

NAIST-IS-MT1351070

修士論文

対話システムに基づく嘘検出に関する研究

角森 唯子

2015 年 3 月 10 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.

角森 唯子

審査委員：

中村 哲 教授 (主指導教員)

小笠原 司 教授 (副指導教員)

戸田 智基 准教授 (副指導教員)

Graham Neubig 助教 (副指導教員)

Sakriani Sakti 助教 (副指導教員)

対話システムに基づく嘘検出に関する研究*

角森 唯子

内容梗概

人間が嘘を見破るための技術として、嘘の特徴を見分けることと、嘘の特徴を引き出す質問を行うことが挙げられる。従来研究として、嘘の特徴を見分けるために、機械学習を用いた嘘の自動検出が行われている。本論文では、もう1つの技術に焦点を当て、嘘の特徴を効率的に引き出せる質問について分析を行い、嘘を検出する対話システムを構築する。まず、日本語の嘘を含むコーパスの収集を行い、同様の収録条件で収集された英語のコーパスを用いて、嘘の検出実験の比較分析を行った。その結果、チャンスレートに対して有意な差が確認でき、嘘検出には言語特徴量よりも音響特徴量が日英ともに有効であることが分かった。さらに、嘘検出率の高かった質問についての分析を行い、確認タイプの質問が最も有効であることを確認した。これらの分析をもとに、HMMで質問者の発話の流れを学習し、対話制御のためのルールを作成し、対話システムを構築した。構築したシステムに対して評価実験を行い、有効性を確認した。

キーワード

嘘の自動検出, 対話システム, 質問法, 言語特徴量, 音響特徴量

*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, NAIST-IS-MT1351070, 2015年3月10日.

A study of dialogue system based deception detection *

Yuiko Tsunomori

Abstract

When humans attempt to detect deception, they perform two actions: looking for telltale signs of deception, and asking questions to attempt to unveil a deceptive conversational partner. There has been significant prior work on automatic deception detection that attempts to learn signs of deception. On the other hand, we focus on the second action, envisioning a dialogue system that asks questions to attempt to catch a potential liar. In this thesis, we describe the results of an initial analysis towards this goal, attempting to make clear which questions make the features of deception more salient, and based on this analysis create a deception detection dialogue system. We first collect a deceptive corpus in Japanese, our target language, perform an analysis of this corpus comparing with a similar English corpus, and perform an analysis of what kinds of questions result in a higher deception detection accuracy. Based on these analyses, we make rules using a model constructed by learning an HMM based on the flow of the interviewer's questions. We performed experiments using this system and confirmed its effectiveness.

Keywords:

Detection deception, Dialogue system, Question techniques, Lexical features, Acoustic/prosodic features

*Master's Thesis, Department of Information Science, Graduate School of Information Science, Nara Institute of Science and Technology, NAIST-IS-MT1351070, March 10, 2015.

目次

1. 緒言	1
1.1 背景	1
1.2 目的	2
1.3 本論文の構成	3
2. 関連研究	4
2.1 嘘について	4
2.2 ポリグラフを用いた嘘の検出法	6
2.3 嘘の自動検出法に関する研究	7
2.4 嘘検出に有効な質問法	9
2.5 嘘を含むコーパス	10
2.6 本章のまとめ	10
3. 日本語の嘘を含むコーパスの収録とラベル付与	12
3.1 収録内容	12
3.2 収録環境	14
3.3 真実・嘘ラベルの付与	14
3.4 本章のまとめ	14
4. 嘘の検出法の実験的評価	16
4.1 言語，音響，個人特徴量を使用した特徴量	16
4.2 分類手法	17
4.3 実験条件	18
4.4 結果と考察	19
4.5 学習していない対象者における嘘検出の評価	21
4.6 特徴量選択	21
4.7 本章のまとめ	22
5. 嘘検出に有効な質問の分析	24

5.1	対話行為ごとの分析	24
5.1.1	対話行為タグの付与	25
5.1.2	結果と考察	25
5.2	質問長ごとの分析	26
5.3	本章のまとめ	27
6.	嘘を検出する対話システム	29
6.1	音声対話システム	29
6.2	JDC に基づくルールを用いた対話制御	32
6.2.1	JDC に特化した対話行為タグ	32
6.2.2	HMM を用いた質問者の質問の流れの分析	32
6.2.3	ルール作成	36
6.3	Wizard of Oz 法	36
6.4	評価実験	36
6.4.1	実験条件	37
6.4.2	評価用システム	38
6.5	結果と考察	39
6.6	本章のまとめ	40
7.	結言	41
7.1	本論文のまとめ	41
7.2	今後の課題	42
	謝辞	44
	参考文献	45

目 次

1	本研究の流れ	3
2	ポリグラフを使用した取り調べの様子（出典：沖縄県警察 HP）	6
3	収録環境	13
4	Bagging の構成	18
5	提案システムの構成	30
6	音声対話システムの典型的な構成	31
7	FSA による対話の記述例	31
8	JDC の質問者の発話から生成した HMM	34
9	FSA による本システムのルール記述	35
10	実験の構成	37

表 目 次

1	文献 [7] で使用している特徴量	8
2	文献 [12] で使用している特徴量	8
3	対話例 (I : 質問者, S : 対象者)	13
4	音響および言語特徴量と個人特徴量の情報	17
5	嘘の検出率 : 音響特徴量 (AP), 言語特徴量 (L), 個人特徴量 (S) .	19
6	被験者ごとの嘘の検出率	19
7	対話例 (F : 高検出率, A : 低検出率)	19
8	1 対話抜き交差検定を行った際の結果	20
9	特に有効な特徴量	21
10	クラスごとの嘘検出の詳細	25
11	クラスごとの対象者 (返答) の平均 SU 長	25
12	質問の SU 長ごとの嘘検出の詳細	26
13	質問の平均 SU 長	27
14	JDC に特化した発話の内容カテゴリ	33
15	システムの嘘検出の F 値	39

1. 緒言

1.1 背景

嘘をついたりつかれたりする行為は決して非日常的なものではなく、人間は1日に平均2回の嘘をついているという報告もある [1]。日常生活の中で使用される嘘の多くは心理的な不利益の回避を目的としており、コミュニケーションの潤滑油としての機能も果たしている [2]。このような嘘は「嘘も方便」で示される嘘にあたり、話し手が嘘をついているかどうかは重要な問題ではない。しかしながら、話し手が嘘をついているかどうかを見極めなければならない場面も存在する。例えば、取り調べなどでは容疑者が嘘をついているのか疑わしい場面は多く、正確にそれを判別することは重要な課題である。そのような場面において嘘を的確に検出することができれば、取り調べは容易になり、さらに冤罪事件の減少も期待できる。

Hopper は、嘘を“事実とは異なる言語的陳述”と定義している [3]。このような対話中の嘘を見分けることは決して容易ではなく、刑事などは嘘を見抜く高度なテクニックを身につけている [4]。具体的には、嘘の特徴を見分けることと、嘘の特徴が露呈しやすいように質問を行うことが挙げられる [5]。

嘘を検出するために、これまで様々な方法について研究が行われてきた。呪術やトリックを用いた嘘検出法といった非科学的な方法までも存在するが、一般的なものとしては、ポリグラフなどの生理計測に基づいた検出法が挙げられる。ポリグラフを用いた嘘検出にはいくつかの質問法があるが、日本で主に用いられているのは裁決質問法 [6] と呼ばれるものである。裁決質問法の概要について以下に示す。

1. 事件に関する項目について尋ねる裁決質問と、それと類似しているが事件と関係しない項目について尋ねる非裁決質問を被検査者に提示する。
2. それらの質問に対して生じる心拍数、呼吸運動、皮膚電気活動、末梢血管抵抗などの自律系活動の変化を測定する。
3. 裁決・非裁決質問法で、自律系反応に違いがあれば事件に関する項目を認

識していると判定する．自律系反応に違いがなければ，事件に関する項目を認識していないと判定する．

しかしながら，このような生理計測に基づく嘘発見器は，上記のような決められた質問に対する返答の真偽を判別するという限定された状況下での嘘検出を目的としており，自由なコミュニケーション中に用いることは困難である．近年では，機械学習を用いた嘘の自動検出についての研究が行われており，自由なコミュニケーション中の嘘に対して一定の成果が得られている [7]．文献 [7] では，英語話者によって収録された嘘を含むコーパス (以下，CSC コーパス) に対して，音響特徴量と言語特徴量を用いて嘘の自動検出を行った結果，チャンスレート (60.2%) を有意に上回る分類精度 (66.4%) が得られることが確認された．一方で，言語間，文化間によって嘘の特徴に差異があることも指摘されている [8]．

これまでの研究は，既に終了した対話を対象にして嘘の検出を行っており，コミュニケーション中にリアルタイムで嘘の検出を行うことは困難である．さらに，人間が行う“嘘の特徴を見分けるテクニック”のみに焦点を当てている．これは嘘を検出するテクニックの1つにすぎず，嘘の特徴が露呈しやすい質問を行うことも同様に必要不可欠である．しかしながら，嘘の自動検出に関する先行研究では，このような質問に関するテクニックについては扱われていない．嘘の特徴が露呈しやすい質問を行うことで，より効率的に嘘の自動検出ができると期待される．

1.2 目的

本研究では，嘘の特徴を見分けるだけでなく，嘘の特徴が露呈しやすい質問を行い嘘を検出する対話システムの構築を目標とする．本研究の流れを，図 1 に示す．日本語を対象とした対話システムを構築するため，日本語の嘘を含むコーパスを収集し，日本語における嘘検出に有効な特徴量の種類を確認する．次に，嘘の特徴を露呈しやすくさせる質問の種類を確認する．これらを明白にすることによって，嘘の特徴を効果的に誘発する質問に焦点を当てた対話システムの構築が可能となる．最後に，以上の分析に基づいて嘘を検出する対話システムを構築して実験的評価を行い，本システムの有効性を検証する．

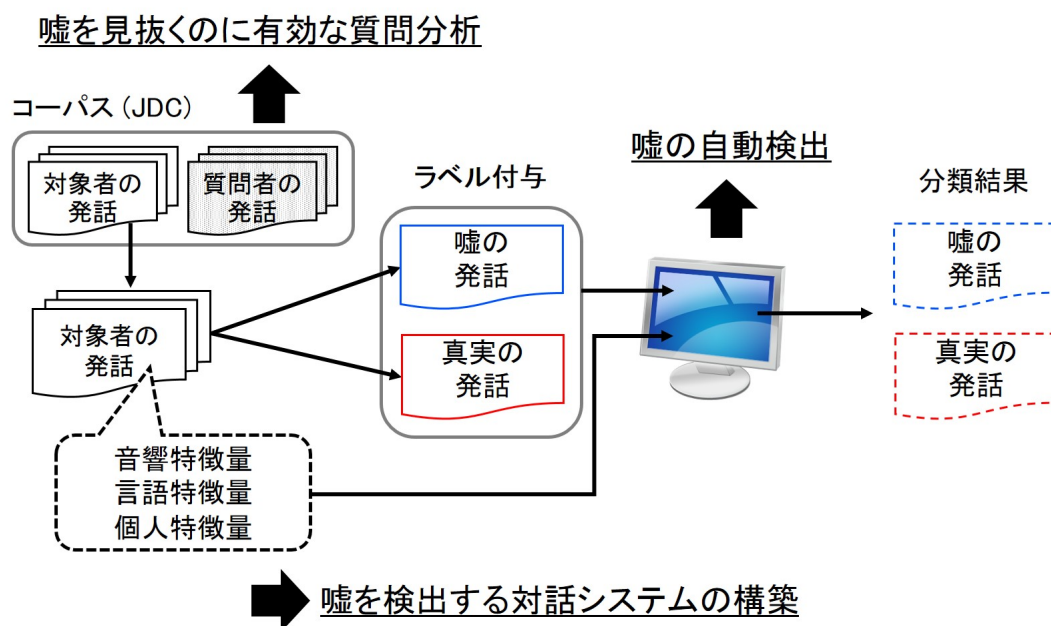


図 1 本研究の流れ

1.3 本論文の構成

本論文の構成について述べる．第2章で本研究に関連する研究として，嘘について述べたのち，嘘の検出法についての研究を説明する．第3章では，本研究で使用する日本語偽言コーパスの収録と真実・嘘ラベルの付与について述べる．第4章では，まず嘘検出に使用する素性である言語的特徴量，音響的特徴量，個人特徴量について説明する．次に，それぞれの素性を組み合わせて嘘の検出実験を行った際の検出率を評価する．さらに，特徴量選択を行い，日本語において有効な特徴量を選別する．第5章では，質問の種類ごとと質問長ごとの嘘の再現率を評価し，効率良く嘘を引き出すことができる質問について分析する．第6章では，第4章や第5章の分析結果に基づき，嘘を検出する対話システムを構築し評価を行う．第7章で本論文の成果をまとめ，今後の課題を示す．

