

修士論文

多様な応答を可能にするニューラル対話生成

中村 良

2019 年 3 月 6 日

奈良先端科学技術大学院大学
情報科学研究科 情報科学専攻

本論文は奈良先端科学技術大学院大学情報科学研究科に
修士（工学）授与の要件として提出した修士論文である。

中村 良

審査委員：

中村 哲 教授	（主指導教員）
松本 裕治 教授	（副指導教員）
須藤 克仁 准教授	（副指導教員）
吉野 幸一郎 助教	（副指導教員）

多様な応答を可能にするニューラル対話生成*

中村 良

内容梗概

ニューラルネットワークを用いた対話応答生成技術が急速に進歩しているが、既存技術に基づく応答文は低い多様性を示す。この問題は Softmax Cross-Entropy (SCE) 損失に起因する。現実の対話では応答文は一意に定まらず確率的・文脈依存な応答候補が複数考えうる。しかしニューラル対話生成では訓練時の Softmax Cross-Entropy 損失が頻出応答に対して多くの訓練ペナルティをもたらす過度な生起確率を与え、評価時の MAP 予測が文脈依存性が少ない頻出応答を選択する。その結果、応答の多様性が低下すると考えられる。本論文では SCE 損失にトークン頻度の逆数に基づく重みをかけた Inverse Token Frequency (ITF) 損失という新たな損失関数を提案する。この損失関数は頻出語クラスの損失を減少することで、SCE 損失が頻出応答に対して過度な生起確率を与える現象を抑制する。また ITF 損失は重みつき SCE 損失という単純な機構であるのでモデルや訓練を難しくしない利点がある。日本語の Twitter データセットを利用した実験では、「魅力」及び「首尾一貫」に関する人手評価において提案手法は先行研究の手法よりも高いスコアを達成した。また品質評価メトリック BLEU-1/2 及び多様性評価メトリック DIST-1/2 に関する自動評価においても提案手法は高いスコアであった。

キーワード

ニューラルネットワーク, 雑談対話システム, 多様性促進, 最大相互情報量, Sampling 探索, Softmax Cross-Entropy 損失, Inverse Token Frequency 損失

*奈良先端科学技術大学院大学 情報科学研究科 情報科学専攻 修士論文, 2019 年 3 月 6 日.

Neural Dialogue Generation Enabling Diverse Responses*

Ryo Nakamura

Abstract

Dialogue response generation technology using a neural network is advancing rapidly, but generated responses based on existing methods show low diversity. This problem is caused by Softmax Cross-Entropy (SCE) loss. In actual dialogues, responses are not uniquely determined, and various context-dependent candidates exist. However, in neural dialogue generation, the SCE loss at training gives more training penalty to frequent responses and gives excessive occurrence probability to them, and MAP prediction at evaluation selects frequent responses with less context dependency. As a result, the diversity of responses decreases. In this paper, we propose a new loss function called Inverse Token Frequency (ITF) loss, which individually downscales the loss for frequent token classes to suppress the phenomenon that the SCE loss gives excessive occurrence probabilities to frequent responses. Since the ITF loss is a simple mechanism as weighted the SCE loss, it does not complicate the model. The training with the ITF loss remains stable as that with the SCE loss. In human evaluation using a Japanese Twitter dataset, our proposed method achieved a higher score than previous research on “engagement” and “consistency”. In automatic evaluation, our proposed method also had a high BLEU-1/2 quality scores and DIST-1/2 diversity scores.

Keywords:

*Master’s Thesis, Department of Information Science, Graduate School of Information Science, Nara Institute of Science and Technology, March 6, 2019.

Neural Network, Chit-Chat Dialogue System, Diversity-Promoting, Maximum Mutual Information, Sampling Search, Softmax Cross-Entropy Loss, Inverse Token Frequency Loss

目次

第 1 章	背景	1
1.1	対話システム	1
1.2	低い多様性問題	2
第 2 章	既存技術	5
2.1	用語定義	5
2.2	基幹技術	6
2.2.1	Sequence-to-Sequence	6
2.2.2	Attention 及び Transformer	7
2.2.3	MemN2N 及び HRED	8
2.2.4	Greedy 探索及び Beam 探索	9
2.3	先行研究 (ベースライン)	10
2.3.1	Sampling 探索及び Top10 探索	10
2.3.2	Maximum Mutual Information	11
2.4	先行研究 (その他)	12
2.4.1	強化学習	12
2.4.2	生成モデル	12
第 3 章	損失関数	13
3.1	既存手法 : Softmax Cross-Entropy 損失	13
3.2	提案手法 : Inverse Token Frequency 損失	13
第 4 章	実験設定	16
4.1	データセット	16
4.2	ベースライン	16
4.3	訓練設定	17
4.4	人手評価	17
4.5	自動評価	18
4.5.1	Perplexity	18

4.5.2	BLEU	20
4.5.3	DIST	21
4.5.4	Repeat	21
第 5 章	結果	22
5.1	人手評価	22
5.2	自動評価	23
5.2.1	ハイパパラメータ	24
5.2.2	モデル	26
5.3	生成例	28
第 6 章	結論	32
	謝辞	33
	参考文献	34
	発表リスト	40

第 1 章 背景

1.1 対話システム

対話システムは自然言語を介して人間とコミュニケーションできるプログラムである。対話システムの先駆けは 1966 年に発明された ELIZA (Weizenbaum, 1966) である。ELIZA は精神療法におけるセラピストの発話をシミュレーションした対話システムであり、非常に単純なプログラムであったが、驚くほど人間らしい対話が可能であった。今日では Microsoft 社の「りんな」、Apple 社の「Siri」、Google 社の「Google Assistant」など、対話内容を考慮した様々な対話システムがスマートフォン等のアプリケーションとして提供されている。この背景には、GPU (Owens et al., 2008) 等の計算機技術の発展と深層学習 (LeCun et al., 2015) 等の機械学習技術の進歩がある。

対話システムに求められる対話は概ね (1) 質疑応答, (2) タスク指向, (3) 雑談に大別できる。(1) の質疑応答対話システムはユーザーの質問に正しく回答する対話システムである。(2) のタスク指向対話システムは映画館やレストランの予約案内のように特定の業務タスクを遂行する対話システムである。対話システムは限定的な目的及び状況下でスムーズに業務をこなす必要がある。(3) の雑談対話システムは会話や談話のようにフラットな言葉のやりとりが可能な対話システムである。対話システムはユーザーが飽きないように会話を持続させ、ユーザーとの友好関係を築く必要がある。

本研究では以下に述べる 3 つの理由から雑談対話システムに取り組み、技術進展を狙う。(a) 雑談は人々の生活の長時間を占めるため、人間の対話を代替ないし拡張する雑談対話システムの創造は人々の生活に大きな影響を与える。(b) 雑談は目的や状況が限定されない困難かつオープンな問題設定であるため、優れた雑談対話システムの創造は汎用人工知能の開発に一役買う。(c) 雑談はユーザーが持つ個性や情動のモデル化などに関連するため、人間らしい対話システムの開発及び各ユーザーにパーソナライズする対話システムの開発に一役買う。

対話システムは様々な技術で開発できるが、今日では (i) 検索ベース (Retrieval-based) と (ii) 生成ベース (Generation-based) の対話システムが主流となっている

る。(i) の検索ベース対話システムはあらかじめ用意されたテンプレート (人間の発話文) をキーワードマッチング等の検索アルゴリズムで選択し (場合によってはテンプレートにキーワードを埋め込み) 応答する。検索ベースの応答文は流暢であるが、柔軟性・応答多様性・文脈的一貫性に問題が生じやすい。ここで流暢とは構文的に正しいこと、柔軟とは機敏に応答を変容できることである。(ii) の生成ベース対話システムは逐次的な単語選択で応答文を生成する。生成ベースの応答文は柔軟であるが、流暢さ・応答多様性・文脈的一貫性に問題が生じやすい。特に生成のためにニューラルネットワークを使用した生成ベース対話システムをニューラル対話生成 (Neural Dialogue Generation) と呼ぶ。本研究では主にニューラル対話生成の応答文が極めて低い多様性を示す問題に取り組む。低い多様性問題は部分的に流暢さや文脈的一貫性を損う原因であるため、この問題の解決は全体的な問題解決の一助となり、個々の問題がより明確になる。

1.2 低い多様性問題

雑談におけるニューラル対話生成はしばしば「そうですね」のような一般的な応答を生成しやすいという問題に直面する。この現象はニューラル対話生成の初期の研究ですでに観測 (Vinyals and Le, 2015; Sordoni et al., 2015) され、Low Diversity や Generic Response (又は Dull Response) と呼ばれる。

対話生成と機械翻訳は系列変換という点で類似したタスクであるが、機械翻訳は原則的に内容変化が起こらない 1 対 1 変換タスクであるのに対し、対話生成は生成文が一意に定まらず確率的・文脈依存な内容変化が起こる多対多変換タスクである。すなわち、機械翻訳は内容が正しければ生成文の表現に多様性がなくても問題ないが、対話生成は (i) 1 つの発話文に対して文脈依存の多様な応答文が考えられ、その逆に (ii) 異なる発話文に対して同じ応答文が考えられる。(i) の例として「好きな食べ物はなんですか？」という発話文に対して、「カレーですね」や「寿司と焼肉が好きです」といった複数の応答候補がある。(ii) の例として「最近忙しいですね」や「寒くなってきましたね」という異なる発話文に対して、「そうですね」という同じ応答候補がある。

これは重大な問題を生じさせる。ニューラル対話生成モデルの訓練時では「そう

ですね」のような学習データに頻出フレーズは希少フレーズよりも多くの訓練ペナルティが供給され、頻出フレーズは大きな生起確率が割り当てられやすくなる (図 1.1).

