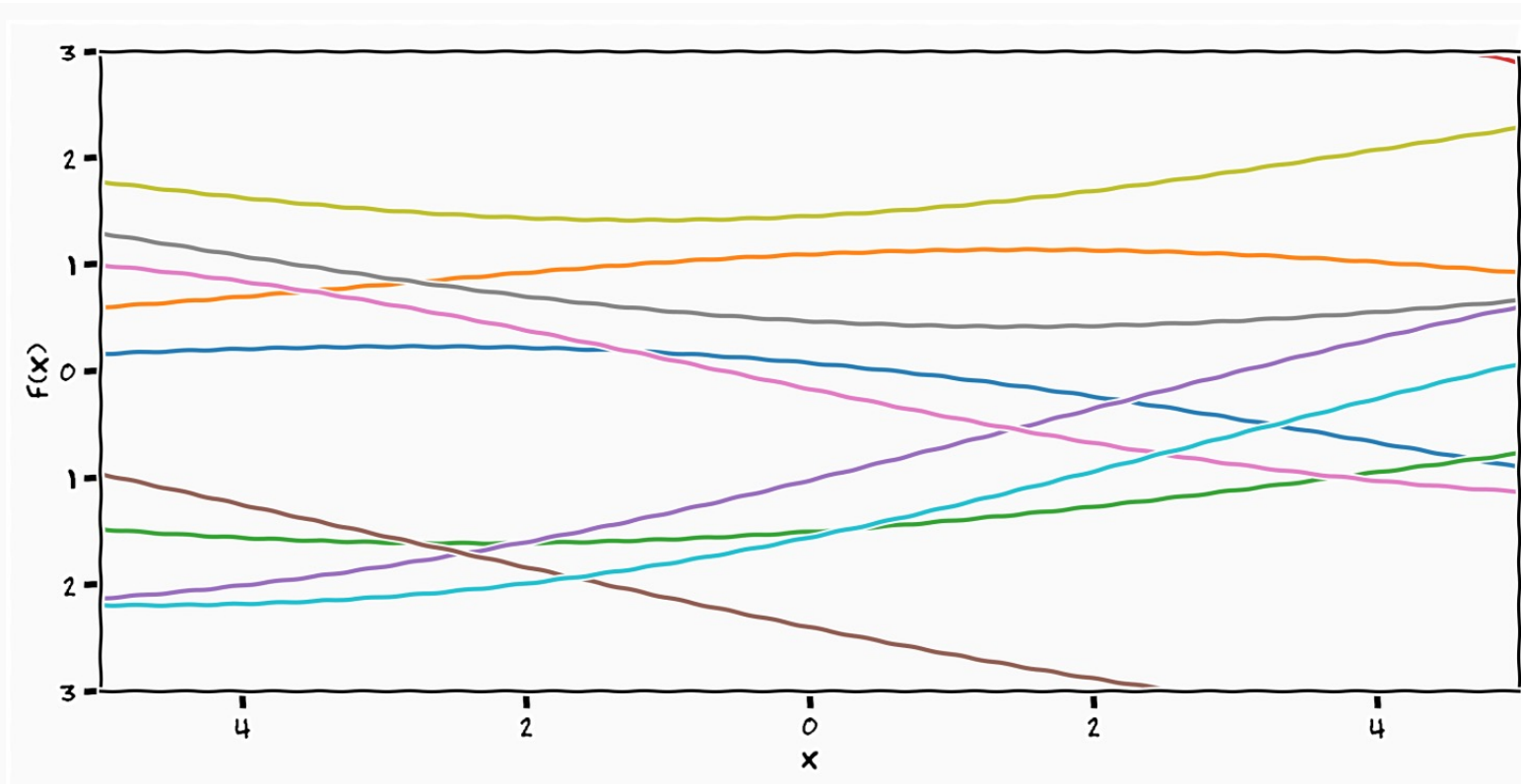
$$\text{Cov}(y(x), y(x')) = k(x, x')$$



高斯过程：函数的高斯分布



- 每个时间点比较相关, $k(x, x')$ 小



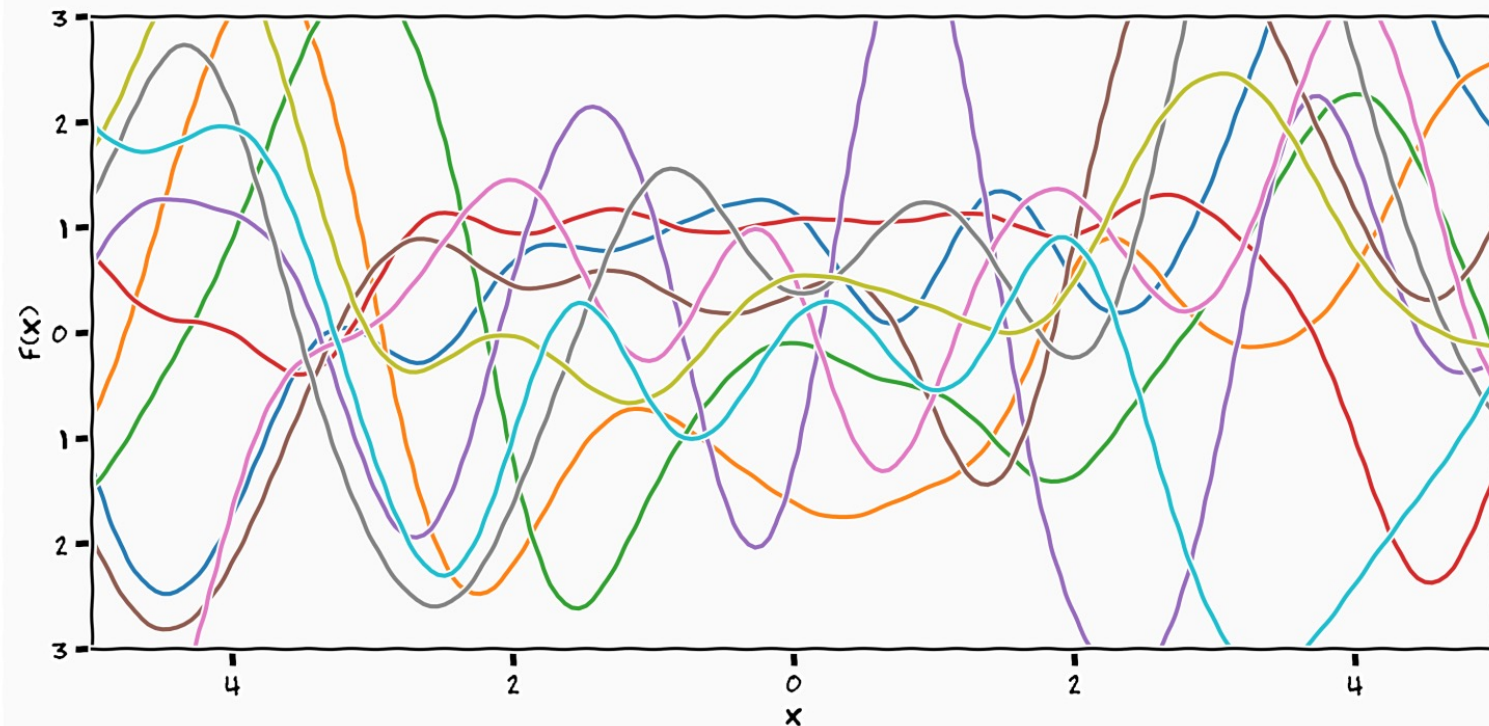
$$k(x_i, x_j) = 3 \cdot e^{-\frac{(x_i - x_j)^2}{150}}$$



高斯过程：函数的高斯分布



- 每个时间点不太相关, $k(x, x')$ 大

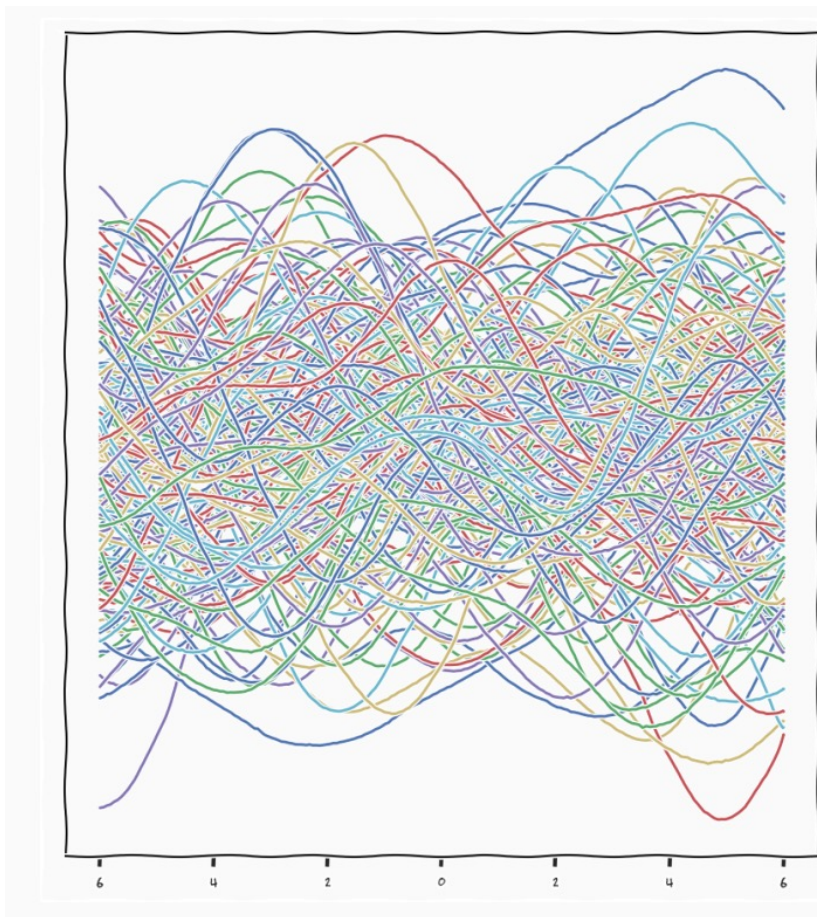


$$k(x_i, x_j) = 3 \cdot e^{-\frac{(x_i - x_j)^2}{1}}$$



高斯过程的后验更新

先验假设



后验假设



二元高斯分布的条件更新



Suppose

$$Y = \begin{pmatrix} Y_1 \\ Y_2 \end{pmatrix} \sim N(\mu, \Sigma)$$

where

$$\mu = \begin{pmatrix} \mu_1 \\ \mu_2 \end{pmatrix} \quad \Sigma = \begin{pmatrix} \Sigma_{11} & \Sigma_{12} \\ \Sigma_{21} & \Sigma_{22} \end{pmatrix}$$

Then

$$Y_2 \mid Y_1 = y_1 \sim N(\mu_2 + \Sigma_{21}\Sigma_{11}^{-1}(y_1 - \mu_1), \Sigma_{22} - \Sigma_{21}\Sigma_{11}^{-1}\Sigma_{12})$$



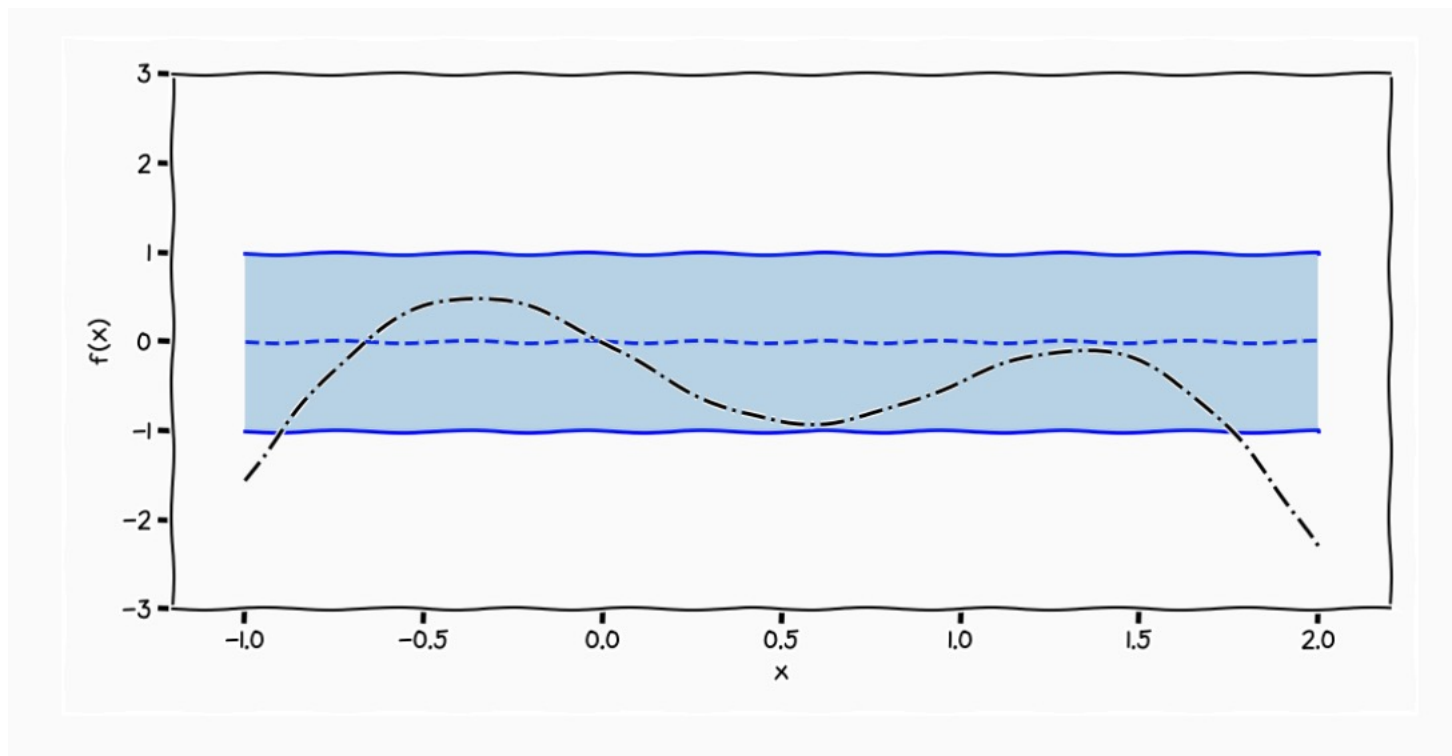
- 利用贝叶斯法则，利用所有在 $\mathbf{x}$ 观测到的 $\mathbf{f}$ ， $\mathbf{f}_*$ 为未观测的点：

$$p(\mathbf{f}_* | \mathbf{f}) = \frac{p(\mathbf{f}, \mathbf{f}_*)}{p(\mathbf{f})} = \frac{p(\mathbf{f}, \mathbf{f}_*)}{\int p(\mathbf{f}, \mathbf{f}_*) d\mathbf{f}_*}$$

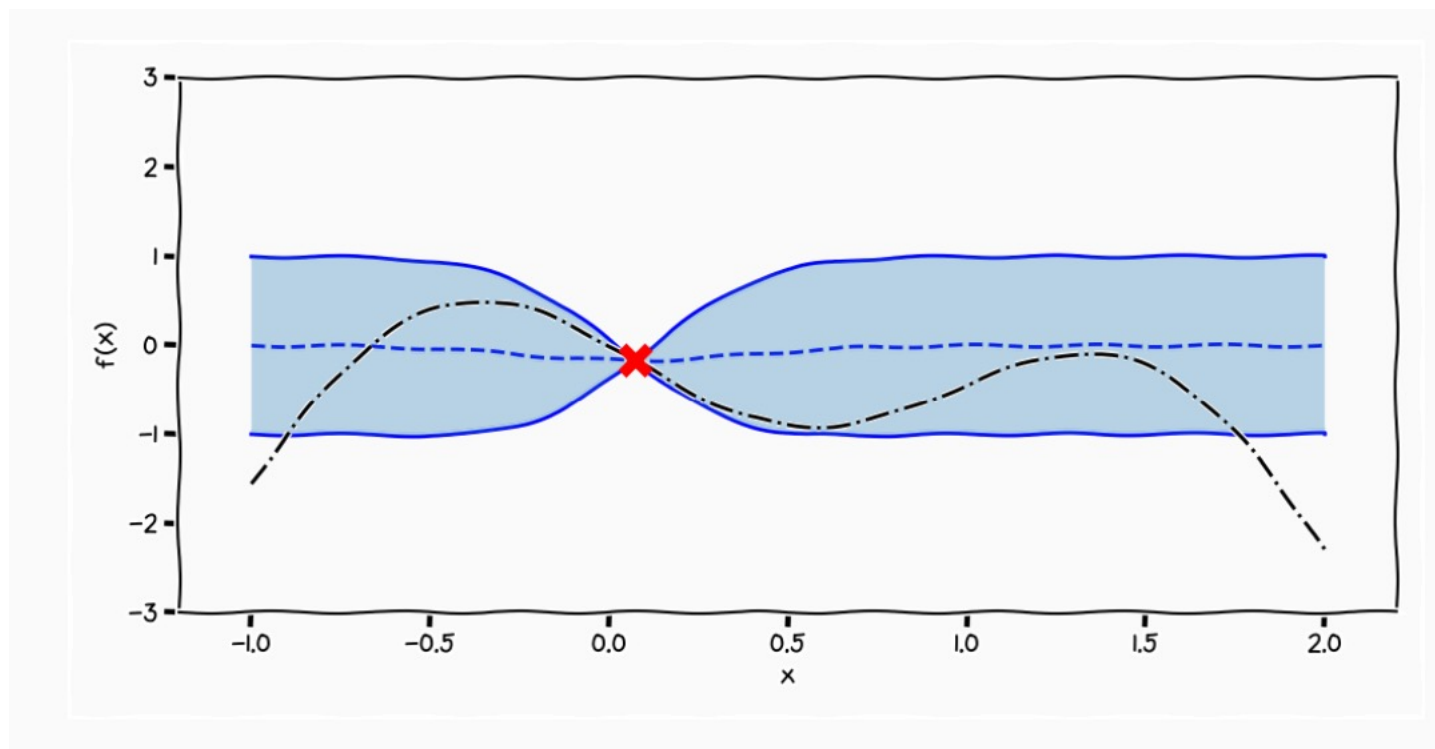
- 在观测 $\mathbf{f}$ 之后，仍然符合高斯假设：

$$p(f_* | \mathbf{x}_*, \mathbf{x}, \mathbf{f}) = \mathcal{N}(k(\mathbf{x}_*, \mathbf{x})^T k(\mathbf{x}, \mathbf{x})^{-1} \mathbf{f}, \\ k(\mathbf{x}_*, \mathbf{x}_*) - k(\mathbf{x}_*, \mathbf{x})^T k(\mathbf{x}, \mathbf{x})^{-1} k(\mathbf{x}, \mathbf{x}_*))$$

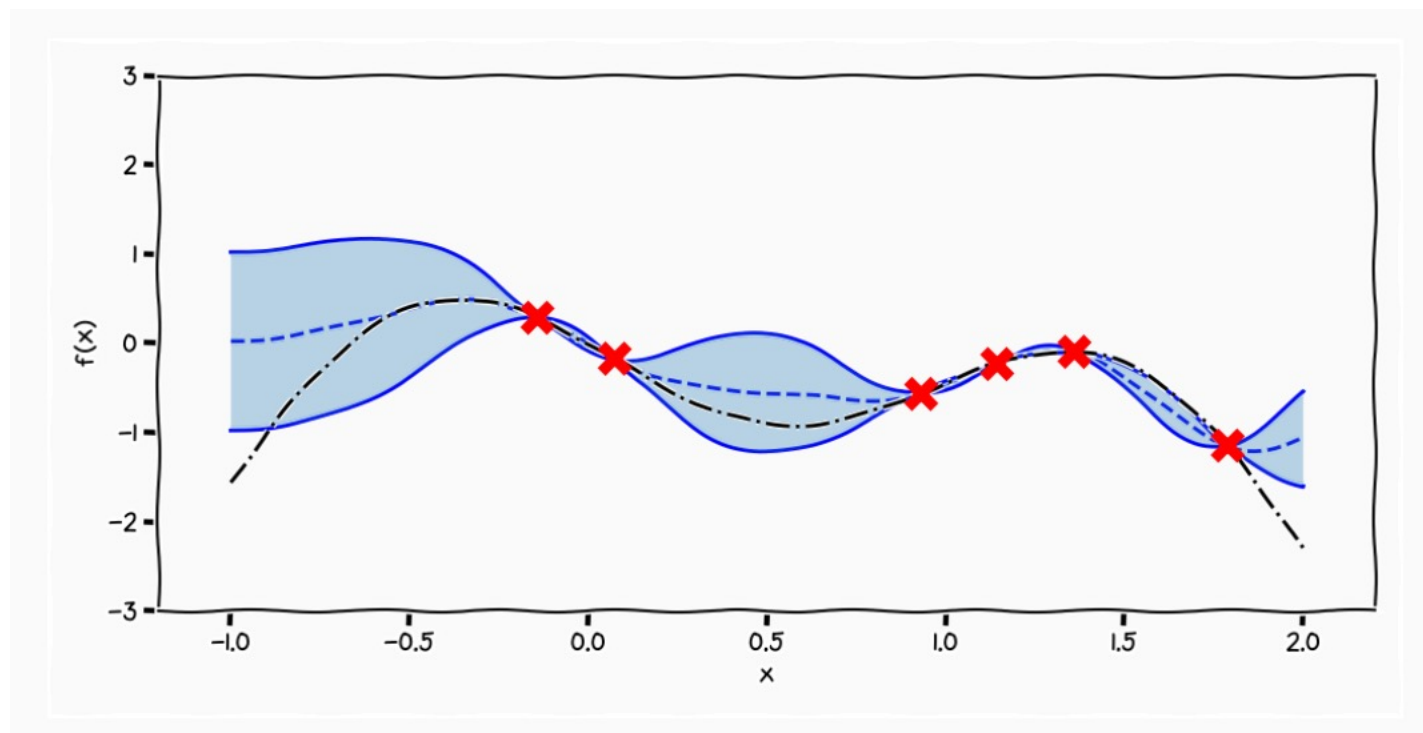
高斯过程的条件更新



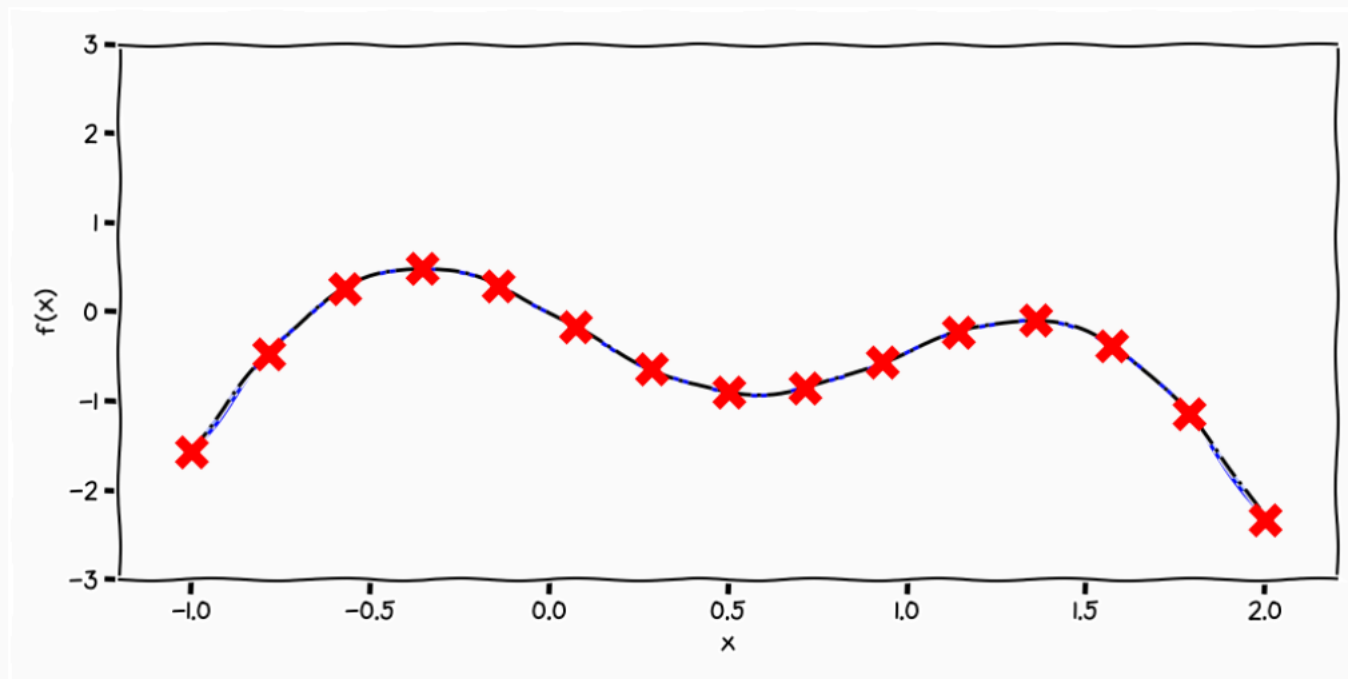
高斯过程的条件更新



高斯过程的条件更新



高斯过程的条件更新



- 假设观测也是有噪声的

$$\begin{bmatrix} \mathbf{f} \\ \mathbf{f}_* \end{bmatrix} \sim \mathcal{N} \left(\begin{bmatrix} \mathbf{0} \\ \mathbf{0} \end{bmatrix}, \begin{bmatrix} k(\mathbf{x}, \mathbf{x}) & k(\mathbf{x}, \mathbf{x}_*) \\ k(\mathbf{x}_*, \mathbf{x}) & k(\mathbf{x}_*, \mathbf{x}_*) \end{bmatrix} \right)$$

$$p(f_* | \mathbf{x}_*, \mathbf{x}, \mathbf{f}) = \mathcal{N}(k(\mathbf{x}_*, \mathbf{x})^\top k(\mathbf{x}, \mathbf{x})^{-1} \mathbf{f}, \\ k(\mathbf{x}_*, \mathbf{x}_*) - k(\mathbf{x}_*, \mathbf{x})^\top k(\mathbf{x}, \mathbf{x})^{-1} k(\mathbf{x}, \mathbf{x}_*))$$



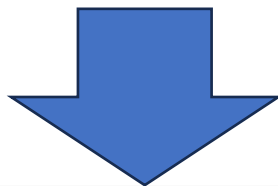
$$\begin{bmatrix} \mathbf{y} \\ \mathbf{f}_* \end{bmatrix} \sim \mathcal{N} \left(\begin{bmatrix} \mathbf{0} \\ \mathbf{0} \end{bmatrix}, \begin{bmatrix} k(\mathbf{x}, \mathbf{x}) + \sigma^2 \mathbf{I} & k(\mathbf{x}, \mathbf{x}_*) \\ k(\mathbf{x}_*, \mathbf{x}) & k(\mathbf{x}_*, \mathbf{x}_*) \end{bmatrix} \right)$$

$$p(f_* | \mathbf{x}_*, \mathbf{x}, \mathbf{y}, \boldsymbol{\theta}) = \mathcal{N}(k(\mathbf{x}_*, \mathbf{x})^\top (K(\mathbf{x}, \mathbf{x}) + \sigma^2 \mathbf{I})^{-1} \mathbf{y}, \\ k(\mathbf{x}_*, \mathbf{x}_*) - k(\mathbf{x}_*, \mathbf{x})^\top (K(\mathbf{x}, \mathbf{x}) + \sigma^2 \mathbf{I})^{-1} K(\mathbf{x}, \mathbf{x}_*))$$

- 假设观测也是有噪声的（一个超参数）

$$\begin{bmatrix} \mathbf{f} \\ \mathbf{f}_* \end{bmatrix} \sim \mathcal{N} \left(\begin{bmatrix} \mathbf{0} \\ \mathbf{0} \end{bmatrix}, \begin{bmatrix} k(\mathbf{x}, \mathbf{x}) & k(\mathbf{x}, \mathbf{x}_*) \\ k(\mathbf{x}_*, \mathbf{x}) & k(\mathbf{x}_*, \mathbf{x}_*) \end{bmatrix} \right)$$

