

RNN (Recurrent Neural Networks)

回帰結合型ニューラルネットワーク

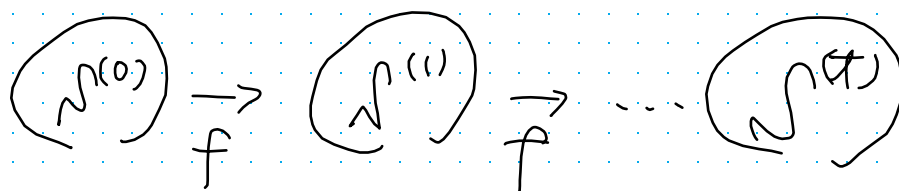
可変長の系列データの処理を可能

計算グラフの展開 (unfold)

回帰の計算を繰り返して持つ計算グラフを表現

$$n^{(t)} = f(n^{(t-1)}; \theta)$$

時刻 t における n の定義が、 $t-1$ における
同じ定義で ref するのを回帰的

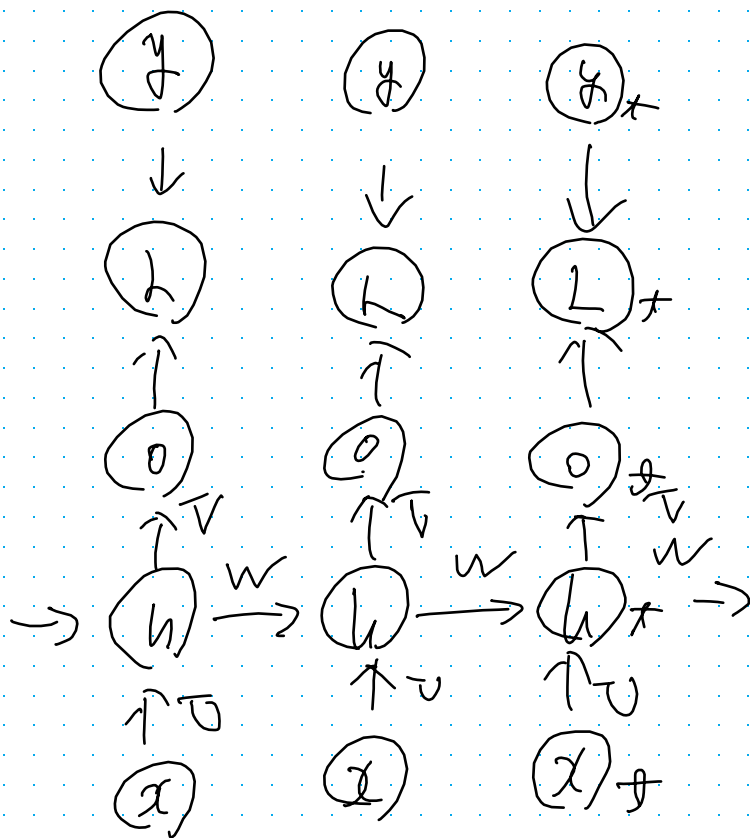


外部信号 $x^{(t)}$ が作用する動的システム

$$n^{(t)} = f(n^{(t-1)}, x^{(t)}; \theta)$$

状態 n は過去の全系列の情報を含む





$$a^t = b + wh^{t-1} + vx^t$$

$$h^t = \tanh(a^t)$$

$$o^t = c + \nabla h^t$$

$$\hat{y}^t = \text{Softmax}(o^t)$$

$$J^{(H)} = \sum_t J^{(H)}(t)$$

$$= - \sum_t \log p_{\text{model}}(y^{(t)} | \underbrace{x^{(1)}, \dots, x^{(t)}})$$

$x^{(1)}, \dots, x^{(t)}$ までの入力 = F での $y^{(t)}$ の

負の対数尤度を $J^{(H)}(t)$ とする

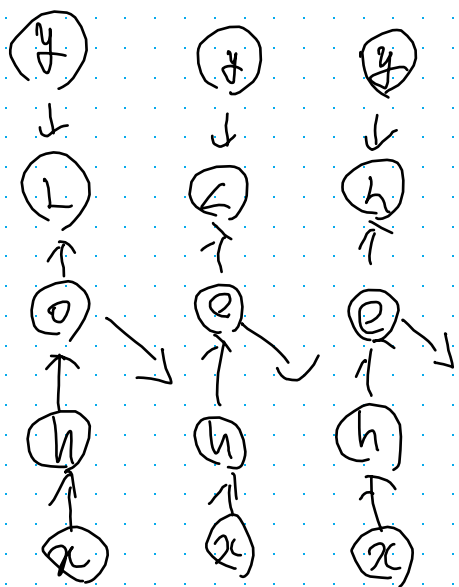
BPTT (Backpropagation Through Time)

