

PLSTMによるチャットボット対話の精度検証

Evaluations for personalized chatbot based on LSTM

奥井 颯平

Sohei OKui

中辻 真

Makoto Nakatsuji

NTT レゾナント

NTT Resonant, Inc

LSTM-based chatbot systems are now commercially accepted. Existing systems, however, do not focus on individual user's utterance history, thus it can not select a response according to the individual's characteristics. To solve this problem, we proposed a Personalized LSTM model (PLSTM) that considers the sequence of user's previous utterances as well as his/her current one and then selects the response targetted to him/her. It extracts the context vector by applying self attention mechanism in encoding user's utterances; it updates user's current utterance vector by reflecting the context vector to his/her current vector. According to the evaluations based on conversation dataset extracted from actual Q&A service, Oshiete-goo, we confirmed that PLSTM can predict the response with high accuracy by using the individual's utterance history. This method has already been implemented to the love advice chatbot, Oshiel, existed in actual AI service, Mydaiz.

1. はじめに

LSTMを用いた対話応答学習がチャットボットサービスへの応用に有用であると注目されている。例えば、ドラマキャラと会話できるチャットボット「AI 家売るオンナ」^{*1}、「AI カホコ」や、^{*2}、女子高生 AI チャットボット「AI りんな」、恋愛相談 AI「AI オシエル」^{*3} などが有名である。これらのサービスでは、インスタントメッセージサービスや Q&A コミュニティサービス上に AI キャラクタが登場し、ユーザと雑談や恋愛相談といったコミュニケーションを行う。特に、AI カホコ等のチャットボットのように、LINE 等のインスタントメッセージサービス上に登場した AI キャラクタとユーザは、その場限りではなく、日をまたがって継続的にコミュニケーションを楽しむことが多くなっている。そうした中で、各ユーザに対して、個人化された AI 応答を提供することが今後のユーザ満足度向上に繋がる可能性があると考えられる。しかし、現状のチャットボットはユーザの名前や性別などの登録されたユーザ情報や、GPS から把握される生活エリアに基づいたルールベースでの個人化しか実現できていない。つまり、ユーザの現在の発話を基にして次の AI 応答を選択しているのみであり、ユーザと AI キャラクタとの過去のコミュニケーションを活かしきれていない。そのため、実際に人間の間でコミュニケーションを円滑にしたり、親しみを感じさせるために重要となる過去のやり取りから推定されるユーザ興味や性格に応じた気の利いた AI 応答は実現できていない。

この問題に対し、我々は個人化された AI 応答を実現する PLSTM [奥井 18] を提案してきた。PLSTM はユーザと AI の過去のコミュニケーション履歴を現在のユーザ発話への AI 応答に用いる。特に時時刻刻と変化するユーザ興味や状況に対応するため、PLSTM は Encoder-Decoder モデルで利用されている Attention メカニズムを応用し、ユーザの過去の発話列の中で、現在のユーザ発話と関連性の高いトピックを抽出し、現在のユーザ発話へ反映させる。このように過去の関連するトピ

クを現在の発話に反映させることで、より個人化されたユーザ満足度の高い AI 応答選択を実現する。本稿では「AI オシエル」で用いた時系列情報を含むチャットボット対話データを用いて PLSTM の精度を検証する。

2. 関連研究

深層学習を用いて対話を個人化する試みは、深層学習を用いた情報推薦の研究 [Elkahky 15] に近い。深層学習ではない試みとしては、明示的な概念体系であるセマンティクスを用いて、直近のユーザ行動と興味概念が体系的に近い概念に属する過去のユーザ行動を組合せ、次のユーザへの推薦を決定する研究がある [Nakatsuji 12]。深層学習に基づく手法は、特徴量の人手調整や、言語 ツール、外部リソースを取得する手順を必要としないメリットがある。例えば、長文回答構築手法 [中辻 19] では質問の背後にあるカテゴリやタイトルなどの情報を単語ベクトルに埋め込むことで質問に使用された単語の背景に合った回答を抽出し、結論を表す文と結論に対する補足を表す文の表現ベクトルを計算する。PLSTM は過去のユーザ発話を、潜在的に関連する過去のユーザ発話トピックをとって用い、外部リソースを活用せず個人の特性に応じた特徴ベクトルを学習する。QA システムにおいて AI の応答を決定するにあたり、生成と選択のアプローチがあるが、本稿における著者らの一次検証において、応答生成は応答選択タスクに比べて性能が劣る傾向が見られた。そのため、本稿では QA-LSTM のような応答選択型のフレームワークと PLSTM を比較する。

