

Linear $T(\alpha\vec{u} + \beta\vec{v}) = T(\alpha\vec{u}) + T(\beta\vec{v})$ *

Invariant to f $T(f(u)) = T(u)$

Equivariant to f $T(f(u)) = f(T(u))$ *

$\star I(i, j) = \sum_{m=-\infty}^{\infty} \sum_{n=-\infty}^{\infty} K(m, n) I(i+m, j+n)$ *

$\star I(i, j) = \sum \sum I(i-m, j-n) K(m, n)$ ♥

$\star$ Commutative $(I \star K)(i, j) = (K \star I)(i, j)$

$\star \star \star$ iff $K(m, n) = K(-m, -n)$

Discrete Convolution as Matrix Mult $I * K = \begin{bmatrix} K_1 & 0 & \dots & 0 \\ K_2 & K_1 & & \\ K_3 & K_2 & & \\ \vdots & \vdots & & \\ 0 & \dots & K_{m-1} & K_m \end{bmatrix} \begin{bmatrix} I_1 \\ I_2 \\ \vdots \\ I_n \end{bmatrix}$

Fwd $Z_{i,j}^{[L]} = \sum_m \sum_n w_{m,n}^{[L]} z_{i-m, j-n}^{[L-1]} + b$

Bwd $\delta_{i,j}^{[L-1]} = \partial L / \partial z_{i,j}^{[L]} \leftarrow \Theta \text{ update w.r.t. } w_{m,n}^{[L]}$

$\sum_i \sum_j (\partial L / \partial z_{i,j}^{[L]}) (\partial z_{i,j}^{[L]} / \partial z_{i,j}^{[L-1]})$

$= \sum_i \sum_j \delta_{i,j}^{[L]} w^{[L]}(m,n)$

$= \delta^{[L]} * \text{ROT}_{180}(w^{[L]})$

Diagram illustrating the forward pass of a neural network layer:

- Input matrix X (4x4) is multiplied by weight matrix K (3x4) to produce the hidden layer output $Y = \text{RELU}(X * K)$ (3x4).
- The output Y is passed through an activation function (represented by a box with '1') to produce the final output Z (3x4).
- The diagram also shows the corresponding Jacobian matrices for each step: $\frac{\partial E}{\partial Y}$ (3x4), $\frac{\partial E}{\partial K}$ (4x4), and $\frac{\partial E}{\partial X}$ (4x4).

Use θ -update and Backward Rule
 where $\partial E / \partial x = \partial E / \partial y * ROT_{180}(K)$
 $= \partial E / \partial y * K$. Slide K over $\partial E / \partial y$
 For $\partial E / \partial K = \frac{\partial E}{\partial y} * X = ROT_{180}(X * \partial E / \partial y)$
 $ROT_{180}(\text{Slide } \partial E / \partial y \text{ over } X)$

Max Pooling Fwd: $z_{ij}^{[L]} = \max\{z_i^{[L-1]}\}$
 Bwd: $\delta^{[L-1]} = \{\delta^{[L]}\}_{i^*, j^*}$ since
 $\frac{\partial z_{ij}^{[L]}}{\partial z_i^{[L-1]}} = \mathbb{I}\{i=i^*\}$ No learnable param

1	1	2	4
5	6	7	8
3	2	1	0
1	2	3	4

→

6	8
3	4

 Max pool, 2x2 filter
 Stride 2

Upsampling Increase Spatial feature size.
Nearest Neighbor, Bed of Nails, Max-unpool.

Transposed Convolutions

1. $Z = S - 1$; $p' = k - p - 1$
2. Insert Z rows between rows & columns
3. Add p' number of zeros around image
4. Perform convolution, Remember Flip K !

U-Net FCNN, Combine global and local feature maps by Copying corresponding tensor from earlier stages in each upsampling stage. $\rightarrow$ Captures global, local ctx, $\rightarrow$ acc

Res.Conn. Help maintain local features as things not completely downsampled every stage

Parameters Let I : length of input vol size
 F : length of filter size, P : amount of zero padding
 S : Stride, then O : output size

$$0 = \left[\frac{I - F + P_{\text{start}} + P_{\text{end}}}{S} \right] + 1$$

often $P_{\text{start}} = P_{\text{end}} \triangleq P$

Complexity	In	Out	Params
CONV: $F \begin{bmatrix} \square \\ \otimes \end{bmatrix} \times k \begin{matrix} \Delta \\ \Delta \end{matrix}$	$I \times I \times C$	$O \times O \times K$	$(F \times F \times C \times K) \times 4$
POOL: $F \begin{bmatrix} \square \\ \max \end{bmatrix}$	$I \times I \times C$	$O \times O \times C$	ZERO
FCNN: 	N_{in}	N_{out}	$(N_{in}+1) \times N_{out}$

K: how many filters

$$\begin{aligned} \hat{y}_t &= \text{Softmax}(Vh_{t+1}), \quad h_t = \tanh(W h_{t-1} + U x_t + b) \\ \mathcal{L}_t &= -\sum_k y_{t,k} \log \hat{y}_{t,k} \quad \mathcal{L} = \sum_t \mathcal{L}_t \\ \frac{\partial \mathcal{L}}{\partial \vec{w}} &= -\sum_k \frac{\partial \mathcal{L}_t}{\partial \hat{y}_{t,k}} \frac{\partial \hat{y}_{t,k}}{\partial \vec{w}} \quad \text{Compute element-Wise} \\ &\quad \text{drop } t \text{ for brevity} \end{aligned}$$

$$\begin{aligned} \frac{\partial \hat{y}_k}{\partial \theta_i} &= \frac{\partial}{\partial \theta_i} \left(\frac{\exp(\theta_k)}{\sum_j \exp(\theta_j)} \right) [i=k] = \hat{y}_i (1 - \hat{y}_i) \\ &= \frac{y_i}{\hat{y}_i} \hat{y}_i (1 - \hat{y}_i) - \sum_{i \neq k} \frac{y_k}{\hat{y}_k} \hat{y}_k \hat{y}_i = \hat{y}_i - y_i \\ \frac{\partial \mathcal{L}}{\partial \mathbf{h}_t} &= \frac{\mathcal{L}}{\mathcal{L}_t} \frac{\partial \mathcal{L}_t}{\partial \mathbf{h}_t} = \frac{\partial \mathcal{L}_t}{\partial \mathbf{h}_t} + \sum_{l=t+1}^T \frac{\partial \mathcal{L}_l}{\partial \mathbf{h}_t} \\ \frac{\partial \mathcal{L}_t}{\partial \mathbf{h}_t} &= \frac{\partial \mathcal{L}_t}{\partial \mathbf{h}_t} \prod_{k=t}^{t-1} \frac{\partial \mathbf{h}_{k+1}}{\partial \mathbf{h}_k} \leftrightarrow \text{diag}(1 - \tanh^2(a_{k+1})) \mathbf{W} \end{aligned}$$

$$\frac{\partial \mathcal{L}_t}{\partial \mathbf{V}} = \sum_i \frac{\partial \mathcal{L}_t}{\partial \mathbf{O}_t} \frac{\partial \mathbf{O}_t}{\partial \mathbf{V}} \leftrightarrow \sum_i \frac{\partial \mathcal{L}_t}{\partial \mathbf{O}_{t,i}} \frac{\partial \mathbf{O}_{t,i}}{\partial \mathbf{V}}$$

$$\mathbf{O}_{t,i} = \mathbf{v}_i \cdot \mathbf{h}_{t,1} + \dots + \mathbf{v}_{i,m} \cdot \mathbf{h}_{t,m} + \mathbf{c}$$

$$\Rightarrow \frac{\partial \mathcal{L}_t}{\partial \mathbf{O}_{t,i}} \frac{\partial \mathbf{O}_{t,i}}{\partial \mathbf{v}_{k,l}} = \mathbf{h}_{t,l} \mathbb{I}\{i=k\} = \mathbf{M}$$