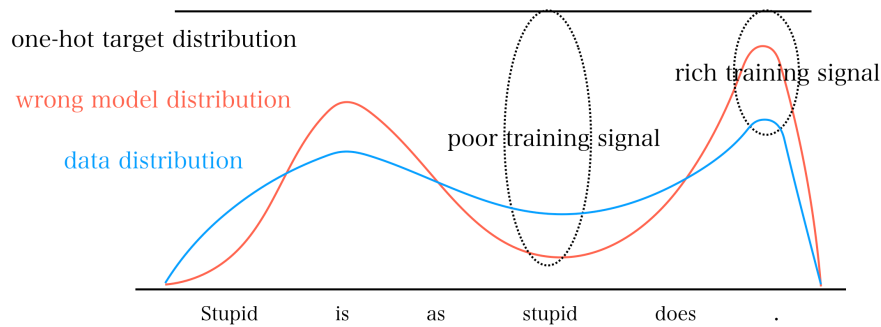


図 1.1 訓練時に頻出する単語「is」や「.」は希少な単語「stupid」よりも多くの訓練ペナルティを供給するため、過度な生起確率が与えられる。

一方、評価時では「そうですね」のようななどの文脈でも通用する文脈依存性が少ない頻出フレーズは尤度が高い生成候補となるが、Greedy 探索や Beam 探索は常にそのような最尤候補のみ選択する (図 1.2).

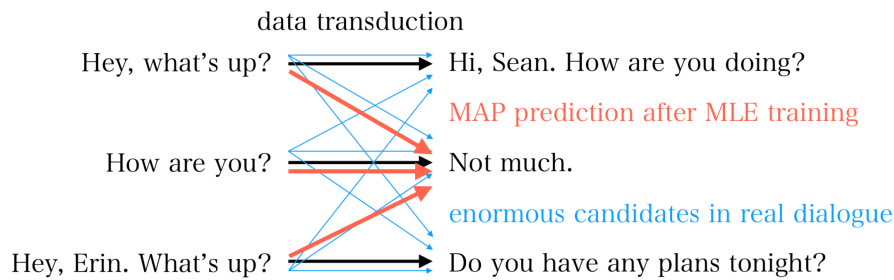


図 1.2 評価時に生起確率が最大のフレーズ「Not much.」のみ選択するため、生成文の多様性が失われる。

これらの問題は 2 点に要約できる。

(1) 第 1 の問題は訓練時に頻出フレーズが希少フレーズよりも多くの訓練ペナルティを供給することである。この問題は学習データの不均衡に起因するが、不均衡

を考慮せずすべてのトークン及びフレーズを同等に扱う Softmax Cross-Entropy (SCE) 損失もこの問題に一役買う。

(2) 第 2 の問題は評価時に生起確率が最大のフレーズのみ選択することである。この問題は尤度の高い選択肢を常に優先する Greedy 探索や Beam 探索などの MAP (Maximum A-Posteriori) 予測に起因し、これを回避するために尤度の高い選択肢以外の選択を促す手法が提案されている (Li et al., 2016a; Serban et al., 2016)。

本研究では (1) を改善するために新しい損失関数 Inverse Token Frequency (ITF) 損失を提案する。学習データの不均衡を考慮するために、ITF 損失は Softmax Cross-Entropy (SCE) 損失の各単語クラスを単語頻度の逆数に基づいて重み付けした損失関数であり、頻出単語の訓練ペナルティ量を抑制する。

ITF 損失で訓練したモデルの生成文は (2) の問題を引き起こす MAP 予測を使用した場合でも多様性が促進される。また ITF 損失のコンセプトは非常に明確であり、簡単に訓練スクリプトに組み込めるという利点がある。本研究では ITF 損失と先行研究で取り組まれた (2) を改善する幾つかの手法との比較実験を行い、ITF 損失の有効性を確認した。

第 2 章 既存技術

この章ではニューラル対話生成及び多様性の促進のために利用される既存技術について解説する。

2.1 用語定義

対話 (dialogue) は 2 人の話者 (speaker) 間で行われ、やや形式的な言葉のやりとり (例えば質疑応答や目的指向タスク) である。会話・談話 (conversation) は 2 人以上の話者間で行われ、よりフラットな言葉のやりとり (例えば冗談や喜怒哀楽) である。特に非目的指向を強調する場合は雑談 (chit-chat) が使われる。対話 (dialogue) は連続した発話 (utterance) の集合と定義される。また発話の単位はターン (turn) が使われ、対話の単位はエピソード (episode) が使われる。特にある話者の発話に呼応する別の話者の発話を応答 (response) という。Twitter では応答をリプライ (reply) と表記され、質疑応答では回答 (answer) と表記される。

対話生成と機械翻訳は類似したタスクである。そのためモデルや訓練手法の流用が頻繁になされ、互いにコミュニティの発展に寄与し合っている。

機械翻訳では翻訳モデルに入力される文をソース (source)、モデルが出力すべき文をターゲット (target) もしくは参照 (reference)、実際にモデルが出力する文を仮説 (hypothesis) という。また文 (sentence) は単語 (word) から構成される。単語の代わりに文字 (character) やサブワード (sub-word) (Sennrich et al., 2016b) から構成される場合もあるため、近年は文を構成する最小単位をトークン (token) と表記する場合が多い。また連続した文の集合を文書 (document) という。

本紙では特に断りがない限り、モデルに入力される文をソース文、モデルが出力すべき文をターゲット文、実際にモデルが出力する文を生成文と表記する。

対話生成は与えられたソース文からターゲット文を生成する系列変換問題として定式化できる。言い換えればソース文 $X = \{x_1, x_2, \dots, x_m\}$ が与えられる条件下での「生成文 $\hat{Y} = \{\hat{y}_1, \hat{y}_2, \dots, \hat{y}_n\}$ を与える明示的なモデル分布」を「ターゲット文 $T = \{t_1, t_2, \dots, t_n\}$ を与える暗黙的なデータ分布」に近似させる問題である。近似には様々な方法があり、以降のセクションで詳細を述べる。

2.2 基幹技術

このセクションではニューラル対話生成の基幹技術について解説する。

2.2.1 Sequence-to-Sequence

機械翻訳においてエンコーダとデコーダにリカレント層を用いた系列変換モデルを RNN Encoder-Decoder (Cho et al., 2014) もしくは Sequence-to-Sequence (Seq2Seq, S2S) (Sutskever et al., 2014) と呼ぶ (図 2.1)。一般にソース文をエンコードした LSTM の各層の出力 (Hidden ベクトルと Cell ベクトル) を生成文をデコードする LSTM の各層の初期値とする。

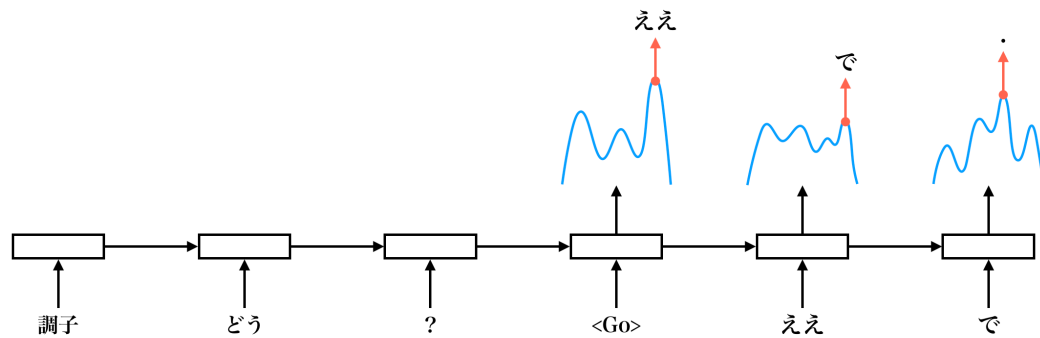


図 2.1 RNN エンコーダにトークン系列「調子」「どう」「?」を入力し, RNN デコーダからトークン系列「ええ」「で」「。」を出力する様子. 一般に訓練時の RNN デコーダは直近のターゲットトークンを入力 (Teacher Forcing) し, 評価時の RNN デコーダは直近の生成トークンを入力する.

エンコーダに双方向 (Bidirectional) RNN を使用することで文全体の特徴表現を抽出する性能が向上される. またリカレント (Recurrent) 層より畳み込み (Convolution) 層を用いた ConvS2S (Gehring et al., 2017) や自己注意 (Self-Attention) 層を用いた Transformer (Vaswani et al., 2017) が高い性能を示すことが報告されている. いずれのモデルも深層であるほど性能が高く, 層間に残差接続 (Residual Connection) を追加することで深層ニューラルネットワークの勾配消失が緩和される. このような Seq2Seq はニューラル対話生成においても広く用いら

れる (Vinyals and Le, 2015; Sordoni et al., 2015; Shang et al., 2015). しかし低い多様性や対話履歴の考慮といった対話生成に固有の問題が見出され, 単にデータを入れ替えるだけではうまく機能しないことが報告されている (Li et al., 2016a; Serban et al., 2016).

2.2.2 Attention 及び Transformer

機械翻訳においてエンコーダの出力をメモリに保持し, デコーダがメモリを参照する機構を Source-Target Attention (ST-Attn) もしくは単に Attention (Attn) (Bahdanau et al., 2014) と呼ぶ (図 2.2).

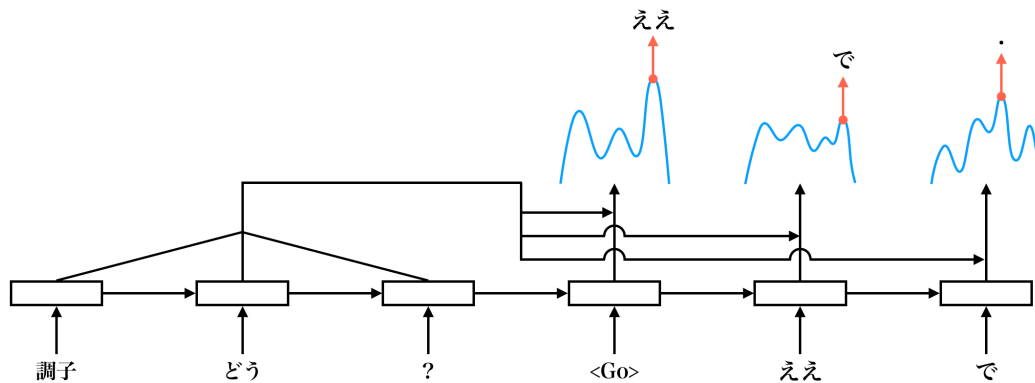


図 2.2 RNN がエンコードしたトークン系列「調子」「どう」「?」をメモリと定義し, RNN デコーダが各タイムステップでメモリを参照し, 文脈情報を取り出し内部表現に追加する. デコーダはソース文の情報を Seq2Seq のような単一のベクトルではなく, トークン系列分の行列として受け取れるので, より柔軟に生成できる.

