

Neural Turing Machine with GRU Controller

日大生産工 (学部) ○石橋 駿 日大生産工 山内 ゆかり

1 まえがき

Recurrent Neural Network(RNN) [1]は隠れ層に再帰的な構造を持っていて、他の Neural Networkとは違い時系列データを扱うことができるモデルである。そのため自然言語処理や音声認識などのタスクで用いられている。

シンプルなRNNは時系列データが長くなってくると学習が難しくなる。そのためRNNを拡張したLong Short-Term Memory (LSTM) [2]やGated Recurrent Unit(GRU) [3]といったモデルがある。

LSTMはRNNの隠れ層のユニットをLSTM blockと呼ばれるメモリと3つのゲートを持つブロックに置き換えたものである。3つのゲートは入力ゲート、出力ゲート、忘却ゲートと呼ばれる。入力ゲート及び出力ゲートは入力重み衝突・出力重み衝突を解決し必要な誤差信号だけを適切に伝搬するためのもの。忘却ゲートはメモリに記憶した内容を忘れることを学習する。

GRUはLSTMが3つのゲートで構成されていたのに対して、更新ゲートとリセットゲートの2つのゲートによって構成される。更新ゲートは最後の状態からどれだけの情報を保持し、前の層からどれだけの情報を取り込むかを決定し、リセットゲートはLSTMの忘却ゲートのように機能する。基本的にはLSTMのように機能するがGRUの方が学習の収束が速い。

一般的にコンピュータプログラムは演算、論理的な制御、外部メモリの3つの基本的なメカニズムを有する。しかし Neural Networkをはじめとした機械学習は論理的な制御と外部メモリの使用をほとんど無視してきた。

そこでDeepMindはRNNと相互作用できる外部メモリを結合したNeural Turing Machine(NTM) [4]を考案した。

NTMはチューリングマシンやノイマン型コンピュータと似た構造をしているが微分可能なエンドツーエンドであり、勾配降下法に

よって訓練することが可能である。

従来研究ではNeural NetworkにLSTMを用いているが、本研究ではGRUを用いることで学習の収束が速くなることを期待する。

2 従来研究

2.1 NTMにはニューラルネットワークコントローラと外部メモリの二つの基本コンポーネントが含まれている。図1はNTMの高レベルの図を示したものである。ほとんどのニューラルネットワークと同様にコントローラは入力ベクトルと出力ベクトルを介して外部世界と相互作用する。また標準的なニューラルネットワークとは異なり、選択的な読み出し及び書き込み動作を使用して外部メモリと相互作用する。これらの操作はチューリングマシンと同様にヘッドによって実現する。

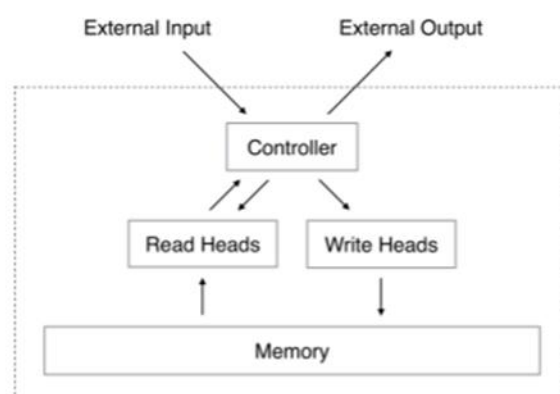


図 1 NTM

2.2 読み取り・書き込み

外部メモリのサイズは $N \times M$ である。Nはメモリ位置を表しMは各位置におけるベクトルサイズである。また外部メモリへアクセスはメモリのどこに注目すべきかを表す分布である attention との加重合計によって求められる。

読み取り時のベクトル r_t は、attention w_t

とメモリ内の行ベクトル M_t を使って式(1)で求められる。

$$r_t \leftarrow \sum_i w_t(i) M_t(i) \quad (1)$$

書き込みは消去と追加の二つに分解することによって可能としている。

消去は消去ベクトル e_t 、attention w_t 、前の時間ステップでの行ベクトル M_{t-1} を使って式(2)によって求められる。

$$M_t \leftarrow M_{t-1}(i)[1 - w_t(i)e_t] \quad (2)$$

追加は追加ベクトル a_t 、attention w_t 、消去後の外部メモリの行ベクトル M_t を使って式(3)で求められる。

$$M_t(i) \leftarrow M_t(i) + w_t(i)a_t \quad (3)$$

2.3. attention 調整

attention はコンテンツベースのアドレッシングと位置ベースのアドレッシングの2つのアドレッシングを組み合わせることによって求められる。またパラメータはコントローラの出力する5つのパラメータと直前のメモリの出力する5つのパラメータと直前のメモリ、ヘッド重みを用いる。図2はattention調整のステップを表す。