2. 関連研究

本章では、まず 2.1 節で本研究で扱う嘘について説明する。嘘の検出の代表的な手法であるポリグラフを用いた手法について 2.2 節で説明し、2.3 節で本研究で扱う嘘の自動検出に関する研究について述べる。2.4 節で、嘘を効率的に検出するための質問法に関する研究について説明する。さらに、2.5 節で、嘘の検出実験などで用いられる嘘を含むコーパスについて紹介する。

2.1 嘘について

これまで、嘘については様々な観点から議論がなされてきた。しかしながら、嘘を示す決定的な手がかりは未だに確立されていない。Ekman [4] は、漏洩（話の途中で部分的に真実が露呈）と矛盾に起因した嘘の手がかりに基づく嘘の見破り方について報告している。感情を隠す目的ではない嘘も、往々にして感情が関与してくる場合が多く、いったん関与してしまうと嘘が露呈しないように隠し通す必要がある。いかなる感情も影響を及ぼすが、特に、見破られるのではないかという不安、嘘をついたことによる罪悪感、そして、うまく人を騙せたという喜びの 3 つの感情が欺瞞行為と密接に関係している。それらの感情が大きくなる場合について、Ekman [4] は以下のように述べている。

- 見破られるのではないかという不安
 - － 嘘をつく相手が、騙すには手強い相手であるとの評判がある場合。
 - － 嘘をつく相手が、疑いをもち始めた場合。
 - － 虚言者の側に実戦経験がほとんどなく、これまで嘘に成功したためではない場合。
 - － 虚言者が、見破られるのではないかという不安に特に敏感な場合。
 - － 見つかる危険性が大きい場合。
 - － 報酬と処罰の両方が問題になる場合。あるいは、どちらか一方だけが問題な場合には、処罰が問題になる場合。

- － 嘘を見破られた場合に、加えられる処罰が大きい場合。あるいは、嘘の内容に対する処罰が非常に重いために、白状する気になれない場合。
- － 嘘をつく相手が、その嘘から少しも恩恵を受けない場合。

- 欺瞞による罪悪感

- － 嘘をつく相手が不本意に騙される場合。
- － 欺瞞が全く利己的で、嘘をつく相手には何の利益もなく、おまけに虚言者の得る利益相当分の、あるいはそれ以上の損失を嘘をつく相手が被る場合。
- － 欺瞞が正当と認められず公認されない場合、また、その状況では、正直に振る舞うことが正しいとみなされている場合。
- － 虚言者が長い間欺瞞行為を実行していなかった場合。
- － 虚言者と嘘をつく相手の間で、社会的価値が共有されている場合。
- － 虚言者がそのターゲットと個人的な知己関係にある場合。
- － 嘘をつく相手のあら探しをしても、容易に卑劣な人間であるとか、騙されやすい人間とはみなせない場合。
- － 嘘をつく相手には、騙されるのではないかと考えるに足る理由がある場合、またそれとは正反対に、虚言者が自分を信用できる人間にみせかけて、嘘をつく相手の信用を得ようとした場合。

- 騙す喜び

- － 嘘をつく相手が騙しにくいことで有名であり、手強い相手の場合。
- － 隠蔽しなければならないそれなりの理由があって、あるいは、でっち上げるべき性質のものであるがゆえに、嘘をつくことがやりがいのある難題の場合。
- － 他人が、その嘘を見ていたり、それを知っていて、虚言者の見事な手際を高く評価する場合。



図 2 ポリグラフを使用した取り調べの様子（出典：沖縄県警察 HP）

このような感情は話し方，声の調子，体の動き，あるいは表情に現れるため，嘘を検出するためにはそれらについて知っておく必要がある。

Frank[9] は，筋書き（嘘の項目，賞金），嘘の種類（隠蔽や偽造），動機（自衛本能，自己プレゼンテーション）などを含む典型的な嘘の全ての側面を考慮し，嘘の構造について報告している．Enos [7] はこの報告をもとに，嘘の自動検出を行っており，それについては 2.3 節で説明する。

2.2 ポリグラフを用いた嘘の検出法

1.1 節でも述べた通り，嘘を検出する技術で最も一般的なものはポリグラフを使用する方法である．図 2 に，ポリグラフを使用した取り調べの様子を示す．その一般的な方法としては，まずはじめに，1 組のトランプから容疑者が引いたカードをポリグラフの機械が当てられることを示す．これは，ポリグラフの機械は決してミスを犯さないと容疑者に信じさせ，見破られるのではないかという不安を抱かせるためである．次に，身体のような部分センサを取り付け，1.1 節で述べたような質問を行った際の心拍数，呼吸の速度，血圧，発汗などの生体信号の変化を記録する．これらの生体信号の変化は，感情の変化の現れでもある．このように，ポリグラフによる検出は嘘そのものを検出するのではなく，嘘によって生

じる感情を検出している。しかしながら、嘘を検出するためにポリグラフは幅広く浸透しているものの、その正確さについてはほとんど科学的な証明はされていない [4]。また、ポリグラフの使用の賛否を煽っている原因の 1 つに、以下のような事例がある [10]。

1. 被害者の死ぬ前の供述により、A は強盗殺人の罪で逮捕された。
2. 警察は A にポリグラフテストを受けるように勧め、嘘を示す反応が出た場合は、その結果を証拠として採用することを認めるよう要求した。逆の場合は、起訴を取り下げるとした。
3. A は承諾してテストを受けた結果、嘘を示す反応が出た。2 回目のテストでも同様の結果が得られ、終身刑が言い渡された。
4. その 2 年後、真犯人が逮捕された。

この事例では、ポリグラフテストでは嘘を示す反応が出ているが、A は嘘をついていない。誤審の理由として、A はポリグラフに対する恐怖や罪悪感などで 2.1 節で述べた感情に変化が表れたと考えられる。このように、ポリグラフテストを受ける人が感じる感情は、その人の性格などに依存する。それゆえ、テスト結果も性格に左右されると言える。

2.3 嘘の自動検出法に関する研究

近年、記録した対話に対して、機械学習を用いた嘘の自動検出について研究が行われている。本節では、機械学習を用いた嘘の検出法の関連研究を 2 つ紹介する。それぞれの研究で使用しているコーパスについては、2.5 節で詳しく説明する。まず、英語に対して嘘の自動検出を行った研究について説明する。Enos [7] は音声特徴量と言語特徴量を用いて、CSC コーパス中の SU¹ に対して嘘の自動検出を行っている。使用した特徴量を表 1 に示す。分類器は Ripper [11] を用いて、嘘と真実の 2 値分類を行った。その結果、チャンスレート（60.2%）を有意に上回る分類精度（66.4%）が得られることが確認されている。

¹ 発話を句読点や休止で区切ったスラッシュユニットの略称

表 1 文献 [7] で使用している特徴量

カテゴリ	説明
一般記述	テスト項目
文構造	代名詞, 否定, YesNo 動詞の原形, 遊び言葉, 合図句, 質問, ポジティブ・ネガティブ単語
パラ言語情報	言いよどみ, 笑い, 雑音, 言い間違い
F_0	中央値, SU 長に対する中央値の割合
音素継続長	母音, 平均, 最大
パワー	平均, SU の第 1・最終フレーム
個人性	性別, 言いよどみ・合図句の頻度

表 2 文献 [12] で使用している特徴量

カテゴリ	説明
視線	話しかけている人の顔を見ている割合 話しかけていない人の顔を見ている割合 話しかけている人のカードを見ている割合 話しかけていない人のカードを見ている割合 顔でもカードでもない場所を見ている割合 注視した対象が推移した回数 注視した対象が推移した割合
韻律	声の高さ (前半, 後半, 変化) 声の大きさ (前半, 後半, 変化)
表情	目元よりも口元が早く変化した 目元が変化していた 口元が変化していた

次に, 日本語に対して嘘の自動検出を行った研究について説明する. 大本ら [12] は韻律と視線情報を用いて, インディアンポーカーコーパスに対して嘘の自動検

出を行っている．使用した特徴量を表 2 に示す．表 2 の変数 16 個を独立変数とし，発話が嘘であるかどうかを従属変数として，判別分析を行った．その結果，約 65% の割合で，嘘を正しく嘘と判別することに成功している．

ポリグラフによる嘘検出法は，決められた質問に対する返答の真偽を判別する限定された状況下での嘘検出を目的としているため，被験者への心理的抑圧がある．一方，Enos や大本らが提案する嘘検出法は自由なコミュニケーションで使うことができ，被験者への心理的抑圧も少なく，より正確な嘘の検出ができると考えられる．しかしながら，既存のデータに対しての使用に限定されており，さらに嘘を引き出すための質問については考慮していないため，こちらが知りたいと思う情報に対して嘘の検出を行うことはできない．

2.4 嘘検出に有効な質問法

嘘を検出するもう 1 つの技術である，嘘の特徴を露呈させる質問のテクニックについて紹介する．この方法は嘘を見抜く側が嘘のつき手に質問を行い，手がかかりが有効に機能するように状況を整えることが目的である．取り調べで用いられる方法として，ポリグラフ使用時に用いられた裁決質問法の他に，行動分析質問 (Behaviour Analysis Interview: BAI) がある [13]．これは，容疑者に対して 16 項目の行動誘発質問を行う．これらの質問に対して，嘘を付いている場合は行動誘発質問に対して，足を組む，身体を触るなどの非言語的特徴が現れる [13]．言語的特徴としては，捜査に非協力的で，事件の発生自体を否定する傾向が現れると言われている [14]．しかしながら，BAI は科学的な根拠については証明されておらず，さらに，Vrij ら [15] の報告によると，嘘をついていない人の方が，上記のような特徴が現れることが示されている．

裁決質問法や BAI は，「はい」か「いいえ」で答えられる質問を多用する傾向にある．これに対して，オープンな質問によってさらに詳細な情報を引き出す質問を“情報収集型”と言う．取り調べなどでは，容疑者を責めるような“問責型”の方法よりも“情報収集型”の方が有効であることがわかっている [16]．この質問では質問者から情報は与えず，容疑者の話を促し，より多くの話をさせることが目的である．“情報収集型”の質問を行うと，容疑者の非言語的・言語的な嘘

の特徴が露呈しやすい。一方，“問責型”の質問を行うと，容疑者は攻防的になり，嘘の特徴が露呈しにくいと言われている [17]。

次に，真実と嘘で異なる特徴を引き出すために，質問を戦略的に行う方法もある。嘘のつき手が全く予想していない質問を突如として行う方法 [18] や，証拠を掴んでいることが前提で，証拠と異なる情報を引き出すための SUE (Strategic Use of Evidence) テクニック [19] などがある。