Attention の利点はモデルの性能を改善するだけでない. 一般に Seq2Seq はデコーディングの処理が進むほどソース文の制約が弱くなり, 言語モデルのように振る舞う. Attention の逐次的なメモリ参照は潜在表現を大きく変容させるので, デコーダが言語モデルのように振る舞うことを防ぎ多様性を促進する働きがある (Zhou et al., 2017; Shao et al., 2017).

一方、下層の出力をメモリに見立て参照する Attention を特に Self-Attention と呼び、また Attention を複数の部分空間に分解した機構を Multi-head Attention と呼ぶ。Self-Attention 及び Source-Target Attention を Multi-head Attention 化したブロックを多層化した Transformer (Vaswani et al., 2017) は機械翻訳や多数の自然言語処理タスク (Devlin et al., 2018) で高い性能が認められる。

2.2.3 MemN2N 及び HRED

対話生成では対話履歴（端的には複数ターンのソース文）を考慮した文脈的な一貫性が要求される。代表的な手法として (i) 対話履歴をメモリに格納する方法と (ii) 対話履歴を階層的ネットワークで処理する方法の 2 通りがある。

(i) は Memory Network (MemN2N) (Sukhbaatar et al., 2015; Miller et al., 2016) と Neural Turing Machine (NTM) (Graves et al., 2014) に大別され、MemN2N は複数のソース文を個別に文ベクトルにエンコードしてそれぞれメモリスロットに格納し Attention でメモリから文脈情報を抽出する。この操作は多層化可能であるが、エンコードの方法次第で計算量が過度に大きくなる。Sukhbaatar et al. (2015) らは Positional Encoding を用いた加重平均で語順を考慮したが、より洗練された方法として RNN でエンコードする方法が検討できる (Nakamura et al., 2018)。また Temporal Encoding を用いてソース文の時系列を考慮できる。一方、NTM は Attention で逐次的に各トークンをメモリにエンコードし、別の Attention で逐次的にメモリから文脈情報を抽出する。

(ii) は Hierarchical Recurrent Encoder-Decoder (HRED) (Serban et al., 2016; 2017b) として知られ、Encoder RNN で複数のソース文を個別に文ベクトルにエンコードして、Context RNN で複数の文ベクトルを 1 つの文脈ベクトルを得る。

対話履歴を利用することで品質は大きく向上し、多様性もやや改善する。また外部資源を利用することでも品質及び多様性の改善が報告されている (Wen et al., 2015b; Ghazvininejad et al., 2017)。

2.2.4 Greedy 探索及び Beam 探索

評価時のデコーダで生起確率が最大のトークンやフレーズを選択することを MAP (Maximum A-Posteriori) 予測と呼ぶ。MAP 予測は最大事後確率における系列予測 (prediction) でありパラメータ推定 (estimation) とは異なる。Greedy 探索及び Beam 探索は MAP 予測をする代表的な手法である。MAP 予測はモデル分布 (デコーダが出力する確率分布) の頂点のみ選択するので、尤度の高い生成文を探索できるが、真のデータ分布を適切にカバーしない規則的かつ変哲もない生成文になる問題がある (図 2.3)。

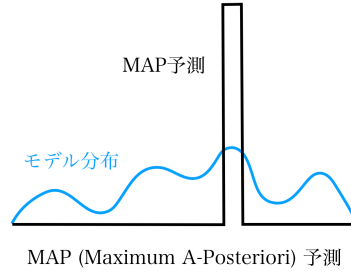


図 2.3 MAP 予測はモデル分布の頂点のみ選択する。

Greedy 探索は argmax 関数によって生起確率が最大のトークンを逐次的に選択する。

$$\hat{y}_i = \text{argmax} Pr(y_i | \hat{y}_{<i}, X) \quad (2.2.1)$$

一方, Beam 探索は topk 関数によって同時確率が上位 k 以内のトークン系列を保持し, 最終的に尤度が最大のトークン系列を選択し生成文とする。

$$\hat{y}_{ik} = \text{topk} \left(\prod_{j=1}^i Pr(y_j | \hat{y}_{<jk}, X) \right) \quad (2.2.2)$$

Beam 探索では同時確率以外のスコア関数を用いて多様性の高い生成文をランキングで選択する方法が報告されている (Li et al., 2016c; Vijayakumar et al.,

2018). また複数の生成文を生成しリランキングで選択する方法も報告されている (Wen et al., 2015a; Li et al., 2016a; Serban et al., 2017a).

2.3 先行研究 (ベースライン)

このセクションでは多様性の促進にフォーカスした先行研究について述べる.

2.3.1 Sampling 探索及び Top10 探索

MAP 予測はモデル分布の頂点のみ選択し, 真のデータ分布を適切にカバーしないため多様性が低下する問題があった. 尤度の低い生成文であっても, データ分布を適切にカバーすることを Non-MAP 予測と呼ぶ. Sampling 探索及び Top10 探索は Non-MAP 予測をする代表的な手法である (図 2.4).

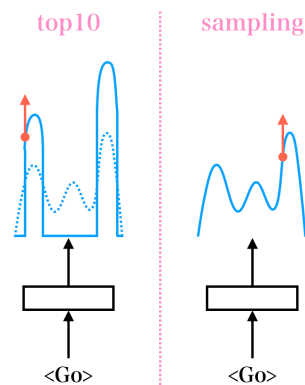


図 2.4 Sampling 探索はモデル分布からサンプリングし, Top10 探索はモデル分布から上位 10 トークンを取り出し正規化後にサンプリングする.

Sampling 探索 (Serban et al., 2016) はモデル分布の生起確率で重み付けしたサンプリングを用いて確率的にトークンを選択する. Softmax 関数の温度 (temperature) によってモデル分布の滑らかさを調節できる.

$$\hat{y}_i \sim Pr(y_i | \hat{y}_{<i}, X) \quad (2.3.1)$$

モデル分布が真のデータ分布に近似していれば、「モデル分布からサンプリングした生成文」は「真のデータ分布からサンプリングされたターゲット文」と同程度の多様性をもつことが予想される。しかし、サンプリングは確率的であるため不適切なトークンが選ばれやすくなるので構文の破綻が懸念される。

Top10 探索 (Edunov et al., 2018) は topk 関数によって生起確率が上位 10 のトークンを抜き出し、Softmax 関数で確率分布に正規化し、Sampling 探索を行う。よって MAP 予測と Sampling 探索の中庸と言える。

$$\hat{y}_i \sim Pr(\text{topk}(y_i k) | \hat{y}_{<i}, X) \quad (2.3.2)$$

機械翻訳において逆翻訳モデルで擬似対訳文 (ターゲット文と擬似ソース文) を生成し順翻訳モデルの学習データを拡張する手法が提案されている (Sennrich et al., 2016a)。このとき、逆翻訳モデルの評価時に Sampling 探索を用いて擬似ソース文の多様性を促進することで、順翻訳モデルの性能が大幅に向上することが報告されている (Imamura et al., 2018; Edunov et al., 2018)。

2.3.2 Maximum Mutual Information

Li et al. (2016a) らは最大相互情報量 (Maximum Mutual Information, MMI) に基づく目的関数を提案し、初めて低い多様性問題にフォーカスして取り組んだ。MMI はソース文と生成文の間の相互情報量を最大化する目的関数であり、実践的には評価時に言語モデル的に妥当な生成トークンをペナライズすることで多様性を促進する。MMI-antiLM は次式によって表される。

$$\hat{y}_i = \text{argmax} \{Pr(y_i | \hat{y}_{<i}, X) - \lambda Pr(y_i | \hat{y}_{<i})\} \quad (2.3.3)$$

ここで $Pr(y_i | \hat{y}_{<i}, X)$ はソース文 X を与えられた時の条件付き確率であり、 $Pr(y_i | \hat{y}_{<i})$ はソース文 X を与えられない時の言語モデルによる確率である。

MMI-antiLM はこのように言語モデルの項を減算することで、言語モデルのような生成を抑制する。

2.4 先行研究 (その他)

2.4.1 強化学習

持続的な対話生成のために、(Li et al., 2016d) らは強化学習ベースの手法を報告している。多様性を促進するために、ソース文と生成文の間の負のコサイン類似度を報酬として与えている。ただし、実験結果から多様性の促進が微々たるものであることが窺える。これまでに様々な報酬が検討されている。例えば Alexa のユーザースコア (Serban et al., 2017a), トークン頻度 (Elbayad et al., 2018), IT-IDF (Yi et al., 2018) などの報酬を用いた強化学習によって多様性の促進が報告されている。

2.4.2 生成モデル

近年は様々な生成モデルが低い多様性問題に取り組んでいる。Kingma and Welling (2013) らは Variational Auto-Encoder (VAE) を提案し、画像生成という分野を開拓した。その後、テキスト生成 (Bowman et al., 2016) や対話生成 (Serban et al., 2017b; Cao and Clark, 2017; Zhao et al., 2017) に適用されている。

Goodfellow et al. (2014) らは Generative Adversarial Network (GAN) を提案し、画像生成の高画質に貢献した。その後、テキスト生成 (Yu et al., 2017) や対話生成 (Li et al., 2017; Xu et al., 2018; Olabiyi et al., 2018) に適用されている。現在、GAN による対話生成は学習が極めて不安定であり、一般に事前学習 (Pre-training) を必要とする。ただし文字レベルのテキスト生成では言語モデルのように識別器を用いることで事前学習をせず訓練できることが報告されている (Press et al., 2017)。

第 3 章 損失関数

一般的に、損失関数はすべてのタイムステップにわたってモデル分布のターゲットトークンクラスの確率 $Pr(t_i|t_{<i}, X)$ とターゲットトークン t_i の間の損失を個別に計算し、損失の総和を求める。このセクションでは任意のタイムステップ i に関する損失についてのべる。