→ 順伝播での計算をした状態を、逆伝播で再利用
 していき、保持する必要がある

計算量が大きいと並列化できない

⇒ 回帰的な場合を t -step での出力と t step での状態

(隠れユニット) F 式に持つ



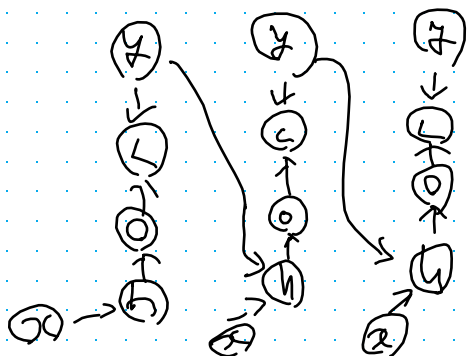
出力ユニットで NN が

未来を予測するために使用する

過去の情報を全て持つ必要

$h \rightarrow 0$ が必要で圧縮されないとダメ

Teacher Forcing



訓練中に t までしか予測しない

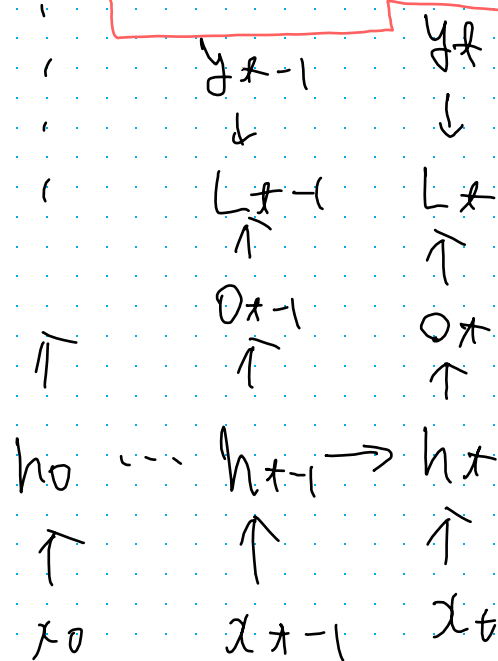
真の出力 y_{t-1} を受ける

↓

条件付で尤度

Teacher Forcing

最終ステップで教師の出力を強制的に用いる BPTT を用いる



$$L = \sum_t L_t$$

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

RNN の

勾配計算

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

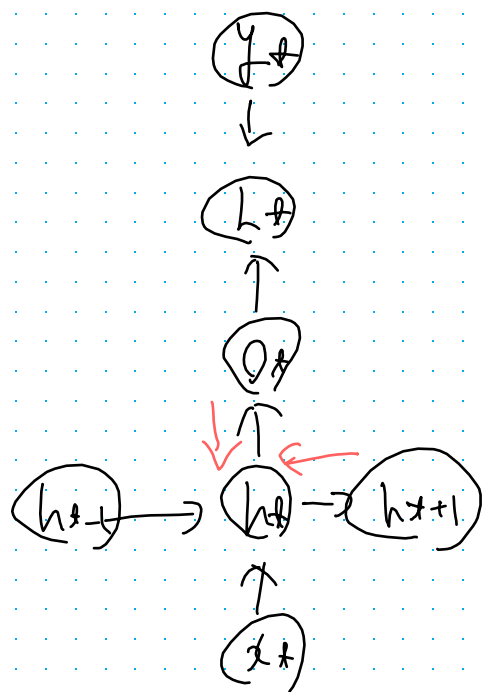
$$= \frac{\partial L_t}{\partial o_t} = \nabla_{o_t} L_t$$

最終ステップで教師の出力を強制的に用いる BPTT を用いる

$$\frac{\partial o_t}{\partial h_t} = V \Rightarrow \frac{\partial L}{\partial h_t} = \frac{\partial o_t}{\partial h_t} \frac{\partial L_t}{\partial o_t} \frac{\partial L}{\partial L_t}$$

$$o_t = C + V h_t$$

$$\frac{\partial L}{\partial h_t} = V^T \nabla_{o_t} L$$



h_t は $\mathbb{R}^n$ - ベクトル $0_t \leq h_{t+1}$ の 両方とも $\mathbb{R}^n$

$$\frac{\partial L}{\partial h_t} = \frac{\partial \sigma}{\partial h_t} \frac{\partial h}{\partial \sigma} + \frac{\partial h_{t+1}}{\partial h_t} \frac{\partial h}{\partial h_{t+1}}$$