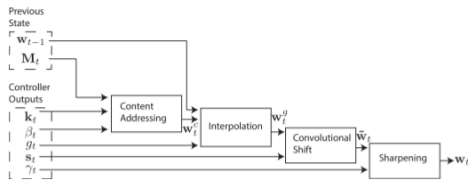


図 2 attention 調整

最初にコンテンツベースのアドレッシングを行う。ここでは式(4)のようにコントローラから与えられたベクトル k_t とメモリ内の各ベクトルの類似度を計算し、その類似度をもとに β_t で正規化してヘッドの attention を初期化する。

$$w_t^c(i) \leftarrow \frac{\exp(\beta_t K[k_t, M_t(i)])}{\sum_j \exp(\beta_t K[k_t, M_t(j)])} \quad (4)$$

次に行うのが位置ベースでのアドレッシングである。ここでは式(5)のようにコンテンツベースのアドレッシングで出力された attention と直前の attention に対してパラメータ g_t を掛け合わせ attention を更新する。 g_t が 0 の場合、コンテンツベースの attention が無視され、1 の場合は直前の attention が無視される。これにより直前の attention をどれだけ引き継ぐか新しい attention をどれだけ採用するかを調整する。

$$w_t^g \leftarrow g_t w_t^c + (1 - g_t) w_{t-1} \quad (5)$$

Convolution Shift では位置ベースのアドレッシングで出力された重みと方向パラメータ s_t を畳み込むために式(6)を使って attention を更新する。これによりコントローラが集中する位置を変更することができる。

$$\tilde{w}_t(i) \leftarrow \sum_{j=0}^{N-1} w_t^g(j) s_t(i-j) \quad (6)$$

Sharpening では Convolution Shift で出力された重みに対して、集中度パラメータ γ_t を用いて式(7)により鋭くして最終的な attention を生成する。

$$w_t(i) \leftarrow \frac{\tilde{w}_t(i)^{\gamma_t}}{\sum_j \tilde{w}_t(j)^{\gamma_t}} \quad (7)$$

2.4. コントローラ

NTM における最も重要なアーキテクチャの選択はコントローラに使用されるニューラルネットワークのタイプである。LSTM などのリカレントコントローラには外部メモリより大きなメモリを補うことができる独自のメカニズムがあり、コントローラをコンピュータの中央処理装置、外部メモリを RAM とするとリカレントネットワークの隠れ層はレジスタに似ている。これによってコントローラは複数の時間ステップで情報を混合することができる。またフィードフォワードコントローラは同時読み書きヘッドの数が NTM が実行できる計算のタイプにボトルネックを課してしまうが、リカレントコントローラは、以前のタイムステップからの読み出しベクトル

を内部的に格納することができるのでその制限を受けない。

3 提案手法

本研究では、コントローラに使用する Neural Network を GRU としたモデルを提案する。

GRU は LSTM に比べて、パラメータが少なく計算量も少ない、そのため学習の収束が速いモデルである。

そのため LSTM Controller よりも学習の収束が速くなることを期待するとともに同等のパフォーマンスを期待する。

4 実験環境

本実験では、8ビットのランダムベクトルのシーケンスに対するコピータスクを行う。シーケンスの長さは1から5または1から10の間で決める。

実験対象とするモデルは従来研究の LSTM をコントローラとしたものと提案手法の GRU をコントローラとしたものとする。

学習はバッチサイズを10とし、100回学習するごとに誤差と正解率を求めた。また5000回学習するごとにNTMの重みを保存する。全体的な学習回数としては20000回とする。

実験結果には LSTM Controller と LSTM Controller に対して誤差 (loss) と正解率 (accuracy) をそれぞれシーケンス長が5の場合と10の場合に分けてグラフにする。またグラフは5ステップずつ平均化する。

5. 実験結果

5.1. シーケンス長が5の場合(loss)

図3の LSTM Controller と GRU Controller の誤差の学習曲線を比較すると学習の初期は GRU の方が収束が速かった。

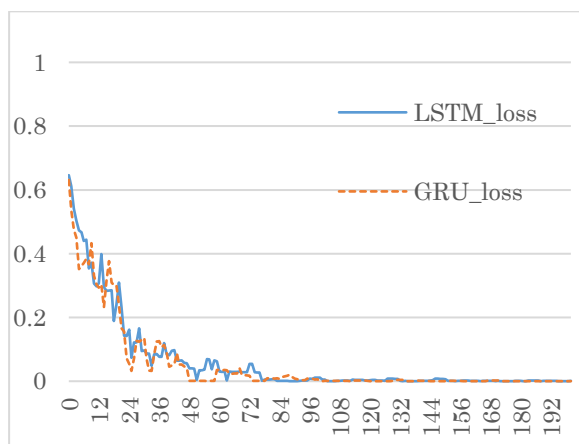


図 3 シーケンス長が 5 の loss

5.2. シーケンス長が5の場合(accuracy)

