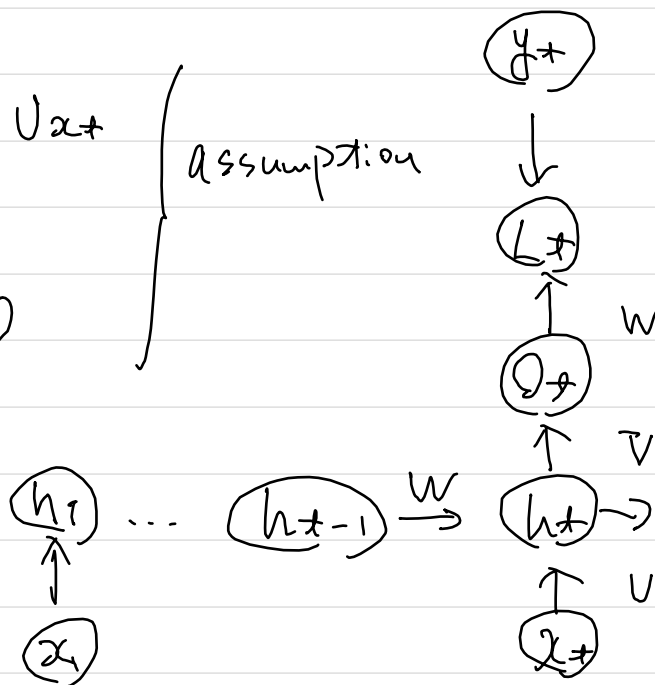


RNN Deep Learning Book 第10章

↙ Recurrent Neural Networks
 回帰結合型 = 2-2, 1, 2-1

$$\begin{aligned} a_t &= b + W h_{t-1} + U x_t \\ h_t &= \tanh(a_t) \\ o_t &= C + V h_t \\ \hat{y}_t &= \text{softmax}(o_t) \end{aligned}$$

} assumption



$$h_t = f(h_{t-1}, x_t, \theta)$$

$$o_t = g(h_t, \theta)$$

$$L_t = L(o_t, y_t)$$

↑
 L の内部で softmax(o_t)

↓
 $\hat{y}_t = \{z_i\}$

$$\begin{aligned} & (\log p(y|x))^{-1} \\ &= \frac{1}{p(y|x)} \end{aligned}$$

x の系列と y の値の系列が与えられたとき、

L とする $x_1, \dots, x_T$ に対して T 回の $y_1, \dots, y_T$ の

計算を繰り返す。

$$L(x_1, \dots, x_T, y_1, \dots, y_T)$$

$$= \sum_t L_t$$

$$= - \sum_t \log p_{\text{model}}(y_t | x_1, \dots, x_t)$$

逐時的な誤差逆伝播法 (back-propagation through time: BPTT)

$$\frac{\partial L}{\partial L_t} = 1$$

時間 step t に与えられた出力の勾配 $(\nabla_{\theta_t} L)_i$:

$$(\nabla_{\theta_t} L)_i = \frac{\partial L}{\partial \theta_{ti}} = \underbrace{\frac{\partial L}{\partial L_t}}_{=1} \frac{\partial L_t}{\partial \theta_{ti}} = \hat{y}_{ti} - 1 \cdot y_t$$

系列の最後 T step から始める

$$\nabla_{\theta} L = \nabla^T \nabla_{\theta_T} L$$

$$\mathbb{V}_{h+1} L = \left(\frac{\partial h_{t+1}}{\partial h_t} \right)^T (\mathbb{V}_{h_{t+1}} L) \\ + \left(\frac{\partial o_t}{\partial h_t} \right)^T (\mathbb{V}_{o_t} L)$$

$$\tanh(x) \text{ の 微分 } (= \frac{d}{dx} \tanh(x)) = 1 - \tanh^2(x)$$

$$\tanh(x) = \frac{\sinh(x)}{\cosh(x)}$$

$$= \frac{\frac{e^x - e^{-x}}{2}}{\frac{e^x + e^{-x}}{2}} = \frac{e^x - e^{-x}}{e^x + e^{-x}}$$

$$= (e^x - e^{-x}) \times (e^x + e^{-x})^{-1}$$

$$= (e^x + e^{-x})(e^x + e^{-x})^{-1} + \frac{(e^x - e^{-x}) \times (-1)}{(e^x + e^{-x})^2 \times (e^x + e^{-x})}$$

$$= 1 + \frac{(e^x - e^{-x}) \times (-1) \times (e^x - e^{-x})}{(e^x + e^{-x})^2}$$

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