3. QA-LSTM

本章では、提案手法の土台となる QA-LSTM [Tan16] フレームワークについて説明する。

はじめに、LSTM について説明するとともに、本稿における用語を定義する。入力文を $\mathbf{X} = \{\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_N\}$ とする。ここで、 $\mathbf{x}_t$ は t 番目の単語ベクトルである。 t 番目の隠れベクトル $\mathbf{h}_t$ は以下の式で更新される。

連絡先: 中辻 真, nakatsuji@nttr.co.jp

*1 <https://www.ntv.co.jp/ieuru-gyakushu/AI>

*2 <http://www.ntv.co.jp/kahogo-kahoko/special/index.html>

*3 <https://oshiete.goo.ne.jp/ai>

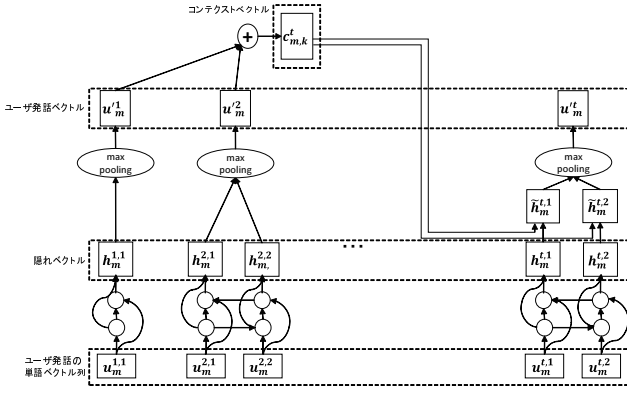


図 1: Context QA-LSTM

$$\begin{aligned}
 \mathbf{i}_t &= \sigma(\mathbf{W}_i \mathbf{x}_t + \mathbf{U}_i \mathbf{h}_{t-1} + \mathbf{b}_i) \\
 \mathbf{f}_t &= \sigma(\mathbf{W}_f \mathbf{x}_t + \mathbf{U}_f \mathbf{h}_{t-1} + \mathbf{b}_f) \\
 \mathbf{o}_t &= \sigma(\mathbf{W}_o \mathbf{x}_t + \mathbf{U}_o \mathbf{h}_{t-1} + \mathbf{b}_o) \\
 \tilde{\mathbf{c}}_t &= \tanh(\mathbf{W}_c \mathbf{x}_t + \mathbf{U}_c \mathbf{h}_{t-1} + \mathbf{b}_c) \\
 \mathbf{c}_t &= \mathbf{i}_t * \tilde{\mathbf{c}}_t + \mathbf{f}_t * \mathbf{c}_{t-1} \\
 \mathbf{h}_t &= \mathbf{o}_t * \tanh(\mathbf{c}_t)
 \end{aligned}$$

LSTM ブロックは3つのゲート(入力ゲート $\mathbf{i}_t$, 忘却ゲート $\mathbf{f}_t$, 出力ゲート $\mathbf{o}_t$)とメモリセル $\mathbf{c}_t$ から構成される。 σ はシグモイド関数である。 $\mathbf{W} \in R^{H \times N}$, $\mathbf{U} \in R^{H \times H}$, $\mathbf{b} \in R^{H \times 1}$ は本ネットワークにおける学習パラメータである。単方向の LSTM の場合、各トークンより以前の情報を引き継いで文脈情報を学習していく一方で、各トークンの後に現れるトークンからの文脈情報を考慮することが難しい。BiLSTM は、前後の文脈情報を双方向に処理し、2つの出力ベクトルを得る。BiLSTM ブロックの出力は、前後方向の出力2つを以下の式のように連結したものである $\overrightarrow{h(t)} = \overrightarrow{h(t)} \parallel \overleftarrow{h(t)}$ 。