3.1 既存手法 : Softmax Cross-Entropy 損失

MLE で Seq2Seq を訓練する際に一般的に使用される Softmax Cross-Entropy (SCE) 損失は次式によって表される。

$$L_{\text{sce}} = -\log \left(\frac{\exp(o_c)}{\sum_k^{|V|} \exp(o_k)} \right) \quad (3.1.1)$$

ここで o_k はデコーダの出力層の出力 $o \in \mathbb{R}^{|V|}$ の k 番目の要素であり、 c はターゲットトークンクラスのインデックスである。SCE 損失は各トークンクラスを同等に扱うが、頻出トークンはモデルに多くの訓練ペナルティをもたらすため、頻出トークンの生成確率は過度に大きくなり、希少トークンの生成確率は過度に小さくなるバイアスが生じる。その結果、膨大なトークン候補が考えうるにもかかわらず、モデルは頻出トークンばかり選択してしまう。

3.2 提案手法 : Inverse Token Frequency 損失

SCE 損失のバイアスに対処し応答多様性を促進するために、SCE 損失にトークン頻度の逆数に基づく重みをかけた Inverse Token Frequency (ITF) 損失という新たな損失関数を提案する。この損失関数は頻出語クラスの損失を減少する (図 3.1) ことで、SCE 損失が頻出応答に対して過度な生起確率を与える現象を抑制する。

ITF 損失は MMI 推論のように直接モデル分布を操作するのではなく、各トークンクラスの勾配 (訓練ペナルティの供給量) を調節し、希少トークンの学習を加速

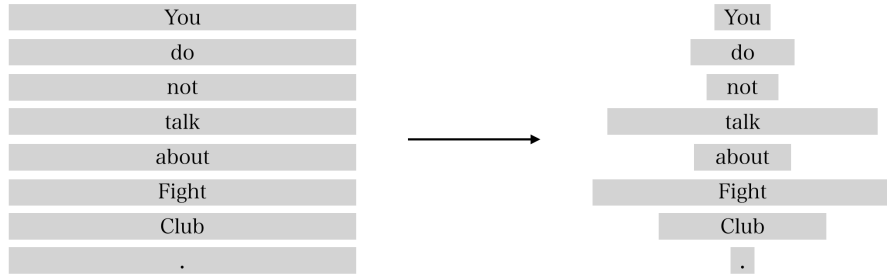


図 3.1 各損失関数が「You do not talk about Fight Club.」というトークン系列に対して重み付けした例である．左は SCE 損失の各トークンクラスの重みであり，全てのトークンクラスで等しい（通常すべての重みは 1 である）．右は ITF 損失の各トークンクラスの重みであり，頻出トークンほど重みが小さい．

させ，頻出トークンの学習を減速させる．よって，ITF 損失は SCE 損失で得られるモデル分布よりも緩やかなモデル分布を描くので，Greedy 探索でも十分に多様な応答生成が可能になる．

ITF 損失は次のような利点もうむ．

- ITF 損失のコンセプトは非常に明確で，理解しやすい．
- ITF 損失は重みつき SCE 損失という単純な機構であるので簡単に訓練スクリプトに組み込むことができる．一方で，MMI, GAN, VAE, RL はモデルを修正または追加する手間が必要があり実装が複雑になる．
- ITF 損失の訓練は MLE の訓練と同程度に安定する．GAN や RL の訓練は通常不安定であり，しばしば MLE の事前訓練 (Pre-training) が必要である．

ITF 損失は SCE 損失のトークン頻度重み付け版と捉えられ，次式によって表される．

$$L_{\text{ITF}} = w_c L_{\text{sce}} \quad (3.2.1)$$

$$w_c = \frac{1}{\text{freq}(\text{token}_c)^\lambda} \quad (3.2.2)$$

ここで w_c は重み $w \in \mathbb{R}^{|V|}$ のうちクラス c に対応する要素であり， token_c はクラス c に対応するトークンである．関数 $\text{freq}(\text{token}_c)$ は token_c が訓練セットに登場

する回数である． λ は頻度の程度を調整するハイパパラメータである． $\lambda = 0$ の時，ITF 損失は SCE 損失と同値である．ソフトマックス層が出力するモデル分布は訓練時と推論時で同じである（対照的に MMI は訓練時と推論時で異なる）． $\langle \text{SoS} \rangle$ や $\langle \text{EoS} \rangle$ (starting of sentence と ending of sentence の略記) などの特殊記号は他の通常トークンと同様に処理する．すなわち， $\langle \text{SoS} \rangle$ や $\langle \text{EoS} \rangle$ は訓練セットの全文で出現するので，それらのトークンの ITF 損失の重みは極めて小さな値をとる．

第 4 章 実験設定

日本語 Twitter データを用いて提案手法の ITF 損失と先行研究の手法の生成文の品質及び多様性を比較する実験を行った．先行研究の手法は Sampling 探索, Top10 探索, MMI 推論の 3 手法である．この章では実験に使用したデータセット, ベースラインの詳細, 訓練設定, 評価方法について説明する．

4.1 データセット

本研究では日本語の Twitter リプライを収集しデータセットを構築した．事前に自己リプライ対話, ボット同士の対話, 過度に長い対話をデータセットから除外した．訓練セットは 4.5M 発話文と 0.7M 対話を含み, 1 対話は平均約 6 発話文から構される．また, 1 epoch あたり 3.8M 標本数で訓練する．検証セットとテストセットはそれぞれ 1024 標本数を用意した．

4.2 ベースライン

提案手法の ITF 損失と比較するベースライン手法として, 人間, Greedy 探索, Sampling 探索, Top10 探索, MMI 推論を使用した．

各手法で共通のニューラル対話生成システムとして, 直近の 8 発話以内を連結した長文を受け取る Long Seq2Seq (LS2S) を使用した．エンコーダのすべての層を Bidirectional LSTM (6 層 x 2 方向), デコーダのすべての層を LSTM (6 層) を採用し, すべての層間に残差接続, 比率 0.1 の Dropout, 層正規化 (Layer Normalization) を適用した．トークンの埋め込みには Bert (Devlin et al., 2018) で提案された埋め込み層 (Token Embedding, Position Embedding, Segment Embedding の加算) を採用した．Token Embedding はトークンに応じた埋め込みベクトル, Position Embedding はトークンの位置に応じた埋め込みベクトル, Segment Embedding は文の位置に応じた埋め込みベクトルを提供する．

MMI は訓練時に使用しても多様性は促進されない．Li et al. (2016a) らの方法に倣い, 訓練時は MLE を使用し, 評価時に MMI 推論を使用する．MMI 推論は次

式によって表される.

$$\hat{y}_i = \operatorname{argmax} \{ \log \operatorname{softmax}(o_i - \lambda u_i) \} \quad (4.2.1)$$

ここで $o \in \mathbb{R}^{|V|}$ はソース文が与えられた時のデコーダの出力層の出力であり, $u \in \mathbb{R}^{|V|}$ はソース文が与えられなかった時のデコーダの出力層の出力である. 実践的には, デコーダの LSTM の隠れ状態を 0 で初期化した. 事前実験では $\log \operatorname{softmax}(o_i) - \log \operatorname{softmax}(\lambda u_i)$ や $\log(\operatorname{softmax}(o_i) - \operatorname{softmax}(\lambda u_i))$ などの他の形式では正常に動作しなかった. 係数 λ は anti-LM 項の強弱を調整するハイパパラメータである. Li et al. (2016a) らに倣い, デコーダの最初の 5 タイムステップのみ MMI 推論を適用した ($\gamma = 5$).

4.3 訓練設定

ニューラル対話生成システムの訓練前に, 訓練セットで 32k 語彙数の Sentencepiece(Kudo and Richardson, 2018) を学習し, Sentencepiece モデルを用いて訓練セットのソース文とターゲット文を subword にトークナイズした. またモデルが生成したトークン系列をデトークナイズして生成文を作成した.

断りがない限り, 我々はすべての手法で次のハイパパラメータを設定した. ハイパパラメータはバッチサイズ: 32, 語彙数: 32k, オプティマイザ: Adam, 学習率: $3e-4$, 隠れ層サイズ: 256, 最大系列長: 32 トークン x8 発話, ソース文数 (メモリサイズ): 8, ヘッド数 (Multi-head Attention): 8, 層数: 6, ホップ数 (MemN2M): 1, Dropout 率: 0.1 とした.

4.4 人手評価

機械翻訳の自動評価では一般的に BLEU (Papineni et al., 2002) が使用される. しかし対話生成では現時点で一般的に使用される評価指標が決まっていない. 特に対話生成では BLEU と人手評価に相関がみられないことが報告されている (Liu et al., 2016). 一般に応答が一意に定まらない対話生成ではターゲット文との n-gram 適合率をもって生成文の良し悪しを判断することが難しい. 本研究の事

前実験では BLEU-4 は 2 以下の著しく低いスコアを示し、モデルの初期値やミニバッチのサンプリングによってスコアが大きく変動した。そこで本研究では以下の実験手順で人手評価を実施した。図 4.1 は実際に使用した評価実験のフォーマットである。

15 名の評価者に Twitter における 2 人の話者間の対話を提示した。各対話の発話は交互に行われる。評価者は対話の最後続く応答文として (1) もっとも魅力的な応答だと思うもの、(2) もっとも首尾一貫した応答だと思うものを 6 つの選択肢から選択した。ただし (1), (2) の選択は重複が許容されている。6 つの選択肢のうち 1 つは人間の発話文、他の 5 つは異なる手法の生成文である。5 手法は Greedy 探索, Sampling 探索, Top10 探索, MMI 推論, ITF 損失を採用した。評価者は全部で 10 対話を評価した。各対話はシャッフルされている。10 対話はテストセットからランダムに選定した。ただし人間の応答が「おはよう」や「ありがとう」のような頻出フレーズになる対話は除外した。

4.5 自動評価

本研究では日本語 Twitter データに最適なトークナイザーを発見できなかったため、訓練セットで学習した Sentencepiece モデルを用いてトークナイズし、subword 系列に対して自動評価スコアを計算した。なお、評価時は生成文とターゲット文から $\langle \text{SoS} \rangle$, $\langle \text{EoS} \rangle$, $\langle \text{Pad} \rangle$, $\langle \text{Unk} \rangle$ (padding と unknown の略記) 等の特殊記号を除外した。我々は生成文の自動評価メトリックスとして流暢さ評価に Perplexity (PPL), 品質評価に BLEU, 多様性評価に DIST, 系列長評価に Length, 反復度評価に Repeat を採用した。Length (すべての生成文の系列長の平均値) 以外のメトリックスの詳細は以下のセクションで述べる。