しかしながら，どのような聞き方でどのような長さで“情報収集型”の質問を行うのかといった，詳細な質問の種類までは明らかにされていない。さらに，これらの質問を行う際には質問者が何らかの情報を掴んでいるという前提条件がある場合が多く，証拠が全く得られていない場合は使用が困難である。

2.5 嘘を含むコーパス

嘘の検出を行うために嘘を含むコーパスが必要であり，いくつかのコーパスが先行研究で構築されている。CSC コーパス [7] は，対象者が質問者に嘘をつくように促されたインタビューを収録し，騙しに成功した場合の報酬によって動機づけされている。インタビューは英語で行われ，25～50 分の合計 22 のインタビューが収録されている。その他にも，Idiap Wolf Corpus [20] がある。これは，ロールプレイングゲーム（人狼）に参加する被験者の会話を収録した視覚情報を含む多人数会話コーパスである。どちらのコーパスも，使用言語は英語である。

英語でコーパスが豊富な一方で，他言語のコーパスは数少ない。日本語のものとしては，インディアンポーカーコーパス [12] がある。これは，インディアンポーカーゲームの会話を収録した視覚情報を含む多人数会話コーパスである。

2.6 本章のまとめ

本章では，2.1 節では，嘘の特徴と嘘をついているときに生じる感情について説明した。見破られるのではないかという不安，嘘をついたことによる罪悪感，そして，うまく人を騙せたという喜びの 3 つの感情が嘘には大きく関係している。2.2 節で，嘘の検出の代表的な手法であるポリグラフを用いた手法について説明

した．ポリグラフは2.1節で述べた感情を検出しており，テスト結果はこれらの感情への耐性を含む被験者の性格に依存するという問題点が挙げられる．2.3節で，本研究で扱う嘘の自動検出についてについての研究について述べた．2.4節で，嘘を効率的に検出するための質問法に関する研究について説明した．これらの質問法は証拠を掴んでいるという前提条件がある場合が多く，さらに詳細な質問の種類などは明らかにされていない．2.5節で，嘘の検出実験などで用いられる嘘を含むコーパスについて紹介した．

3. 日本語の嘘を含むコーパスの収録とラベル付与

2.5 節で、嘘の検出実験などで用いられる嘘を含むコーパスについて紹介した。本研究は一对一の対話システムを構築することを目標にしているため、インディアンポーカーコーパスのような多人数話者によるコーパスを用いるのは適切ではない。本章では、言語と文化間で嘘検出の研究を比較し、日本語の嘘検出対話システムを構築するために、日本語の嘘を含むコーパスを収集する。3.1 節で収録内容について紹介し、実際に収録した対話の内容と結果をまとめる。3.2 節で、対話コーパスの収録環境について述べる。3.3 節では、収録した対話コーパスの対象者の発話に真実・嘘ラベルを付与する手法について述べる。

3.1 収録内容

CSC コーパス [7] と同様の収録設定の下、日本語で対話収録を行う。嘘が誘発されやすい場面の一例として、面接 [21] における質問者と対象者の 2 名の対話を想定する。収録の手順を以下に示す。

1. 対象者が 6 項目（政治、音楽、地理、食べ物、インタラクティブ、サバイバル）に対する、ある“目標プロフィール”に適合するかどうかを調べるための実験であると対象者に伝える。
2. 対象者に 6 項目のテストを収録前に受けてもらう。
3. 実験者は、2 項目が適合、4 項目が不適合となるように結果を操作し、その結果を対象者に伝える。
4. この実験の本当の目的は“目標プロフィールに適合している”と主張し、質問者を納得させることができる人物を探すことであり、上手く納得させれば賞金があると対象者に伝える。
5. 対象者は、全項目のテスト結果において“目標プロフィールに適合している”と、面接で質問者に主張する。質問者は、目標プロフィールやテスト内容についての知識はなく、いかなる質問をしても構わないとする。

表 3 対話例 (I : 質問者, S : 対象者)

話者	書き起こし文	ラベル
I	音楽に関して、あなたはマッチしていましたか？	
S	はい、マッチしていました。	嘘
I	それはなぜだと思いますか？	
S	えーと、そこそこ答えれたからです。	真実
S	小さい頃からずっとピアノをやっていたので。	嘘

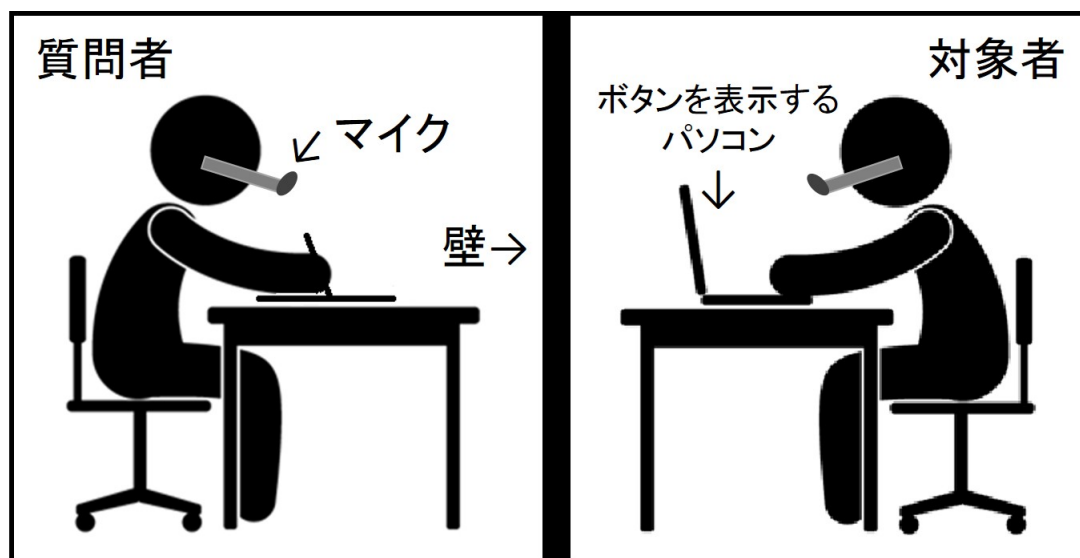


図 3 収録環境

20代の男女計2名を質問者とし、20代の男女計9名を対象者として対話収録を行った。収録した総対話数は9対話であり、総時間は約130分である。対象者のSUの総数は1545である。収録した日本語偽言コーパス(JDC)²の一例を表3に示す。

²<http://ahclab.naist.jp/resource/ja-deception/> : 研究目的に限り使用することができる。

3.2 収録環境

JDC を収録するために、質問者と対象者は別々の部屋に入って対話を行う。それぞれの部屋は防音室になっており、自身以外が出す雑音や反響音は極めて少ない状況で収録する。互いの顔が見えない状態かつ、互いの声は直接聞こえない状態とする。本コーパスは、音響・言語特徴量を用いての嘘の検出実験や分析に使用する。質問者の発話が視覚情報（対象者の表情や視線）に左右されることを防ぐため、互いの顔が見えない状態での収録を行った。質問者と対象者は、マイク (CROWN 社製 CONDENSER MICROPHONE) とイヤフォンを装着した状態で、対話収録を行う。対話収録時の簡略な様子を図 3 に示す。対話音声は、同時に 2 チャンネルで独立に収録し、サンプリングレートを 48 kHz、量子化ビット数を 16 bit とする。

3.3 真実・嘘ラベルの付与

対象者の SU に対してラベル付与を行う方法として“真実”と“嘘”のボタンを用意し、対象者には対話収録中の SU ごとにどちらかのボタンを押してもらうように指示する。なお、一部分でも嘘が含まれる SU は、嘘として定義する。それぞれのボタンが押された時間と書き起こしから、SU 毎にラベル付けを行う。1545 SU のうち真実ラベルが付与された SU は 1296 で、嘘ラベルが付与された SU は 249 である。

3.4 本章のまとめ

本章では、日本語における嘘を含むコーパスの収集を行った。3.1 節で CSC コーパス [7] と同様の収録条件を説明し、実際に収録した対話の内容と結果をまとめた。質問者 2 名、対象者 9 名を被験者とし、合計 9 対話の収録を行った。3.2 節で、対話コーパスの収録環境について説明した。質問者と対象者の 2 名の対話を想定し、互いに顔が見えない環境下で収録を行った。3.3 節では、収録した対話コーパスの対象者の発話に真実・嘘ラベルを付与する手法について述べた。コーパス

中のラベル比は，真実ラベルが約 8 割を占め，嘘ラベルが約 2 割であった．これらの条件のもと，日本語偽言コーパス（JDC）を構築した．

4. 嘘の検出法の実験的評価

本章では，日本語における嘘検出に有効な特徴量を調べるために，Enos [7] が英語のコーパスに対して行った実験に則り，いくつかの検証を行う．4.1 節では，まず日本語コーパスに対して行う実験で用いる特徴量について述べる．4.2 節では，“真実”と“嘘”に2値分類するための機械学習手法である Bagging について説明する．4.3 節で実験条件について説明し，4.4 節で実験の結果と考察を述べる．4.5 節では，学習を行っていない話者に対して嘘の検出実験を行った際の結果を示す．4.6 節では特徴量選択を行い，日本語において嘘の検出に有効な特徴量を選別し，日英間における有効な特徴量の差異を調べる．

4.1 言語，音響，個人特徴量を使用した特徴量

嘘の検出実験を行うために，嘘の手がかりを示す特徴量を抽出する．本研究は Enos ら [7] の報告に基づき，SU 毎に言語特徴量と音響特徴量に着目する．抽出した特徴量を表 4 に示す．

- 言語特徴量

人手による書き起こし文に対して MeCab [22] で形態素解析を行い，得られた形態素から言語特徴量を抽出する．ポジティブ・ネガティブ単語の頻度の抽出には，単語感情極性対応表 [23] を用いる．英語の先行研究 [24] で用いられた特徴量に加えて，“終助詞の数”を新たに加える．

- 音響特徴量

音響特徴量として，基本周波数 (F_0) とパワー，音素継続長を用いる．文献 [7] とは異なるが，パラ言語情報についてもこちらに含むこととする．雑音は，咳やマイクとの接触音など対象者によって生成されたものに限定する． F_0 の抽出には Snack Sound Toolkit [25] を用いる．Kaldi [26] を用いて，音素列を与えてアライメントを行い，音素継続長を抽出する．

- 個人特徴量

対象者の性質に依存する特徴量を抽出する．対象者の性別と，1 対話中の全

表 4 音響および言語特徴量と個人特徴量の情報

カテゴリ	説明
一般記述	テスト項目
文構造	代名詞, 否定, YesNo, 終助詞, 動詞の原形, 遊び言葉, 合図句, 質問, ポジティブ・ネガティブ単語
パラ言語情報	言いよどみ, 笑い, 雑音, 言い間違い
F_0	中央値, SU 長に対する中央値の割合
音素継続長	母音, 平均, 最大
パワー	平均, SU の第 1・最終フレーム
個人特徴量	性別, 言いよどみ・合図句の頻度

SU 数に対する合図句・言いよどみの頻度を個人特徴量とする.

4.2 分類手法

本研究では, 2 値分類するための機械学習手法として Bagging [27] を用いる³. Bagging とは複数の弱分類器を組み合わせて集団学習を行い, 結果推定を行うコミッティ学習手法 [29] の 1 つである.

図 4 に, Bagging の構成を示す. 以下に Bagging のアルゴリズムを示す.