$\mathbf{K} \times \mathbf{M}$ zero except $\mathbf{M}_{i,i}$, so $= \mathbf{h}_t^T \frac{\partial \mathcal{L}_t}{\partial \mathbf{O}_t}$

$$\frac{\partial \mathcal{L}_t}{\partial \mathbf{W}} = \frac{\partial \mathcal{L}_t}{\partial \mathbf{O}_t} \frac{\partial \mathbf{O}_t}{\partial \mathbf{h}_t} \frac{\partial \mathbf{h}_t}{\partial \mathbf{W}} \leftrightarrow \sum_{k=1}^t \frac{\partial \mathcal{L}_t}{\partial \mathbf{h}_k} \frac{\partial \mathbf{h}_k}{\partial \mathbf{W}}$$

$$\frac{\partial h_t}{\partial h_{t-k}} = \prod_{t'=k+1}^t \frac{\partial h_{t'}}{\partial h_{t'-1}} = \prod_{t'=k+1}^t \text{diag}(1 - h_{t'}^2) W$$

Vanishing / Exploding Gradient $W = Q\Lambda Q^T$

$$(W^T)^{t-k-1} = (Q^T \Lambda Q)^{t-k-1} = Q^T \Lambda^{t-k-1} Q$$

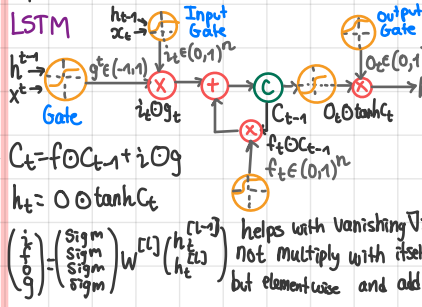
$$\Rightarrow \|\text{diag}(f(h_{t-1}))\| < \gamma \quad f = \sigma \text{ OR } \tanh$$

$\lambda_1 < \frac{1}{\gamma}$ vanish, $\lambda_1 > \frac{1}{\gamma}$ explodes as $t \rightarrow \infty$

$$\lambda_1 < \frac{1}{\gamma}: \left\| \frac{\partial h_t}{\partial h_{t-1}} \right\| \leq \|W\| \|\text{diag}(f(h_{t-1}))\| < \frac{1}{\gamma} < 1$$

Let $\eta \in \mathbb{R}$ s.t. $\left\| \prod_{i=1}^t \frac{\partial h_i}{\partial h_{i-1}} \right\| < \eta^{t-k} \rightarrow 0$ as $t \rightarrow \infty$

by induction



Gradients can still explode in LSTM

Gradient Clipping $\hat{g} \leftarrow \partial \mathcal{L} / \partial \theta$ When $\hat{g}$
 $\theta \leftarrow \theta - \lambda \hat{g} \tau / \|\hat{g}\|$ if $\|\hat{g}\| > \tau$ explode

Implicit	Explicit
- MC (GSN) - Direct -GAN	- Tractable [Autoregressive, NF] - Approximate -MCMC (Boltzmann Machine) -Variational (VAE, Diffusion)

Efficient Data Representation,
Desirable: Continuity & Interpolation

Original Space $X \xleftrightarrow[\text{g}]{f}$ Latent Space Z
 $[g \circ f]$ Approximates $\Pi(\cdot)$ on the data
 Linear AE finds Linear embedding/code z
 Non-Linear AE $\hat{\theta}_f, \hat{\theta}_g = \arg \min_{\theta_f, \theta_g} \mathcal{L}$ where
 $\mathcal{L} = \sum_{n=1}^N \|x_n - \hat{x}_n\|^2 = \sum_{n=1}^N \|x_n - g(f(x_n))\|^2$
 Undercomplete $\dim(z) < \dim(x)$ $\oplus$ to extract features to train, $\ominus$ Out of distribution Samples
 OverComplete $\dim(z) > \dim(x)$ Perf Reconstruction
 Denoising AE $\mathcal{L}$ evaluated based on Com w/ Clean img. Learns transform. to rm. noise
 AE $\ominus$ Lack of continuity in Z , Struggles to generate high quality samples. $\oplus$ Good for reconstructions.

Objective $P_\theta(x) = \int_Z p(x|z)p(z)dz$. Int
Approx. posterior $P_\theta(z|x)$ with $q_\phi(z|x)$
 $\log P_\theta(x) = \mathbb{E}_{z \sim q_\phi(z|x)} [\log p(x)]$
 $= \mathbb{E}_z \left[\log \frac{p(x|z)p(z)}{p(z|x)} \right] = \mathbb{E}_z \left[\log \frac{p(x|z)p(z)}{p(z|x)} \frac{q(z|x)}{q(z|x)} \right]$
 $= \mathbb{E}_z [\log p(x|z)] - \mathbb{E}_z \left[\log \frac{q(z|x)}{p(z)} \right] +$
 $\mathbb{E}_z [\log q(z|x) / p(z|x)] \quad \downarrow \text{ELBO VA}$
 $\propto \mathbb{E}_z [\log p(x|z)] - \text{KL}(q(z|x) || p(z))$
Generate new Samples by max data likelihood

$\log p_\theta(x_i) \geq \mathcal{L}(x_i, \theta, \phi)$
 $\theta^*, \phi^* = \arg \max_{\theta, \phi} \sum_{i=1}^N \mathcal{L}(x_i, \theta, \phi)$

FWD $\hat{x}$ ① Maximize Reconstruction Likelihood, encourage samples of same category to be close in Z . $\text{0 Var} \equiv \text{AI}$

$\text{Sample } x|z \sim \mathcal{N}(\mu, \Sigma)$
 $\mu(x|z)$ $\Sigma(x|z)$
 z

$\text{Sample } z|x \sim \mathcal{N}(\mu, \Sigma)$ ② Make approx similar to prior encourages encoder to project latent space repr. evenly around center of latent space. At

$\mu(z|x)$ $\Sigma(z|x)$
 x

runtime, Only use decoder network.
 $z \sim p(z)$.
 $KL Z0 - KL = - \int p(x) \log \frac{p(x)}{q(x)} dx$
 $= \int p(x) \log \frac{q(x)}{p(x)} dx \leq \log \int p(x) \frac{q(x)}{p(x)} dx = 0$
Reparam Trick (To allow VAE backprop)
 Instead of $z \sim N(\mu, \sigma^2)$: $z = \mu + \sigma \epsilon$ $\epsilon \sim N(0,1)$
VAE-Backprop $E_z [\log p(x|z)] -$
 $KL(q(z|x) || p(z)) =$
 $E_z [\log p(x|z)] + E_z \left[\log \frac{p(z)}{q(z|x)} \right]$
 $= E_z \left[\log \frac{p_\theta(x, z)}{q_\phi(z|x)} \right]$ taking gradient w.r.t θ, ϕ
 $= \nabla_{\theta, \phi} E_{z \sim q_\phi} \left[\log \frac{p_\theta(x, f(x, \epsilon, \theta))}{q_\phi(f(x, \epsilon, \theta) | x)} \right]$
 $= E_\epsilon \left[\nabla_{\theta, \phi} \log \frac{p_\theta(x, f(x, \epsilon, \theta))}{q_\phi(f(x, \epsilon, \theta) | x)} \right]$
 $\approx \frac{1}{K} \sum_{k=1}^K \nabla_{\theta, \phi} \log(p_\theta(x, f) / q_\phi(f|x))$
*** Monte-Carlo Gradient Estimator**
 $\nabla_\psi E_{p_\psi(z)}[f(z)] = \nabla_\psi \int p_\psi(z) f(z) dz$
 f differentiable. Given $z = g(\epsilon, \psi)$
 Via Reparam Trick and change of variables $p(z) | dz| = p(\epsilon) | d\epsilon|$ so
 $\int p_\psi(z) f(z) dz = \int p(\epsilon) f(z) d\epsilon =$
 $p(\epsilon) f(g(\epsilon, \psi)) d\epsilon$ Hence
 $E_{p(\epsilon)} [\nabla_\psi f(g(\epsilon, \psi))] =$
 $\nabla_\psi \int p(\epsilon) f(g(\epsilon, \psi)) d\epsilon = \nabla_\psi E_{p(\epsilon)} [f(g(\epsilon, \psi))]$
 $= \frac{1}{S} \sum_{s=1}^S [\nabla_\psi f(g(\epsilon^{(s)}, \psi))] \epsilon^{(s)} \sim p(\epsilon)$
Analytical KL $\forall p(z) = N(\mu_p, \sigma_p^2 I)$ and
 $q(z) = N(\mu_q, \sigma_q^2 I)$. $\int p(z) \log(q(z)) dz$
 $= -\frac{J}{2} \log(2\pi) - \frac{1}{2} \sum_{j=1}^J \log \sigma_{q,j}^2 -$
 $\frac{1}{2} \sum_{j=1}^J (\sigma_{p,j}^2 + (\mu_{p,j} - \mu_{q,j})^2) / \sigma_{q,j}^2$
Hierarchical Latent Variable Models
 $(x) \rightarrow [h_1] \dots [h_L] (z_L) \dots (z_1) \rightarrow (x)$
 $(z_L) \rightarrow (z_1)$
Stoch $VAE \Rightarrow L=1$
 $q_\phi(z_1, \dots, z_L | x) = q_\phi(z_1 | x) \dots q_\phi(z_L | x)$
 $p_\theta(z_1, \dots, z_L) = p_\theta(z_L) p_\theta(z_{L-1} | z_L) \dots p_\theta(z_1 | z_2)$

