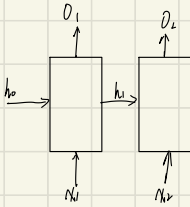




• RNN Auto-regressive Model (Pixel RNN)

Model $p(x^t | x_{1:t-1}; o^t)$

keep a summary of the history and recursively update it



summary update rule: $h_{t+1} = \tanh(W_h h_t + W_o x_{t+1})$

prediction: $o_{t+1} = W_y h_{t+1}$ then pass through a non-linearity

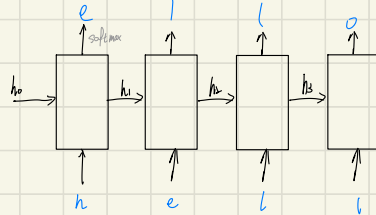
summary initialization: $h_0 = b$

hidden state h_t is a summary of inputs seen till t

o_t specifies parameters for conditional $p(x_t | x_{1:t-1})$

Parameter: b_0, W_h, W_o, W_y constant number of params w.r.t n

Autoregressive: $p(x = \text{hello}) = p(x_1=h) \cdot p(x_2=e | x_1=h)$
 $p(x_1=h | x_1=h, x_2=e) \dots$



pros:

arbitrary length

cons:

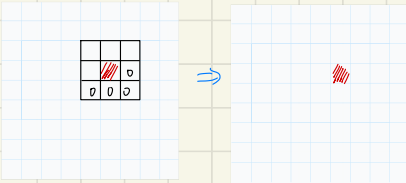
still requires ordering

sequential likelihood function (slow for training)

sequential generation (unavoidable for autoregressive model)

• CNN Auto-regressive Models (Pixel CNN)

use convolutional architecture to predict next pixel given context (a neighborhood of pixels)



mask out the lower-right part of filter
 (a pixel is not conditioned on the pixels 'behind' it)

assume that the domain is governed by some underlying distribution P_{data}

given a dataset D of m samples from P_{data}

assume samples in D is iid

given a family of models $\mathcal{M}$. Want to learn a "good" model $\hat{M} \in \mathcal{M}$ that defines a distribution $P_{\hat{M}}$

the definition of "best model" depends on what we want to do:

density estimation: want the full distribution

prediction: want $P(Y|X)$

structure or knowledge discovery: interested in the model itself

KL-Divergence

$$D(P||Q) = \sum_x P(x) \log \frac{P(x)}{Q(x)} = \mathbb{E}_{P \sim P} \left[\log \frac{P(x)}{Q(x)} \right] \quad \left(\int P(x) \log \frac{P(x)}{Q(x)} dx \right)$$

$$D(P||Q) \geq 0, \quad = 0 \text{ iff } P=Q$$

$$\mathbb{E}_{P \sim P} \left[\log \frac{P(x)}{Q(x)} \right] = -\mathbb{E}_{P \sim P} \left[\log \frac{Q(x)}{P(x)} \right] \geq -\log \mathbb{E}_{P \sim P} \left[\frac{Q(x)}{P(x)} \right] = -\log 1 = 0$$

KL is asymmetric $D(P||Q) \neq D(Q||P)$

$$D(P_{\text{data}} || P_{\theta}) = \mathbb{E}_{P_{\text{data}} \sim P_{\text{data}}} \left[\log \frac{P_{\text{data}}(x)}{P_{\theta}(x)} \right] = \sum_x P_{\text{data}}(x) \log \frac{P_{\text{data}}(x)}{P_{\theta}(x)} = \sum_x P_{\text{data}}(x) \log P_{\text{data}}(x) - P_{\text{data}}(x) \log P_{\theta}(x)$$

$$\arg \min_{\theta} - \sum_x P_{\text{data}}(x) \log P_{\theta}(x) = \arg \max_{\theta} \mathbb{E}_{P_{\text{data}} \sim P_{\text{data}}} \left[\log P_{\theta}(x) \right] \quad \text{minimizing the KL} = \text{maximizing the Expected log-likelihood}$$

Although we can compare models, since we know $H(P_{\text{data}}) = -\sum_x P_{\text{data}}(x) \log P_{\text{data}}(x)$,

we do not explicitly compare $D(P||Q)$, thus we don't know how far we are from optimal (we approximate $\mathbb{E}_{P_{\text{data}}} [\log P_{\theta}(x)]$)

we approximate $\mathbb{E}_{P_{\text{data}}} [\log P_{\theta}(x)]$ with empirical log-likelihood $\mathbb{E}_D [\log P_{\theta}(x)] \stackrel{!}{=} \frac{1}{|D|} \sum_{x \in D} \log P_{\theta}(x)$

Monte Carlo Estimation

$$\mathbb{E}_{P \sim P} [g(x)] \approx \mathbb{E}_D [g(x)] = \frac{1}{|D|} \sum_{x \in D} g(x) \quad \text{where } x \in D \text{ are iid sampled from } P$$

MC estimation is unbiased:

$$\mathbb{E}_P \left[\frac{1}{|D|} \sum_{x \in D} g(x) \right] = \frac{1}{|D|} \cdot \mathbb{E}_P \left[\sum_{x \in D} g(x) \right] = \mathbb{E}_P [g(x)]$$

Convergence: by law of large number

$$\frac{1}{|D|} \sum_{x \in D} g(x) \rightarrow \mathbb{E}_P [g(x)] \quad \text{as } |D| \rightarrow +\infty$$

Variance

$$\text{Var}_f\left(\frac{1}{|D|} \sum_{x \in D} g(x)\right) = \frac{1}{|D|} \text{Var}(g(x))$$

larger $|D|$, smaller variance