4.5.1 Perplexity

Perplexity (PPL) は Teacher forcing 時のモデル分布のうちターゲットトークンクラス c の正規確率 p_c の負の対数 $-\log p_c$ をすべてのタイムステップに渡って平均した値 $-\sum^N \frac{1}{N} \log p_c$ (これをエントロピー H といい, Cross-Entropy Loss と

雑談対話システムの評価実験

M2の中村良です。これは修士論文のための評価実験です。
これからTwitterにおける2人の話者間の対話を提示します。発話は交互に行われます。
対話の最後に続く応答文として(1) もっとも「魅力的な応答」だと思うもの、(2) もっとも「首尾一貫した応答」だと思うものを6つの選択肢から選択してください。ただし(1),(2)の選択は重複しても構いません。
6つの選択肢のうち1つは人間の発話文、他の5つは異なる対話システムの生成文です。
対話は全部で10対話です。各対話で選択肢はシャッフルしています。
評価実験終了後、人間の応答の選択肢を公表します。
お忙しいところ恐縮ですが、よろしくお願いします。

対話1

A 「バイトにバイクで、行く」
B 「おはようございます😊 行ってらっしゃい!頑張ってくださいね!」
A 「おはようございます!バイトで家出るのは13時半ぐらいですが、頑張ってきます!」
B 「では午前中はゆっくり休んでください、」

もっとも魅力的な応答を選択してください。 *

- ☐ A 「明日は学校なので、帰りたくないんですよねwあと6時ですね!」
- ☐ A 「ありがとうございます!」
- ☐ A 「はい!ゆっくり寝たりしときますw」
- ☐ A 「夜勤なので今から帰る日に目が覚めてきました!」
- ☐ A 「午前中まで仕事です。」
- ☐ A 「ありがとうございます!午前中は、仕事なのでまだ時間がないので...」

図 4.1 実際に使用した人手評価のフォーマットである。評価者は各対話ごとに「もっとも魅力的な応答」と「もっとも首尾一貫した応答」を選択した。

同値) を求め、すべての生成文のパープレキシティ $\exp(H)$ の平均値を計算する。ここで Teacher forcing は条件付き確率の役割を果たす。Perplexity が意味するところは、ターゲットトークンはもっとも流暢であると仮定し、モデル分布がター

ゲットトークンクラスに高い正規確率を割り当てるほど生成文は流暢ということである。したがって Perplexity のスコアが低いほど流暢であると言える。

ちなみに、言語モデルが損失関数としてしばしば Softmax Cross-Entropy 損失を採用する理由はパープレキシティに対して最適化するためである。裏を返せば ITF 損失は通常のパープレキシティに対して最適化せず、トークン頻度の逆数で重み付けしたパープレキシティに対して最適化する。またパープレキシティからわかるフィッティング性能と生成文の良し悪しは必ずしも相関しないことに留意されたい。

4.5.2 BLEU

BLEU (Papineni et al., 2002) はすべての生成文とすべてのターゲット文の n -gram 適合率を計算する。すなわち、BLEU-1 は unigram 適合率、BLEU-2 は unigram 適合率の 0.5 乗と bigram 適合率の 0.5 乗の積を測り、スコアが高いほど品質が良いと言える。 n -gram 適合率はその名の通り適合率 (precision) のみ評価し、再現率 (recall) を考慮しないため短文を生成する方が有利であるので、系列長に応じてスコアをペナライズする Brevity Penalty (BP) を適用した。また適合数が 0 の n -gram 項に 0.1 を加えるスムージングを適用した。

いくつかの対話生成の研究では機械翻訳で一般的な BLEU-4 を使用するが、我々の事前実験ではモデルや学習手法に関わらず BLEU-4 は 2 以下の著しく低いスコアを示した。これは対話生成は機械翻訳に比べて考えうる生成候補が膨大であるため、ターゲット文と高階の n -gram において適合することが困難であるからで、我々はモデルの初期値やミニバッチのサンプリング差異に依存して容易にスコアが上下する現象を確認した。すなわち、BLEU-4 は対話生成の評価として不適切である。

また前述の通り、対話生成では BLEU と人手評価に相関がみられないことが報告されている (Liu et al., 2016)。

4.5.3 DIST

DIST (Li et al., 2016a) はすべての生成文に含まれる n -gram のうち異なる n -gram の割合を計算する。すなわち、DIST-1 は異なる unigram の割合、DIST-2 は異なる bigram の割合を測り、スコアが高いほど多様性が高いと言える。生成トークンの数が少ないほど n -gram が重複する可能性は低くなるので、DIST は生成する標本数が少ないほど有利かつ短文を生成するほど有利であるが、特に対応策を講じていないことを留意されたい。

4.5.4 Repeat

Repeat はすべての生成文のうち重複するトークンを含む生成文の割合を計算する。

事前実験ではデコーダが繰り返し同じフレーズを生成する現象が確認された。この現象は生成文の主観的な品質を大幅に悪化させる。そこで全ての手法で反復サプレッサー (Repetition Suppressor) を適用した。サプレッサーは生成済みのトークンの生起確率を減衰することで、文中に同じトークンやフレーズが再出現しやすい現象を抑制する。サプレッサーは次式によって表される。

$$\text{suppressor}(o_k) = \frac{1}{\{1 + \text{count}(\text{token}_k)\}^\lambda} o_k, \quad (4.5.1)$$

ここで o_k はデコーダの出力層の出力 $o \in \mathbb{R}^{|V|}$ の k 番目の要素であり、 $\text{count}(\text{token}_k)$ は推論時のデコーダが生成済みのトークンに token_k が含まれる回数である。

第 5 章 結果

5.1 人手評価

「もっとも魅力的な応答」の結果を図 5.1, 「もっとも首尾一貫した応答」の結果を図 5.2 に示した.

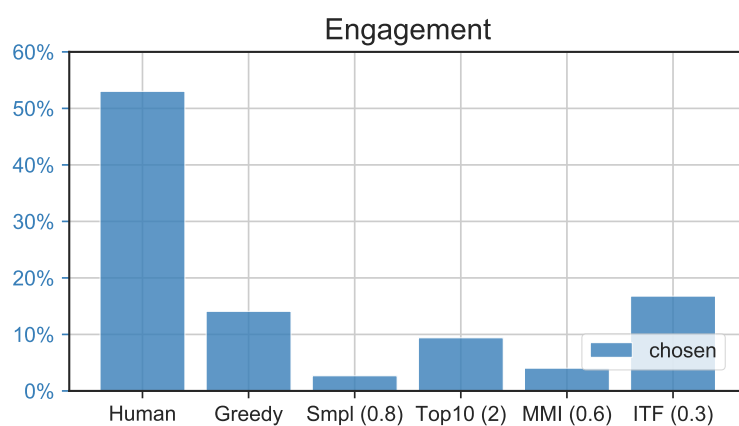


図 5.1 もっとも魅力的な応答

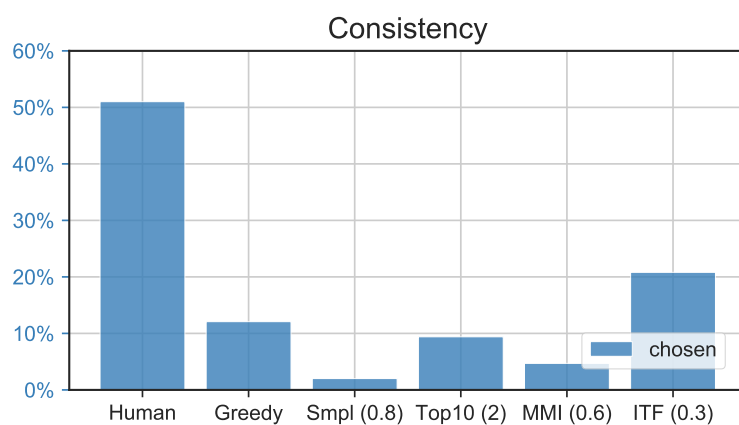


図 5.2 もっとも首尾一貫した応答

それぞれ対話 1~10 のすべての回答を総合した上に人間及び 5 手法の応答が選択された割合である．5 手法は Greedy 探索，Sampling 探索，Top10 探索，MMI 推論，ITF 損失である．

全体では人間の応答が約半数ほど選択され，手法側は提案手法の Inverse Token Frequency (ITF) 損失がもっともよく選択され，次いで Greedy 探索及び Top10 探索がよく選択された．Sampling 探索は後述の DIST-1/2 にて人間の応答並みの n-gram 多様性を示したが，評価者はほとんど選択しなかった．また「もっとも魅力的な応答」と「もっとも首尾一貫した応答」には強い相関があった．

選択肢の応答が短い対話では評価者の選択がばらける傾向があり，選択肢の応答が長い対話では 8 割の評価者が人間の応答を選択する傾向があった．実験後に実施したアンケートによると，長い対話では手法側に明らかな文法的間違いや齟齬が頻繁に起こっているため，評価者はそれらが人間の応答ではないことを見破っていることがわかった．また「魅力」と「首尾一貫」の定義を明確にされなかったため答えづらかったという意見があった．