$$\mathcal{L}(\theta, \phi; x) = \mathbb{E}_{q_{\phi}(z|x)} [\log p_{\theta}(x|z)] - \text{KL}(q_{\phi}(z|x) \| p_{\theta}(z))$$

$$= \mathbb{E}_{q_{\phi}(z_1, \dots, z_L|x)} [\log p_{\theta}(x|z_1, \dots, z_L)] - \text{KL}(q_{\phi}(z_1, \dots, z_L|x) \| p_{\theta}(z_1, \dots, z_L))$$

$$= \mathbb{E}_{q_{\phi}(z_1|x)} [\log p_{\theta}(x|z_1)] - \text{KL}(q_{\phi}(z_1, \dots, z_L|x) \| p_{\theta}(z_1, \dots, z_L))$$

$$= \langle \text{Only looking at 2nd Term} \rangle + \int q_{\phi}(z_1, \dots, z_L|x) \log p_{\theta}(z_1, \dots, z_L) / q_{\phi}(z_1, \dots, z_L|x) dz_1, \dots, z_L$$

$$= \int \prod_{j=1}^L q_{\phi}(z_j|x) \log (\prod_{i=1}^L p_{\theta}(z_i|z_{i:n}) / q_{\phi}(z_i|x) dz_1, \dots, z_L$$

$$= \sum_{i=1}^L \int q_{\phi}(z_{i:n}|x) q_{\phi}(z_i|x) \log q_{\phi}(z_i|x) dz_{i:n}$$

$$= \text{KL}(q_{\phi}(z_i|x) \| p_{\theta}(z_i)) - \sum_{i=1}^L \mathbb{E}_{z_{i:n} \sim q_{\phi}(z_{i:n}|x)} [\text{KL}(q_{\phi}(z_i|x) \| p_{\theta}(z_i|z_{i:n}))]$$

β -VAE Learn disentangled repr. Unsup. (distinct factors of variation)
 $p(\tilde{x}|\tilde{z}) \approx p(\tilde{x}|\tilde{v}, \tilde{w})$, Cond. Ind $V \in \mathbb{R}^K$
 Cond. dep, $w \in \mathbb{R}^H$ $\log p(\tilde{v}|\tilde{x}) = \sum_k \log p(V_k|\tilde{x})$
 $z \in \mathbb{R}^M$ $M \geq K$. **Objective**

$$\max_{\theta, \phi} \mathbb{E}_{x \sim D} [\mathbb{E}_{z \sim q_{\phi}(z|x)} [\log p_{\theta}(x|z)]]$$

Subject to $\text{KL}(q_{\phi}(z|x) \| p_{\theta}(z)) < \delta$
 Via KKT: $\tilde{\mathcal{L}}(\theta, \phi, \beta) = \mathbb{E}_{z \sim q_{\phi}(z|x)} [\log p_{\theta}(x|z)] - \beta [\text{KL}(q_{\phi}(z|x) \| p_{\theta}(z)) - \delta]$, $\beta, \delta \geq 0$
 $\geq \mathbb{E}_z [\log p_{\theta}(x|z)] - \beta \text{KL}(q_{\phi}(z|x) \| p_{\theta}(z))$

$$\mathcal{L}_{\text{beta}}(\theta, \phi, \beta) = -\mathbb{E}_z [\log p_{\theta}(x|z)] \quad \beta=1: \text{VAE}$$

$$+ \beta \text{KL}(q_{\phi}(z|x) \| p_{\theta}(z)) \quad \beta>1 \text{ Strong Const. on } z$$

AUTOREGRESSIVE MODELS

Autoregression for timeseries $x_t = b_0 + b_1 x_{t-1} + b_2 x_{t-2}$ Uses same input variable at previous time steps.

Tabular $p(\tilde{x}) = \prod_{i=1}^n p(x_i | \tilde{x}_{<i})$

$\oplus$ Represents any possible distribution of n RV $\Theta = (\sum_{i=1}^n z_i^{-1}) \in \mathcal{Z}^{(n)}$ param

Independence Assump. $p(\tilde{x}) = p(x_1) \dots p(x_n)$

$\oplus$ n params Θ "Random Samly unrelated"

Conditionals w/ fixed number of params

$\forall$ position i ; learn $f_i: \{0,1\}^{i-1} \rightarrow [0,1]$
 Params: $\sum_{i=1}^n | \partial_i |$, if we carefully Set dimensions $\forall \{ \partial_i \}_{i=1}^n$ controllable params.
Fully Visible Sigmoid Belief Network
 $f_i(x_1, \dots, x_{i-1}) = \sigma(\alpha_0^i + \dots + \alpha_{i-1}^i x_{i-1}) \leq \mathbb{R}$
 # Params = $\sum_{i=1}^n i = (n^2+n)/2$, Stack layers $\rightarrow$

Neural Autoregressive Density Estimator
 AE-like NN to learn $p(x_i=1|\tilde{x}_{<i})$
 $\tilde{h}_i = \sigma(\tilde{b} + W[\cdot, i] \tilde{x}_{<i})$ first i columns
 $\hat{x}_i = \sigma(C_i + \tilde{V}[\cdot, i] \tilde{h}_i)$ i th row
 Each Conditional Modeled by same NN
Train maximizing $\frac{1}{n} \sum_{t=1}^n \log(p(\tilde{x}_t))$
 $= \frac{1}{n} \sum_{t=1}^n \sum_{i=1}^N \log p(x_i^{(t)} | \tilde{x}_{<i}^{(t)})$

Ground Truth are used for conditioning
 $\oplus$ Complete O(1D), could make use of 2nd order optimizers, extendable to other observ.

Masked AE Distribution Estimator
 Constrain AE s.t. Outputs can be used as conditionals $p(x_i | \tilde{x}_i)$ $\oplus$ Training has same complexity as AE, Criteria is NLLK for bin $\tilde{x}$, $p(\tilde{x})$ is just fwd.
 (but Sampling D fwd, very large h_t needed.)
Complexity at inference larger than AE)

PixelCNN Masked Convolutions for Conditionals
Train Maximizes likelihood of train imgs
 $p(\tilde{x}) = \prod_{i=1}^n p(x_i | x_1, \dots, x_{i-1})$

Predicts $\forall$ pixels, a distribution $0 \rightarrow 255$

Masking $p(x_i | \tilde{x}_{<i}) = p(x_{i,R} | \tilde{x}_{<i}) \cdot p(x_{i,B} | x_{i,R}, \tilde{x}_{<i})$

Stacking layers of masked convolutions
 Creates a blindspot $\rightarrow$ Use 2 stacks of convolution [vertical, horizontal stack]

WaveNet Use dilated convolution to increase receptive field without increasing num. parameters for Audio.

