

マルチタスク学習による雑談対話システムへの知識付与*

☆麻生大聖, 滝口哲也, 有木康雄 (神戸大)

1 はじめに

近年, Twitter などの SNS から大量に対話データを収集し, sequence to sequence モデル (seq2seq)^[1] を学習することにより, ユーザ発話に対する応答を生成することができるようになった。

しかし, seq2seq は「そうですね」などの短いよくある応答を生成する傾向がある。また, 相手の発話内容を掘り下げた返事を行う Twitter ユーザは少数であり, システム応答の情報量が少なくなってしまうという問題がある。

本研究では, Twitter に加えてインターネット百科事典 Wikipedia の記事を知識源として使用し, ユーザ発話を掘り下げよう, 知識を含んだ応答生成を行うモデルの構築を目的としている。Wikipedia 記事を入力して同じものを生成するオートエンコーダタスクと, Twitter ユーザ発話を入力してそれに対するリプライ発話を生成する対話学習タスクを同時に学習するマルチタスク学習^[2]を行った。

2 提案手法

LSTM Encoder と LSTM Decoder で構成される seq2seq モデルを使用する。単語の埋め込み表現には, Wikipedia 記事を用いて Skip-Gram モデルで学習した 256 次元の word2vec を使用する。

2.1 学習

対話学習タスクとオートエンコーダタスクが, Encoder と Decoder を共有する model A (Fig. 1) と, Decoder のみを共有する model B (Fig. 2) を学習した。

各タスクはバッチサイズ 256 で, クロスエントロピー誤差が収束するまで学習を行った。1 バッチごとにタスクを切り替えながら交互にパラメータ更新を行ってマルチタスク学習をした。

使用した Encoder と Decoder はすべて共通して, 隠れ層は 2 層で, 各層のノード数は 512 とした。最

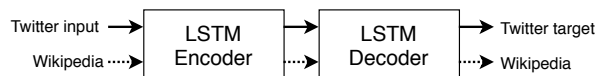


Fig. 1 Model A

適化手法には Adam を使用し, 初期学習率は $5e-4$ とした。

2.2 システム応答生成

システム応答を生成する 3 つの手法を比較した。すべての生成手法において, オートエンコーダタスクで学習した Encoder を使用してベクトル化した Wikipedia 記事を知識源とした。

生成手法 1 (Fig. 3) では, model A を使用して応答を生成する。対話学習タスクで学習した Encoder を使用してユーザ入力をベクトル化し, 知識源から計算した近傍知識ベクトルと合成して得られたベクトルを, Decoder に入力して応答を生成する。

生成手法 2 (Fig. 4) では, model B を使用して, 生成手法 1 と同様に応答を生成する。生成手法 3 (Fig. 5) では, model B を使用するが, 両タスクで学習した Encoder を使用してユーザ入力をベクトル化する。そのうちオートエンコーダタスクで学習した Encoder を使用したベクトルを近傍知識ベクトルの計算に使用する。

ユーザ入力ベクトルを u , 知識源のベクトルを $v_1, v_2, \dots, v_n$ とする。距離関数 d を (1) 式により定義する。このとき, 近傍知識ベクトル k は (2) 式により計算される。ここで, $N(u, n)$ は $d(u, v_i)$ が小さくなる順に i の値を n 個並べた集合である。最終的に Decoder への入力となるベクトル V は (3) 式により計算される。

$$d(u, v) := \|u - v\|_2 \quad (1)$$

$$k = \sum_{i \in N(u, n)} \frac{v_i}{d(u, v_i)} / \sum_{i \in N(u, n)} \frac{1}{d(u, v_i)} \quad (2)$$

$$V = \alpha k + (1 - \alpha) u, \alpha \in [0, 1] \quad (3)$$

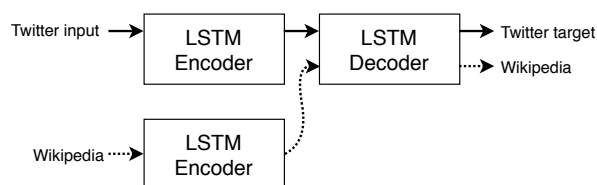


Fig. 2 Model B

*Knowledge embedding to non-task-oriented dialogue system by multitask learning. by ASO, Taisei, TAKIGUCHI, Tetuya, ARIKI, Yasuo (Kobe University)