5.2 自動評価

自動評価メトリックスのスコアを図 5.3 及び図 5.4 に示した．

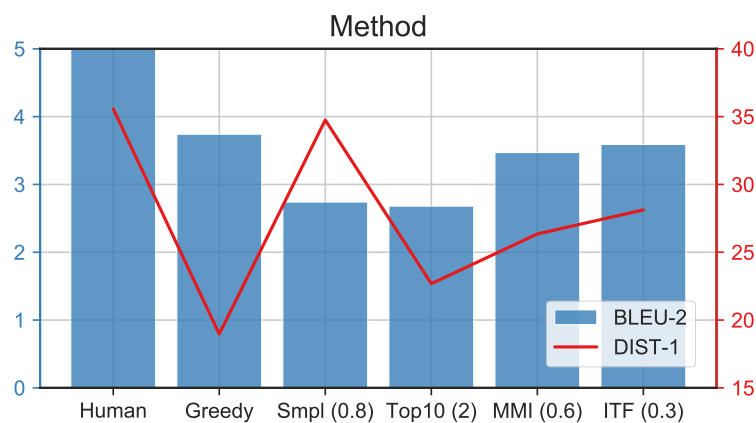


図 5.3 人間及び 5 手法の応答を BLEU-2 及び DIST-1 で評価した結果である．

	Type	PPL	BLEU-1	BLEU-2	DIST-1	DIST-2	Length	Repeat
0	Human	1.00	100.00	100.00	35.58	81.70	13.11	41.70
1	Greedy	130.06	12.29	3.74	18.97	47.41	7.17	3.42
2	Smpl (0.8)	130.06	12.95	2.74	34.75	80.04	9.94	6.84
3	Top10 (2)	130.06	13.89	2.68	22.68	69.14	11.36	10.94
4	MMI (0.6)	130.06	12.64	3.47	26.35	63.46	7.82	4.39
5	ITF (0.3)	215.73	12.30	3.59	28.12	60.13	8.51	9.28

図 5.4 人間及び 5 手法の応答を複数の評価メトリックスで評価した結果である。

Sampling 探索は人間並みの DIST-1/2 を示したが、一方で Greedy 探索より 1 ポイントも低い BLEU-2 を示した。Top10 探索はやや DIST-1/2 を改善し、BLEU-1 はもっとも高かったが、BLEU-2 は Sampling 探索と同程度であった。MMI 推論と ITF 損失は同程度の BLEU-1/2 及び DIST-1/2 を示した。

人手評価と自動評価の相関係数を計算した。

表 5.1 人手評価と自動評価の相関係数

	BLEU-1	BLEU-2	DIST-1	DIST-2
Engagement	-0.38	+0.58	-0.53	-0.72
Consistency	-0.39	+0.55	-0.35	-0.58

人手評価と BLEU-1 には負の相関があり、BLEU-2 には正の相関があった。人手評価と DIST-1 及び DIST-2 には負の相関があった。また人手評価と Length 及び Repeat にはほとんど相関がなかった。よって本研究での自動評価メトリックスには BLEU-2 を参考にするのが適切だと考えられる。

5.2.1 ハイパパラメータ

本研究では各手法の最適なハイパパラメータを探るための比較実験も行った。Sampling 探索及び Top10 探索では次式のように Softmax 関数の温度 (temperature) によってモデル分布の滑らかさを調節できる。温度が高いほどモデル分布は滑らかになる。

$$Pr(y_i) = \text{softmax}(o_i / \text{temperature}) \quad (5.2.1)$$

MMI 推論はハイパパラメータ λ で antiLM の強度を調節でき、ITF 損失はハイパパラメータ λ で ITF の強度を調節できる。 λ が高いほど手法が強く適用される。結果を図 5.5 及び図 5.6 に示した。

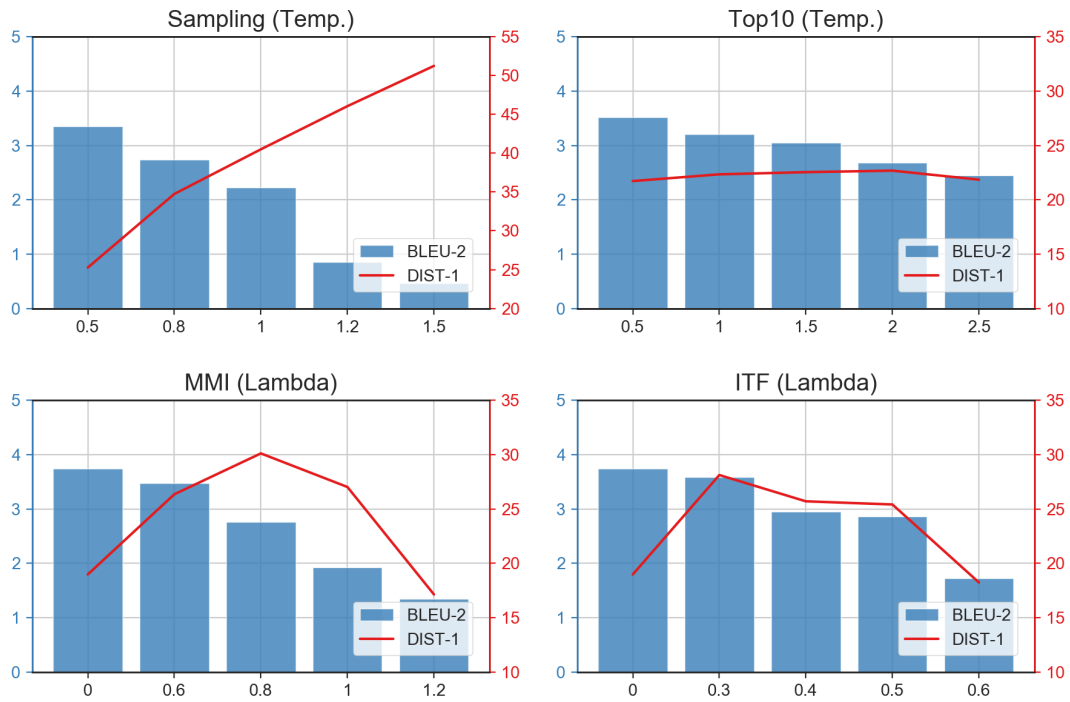


図 5.5 4 手法の異なるハイパパラメータの応答を BLEU-2 及び DIST-1 で評価した結果である。

Sampling 探索は λ が高いほど BLEU-1/2 が低く DIST-1/2 が高かった。すなわち、BLEU と DIST は完全なトレードオフである。一方、Top10 は温度が 2 のとき BLEU-1 が最も高く、温度が高いほど BLEU-2 が低かった。また温度に関わらず DIST-1 は一定で、温度が高いほど DIST-2 が高かった。MMI 推論は λ が高いほど BLEU-1/2 が低く、 λ が 0.8 付近のとき DIST-1/2 が最も高かった。ITF 損失は λ が 0.4 のとき BLEU-1 が最も高く、 λ が高いほど BLEU-2 が低かった。 λ

Sampling								
	Type	PPL	BLEU-1	BLEU-2	DIST-1	DIST-2	Length	Repeat
0	0.5	130.06	12.63	3.35	25.26	63.66	8.00	4.30
1	0.8	130.06	12.95	2.74	34.75	80.04	9.94	6.84
2	1	130.06	12.41	2.22	40.50	87.38	12.55	8.30
3	1.2	130.06	8.51	0.85	46.06	93.96	17.21	10.64
4	1.5	130.06	5.04	0.46	51.25	98.31	24.44	9.96
Top10								
	Type	PPL	BLEU-1	BLEU-2	DIST-1	DIST-2	Length	Repeat
0	0.5	130.06	12.58	3.52	21.72	57.43	7.74	4.30
1	1	130.06	13.20	3.21	22.34	63.03	9.15	6.54
2	1.5	130.06	13.63	3.05	22.54	67.20	10.29	9.96
3	2	130.06	13.89	2.68	22.68	69.14	11.36	10.94
4	2.5	130.06	13.76	2.45	21.85	70.75	12.65	13.77
MMI								
	Type	PPL	BLEU-1	BLEU-2	DIST-1	DIST-2	Length	Repeat
0	0	130.06	12.29	3.74	18.97	47.41	7.17	3.42
1	0.6	130.06	12.64	3.47	26.35	63.46	7.82	4.39
2	0.8	130.06	11.01	2.76	30.11	69.22	7.94	3.81
3	1	130.06	7.65	1.92	27.02	70.28	7.75	4.49
4	1.2	130.06	4.98	1.34	17.13	54.37	7.65	4.00
ITF								
	Type	PPL	BLEU-1	BLEU-2	DIST-1	DIST-2	Length	Repeat
0	0	130.06	12.29	3.74	18.97	47.41	7.17	3.42
1	0.3	215.73	12.30	3.59	28.12	60.13	8.51	9.28
2	0.4	308.02	12.42	2.95	25.70	56.25	9.73	6.64
3	0.5	428.81	12.29	2.86	25.41	57.83	12.37	3.22
4	0.6	702.85	8.71	1.72	18.23	46.29	16.22	8.01

図 5.6 4 手法の異なるハイパパラメータの応答を複数の評価メトリックスで評価した結果である。

が 0.3 のとき DIST-1/2 が最も高かった。

5.2.2 モデル

本研究では直近の 8 発話以内を連結した長文を受け取る Long Seq2Seq (LS2S) を使用したが、様々なモデルの性能比較も行った。比較した 7 モデルは Seq2Seq (S2S), Long Seq2Seq (LS2S), LS2S with Attention (LAttn), Long Transformer (LTrm), Memory Network (Mem), Hred, Hierarchical Transformer (HTrm) で

ある．全てのモデルで Sampling 探索を使用した．

LAttn には Transformer (Vaswani et al., 2017) で提案された Multi-head Attention (8 heads) を採用し，デコーダのすべての層，すべてのタイムステップに適用した．Multi-head Attention は Query, Key, Value, Output ベクトルの線形変換，残差接続，Dropout，Layer 正規化を含む．

Mem は計算量を抑えるため 1 hop とし，直近の発話以外のソース文を Bidirectional LSTM (6 層 x 2 方向) と Temporal Encoding でエンコードしてメモリ (最終層の Hidden と Cell) を形成した．直近の発話のソース文を Bidirectional LSTM でエンコードした Hidden と Cell を Query とし，メモリに対する Attention を個別に計算して Hidden と Cell の Context を得た．最後に Query と Context の和をデコーダの LSTM (6 層) の各層の初期値にした．

Hred は Encoder RNN に Bidirectional LSTM (6 層)，Context RNN にも Bidirectional LSTM (6 層) を採用し，Context RNN の各層の出力をデコードする LSTM (6 層) の各層の初期値にした．