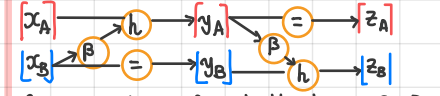
Self-Attention $\tilde{x}_t = f(\tilde{x}_{t-1}, \tilde{x}_{t-2}, \dots, \tilde{x}_1)$
 $W_K, W_Q, W_V \in \mathbb{R}^{D \times D}$ $\tilde{x} \in \mathbb{R}^D$ (learnable)
 $K = XW_K, V = XW_V, Q = XW_Q$
 $\tilde{\alpha} = \text{Softmax}(QK^T/\sqrt{D}) \in \mathbb{R}^{1 \times T}$
 $X = \text{Softmax}(QK^T/\sqrt{D} + M) V$

Transformer MHSA Split attention into Multiple heads s.t. each head calcul. a chunk of the representation.
P.E Inject sinusoidal encodings Unique to indices to preserve input ordering. $= \sin(w_k \cdot t)$ if $i = 2k$
 $\cos(w_k \cdot t)$ d.w. $w_k = 10000^{-2k/d}$
Encoder is Autoregressive

Complexity	Per layer	seq ops	max path len
Self-Attention	$O(n^2 d)$	$O(1)$	$O(1)$
Recurrent	$O(nd^2)$	$O(n)$	$O(n)$
Convolutional	$O(knd^2)$	$O(1)$	$O(\log_k n)$
SA Restricted	$O(rnd)$	$O(1)$	$O(n/r)$

VRNN $p_{\theta}(\tilde{z}) = \prod_{t=1}^T p_{\theta}(z_t | \tilde{z}_{<t}, \tilde{x}_{<t})$
 $q_{\phi}(\tilde{z} | \tilde{x}) = \prod_{t=1}^T q_{\phi}(z_t | x_t, \tilde{x}_{<t}, \tilde{z}_{<t})$
 $h_t = f_{\theta}(\tilde{\alpha}_t, \tilde{z}_t, h_{t-1})$, $p_{\theta}(\tilde{x}_t | \tilde{z}_t) = p_{\theta}(\tilde{x}_t | \tilde{z}_t) p_{\theta}(\tilde{z}_t)$
 $= \prod_{t=1}^T p_{\theta}(z_t | \tilde{z}_{<t}, \tilde{x}_{<t}) p_{\theta}(x_t | z_t, \tilde{z}_{<t}, \tilde{x}_{<t})$
 $\mathcal{L}(\theta, \phi; x) = \mathbb{E}_{q_{\phi}(z_1, x_1, \tilde{z}_1, x_2, \tilde{z}_2)} [\log p_{\theta}(x_1 | z_1, \tilde{z}_1, x_2)] - \text{KL}(q_{\phi}(z_1 | x_1, \tilde{z}_1, x_2) \| p_{\theta}(z_1 | z_1, x_2))$

NORMALIZING FLOWS



f : Invertible, differentiable, dim preserving
 $f_0: \mathbb{R}^d \rightarrow \mathbb{R}^d, X = f_0(Z), Z = f_0^{-1}(X)$

Model $P_X(\tilde{x}, \theta) = P_Z(f^{-1}(\tilde{x})) | \det \frac{\partial f^{-1}(\tilde{x})}{\partial \tilde{x}} |$

Training Maximize exact Log-Likelihood
 $\log P_X(D) = \sum_{x \in D} (\log P_Z(f^{-1}(x)) + \sum_k \log |\det \frac{\partial f^{-1}(x)}{\partial x}|)$

Inference to generate $\tilde{x}$: $\tilde{z} \sim p_{\tilde{z}} \Rightarrow \tilde{x} = f(\tilde{z})$
 Probability of $\tilde{x}$: $\tilde{z} = f^{-1}(\tilde{x}) \Rightarrow p_{\tilde{z}}(\tilde{z})$

Planar NF $\tilde{x} = f(\tilde{z}) = \tilde{z} + \tilde{u} h(\tilde{w}^T \tilde{z} + \tilde{b})$
 $p_z(f^{-1}(\tilde{x})) = p_z(\tilde{z}) \propto \exp(-\tilde{z}^2/2)$
 $\det(\frac{\partial f(\tilde{z})}{\partial \tilde{z}}) = \det(I + h' \tilde{u} \tilde{u}^T) = (1 + h' \tilde{u}^T \tilde{u}) \Rightarrow \log p(\tilde{x}) = -\frac{\tilde{z}^2}{2} - \log(1 + h' \tilde{u}^T \tilde{u})$

Conditional Coupling NF
 $p(\tilde{x}|\tilde{c}) = p(\tilde{z}|\tilde{c}) |\det(\frac{\partial f(\tilde{z})}{\partial \tilde{z}})|^{-1}$
 $f(\tilde{z}) = \begin{bmatrix} \tilde{z}_0 \\ \tilde{z}_1 \end{bmatrix} = \begin{bmatrix} \tilde{z}_0 \\ \gamma(\tilde{z}_0) + \tilde{z}_1 \odot \sigma(\tilde{A} \tilde{z}_0 + \tilde{b}) \end{bmatrix}$
 $\tilde{z} = [\tilde{z}_0, \tilde{z}_1]$, $\tilde{x} = [\tilde{x}_0, \tilde{x}_1]$ $\tilde{f} = [f_1, f_2]$
 $\frac{\partial f(\tilde{z})}{\partial \tilde{z}} = J_{\tilde{f}}(\tilde{z}) = \begin{bmatrix} \partial f_1(\tilde{z})/\partial \tilde{z}_0 & \partial f_1(\tilde{z})/\partial \tilde{z}_1 \\ \partial f_2(\tilde{z})/\partial \tilde{z}_0 & \partial f_2(\tilde{z})/\partial \tilde{z}_1 \end{bmatrix}$
 $\det(\frac{\partial f(\tilde{z})}{\partial \tilde{z}}) = \prod_{j=1}^D \sigma(\tilde{A} \tilde{z}_0 + \tilde{b})_j$

Squeeze Reduce Spatial Dim, preserves dimensionality. One per level. $4 \times 4 \times 1 \Rightarrow 2 \times 2 \times 4$

Step Off Flow 1. Act+Norm 2. Invertible 1x1 Convolution 3. C. Coupling layer

1. Data-dependent Initialization
 $\hookrightarrow$ Run one iter, init layer $\sim N(0, I)$

Fwd: $\forall i, j$, $\tilde{y}_{i,j} = \tilde{z} \odot \tilde{x}_{i,j} + \tilde{b}$
 Rev: $\forall i, j$, $\tilde{x}_{i,j} = (\tilde{y}_{i,j} - \tilde{b}) / \tilde{z}$
 log-det: $H \cdot W = \sum (\log |\tilde{z}|)$

2. Permutation in the Channel dimension
 $\tilde{H} \in \mathbb{R}^{H \times W \times C}$ $\tilde{W}$ Orthogonal $\det(\tilde{W}) = 1 \in \mathbb{R}^{C \times C}$

Compute in $O(C^3)$ $\log |\det(\frac{\partial f}{\partial h} \text{conv}(h; w))| = H \cdot W \cdot \log |\det(W)|$ Parameterize $W = PL(U + \text{diag}(\tilde{z}))$ $L: \text{diag} = 1$ $U: \text{diag} = 0$
 in $O(C)$ $\log |\det(W)| = \sum (\log(\tilde{z}))$

SRFlow Conditioning on low-res image
 1, 2, Affine Injector A pre-trained CNN encodes a low-res img $\tilde{u} = g_{\theta}(\tilde{x})$

Activation map $\tilde{h}_n$ modulated with Spatial feature maps $\tilde{u}$ $\forall$ channels, locations
 Fwd: $h^{(n+1)} = \exp(\beta_{\theta, S}^n(\tilde{u})) \cdot h^{(n)} + \beta_{\theta, b}^n(u)$
 Inv: $h^{(n)} = \exp(-\beta_{\theta, S}^n(\tilde{u})) (h^{(n+1)} - \beta_{\theta, b}^n(u))$
 log-Det: $\sum_{i,j,k} \beta_{\theta, S}^n(u)_{i,j,k}$

Style Flow NF for Disentanglement
C-Flow Two Flows, One Conditioned on other
 $\rightarrow$ Multimodal img-img mapping, style transfer, img manip, 3D point cloud recon.
 Flow B: Compute transform. params. | Flow A