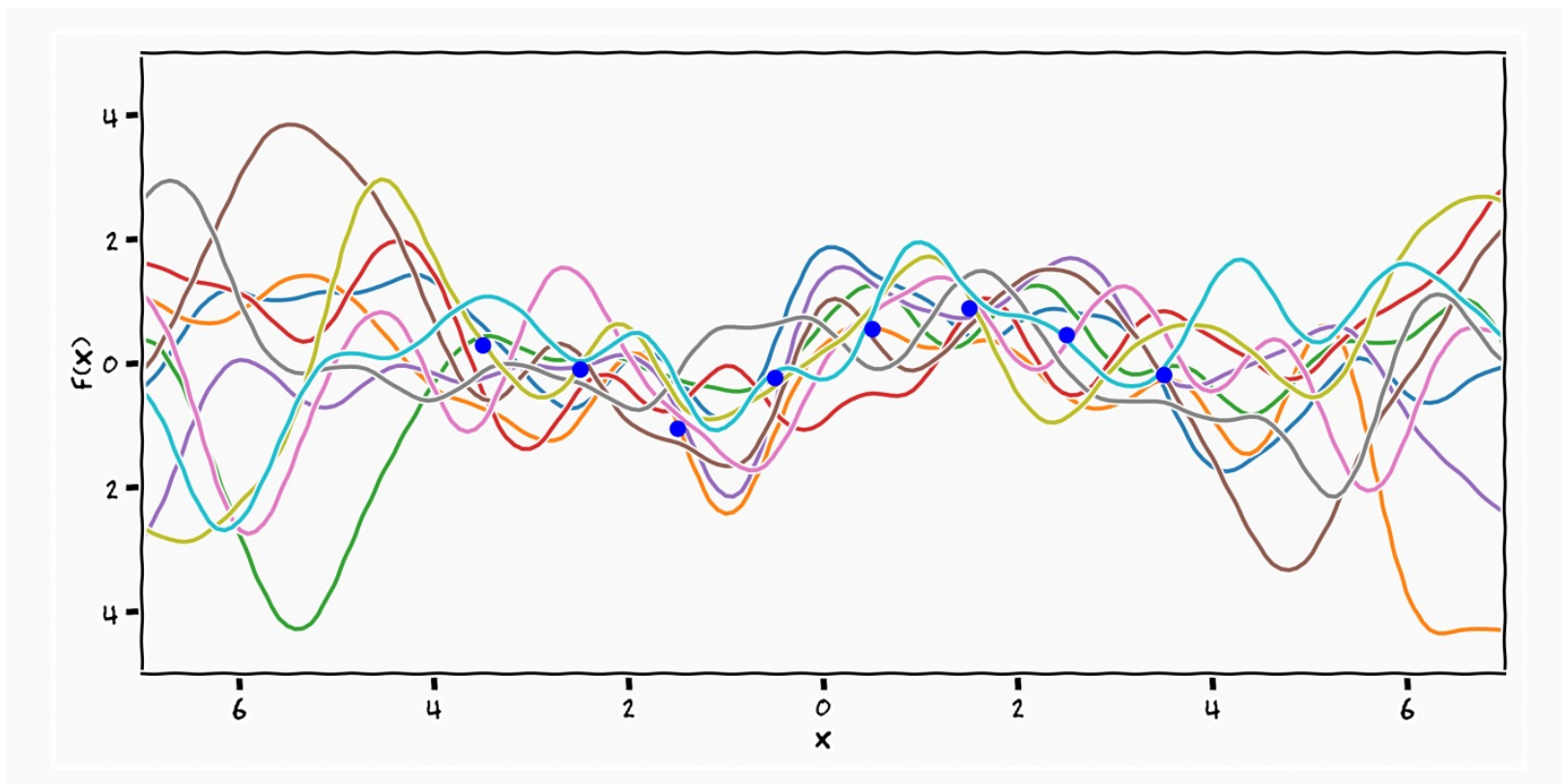
$$p(f_* | \mathbf{x}_*, \mathbf{x}, \mathbf{f}) = \mathcal{N}(k(\mathbf{x}_*, \mathbf{x})^\top k(\mathbf{x}, \mathbf{x})^{-1} \mathbf{f}, \\ k(\mathbf{x}_*, \mathbf{x}_*) - k(\mathbf{x}_*, \mathbf{x})^\top k(\mathbf{x}, \mathbf{x})^{-1} k(\mathbf{x}, \mathbf{x}_*))$$



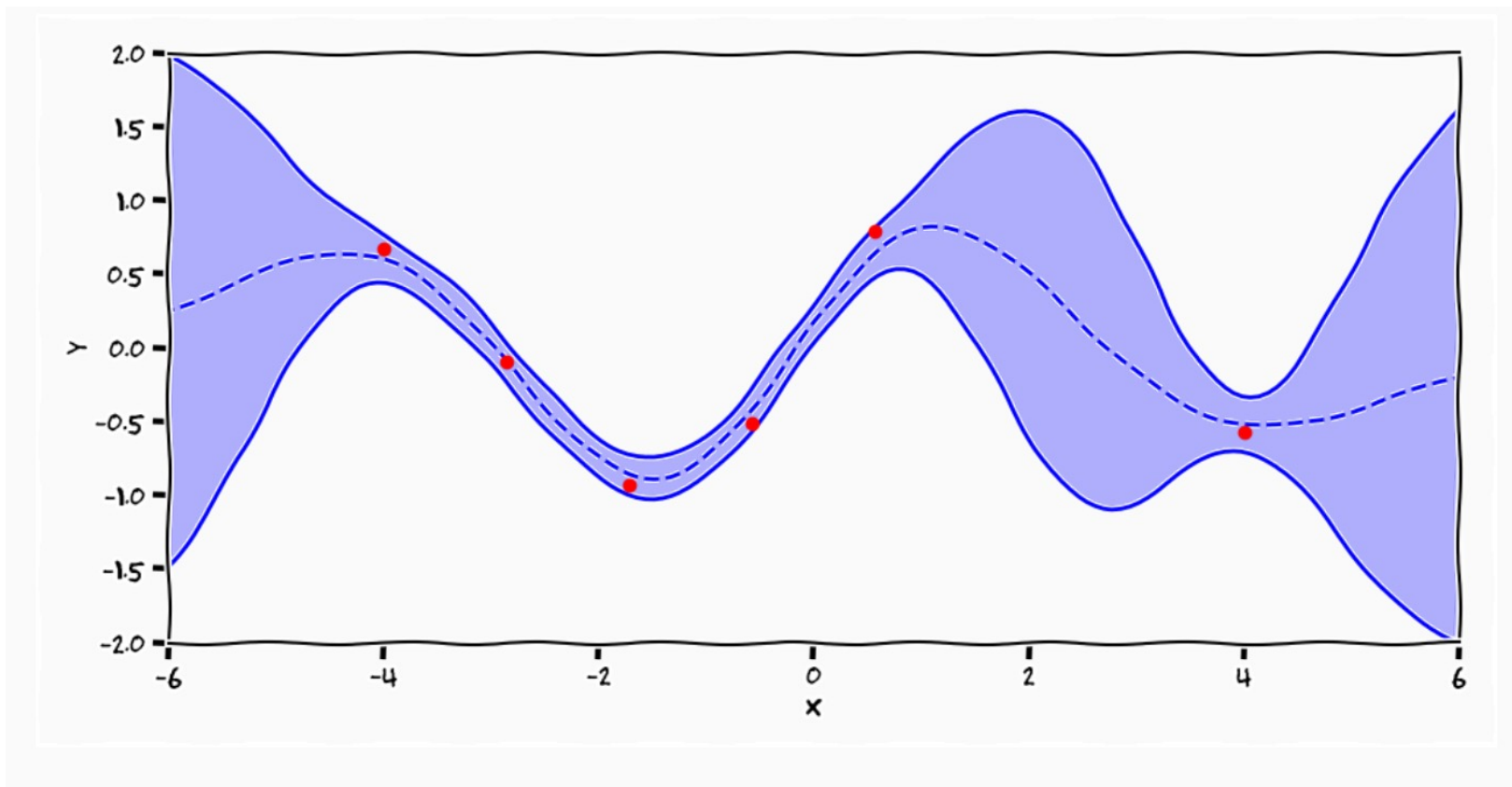
$$\begin{bmatrix} \mathbf{y} \\ \mathbf{f}_* \end{bmatrix} \sim \mathcal{N} \left(\begin{bmatrix} \mathbf{0} \\ \mathbf{0} \end{bmatrix}, \begin{bmatrix} k(\mathbf{x}, \mathbf{x}) + \sigma^2 \mathbf{I} & k(\mathbf{x}, \mathbf{x}_*) \\ k(\mathbf{x}_*, \mathbf{x}) & k(\mathbf{x}_*, \mathbf{x}_*) \end{bmatrix} \right)$$

$$p(f_* | \mathbf{x}_*, \mathbf{x}, \mathbf{y}, \boldsymbol{\theta}) = \mathcal{N}(k(\mathbf{x}_*, \mathbf{x})^\top (K(\mathbf{x}, \mathbf{x}) + \sigma^2 \mathbf{I})^{-1} \mathbf{y}, \\ k(\mathbf{x}_*, \mathbf{x}_*) - k(\mathbf{x}_*, \mathbf{x})^\top (K(\mathbf{x}, \mathbf{x}) + \sigma^2 \mathbf{I})^{-1} K(\mathbf{x}, \mathbf{x}_*))$$

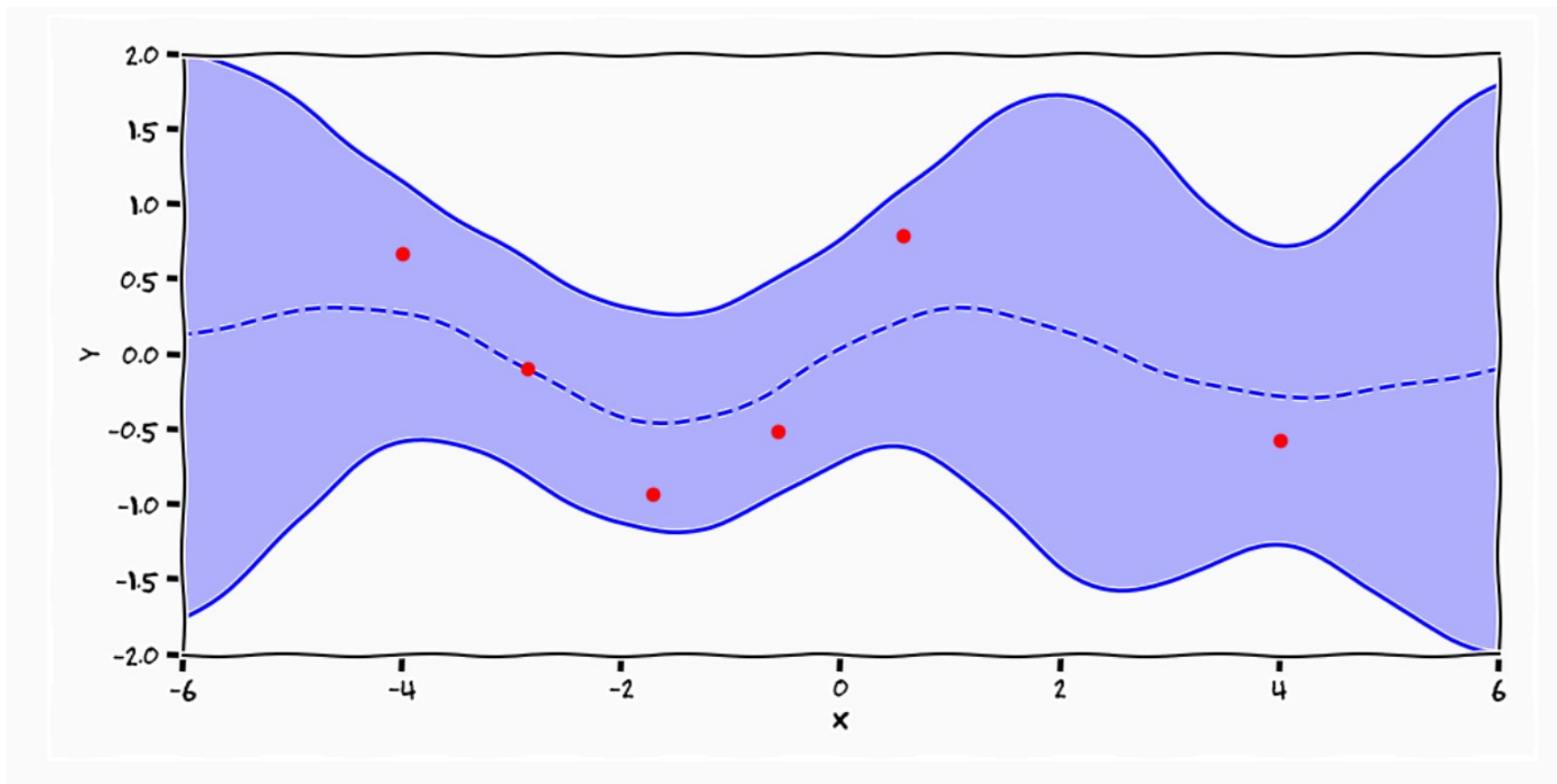
- 假设观测也是有噪声的



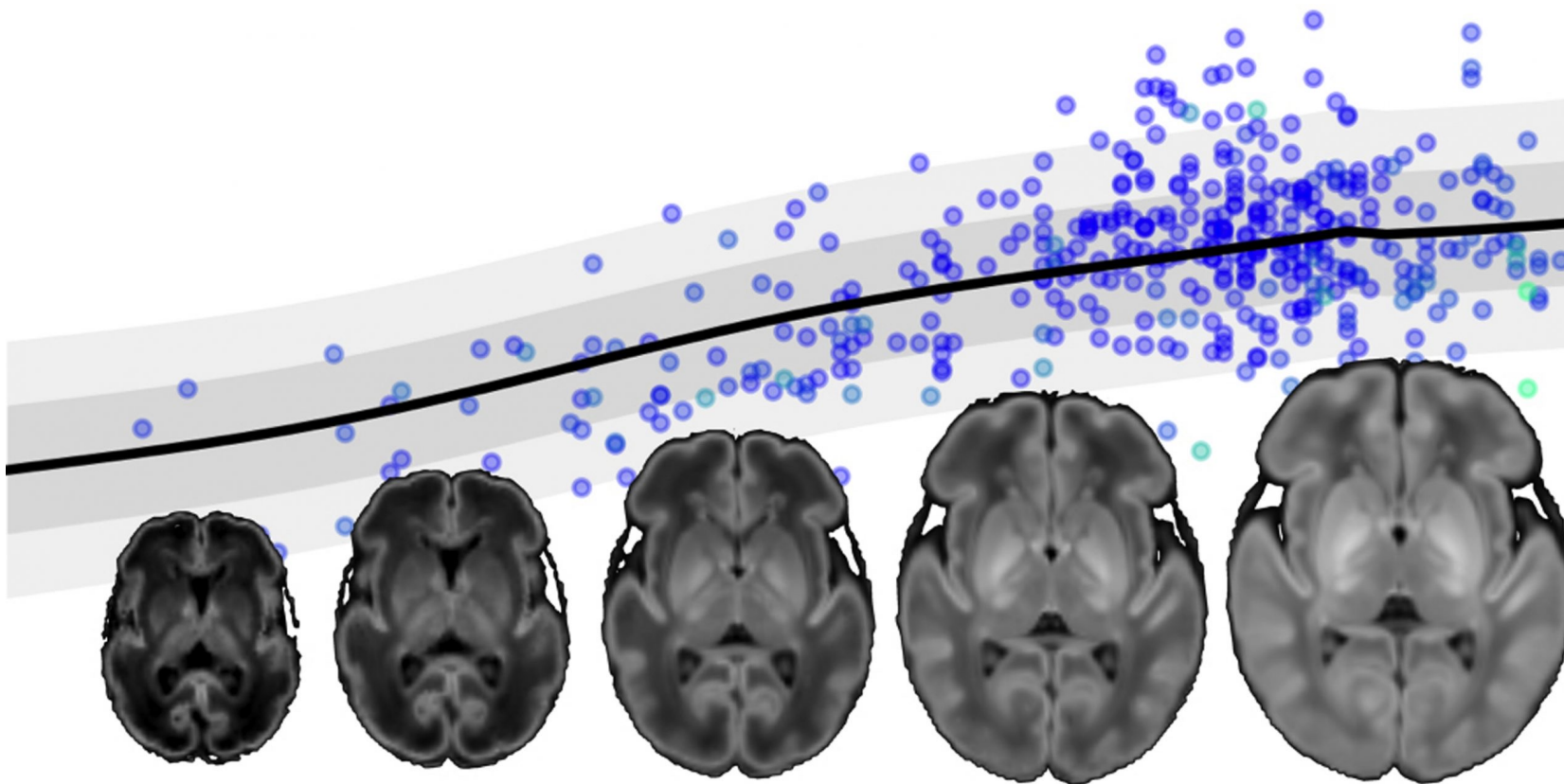
- 噪声比较小时



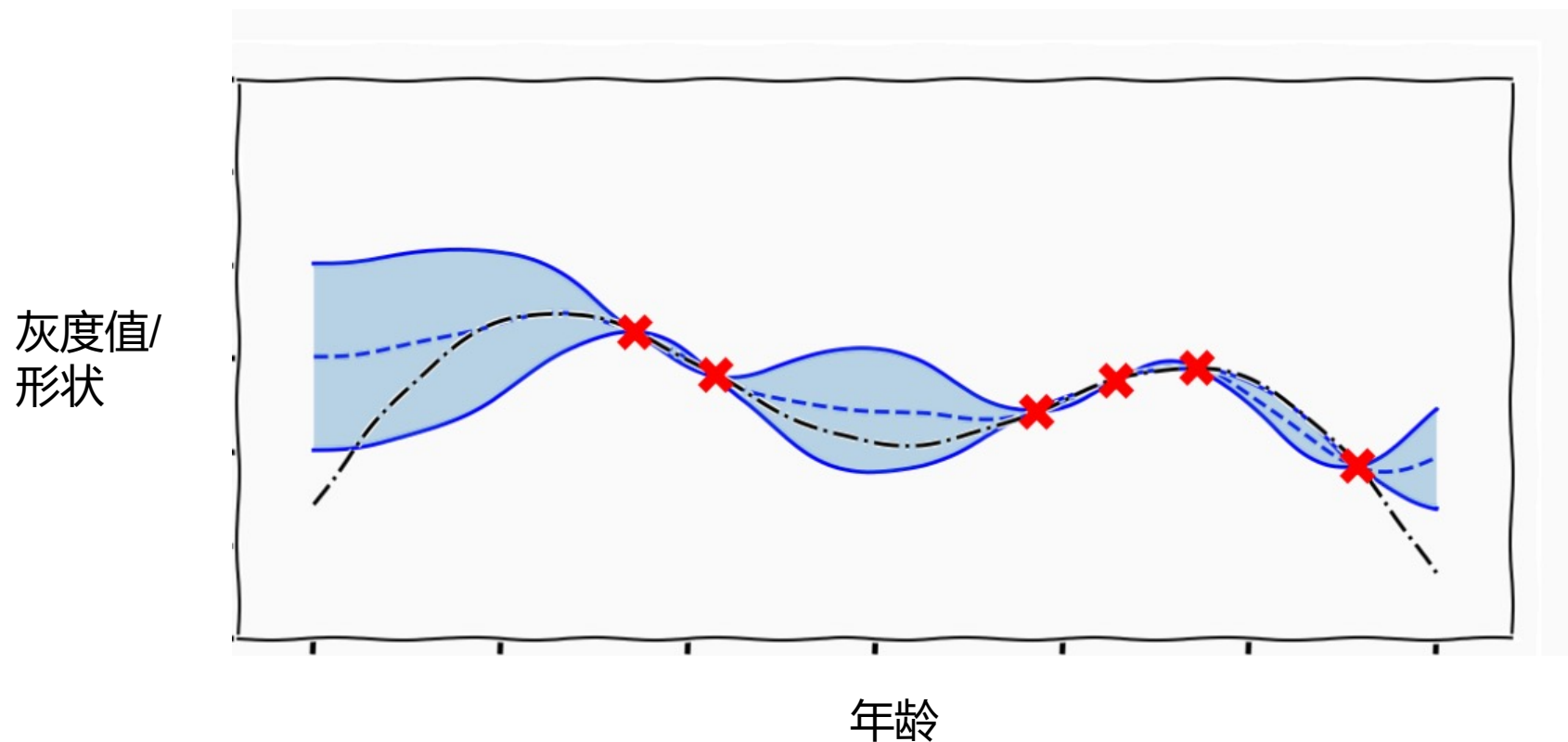
- 噪声比较大时



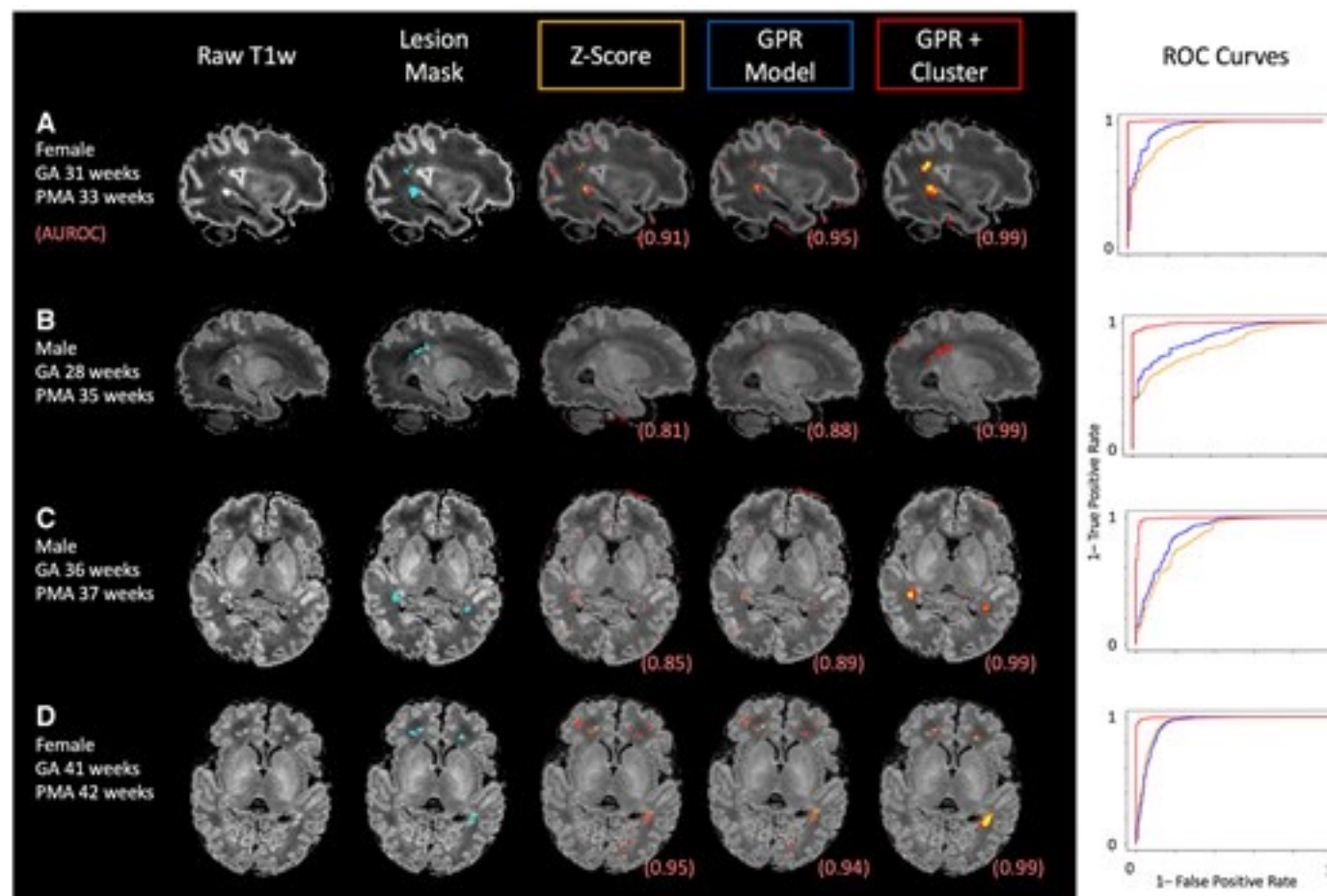
- 非常适合时序、空间连续建模



- 利用高斯过程来对脑成长过程建模



- 与真实结果对比，检测白质病变



Modelling brain development to detect white matter injury in term and preterm born neonates