HTrm は HRED の RNN を Transformer に置き換えたモデルである．

結果を図 5.7 及び図 5.8 に示した．

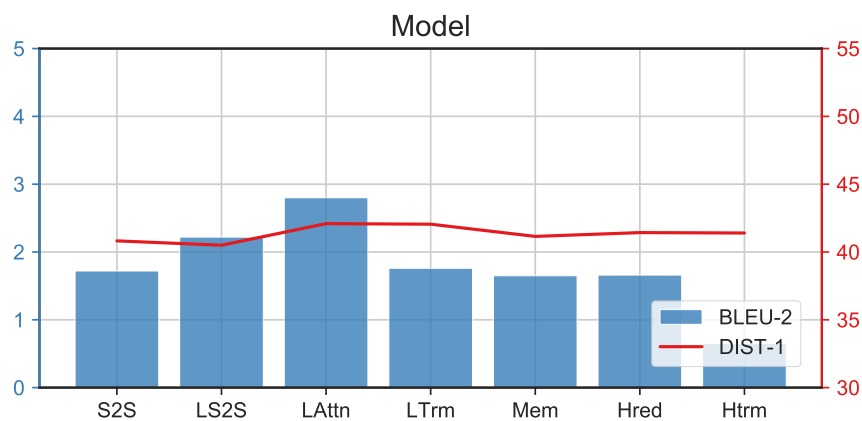


図 5.7 異なるモデルの応答を BLEU-2 及び DIST-1 で評価した結果である．

LAttn はもっとも高い BLEU-1/2 を示し，次いで LS2S が高い BLEU-1/2 を示した．LS2S は S2S より高い BLEU-1/2 を示し，複数発話を考慮することが有意

	Type	PPL	BLEU-1	BLEU-2	DIST-1	DIST-2	Length	Repeat
0	S2S	140.73	11.08	1.72	40.82	86.87	12.62	8.20
1	LS2S	130.06	12.41	2.22	40.50	87.38	12.55	8.30
2	LAttn	116.28	13.02	2.80	42.09	87.58	11.85	2.15
3	LTrm	148.70	11.83	1.76	42.05	88.16	12.40	4.30
4	Mem	141.79	11.85	1.65	41.15	87.66	12.02	10.45
5	Hred	134.04	12.00	1.66	41.43	87.33	12.17	5.27
6	Htrm	192.39	10.03	0.65	41.40	87.19	11.72	7.52

図 5.8 異なるモデルの応答を複数の評価メトリックスで評価した結果である。

義であるとわかった。Sampling 探索の場合、モデル選択に関わらず、人間と同程度の DIST-1/2 を示すことがわかった。

予想に反して LTrm は LS2S より低い BLEU-1/2 を示した。これは Transformer が本研究で使用した日本語 Twitter データ (もしくは多対多変換タスク全般) にそぐわないのか、本研究のように平均系列長が短いデータ全般にそぐわないのか、実装上のミスがあったのか、様々な原因が考えられる。今後、WMT 等の機械翻訳データセットでも評価する必要がある。

Mem 及び Hred は LS2S より低い BLEU-1/2 を示した。すなわち、複数発話を考慮するために Memory Network や HRED 等の専用モデルを用意せずとも、複数発話の 1 文に連結して Segment Embedding 付きで Seq2Seq に入力するだけで十分である可能性が高い。とりわけ Mem や Hred ははじめに文ベクトルを形成し、それから更に文脈ベクトルを形成するため、重要なトークン情報が失われている可能性がある。

5.3 生成例

最後に人手評価で用いた各手法の生成文を幾つか例示する。

対話例 1

A 「今日 cs?」

B 「そす」

A「今大阪いて今日暇だから遊びに行こうかな」
B「おおさかにいてるのはわらう」
A「昨日オタクライブだったので... 何時から?」
B「じぶんは9時くらいに駅すあと、まえと場所が違うす」

human : A「まじ?最寄りどこ?」
greedy : A「なるほど!」
sampling: A「マジか...」
top10 : A「あ!じゃあ今日、1月か8日に帰る?w」
mmi : A「何時に着く?」
itf : A「9時ならワンチャンありますね」

対話例 2

A「そろそろ吐くまで飲みたいぞ」
B「宅飲みはいつ...<kaomoji>←」
A「いつでもいい!明日ね!明日新河岸集合ね!←」
B「そいつぁ無理だ<kaomoji>」

human : A「私も昨日は美容行ってたから無理だったわ (笑) ほんとそろそろやろーよー」
greedy : A「じゃあ、明日か明後日の夜祭で!」
sampling: A「みんな寝るの早いんだよね<kaomoji>)」
top10 : A「お、じゃあ今度ご飯行ってやろう!←」
mmi : A「何時に帰る?」
itf : A「じゃあ、今から行くね!」

対話例 3

A「14<kaomoji>」

B「歳下なの?!え、てっきり年上かと... 笑笑」
A「ええだよねええまだ中2だよ<kaomoji>あかねは <emoji>」
B「中2に見えなっがあかね16歳!高1!」
A「高1とかあーなんか青春してそうだねうん」
B「なーに言ってんの。青春なんてしてないよ泣ける」
A「いやいやワシ大学まで女子校未来ない<kaomoji>」

human : B「なに、お嬢様学校なの?」
greedy : B「え、なにそれ<kaomoji>」
sampling: B「え、まじ!?」
top10 : B「大学受験と高校合格して学校で卒業した!爆笑てか、女子の学年2年
目笑中学3年生に進学するから春休み中。笑笑」
mmi : B「え笑笑女子校?」
itf : B「女子校なのね!」

対話例 4

A「モンダミンの匂い嗅いだらおえってなる」
B「お口くちゅくちゅしろ」
A「おくち... くちゅくちゅ... もんだみん...//」
B「悲しいくらい萌えない...」
A「まじんこ!?!超絶いろっぽいやん!?!」
B「今度会った時やって」
A「出会い厨きつつう!<emoji>」
B「どこ住?w てか LINE やってる?ww」

human : A「らいんじゃなくて、カカオならいいよ><」
greedy : A「え、まって ww」
sampling: A「たじばは公園で北海道参戦する!」
top10 : A「今家!」

mmi : A「んー LINE!」

itf : A「埼玉だよ!」

Greedy 探索は「なるほど!」や「え、まって ww」など文脈依存性が少ない Generic Response を生成する傾向があった。Sampling 探索と Top10 探索は「公園で北海道参戦する!」や「大学受験と高校合格して学校で卒業した!」など意味的に齟齬を含んだ応答を生成する傾向があった。またソース文との繋がりが他の手法に比べて不自然である。MMI 推論はややソース文との繋がりが不自然である。全体的に、MAP 予測以外の手法は文脈的一貫性が悪化する傾向にあった。ITF 損失は Generic Response を避けつつ、Greedy 探索であるためソース文との繋がりが比較的自然であった。

第 6 章 結論

本研究ではニューラル対話生成における低い多様性問題にフォーカスし、トークン頻度の逆数に基づく損失関数を適用することで、生成文の多様性を促進し、先行研究よりも主観的魅力及び文脈一貫性を向上することができた。しかし全ての手法で言えるが、長い応答文の対話では明らかな文法的間違いや齟齬が多くあった。今後、事前学習やデータ拡張によってデータ量を増やし、Transformer の性能を機械翻訳並みに発揮するために WMT 等のデータセットで自動評価を行いたい。また人手評価と BLEU-2 以外の自動評価メトリックスには相関がない、もしくは負の相関があることが確かめられ、今後より詳細な人手評価を実施することを検討したい。

今後、品質を維持しつつ多様性を促進するより洗練された損失関数を模索する予定である。また対話履歴のみならずユーザーの個性を考慮した対話生成、並びに対話システム自体に個性を付加した対話生成 (Li et al., 2016b) などでの「より条件づけられた多様性」への発展を検討している。

謝辞

本論文は筆者が奈良先端科学技術大学院大学情報科学研究科情報科学専攻博士前期課程に在籍中の研究成果をまとめたものである。同専攻教授中村哲先生には指導教員として本研究の実施の機会を与えて頂き、その遂行にあたって終始、ご指導を頂いた。ここに深謝の意を表する。同専攻教授松本裕治先生、並びに同専攻准教授須藤克仁先生、並びに同専攻助教吉野幸一郎先生には副査としてご助言を頂くとともに本論文の細部にわたりご指導を頂いた。ここに深謝の意を表する。本専攻知能コミュニケーション研究室の各位には研究遂行にあたり日頃より有益なご討論ご助言を頂いた。ここに感謝の意を表する。

参考文献

- Dzmitry Bahdanau, Kyunghyun Cho, and Yoshua Bengio. Neural machine translation by jointly learning to align and translate. *arXiv preprint arXiv:1409.0473*, 2014.
- Samuel R Bowman, Luke Vilnis, Oriol Vinyals, Andrew Dai, Rafal Jozefowicz, and Samy Bengio. Generating sentences from a continuous space. In *Proceedings of The 20th SIGNLL Conference on Computational Natural Language Learning*, pages 10–21, 2016.
- Kris Cao and Stephen Clark. Latent variable dialogue models and their diversity. In *Proceedings of the 15th Conference of the European Chapter of the Association for Computational Linguistics (Volume 2, Short Papers)*, volume 2, pages 182–187, 2017.
- Kyunghyun Cho, Bart van Merriënboer, Caglar Gulcehre, Dzmitry Bahdanau, Fethi Bougares, Holger Schwenk, and Yoshua Bengio. Learning phrase representations using rnn encoder–decoder for statistical machine translation. In *Proceedings of the 2014 Conference on Empirical Methods in Natural Language Processing*, pages 1724–1734, 2014.
- Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. Bert: Pre-training of deep bidirectional transformers for language understanding. *arXiv preprint arXiv:1810.04805*, 2018.
- Sergey Edunov, Myle Ott, Michael Auli, and David Grangier. Understanding back-translation at scale. In *Proceedings of the 2018 Conference on Empirical Methods in Natural Language Processing*, pages 489–500, 2018.
- Maha Elbayad, Laurent Besacier, and Jakob Verbeek. Token-level and sequence-level loss smoothing for rnn language models. *arXiv preprint arXiv:1805.05062*, 2018.
- Jonas Gehring, Michael Auli, David Grangier, Denis Yarats, and Yann N Dauphin. Convolutional sequence to sequence learning. *arXiv preprint arXiv:1705.03122*, 2017.