1. 1 つの学習データセット $\mathbf{L}$ に対して Bootstrap [30] と呼ばれるリサンプリング法を適応し, 大きさ M の複数の学習データセット $\mathbf{L}_k = \{\langle \mathbf{x}_1, y_1 \rangle, \dots, \langle \mathbf{x}_M, y_M \rangle\}, k = 1, \dots, N$ を生成する.
2. N 個の学習データセットをそれぞれ弱分類器に与えて関数 ϕ_k を学習する. これを用いて, テストデータ $\mathbf{x}$ を入力した際の推定結果 $\phi_k(\mathbf{x}), k = 1, \dots, N$ を得る.

³予備実験として, SVM [28] や Ripper [11] を使用して分類を行ったが, Bagging を使用した場合の検出率が最も高かったため, Bagging を使用する.

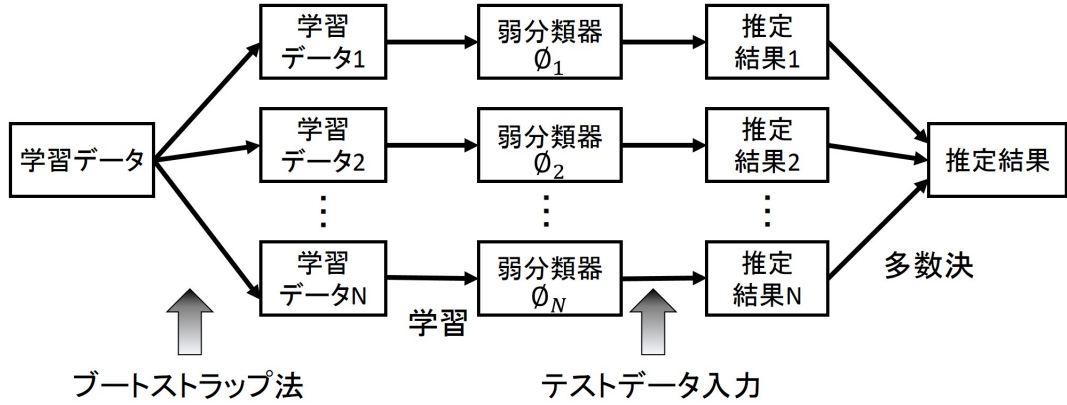


図 4 Bagging の構成

3. 各推定結果 $\phi_k(\mathbf{x})$ に基づき，最終推定結果 $\phi(\mathbf{x})$ の決定を行う．最終推定結果は， y が数値の場合は以下に従って平均を取る．

$$\phi(\mathbf{x}) = \frac{1}{N} \sum_{k=1}^N \phi_k(\mathbf{x}) \quad (1)$$

y がクラスの場合は，以下に従い多数決を取る．本研究では“真実”と“嘘”の2クラスに分類するため，こちらに従う． $I(\cdot)$ は指示関数を表す．

$$\phi(\mathbf{x}) = \arg \max_j \sum_{k=1}^N I(\phi_k(\mathbf{x}) = j) \quad (2)$$

4. 最終推定結果 $\phi(\mathbf{x})$ を返し，テストデータの分類クラスとする．

4.3 実験条件

JDC に対しては，4.2 節で述べた音響特徴量と言語特徴量，個人特徴量を合わせた特徴量（表 4）を用いて，SU に対して嘘の検出を行う．CSC に対しては，表 1 の特徴量を用いる．分類には 4.2 節で述べた Bagging を用い，弱分類器には REPTree [31] を使用する．嘘の検出率の評価には，1544 SU を学習用のデータにし，残り 1 SU をテスト用のデータにする一個抜き交差検定を用いる．

表 5 嘘の検出率：音響特徴量 (AP)，言語特徴量 (L)，個人特徴量 (S)

特徴量	日本語		英語	
	検出率 (%)	F 値 (%)	検出率 (%)	F 値 (%)
チャンスレート	83.8	0.0	71.4	0.0
AP	90.7	62.5	86.9	74.5
L	83.7	9.4	71.5	15.2
AP+S	91.8	67.4	88.1	77.7
L+S	85.2	34.4	76.5	51.9
AP+L	90.5	61.2	86.8	74.6
AP+L+S	91.1	63.7	87.8	77.2
人間	83.0	28.4		

表 6 被験者ごとの嘘の検出率

被験者	A	B	C	D	E	F	G	H	I
チャンスレート (%)	89.4	92.6	76.9	78.7	75.9	64.9	82.9	73.2	83.9
検出率 (%)	89.4	93.2	80.2	86.5	88.6	84.7	87.4	73.1	93.7

表 7 対話例 (F：高検出率，A：低検出率)

対象者	書き起こし文
F	(SN) は一，そうですね，まあ (MP) 七八割は答えれたかなっていう位ですね.
A	たぶん大丈夫だと思います.

4.4 結果と考察

表 5 に，各特徴量を用いた際の嘘の検出率と F 値を示す．“日本語”は本研究における日本語の嘘検出率，“英語”は Enos [7] の報告に基づき，CSC コーパスの一部を用いて英語の再現実験（日本語で加えた特徴量は加えていない）を行った際の嘘検出率である．“人間による検出”とは，被験者と別の人間に同じ対話データを聞いてもらい，嘘の検出を行ってもらった結果である．なお，検出の際に文

表 8 1 対話抜き交差検定を行った際の結果

被験者	A	B	C	D	E	F	G	H	I
チャンスレート (%)	89.4	92.6	76.9	78.7	75.9	64.9	82.9	73.2	83.9
検出率 (%)	89.4	86.9	64.8	79.2	78.5	67.6	82.9	75.6	74.9

脈は考慮しない。

日本語では音響特徴量+個人特徴量を用いた場合が最も精度が高く、言語特徴量のみと比べて嘘検出率が約8%高かった。全特徴量を用いた場合でも、音響特徴量+個人特徴量、音響特徴量のみを検出率より低い精度であった。英語でも音響特徴量+個人特徴量を用いたものが最も精度が高く、チャンスレートに比べて約17%高かった。人間による検出はチャンスレートとほぼ同等の精度であり、どの特徴量を用いた検出もこれを上回る精度であった。また、日英ともに言語特徴量を用いた場合の検出率はチャンスレートとほぼ同等であった。さらに、言いよどみと合図句の頻度、性別などを追加して検出を行った場合に精度の向上が見られたことから、日英ともに嘘検出には個人特徴量が影響を及ぼすことがわかった。なお、フィッシャーの正確確率検定の1%水準を用いて特徴量間の違いを比較したところ、日英ともにチャンスレートと音響特徴量、音響特徴量+個人特徴量、音響特徴量+言語特徴量、音響特徴量+言語特徴量+個人特徴量で p 値が 0.01 以下となり、有意差が見られた。

表 6 に、話者ごとの嘘の検出率を示す。検出率が低いほど、嘘が見破られにくい被験者ということを表している。表 7 に、嘘の高検出率と低検出率の対象者の対話例を示す。SN は雑音、MP は言い間違いのことを指す。質問者がテストの出来について質問した際の対話であり、チャンスレートに対して F は検出率が高かった話者、A は低かった話者である。F は雑音や言い間違いが多く、語尾が上ずっていた。その一方で、A は真実の発話に比べて際立った変化は見られなかった。

表 9 特に有効な特徴量

カテゴリ	英語	日本語
言語	雑音, 代名詞, YesNo	動詞の原形
個人特徴量	合図句の頻度	
F_0	中央値	中央値
音素継続長	平均, 母音	母音
パワー	第 1・最終フレーム, 平均	最終フレーム

4.5 学習していない対象者における嘘検出の評価

上記の実験では、1544 SU を学習用のデータにし、残り 1 SU をテスト用のデータにする一個抜き交差検定を用いている。本節では、学習データに用いていない対象者に対して嘘の検出実験を行う。つまり、8 対話を学習用のデータにし、残り 1 対話をテスト用のデータにする一個抜き交差検定を用いて評価を行う。表 8 に、本人のデータで学習に用いていない対象者をテストデータに用いた際の嘘の検出率を示す。ほとんどの対象者において、嘘の検出率がチャンスレートを上回らないことがわかる。このことから、システムに実装する際は嘘の検出を行う前に、その話者の“嘘”と“真実”の発話について予め学習しておく必要があることがわかる。

4.6 特徴量選択

嘘を検出する対話システムで使用するために、4.4 節で用いた特徴量のうち、日本語において特に有効な特徴量を選別する必要がある。言語間、文化間において嘘の特徴には差異があることが指摘されていることから [8]、4.4 節で用いた特徴量のうち、特に有効な特徴量は日英間で差異があると考えられる。本節では、嘘の検出に有効な特徴量を選別するために、日本語と英語それぞれについて特徴量選択を行う。最良優先探索を用いて、分類率が学習データに対して最大となるように特徴量選択を行った。

表9に、特徴量選択の結果、特に有効とされた特徴量を示す。日本語では、言語特徴量としては動詞の原形、音響特徴量としては F_0 の中央値、母音の平均継続長、パワーの最終フレームが有効であるとされた。次に、特徴量選択の結果から日英間の差異を考察する。言語特徴量は、日英で有効な特徴量が大きく異なっており、英語では、代名詞と YES・NO の有無が有効であるとされた。一方で、日本語では有効とされた言語特徴量は動詞の原形の有無のみであった。音響特徴量に関しては、 F_0 の中央値、母音の平均継続長、パワーの最終フレームが日英ともに有効であることが確認された。これらの特徴量が日英ともに有効であるとされた理由は、言語に関係なく以下の特徴を捉えているためと考えられる。

- パワーの最終フレーム

一般に、感情の変化（不安など）は発話の語尾に現れると言われている。

- F_0 の中央値

嘘をつく際は、声の高さに変化が出る [21]。

- 母音の平均継続長

嘘を話す場合と真実を話す場合を比較すると、話す速度が異なる [32]。

ほぼ同じ特徴量が日英ともに有効であることから、音響特徴量に関しては日英間で大きな差異がないと言える。一方で、有効な言語特徴量は日英間において大きく異なっており、英語で有効であるとされた言語特徴量について、日本語では有効性を確認できなかった。

4.7 本章のまとめ

本章では、日本語における嘘検出に有効な特徴量を調べるために、Enos が [7] 英語のコーパスに対して行った実験に則り、いくつかの検証を行った。4.1 節では、まず日本語コーパスに対して行う実験で用いる特徴量について述べた。4.2 節では、“真実”と“嘘”に2値分類するための機械学習手法である Bagging について説明し、4.3 節ではこれを用いた実験条件について述べ、嘘の検出実験を行った。4.4 節では、実験の結果、日本語においても Enos らの報告 [7] とほぼ同

等の結果が得られたと述べた。4.5 節では，学習されていない話者に対しての嘘検出実験を行った。その結果，学習されていない話者の嘘は検出できないことが確認された。4.6 節では，日本語において嘘の検出に有効な特徴量を選別し，日英間における有効な特徴量の差異を調べた。英語と日本語で嘘の検出に有効な特徴量の差異を調べた結果，英語で有効とされた言語特徴量は，日本語では有効性を確認できなかった。一方，音響特徴量は日英間でほぼ同じ特徴量が有効であることを確認した。

5. 嘘検出に有効な質問の分析

本章では、4章で着目した嘘を見分けるテクニックではなく、もう1つのテクニックである嘘の特徴を露呈させやすい質問について分析を行う。検出には、音響特徴量、言語特徴量、個人特徴量の全てを用いる。5.1節では、有効な質問の種類について分析を行う。5.2節では、有効な質問の長さについて分析を行う。