QA-LSTM のフレームワークは以下の通りである。まず質問 q , 回答候補 a からなる、与えられた入力ペア (q, a) に対し、 q , a それぞれについて単語ベクトル化する。次に、それぞれ別々の BiLSTM ブロックに単語ベクトル化したシークエンスを入力として与える。BiLSTM から出力されたベクトルに対して max pooling を行い、双方向分の2つの出力ベクトルを連結することにより、 q においては $\mathbf{o}_q$, また a においては $\mathbf{o}_a$ として固定長の分散表現にする。最後に、入力ペア (q, a) に対してコサイン類似度 $\cos(\mathbf{o}_q, \mathbf{o}_a)$ によってスコアリングを行う。

学習要素は以下のヒンジロス関数によって定義する。

$$\mathcal{L} = \max\{0, M - \cos(\mathbf{o}_q, \mathbf{o}_a^+) + \cos(\mathbf{o}_q, \mathbf{o}_a^-)\} \quad (1)$$

ここで $\mathbf{o}_a^+$ 本来の正解の回答の出力ベクトルであり、 $\mathbf{o}_a^-$ は全回答空間からランダムに選出された不正解の回答の出力ベクトルである。また、 M は損失関数におけるマージンである。複数の正解回答を持つ質問については、その質問に紐づく各回答についてそれぞれ複数の別のペアとして扱う。

4. PLSTM

本章ではユーザの過去の発話列から、現在の発話と関係付けられるトピックをアテンションモデルを用いて抽出する PLSTM について説明する。まず、アテンションメカニズムによるユー

ザの過去の発話列を現在の発話と関係付ける方法について述べ、その次に、AI 応答の個人化により選択された応答例について述べる。

4.1 アテンションメカニズム

本稿で提案するアテンションメカニズムについて説明する。図 1 に提案するアテンションメカニズムのイメージを示す。

まず、用語を定義する。

ユーザの数を M とし、 m 番目のユーザの発話を $\mathbb{U}_m = \{\mathbf{u}_m^1, \mathbf{u}_m^2, \dots, \mathbf{u}_m^t\}$ とする。ここで、 $\mathbf{u}_m^t$ は t 番目の発話ベクトルである。各発話は、複数の単語トークンからなる。例えば t 番目の発話列は $\{\mathbf{u}_m^{t,1}, \mathbf{u}_m^{t,2}, \dots, \mathbf{u}_m^{t,n_m^t}\}$ と表される。ここで、 n_m^t は m 番目のユーザの t 番目のユーザ発話に含まれる単語長であり、 $\mathbf{u}_m^{t,k}$ は m 番目のユーザの t 番目の発話の中の k 番目の単語ベクトルである。なお、ユーザ発話 $\mathbf{u}_m^t$ に対応する応答は、 $\mathbf{a}_m^t$ であり、その発話ベクトルは $\mathbf{a}_m^t$ 、発話列は $\{\mathbf{a}_m^{t,1}, \mathbf{a}_m^{t,2}, \dots, \mathbf{a}_m^{t,n_m^{t,a}}\}$ と表される。ここで、 $n_m^{t,a}$ は m 番目のユーザの t 番目のユーザ発話への AI 応答に含まれる単語長であり、 $\mathbf{a}_m^{t,k}$ は m 番目のユーザの t 番目の発話への AI 応答の中の k 番目の単語ベクトルである。

次にアテンションを現在の発話ベクトル $\mathbf{u}_m^t$ に反映する手順について述べる。

(1) データセットからあるユーザ m の現在の発話 $\mathbf{u}_m^t$ を一つピックアップする。

(2) ユーザ m の過去の発話系列 $\{\mathbf{u}_m^1, \mathbf{u}_m^2, \dots, \mathbf{u}_m^{t-1}\}$ の各発話に対し、QA-LSTM と同様の方法で発話ベクトルを計算する。例えば $t-1$ 番目の発話ベクトルは、QA-LSTM と同様に、その発話内の単語列に BiLSTM を適用し、出力される n_m^{t-1} 個の隠れベクトル $\{\mathbf{h}_m^{t-1,1}, \dots, \mathbf{h}_m^{t-1,n_m^{t-1}}\}$ に対し max-pooling を適用することで、発話ベクトル $\mathbf{u}_m^{t-1}$ を獲得する。