$$= \frac{\partial \sigma}{\partial h_t} \nabla_{\sigma} L + \frac{\partial h_{t+1}}{\partial h_t} \nabla_{h_{t+1}} L$$

$$= U^T \nabla_{\sigma} L + W^T \text{diag}(1 - (h_{t+1})^2) \nabla_{h_{t+1}} L$$

$$\sigma = \sigma(h_t) = \sigma$$

$$h_{t+1} = \tanh(b + wh_t + Ux_t)$$

$\therefore \text{diag}(1 - (h_{t+1})^2)$ は 要素 $(1 - |h_{t+1, i}|^2)$ だけ、
対角行列

($\vec{u}$)

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

$$f'(x) = (1 - (f(x))^2)$$

$$\frac{\partial}{\partial x} \tanh(b + wh_t + Ux_t)$$

$$= 1 - (\tanh(b + wh_t + Ux_t))^2 U$$

$$= 1 - (h_{t+1})^2 U$$

$$\frac{\partial L}{\partial w} = \sum_t \frac{\partial h_t}{\partial w} \frac{\partial L}{\partial h_t}$$

$$= \sum_t \text{diag}((-h_t)^2) h_{t-1} \nabla_{h_t} L$$

$$\frac{\partial L}{\partial u} = \sum_t \frac{\partial h_t}{\partial u} \frac{\partial L}{\partial h_t}$$

$$= \sum_t \text{diag}(1 - h_t^2) x_t \nabla_{h_t} L$$

カシワタシハズレ
算出スベシ

10.2.3 Directed Graphical Model との Recurrent Neural Network

$$P(\mathbf{Y}) = \prod_{t=1}^T P(y^{(t)} | y^{(t-1)} \dots y^{(1)})$$

Graphical Model の Edge は他の変数に直接依存

に、子か父か区別せず e.g. $z_t \neq 7$ 仮説は

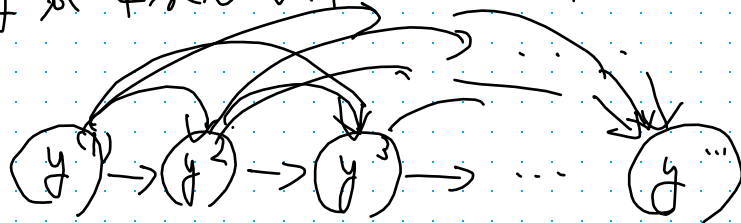
過去の history から $y^{(t)}$ に edge が付くことは不可能

$\{y^{(t-1)}, \dots, y^{(1)}\}$ から $y^{(t)}$ に edge

が付く

① RNNを完全グラフとして解釈

y^t が全ての元へつながる $\Rightarrow \text{Param } O(k^2)$

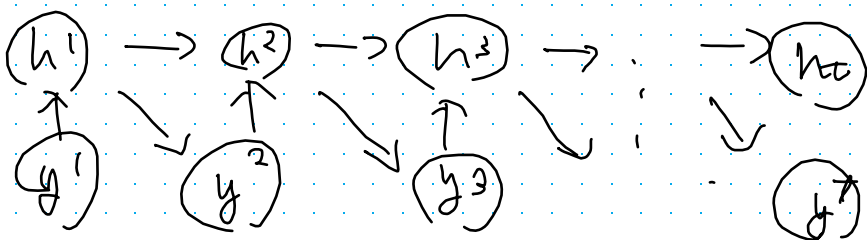


非効率

② 隠れunit h_t を random variableとして扱う

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

同じパラメータ θ が全ての t に対して共通 $O(1)$



このようにして h_t を state とする $y_1 \rightarrow h_1 \rightarrow y_2$

のように past future の両方に関与する

h_t は時間の変数 $h_t \rightarrow h_{t+1}$

影響を及ぼす

Recurrent NNでは

- ① 1つの x -が共有されている場合, 同一 param が異なる時間 step で使用されるという前提に基づいて

t に x が変化する時刻 $t+1$ に x が変化する conditional dist. は stationary である前提

① ↓ ②

- (2) ある step とその次の step の関係が x に依存しない

→ 早期に減衰! ? LR schedule の

step decay などは

無理そう

= できる?