5.1 対話行為ごとの分析

まず、質問者の発話の種類（対話行為）が嘘検出率に影響を及ぼすと仮定する。この仮説を検証するために、返答を分類するクラスとして、質問の対話行為を使用する。ISO国際標準により定められた一般目的機能（general-purpose functions: GPF）を利用し、質問者の発話にGPFタグを付与する（ISO24617-2, 2010）。GPFタグの定義を以下に述べる。

- **CheckQ**

聞き手によって与えられた命題の真偽について、話し手が確信が持てていない場合に使用する。話し手が命題の真偽を確認するため、聞き手に対して情報提供を促す。

E.g. 「打ち合わせは10時から始まりますよね？」

- **ChoiceQ**

話し手が、聞き手が知っていると仮定する命題のリストの中における真の要素を知るために、聞き手に対して情報提供を促す。

E.g. 「打ち合わせは10時と11時、どちらの時間から始まりますか？」

- **ProQ**

話し手が、聞き手が知っていると仮定する命題の真偽を得るために、聞き手に対して情報提供を促す。

E.g. 「打ち合わせは10時から始まりますか？」

表 10 クラスごとの嘘検出の詳細

クラス	嘘の再現率 (%)	検出率 (%)	チャンスレート (%)
CheckQ	83.3	95.3	77.7
ChoiceQ	80.0	96.4	78.6
ProQ	69.1	91.5	66.7
SetQ	73.7	94.0	75.8

表 11 クラスごとの対象者（返答）の平均 SU 長

	クラス				
	CheckQ	ChoiceQ	ProQ	SetQ	平均
SU 長	6.4	12.2	11.8	18.1	12.3

- **SetQ**

話し手が、聞き手が知っていると仮定する集合のどの要素が固有の特性を持つかを知るために、聞き手に対して情報提供を促す。

E.g. 「打ち合わせは何時から始まりますか？」

5.1.1 対話行為タグの付与

本研究では、JDC のうち同一質問者で収録された 7 対話に対して、人手で GPF タグの付与を行う。JDC の 10% に対し 2 人のアノテータがタグの付与を行った結果、一致率は約 80% であった。これより、第 1 アノテータが付与した GPF タグに従う。

5.1.2 結果と考察

表 10 に、GPF タグが付与された質問のクラスごとの嘘検出率を示す。ユーザが嘘をついている際にシステムが正しく検出を行うために、嘘の再現率に着目する。最も正しく嘘を検出できているのは、CheckQ のクラスである。この対話行為は、JDC において、SetQ の質問から得た返答に対して、質問者が疑いを持っ

表 12 質問の SU 長ごとの嘘検出の詳細

SU 長	嘘の再現率 (%)	検出率 (%)	チャンスレート (%)
1 ~ 10	83.6	95.3	75.4
11 ~ 20	68.2	90.6	70.4
21 ~	69.7	93.3	78.0

た場合に使用されることが多かった。また、CheckQ への返答において、対象者は再度同じ話をする傾向にあった。対象者の自信を揺るがして嘘を露呈させるために、再度同じ話をさせることが有効であると [32] で述べられていることから、CheckQ が嘘の検出に有効であることがわかる。その一方で、最も検出率が低かったのは ProQ である。ProQ はイエス・ノーで返答することから、対象者への心理的抑圧が少なく、嘘の特徴が露呈しにくかったと考えられる。

さらに、表 11 に、各クラスの返答の SU あたりの単語数（平均 SU 長）を示す。これより、CheckQ のクラスの返答の平均 SU が最も短いことがわかる。極端に短い発話において嘘が露呈しやすいと [32] で述べられていることから、CheckQ は嘘検出に有効な質問であると言える。

5.2 質問長ごとの分析

嘘を見破るために、嘘の特徴を露呈しやすくさせる質問を行うことは重要であることは既に述べた。特に、話の細部について質問し、相手の自信を揺るがすことで、嘘の特徴を露呈させやすい [32] ことは前節で確認した。その際、CheckQ のような再度同じ話をさせる質問に対して、対象者は作り話の細かい点にまで答えざるを得ない場合は多い。このような場合において、質問者が質問を行っている時間に対象者が作り話を考えるという点から、質問長は重要な要素であると考えられる。ここで、質問長が短いと対象者が嘘を考える時間が短くなるために下手な嘘をつき、質問長が長いと巧みな嘘をつくことができると仮定する。これを調べるために、質問長ごとの嘘の検出率について検証する。

表 12 に、質問の SU 長ごとの嘘検出の詳細を示す。1 ~ 10 は、SU 長が 1 ~ 10

表 13 質問の平均 SU 長

	質問				
	CheckQ	ChoiceQ	ProQ	SetQ	平均
SU 長	13.3	25.0	16.3	13.6	15.9

単語の質問に対する対象者の返答に対して嘘検出を行った際の嘘の再現率を指す。SU 長が 1 ～ 10 単語が最も再現率が高いことから、SU 長が短い質問が嘘の検出には有効であるとわかる。短い質問時の嘘の再現率が高い理由として、質問長が短い場合は下手な嘘をつき、質問長が長い場合は巧妙な嘘をついていると考えることができる。さらに、SU 長が短い質問を行い対話の展開を早く行うことで、対象者の嘘が露呈しやすいと考えられる。また、フィッシャーの正確確率検定の 5% 水準を用いて SU 長間の違いを比較したところ、1 ～ 10 単語の質問と 11 ～ 20 単語の質問で p 値が 0.05 以下となり有意差が確認された。1 ～ 10 単語の質問と 21 ～ 単語の質問では、 p 値が 0.1 以下となり、有意傾向が見られた。

なお、本節で有効とされた 1 ～ 10 単語の質問に、5.1 節で有効とされた CheckQ の質問が集中している可能性がある。そこで、表 13 に、各質問ごとの平均 SU 長を示す。いずれの質問の平均 SU 長も 1 ～ 10 単語より長く、CheckQ が 1 ～ 10 単語の質問長に集中しているわけではない。以上のことから、質問のタイプに関わらず、短い質問が嘘の検出には有効であると言える。

5.3 本章のまとめ

本章では、もう 1 つのテクニックである嘘の特徴を露呈させやすい質問について分析を行った。5.1 節では、有効な質問の種類について分析を行った。その結果、CheckQ 型の対象者に再度同じ話をさせる CheckQ の質問が有効であることを確認した。CheckQ に対する返答の SU 長が極端に短かったことから、対象者が動揺していることが伺える。5.2 節では、質問長が短いと嘘を考える時間が短くなるために下手な嘘をつき、質問長が長いと巧妙な嘘をつくことができるという仮定に基づき、有効な質問の長さについての分析を行った。その結果、質問の

種類に関わらず，短い質問が有効であることが確認された．

6. 嘘を検出する対話システム

3章で収集したJDCの分析に基づき，対話制御部を実装する．5章で有効とされた質問を行い，4章で有効とされた音響特徴量を用いて嘘の検出を行う．図5に，提案システムの構成を示す．

本章では，最初に対話システムの概要について6.1節で説明する．本システムの対話制御部については，6.2節で説明する．本研究で用いるシミュレーション手法であるWoZ法について6.3節で説明し，6.4節の条件に従い実験を行い，6.5節で結果と考察についてまとめる．

6.1 音声対話システム

音声対話システムとは，一連の音声による入力発話を“理解”し，“適切に応答”する（と人間が感じる）システムである．一般的な音声対話システムは，音声認識部，言語理解部，対話制御部，応答文生成部，音声合成部の5つのモジュールから構成されている．音声対話システムの典型的な構成を図6に示す．

- 音声認識部

入力された音声を，文字列・単語列に変換する処理を行う．音声認識エンジンは，音響モデル，発音辞書，言語モデルを用いて，入力された特徴量から発話された言葉を推定する．音声認識の誤りは，その後の全ての処理に影響するため，ユーザ満足度の最大のボトルネックと言われている [33]．

- 言語理解部

認識した文字列から，意図を抽出する処理を行う．曖昧性のない論理式のように，対話制御部において処理が可能な意味表現に変換する．意味ネットワーク [34] やワードスポッティング [35] といった手法がよく用いられている．

- 対話制御部

入力された情報から意図を抽出し，システムが次に何をするか決定する処理を行う．システムの応答を決定する枠組みとして，様々な手法が提案さ

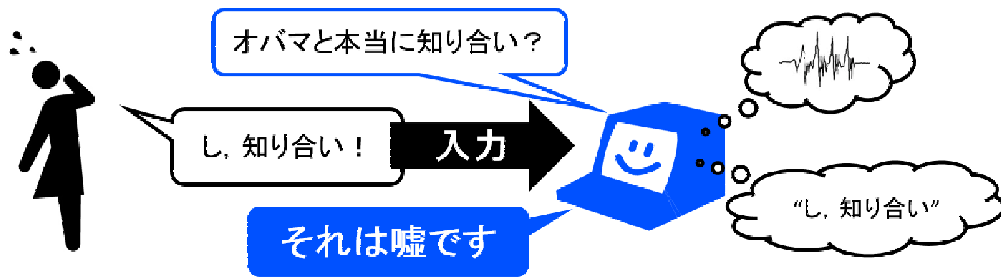


図 5 提案システムの構成

れている．有限状態オートマトン (Finite State Automaton: FSA) [36] に基づく対話制御法では，ユーザからの入力に対するシステムの応答規則を，FSA の形式で明示的に記述する．FSA による対話の記述例を図 7 に示す．全く文脈と合わないような制御を行う可能性が低く，対話破綻が起こる可能性は低い．また，対話の文脈の細かな部分に応じて制御を行うことができ，対話状態を考慮した対話制御が可能となる一方で，人手によりシステムの応答ルールを決定する必要がある．そのため，応答が単調になったり，複雑な文脈に対する細かいルールの記述が多くなりすぎるために，作成・保守が困難になるという問題がある．この問題に対して，マルコフ決定過程 (Markov Decision Process: MDP)[37] は，対話状態に対する最適なシステムの応答ルールを強化学習を用いて決定することが可能である．しかしながら，対話システムは音声認識部や言語理解部の誤りにより生じる曖昧性のために，ユーザの発話を一意に決定することは難しい．対話状態を決定的に求める FST や MDP に基づく対話制御法では，それらの問題に対処することができない．そこで，対話状態の不確定性を考慮するために，ベイジアンネットワーク (Bayesian Network: BN) [38] を用いて，対話状態を信念状態として確率的にモデル化し，信念状態に基づいた応答ルールを人手で記述する枠組みも提案されている．さらに，これらの手法を組み合わせたものとして，部分観測マルコフ決定過程 (Partially Observable MDP: POMDP)[39] に基づく対話制御法も提案されている．POMDP では，対話状態を信念状態として確率的にモデル化することで不確定性を考慮しつつ，強化学習による最適なシステム応答ルールを決定することが可能である．