- Marjan Ghazvininejad, Chris Brockett, Ming-Wei Chang, Bill Dolan, Jianfeng Gao, Wen-tau Yih, and Michel Galley. A knowledge-grounded neural conversation model. *arXiv preprint arXiv:1702.01932*, 2017.
- Ian Goodfellow, Jean Pouget-Abadie, Mehdi Mirza, Bing Xu, David Warde-Farley, Sherjil Ozair, Aaron Courville, and Yoshua Bengio. Generative adversarial nets. In *Advances in neural information processing systems*, pages 2672–2680, 2014.
- Alex Graves, Greg Wayne, and Ivo Danihelka. Neural turing machines. *arXiv preprint arXiv:1410.5401*, 2014.
- Kenji Imamura, Atsushi Fujita, and Eiichiro Sumita. Enhancement of encoder and attention using target monolingual corpora in neural machine translation. In *Proceedings of the 2nd Workshop on Neural Machine Translation and Generation*, pages 55–63, 2018.
- Diederik P Kingma and Max Welling. Auto-encoding variational bayes. *arXiv preprint arXiv:1312.6114*, 2013.
- Taku Kudo and John Richardson. Sentencepiece: A simple and language independent subword tokenizer and detokenizer for neural text processing. In *Proceedings of the 2018 Conference on Empirical Methods in Natural Language Processing: System Demonstrations*, pages 66–71, 2018.
- Yann LeCun, Yoshua Bengio, and Geoffrey Hinton. Deep learning. *nature*, 521 (7553):436, 2015.
- Jiwei Li, Michel Galley, Chris Brockett, Jianfeng Gao, and Bill Dolan. A diversity-promoting objective function for neural conversation models. In *Proceedings of the 2016 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*, pages 110–119, 2016a.
- Jiwei Li, Michel Galley, Chris Brockett, Georgios Spithourakis, Jianfeng Gao, and Bill Dolan. A persona-based neural conversation model. In *Proceedings of the 54th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, volume 1, pages 994–1003, 2016b.

- Jiwei Li, Will Monroe, and Dan Jurafsky. A simple, fast diverse decoding algorithm for neural generation. *arXiv preprint arXiv:1611.08562*, 2016c.
- Jiwei Li, Will Monroe, Alan Ritter, Michel Galley, Jianfeng Gao, and Dan Jurafsky. Deep reinforcement learning for dialogue generation. In *Proceedings of the 2016 Conference on Empirical Methods in Natural Language Processing*, pages 1192–1202, 2016d.
- Jiwei Li, Will Monroe, Tianlin Shi, Sébastien Jean, Alan Ritter, and Dan Jurafsky. Adversarial learning for neural dialogue generation. In *Proceedings of the 2017 Conference on Empirical Methods in Natural Language Processing*, pages 2157–2169, 2017.
- Chia-Wei Liu, Ryan Lowe, Iulian Serban, Mike Noseworthy, Laurent Charlin, and Joelle Pineau. How not to evaluate your dialogue system: An empirical study of unsupervised evaluation metrics for dialogue response generation. In *Proceedings of the 2016 Conference on Empirical Methods in Natural Language Processing*, pages 2122–2132, 2016.
- Alexander Miller, Adam Fisch, Jesse Dodge, Amir-Hossein Karimi, Antoine Bordes, and Jason Weston. Key-value memory networks for directly reading documents. In *Proceedings of the 2016 Conference on Empirical Methods in Natural Language Processing*, pages 1400–1409, 2016.
- Ryo Nakamura, Katsuhito Sudoh, Koichiro Yoshino, and Satoshi Nakamura. Another diversity-promoting objective function for neural dialogue generation. 2018.
- Oluwatobi Olabiyi, Alan Salimov, Anish Khazane, and Erik Mueller. Multi-turn dialogue response generation in an adversarial learning framework. *arXiv preprint arXiv:1805.11752*, 2018.
- John D Owens, Mike Houston, David Luebke, Simon Green, John E Stone, and James C Phillips. Gpu computing. *Proceedings of the IEEE*, 96(5):879–899, 2008.
- Kishore Papineni, Salim Roukos, Todd Ward, and Wei-Jing Zhu. Bleu: a method for automatic evaluation of machine translation. In *Proceedings of*

- the 40th Annual Meeting of the Association for Computational Linguistics*, pages 311–318, 2002.
- Ofir Press, Amir Bar, Ben Bogin, Jonathan Berant, and Lior Wolf. Language generation with recurrent generative adversarial networks without pre-training. 2017.
- Rico Sennrich, Barry Haddow, and Alexandra Birch. Improving neural machine translation models with monolingual data. In *Proceedings of the 54th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, volume 1, pages 86–96, 2016a.
- Rico Sennrich, Barry Haddow, and Alexandra Birch. Neural machine translation of rare words with subword units. In *Proceedings of the 54th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, volume 1, pages 1715–1725, 2016b.
- Iulian V Serban, Alessandro Sordoni, Yoshua Bengio, Aaron Courville, and Joelle Pineau. Building end-to-end dialogue systems using generative hierarchical neural network models. In *Thirtieth AAAI Conference on Artificial Intelligence*, 2016.
- Iulian V Serban, Chinnadhurai Sankar, Mathieu Germain, Saizheng Zhang, Zhouhan Lin, Sandeep Subramanian, Taesup Kim, Michael Pieper, Sarath Chandar, Nan Rosemary Ke, et al. A deep reinforcement learning chatbot. *arXiv preprint arXiv:1709.02349*, 2017a.
- Iulian Vlad Serban, Alessandro Sordoni, Ryan Lowe, Laurent Charlin, Joelle Pineau, Aaron C Courville, and Yoshua Bengio. A hierarchical latent variable encoder-decoder model for generating dialogues. In *Thirty-First AAAI Conference on Artificial Intelligence*, 2017b.
- Lifeng Shang, Zhengdong Lu, and Hang Li. Neural responding machine for short-text conversation. In *Proceedings of the 53rd Annual Meeting of the Association for Computational Linguistics and the 7th International Joint Conference on Natural Language Processing (Volume 1: Long Papers)*, volume 1, pages 1577–1586, 2015.

- Yuanlong Shao, Stephan Gouws, Denny Britz, Anna Goldie, Brian Strope, and Ray Kurzweil. Generating high-quality and informative conversation responses with sequence-to-sequence models. In *Proceedings of the 2017 Conference on Empirical Methods in Natural Language Processing*, pages 2210–2219, 2017.
- Alessandro Sordoni, Michel Galley, Michael Auli, Chris Brockett, Yangfeng Ji, Margaret Mitchell, Jian-Yun Nie, Jianfeng Gao, and Bill Dolan. A neural network approach to context-sensitive generation of conversational responses. In *Proceedings of the 2015 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*, pages 196–205, 2015.
- Sainbayar Sukhbaatar, Jason Weston, Rob Fergus, et al. End-to-end memory networks. In *Advances in neural information processing systems*, pages 2440–2448, 2015.
- Ilya Sutskever, Oriol Vinyals, and Quoc V Le. Sequence to sequence learning with neural networks. In *Advances in neural information processing systems*, pages 3104–3112, 2014.
- Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N Gomez, Łukasz Kaiser, and Illia Polosukhin. Attention is all you need. In *Advances in Neural Information Processing Systems*, pages 5998–6008, 2017.
- Ashwin K Vijayakumar, Michael Cogswell, Ramprasath R Selvaraju, Qing Sun, Stefan Lee, David Crandall, and Dhruv Batra. Diverse beam search: Decoding diverse solutions from neural sequence models. 2018.
- Oriol Vinyals and Quoc Le. A neural conversational model. *arXiv preprint arXiv:1506.05869*, 2015.
- Joseph Weizenbaum. Eliza—a computer program for the study of natural language communication between man and machine. *Communications of the ACM*, 9(1):36–45, 1966.
- Tsung-Hsien Wen, Milica Gašić, Dongho Kim, Nikola Mrkšić, Pei-Hao Su, David

- Vandyke, and Steve Young. Stochastic language generation in dialogue using recurrent neural networks with convolutional sentence reranking. In *16th Annual Meeting of the Special Interest Group on Discourse and Dialogue*, page 275, 2015a.
- Tsung-Hsien Wen, Milica Gasic, Nikola Mrksic, Pei-Hao Su, David Vandyke, and Steve Young. Semantically conditioned lstm-based natural language generation for spoken dialogue systems. In *Proceedings of the 2015 Conference on Empirical Methods in Natural Language Processing*, pages 1711–1721, 2015b.
- Jingjing Xu, Xu Sun, Xuancheng Ren, Junyang Lin, Binzhen Wei, and Wei Li. Dp-gan: Diversity-promoting generative adversarial network for generating informative and diversified text. *arXiv preprint arXiv:1802.01345*, 2018.
- Xiaoyuan Yi, Maosong Sun, Ruoyu Li, and Wenhao Li. Automatic poetry generation with mutual reinforcement learning. In *Proceedings of the 2018 Conference on Empirical Methods in Natural Language Processing*, pages 3143–3153, 2018.
- Lantao Yu, Weinan Zhang, Jun Wang, and Yong Yu. Seqgan: Sequence generative adversarial nets with policy gradient. In *Thirty-First AAAI Conference on Artificial Intelligence*, 2017.
- Tiancheng Zhao, Ran Zhao, and Maxine Eskenazi. Learning discourse-level diversity for neural dialog models using conditional variational autoencoders. In *Proceedings of the 55th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, volume 1, pages 654–664, 2017.
- Ganbin Zhou, Ping Luo, Rongyu Cao, Fen Lin, Bo Chen, and Qing He. Mechanism-aware neural machine for dialogue response generation. In *Thirty-First AAAI Conference on Artificial Intelligence*, 2017.

発表リスト

<査読付き国際学会発表>

[1] Ryo Nakamura; Katsuhito Sudoh; Koichiro Yoshino; Satoshi Nakamura. Another Diversity-Promoting Objective Function for Neural Dialogue Generation. In Thirty-Third AAAI Conference on Artificial Intelligence, The Second AAAI Workshop on Reasoning and Learning for Human-Machine Dialogues (DEEP-DIAL 2019). 2019 年 1 月 (7pp. オーラル発表)