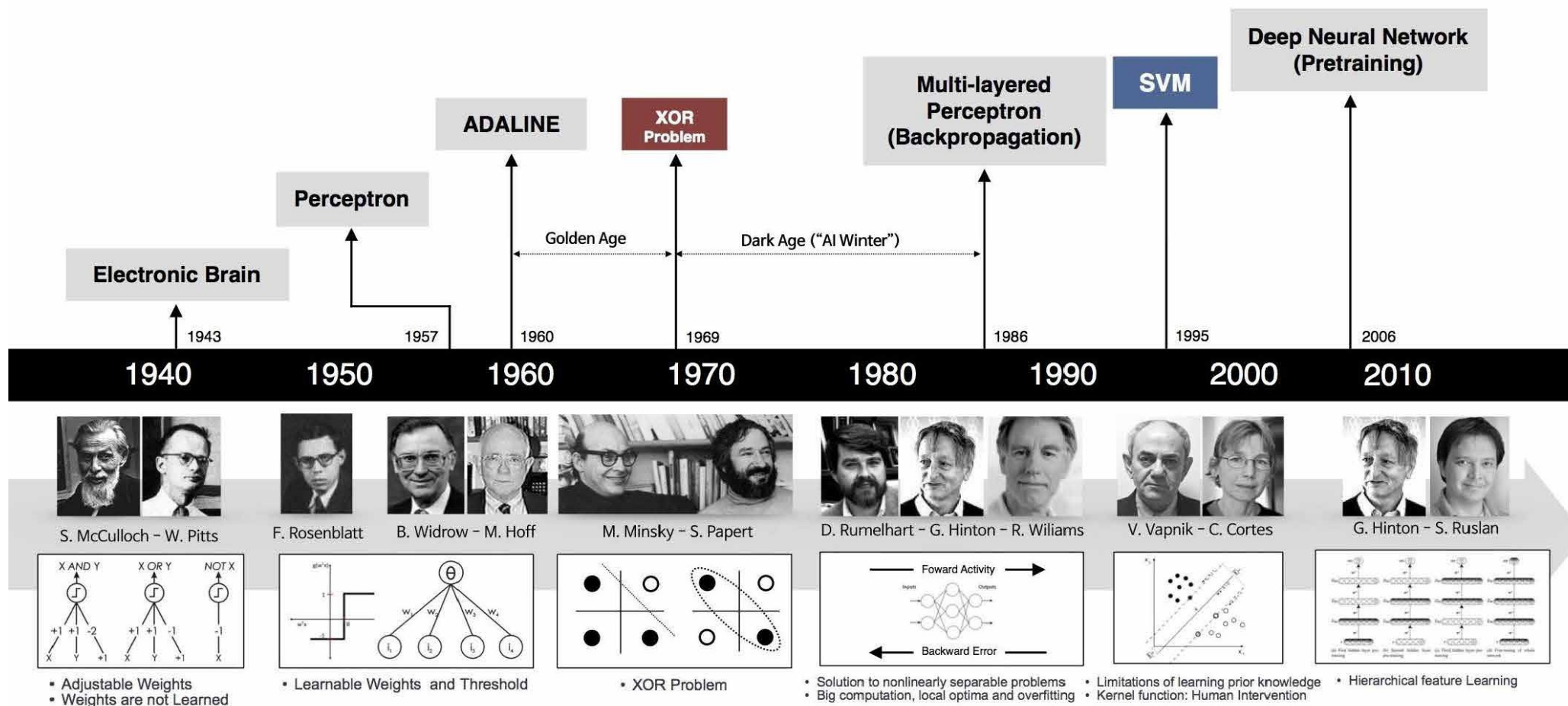
目录

1 高斯过程

2 CNN

3 医学图像分割

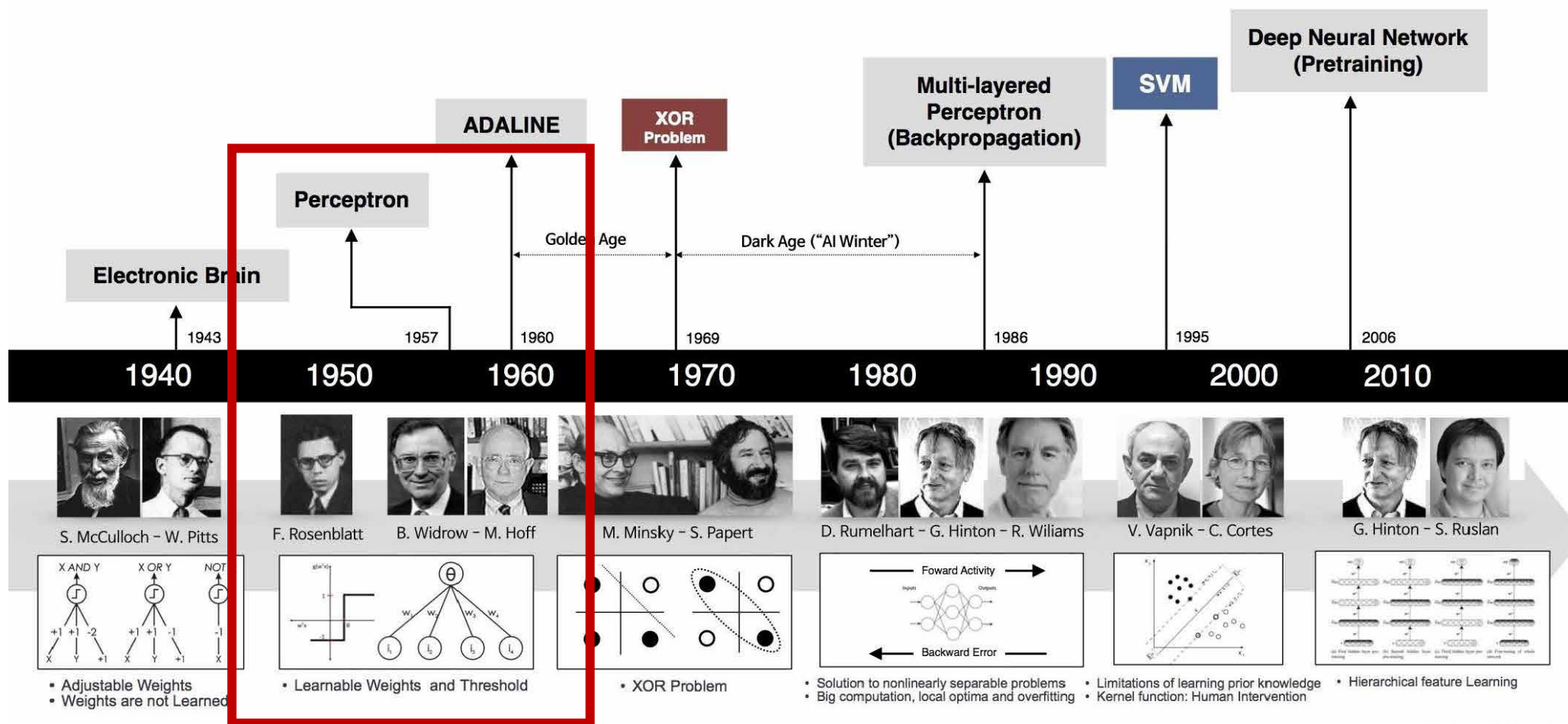
• 人工神经网络就是AI的历史



- 生物神经网络 (Biological Neural Network)
- 神经元 (Neuron) : 主体
- 轴突 (Axon) : 传递冲动
- 突触 (Synapse) : 神经元之间传递冲动
- 树突 (Dendrite) : 接受冲动

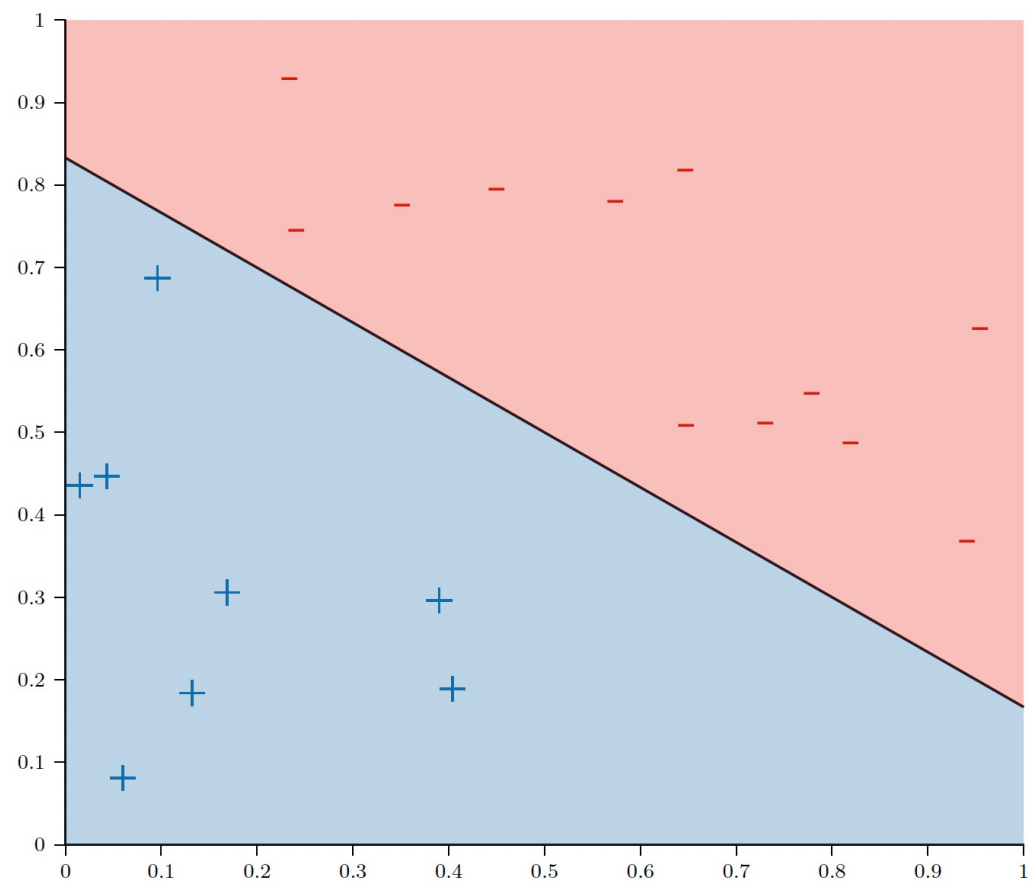


• 人工神经网络就是AI的历史

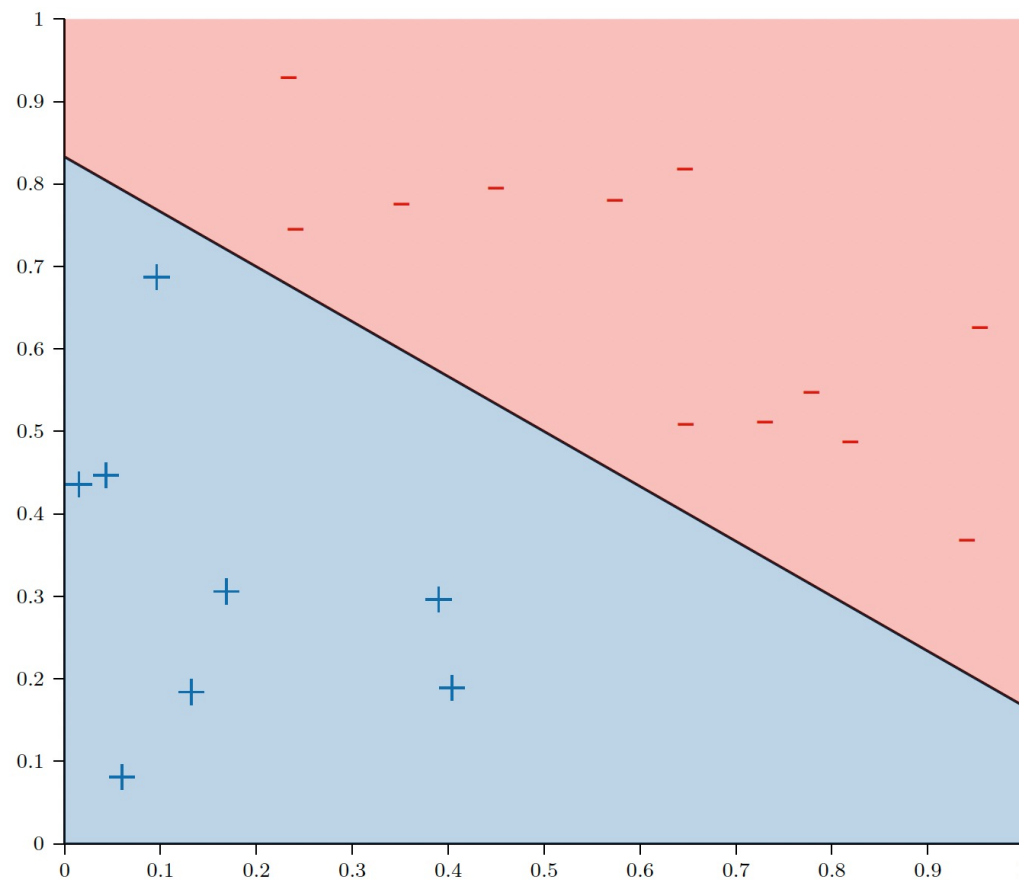
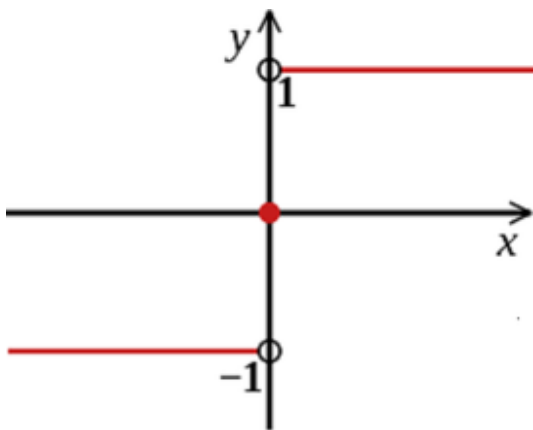


- 基础线性模型

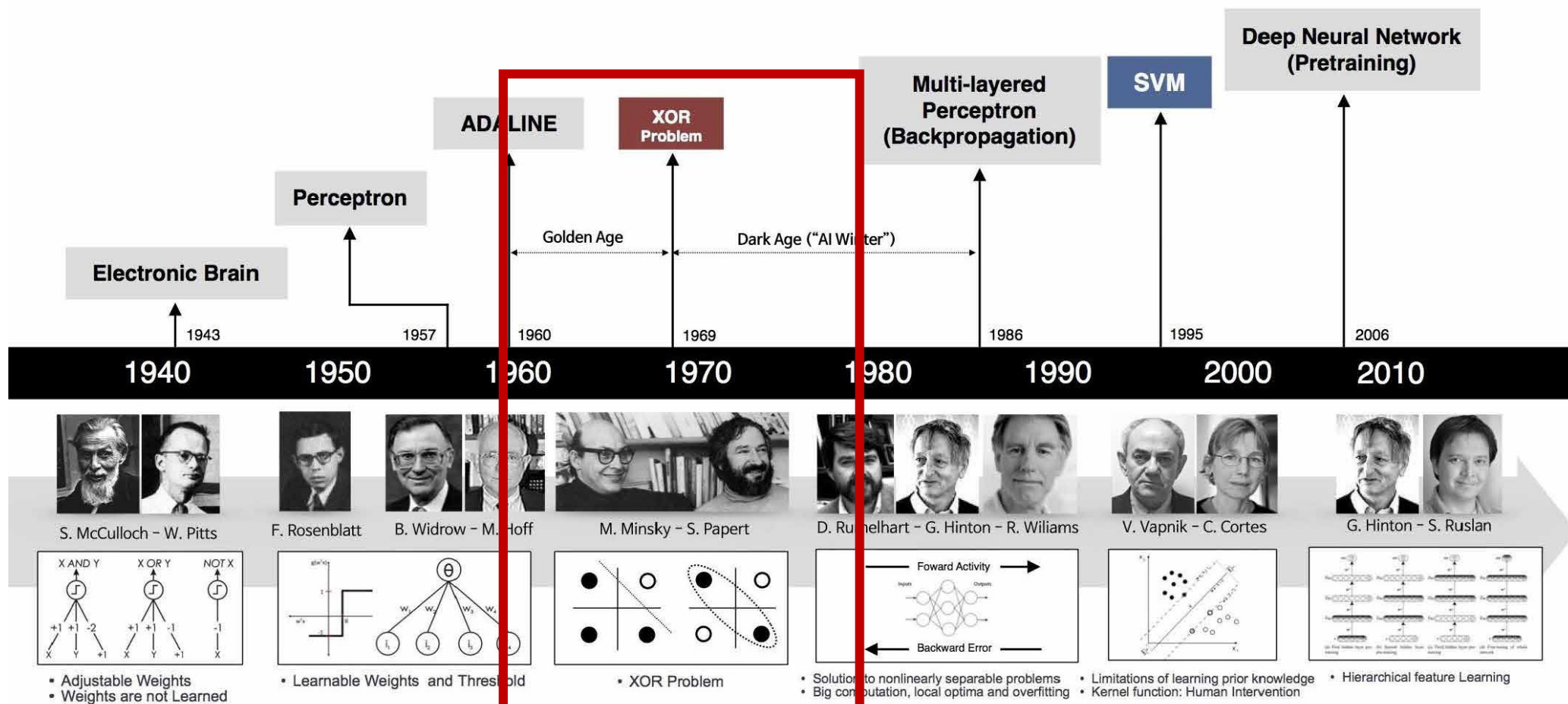
- $w^T x = b$



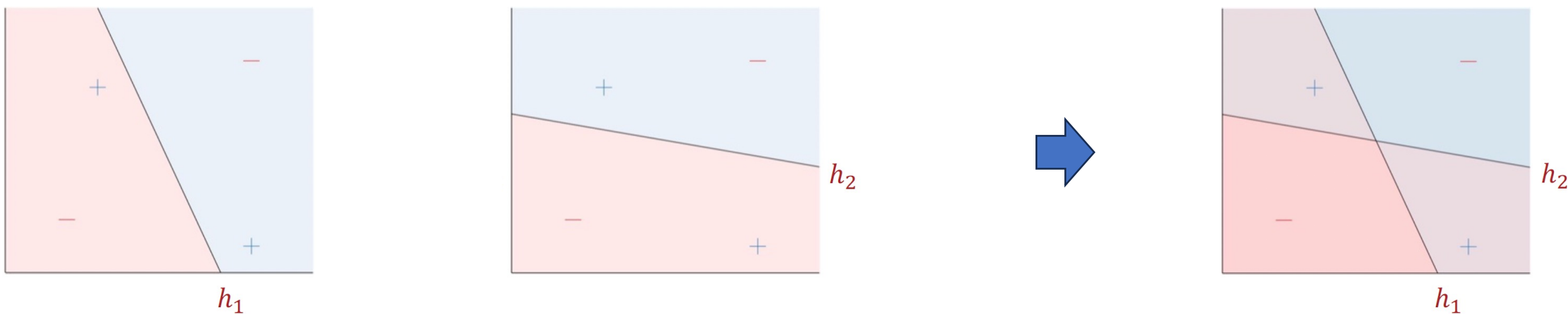
- 决策边界为 $h(x) = \text{sign}(w^T x + b)$
- 感知器 (Perception)
- $\text{sign}(x)$



• 人工神经网络就是AI的历史

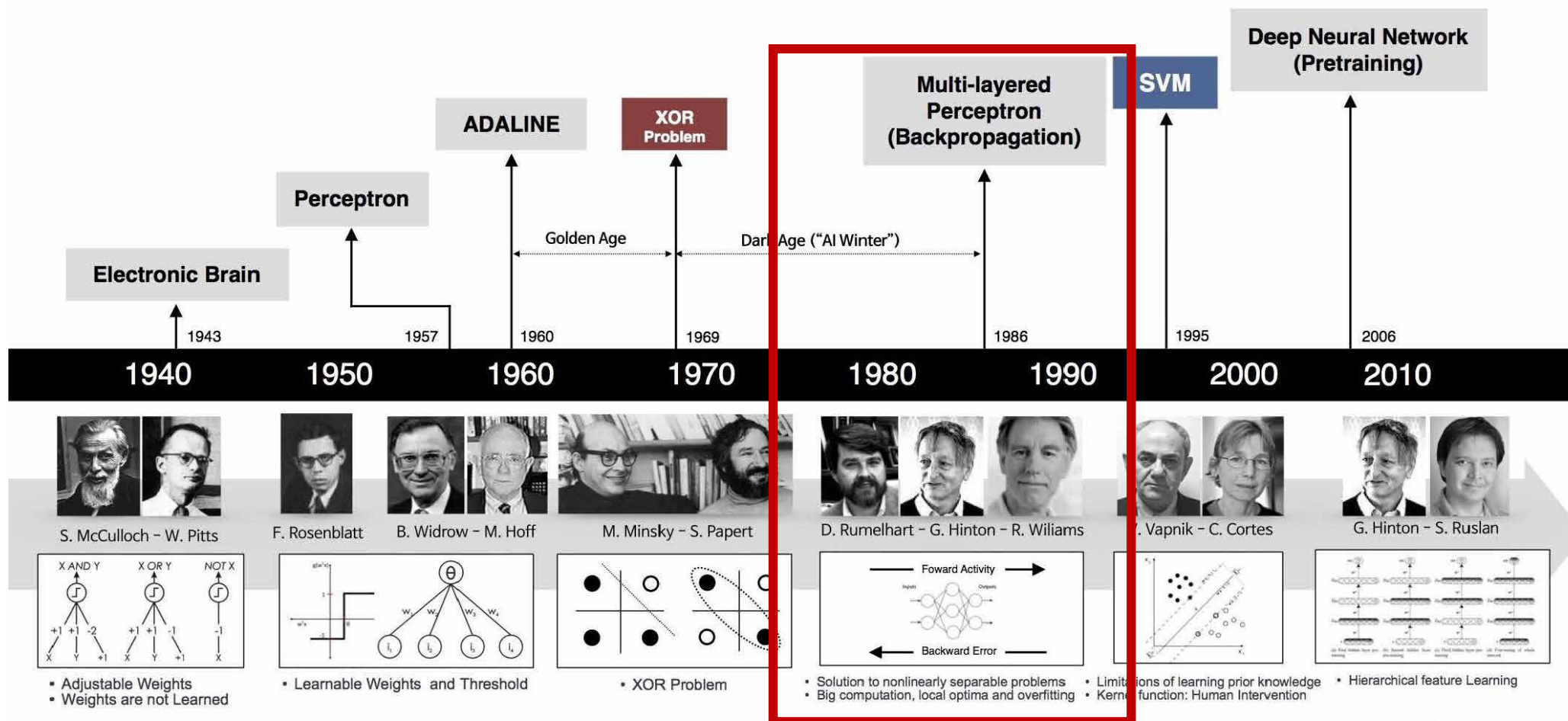


- 多个线性模型



$$h(x) = \begin{cases} +1 & \text{if } (h_1(x) = +1 \text{ and } h_2(x) = -1) \text{ or } (h_1(x) = -1 \text{ and } h_2(x) = +1) \\ -1 & \text{otherwise} \end{cases}$$

• 人工神经网络就是AI的历史

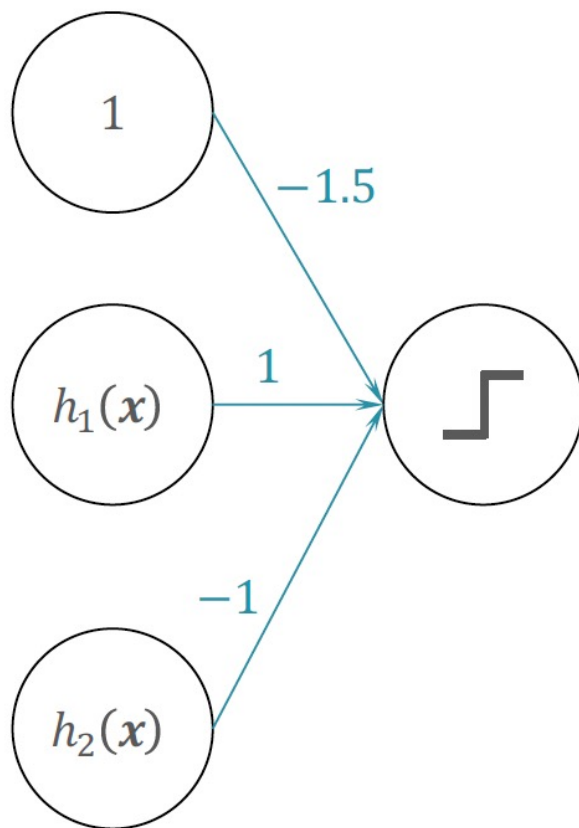


$$h(\mathbf{x}) = OR \left(AND(h_1(\mathbf{x}), \neg h_2(\mathbf{x})), AND(\neg h_1(\mathbf{x}), h_2(\mathbf{x})) \right)$$

多层感知器 Multi-Layer Perceptron (MLP)



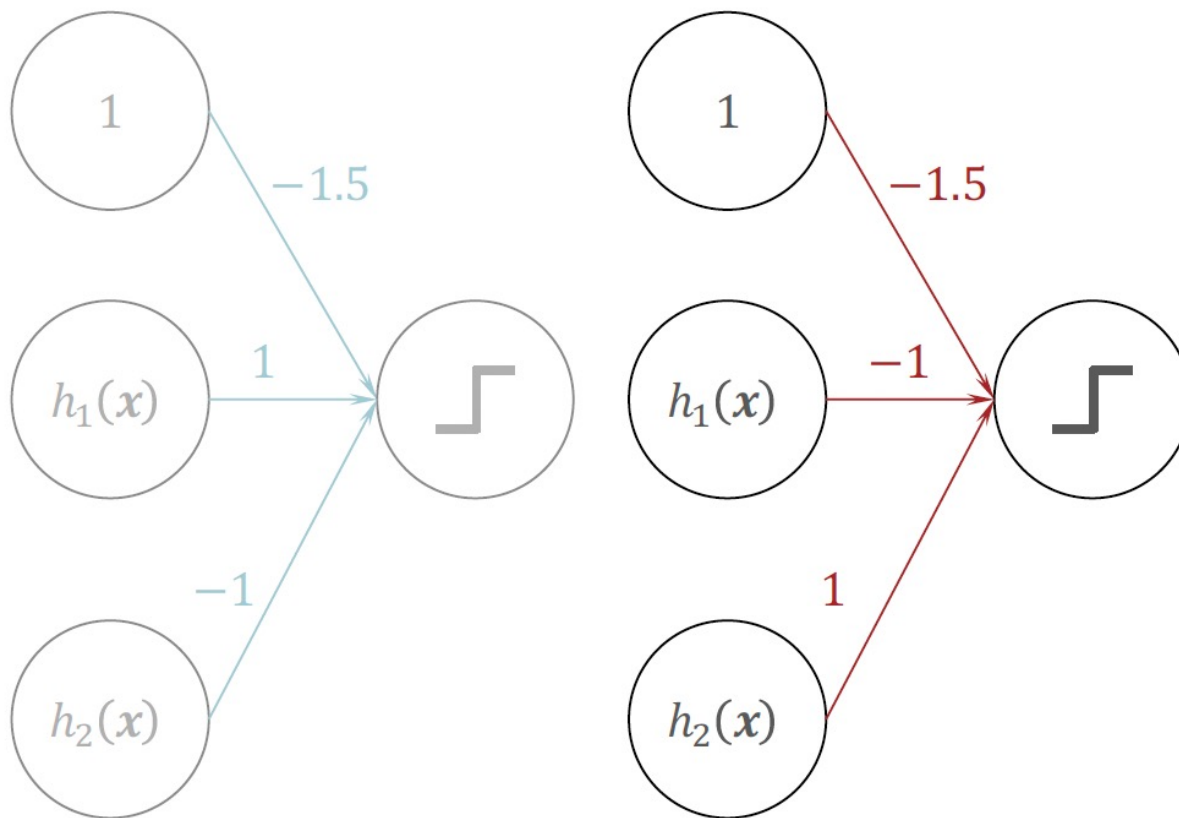
$$h(x) = OR \left(AND(h_1(x), \neg h_2(x)), AND(\neg h_1(x), h_2(x)) \right)$$



多层感知器 Multi-Layer Perceptron (MLP)



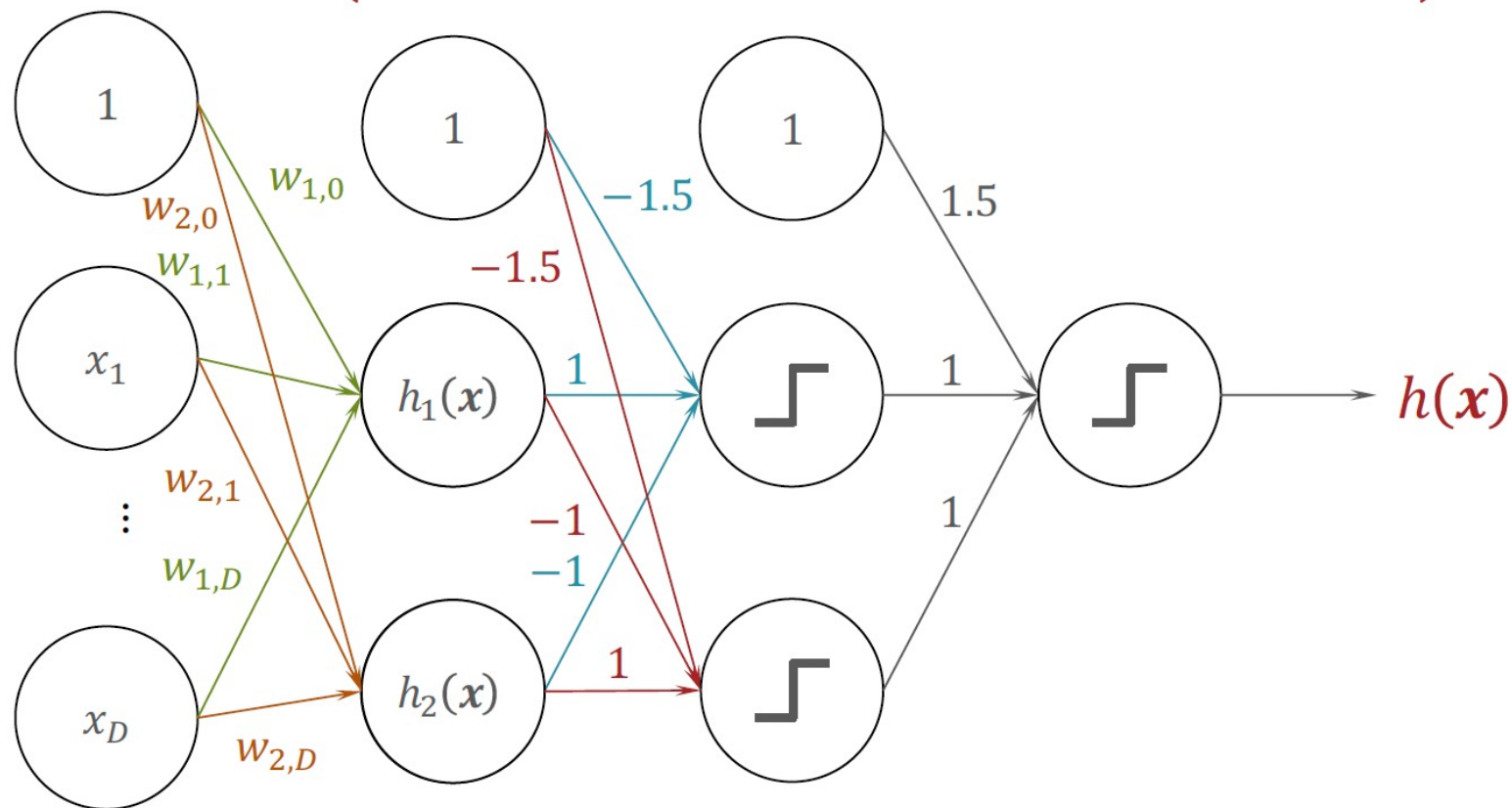
$$h(x) = OR\left(AND(h_1(x), \neg h_2(x)), AND(\neg h_1(x), h_2(x))\right)$$



