



Title	対話システムにおける人間とのコミュニケーションの円滑化に関する研究
Author(s)	高山, 隼矢
Citation	大阪大学, 2022, 博士論文
Version Type	VoR
URL	https://doi.org/10.18910/88153
rights	
Note	

The University of Osaka Institutional Knowledge Archive : OUKA

<https://ir.library.osaka-u.ac.jp/>

The University of Osaka

論文内容の要旨

氏 名 （ 高山 隼矢 ）	
論文題名	対話システムにおける人間とのコミュニケーションの円滑化に関する研究
論文内容の要旨 言語能力を持つ人間にとって、自然言語を用いた対話は、円滑にコミュニケーションを行うための容易かつ効果的な手段であると言える。そのため、対話によって人間と情報システムのコミュニケーションを媒介する対話システム技術への注目度は高い。近年では、対話応答生成モデルと呼ばれる深層学習に基づく対話システム構築手法の登場により、大量の対話データを用いて容易に対話システムを構築できるようになった。 対話応答生成モデルは流暢な応答を生成できる利点を持つ一方で、抽象的な発話に対してその意図や要求を解釈できないという問題がある。また、しばしば「そうですね」や「わかりません」といったどのような発話に対しても成り立つ汎用かつ具体的でない応答を生成してしまう問題、さらに対話を破綻させるような応答を生成してしまうなどの問題がある。これらの問題は人間との円滑なコミュニケーションを妨げるものとなる。本論文では発話や応答の具体性に関するこれらの課題を解決するために次の 3 つの研究に取り組む：(1) 抽象的な発話の具体的な発話への事前言い換えに基づく対話応答生成、(2) 応答の具体性を制御可能な対話応答生成モデルの構築、(3) 対話破綻検出。 (1) では、人間はしばしば対話において自らの意図や要求を明示しない抽象的な発話を行うことに着目し、抽象的な発話の意図を理解し適切な応答を生成する手法の構築に取り組む。まず抽象的な発話と具体的な発話の対を含む 7 万件規模の対話コーパス DIRECT (Direct and Indirect REsponses in Conversational Text corpus) を構築する。また、対話応答生成モデルへの入力発話を事前により具体的な発話に言い換える手法を提案する。実験においては提案手法の適用によって通常の対話応答生成モデルよりも発話の意図に即した応答をより多く生成できるようになることを確認した。 (2) では、対話応答生成モデルがしばしば具体性の低い発話を生成してしまうという問題に対し、具体性の高い応答を生成可能な技術の実現を目的として研究を行う。人間は常に具体性の高い応答を行うのではなく必要に応じて応答の具体性を使い分けていることに着目し、具体性を制御可能な対話応答生成手法の構築に取り組む。発話に対する応答の具体性を測る簡潔な尺度 MaxPMI を提案し、そのスコアを制御変数として訓練データに付与して訓練することで、具体性の制御を可能にする。評価実験では、提案手法が先行研究に比べてより応答の具体性をより広い範囲で制御できることを、自動評価・人手評価によって検証・確認した。 (3) では、対話システムが生成した破綻応答を検知する対話破綻検出タスクに取り組む。対話破綻を感じる基準は人によって異なることに着目し、破綻ラベルのアノテーション傾向に基づいて訓練データのアノテーションをクラスタリングし、様々なアノテーション傾向を再現する複数の検出モデルを構築してアンサンブルする手法を提案する。実験においては、他のアンサンブル手法や従来の対話破綻検出手法に比べて提案手法が高い検出性能を達成できることを確認した。また、(2) において提案した対話応答生成モデルの生成結果に対して対話破綻検出モデルを適用した場合の性能比較も行い、提案手法の優位性を示した。 本論文では、提案したこれらの手法によって対話応答生成モデルや対話破綻検出器の性能向上を達成し、人間との円滑なコミュニケーションが可能な対話システムの構築に寄与した。	

論文審査の結果の要旨及び担当者

氏 名 (高 山 隼 矢)			
	(職)	氏 名	
論文審査担当者	主査	准教授	荒瀬 由紀
	副査	教授	下條 真司
	副査	教授	原 隆浩
	副査	教授	鬼塚 真
	副査	教授	藤原 融
	副査	教授	松下 康之
	副査	教授	春本 要

論文審査の結果の要旨

対話応答生成モデルと呼ばれる深層学習に基づく対話システム構築手法の登場により、大量の対話データを用いて容易に対話システムを構築できるようになった。対話応答生成モデルは流暢な応答を生成できる利点を持つ一方で、抽象的な発話に対してその意図や要求を解釈できないという問題がある。また、しばしば「そうですね」や「わかりません」といったどのような発話に対しても成り立つ汎用かつ具体的でない応答を生成してしまう問題、さらに対話を破綻させるような応答を生成してしまうなどの問題がある。これらの問題は人間との円滑なコミュニケーションを妨げるものとなる。本論文では発話や応答の具体性に関するこれらの課題を解決するため：(1) 抽象的な発話の具体的な発話への事前言い換えに基づく対話応答生成、(2) 応答の具体性を制御可能な対話応答生成モデルの構築、(3) 対話破綻検出の3つの研究に取り組む。

(1) では、人間はしばしば対話において自らの意図や要求を明示しない抽象的な発話を行うことに着目し、抽象的な発話の意図を理解し適切な応答を生成する手法の構築に取り組む。まず抽象的な発話と具体的な発話の対を含む7万件規模の対話コーパス DIRECT (Direct and Indirect REsponses in Conversational Text corpus) を構築する。また、対話応答生成モデルへの入力発話を事前により具体的な発話に言い換える手法を提案する。実験においては提案手法の適用によって通常の対話応答生成モデルよりも発話の意図に即した応答をより多く生成できるようになることを確認した。

(2) では、対話応答生成モデルがしばしば具体性の低い発話を生成してしまうという問題に対し、具体性の高い応答を生成可能な技術の実現を目的として研究を行う。人間は常に具体性の高い応答を行うのではなく必要に応じて応答の具体性を使い分けていることに着目し、具体性を制御可能な対話応答生成手法の構築に取り組む。発話に対する応答の具体性を測る簡潔な尺度 MaxPMI を提案し、そのスコアを制御変数として訓練データに付与して訓練することで、具体性の制御を可能にする。評価実験では、提案手法が先行研究に比べてより応答の具体性をより広い範囲で制御できることを、自動評価・人手評価によって検証・確認した。

(3) では、対話システムが生成した破綻応答を検知する対話破綻検出タスクに取り組む。対話破綻を感じる基準は人によって異なることに着目し、破綻ラベルのアノテーション傾向に基づいて訓練データのアノデータをクラスタリングし、様々なアノテーション傾向を再現する複数の検出モデルを構築してアンサンブルする手法を提案する。実験においては、他のアンサンブル手法や従来の対話破綻検出手法に比べて提案手法が高い検出性能を達成できることを確認した。また、(2) において提案した対話応答生成モデルの生成結果に対して対話破綻検出モデルを適用した場合の性能比較も行い、提案手法の優位性を示した。

以上のように本論文は対話応答生成モデルや対話破綻検出器の性能向上を達成し、人間との円滑なコミュニケーションが可能な対話システムの構築に貢献したもので、情報科学技術の発展に寄与するところが大い。よって本論文は博士（情報科学）の学位論文として価値のあるものと認める。