系列の長さの決定

任意の RNN に x が

- ① EOF, EOL (シフト・リセット)

↑ の一般的な

- ② 生成 stop 可能な判断をする x の出力を z とし z に追加

- ③ 系列の長さ n を決定 $P(x_1 \dots x_n) = P(x_1) P(x_2 \dots x_n | x_1)$

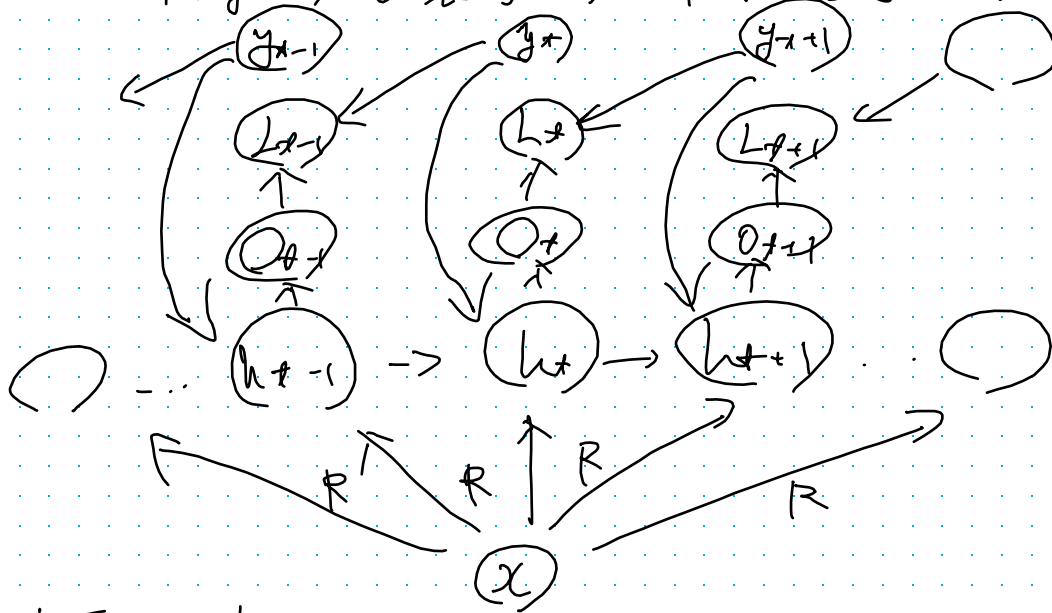
10.2.4 RNN での系列モデリング

$P(y|\theta)$ と $\theta = w$ として $P(y|w)$

と conditional dist を表すようにして解釈できる

$P(y|w)$ のモデルは w 及び x の関数に与えられる

$P(y|x)$ を表すように拡張することが可能



固定長ベクトル

x を系列 y の分布へ写像する RNN

e.g. 画像からの caption 生成

観測された出力系列の各要素 $y^{(t)}$ は,

現在 step の h_t として

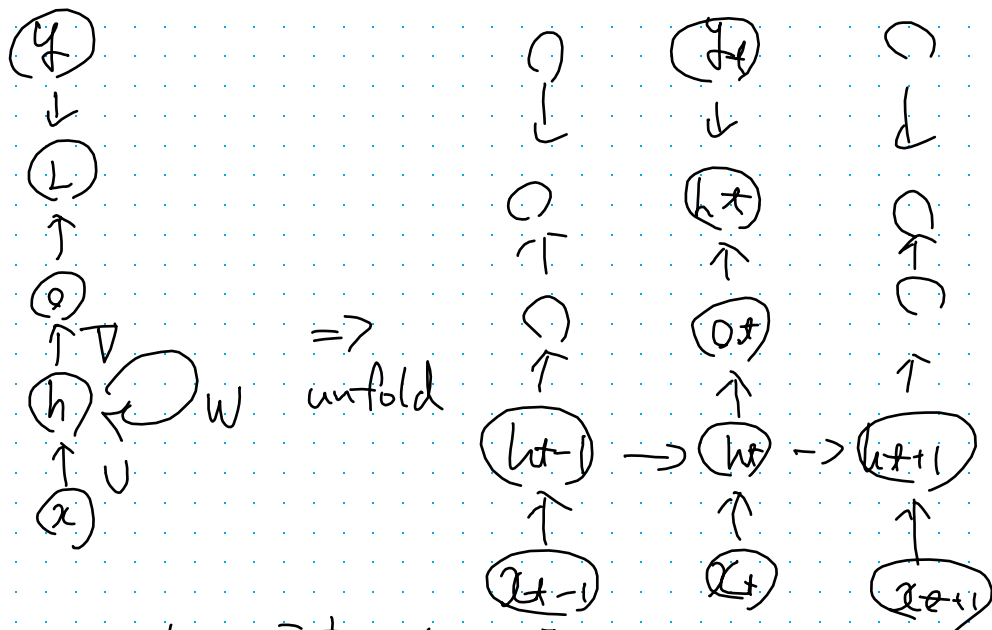
と 前 step の正解として

両者の役割を果たす

重み行列 R と
 h_t と h_{t+1} の相互作用を