多层感知器 Multi-Layer Perceptron (MLP)



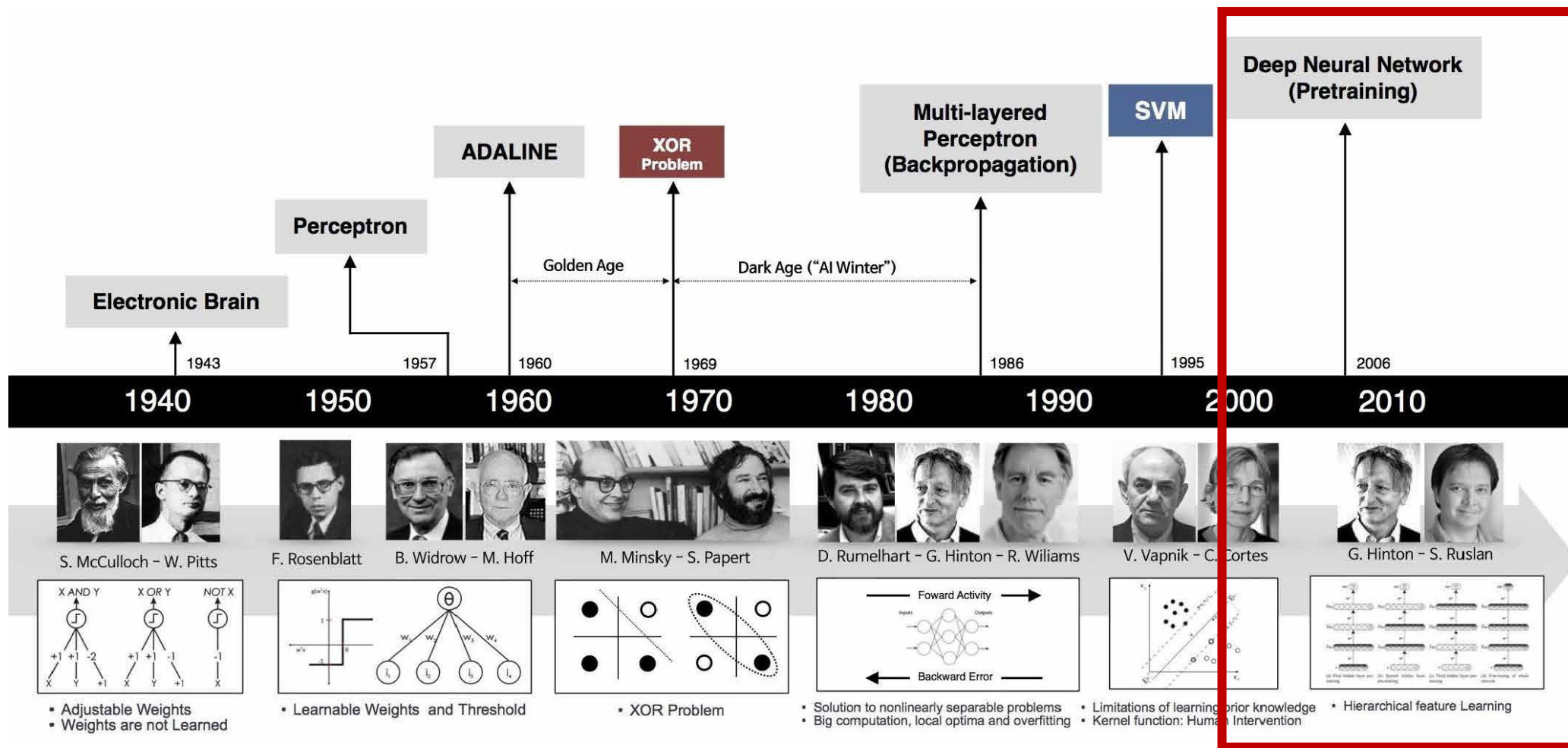
$$h(\mathbf{x}) = OR \left(AND(h_1(\mathbf{x}), \neg h_2(\mathbf{x})), AND(\neg h_1(\mathbf{x}), h_2(\mathbf{x})) \right)$$



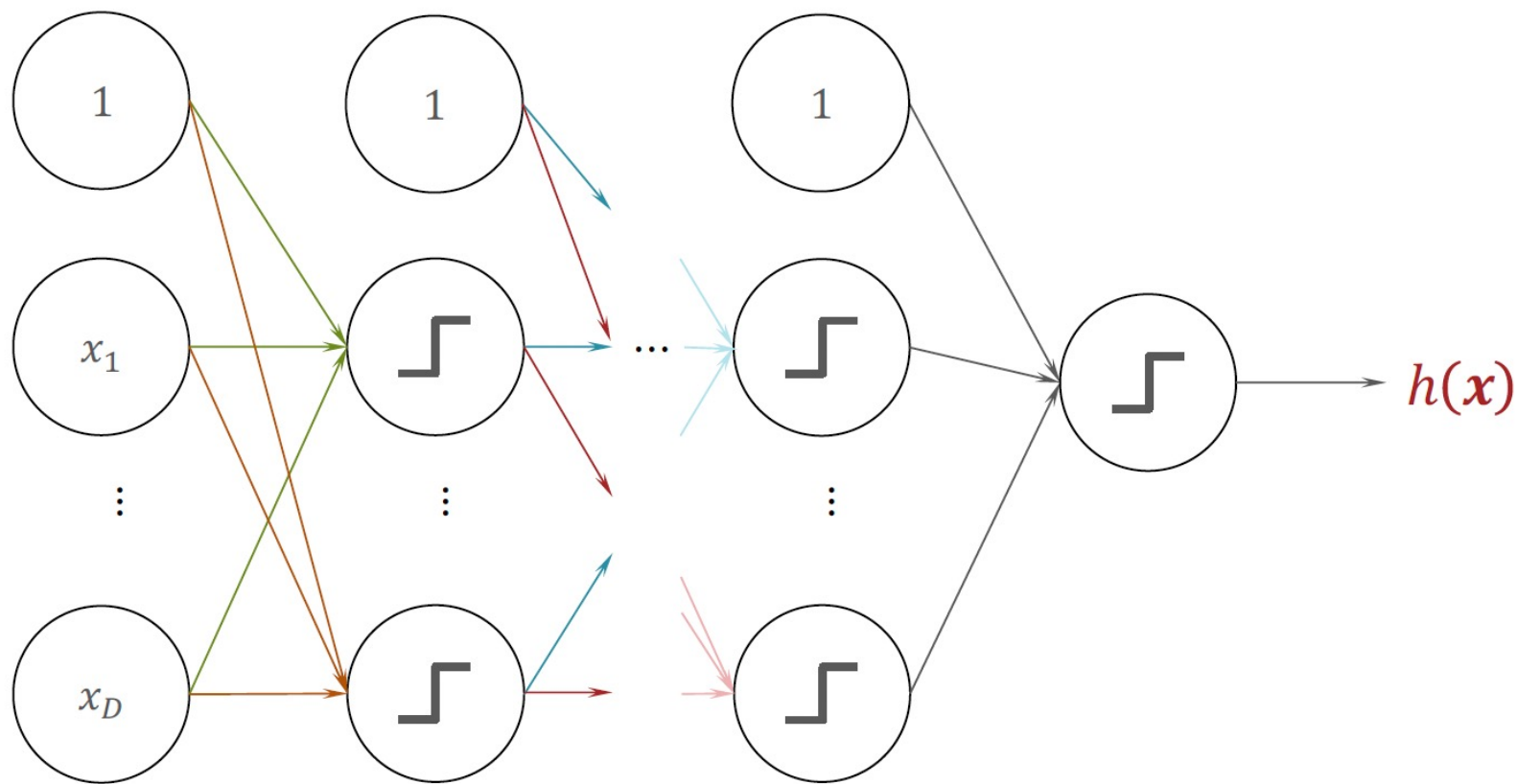
$$h(\mathbf{x}) = \text{sign}(\text{sign}(\text{sign}(\mathbf{w}_1^T \mathbf{x}) - \text{sign}(\mathbf{w}_2^T \mathbf{x}) - 1.5) + \text{sign}(-\text{sign}(\mathbf{w}_1^T \mathbf{x}) + \text{sign}(\mathbf{w}_2^T \mathbf{x}) - 1.5) + 1.5)$$



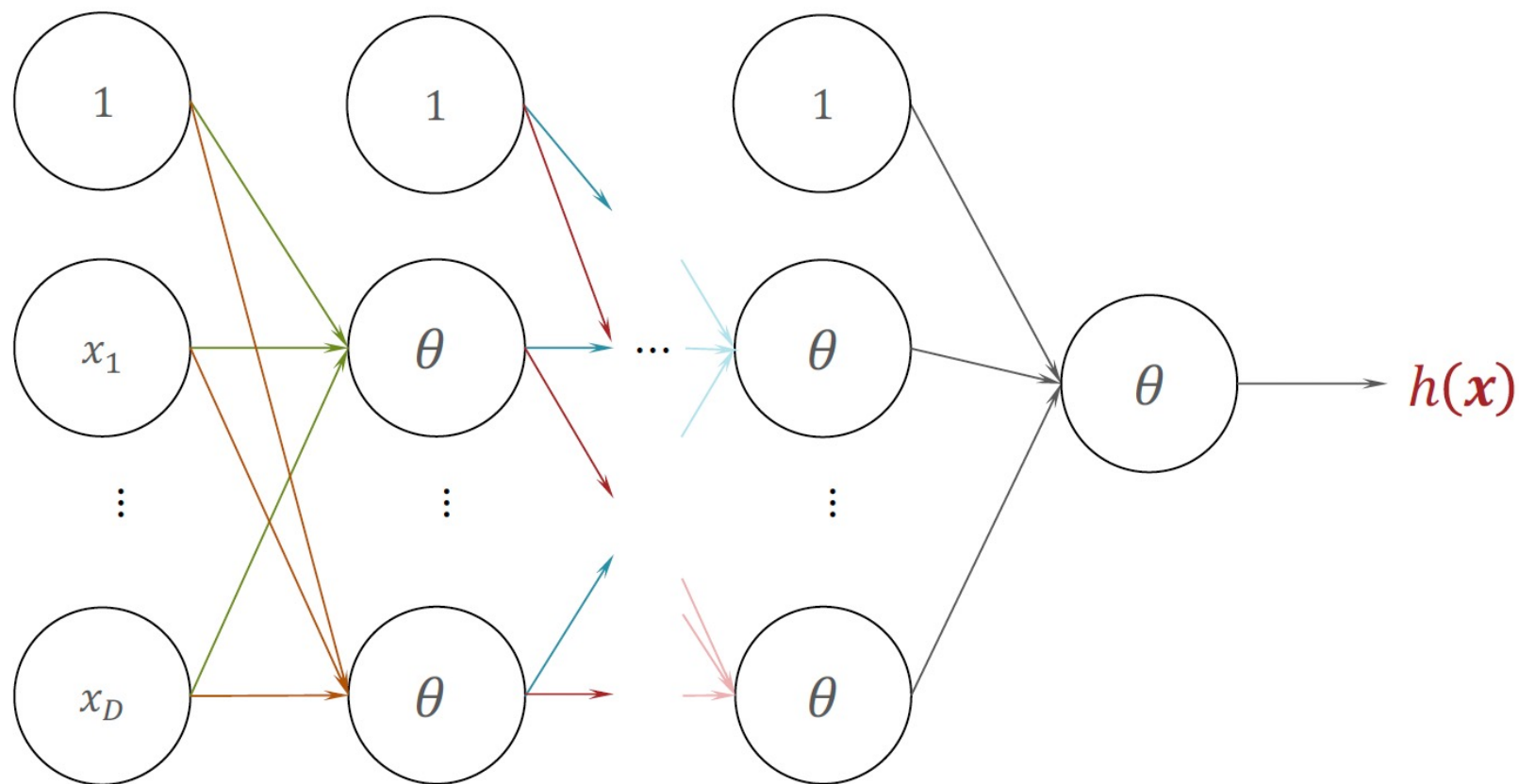
• 人工神经网络就是AI的历史



- 多层感知器



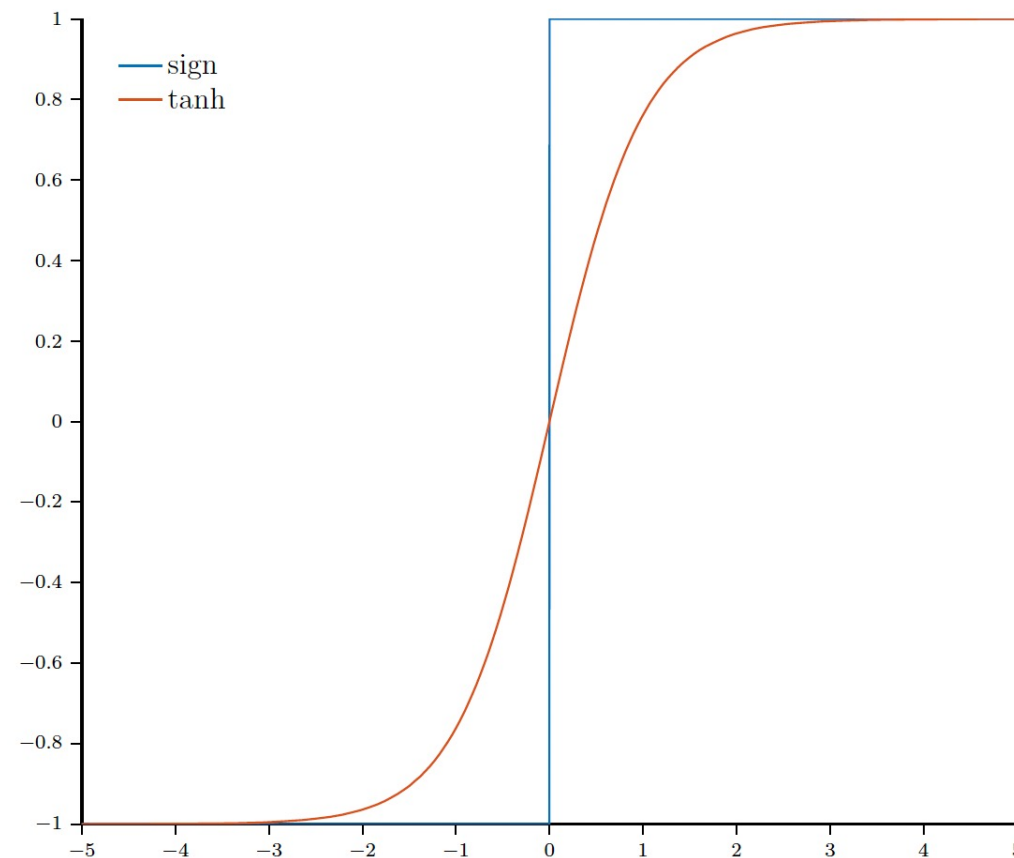
- 神经网络



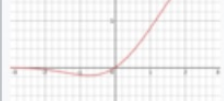
- $\theta(z)$: 激活函数
- 有明确定义的梯度

$$\tanh(z) = \frac{\sinh(z)}{\cosh(z)} = \frac{e^z - e^{-z}}{e^z + e^{-z}}$$

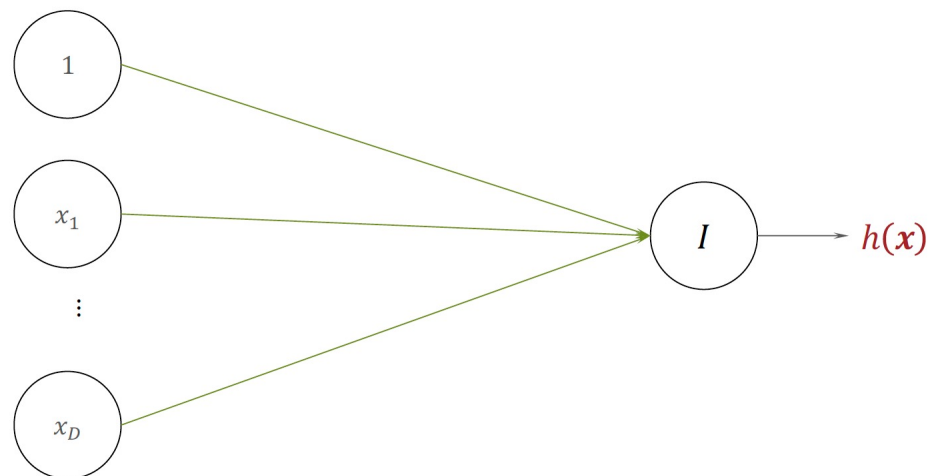
$$\bullet \frac{\partial \tanh(z)}{\partial z} = 1 - \tanh(z)^2$$



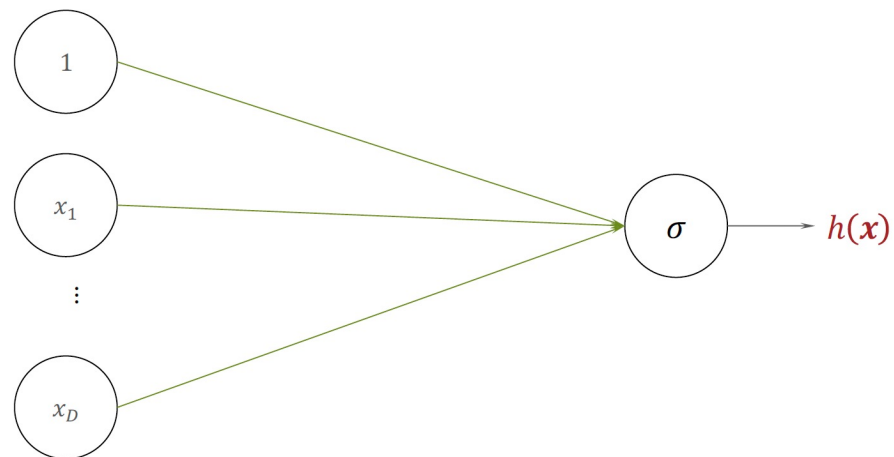
- $\theta(z)$: 其他激活函数

Logistic, sigmoid, or soft step		$\sigma(x) = \frac{1}{1 + e^{-x}}$
Hyperbolic tangent (tanh)		$\tanh(x) = \frac{e^x - e^{-x}}{e^x + e^{-x}}$
Rectified linear unit (ReLU) ^[7]		$\begin{cases} 0 & \text{if } x \leq 0 \\ x & \text{if } x > 0 \end{cases}$ $= \max\{0, x\} = x \mathbf{1}_{x>0}$
Gaussian Error Linear Unit (GELU) ^[4]		$\frac{1}{2}x \left(1 + \operatorname{erf}\left(\frac{x}{\sqrt{2}}\right) \right)$ $= x\Phi(x)$
Softplus ^[8]		$\ln(1 + e^x)$
Exponential linear unit (ELU) ^[9]		$\begin{cases} \alpha(e^x - 1) & \text{if } x \leq 0 \\ x & \text{if } x > 0 \end{cases}$ with parameter α
Leaky rectified linear unit (Leaky ReLU) ^[11]		$\begin{cases} 0.01x & \text{if } x < 0 \\ x & \text{if } x \geq 0 \end{cases}$
Parametric rectified linear unit (PReLU) ^[12]		$\begin{cases} \alpha x & \text{if } x < 0 \\ x & \text{if } x \geq 0 \end{cases}$ with parameter α

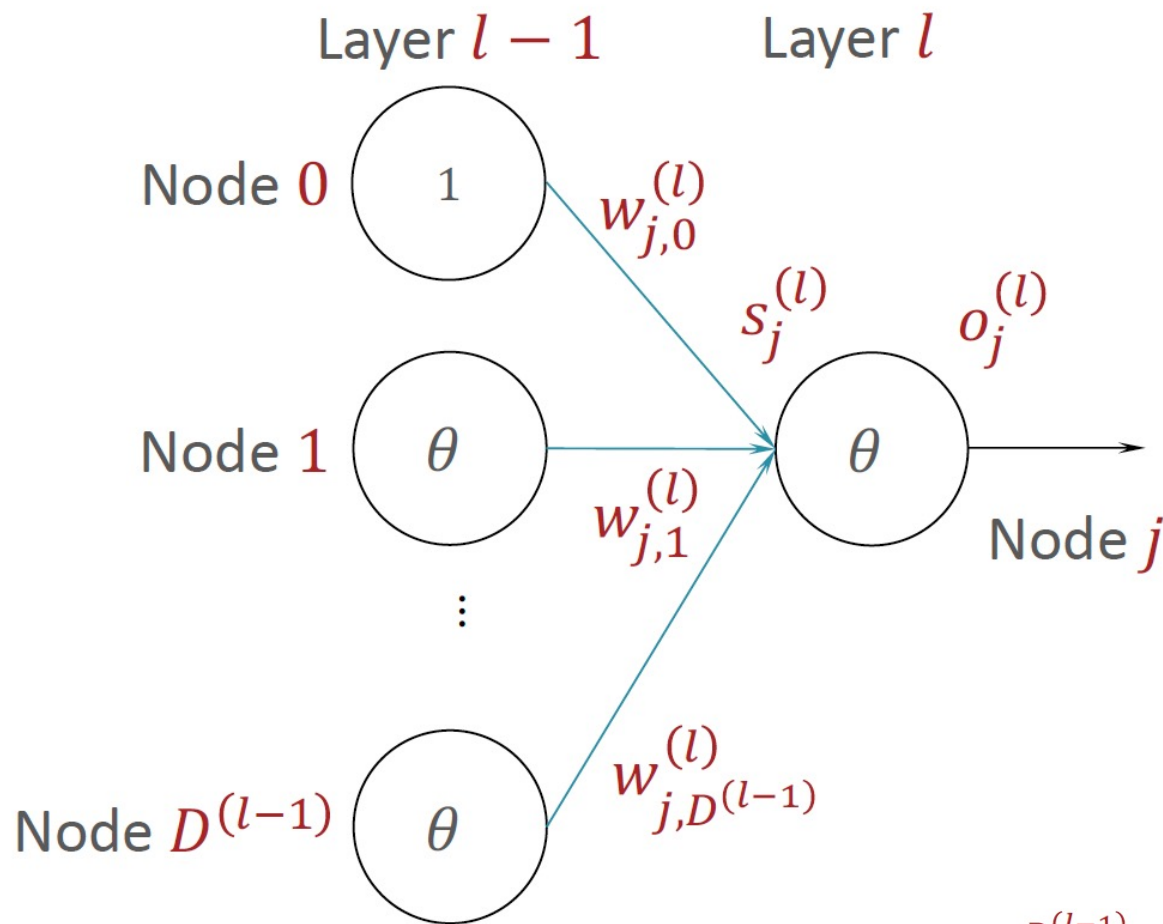
- 线性回归的神经网络实现（回归）



- 逻辑回归的神经网络实现（分类）

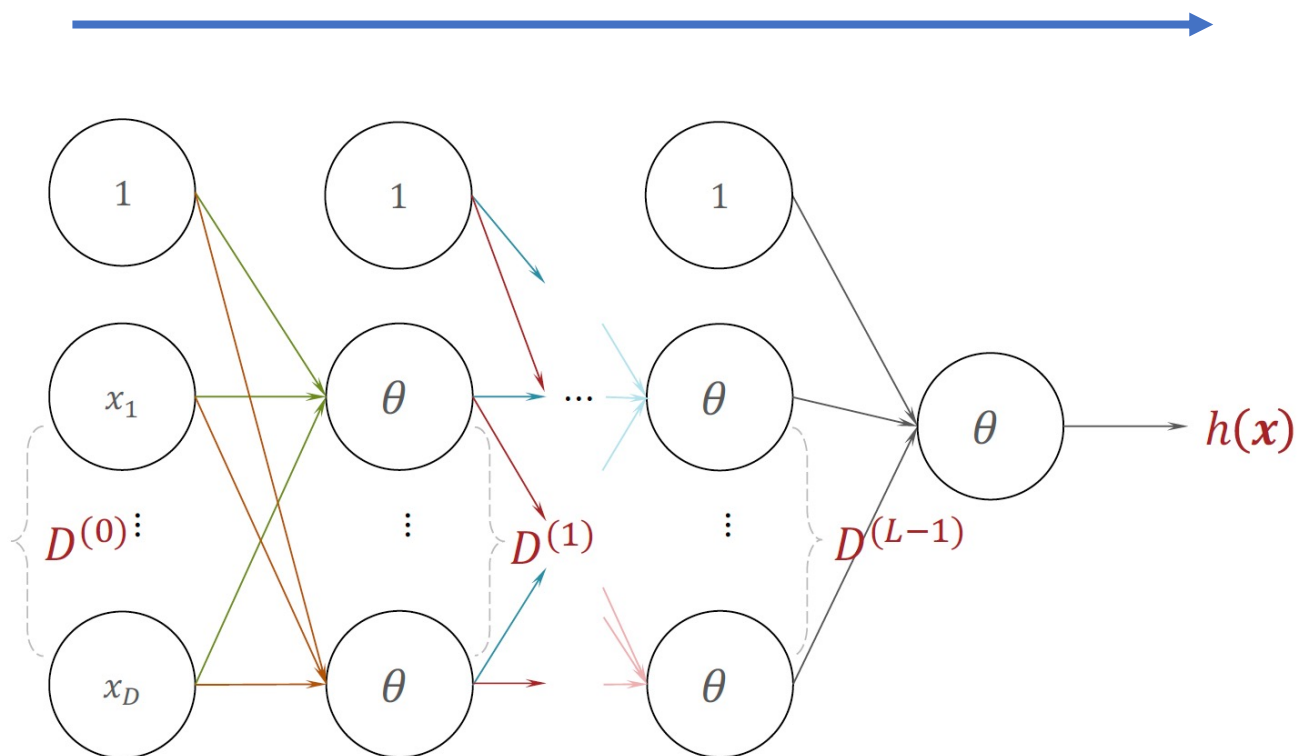


- 前向传播



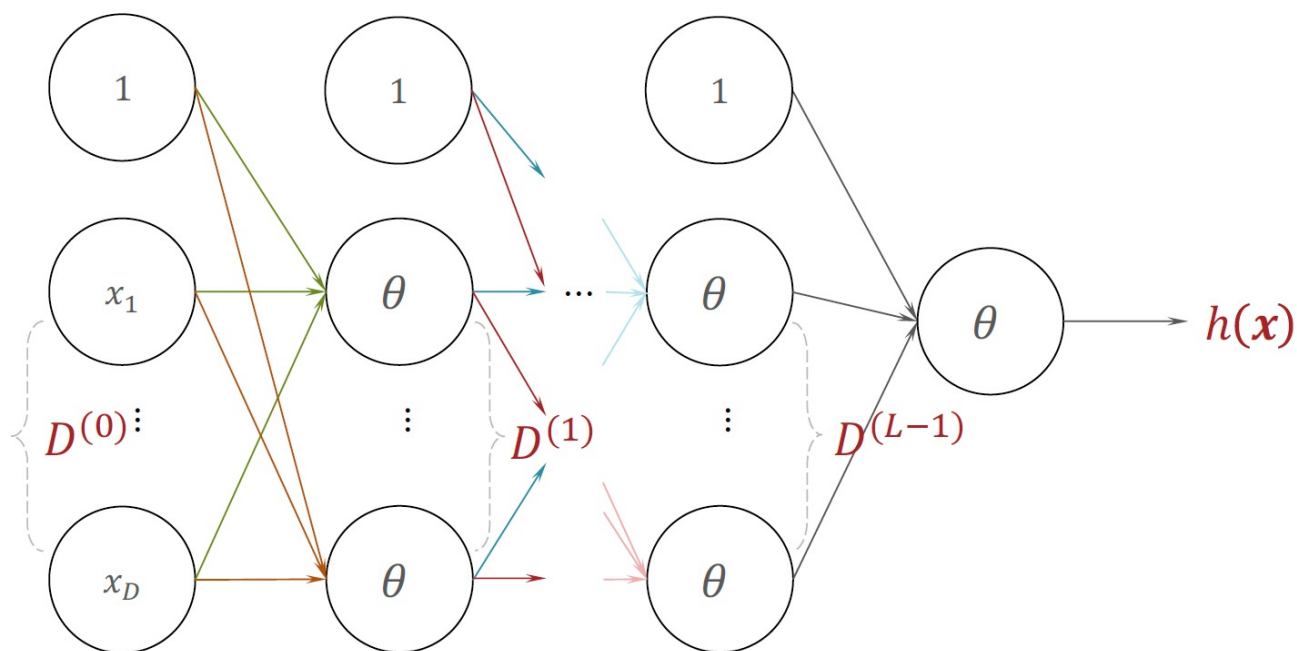
$$s_j^{(l)} = \sum_{i=0}^{D^{(l-1)}} w_{j,i}^{(l)} o_i^{(l-1)} \text{ and } o_j^{(l)} = \theta(s_j^{(l)})$$