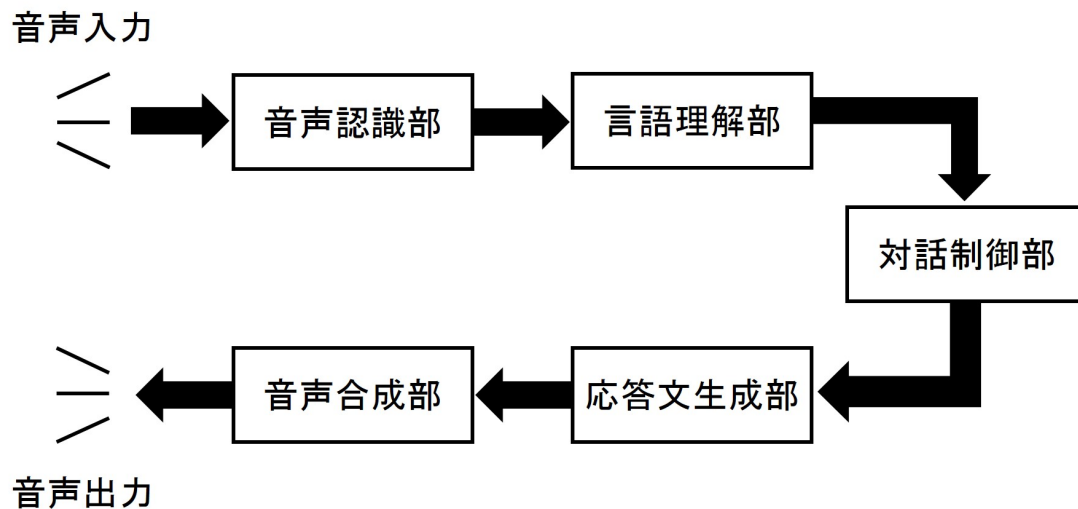


図 6 音声対話システムの典型的な構成

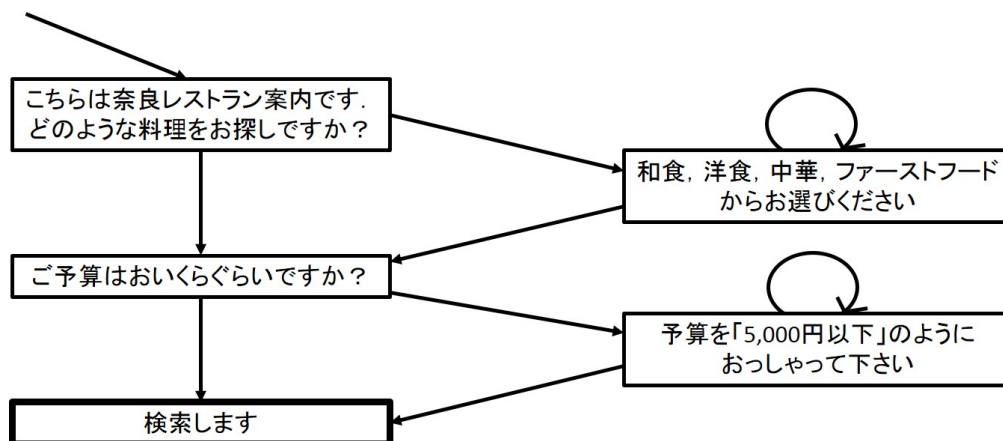


図 7 FSA による対話の記述例

- 応答文生成部

システム応答の文章を生成する処理を行う。テンプレート方式やスロット法などを用いて、文章の生成を行う。

- 音声合成部

生成した文章を音声に変換する処理を行う。テキスト音声合成 [40] を用いる方法が一般的である。

6.2 JDCに基づくルールを用いた対話制御

対話をモデル化する手法として、隠れマルコフモデル（HMM）がよく用いられている。目黒ら [41] は、人同士が傾聴を行っている対話を HMM を用いてモデル化を行い、その分析結果に基づいたルールを用いて対話システムを構築し、有効性を確認している。

本研究では HMM を用いた分析に基づき、ルールを用いて対話制御部を実装する。5 章で使用した 4 種類の GPF タグに JDC に特化した発話の内容を表すカテゴリを組み合わせ、細分化対話行為タグを提案する。細分化対話行為タグを付与した JDC から HMM を学習し、そのモデルに基づいて対話制御のルールを構築し、評価実験を行う。

6.2.1 JDC に特化した対話行為タグ

GPF タグに JDC に特化した発話内容のカテゴリを組み合わせ、細分化した対話行為タグを生成する。表 14 に示す 11 個の発話内容のカテゴリを、5.1 節で使用した質問の GPF タグ 4 種類と組み合わせる。“CheckQ_プロフィール”のように、“GPF タグ_内容カテゴリ”を 1 つの細分化対話行為タグとして使用する。JDC のうち同一質問者で収録された 7 対話に対して、人手で細分化対話行為タグを行う。

6.2.2 HMM を用いた質問者の質問の流れの分析

質問者の発話の流れを分析するために、対話行為の系列を HMM を用いてモデル化する。GPF タグと前節で提案したカテゴリを組み合わせた細分化対話行為タグを用い、これを観測値とする。ただし、5.1 節で有効であるとされた CheckQ の質問は JDC から除外したのち、HMM で学習を行う。JDC において CheckQ の質問は、対象者の発話に疑いを抱いた際に使用されている場合が多いため、対象者の発話の嘘確率に依存すると考えるためである。HMM の学習には EM アルゴリズムを用いる。状態数を 1 ～ 10 個と変化させ、それぞれの状態数で 100 個の HMM ずつ、計 1,000 個の HMM を学習する。これは、初期値に依存して HM

表 14 JDC に特化した発話の内容カテゴリ

クラス	定義
項目	テストの項目に適合していたかどうかについて E.g. 「あなたは音楽の項目は適合していましたか？」
経験	対象者の経験について E.g. 「東京に住んでいたんですか？」
テスト	テストの内容について E.g. 「どういう問題が出たんですか？」
プロフィール	対象者のプロフィール・嗜好について E.g. 「ワインとかよく飲むんですか？」
位置	どの程度で適合したかについて E.g. 「高い得点と低い得点どちらで適合したと思いますか？」
ターゲットの得点	目標プロフィールである起業家の得点について E.g. 「起業家の人たちはこのテスト何点ぐらい取りますか？」
得点	テストにおける得点について E.g. 「何点ぐらい取れたと思いますか？」
得点順位	6 項目中の得点の順位について E.g. 「一番高得点だったと思う項目はどれですか？」
必要性	テストの項目に関する知識の必要性について E.g. 「起業家にとって音楽の知識は必要だと思いますか？」
必要性順位	6 項目中で起業家にとっての必要順位について E.g. 「起業家にとってどれが一番必要だと思いますか？」
その他	その他 E.g. 「最後に質問はありますか？」

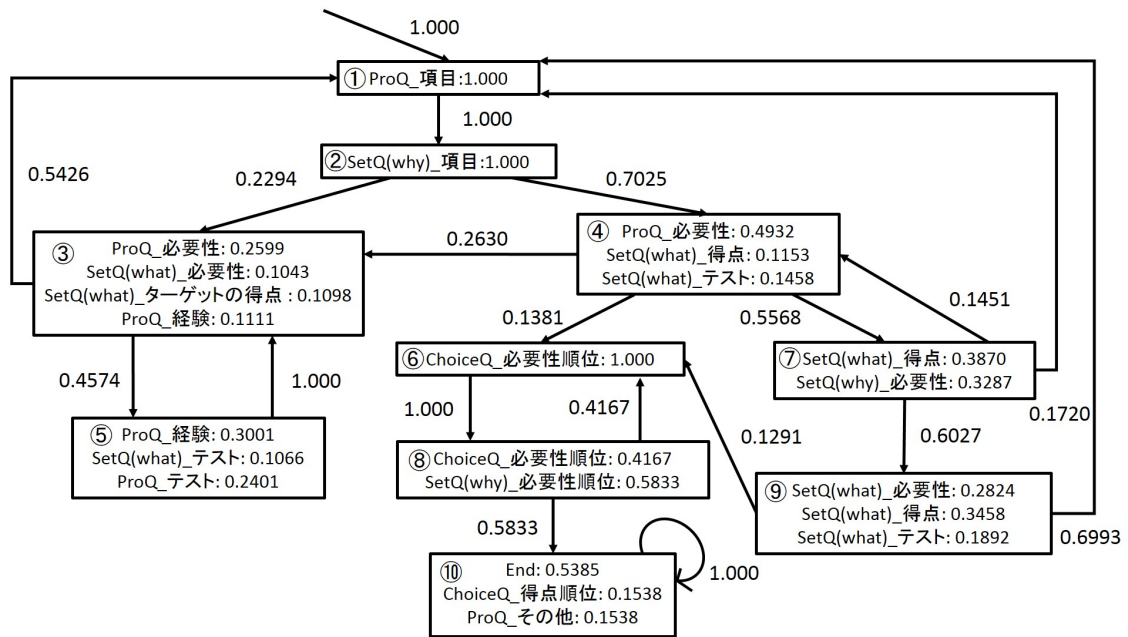


図 8 JDC の質問者の発話から生成した HMM

Mの学習結果が異なるためである．学習した 1,000 個の HMM の中から，最適な HMM を 1 つ選ぶ．選択には，式 (3) に示す Minimum Description Length (MDL) 基準を用いる．

$$MDL = -\ln L + \frac{k \ln n}{2} \quad (3)$$

L は選択した HMM に，その HMM を生成するために用いた学習データを与えて得られた対数尤度の合計である． k は確率が 0 でない遷移確率と観測確率の数の合計， n は学習データの数を表す．

図 8 に JDC から学習した HMM を示す．四角は状態を表し，四角内には観測値，観測確率を示している．状態間の遷移確率は矢印横に示している．観測確率，遷移確率はそれぞれ 0.1 以上のもののみを示している．学習した HMM モデルにおいて，質問者は各テスト項目について適合していたかどうかの質問から始まり，続いて適合していた理由について質問していることがわかる．最初に項目の適合について質問し，その後でテストの得点や対象者の経験についてなど，細かな質問を行っている．

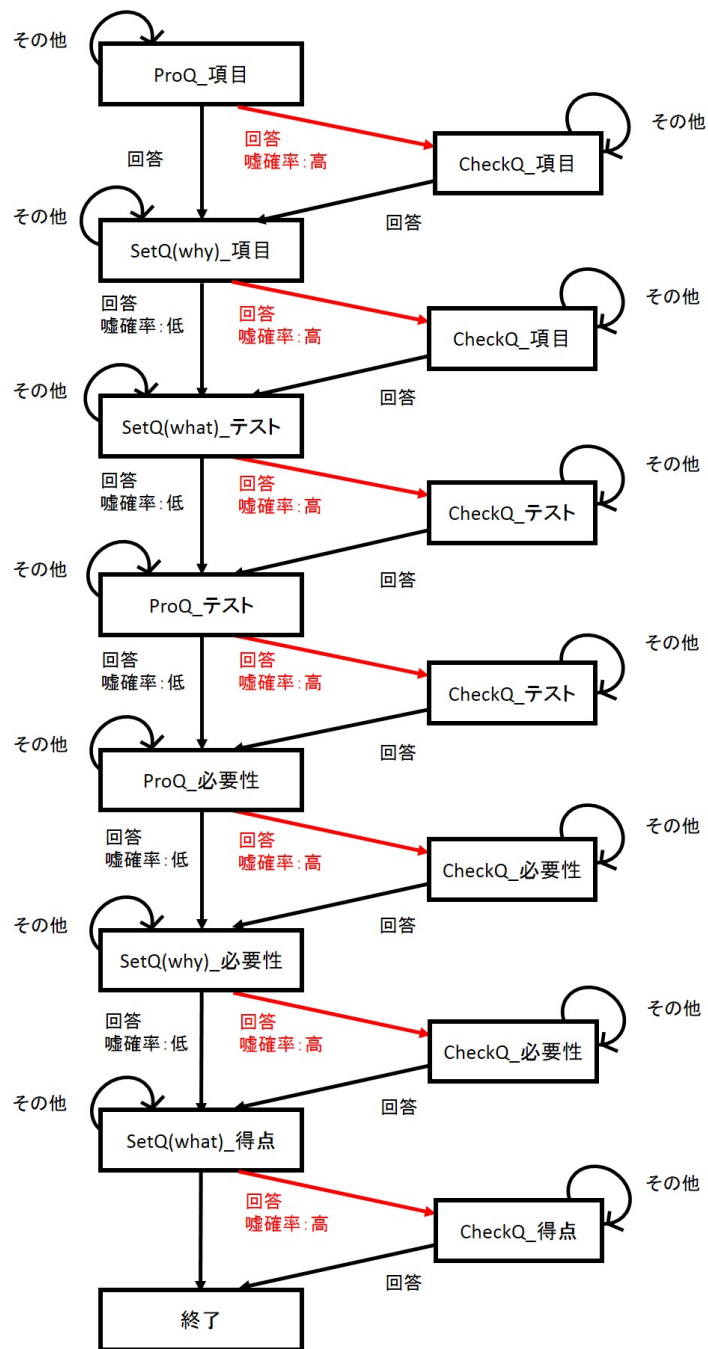


図 9 FSA による本システムのルール記述

6.2.3 ルール作成

前節で、必ず各テスト項目について適合していたかどうかの質問から始まり、続いて適合していた理由について質問するという分析結果を得た。これ以降は、観測確率、遷移確率ともに最大のものを使用し、ルールを生成する。ユーザの発話を質問の内容カテゴリに対する返答であるかどうかで、次の質問に遷移するか繰り返すかを決定するようにする。前節の分析に基づくルールを FSA で記述したものを、図 9 に示す。ユーザの発話ごとに嘘検出を行い、分類されたクラス（嘘、真実）である確率をそれぞれ算出する。嘘である確率が 0.5 以上である発話に対しては、CheckQ の質問を行うようにする。