GENERATIVE ADVERSARIAL NETWORKS
 $\mathcal{L}(D) = -\frac{1}{2N} (\sum_{i=1}^N \log D(x^{(i)}) + \sum_{n=1}^N \log(1 - D(y^{(n)})))$ via BCE loss

$D^* = \arg \min_D -\frac{1}{2} (\mathbb{E}_{x \sim p_d} [\log D(x)] + \mathbb{E}_{z \sim p_z} [\log(1 - D(G(z)))])$ for a fixed G
Objective $m_{\min} \max V(D, G) \downarrow$ highlighted
 $D^*(x) = \arg \max_x \mathbb{E}_{x \sim p_d} [\log D(x)] + \mathbb{E}_{z \sim p_z} [\log(1 - D(G(z)))] = \int_x p_d(x) \log(D(x)) dx + \int_z p_z(z) \log(1 - D(G(z))) dz$
 $= \int_x p_d(x) \log(D(x)) + p_g(x) \log(1 - D(x)) dx$
 $= p_d(x) / (p_d(x) + p_g(x))$ [optimal $D = 0.5$]

$V(G, D^*) = \mathbb{E}_{x \sim p_d} [\log \frac{2 p_d(x)}{p_d(x) + p_g(x)}] + \mathbb{E}_{z \sim p_z} [\log \frac{2 p_g(x)}{p_d(x) + p_g(x)}]$
 $= \text{KL}(p_d(x) \| \frac{p_d(x) + p_g(x)}{2}) + \text{KL}(p_g(x) \| \frac{p_d(x) + p_g(x)}{2})$
 $-\log 4 = 2JS(p_d(x) \| p_m(x)) - \log 4$

$G^* = \arg \min_G V(G, D^*) \wedge \forall x: JS(x) \geq 0 = -\log(4)$

Convergence G, D enough Capacity $\wedge$ at each Step D allowed to reach D^* , and p_g is updated to improve $\sup_x p_g(x) \log(1 - D(x)) dx$
 $\Rightarrow p_g$ converges to p_d (directly opt p_g not θ)

Training Alternate between Gradient-Ascent on D $\max_{\theta} \mathbb{E}_{x \sim p_d(x)} [\log D_{\theta}(x)]$ + $\mathbb{E}_{z \sim p_z} [\log(1 - D_{\theta}(G_{\theta}(z)))]$

and Gradient Ascent on G $\max_{\theta} \mathbb{E}_{z \sim p_z} [\log(D_{\theta}(G_{\theta}(z)))] \neq \log(1 - D(G(z)))$

Saturates in early stages.

Mode Collapse $G \rightarrow$ High Quality, Low Variability

Sol: Unrolled GANs: After k updates, G_k is optimized w.r.t. D after next k steps.

$\hookrightarrow$ Discourages G to exploit local min.

JS-Issues Correlates badly with Sample Quality (don't know when to stop train)
 D tends to be optimal, JS saturates (When no overlap Supports) → Bad Gradients
 Sol: Wasserstein Distance

Gradient Penalty w.r.t gradients of D
 $\lambda ||\nabla_x D(x)||^2$ Stabilizes GAN training

DGGAN Replace pooling layers w/ Strided Conv [D], fractional-strided Conv [G]
 batch Norm [D,G], Remove fully connected h for deeper architectures, RelU expct Tanh on Last layer [G], Leaky RelU w/ L [D]

Men Hat - Men Not Hat + Women No Hat = Women With Hat
StyleGAN pass $\tilde{z}$ through a series of FCNN mapping $\tilde{z} \rightarrow \tilde{w}$, insert $\tilde{w}$ multiple times w/ (actnorm). Learn params in first layer
 layerwise style and noise injection

$P_{ix} 2P_{ix} \mathcal{L}(G,D) = \mathcal{L}_{GAN}(G,D) + \lambda \mathcal{L}_{L1}(G)$ where
 $\mathcal{L}_{GAN}(G,D) = E_{x,y} [\log D(x,y)] + E_{x,z} [1 - \log D(x,g(x,z))]$
 $\mathcal{L}_{L1}(G) = E_{x,y,z} [||y - g(x,z)||_1]$ Img Translation

CycleGAN Unpaired Image Translation
 $\mathcal{L}(G,F,P_x,D_y) = \mathcal{L}_{GAN}(G,D_y,x,y) + \mathcal{L}_{GAN}(F,D_x,y,X) + \lambda \mathcal{L}_{CYC}(G,F)$ where $\mathcal{L}_{CYC} = E_{x,y,p} [||F(G(x)) - x||_1] + E_{y,x,p} [||G(F(y)) - y||_1]$

DIFFUSION MODELS $x_T \sim N(0,1)$
 real img $x_0 \sim q(x_0)$; $q(x_0)$ original data distr.

Diffusion Step $q(x_t|x_{t-1}) = N(x_t; \sqrt{1-\beta_t}x_{t-1} + \beta_t \epsilon)$
 $0 < \beta_1 < \dots < \beta_T < 1$ $x_t = \sqrt{1-\beta_t}x_{t-1} + \sqrt{\beta_t}\epsilon$ Iterative

Direct Let $\alpha_t = 1 - \beta_t$; $\bar{\alpha}_t = \prod_{s=1}^t \alpha_s$
 $x_t = \sqrt{\alpha_t}x_0 + \sqrt{1-\alpha_t}\epsilon$ $\epsilon = \sqrt{\alpha_t}x_0 + \sqrt{1-\alpha_t}\epsilon$
 $\therefore \sqrt{\alpha_t}x_0 + \sqrt{1-\alpha_t}\epsilon$ $q(x_t|x_0) = N(\sqrt{\alpha_t}x_0, (1-\alpha_t)\mathbf{I})$

Linear scheduler img noise too quick → Use cosine
Denoising Step $q_\theta(x_t) = \int q(x_t|x_0)q_\theta(x_0)dx_0$ Intract!
 $p_\theta(x_{t-1}|x_t) = N(x_{t-1}; \mu_\theta(x_t, t), \sigma_\theta^2(x_t, t)) \approx q(x_{t-1}|x_t, x_0)$

Objective $\log p(x) = \log p(x_{0:T}) d x_{1:T}$
 $= \log p(x_{0:T}) \frac{q(x_{1:T}|x_0)}{q(x_{1:T}|x_0)} dx_{1:T} = \log E_q \left[\frac{p(x_{0:T})}{q(x_{1:T}|x_0)} \right]$
 $\geq E_q \left[\log \frac{p(x_{0:T})}{q(x_{1:T}|x_0)} \right] = ELBO = E_q [\log p_\theta(x_{0:T})] - KL(q(x_{1:T}|x_0) || p(x_{1:T}))$
 $\argmin_{\theta} \mathcal{L}_t = \argmin_{\theta} KL(q(x_{t+1}|x_t, x_0) || p_\theta(x_{t+1}|x_t))$
 $= \argmin_{\theta} KL(N(\mu_\theta, \Sigma_\theta(t)) || N(\mu_0, \Sigma_0(t)))$
 $= \argmin_{\theta} \frac{1}{2} [\log \frac{|\Sigma_\theta(t)|}{|\Sigma_0(t)|} - d + \text{tr}(\Sigma_0(t)^{-1} \Sigma_\theta(t)) + (\mu_0 - \mu_\theta)^T \Sigma_0(t)^{-1} (\mu_0 - \mu_\theta)]$
 $\mu_\theta(x_t, x_0) = \frac{1}{\sqrt{\alpha_t}} x_t - \frac{1-\alpha_t}{\sqrt{1-\alpha_t}} \epsilon_0$
 $\Sigma_\theta(x_t, x_0) = \frac{1}{\alpha_t} x_t - \frac{1-\alpha_t}{\sqrt{1-\alpha_t}} \epsilon_0 \hat{\Sigma}_\theta(x_t, t) \mathbf{NN}$