• 前向传播



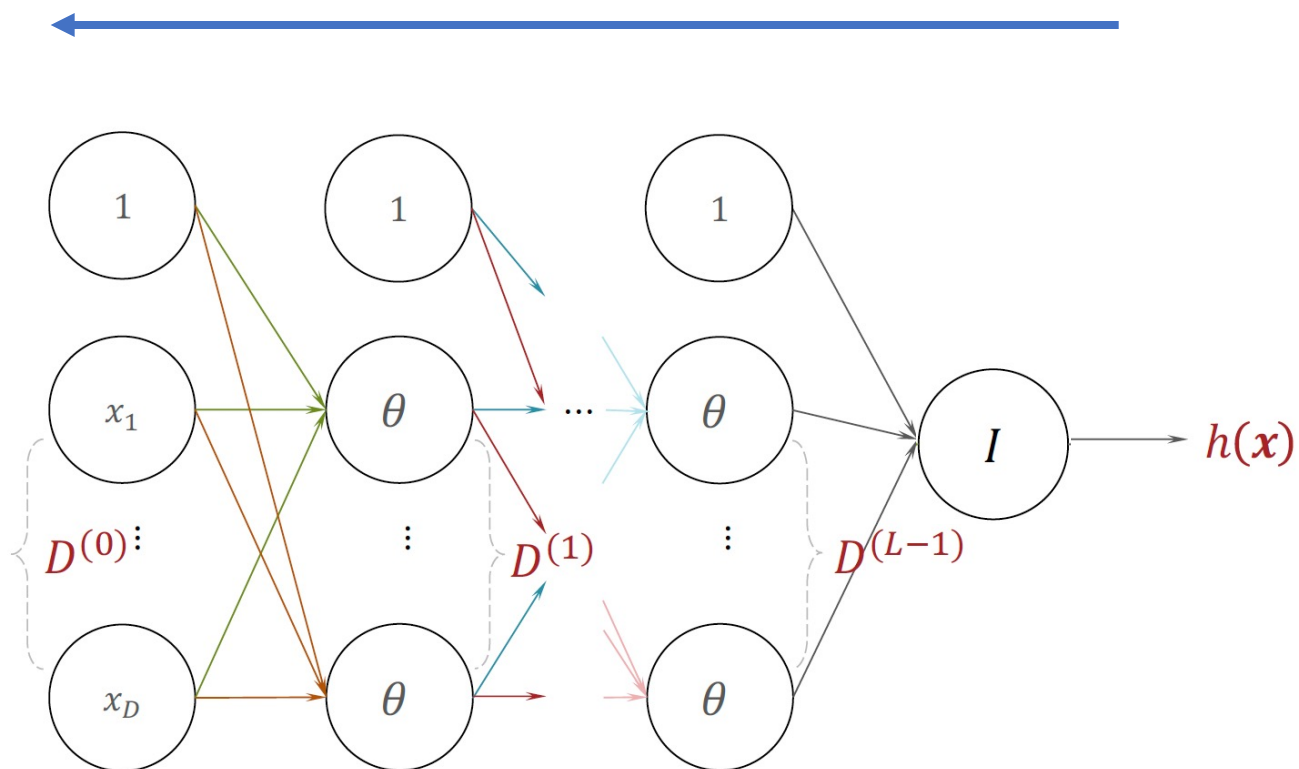
- Input: weights $W^{(1)}, \dots, W^{(L)}$ and a query data point $\mathbf{x}$
- Initialize $\mathbf{o}^{(0)} = \begin{bmatrix} 1 \\ \mathbf{x} \end{bmatrix}$
- For $l = 1, \dots, L$
 - $\mathbf{s}^{(l)} = W^{(l)} \mathbf{o}^{(l-1)}$
 - $\mathbf{o}^{(l)} = \begin{bmatrix} 1 \\ \theta(\mathbf{s}^{(l)}) \end{bmatrix}$
- Output: $h_{W^{(1)}, \dots, W^{(L)}}(\mathbf{x}) = \mathbf{o}^{(L)}$

• 模型优化：梯度下降



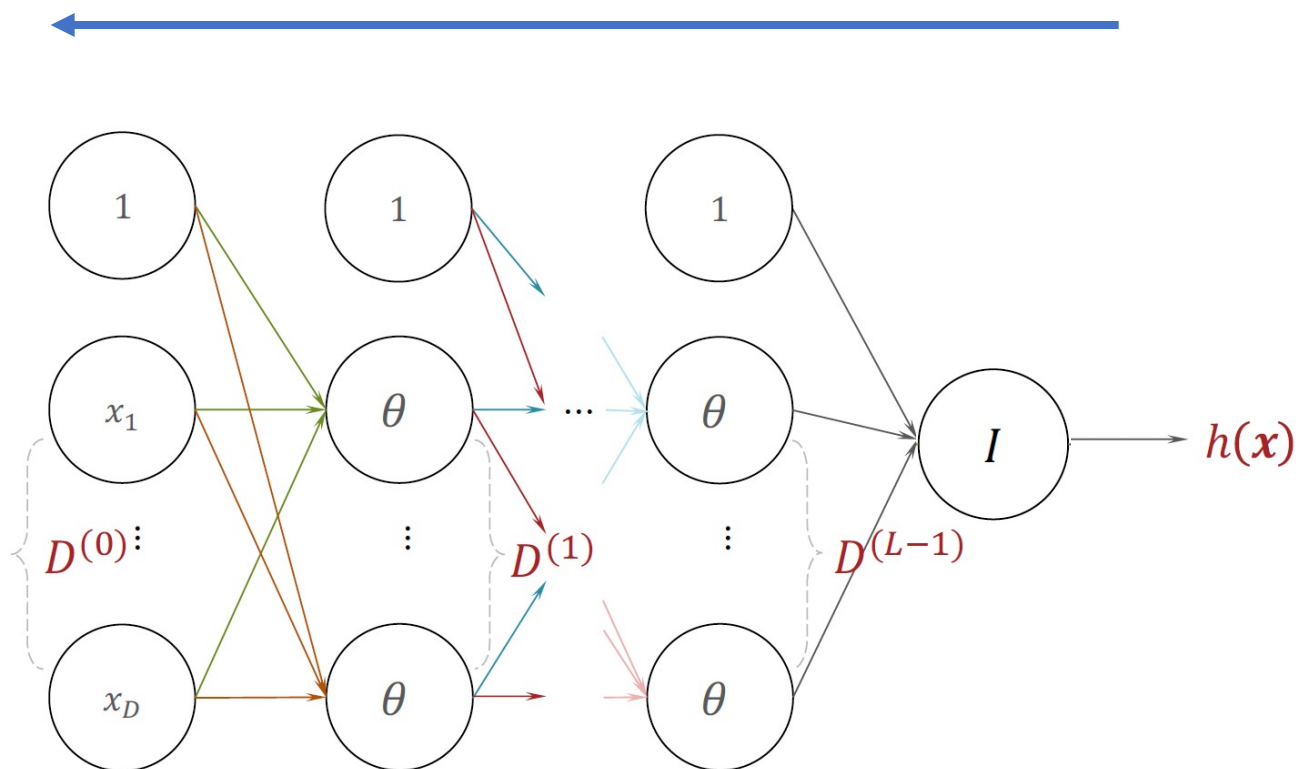
- Input: $\mathcal{D} = \{(\mathbf{x}^{(n)}, y^{(n)})\}_{n=1}^N, \eta^{(0)}$
- Initialize all weights $W_{(0)}^{(1)}, \dots, W_{(0)}^{(L)}$ to small, random numbers and set $t = 0$
- While TERMINATION CRITERION is not satisfied
 - For $i \in \text{shuffle}(\{1, \dots, N\})$
 - For $l = 1, \dots, L$
 - Compute $G^{(l)} = \nabla_{W_{(l)}} \ell^{(i)}(W_{(t)}^{(1)}, \dots, W_{(t)}^{(L)})$
 - Update $W^{(l)}$: $W_{(t+1)}^{(l)} = W_{(t)}^{(l)} - \eta^{(0)} G^{(l)}$
 - Increment t : $t = t + 1$
- Output: $W_{(t)}^{(1)}, \dots, W_{(t)}^{(L)}$

• 模型优化常规实现：反向传播



- Input: $W^{(1)}, \dots, W^{(L)}$ and $(\mathbf{x}^{(i)}, y^{(i)})$
- Run forward propagation with $\mathbf{x}^{(i)}$ to get $\mathbf{o}^{(1)}, \dots, \mathbf{o}^{(L)}$
- (Optional) Compute $\ell^{(i)} = (o^{(L)} - y^{(i)})^2$
- Initialize: $\delta^{(L)} = 2(o_1^{(L)} - y^{(i)})$

• 模型优化常规实现：反向传播

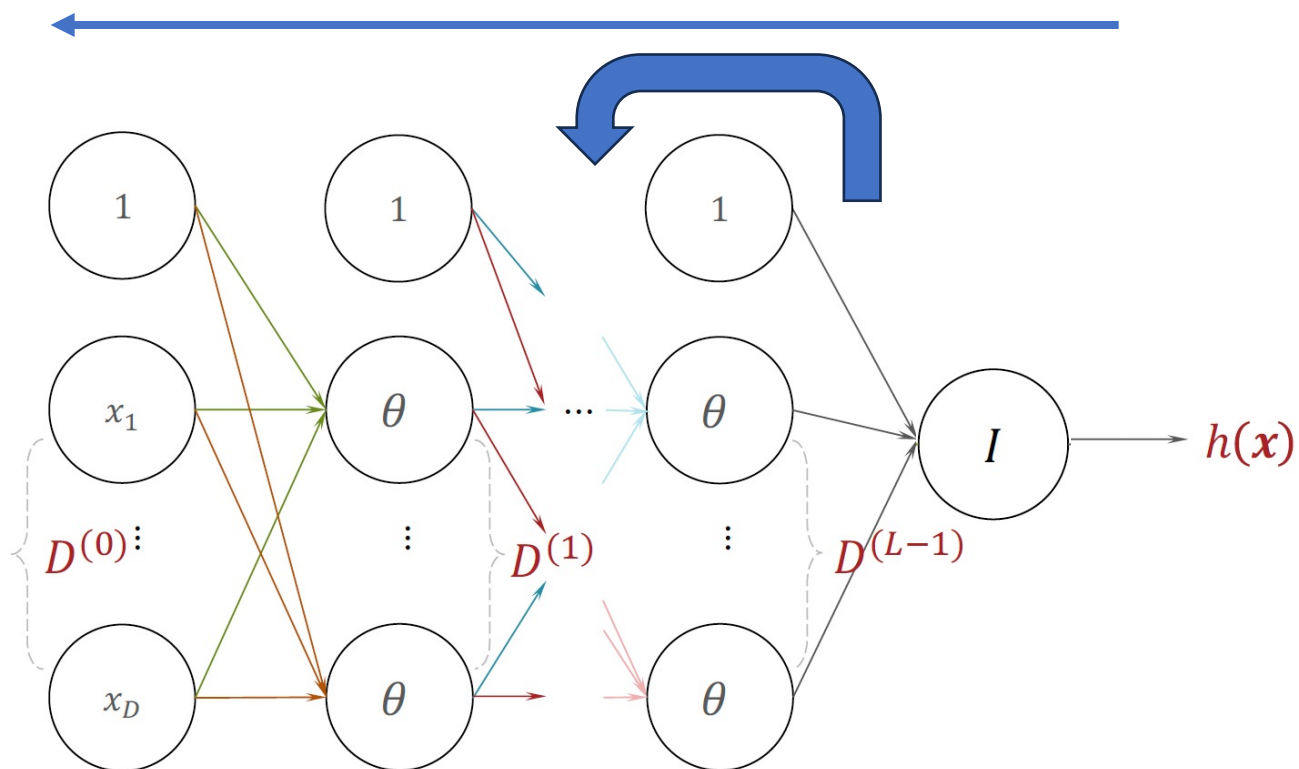


- Input: $W^{(1)}, \dots, W^{(L)}$ and $(\mathbf{x}^{(i)}, y^{(i)})$
- Run forward propagation with $\mathbf{x}^{(i)}$ to get $\mathbf{o}^{(1)}, \dots, \mathbf{o}^{(L)}$
- (Optional) Compute $\ell^{(i)} = (o^{(L)} - y^{(i)})^2$
- Initialize: $\delta^{(L)} = 2(o_1^{(L)} - y^{(i)})$

$$\delta^{(l+1)} = \frac{\partial \ell}{\partial \mathbf{z}^{(l+1)}}$$

$$\mathbf{z}^{(l+1)} = W^{(l+1)} \mathbf{o}^{(l)}$$

- 模型优化常规实现：反向传播



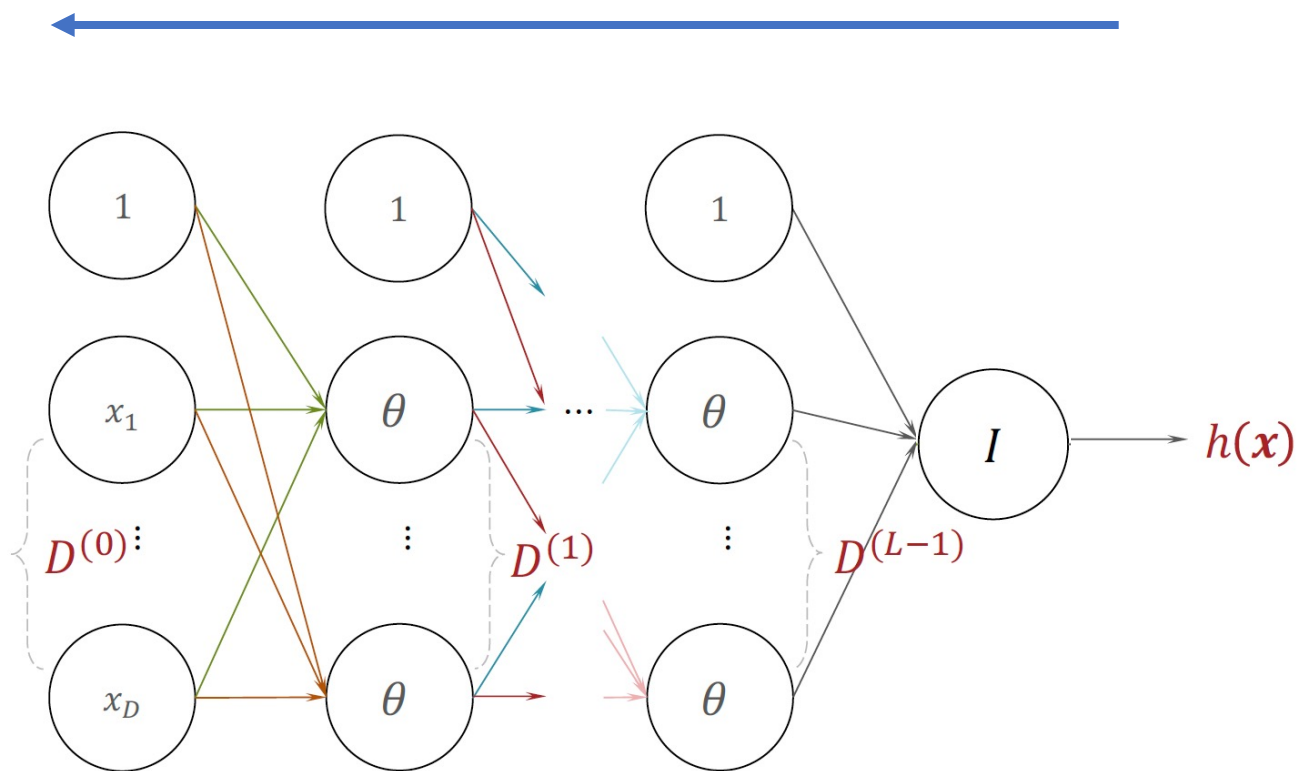
$$\delta^{(l)} = \frac{\partial \ell}{\partial z^{(l)}} = \left(\frac{\partial \ell}{\partial o^{(l)}} \right) \odot \left(\frac{\partial o^{(l)}}{\partial z^{(l)}} \right)$$

- Compute $\delta^{(l)} = \underline{W^{(l+1)T} \delta^{(l+1)}} \odot \underline{(1 - o^{(l)}) \odot o^{(l)}}$

$$\frac{\partial \ell}{\partial o^{(l)}} = W^{(l+1)T} \delta^{(l+1)}$$

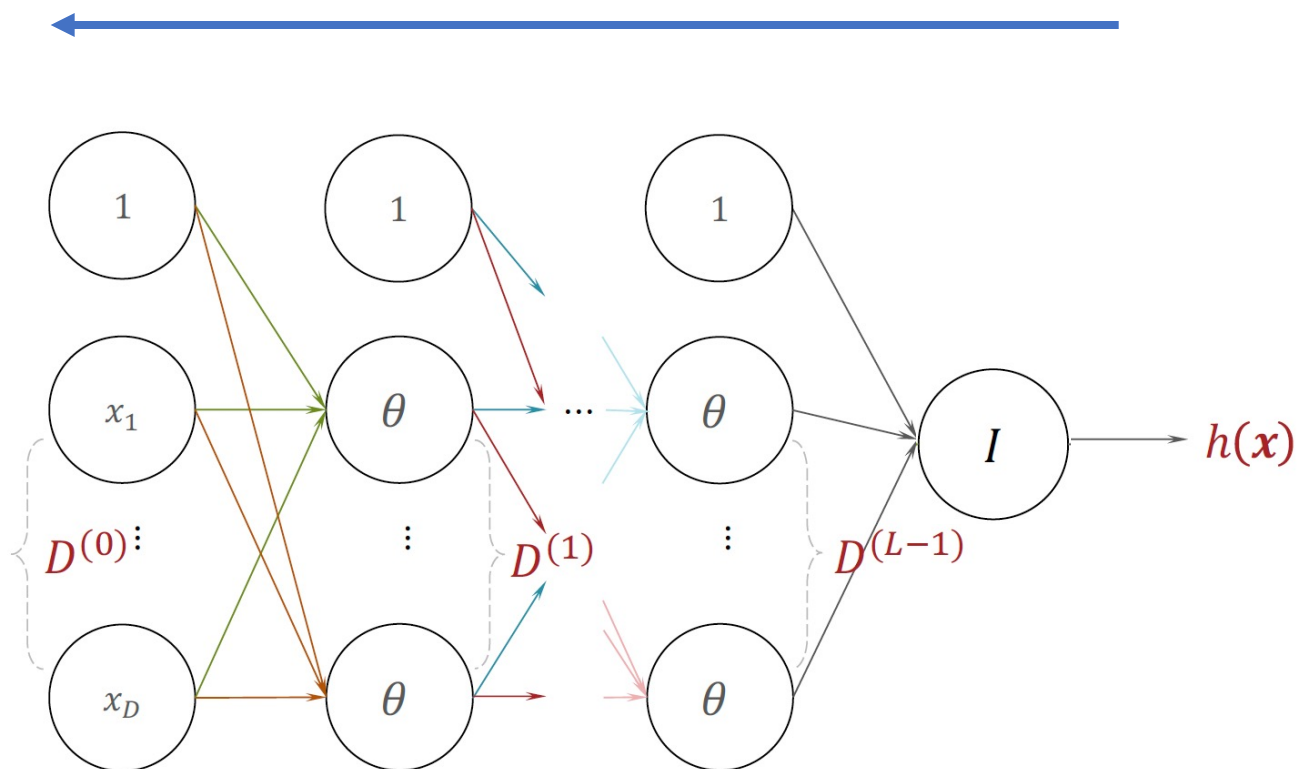
$$\cdot \frac{\partial \tanh(z)}{\partial z} = 1 - \tanh(z)^2$$

• 模型优化常规实现：反向传播



- Input: $W^{(1)}, \dots, W^{(L)}$ and $(\mathbf{x}^{(i)}, y^{(i)})$
- Run forward propagation with $\mathbf{x}^{(i)}$ to get $\mathbf{o}^{(1)}, \dots, \mathbf{o}^{(L)}$
- (Optional) Compute $\ell^{(i)} = (o^{(L)} - y^{(i)})^2$
- Initialize: $\delta^{(L)} = 2(o_1^{(L)} - y^{(i)})$
- For $l = L - 1, \dots, 1$
 - Compute $\delta^{(l)} = W^{(l+1)T} \delta^{(l+1)} \odot (1 - \mathbf{o}^{(l)} \odot \mathbf{o}^{(l)})$
 - Compute $G^{(l)} = \delta^{(l)} \mathbf{o}^{(l-1)T}$

• 模型优化常规实现：反向传播

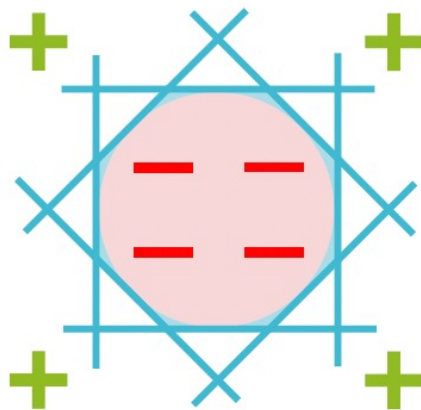


`torch.Tensor.backward`

`Tensor.backward(gradient=None, retain_graph=None, create_graph=False, inputs=None)`

- Input: $W^{(1)}, \dots, W^{(L)}$ and $(x^{(i)}, y^{(i)})$
- Run forward propagation with $x^{(i)}$ to get $o^{(1)}, \dots, o^{(L)}$
- (Optional) Compute $\ell^{(i)} = (o^{(L)} - y^{(i)})^2$
- Initialize: $\delta^{(L)} = 2(o_1^{(L)} - y^{(i)})$
- For $l = L - 1, \dots, 1$
 - Compute $\delta^{(l)} = W^{(l+1)T} \delta^{(l+1)} \odot (1 - o^{(l)} \odot o^{(l)})$
 - Compute $G^{(l)} = \delta^{(l)} o^{(l-1)T}$
- Output: $G^{(1)}, \dots, G^{(L)}$, the gradients of $\ell^{(i)}$ w.r.t $W^{(1)}, \dots, W^{(L)}$

- 理论上, 无限宽的3层MLP可以拟合任何函数
 - Theorem: any function that can be decomposed into perceptrons can be modelled exactly using a 3-layer MLP
 - Any smooth decision boundary can be approximated to an arbitrary precision using a finite number of perceptrons



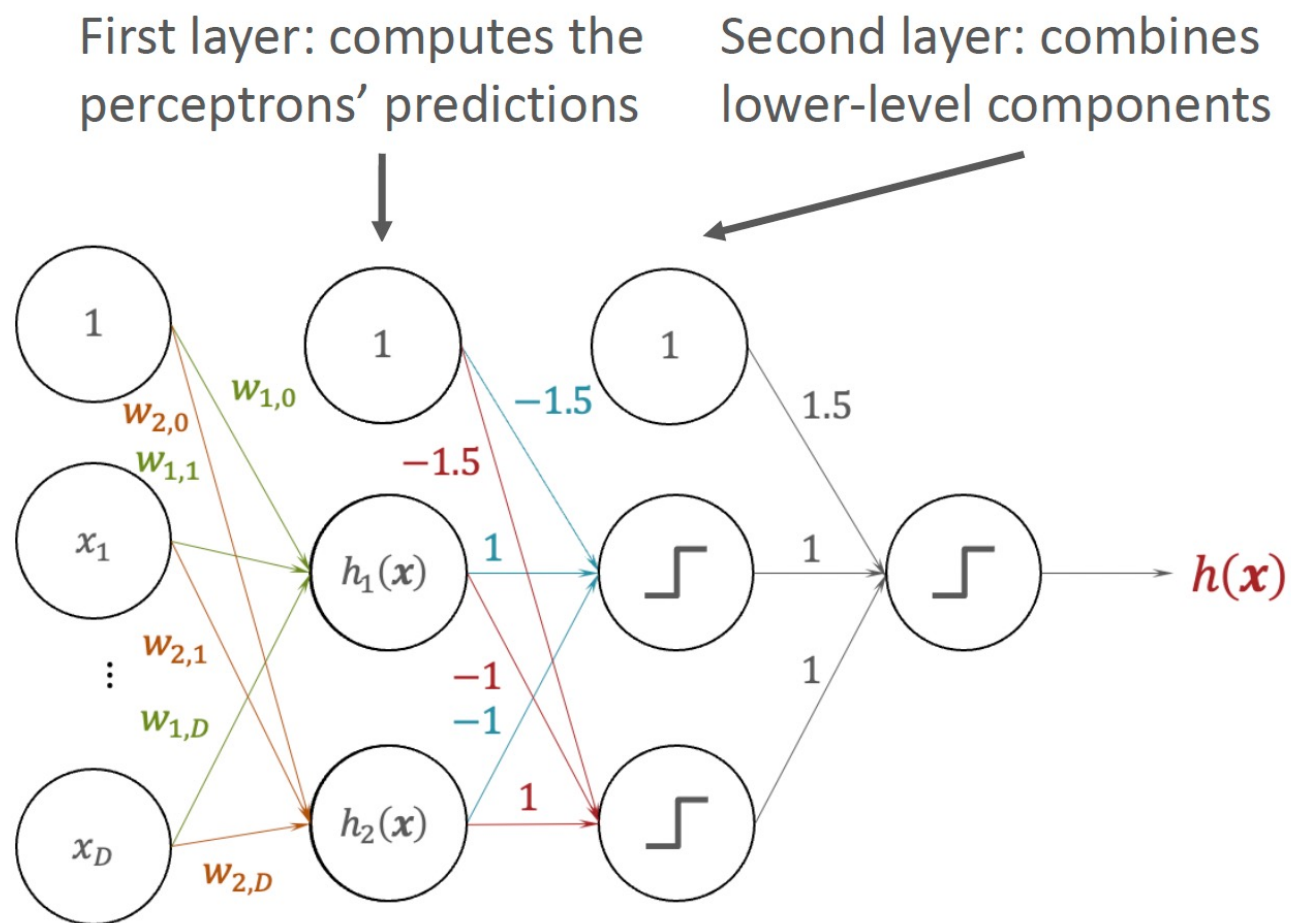
- 深度学习的定义：more than one layer

Definition [\[edit \]](#)

Deep learning is a class of [machine learning algorithms](#) that^[11](pp199–200) uses multiple layers to progressively extract higher level features from the raw input. For example, in image processing, lower layers may identify edges, while higher layers may identify the concepts relevant to a human such as digits or letters or faces.