parametrized

条件付き独立性

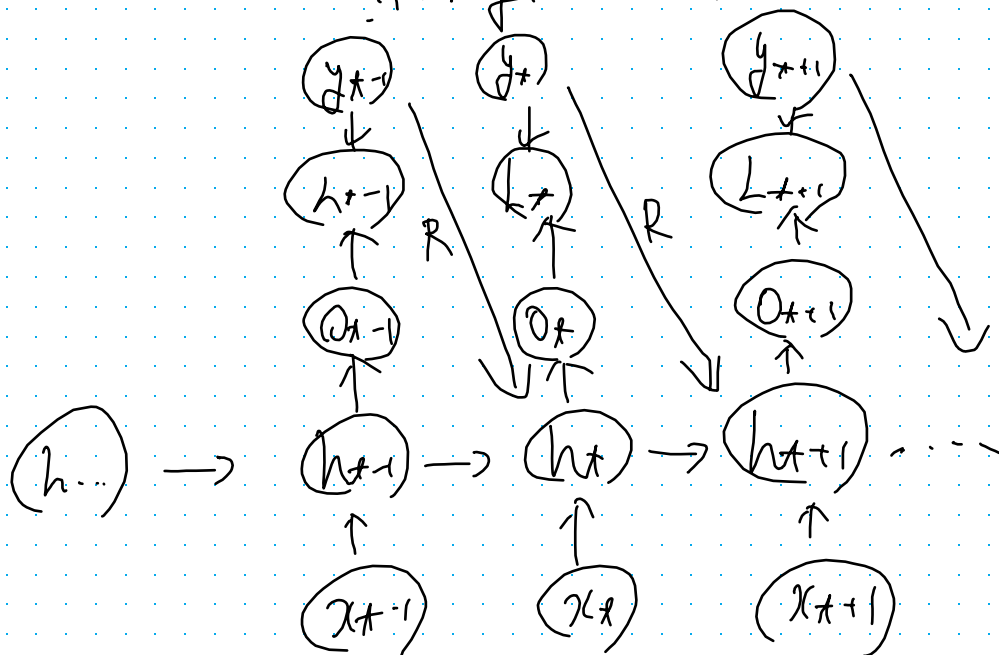
$$P(y_1 \dots y_T | x_1 \dots x_T) = \prod_t P(y_t^* | x_1 \dots x_T)$$



この独立を仮定しなさい

$\Rightarrow$ 直前の出力から現在の状態への接続を省略する

系列 y_1 に対して任意の分解を元でなさい



入力系列全体に依存する y^t の予測を出力したいことがある

e.g. 音声認識

(未来と過去を見ることが必要な場合)

10.4 Encoder - Decoder

可変長 $\rightarrow$ 可変長系列の学習を行う RNN

e.g. 音声認識・機械翻訳・Q&A

Encoder は 文脈 C を出力

入力系列 $X = \{x^1, \dots, x^{(n_x)}\}$ を要約した 1 つの
or 1 つの系列

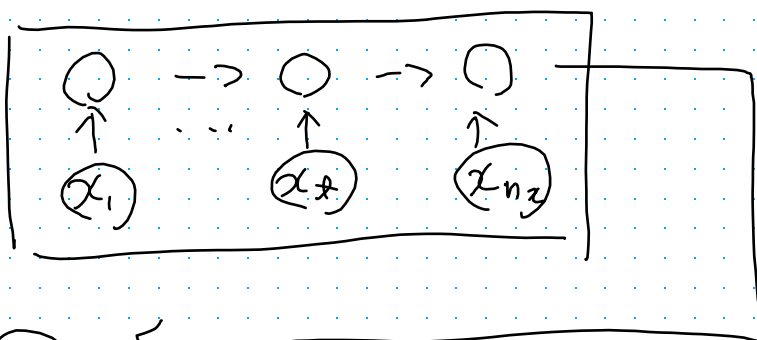
Decoder は 出力系列 $y = \{y^1, \dots, y^{(n_y)}\}$ を生成するため
文脈 C に よる 学習が行われる