(1) Learns to predict total added noise $\epsilon_0 \sim N(0, \mathbf{I})$
 that was used from x_0 to x_t . Pred ϵ_0 better x_0 or μ_θ
 $||\epsilon - \epsilon_0(x_t, t)||^2 - ||\epsilon - \epsilon_0(\sqrt{\alpha_t}x_0 + \sqrt{1-\alpha_t}\epsilon, t)||^2$

Train while τ Conv. $[x_0, q(x_0)]$; $t \sim \mathcal{U}\{1, \dots, T\}$
 $t \sim N(0, \mathbf{I})$; GD step $\nabla_{\theta} ||\epsilon - \epsilon_0(\sqrt{\alpha_t}x_0 + \sqrt{1-\alpha_t}\epsilon, t)||^2$

Sampling $x_t \sim N(0, \mathbf{I})$ for $(t=T, \dots, 1)$ $[z \sim N(0, \mathbf{I})]$
 if $t > 1$ else 0; Given $\beta_t = \beta_t^z$: Predicted Noise
 $x_{t-1} = \frac{1}{\sqrt{\alpha_t}} (x_t - \frac{1-\alpha_t}{\sqrt{1-\alpha_t}} \epsilon_0(x_t, t)) + \beta_t z$

Reverse Distribution $q(x_{t-1}|x_t, x_0) \propto$ Bayes
 $q(x_t|x_{t-1}, x_0) q(x_{t-1}|x_0) \propto$
 $\exp \left\{ -\frac{(x_t - 2x_{t-1}\sqrt{\alpha_t} - \beta_t x_{t-1} + (1-\beta_t)x_t)^2}{2\beta_t^2} - \frac{(x_{t-1} - 2x_{t-1}\sqrt{\alpha_t}x_0 + \bar{\alpha}_t x_0^2)/2(1-\alpha_t)}{\sigma_{x_t}^2} \right\}$
 $\sigma_{x_t}^2 = (1-\beta_t)/\beta_t + 1/(1-\alpha_{t-1})$
 $\mu_{x_t} = \sqrt{1-\beta_t} x_{t-1} \bar{\sigma}_t^2 / \beta_t + \sqrt{\alpha_t} x_0 \bar{\sigma}_t^2 / (1-\alpha_t)$
 Since $q(x_{t-1}|x_t, x_0) \sim N(x_{t-1}; \mu_{x_t}^2, \sigma_{x_t}^2) \propto$
 $\exp \left\{ \frac{(x_{t-1} - \mu_{x_t})^2}{2\sigma_{x_t}^2} \right\}$
 $= \exp \left\{ \frac{\alpha_t^2 - 2\alpha_t \bar{\alpha}_t + \bar{\alpha}_t^2}{2\sigma_{x_t}^2} \right\}$

NEURAL IMPLICIT REPRESENTATIONS

Voxels $O(n^3)$ Memory, Resolution Limited
Points discretize surface → 3D pts, No connectivity
Meshes discretize → Vertices, faces, approx error
Implicit functions Learn analytic f which represents the 3D-surface. NO approx error
 ⇒ Smooth Continuous. No graphical visual

hard to obtain high frequency details
Level-Set of a continuous function
 to repr. Surfaces $f(x) = x_1^2 + x_2^2 + x_3^2 - r^2$
 $S = \{x | f(x) = 0\}$ (pnb inside S ↓)

Occupancy Networks $f_\theta: \mathbb{R}^3 \times \mathcal{X} \rightarrow [0,1]$
DeepSOF $f_\theta: \mathbb{R}^3 \times \mathcal{X} \rightarrow \mathbb{R}$ [-in, +out]

NIR Implicit ⇒ Arbitrary topology, resolu.
 Low Memory footprint. Learning from:

Watertight Meshes 1. Randomly Sample 3D points in space 2. Query GT occupancy Or SDF from GT meshes 3. Train with CE.

Point clouds (Given) $\mathcal{X} = \{x_i\}_{i \in \mathcal{I}} \subset \mathbb{R}^3$
 Compute $\theta, f_\theta(x)$ SDF of $\mathcal{M}$, defined by $\mathcal{X}$ w/o any additional supervised data.

Weak Supervision Want $\mathcal{L}$ vanish train pts
 $\mathcal{L}(\theta) = \sum_{i \in \mathcal{I}} |f_\theta(x_i)|^2 + \lambda E_{x'} [||\nabla_x f_\theta(x')||^2]$

② Want Spatial gradient = 1 ⇒ interpret as geometric surface, encourages smoothness

Ray Marching FIRST zero-crossing of SDF
Secant Method $x_{n-1} - x_{n-2}$

DVR fwd $x_n = x_{n-1} + f(x_{n-1}) - f(x_{n-2})$
 $t_0(\hat{p})$ $f_\theta = \tau$ Vpixels u, find $\hat{p}$ along $\vec{w}$
 Evaluate texture field $t_0(\hat{p})$
 $\vec{r}_0$ Insert Colour $t_0(\hat{p})$ at pixel u

Bwd $\hat{I}$: pred. $\mathcal{L} = \sum_u ||\hat{I}_u - I_u||$
 $\frac{\partial \mathcal{L}}{\partial \theta} = \frac{\partial \mathcal{L}}{\partial I_u} \frac{\partial I_u}{\partial \theta}$
 $\frac{\partial \mathcal{L}}{\partial \theta} = \frac{\partial \mathcal{L}}{\partial \theta} + \frac{\partial \mathcal{L}}{\partial \hat{p}} \cdot \frac{\partial \hat{p}}{\partial \theta}$

$f_\theta(\hat{p}) = \tau$ with $\hat{p} = r_0 + \hat{d} w$
 $\frac{\partial f_\theta(\hat{p})}{\partial \theta} + \frac{\partial f_\theta(\hat{p})}{\partial \hat{p}} \frac{\partial \hat{p}}{\partial \theta} = 0 \Rightarrow w \frac{\partial \hat{p}}{\partial \theta}$
 Solve for $\frac{\partial \hat{p}}{\partial \theta}$, then multiply w back.
 $\frac{\partial \hat{p}}{\partial \theta} = -w \left(\frac{\partial f_\theta(\hat{p})}{\partial \hat{p}} \right)^{-1} \frac{\partial f_\theta(\hat{p})}{\partial \theta}$ Analytic
 No storing

Neural Radiancefield

Can model transparency → worse geometry
Input x, y, z, θ, ϕ (Camera parameters)

Output σ, c (density, RGB value of point)
 $x, y, z \rightarrow 256 \rightarrow \uparrow \rightarrow 256 \rightarrow 128 \rightarrow \text{RGB}$
 (prevent washout) x, y, z θ, ϕ don't want

σ to depend on θ, ϕ just x, y, z
Alpha Compositing $\delta_i = t_i + (1-t_i)\delta_{i-1}$

$\alpha = 1 - \exp(-\sum_i \delta_i)$
 Transmittance $T_i = \prod_{j=1}^{i-1} (1 - \alpha_j)$ $C = \sum_{i=1}^N T_i \alpha_i C_i$ weight

Treat w as probability for new samples
All points used for Volumetric rendering

Train $\min_{\theta} \sum_i ||\text{render}_i(\theta) - I_i||^2$
 Use PE to map continuous input coordinates into higher-dim space ⇒ MLP easily approx higher frequency function.

③ 50+ Calibrated views needed, slow rendering speed for high-res imgs, Only models static scenes.

Sparse Views ⇒ Regularize depth by Smoothness, image by likelihood.