- 本质原因：无限宽不现实，深度可以替代宽度（因为ResNet, BN etc...）

- 深度学习:



- 卷积神经网络（CNN）：加入卷积操作的神经网络
- 老婆饼没有老婆，卷积神经网络没有卷积
- 本质是卷积核的互相关

0	0	0	0	0	0
0	1	2	2	1	0
0	2	4	4	2	0
0	1	3	3	1	0
0	1	2	3	1	0
0	0	1	1	0	0

 $*$

0	1	0
1	-4	1
0	1	0

 $=$

0			

$$(0 * 0) + (0 * 1) + (0 * 0) + (0 * 1) + (1 * -4) + (2 * 1) + (0 * 0) + (2 * 1) + (4 * 0) = 0$$

- 卷积神经网络 (CNN)：加入卷积操作的神经网络
- 老婆饼没有老婆，卷积神经网络没有卷积
- 本质是卷积核的互相关

0	0	0	0	0	0
0	1	2	2	1	0
0	2	4	4	2	0
0	1	3	3	1	0
0	1	2	3	1	0
0	0	1	1	0	0

 $*$

0	1	0
1	-4	1
0	1	0

 $=$

0	-1		

$$(0 * 0) + (0 * 1) + (0 * 0) + (1 * 1) + (2 * -4) \\ + (2 * 1) + (2 * 0) + (4 * 1) + (4 * 0) = -1$$

- 卷积神经网络 (CNN)：加入卷积操作的神经网络
- 老婆饼没有老婆，卷积神经网络没有卷积
- 本质是卷积核的互相关

0	0	0	0	0	0
0	1	2	2	1	0
0	2	4	4	2	0
0	1	3	3	1	0
0	1	2	3	1	0
0	0	1	1	0	0

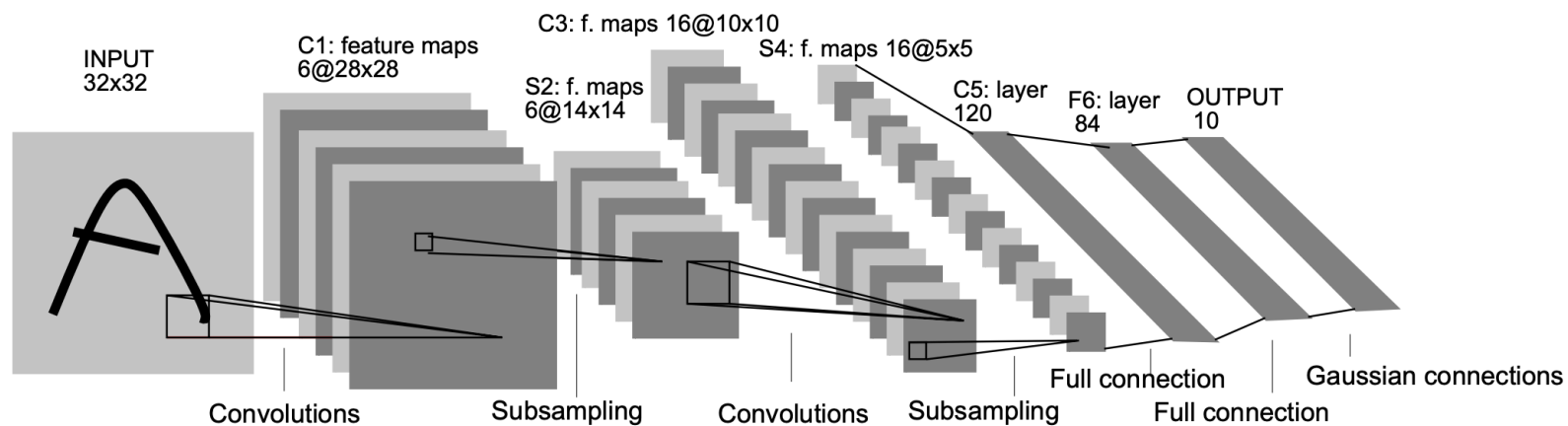
 $*$

0	1	0
1	-4	1
0	1	0

 $=$

0	-1	-1	0
-2	-5	-5	-2
2	-2	-1	3
-1	0	-5	0

- 现有的网络结构30年前就已经完成了



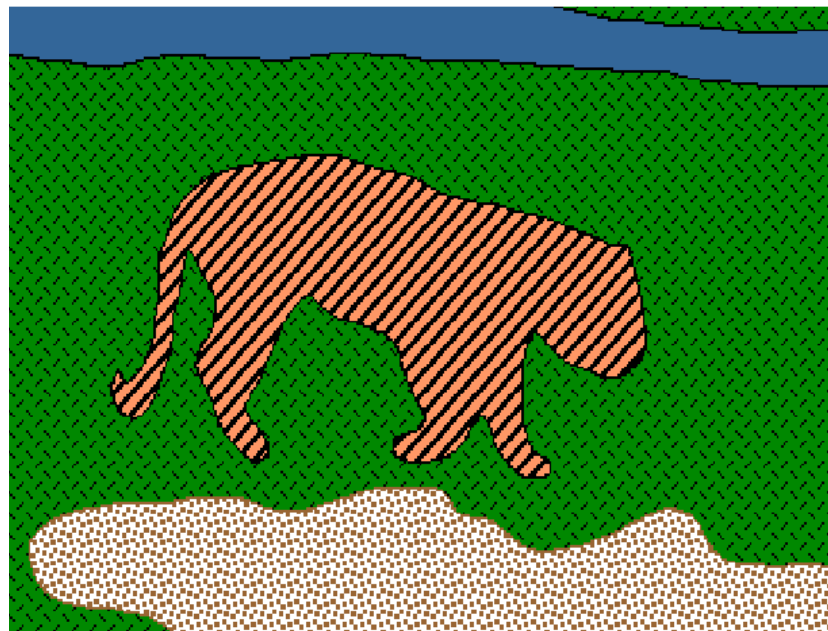
目录

1 高斯过程

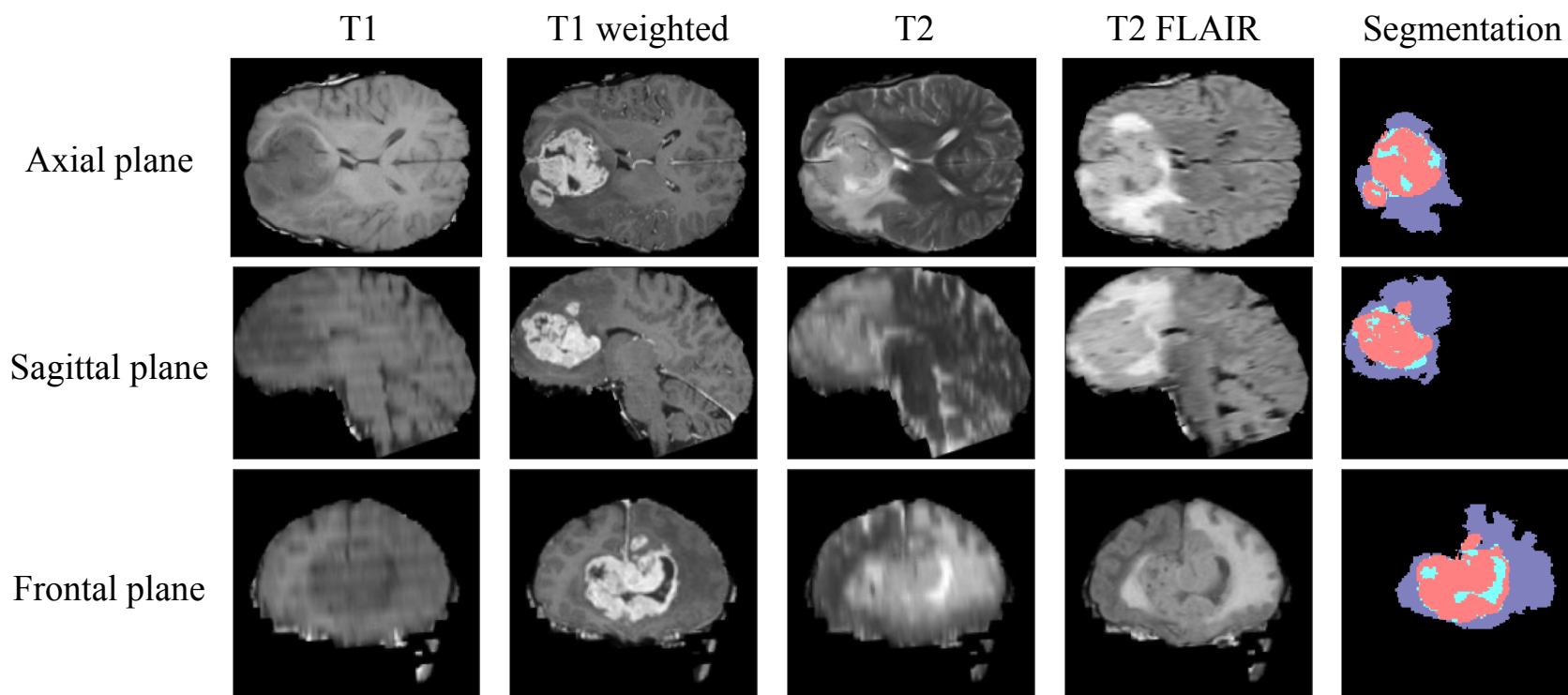
2 CNN

3 医学图像分割

- 将图像划分成多个结构意义的区域

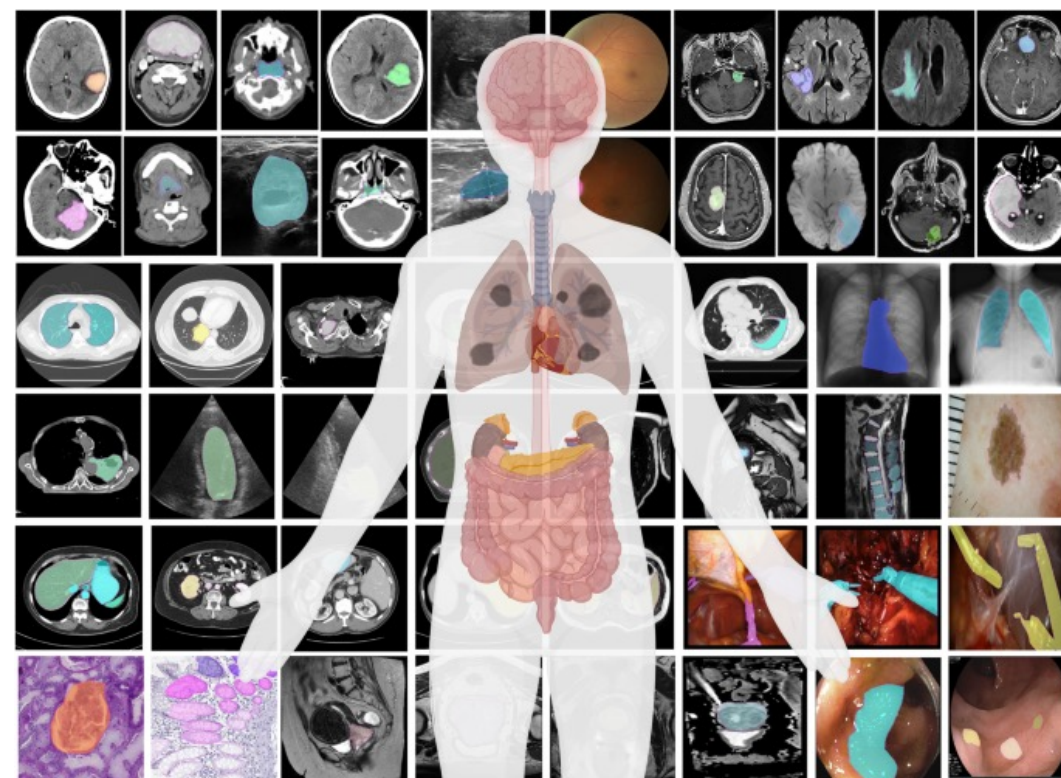


- 将医学图像（磁共振，CT，超声等）划分成多个结构意义的区域

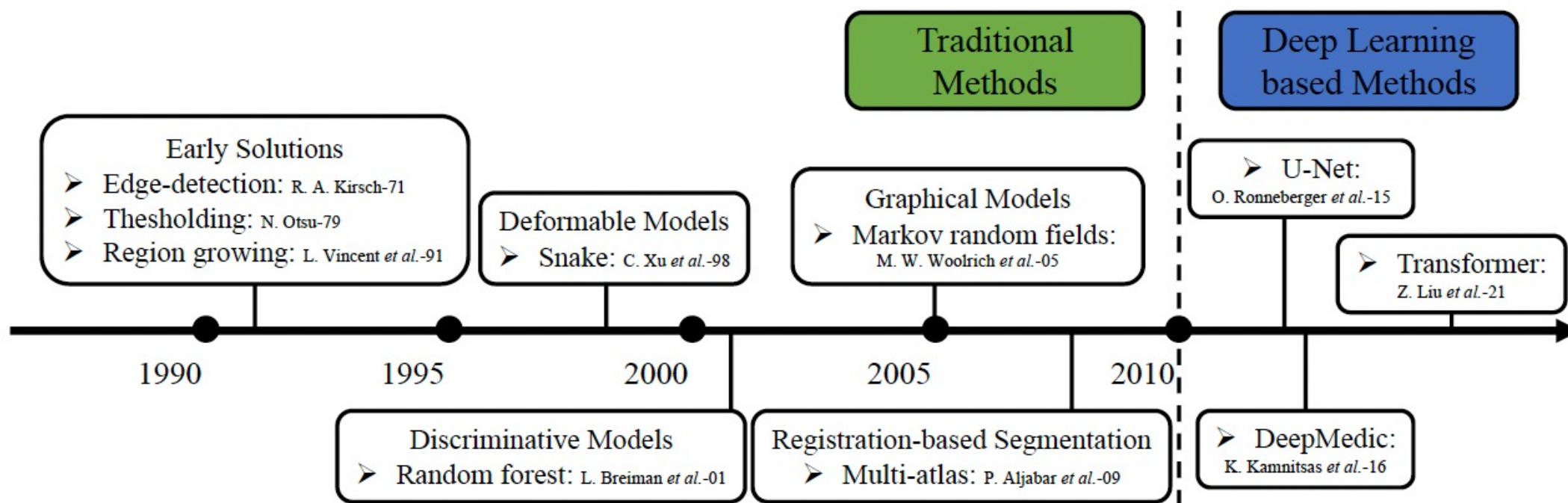


- 多用于简化繁琐的标注过程
- 后续任务的先决条件

- | | |
|---|---|
| ☐ Abdomen 15 | ☐ Lower Limb 4 |
| ☐ Colon 1 | ☐ Knee 4 |
| ☐ Kidney 6 | ☐ Pelvis 10 |
| ☐ Liver 11 | ☐ Cervix 2 |
| ☐ Pancreas 3 | ☐ Prostate 8 |
| ☐ Spleen 3 | ☐ Skin 2 |
| ☐ Cardiac 13 | ☐ Skin 2 |
| ☐ Heart 13 | ☐ Spine 3 |
| ☐ Head and Neck 54 | ☐ Spinal Cord 0 |
| ☐ Brain 41 | ☐ Vertebral Column 3 |
| ☐ Cranium 0 | ☐ Thorax 25 |
| ☐ Retina 12 | ☐ Breast 9 |
| ☐ Teeth 1 | ☐ Lung 18 |
| | ☐ Upper Limb 0 |
| | ☐ Hand 0 |



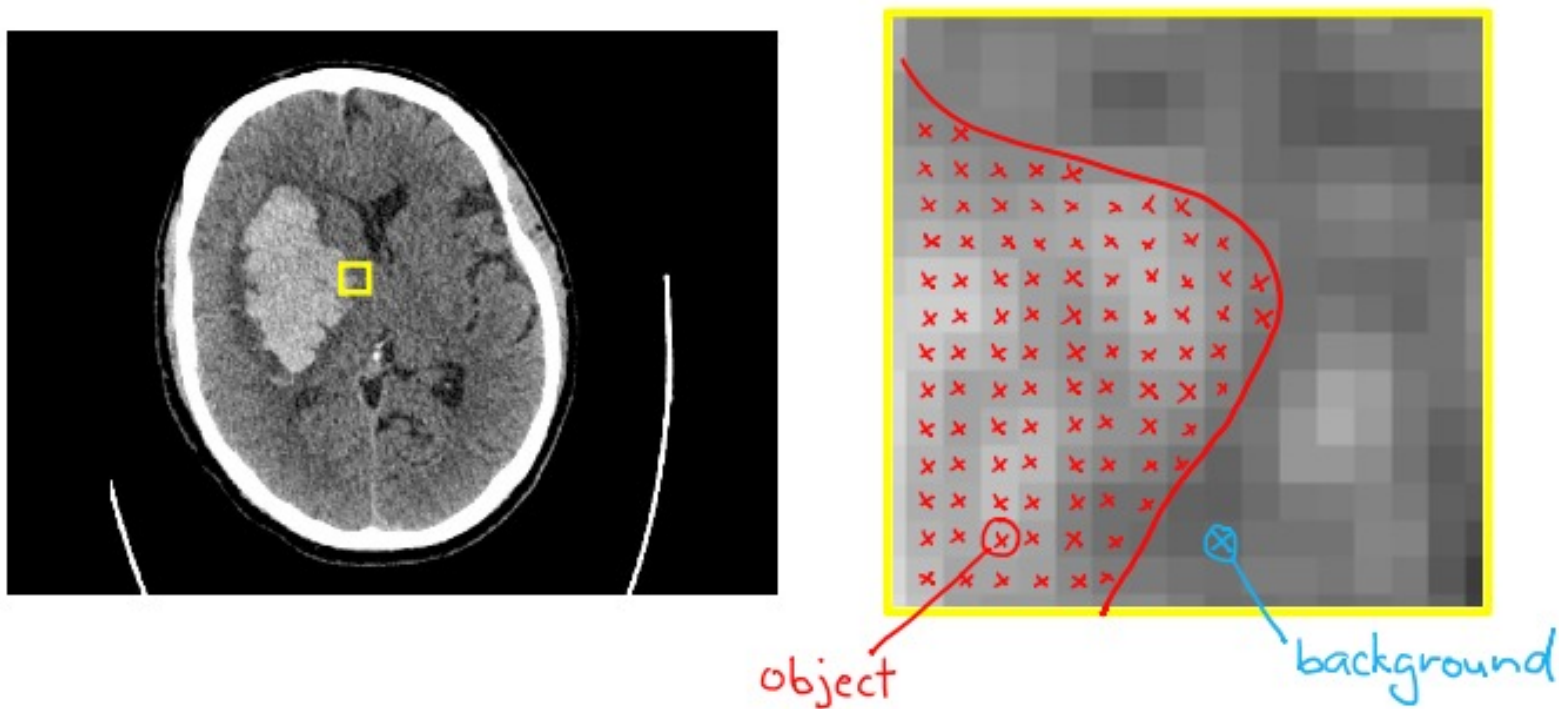
- 研究已有**50年**历史
- 2015年 (U-Net) 以后迅速收敛到神经网络



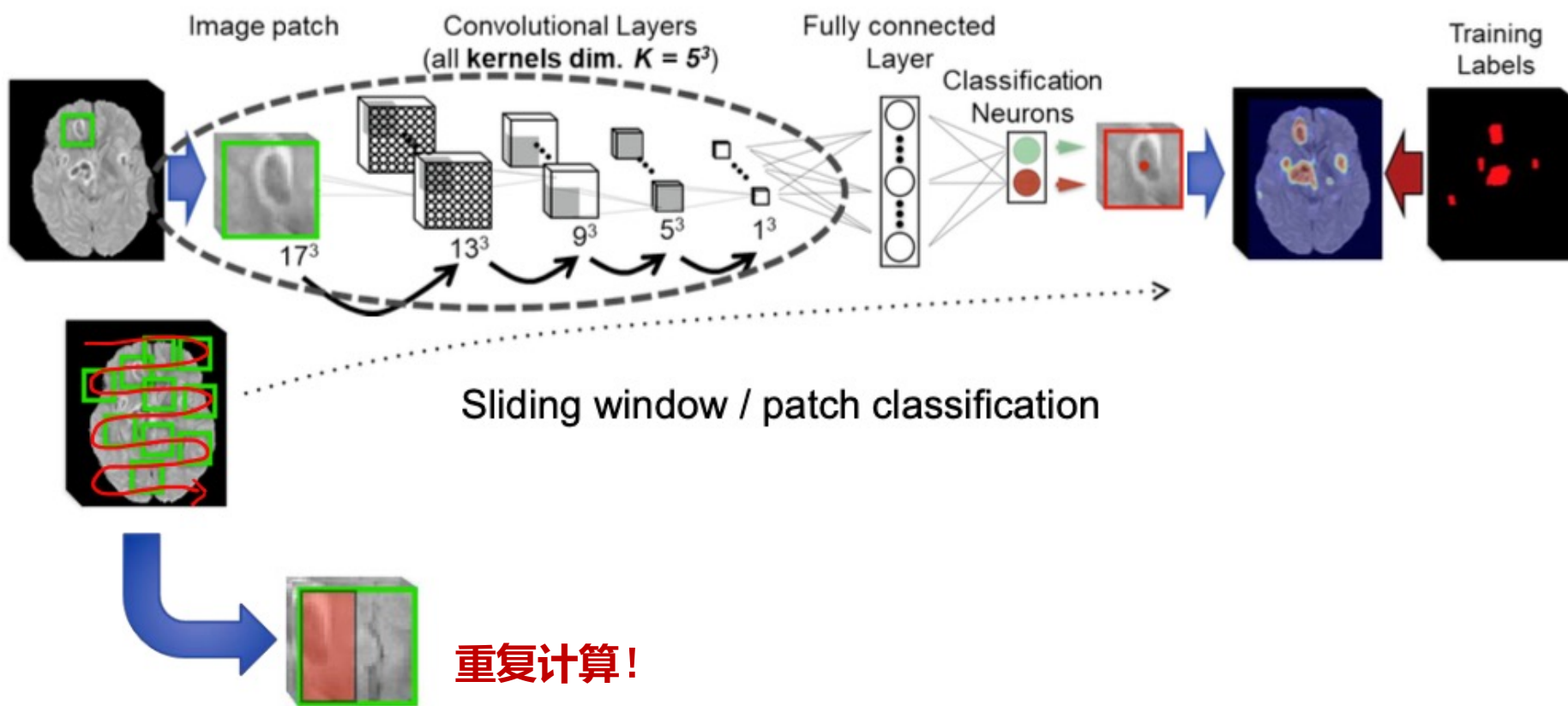
O. Ronneberger, P. Fischer and T. Brox. "U-net: Convolutional networks for biomedical image segmentation", MICCAI, 2015