(3) そして、 t 番目の発話ベクトル $\mathbf{u}_m^t$ を計算する際に、(2) で計算した過去の発話ベクトルを t 番目の現在の発話の k 番目の単語へどれだけ反映するかを決定するコンテキストベクトルを $\mathbf{c}_m^{t,k}$ を以下のように計算する。なお、 $\mathbf{V}_a$, $\mathbf{W}_a$, $\mathbf{U}_a$ は重み行列である。

$$\begin{aligned}
 \mathbf{c}_m^{t,k} &= \sum_{j=1}^{t-1} \alpha_{k,j} \mathbf{u}_m^j, \\
 \alpha_{k,j} &= \frac{\exp(e_{k,j})}{\sum_{l=1}^{t-1} \exp(e_{k,l})}, \\
 e_{k,j} &= \mathbf{V}_a^T \tanh(\mathbf{W}_a \mathbf{h}_m^{t,k} + \mathbf{U}_a \mathbf{u}_m^j)
 \end{aligned}$$

このコンテキストベクトルの計算方法は、[Bahdanau 14]における方法を基にしているが、[Bahdanau 14]は Encoder-Decoder model にアテンションを適用している。一方、提案手法は、同一の発話ストリームの中で、現在の発話に過去の発話からアテンションを計算している点で大きく異なる。さらに、提案手法では、過去のユーザの興味は抽象的に捉えることが重要であると考えている。そのため、複数の単語ストリームからより抽象度の高いユーザ発話ベクトルを計算し、各ユーザ発話ベクトルをどの程度、現在のユーザ発話を構成する単語ストリーム内の隠れベクトルに反映するかを決定している。これは、[Bahdanau 14]のように、一連の単語ストリームの中でアテンションを計算している従来方式とは異なる。

次にコンテキストベクトル $\mathbf{c}_m^{t,k}$ を BiLSTM に以下のように反映させることで、提案手法は、過去のユーザ発話を現在の発

表 1: 実験で用いたデータセットの作成例

	質問例	回答例	ユーザ発話	AI 応答
発話文	以前、合コンで知り合ったある男性とこの5ヶ月ほど毎日連絡を取り合ったり、よく2人で数回遊びに行ったりしています。相手は社会人20代で私も社会人の20代です。	相手のことを意識しすぎじゃないかな？考え過ぎだと思う。脈の有無はそんなに意識しないで、まずは気楽に関係を築くようにしましょう。	u^1 : 合コンで気になる人ができました。	a^1 : そうなんです。その後はどうでしたか？
			u^2 : 5ヶ月間毎日連絡を取り、良く遊びに行きます。	a^2 : 相手はあなたと連絡するのが楽しいのだと思います。それから？
			u^3 : 私は20代の社会人です。	a^3 : そうなのですね。
			u^4 : 相手も20代の社会人です。	a^4 : 年齢が近いのはいいですね。
	これは、脈がもうないということなのでしょうか。		u^t : これは、脈がもうないということなのでしょうか。	a^t : 脈の有無はそんなに気にしなくても良いと思いますよ。

表 2: 回答選択タスクにおける平均適合率の比較

	QA-LSTM	PLSTM-3	PLSTM-5
Top 1	0.05	0.09	0.04
Top 3	0.19	0.21	0.21
Top 5	0.23	0.28	0.24

話の k 番目の隠れベクトル $\mathbf{h}_m^{t,k}$ へ反映することができるようになる。

$$\begin{aligned}
 \mathbf{z}_k &= \sigma(\mathbf{W}_z[\mathbf{u}_m^{t,k-1}, \mathbf{T}] + \mathbf{U}_z \mathbf{h}_m^{t,k-1} + \mathbf{W}_z^a \mathbf{c}_m^{t,k} + \mathbf{b}_z) \\
 \tilde{\mathbf{l}}_k &= \tanh(\mathbf{W}_l[\mathbf{u}_m^{t,k-1}, \mathbf{T}] + \mathbf{U}_l \mathbf{h}_m^{t,k-1} + \mathbf{W}_l^a \mathbf{c}_m^{t,k} + \mathbf{b}_l) \\
 \mathbf{l}_k &= \mathbf{i}_k * \tilde{\mathbf{l}}_k + \mathbf{f}_k * \mathbf{l}_{k-1} \\
 \mathbf{h}_m^{t,k} &= \mathbf{o}_k * \tanh(\mathbf{l}_k)
 \end{aligned}$$

ここで、 z は入力 i , 忘却 f , 出力 o を示すトークンであり、入力ゲート $\mathbf{i}_k$, 忘却ゲート $\mathbf{f}_k$, 出力ゲート $\mathbf{o}_k$ に対応する。 $\mathbf{l}_k$ はセルメモリベクトルである。 $\mathbf{W}_z^a$ と $\mathbf{W}_l^a$ はアテンションパラメータである。