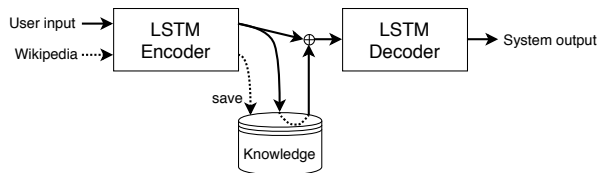


Fig. 3 System Response Generation Method 1

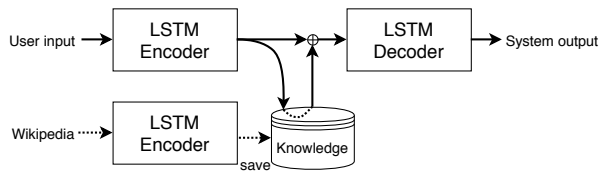


Fig. 4 System Response Generation Method 2

3 実験

3.1 データセット

対話データとして Twitter でのツイートとリプライのペアを収集し、英数字・記号などの特定の文字や、画像や URL などの言外の情報を含むペアを除去し、合計 80 万ペアを使用した。知識源となる非対話データとして Wikipedia 記事を使用した。英数字を含む文や、単語数が 5 未満または 31 以上の文は除去し、特定の文字を除去した合計 260 万文を使用した。

3.2 結果

幅 15 のビームサーチを行って生成したシステム応答の一覧を Table 1 に示す。

生成手法 1 では α によるシステム応答の変化がほとんどなかった。これは、両タスクで同じ Encoder を共有しているため回帰性能が足りず、ユーザ入力に関連する知識がないためであると考えられる。

生成手法 2 では応答文が短く破綻してしまった。これは、知識源を作った Encoder とは別の Encoder でユーザ入力をベクトル化しているため、文のベクトル表現に乖離が生じているためであると考えられる。

生成手法 3 では、応答文が長くなり含まれる語の種類が増えた。生成手法 1,2 と比較して、ユーザ入力と意味的に類似した知識を知識源から取り出すことができているが、複数の知識を合わせることで文の破綻がみられた。取り出すベクトルを 1 個だけにしたときは破綻が少なくなったが、依然として α の値を 0.1 以外にすると大きな破綻がみられた。オートエンコーダタスクに使用する Encoder を、ドロップアウトを行いながら学習することで汎化能力が上がり、複数の知識を合わせて近傍知識ベクトルを計算したときでも、応答文の破綻を低減できていると考えている。

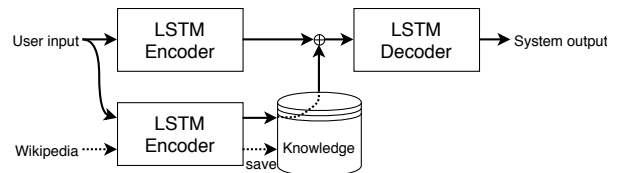


Fig. 5 System Response Generation Method 3

Table 1 System Response

ユーザ入力		日本語には地方ごとに多様な方言がありますよね	
手法	n	α	システム応答
1	5	0.0	そうなんです
		0.5	そうなんです...
		1.0	そうなんです
2	5	0.0	そうなんです
		0.5	ね
		1.0	西部樺は
3	5	0.0	そうなんです
		0.5	にはさらに近いねごとに.
		1.0	スペイン語においては共通語にも近い、登山道という影響がある.
3	1	0.0	そうなんです
		0.5	日本語には多様な方言があり、緑積極が異なりように.
		1.0	日本語には多様な方言があり、地域によってかなりの違いがある.

4 おわりに

本研究では、seq2seq を用いたマルチタスク学習により、ユーザ入力に関連する知識を含む応答生成を行った。応答文は幅広い話題を扱うようになったが、破綻しやすくなった。今後の課題として、ビームサーチにより生成した応答候補の中から破綻の少ない文を選択する方法を研究する予定である。

謝辞 本研究の一部は、JSPS 科研費 JP17K00236, JP17H01995 の助成を受けたものである。

参考文献

- [1] Ilya Sutskever *et al.*: “Sequence to sequence learning with neural networks,” In *Advances in neural information processing systems*, pp. 3104-3112, 2014.
- [2] Minh-Thang Luong *et al.*: “Multi-task sequence to sequence learning,” In *ICLR*, 2016.