Train 系列の x と y について

$\log P(y_1 \dots y_{n_y} | x_1 \dots x_{n_x})$ の平均が最大になるように

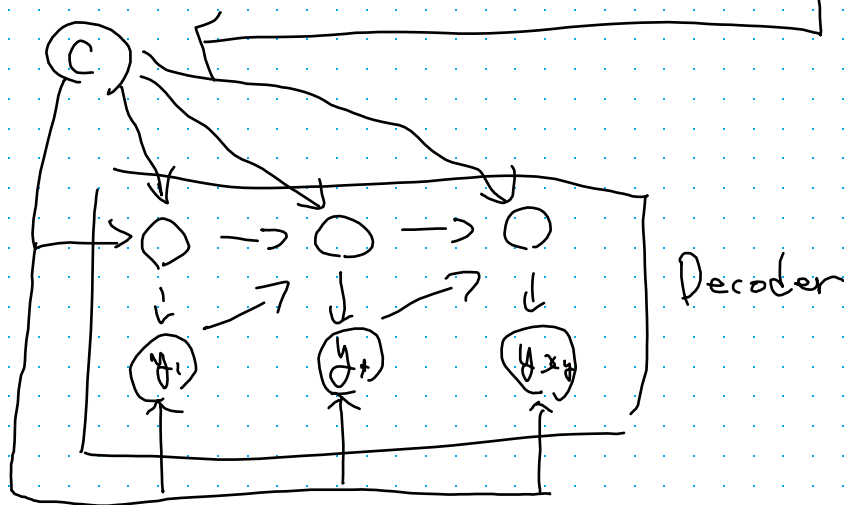
Encoder

2つの RNN を使う



通常

Encoder の最後の状態

 h_{n_x} を Context c として利用

Decoder

c のホールドが低すぎると、長.. 系列情報を適切に要約できない

Solution 1, c を可変長ベクトルにする

2, アテンションメカニズム

RNNのデコを深くする

の複雑さの分解

のデコに個別にMLPを挿入する

の隠れ層 → 既知層の経路に skip 接続する

10.6 再帰型 NN

Recursive NN は Recurrence NN と違って
深層構造の計算グラフを持つ

深さ t から $\log t$ に減らせる

→ 長期依存を扱うのに有利

最大の eigen vec
を保持する h_0 の
要素は最終的に
破棄される

最適な木構造自体は未解明

10.7 長期依存性の課題

$$h_t = W^T h_{t-1}$$

$$h_t = (W^t)^T h_0$$

$$\therefore W = Q \Lambda Q^T$$

$$h_t = Q^T \Lambda^t Q^T h_0$$

※

固有値は t 乗する → 未満は 0 に減衰し、1 以上になる
爆発

慎重にスケールが選ばれた FFN を上記が出来る

スケールが広い領域に存在するとしても、小規模なモデルに Robust になる
(36)

10.8 Echo State Networks (ESNs)

reservoir computing

① 回帰重みと入力重みの学習の難いことを回避するアプローチ

liquid state machine (ESN の 隠層 unit
→ 隠層の)

τ = 隠出力の spiking
= ニューロン (使用)

学習が楽

Problem

RNN にたいして $h_t \rightarrow h_{t+1}$ $\left\{ \begin{array}{l} \text{この学習重みは} \\ x_t \rightarrow h_t \end{array} \right.$ 学習が難しい

②

ヤコビ行列の spectral 半径に固定

③ 重みをうまく初期化する

10.9 複数時間スケールのための leaky unit と他手法

1. 時間方向に skip 接続を追加

2. leaky unit

μ_t の移動平均 μ_t

$$\mu_t \leftarrow \alpha \mu_{t-1} + (1 - \alpha) \mu_t$$

α が 1 に近いと μ_t は 過去の情報を長期間記憶し、

α が 0 に近いと 過去の情報は急速に破棄される

3. 接続の削除

長期的に情報を削除して、より短い接続に replace

10.10 Long Short Term Memory & RNN w/ Gate

gated RNN

e.g. LSTM

gated recurrent unit (GRU)

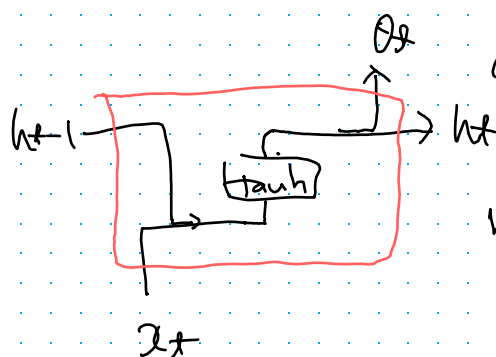
これを行うためには、移動で選択した定数
また、パラメータによる接続重みを使った。