6.3 Wizard of Oz 法

Wizard of Oz 法 (WoZ) [42] は、Wizard (人間) がシステムのふりをして対話を行うシミュレーション手法である。対話システムの機能の一部を Wizard が代行することにより、システムにおける不完全な箇所を補完することが可能となる。一般的な WoZ 法では、Wizard は応答文を瞬時に作成し、音声合成を用いて応答を行う。特に、ユーザ満足度の最大のボトルネックとも言われている音声認識部は、音声認識誤りの影響を排除するために Wizard が行う場合が多い。音声認識誤りの影響を含めて評価を行う場合は、Wizard は直接音声を聞かず、音声認識の結果を見て応答の作成を行う。

6.4 評価実験

本提案ルールを用いた対話制御部の評価実験を行う。本研究では嘘検出に焦点を当てているため、提案システムが嘘を検出するために有効であるかどうかの評価を行う。そこで、嘘の検出には直接関係のない言語理解部と応答文生成部は Wizard が行うことで、それらの曖昧性に左右されない状態での評価を行う。また、人間がシステムに対して発話を行う際、システムの言語的・音響的特徴に近い発話様式に引き込まれるエンタレイメント現象 [43] が報告されている。本実験

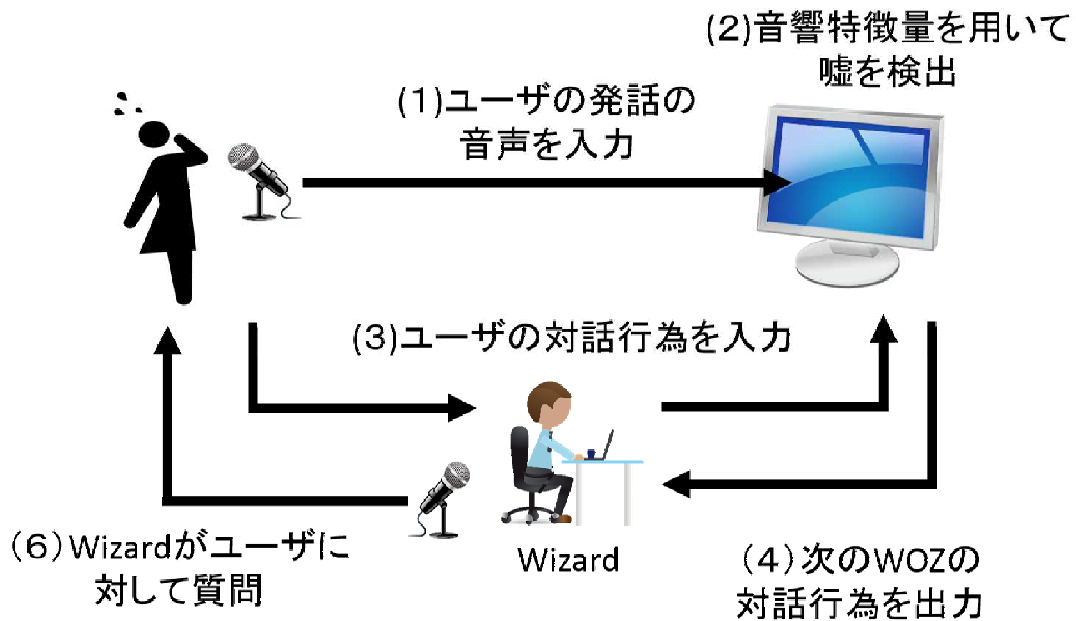


図 10 実験の構成

ではユーザの発話音声から嘘の特徴量を抽出するため、ユーザがエンタテインメント現象の影響を受け、発話が機械的になってはならない。さらに、システムと直接対話を行う場合、2.1 節で述べた嘘をついているときに生じる感情が現れない可能性がある。本実験ではその感情に基づく嘘の特徴量を用いるため、嘘の検出が正しく行えない場合も有り得る。これらを防ぐため、音声合成部は使用せず Wizard が読み上げ、見かけ上は人間同士が対話している仕様にする。

6.4.1 実験条件

嘘の検出には、4 章で有効性を確認できた音響特徴量を使用し、分類器は Bagging [27] を用いる。言語理解部と応答文生成部、音声合成部は、Wizard が行う WoZ 方式の対話実験を行う。応答文生成部では、システムが各質問のテンプレートを表示し、Wizard はそれに沿って応答文を生成する。なお、テンプレートに示す質問長は、5.2 節で有効とされた 1 ～ 10 単語のものとする。比較のためのシステムは、提案システムの他に 3 種類用意する（次節を参照）。それぞれのシステム

は、全く同じ実験インターフェース上で動作する。実験構成は図 10 の通りであり、実験の流れは以下に示す。

1. ユーザが発話を行う。
2. ユーザの発話に対して、システムは嘘の確率を算出する。
3. Wizard はユーザの発話を受け取り、対話行為の定義を基に対話行為タグに変換する。
4. Wizard は変換したユーザの対話行為タグを、システムに入力する。
5. システムは、次に行う対話行為タグを出力する。
6. Wizard は、出力された対話対話タグに従って応答文を生成する。
7. Wizard は、生成した応答文をユーザに向かって発話する。

対話のテーマは、JDC 収録に用いたテストのうち 1 項目を使用する。テストの結果に関わらず、被験者は“適合している”と主張する。

4.5 節で述べた通り、嘘の検出を行うためには、システムを使用するユーザの学習データが必要である。本実験では、学習データが既存であることから、JDC 収録に対象者として参加した 7 名を被験者とし、Wizard として 1 名が参加する。被験者らが使用するシステムの順番はランダムとし、嘘の検出結果に順番による影響がないようにした。被験者は 4 つのシステムとそれぞれ 1 回ずつ対話を行い、計 4 対話を行う。ラベル付与のため、被験者は実験中に発話ごとに“真実”か“嘘”どちらかのボタンを押してもらい、被験者の発話をそれぞれテストデータとし、発話ごとに嘘の検出を行う。それぞれのシステムにおいて、全体の嘘の検出率について評価を行う。

6.4.2 評価用システム

評価用システムとして、以下の 4 つを使用する。

表 15 システムの嘘検出の F 値

	提案ルール	ChoiceQ を含むルール	人間	ランダム
F 値 (%)	81.3	77.8	78.3	75.9

1. 提案ルールベースシステム

6.2.3 節で提案したルールに基づき，対話制御を行うシステムである．

2. CheckQ を含むルールベースシステム

6.2.3 節で提案したルールとは異なり，CheckQ を含めた対話の流れを HMM で学習し，ルールを生成する．このルールに基づき，対話制御を行うシステムである．

3. 人間システム

Wizard は自由に対話行為と対話数を選択し，対話を行う．Wizard は，ユーザの発話の嘘確率を見ながら発話行為を選択することができる．

4. ランダムシステム

等確率でランダムにシステムの対話行為を選択する．1 発話において選ばれる対話行為の個数は 1 である．

6.5 結果と考察

表 15 に，発話ごとに嘘の検出を行った際の F 値をシステムごとに示す．F 値が最も高いことから，嘘の検出には提案ルールベースシステムが最も有効であることがわかる．人間システムよりも F 値が高く，提案システムの方が人間より優れた対話戦略を持つ可能性を示唆している．この理由として，本実験に参加した Wizard は嘘の検出に慣れておらず，有効な対話戦略が決定できなかったためと考えられる．また，CheckQ を含むルールベースシステムよりも，提案システムの方が正しく嘘を検出できている．これより，発話の嘘確率が高い場合に CheckQ の質問を行う方が，嘘の検出に有効であることがわかる．ランダムシステムにお

いては，発話数が極端に短い場合も多く，正しい嘘検出が行えなかったと考えられる．

6.6 本章のまとめ

本章では，最初に対話システムの概要について 6.1 節で説明した．6.2 節では，本システムの対話制御部について説明した．まず，JDC の質問者の発話を HMM で学習し，MDL 基準を用いて最適な HMM を選択した．次に，学習した HMM モデルからルールを生成し，嘘検出率の高い発話に対しては CheckQ の質問を行うように対話制御部を構築した．6.3 節で本研究で用いるシミュレーション手法である WoZ 法について説明し，6.4 節の条件のもと実験を行った．6.5 節では，提案システムの F 値が最も高いという実験結果から，提案システムの有効性を確認した．

7. 結言

7.1 本論文のまとめ

本研究では、自由なコミュニケーション中に嘘を検出する方法として、嘘を検出する対話システムについて提案した。システムを構築するための分析として、日本語において有効な特徴量と、嘘を露呈させやすい質問についての分析を行った。

第2章では、まず2.1節で嘘の特徴について述べた。2.2節では、嘘検出の代表的な手法として、ポリグラフを用いた嘘の検出法とその問題点について説明した。2.3節では、本研究で用いるシステムによる嘘の自動検出手法について説明した。2.4節では、嘘の特徴を露呈させる質問のテクニックについて説明した。2.5節では、既存の嘘を含むコーパスについて紹介した。

第3章では、本研究で新たにコーパスを収集するモチベーションについて述べた。3.1節では、収集するコーパスの概要について説明し、3.2節ではコーパスを収録する環境について述べた。3.3節では、収集したコーパスに対し、本研究で用いる“真実”と“嘘”のラベル付与の基準と方法について説明した。日本語における2話者間の嘘を含むコーパスを収集し、ラベル付与を行った。コーパス中のラベル比は、真実ラベルが約8割を占め、嘘ラベルが約2割であった。これらの条件の下、日本語偽言コーパス（JDC）を構築した。

第4章では、Enos [7] の報告で有効であるとされた特徴量を用いて、日本語のコーパスに対して嘘の検出実験を行った。4.1節では、本研究で使用する音響・言語・個人特徴量について説明した。4.2節では、嘘の検出に用いる機械学習手法である Bagging [27] について説明した。4.3節で実験条件について説明し、4.4節で英語と日本語でほぼ同等の検出率であることを確認した。4.5節では、学習されていない話者に対して嘘の検出が出来るかどうかの実験を行い、学習されていない話者の嘘は検出できないことを確認した。4.6節では、特徴量に関して属性選択を行い、日本語において有効な特徴量を選別した。日本語では、言語特徴量としては動詞の原形、音響特徴量としては F_0 の中央値、母音の平均継続長、パワーの最終フレームが有効であるとされた。さらに、英語と日本語で嘘の検出に有効な特徴量の差異を調べた結果、英語で有効とされた言語特徴量は、日本語で