FastNerf Replace MLP w/ feature grid + shallow MLP, represent grid w/ hash table for memory efficiency

Gaussian Splatting hard to model thin structures w/ isotropic shapes (sphere)

→ Non-isotropic parameterization (ellipsoid)
 3D Gaussians to 2D space Gaussians (splats)

1. Vpixels (approx by tile-based Sorting, sort all Gaussians according to their depth.

2. Calculate opacity of each Gaussian
 $\alpha = 0 \cdot \exp(-0.5(x-\mu)^T \Sigma^{-1}(x-\mu))$

3. Use Volume Rendering Equation
 $C = T_1 \alpha_1 C_1 + T_2 \alpha_2 C_2 + T_3 \alpha_3 C_3$ $T_1 = 1$
 $T_2 = 1 - \alpha_1$ $T_3 = (1 - \alpha_2)(1 - \alpha_1)$ $T_i = \prod_{j=1}^{i-1} (1 - \alpha_j)$

Parameterized via Spherical harmonics
 $Y_0(\theta, \phi) = \frac{1}{\sqrt{4\pi}}$ $Y_1(\theta, \phi) = \sqrt{\frac{3}{4\pi}} \sin \theta \cos \theta$
 $Y_2(\theta, \phi) = \sqrt{\frac{15}{4\pi}} \cos^2 \theta$ $Y_3(\theta, \phi) = \sqrt{\frac{35}{4\pi}} \cos^3 \theta$

$x = r \sin \theta \sin \phi$ $y = r \cos \theta \sin \phi$ $z = r \cos \theta$
 When $\theta = \pi/2$, $z=0$ [3D → 2D] dense as a

Y_0 Constant $r \rightarrow$ Circle w/ 0 center
 $Y_1 = \sqrt{\frac{3}{4\pi}} \frac{x}{r}$ $Y_2 = \sqrt{\frac{5}{4\pi}} \frac{z}{r}$ $Y_3 = \sqrt{\frac{35}{4\pi}} \frac{y}{r}$

Using absolute val of Y_m as new radius
 $x^2 + y^2 + z^2 = a^2 x^2 / r^2$ $x^2 + y^2 + z^2 = a^2 y^2 / r^2$
 $x^2 + y^2 = |a x|$ and $x^2 + y^2 = |a y|$

Circle center at $|a/2|$ $\forall x > 0 - |a/2| < 0$

The Y_m define spheres in 3D, where Y_0 aligned with the z-axis.

PARAMETRIC BODY MODELS

Pictorial Structure Model
 $S(I, L) = \sum_{i \in V} \alpha_i \cdot \phi(I, l_i) + \sum_{i, j \in E} \beta_{ij} \cdot \psi(l_i, l_j)$

① Unary term: how likely on image I will body joint i be located at l_i

② Pairwise term: likelihood body joint i is linked with joint j

Direct Regression Based on deep CNN to directly regress x, y coordinates involving a refiner

Heatmaps Improve human pose estimation with occlusions.

Each Keypoint has a separate binary heatmap (option of $N(\mu, \sigma)$ around it)

2D → 3D Lift directly, no restrictions 2mArm

SMPL Represent via 3D mesh, $V \sim 7000$
 Body defined by Shape and Pose ①

Shape PCA of meshes in canonical pose to estimate directions of maximum shape variation and obtain a low-dimensional subspace (10D-300D) in the canonical pose.

Pose Linear Blend Skinning (4)

Each vertex t_i in rest position is transformed $t_i' = \sum_{k=1}^K w_{ki} G_k(\theta, J) t_i$
 w_{ki} : Blend skinning weights by artist
 G_k : Rigid bone transform
 J : joint locations θ : Pose

Posed vertices are linear combination of transformed template vertices.

⊖ Produces well known artifacts

↳ Augment base shape: Pose blend shape

is a vector of vertex displacements in rest pose.

SIMPL: Learn from data:

$$t_i' = \sum_k w_{ki} G_k(\theta, J(\beta)) (t_i + S_i(\beta) + P_i(\theta))$$

Joints $J(\beta)$ depends on shape β (PCA)

Shape correctives $s(\beta) = s\beta$ can be

Pose correctives $p(\theta) = p\theta$ NN

$$B_s(\beta, S) = \sum_{n=1}^{15} \beta_n S_n$$

$\beta = [\beta_1, \dots, \beta_{15}]^T$ linear shape coefficients

$S_n \in \mathbb{R}^{3N}$ orthonormal PC of shape displacements

$$B_p(\theta, P) = \sum_{n=1}^{9K} (R_n(\theta) - R_n(\theta^*)) P_n$$

23 Joints, $K=3$, $R(\theta)$ length 23×9

$P_n \in \mathbb{R}^{3N}$ vector of vertex displacements

$P = [P_1, \dots, P_{9K}] \in \mathbb{R}^{3N \times 9K}$ all 207 Bp

$M(\vec{\beta}, \vec{\theta}; \vec{\Phi})$ Mesh $J(\vec{\beta}; J, \vec{\tau}, S)$

Update Rule $\frac{\partial L_{reproj}}{\partial \theta} + \frac{\partial L_{prior}}{\partial \theta}$

$$\theta^{t+1} \leftarrow \theta^t - \eta \left(\frac{\partial L_{reproj}}{\partial \theta} + \frac{\partial L_{prior}}{\partial \theta} \right)$$

⊖ Hand-crafted Optimization Routine

⊖ Sensitive to initialization

⊖ Slow convergence

Learned Gradient Descent

$$\theta^{t+1} = \theta^t + F(\frac{\partial L_{reproj}}{\partial \theta}, \theta^t, x)$$

Reconstruction Human Detail hard-poses

Regression-based ✓

Template-based ✓ Needs Temp. ✓

Animatable Neural implicit Surfaces

1. Input Posed 3D Meshes
2. Learned Continuous Canonical Shape and Skinning weights
3. Continuous Implicit Surfaces $\approx$ Unseen poses.

REINFORCEMENT LEARNING

Markov Property $P(S_{t+1}|S_t) = P(S_{t+1}|S_1, \dots, S_t)$

Markov Chain $\langle S, P \rangle$ $P_{SS'} = P[S_t = S' | S_t = S]$

Markov Reward Process $\langle S, P, R, \gamma \rangle$

$R_t = E[R_{t+1} | S_t = S]$ $\gamma \in (0, 1)$

Return G_t total discounted reward from t .

$$G_t = R_{t+1} + \gamma R_{t+2} + \dots = \sum_{k=0}^{\infty} \gamma^k R_{t+k}$$

Markov Decision Process $\langle S, A, P, R, \gamma \rangle$

Bellman $V_{\pi}(s) = E_{\pi}[G_t | S_t = s]$

$$= E_{\pi}[R_{t+1} + \gamma G_{t+1} | S_t = s] = \sum_a \pi(a|s) [R(s,a) + \gamma V_{\pi}(s')]$$

$\sum_s \sum_a p(s', r | s, a) [r + \gamma V_{\pi}(s')]$

Non-linear, No closed form

Action-Value $Q_{\pi}(s, a) = E_{\pi}[G_t | S_t = s, A_t = a]$

$$= E_{\pi}[R_{t+1} + \gamma Q_{\pi}(S_{t+1}, A_{t+1}) | S_t = s, A_t = a]$$

Bellman Optimality $V_*(s) = \max_a V_{\pi}(s)$

$$= \max_a Q_*(s, a) = \max_a \sum_{s', r} p(s', r | s, a) (r + \gamma V_*(s'))$$

$\pi^*(s) = \arg \max_{a \in A} (r(s, a) + \gamma V_{\pi^*}(P(s, a)))$

$V_{\pi^*}(s) \geq V_{\pi}(s)$

Value Iter $V_{new}^*(s) = \max_{a \in A} \sum_{s', r} p(s', r | s, a) (r + \gamma V_{old}^*(s'))$

Policy Iter. Initialize π_0

Until Convergence: iterates while $S \neq$

Compute $V_{\pi}(x)$ $\forall x$

Compute Greedy Policy π_G wrt V_{π}

Set $\pi \leftarrow \pi_G$

Finds exact solution in polynomial # iterations! $\pi^* \in O^*(\frac{1}{1-\gamma})$

Guaranteed to monotonically improve

⊖ needs to know transition matrix

Monte Carlo Sampling

Run episodes w/ given π and compute samples of $\sum_{i=0}^{\infty} \gamma^i r(S_{t+i}, a_{t+i})$ of $V_{\pi}(S_t)$
 ⊕ Unbiased estimate, no need system dynam
 ⊖ high Variance, explore/exploit dilemma need termination state (slow for long eps)

MODEL-FREE RL

On-Policy Compute Q according to π and follows π

Off-Policy Compute Q according to GREEDY π , follows different exploration

TD Learning $\Delta V(s) = r(s, a) + \gamma V(s') - V(s)$

$V(s) \leftarrow V(s) + \alpha \Delta V(s)$ on-policy

follows π to obtain (x, a, r, x')

SARSA $\Delta Q(S, A) = R + \gamma Q(S, A') - Q(S, A)$

$Q(S, A) \leftarrow Q(S, A) + \alpha \Delta Q(S, A)$

on-policy Use current policy to get the rollout of states, actions, rewards

Q-learning off-Policy policy to update Q different π to collect data ϵ

$$\Delta Q(S, A) = R_{t+1} + \gamma \max_{a'} Q(S', a') - Q(S, A)$$

$Q(S, A) \leftarrow Q(S, A) + \alpha \Delta Q(S, A)$

⊕ Less Variance than MC due bootstrap

⊕ More Sample efficient than D.P.