GRUでは、これを時間 step ごとに変更可能な
接続重み
1に一般化
した

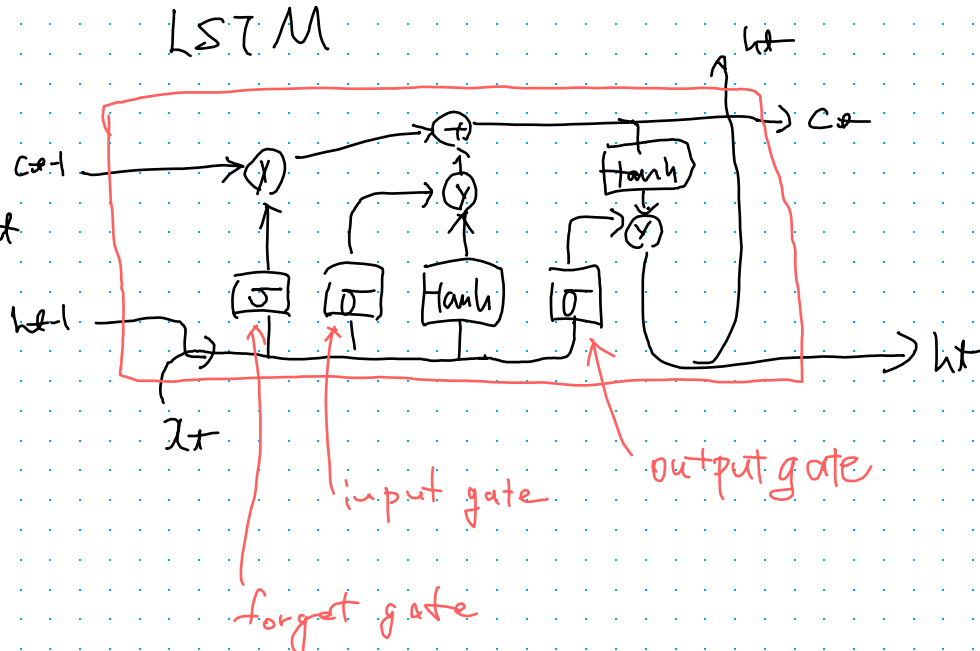
Leaky Unit

と同様に ゲート付き RNN は 消失を爆発を抑制
の配を持ち、これを伝播させる

RNN



LSTM



LSTM 回帰結合型ネットワークは、

入力と回帰ユニットのアフィン変換に要素ごとの非線形変換を単純に単純に適用するユニットの代わりに、通常の RNN の回帰に加えて、内部回帰を持つセルを LSTM セルと言う。

GRU (Gated Recurrence Unit) と呼ばれる

ユニットを持った RNN が最近提案されている。

↓

LSTM と似たような単一のユニットが forget status の Unit の更新の決定を同時に制御する点で異なる。

10.12 長期依存性の option

grad clipping

↳ RNNなどの遅延非同期関数の

非常に大きな or 小さな grad に対して

傾向がある

10.11.2 情報の流れを促進するための

normalize

grad 消失への対策

grad の消失を防ぐために

↳ LSTM の利用 (self loop と gate)

(0.13 PA の F の X と Y)

DONE

Sequence parity function

input: sequence (binary)

output: whether num of 1's seen so far

is odd or even

↓

1

↓

0

e.g. input [0, 1, 0, 1, 1, 0]

the parity sequence is [0, 1, 1, 0, 1, 1]

Design vanilla RNN to implement the parity func

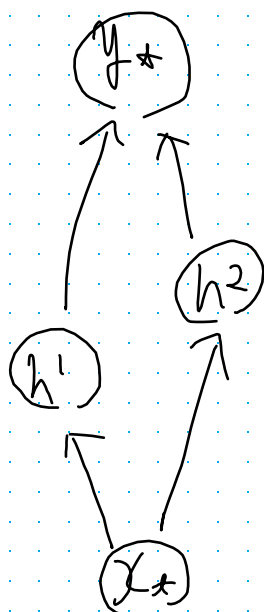
Your implementation should include the equations of your hidden units and the details about

activations w/ diff input at x_t (0 or 1)

y_t Parity Bits 0 1 1 0 1 1 $\rightarrow$
 x_t input 0 1 0 1 1 0

Each parity bit is the XOR of
input and the previous parity bit

Parity is a classic example of a problem
that's hard to solve with shallow
feed-forward net but easy to solve
w/ RNN