図4の LSTM Controller と GRU Controller の正解率の学習曲線を比較すると学習の初期は GRU の方が収束が速かった。

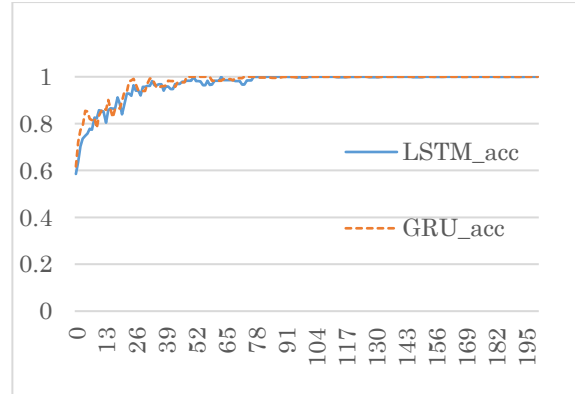


図 4 シーケンス長が 5 の accuracy

5.3. シーケンス長が10の場合(loss)

図5の LSTM Controller と GRU Controller の誤差の学習曲線を比較すると GRU の方がばらつきが大きく、70当たりを境に LSTM の loss が下がって安定していくのに対して GRU は loss が下がっていかなくなった。

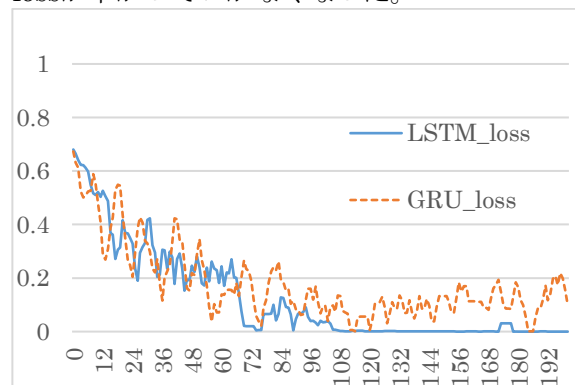


図 5 シーケンス長が 10 の loss

5.4. シーケンス長が10の場合(accuracy)

図6の LSTM Controller と GRU Controller の正解率の学習曲線を比較すると GRU の方がばらつきが大きく、70当たりを境に LSTM の accuracy が上がって安定していくのに対して GRU は accuracy が上がっていかなくなった。

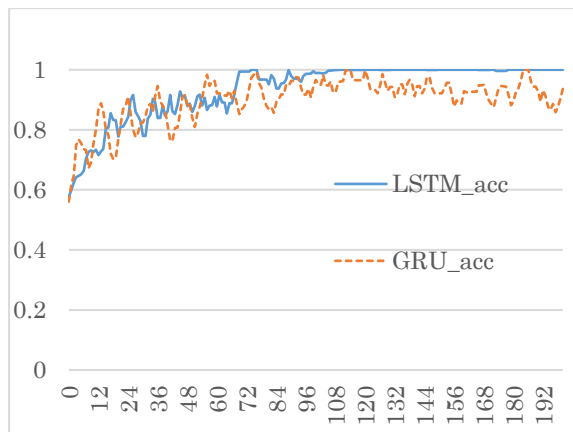


図 6 シーケンス長が 10 の accuracy

6. まとめ

本研究ではNTMのコントローラにLSTMではなくGRUを用いることで学習の収束を速くすることを目的とした。

実験結果を比較した結果、誤差に関してはシーケンス長が5の場合は学習の初期はGRUの方が収束は速かったが、シーケンス長を10とした場合はGRU Controllerはバラツキが大きく誤差が下がりにくくなった。

同様に正解率に関してもシーケンス長が5の場合はあまり差が見られなかったがシーケンス長が5の場合は学習の初期はGRUの方が収束は速かったが、シーケンス長を10とした場合はGRU Controllerはバラツキが大きく正解率が上がりにくくなった。

よってGRUをControllerとするのはLSTMをControllerとするのに比べて長い時系列データに対するタスクを学習するのは難しいものと考えられる。

提案手法の今後の課題としてはなぜGRU Controllerの方が長期間の記憶が難しくなるのかを考察したいと思う。また今回はまだ学習のみなので、テストを行いたいと思う。テストでは学習に用いたシーケンス長より長いシーケンスに対してテストし比較を行いたいと考えている。

参考文献

- 1) Siegelmann, H. T. and Sontag, E. D. (1995). On the computational power of neural nets. *Journal of computer and system sciences*, 50(1):132–150.
- 2) S. Hochreuter, and J. Schmidhuber. Long short-term memory. *Neural Computation*, 9(8), pp.1735-1780, 1997.
- 3) K. Cho, B. Merriehoe, C. Gulcehre, F. Bougares, H. Schwenk, and Y. Bengio. Learning phrase representations using rnn

encoder-decoder for statistical machine translation. *Proceedings of the Empirical Methods Language Processing(EMNLP 2014)*, 2014.

4)