Most cited paper in the MICCAI history

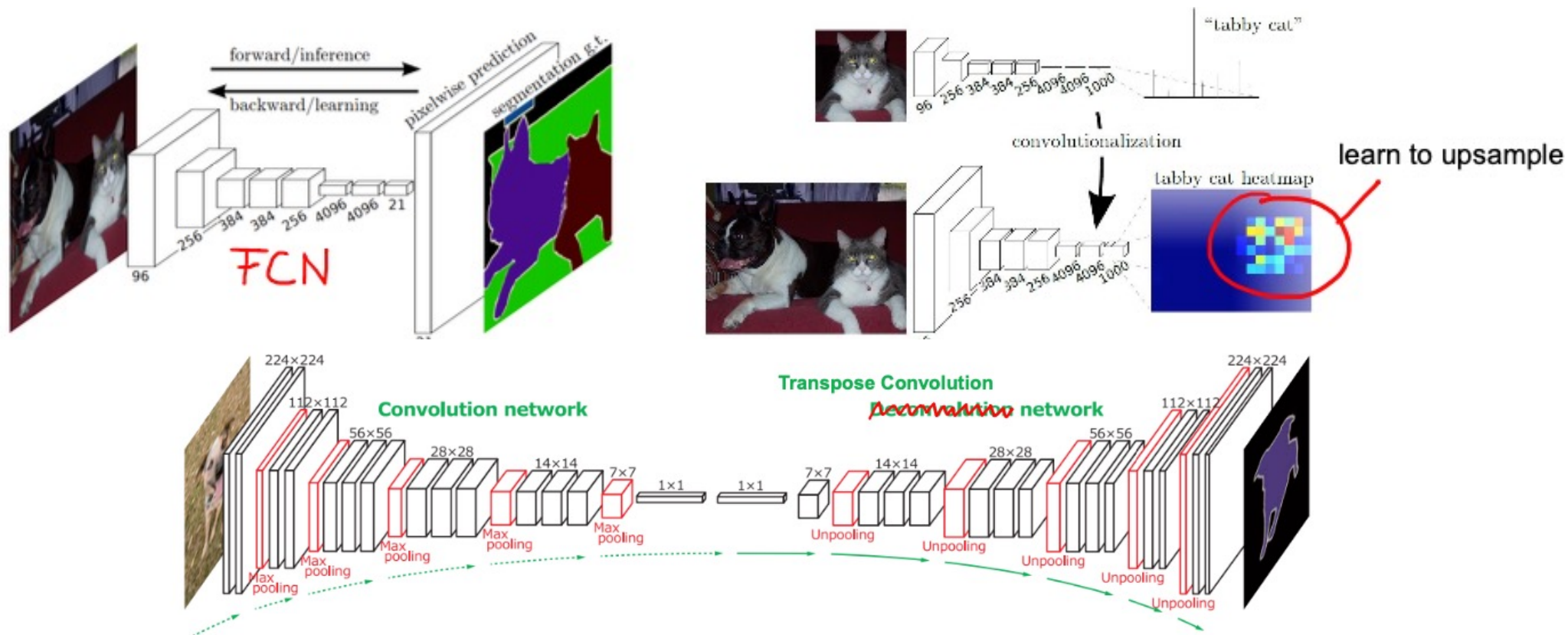
- 基于稠密分类的图像分割



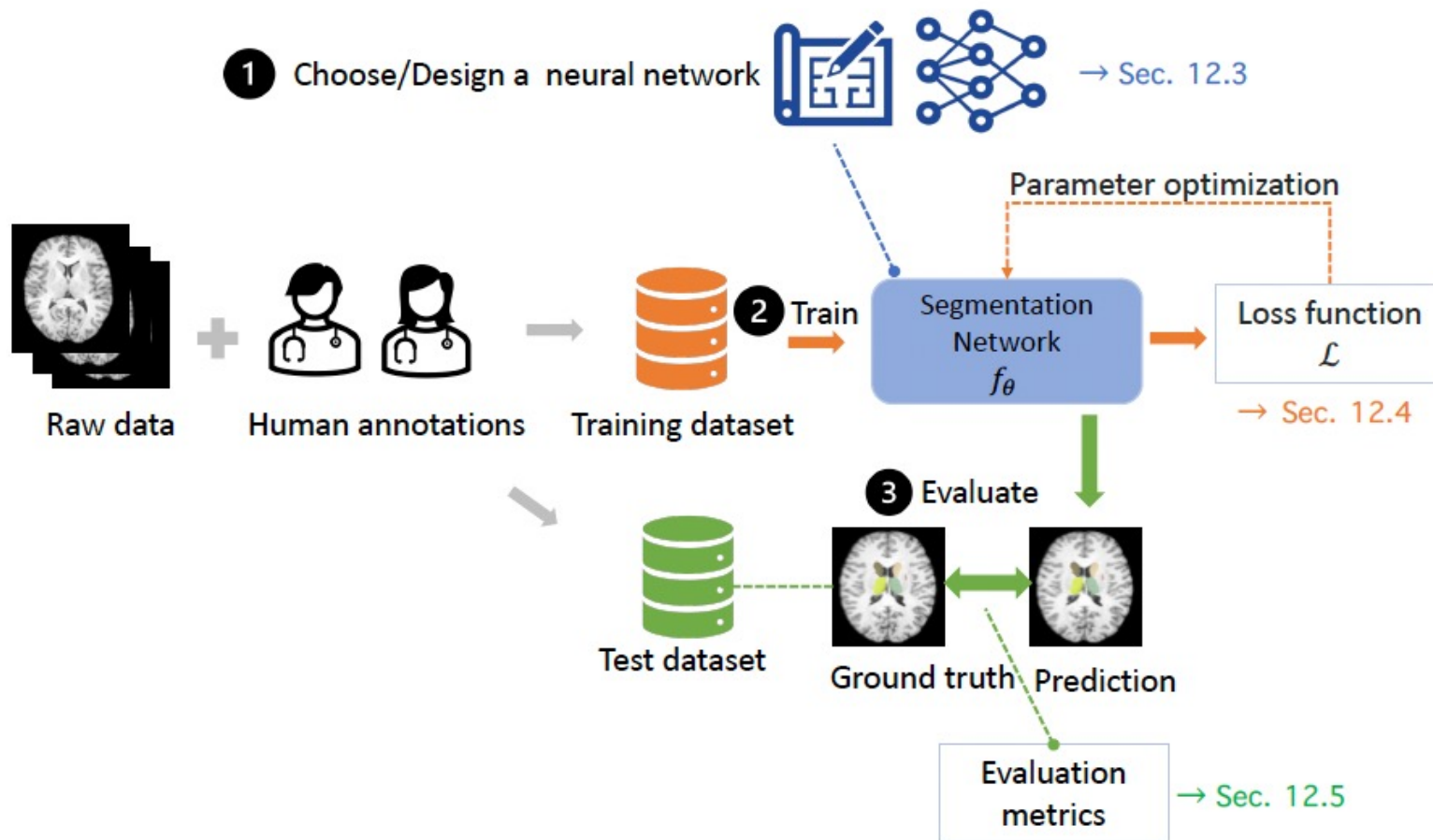
- 基于稠密分类的图像分割



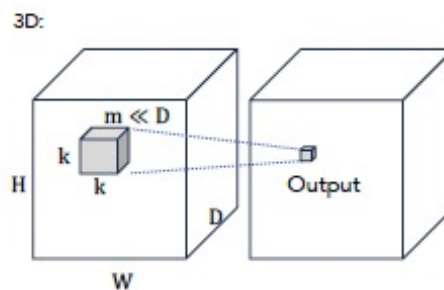
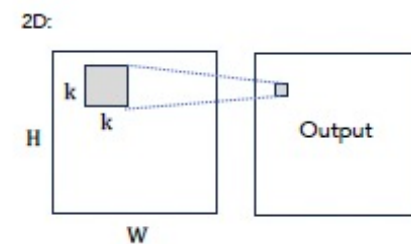
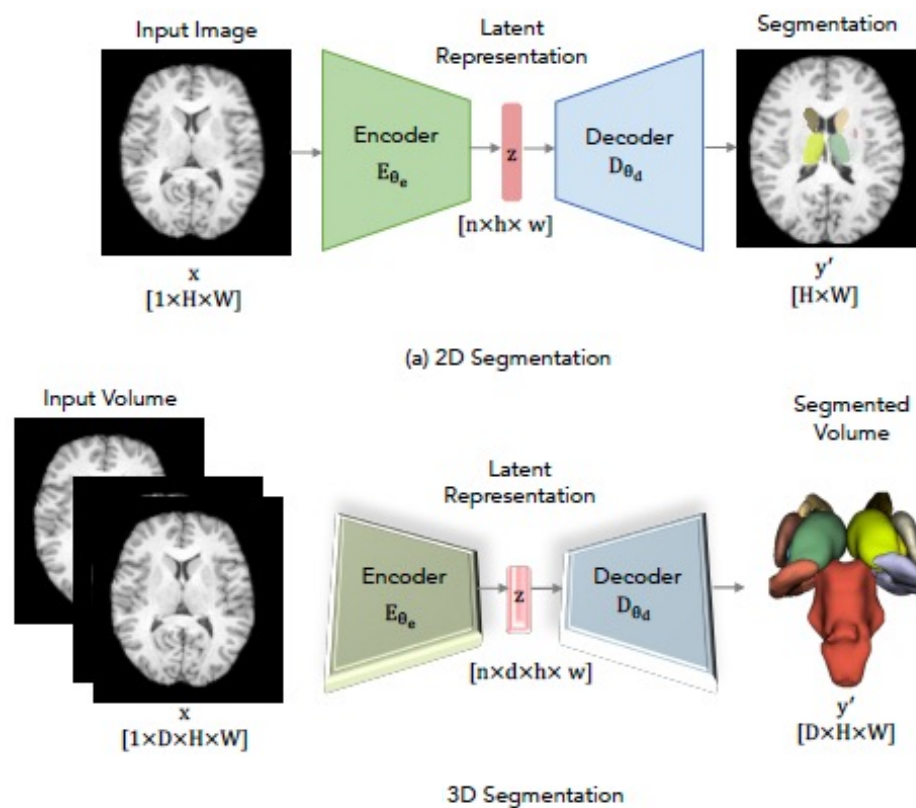
- 全卷积神经网络



- 分割算法的开发过程

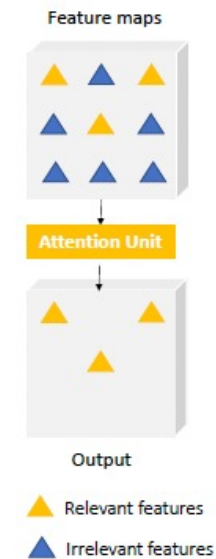
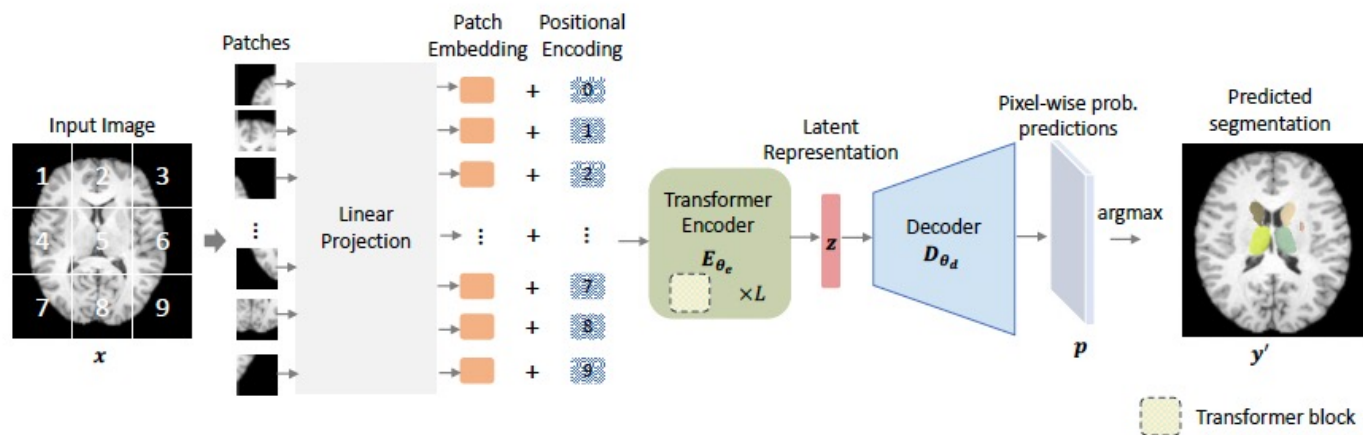


- 与自然图像不同，医学图像大多为3D



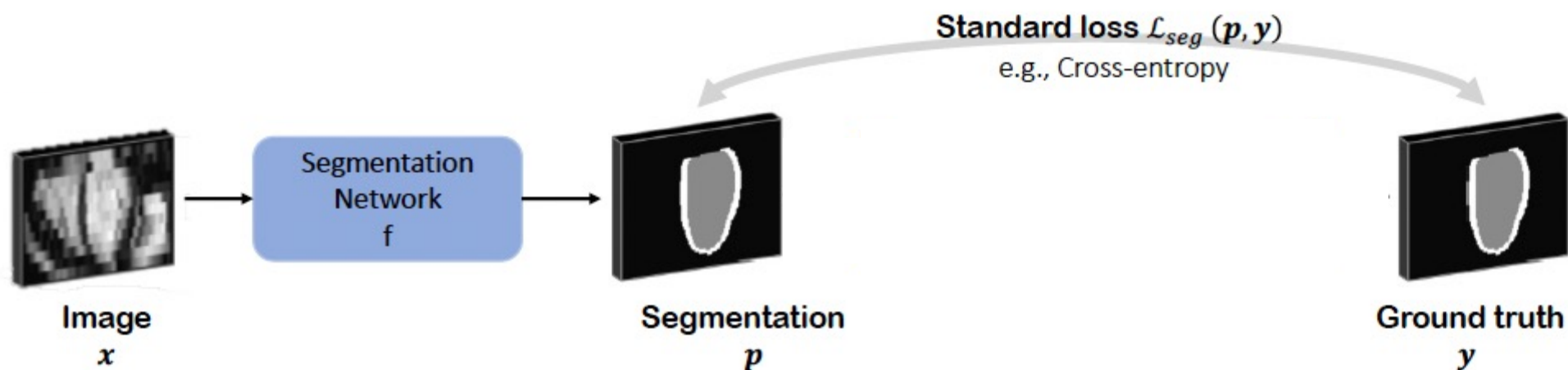
2D vs 3D Operation

- Transformer和注意力机制强调空间相关性



Attention block

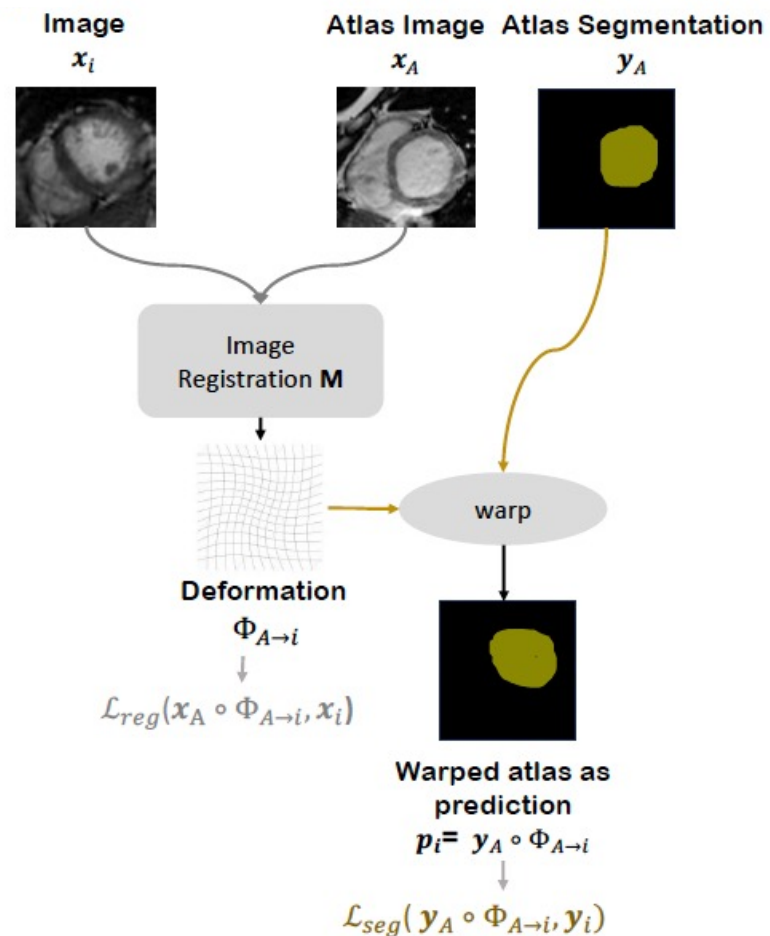
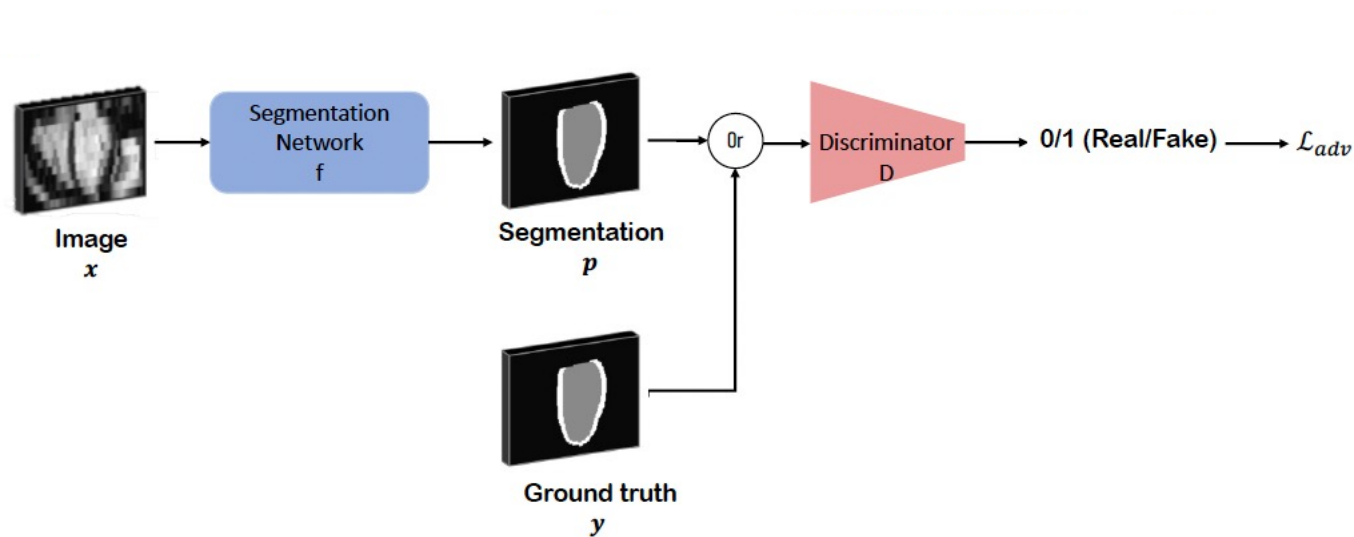
- 全监督学习常规为Cross-entropy和Soft-Dice损失的平均



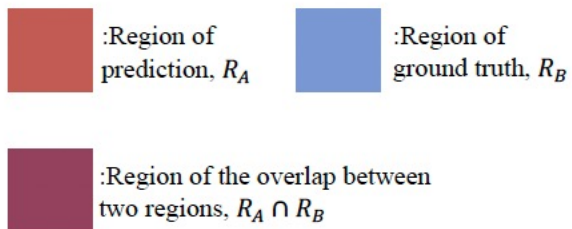
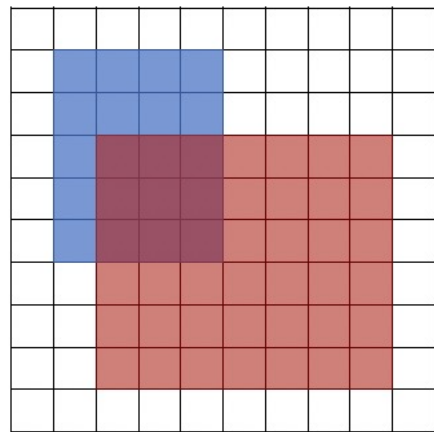
$$\mathcal{L}_{CE} = -\frac{1}{n} \sum_{i=1}^n \sum_{c=1}^C y_i^c \log(p_i^c),$$

$$\mathcal{L}_{Dice} = 1 - \frac{2 \sum_{i=1}^n \sum_{c=1}^C y_i^c p_i^c}{\sum_{i=1}^n \sum_{c=1}^C (y_i^c + p_i^c)}.$$

- 加入分割结果的形状先验知识



- 基于区域的评估方法

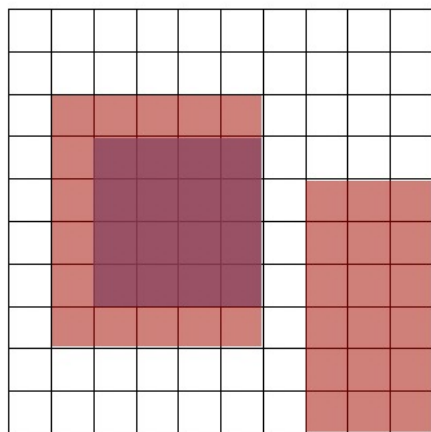


$$\text{DSC}(R_A, R_B) = \frac{2|R_A \cap R_B|}{|R_A| + |R_B|}.$$

- 基于区域的评估方法

(a) Over-Segmentation

DSC = 0.5
Sensitivity = 1.0
Precision = 0.33



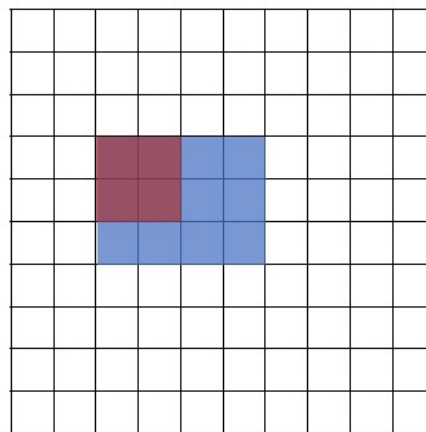
:Region of
prediction, R_A



:Region of
ground truth, R_B

(b) Under-Segmentation

DSC = 0.5
Sensitivity = 0.33
Precision = 1.0

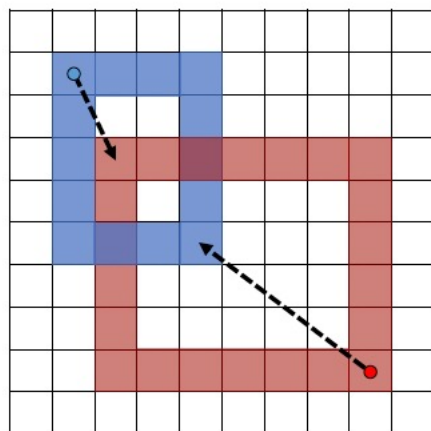


:Region of the overlap between
two regions, $R_A \cap R_B$

$$\text{Sensitivity}(R_A, R_B) = \frac{|R_A \cap R_B|}{|R_B|}$$

$$\text{Precision}(R_A, R_B) = \frac{|R_A \cap R_B|}{|R_A|}$$

- 基于距离的评估方法

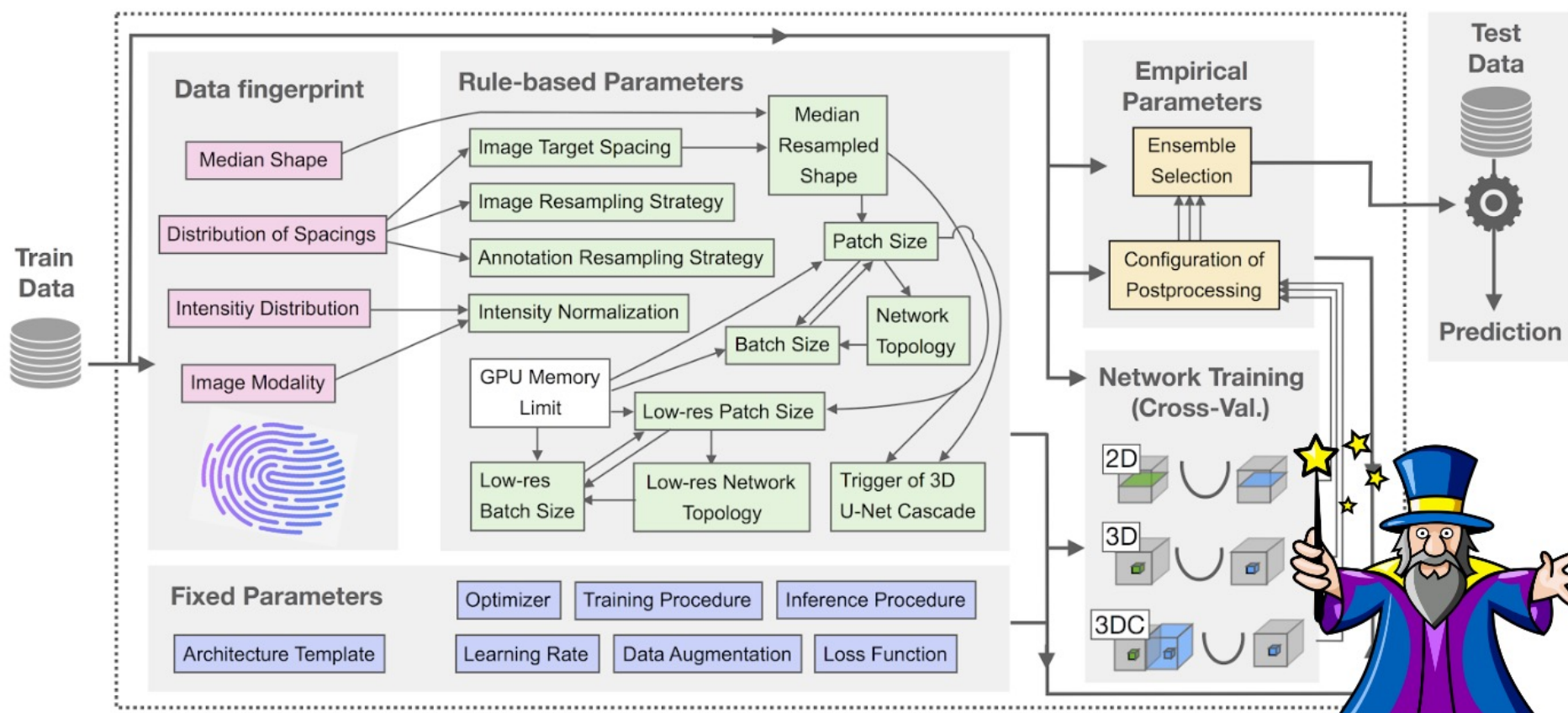


-  :Contour of prediction, S_A  :Contour of ground truth, S_B
-  :The maximum distance from all pixels in S_B to S_A ,
 $\max_{b \in S_B} \min_{a \in S_A} d(a, b)$
-  :The maximum distance from all pixels in S_A to S_B ,
 $\max_{a \in S_A} \min_{b \in S_B} d(a, b)$

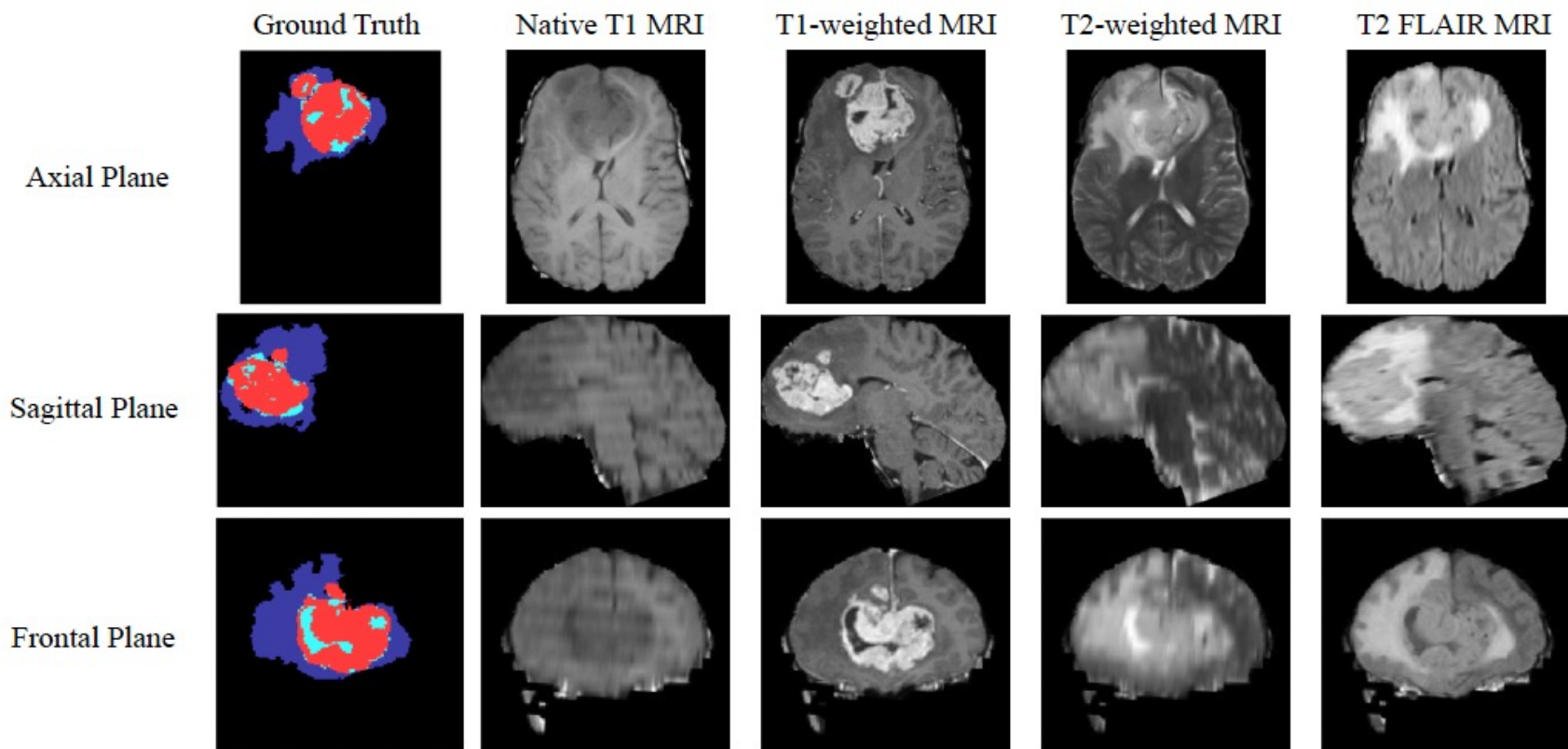
$$HD(S_A, S_B) = \max \left\{ \max_{b \in S_B} \min_{a \in S_A} d(a, b), \max_{a \in S_A} \min_{b \in S_B} d(a, b) \right\}$$