⊕ Don't need to know transition M

⊖ Biased due to bootstrapping, using old estimates as labels

⊖ Explore, exploit dilemma

DEEP REINFORCEMENT LEARNING

Deep Q-Learning NN to learn (s, a)

$$L(\theta) = (R + \gamma \max_{a'} Q_{\theta}(s', a') - Q_{\theta}(s, a))^2$$

Replay Buffer to store generated samples, during updates randomly sample transitions from it. Q-Learning is off-policy

So we can use old samples.

⊖ Limited to discrete action spaces

Policy Gradients

$\pi(a_t | s_t) = \mathcal{N}(a_t | \mu_t, \sigma_t^2)$
 Collect rollout trajectories (explore)
 $p(\tau) = p(s_1, a_1, \dots, s_T, a_T) = p(s_1) \prod_{t=1}^{T-1} \pi(a_t | s_t) p(s_{t+1} | a_t, s_t)$
 Samples action $\forall t$ from $\pi(a_t | s_t)$ on-policy
 Evaluation $J(\theta) = E_{\tau \sim p_{\theta}(\tau)} [\sum_{t=0}^T \gamma^t r(S_t, A_t)]$
 Policy Update Goal: $\theta^* = \arg \max_{\theta} J(\theta)$
 $\theta \leftarrow \theta + \nabla_{\theta} J(\theta)$ where $J(\theta) = E_{\tau} [r(\tau)] = \int p(\tau) r(\tau) d\tau$ so $\nabla_{\theta} J(\theta) = \int \nabla_{\theta} p(\tau) r(\tau) d\tau = \int p(\tau) \nabla_{\theta} \log p(\tau) r(\tau) d\tau$ w/ $\nabla \log f = \frac{\nabla f}{f}$
 $= E_{\tau} [\nabla_{\theta} \log p(\tau) r(\tau)]$ $-b(s; \xi)$
 $= E_{\tau} [\nabla_{\theta} \sum_{t=0}^T \log \pi_{\theta}(a_t | s_t) \sum_{t=0}^T \gamma^t r(s_t, a_t)]$
 REINFORCE to compute E in practice
 $\theta \leftarrow \theta + \eta \sum_t \gamma^t G_t \nabla \log \pi_{\theta}(A_t | x_t; \theta)$
 Unbiased high variance $\rightarrow$ Baseline
 ACTOR-CRITIC use bootstrapping to est. r.
 $\nabla_{\theta} J(\theta) = \frac{1}{N} \sum_i \sum_{t=0}^T \nabla_{\theta} \log \pi_{\theta}(a_i^t | s_i^t) (r(s_i^t, a_i^t) + \gamma V(s_{i+1}^t) - V(s_i^t))$ TD-Error

NEURAL NETWORK BASICS

MLE $\theta_{MLE}^* = \arg \max_{\theta} P_m(y | x, \theta)$ ①

$$= \arg \max_{\theta} \prod_{i=1}^N P_m(y_i | x_i, \theta)$$

$$= \arg \max_{\theta} \sum_{i=1}^N \log P_m(y_i | x_i, \theta)$$

Universal Approx Theorem Let $\phi: \mathbb{R} \rightarrow \mathbb{R}$ non-constant, bounded, continuous,

$I_m = [0, 1]^M$, space of real-valued function on $I_m = (I_m) \forall f \in C(I_m)$ cont, $\forall \epsilon > 0$,

$n \in \mathbb{Z}, v_i, b_i \in \mathbb{R}, \tilde{w}_i \in \mathbb{R}^M, \forall i \in \{1, \dots, n\}$

$$f(x) \approx g(x) = \sum_{i=1}^n v_i \phi(w_i^T x + b_i) \wedge |g(x) - f(x)| \leq \epsilon, \forall x \in I_m$$

Softmax $(\vec{z})_i = \exp(z_i) / \sum_{j=1}^K \exp(z_j)$

Activation functions make layer to layer mappings non-linear.

Sigmoid $\sigma(x) = 1 / (1 + e^{-x})$

$$\sigma'(x) = \sigma(x)(1 - \sigma(x))$$

$$\tanh(x) = (e^x - e^{-x}) / (e^x + e^{-x}) = 2\sigma(2x) - 1$$

$$\tanh'(x) = 1 - \tanh^2(x)$$

$$ReLU(x) = \max(0, x)$$

$$ReLU(x) = \begin{cases} x & x < 0 \\ 0 & x \geq 0 \end{cases}$$

Leaky ReLU $x < 0$

Randomized Leaky ReLU $x \geq 0$

$\alpha \sim \text{Unif}(a, b)$ test $\alpha = E[\alpha] = (a+b)/2$

Ensemble train different M on same d

OR same m on different d and aggregate to decrease variance.

Dropout M share θ , train small % of M

Bagging M independent, train all till conv.

$$\text{Batch Norm } z_i^{\text{norm}} = (z_i - \mu) / \sqrt{\epsilon + \sigma^2}$$

$$\tilde{z}_i = \gamma \tilde{z}_i^{\text{norm}} + \beta$$

At testing, M, ϵ

$x \in \mathbb{R}^d, y, b \in \mathbb{R}^C, w \in \mathbb{R}^{C \times d}$ from layer

$h \in \mathbb{R}^d$ to $z \in \mathbb{R}^K (n+1) \cdot K$

PROBABILITY AND MATH

$$\frac{\partial}{\partial x} [a^T x] = \frac{\partial}{\partial x} [x^T a] = a$$

$$\frac{\partial}{\partial x} [B^T A x] = A^T B$$

$$\frac{\partial}{\partial x} [x^T A] = A^T$$

$$\frac{\partial}{\partial x} [x^T \odot a] = \text{diag}(a)$$

$$KL-Div KL(q || p) = \int q(\theta) \log \frac{q(\theta)}{p(\theta)} d\theta$$

$$\hookrightarrow KL(q || p) \geq 0 \quad KL(q || p) \neq KL(p || q)$$

$$JS-Div JS(q || p) = \frac{1}{2} KL(q || M) + \frac{1}{2} KL(p || M)$$

where $M = \frac{q+p}{2}$

Convolutions $A * \text{ROT}_{180} B = A \star B$

$A \star B = \text{ROT}_{180}(B \star A)$; $A \star B$: stride B on A

2D Change of Vars $\iint_R f(x, y) dx dy = \iint_S f(g(u, v)) |J(u, v)| du dv$ Let $x = f(z)$ f inv, cont, diff,

$$dx = \left| \det \left(\frac{\partial f(z)}{\partial z} \right) \right| dz \approx p_{\mathcal{X}}(x) = p_{\mathcal{Z}}(f^{-1}(x)) \left| \det \left(\frac{\partial f(z)}{\partial z} \right) \right|$$

Matrix Det Lemma $\det(A + uv^T) = (1 + v^T A^{-1} u) \det(A)$

Jensen $g(E[x]) \leq E[g(x)]$ $g(\cdot)$ convex

Chain Rule $p(x_{1:T}) = p(x_1) p(x_2 | x_1) \dots p(x_T | x_{1:T-1})$

Product $p(x, y) = p(y | x) p(x)$

Bayes $p(x | y) P(y) = p(y | x) p(x)$

$$p(x_t | x_{t-1}, x_0) = p(x_{t-1} | x_t, x_0) p(x_t | x_0) / p(x_{t-1} | x_0)$$