(4) t 番目の発話ベクトル $\mathbf{u}_m^t$ は、上記手順で更新された n_t^m 個の隠れベクトル $\{\tilde{\mathbf{h}}_m^{t,1}, \dots, \tilde{\mathbf{h}}_m^{t,n_t^m}\}$ に対し max-pooling を適用することで獲得できる。

4.2 AI 応答の個人化

チャットボットシステムにおいて、AI キャラクタに対する過去のユーザ発話が「大好きだよ」であり、現在の発話が「久しぶりだね」である場合、ユーザは AI キャラクタと親密なやり取りを行う傾向があり、結果として「会いたかった」と親密さを演出する応答を行う。一方、過去発話が「眠い今日も疲れた」であり、現在の発話が「久しぶりだね」である場合、ユーザは日常の愚痴をチャットボットにこぼす傾向があり、ユーザを心配する「元気にしてた？」を応答が期待できる。このように PLSTM ではユーザの過去発話列からアテンションメカニズムを活用し、コンテキストベクトルを抽出することで、AI 応答の個人化を実現することが期待される。

5. 評価実験

PLSTM が時系列情報を用いて個人の特徴を考慮した応答を選択することによる有効性を検証するために QA-LSTM との比較による評価実験を行った。

5.1 データセット

PLSTM の評価に日本の QA コミュニティサイト「教えて!goo」*4 において質問が複雑かつ長く、我々に知見のある「恋愛相談」カテゴリのデータセットを活用した。QA データ

をチャットボット向けの対話データに成形した例を表 1 に示す。教えて!goo!における質問文書は“二人の出会い”、“二人の関係”、“自身のこと”、“意中の相手のこと”、“具体的な悩み”などに細分化することができる。質問文書の内容を手で細分化し、ユーザ発話系列 $\{u^1, \dots, u^t\}$ として成形し、チャットボットの応答は回答文書を参考に、手で応答系列 $\{a^1, \dots, a^t\}$ を作成した。このように QA データにおける質問応答を読み解きながら、AI 応答文を割り当てることで、時系列情報を持った約 2 万 5 千件の対話応答データセットを作成した。そして、このデータセットに対してランダムに選択した 1 割を試験データセットとし、9 割を学習データセットとし、検証を実施した。また、「教えて!goo」に蓄積された「恋愛相談」、「旅行」、「ヘルスケア」等の 16 カテゴリ、約 100 万件の QA データを使用し、word2vec [Mikolov 13] により単語ベクトルを獲得した。

5.2 比較手法

本稿では以下の 3 つの手法間で精度比較を行った。

- QA-LSTM [Tan16]: LSTM により質問応答それぞれの特徴ベクトルを学習する。モデルの性質上ユーザ発話の時系列情報は学習・推論に活用されない。
- PLSTM-3: 本稿が対象とする著者らの提案手法。3 ステップ前までのユーザ発話からコンテキストベクトルを学習し、ユーザ発話ベクトルを導出する。
- PLSTM-5: 本稿が対象とする著者らの提案手法。5 ステップ前までのユーザ発話からコンテキストベクトルを学習し、ユーザ発話ベクトルを導出する。

5.3 評価基準

定量的な評価方法として、選択した応答が元の発話にひもづく応答と同一であれば正答とし、生成した応答を類似度順に並べ平均適合率*5 を算出した。また、定性的な評価方法として検証者が恋愛相談文書を入力し、選択される応答を考察した。

5.4 評価結果

QA-LSTM と PLSTM の定量的な評価結果、及び PLSTM における定性的な評価結果を以下に示す。

5.4.1 定量評価

各手法の平均適合率を表 2 に示す。QA-LSTM と PLSTM-5 との間に精度における差は少ない。一方、QA-LSTM に対して PLSTM-3 は精度が高い。これにより、現時刻の発話から 3 回程度前までの発話列を応答の推論に活用することで、QA-LSTM に対して高い精度で応答を選択できることが読み取れる。一方、現時刻の発話から 5 回程度前までの発話列を応答の推論に活用することの精度への寄与は少ない。よって、直近 3 回前までの発話を参照することが対話応答の精度向上に有用であると考察できる。