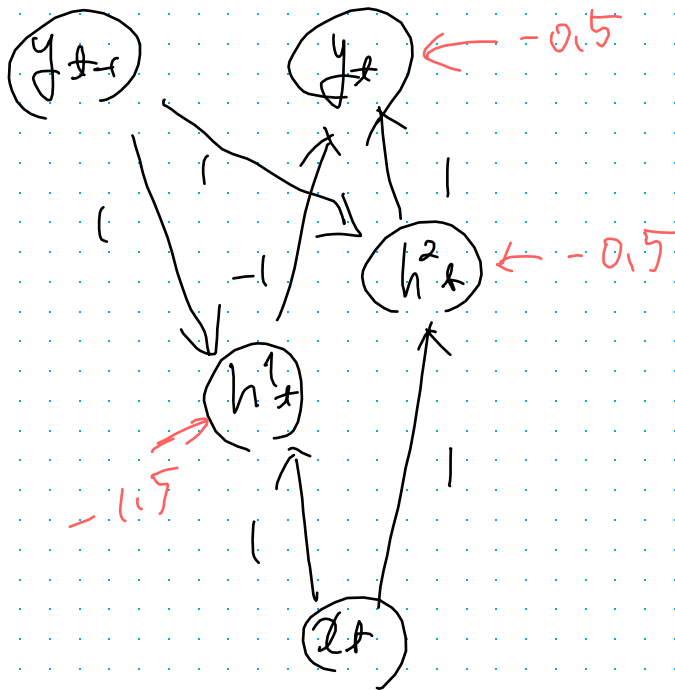
XOR

y_{t-1}	x_t	y_t
0	0	0
0	1	1
1	0	1
1	1	0

linear model cannot compute XOR

$\hookrightarrow$ then use hidden layer

one binary threshold unit



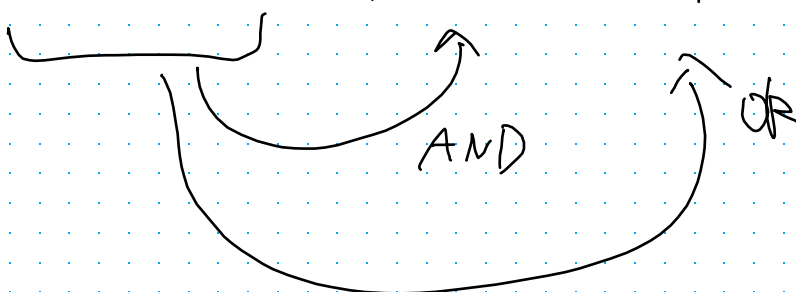
$$h^1_t = f(y_{t-1} + x_t - 1.5)$$

$$h^2_t = f(y_{t-1} + x_t - 0.5)$$

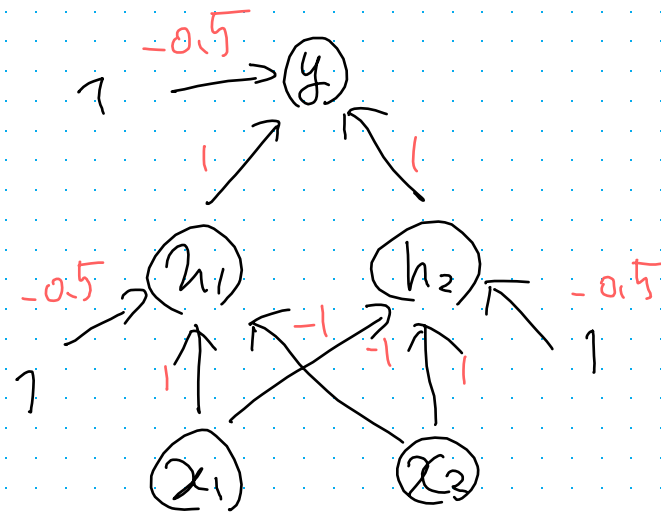
$$y_t = f(h^1_t - h^2_t - 0.5)$$

f is hard threshold
activation
function

y_{t-1}	x_t	AND ↓ h^1_t	OR ↓ h^2_t	y_t
0	0	0	0	0
0	1	0	1	1
1	0	0	1	1
1	1	1	1	0



Designing a network to compute XOR



1 is bias term

$$h_1 = f(x_1 - x_2 - 0.5)$$

$$h_2 = f(x_2 - x_1 - 0.5)$$

$$y = f(h_1 + h_2 - 0.5)$$

x_1	x_2	h_1	h_2	y
0	0	0	0	0
0	1	0	1	1
1	0	1	0	1
1	1	0	1	0

← この表、文字は
fを計算した
手紙。

↑ ↑ ↑
AND OR XOR
zick zick zick

End of Note.

$$f(z) = \begin{cases} z > 0 & \rightarrow 1 \\ z \leq 0 & \rightarrow 0 \end{cases}$$