通用的解决方案

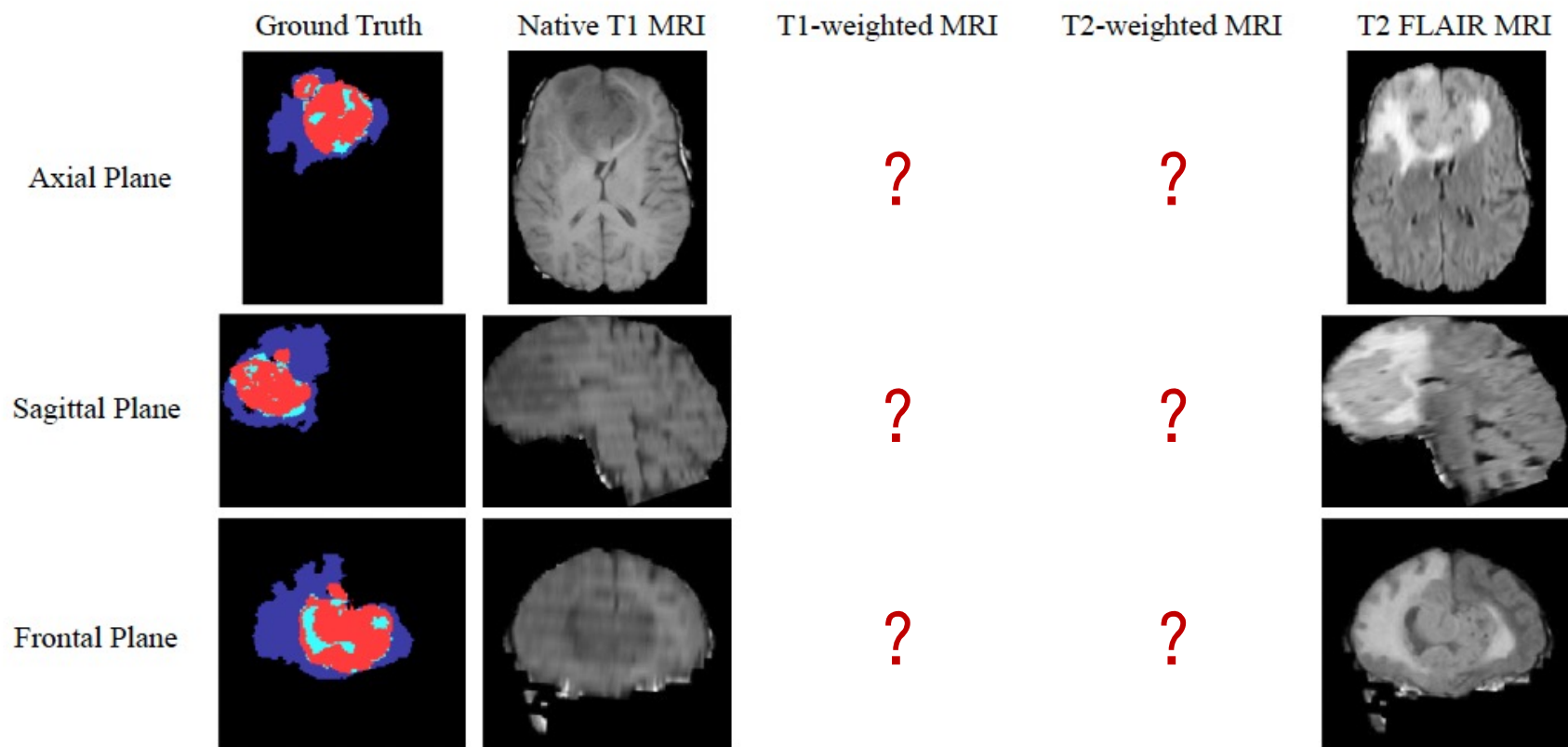
- Q: 是否所有任务都可以用同一个框架解决?
- A: 大部分可以



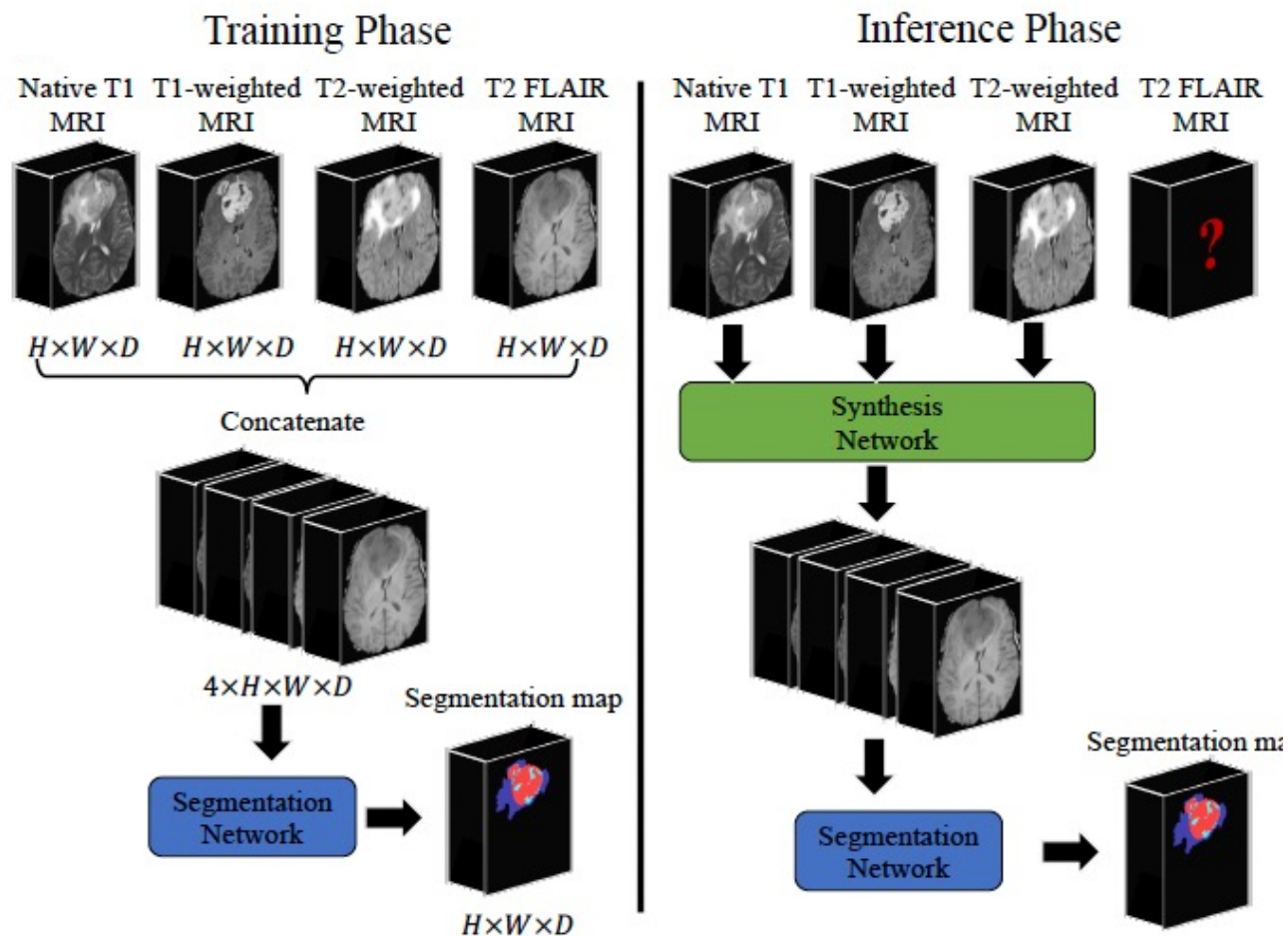
- 磁共振多序列分割



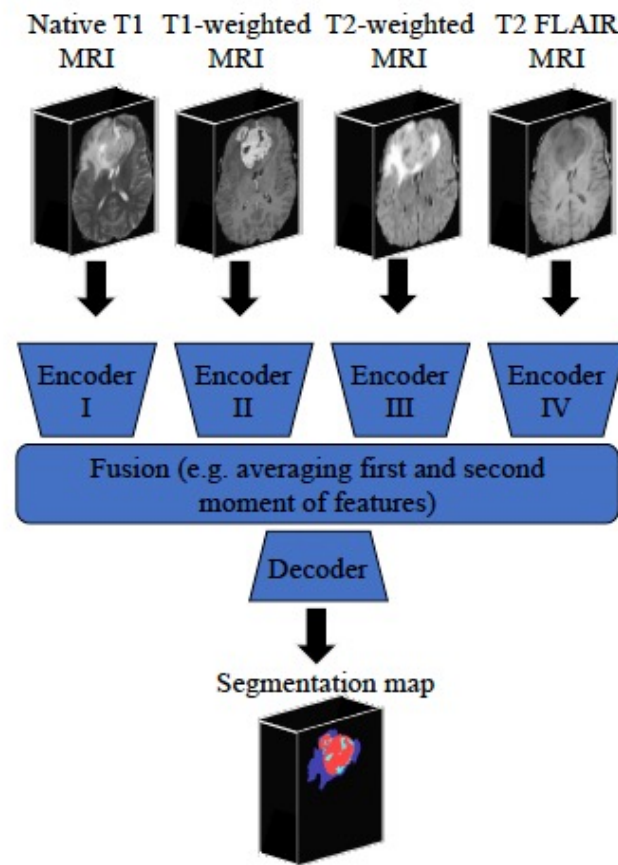
- 问题：测试时**缺少配对序列**



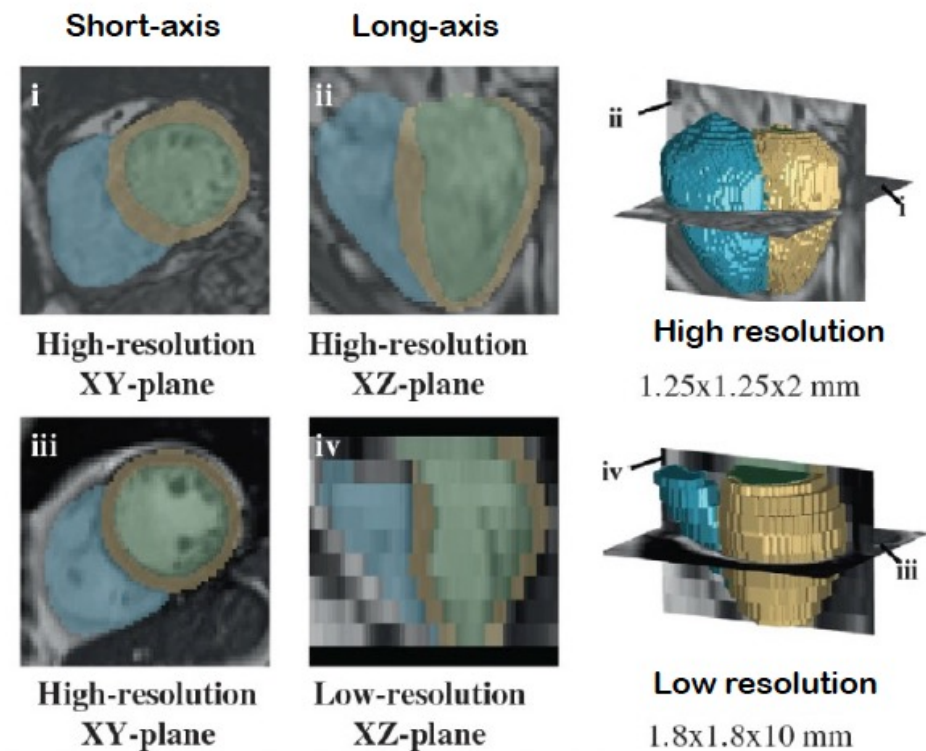
- 基于模态插补的分割方法
- **优势:**
 - 可以利用先进的生成模型技术
 - 从而效果更好



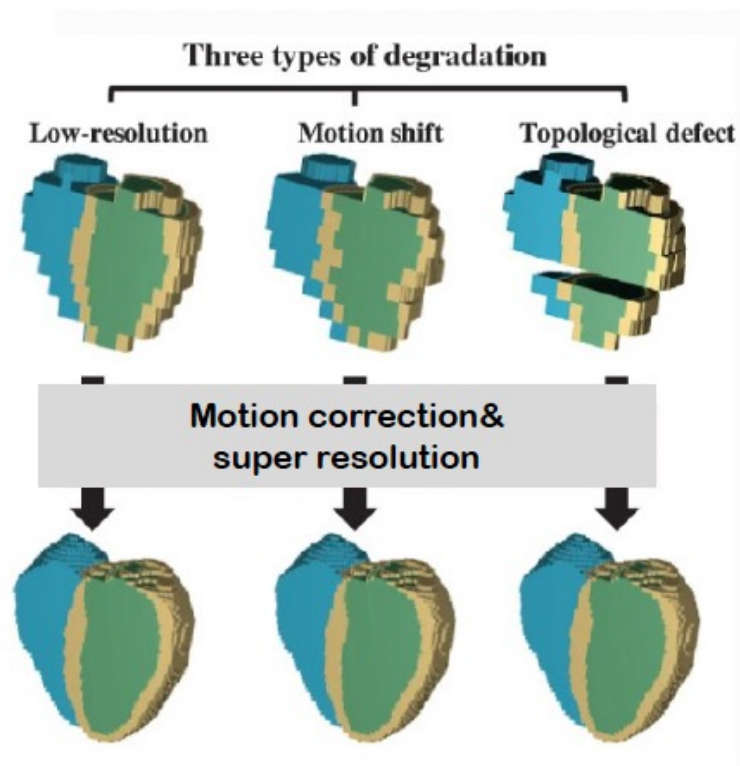
- 基于特征融合的分​​割方法
- **优势:**
 - 可以处理多种模型丢失情况
 - 从而应用更为灵活



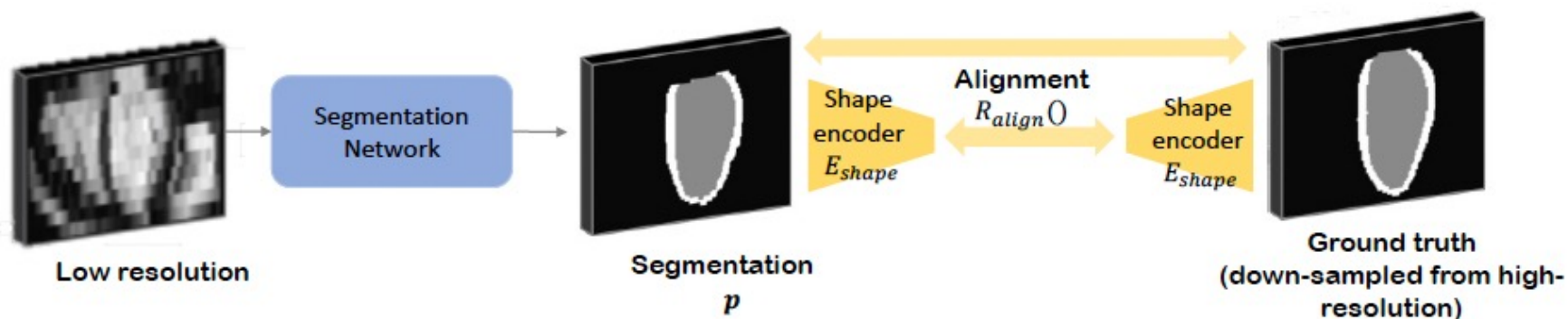
- 心脏在磁共振扫描中变动快，分辨率低



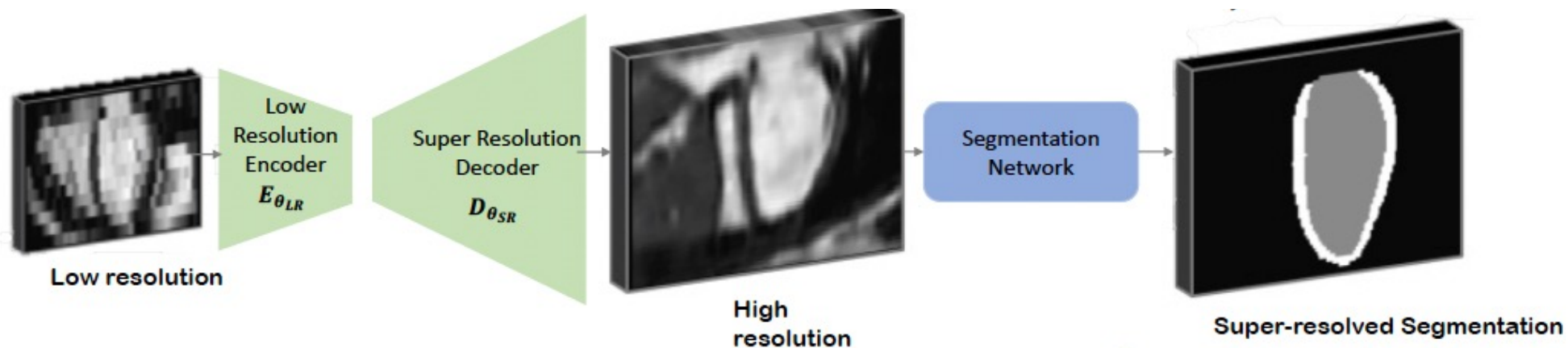
- 问题：心脏**形状**信息容易**丢失**



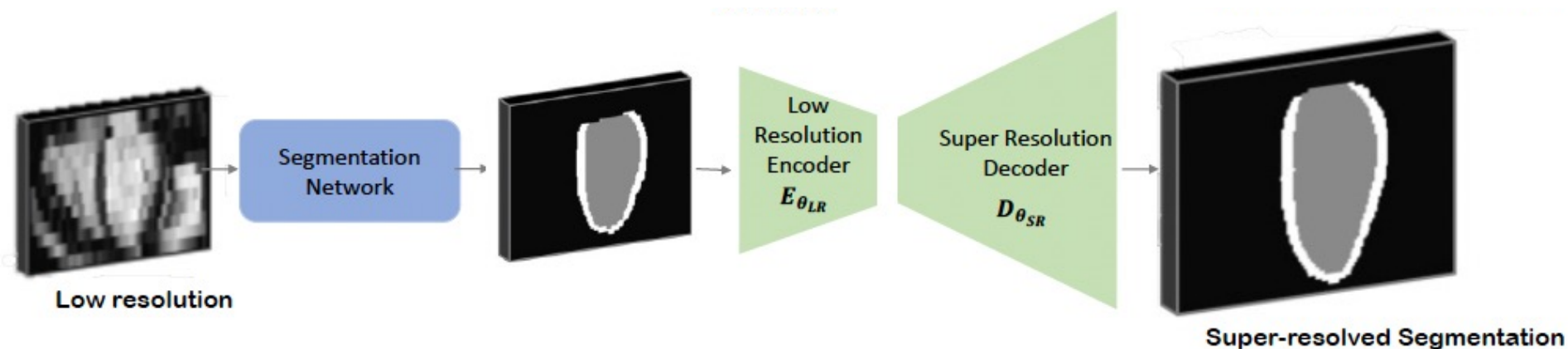
- 基于形状先验的分割方法
- 优势：
 - 可以轻量化改进分割结果



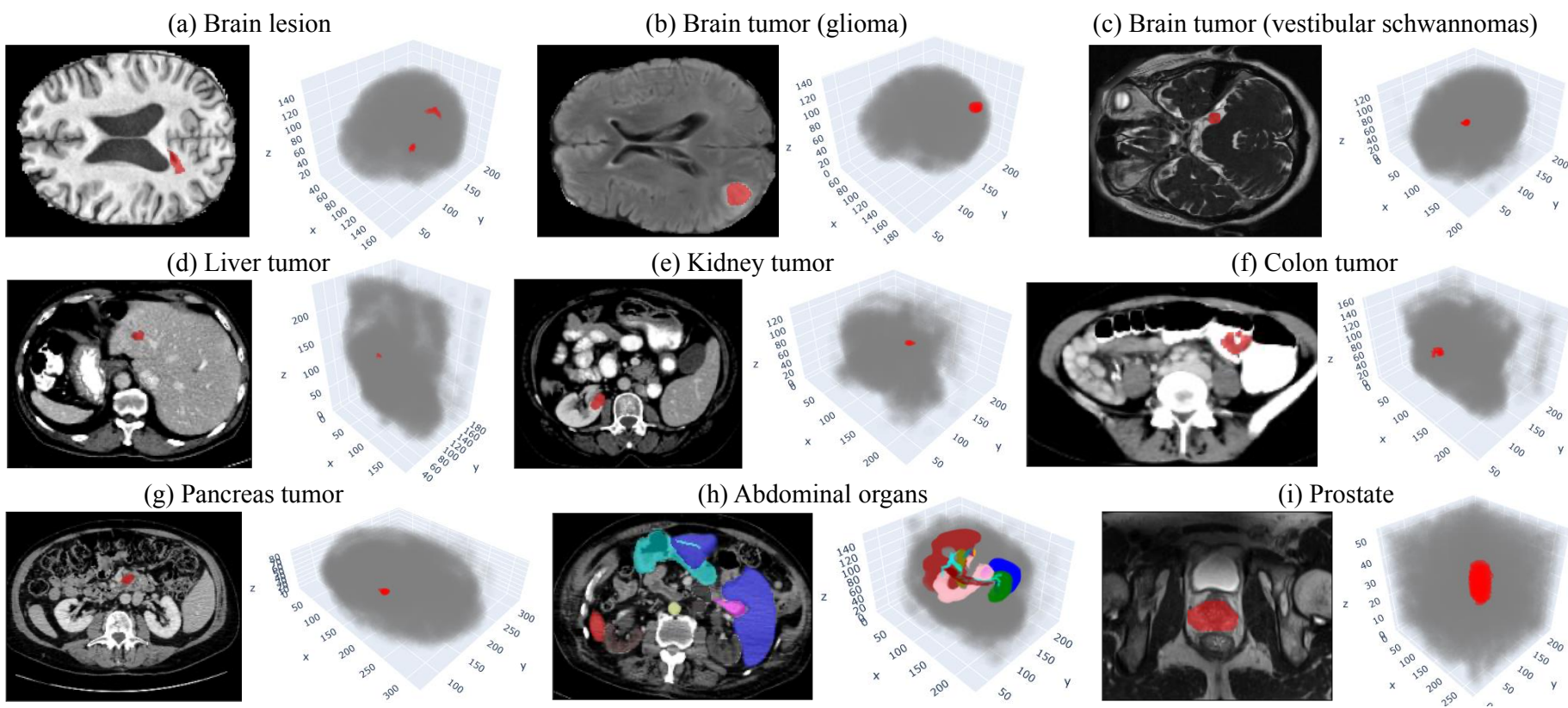
- 基于图像超分辨率的分割方法
- **优势：**
 - 可以利用高分辨率的全部信息



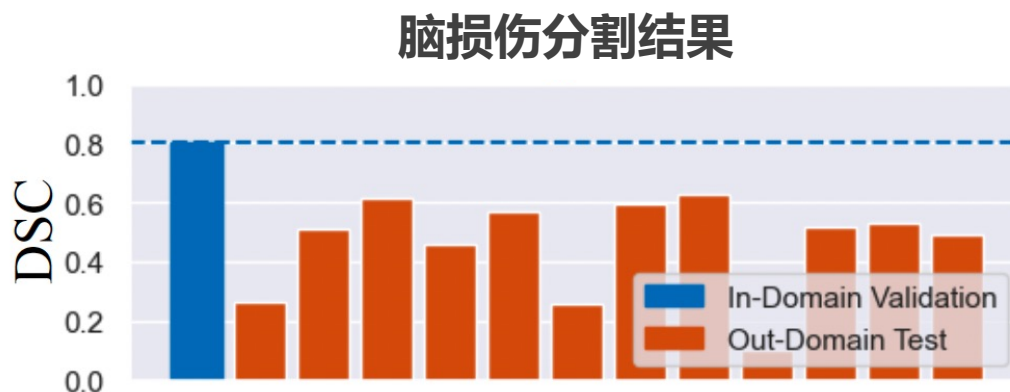
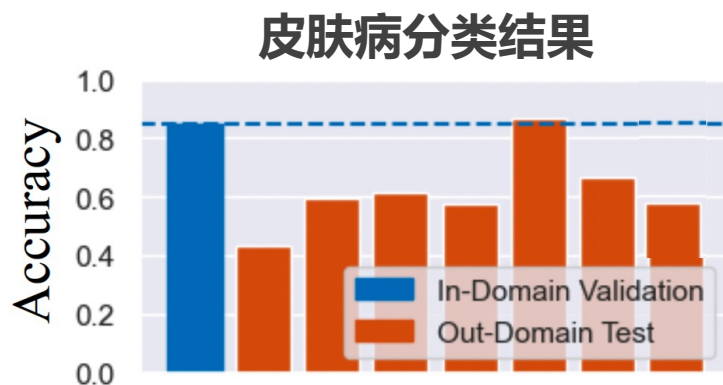
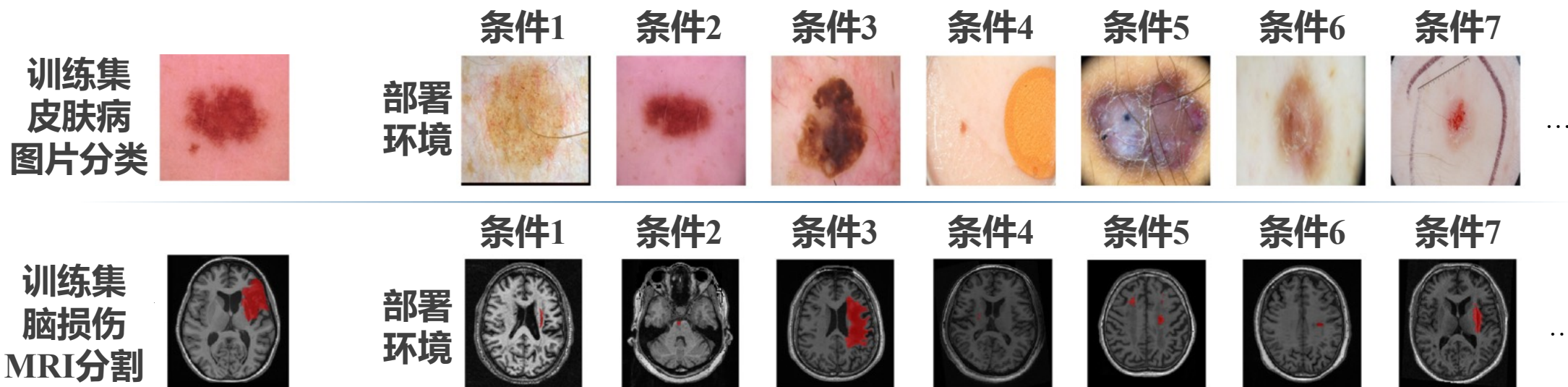
- 基于形状超分辨率的分割方法
- **优势：**
 - 作为一种后处理方法，可以应用于已训练好的分割网络



- 样本不均衡是医学图像中的一个常见问题



- 因为采集仪器不同，深度学习算法在不同成像条件下泛化性差



部署时性能下降
且下降程度未知

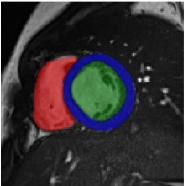
磁共振成像参数对图像质量影响显著



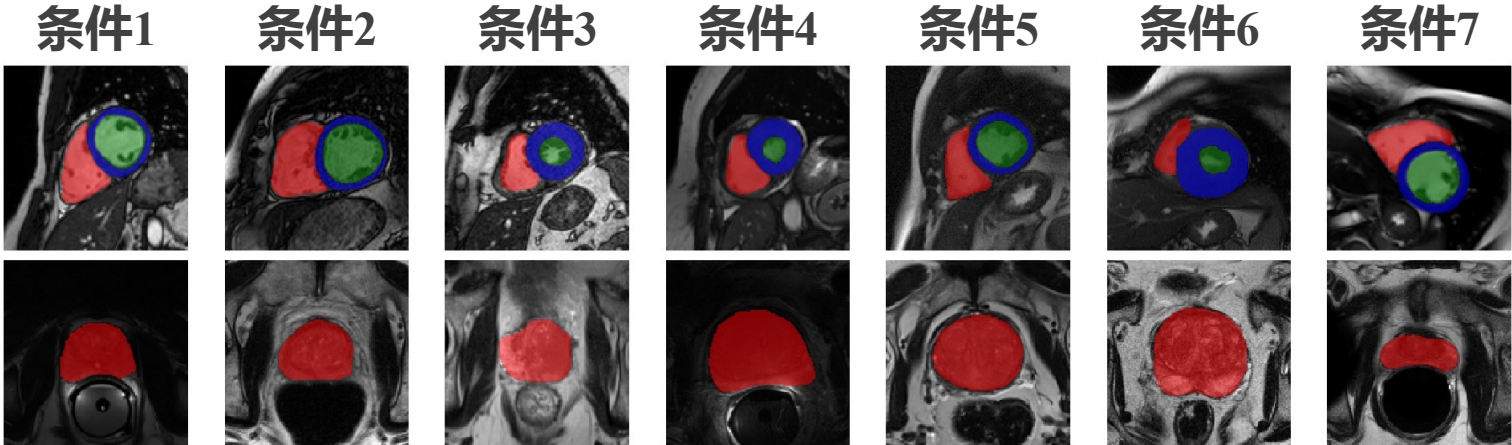
不同采集条件
脑部MRI
灰度直方图



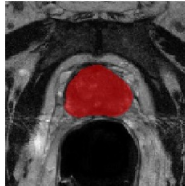
心脏
MRI分割



部署
环境



前列腺
MRI分割



部署
环境



类别不平衡导致的模型误差累积



目标未对齐



模型不可靠

模型表现

部署环境

优化算法

数据分布



可靠性保障

泛化性

模型拟合



- 判别模型和生成模型本质上是一个问题的两个方面
- 高斯过程是一种基于贝叶斯定理的非参数模型，适用于连续空间的建模
- 深度学习的核心结构源于多层感知机
- CNN是一种通过卷积核进行局部连接和参数共享的特化神经网络
- 医学图像分割是医学图像分析领域的核心任务，主流方法基于深度学习技术