*4 <https://oshiete.goo.ne.jp/>

*5 [https://en.wikipedia.org/wiki/Evaluation_measures_\(information_retrieval\)](https://en.wikipedia.org/wiki/Evaluation_measures_(information_retrieval))

表 3: 定性評価の一例：コンテキストにより変化する AI 応答の例

	検証者 X	検証者 Y
過去発話	u_x^1 : 好きな人ができました u_x^2 : その人とはよく食事に行きます。 u_x^3 : とても素敵な人です。 u_x^4 : 告白しようと思うのですが ⁶ , 大丈夫でしょうか	u_y^1 : 好きな人ができました。 u_y^2 : 目も合わせてくれません。 u_y^3 : いつも私に冷たいです。 u_y^4 : 告白しようと思うのですが ⁶ , 大丈夫でしょうか
ユーザ発話	u_x^t : 脈はありますか？	u_y^t : 脈はありますか？
AI 応答	a_x^t : 脈ありですね	a_y^t : 脈はないとは思いませんよ

5.4.2 定性評価

本モデルが推論時にユーザの過去発話の差により、異なる応答を選択した例を表 3 に示す。この検証では検証者 X と検証者 Y が「恋愛相談」カテゴリにおける質問応答を想定して、異なる過去発話と同一の現在発話をそれぞれ作成し、PLSTM における推論結果を比較した。検証者 X、検証者 Y の現在の発話は「脈はありますか？」と同様の発話をしているが、過去に異なる内容の恋愛相談を発話している。検証者 X は相手とよく食事に行くことを発話し、検証者 Y は相手に好印象を抱かれていないことを発話しており、それぞれ異なる AI 応答が選択されている。これにより、PLSTM モデルにおいて AI 応答の個人化が実現されていることを確認した。

6. 結論

本稿では、現在の発話に基づく AI 応答を個人毎に特化したものとするため、ユーザの過去の発話列から重要なトピックを現在のユーザ発話へ反映し、次の AI 応答の選択に活用する PLSTM の精度を検証した。結果として直近 3 回前までのユーザの過去発話列を考慮することで、AI 応答の選択において平均適合率の精度向上を確認した。本手法を活用したチャットボットは株式会社 NTT ドコモ社が提供する my daiz(マイデイズ)^{*6} という実サービス上で提供されており、日々ユーザの質問に回答している。今後はアプリケーション運用により得られた知見をもとに、以下について検証を進めることとする。ユーザ発話と AI 応答の系列において、ユーザの発話特性を考慮することができる。例えばレスポンスの速度やユーザ発話の内容により、対話の品質を評価できる。これらを考慮した損失関数を設計、精度検証を行う。

参考文献

- [Bahdanau 14] Bahdanau, D., Cho, K., and Bengio, Y.: Neural Machine Translation by Jointly Learning to Align and Translate, CoRR, Vol. abs/1409.0473, (2014)
- [Elkahky 15] Elkahky, A. M., Song, Y., and He, X.: A Multi-View Deep Learning Approach for Cross Domain User Modeling in Recommendation Systems, in Proc. WWW '15, pp. 278-288 (2015)
- [Mikolov 13] Mikolov, T., M., Sutskever, I., Chen, K., Corrado, G., and Dean, J.: Distributed Representations of Words and Phrases and their Compositionality, in Proc. NIPS '13, pp. 3111-3119 (2013)
- [Nakatsuji 12] Nakatsuji, M., Fujiwara, Y., Uchiyama, T., and Toda, H.: Collaborative Filtering by Analyzing Dynamic User Interests Modeled by Taxonomy., in Proc. ISWC '12, pp. 361-377 (2012)

[Tan16] Tan, M., Santos, dos C. N., Xiang, B., and Zhou, B.: Improved Representation Learning for Question Answer Matching, in Proc. ACL '16, pp. 464-473 (2016)

[奥井 18] 奥井 颯平, 中辻 真: LSTM を用いたパーソナル対話技術, 人工知能学会第 32 回全国大会 (2018)

[中辻 19] 中辻 真, 奥井颯平, 藤田明久: LSTM を用いた Non-Factoid 型長文回答構築手法, 電子情報通信学会論文誌 Vol.J102-D, No.4 doi:10.14923/transinfj. (2019)

*6 <https://www.nttdocomo.co.jp/service/mydaiz/>