は有効性を確認できなかった。一方で、音響特徴量は日英間でほぼ同じ特徴量が有効であることを確認した。

第5章では、質問の種類と、嘘検出の容易性についての関係の分析を行った。5.1節では、質問の種類観点から分析を行った。その結果、CheckQが効果的に嘘の特徴を引き出すことができ、嘘の検出に最も有効であることを確認した。5.2節では、質問の長さの観点から分析を行った。短いSU長の質問を行い対話の展開を早くすることで、相手の嘘が露呈しやすいという結果を得た。

第6章では、嘘を検出する対話システムを構築し、評価実験を行った。6.1節で、一般的な対話システムの構成について説明した。6.2節で、HMMを用いて質問者の発話を学習し、そのモデルから本システムの対話制御部で用いる応答ルールを作成し、対話システムを構築した。6.3節では、本システムの評価実験で用いる Wizard of Oz 法について説明した。6.4節の条件のもと実験を行い、6.5節では、提案システムのF値が最も高いという結果から、提案システムの有効性を確認した。

7.2 今後の課題

本研究では、ユーザの発話から音声言語特徴量という表層的な情報のみを用いて嘘の検出を行っている。今後は、ユーザの発話内容にまで踏み込み、話の矛盾という観点からも嘘の検出を行う予定である。ユーザが以前に発話した内容との相違のような時系列的な矛盾や、一般常識に基づく矛盾が考えられる。一般常識に基づく矛盾に関しては、ウェブなどから対話中で用いる大規模な用語データベースを構築し、それを基に検出を行う予定である。

また、人間が嘘をつく際には発話全体が嘘の場合よりも、発話中の一部分のみが嘘であるという場合の方が多い。本研究では、SU中に一部分でも嘘が含まれる場合は、そのSU全体に“嘘”のラベルを付与しており、部分的な嘘は扱っていない。JDCに対してさらに細かいラベル付与を行い、特徴量抽出の粒度を細かくすることで、部分的な嘘の検出が可能になると考える。

本研究はシステムの評価をWoZで行っているが、アプリなどを実装する場合も含めて、全自動化する必要がある。そのために、言語理解部、応答文生成部の構

築も目指していく予定である。さらに、本研究では音声合成部を用いず、Wizardが発話まで行う仕様にした。合成音声部を用いてシステムを全自動化するために、まず、質問者の音響特徴量を加え、声質という点からも嘘の検出に有効な質問の分析を行う。次に、質問者のモデルを基に対話システムの合成音声の発話様式を変化させることで、エンタテインメント現象の影響によりユーザの発話が機械的になるのを防ぐことができると考えられる。これにより、2.1節で述べた嘘をついているときに生じる感情が人間と対話しているときと同様に現れ、嘘の検出を正しく行うことができると考える。

本研究は、人手で作成したルールによって対話システムの対話制御部を実現した。しかしながら、全ての対話状態に対応するルールを記述することは難しく、限定的な応答になるという問題点がある。この問題を解決するために、POMDPなどの統計的手法を用いて対話制御部の構築を行い、曖昧性に対しても頑健な対話システムの構築を目指す。

謝辞

本学情報科学研究科の中村哲教授には主指導教官として、熱心なご教示とご助言を賜りました。また、ホンダ・リサーチ・インスティテュート・ジャパンでの貴重なインターンシップの機会を賜りました。心より感謝致します。

本学情報科学研究科の小笠原司教授には副指導教官として、貴重なご助言を賜りました。心より感謝致します。

本学情報科学研究科の戸田智基准教授には、副指導教官として熱心なご教示とご助言を賜りました。特に、対外発表の際には、多くのご助言を賜りました。心より感謝致します。

本学情報科学研究科のGraham Neubig助教には、副指導教官として多くの熱心なご教示とご助言を賜りました。日頃から進捗を気にかけて頂き、優しい言葉で励まして下さいました。また、英語での論文執筆の際には、著者のつたない文章を解説して頂き、大変貴重なご助言を賜りました。心より感謝致します。

本学情報科学研究科のSakriani Sakti助教には、副指導教官として熱心なご教示とご助言を賜りました。特に、音声に関する場面においては、多くのご助言を賜りました。心より感謝致します。

本知能コミュニケーション研究室の諸氏には、研究を通してご討論、ご協力頂きました。心より感謝致します。

最後に、陰ながら支えてくれた父と母、そして祖父母に、心から感謝します。

参考文献

- [1] Bella M DePaulo, Deborah A Kashy, Susan E Kirkendol, Melissa M Wyer, and Jennifer A Epstein. Lying in everyday life. *Journal of personality and social psychology*, Vol. 70, No. 5, p. 979, 1996.
- [2] 村井潤一郎, 島田将喜, 菊池史倫, 佐藤拓, 大幡直也, 武田美亜, 林創, 上宮愛, 野瀬出, 鈴木菜実子. 嘘の心理学, 第4巻. ナカニシヤ出版, 2013.
- [3] Robert Hopper and Robert A Bell. Broadening the deception construct. *Quarterly Journal of Speech*, Vol. 70, No. 3, pp. 288–302, 1984.
- [4] Paul Ekman. *TELLING LIES*. W. W. Norton & Company, 1985.
- [5] Aldert Vrij, Pär Anders Granhag, Samantha Mann, and Sharon Leal. Outsmarting the liars: Toward a cognitive lie detection approach. *Current Directions in Psychological Science*, Vol. 20, No. 1, pp. 28–32, 2011.
- [6] 松田いづみ, 廣田昭久, 小川時洋, 高澤則美, 繁榊算男. ポリグラフ検査 ポリグラフ検査における統計的判定法. 生理心理学と精神生理学, Vol. 27, No. 1, pp. 45–56, 2009.
- [7] Frank Enos. *Detecting deception in speech*. PhD thesis, Citeseer, 2009.
- [8] Verónica Pérez-Rosas and Rada Mihalcea. Cross-cultural deception detection. *Proc. ACL*, pp. 440–445, 2014.
- [9] Mark G Frank. Commentary: On the structure of lies and deception experiments. *Cognitive and social factors in early deception*, pp. 127–146, 1992.
- [10] David T Lykken. Polygraphic interrogation. *Nature*, Vol. 307, pp. 681–684, 1984.
- [11] William W Cohen. Fast e effective rule induction. In *Proceedings of the Twelfth International Conference on Machine Learning, Lake Tahoe, California*, 1995.

- [12] Yoshimasa Ohmoto, Kazuhiro Ueda, and Takehiko Ohno. A method to detect lies in free communication using diverse nonverbal information: Towards an attentive agent. In *Active Media Technology*, pp. 42–53. Springer, 2009.
- [13] Inbau F. E., Reid J. E., Buckley J. P., and Jayne B. C. Criminal interrogation and confessions aspen. *Gaithersburg, MD*, 2001.
- [14] Frank Horvath, Brian Jayne, and Joseph Buckley. Differentiation of truthful and deceptive criminal suspects in behavior analysis interviews. *Journal of Forensic Sciences*, 1994.
- [15] Aldert Vrij. *Detecting lies and deceit: Pitfalls and opportunities*. John Wiley & Sons, 2008.
- [16] Aldert Vrij, Samantha Mann, Susanne Kristen, and Ronald P Fisher. Cues to deception and ability to detect lies as a function of police interview styles. *Law and human behavior*, Vol. 31, No. 5, p. 499, 2007.
- [17] Aldert Vrij, Pär Anders Granhag, and Stephen Porter. Pitfalls and opportunities in nonverbal and verbal lie detection. *Psychological Science in the Public Interest*, Vol. 11, No. 3, pp. 89–121, 2010.
- [18] Aldert Vrij, Sharon Leal, Pär Anders Granhag, Samantha Mann, Ronald P Fisher, Jackie Hillman, and Kathryn Sperry. Outsmarting the liars: The benefit of asking unanticipated questions. *Law and human behavior*, Vol. 33, No. 2, pp. 159–166, 2009.
- [19] Maria Hartwig, Pär Anders Granhag, Leif A Strömwall, and Ola Kronkvist. Strategic use of evidence during police interviews: when training to detect deception works. *Law and human behavior*, Vol. 30, No. 5, p. 603, 2006.
- [20] Hayley Hung and Gokul Chittaranjan. The IDIAP wolf corpus: exploring group behaviour in a competitive role-playing game. In *Proc. The international conference on Multimedia*, pp. 879–882. ACM, 2010.

- [21] Bella M DePaulo, James J Lindsay, Brian E Malone, Laura Muhlenbruck, Kelly Charlton, and Harris Cooper. Cues to deception. *Psychological bulletin*, Vol. 129, No. 1, p. 74, 2003.
- [22] Taku Kudo, Kaoru Yamamoto, and Yuji Matsumoto. Applying conditional random fields to Japanese morphological analysis. *Proc. EMNLP*, Vol. 4, pp. 230–237, 2004.
- [23] Hiroya Takamura, Takashi Inui, and Manabu Okumura. Extracting semantic orientations of words using spin model. *Proc. ACL*, pp. 133–140, 2005.
- [24] Julia Bell Hirschberg, Stefan Benus, Jason M Brenier, Frank Enos, Sarah Friedman, Sarah Gilman, Cynthia Girand, Martin Graciarena, Andreas Kathol, Laura Michaelis, et al. Distinguishing deceptive from non-deceptive speech. *Proc. Eurospeech*, 2005.
- [25] Kimmen Sjolander. Tcl/tk snack toolkit. 2004. <http://www.speech.kth.se/snack/>.
- [26] Daniel Povey, Arnab Ghoshal, Gilles Boulianne, Lukas Burget, Ondrej Glembek, Nagendra Goel, Mirko Hannemann, Petr Motlicek, Yanmin Qian, Petr Schwarz, Jan Silovsky, Georg Stemmer, and Karel Vesely. The Kaldi speech recognition toolkit. *Proc. ASRU*, 2011.
- [27] Leo Breiman. Bagging predictors. *Machine learning*, Vol. 24, No. 2, pp. 123–140, 1996.
- [28] Vladimir Vapnik. *The nature of statistical learning theory*. Springer Science & Business Media, 2000.
- [29] Zijan Zheng. Generating classifier committees by stochastically selecting both attributes and training examples. *PRICAI '98: Topics in Artificial Intelligence*, pp. 12–23. Springer, 1998.

- [30] Bradley Efron. Bootstrap methods: another look at the jackknife. *The annals of Statistics*, pp. 1–26, 1979.
- [31] Yongheng Zhao and Yanxia Zhang. Comparison of decision tree methods for finding active objects. *Advances in Space Research*, Vol. 41, No. 12, pp. 1955–1959, 2008.
- [32] Pamela Meyer. *LIE SPOTTING*. Griffin, 2011.
- [33] 河原達也, 荒木雅弘. 音声対話システム. オーム社, 2006.
- [34] Marvin Lee Minsky, Marvin Minsky, and Marvin Minsky. *Semantic information processing*, Vol. 15. MIT press Cambridge, MA, 1968.
- [35] 坪井宏之. キ-ワ-ドスポッティングに基づく連続音声理解. 信学技報, pp. SP91–95, 1991.
- [36] Roberto Pieraccini and Juan Huerta. Where do we go from here? research and commercial spoken dialog systems. In *6th SIGdial Workshop on Discourse and Dialogue*, 2005.
- [37] Esther Levin, Roberto Pieraccini, and Wieland Eckert. A stochastic model of human-machine interaction for learning dialog strategies. *Speech and Audio Processing, IEEE Transactions on*, Vol. 8, No. 1, pp. 11–23, 2000.
- [38] Stephen G Pulman, et al. Conversational games, belief revision and bayesian networks. In *Computational Linguistics in the Netherlands*. Citeseer, 1996.
- [39] Jason D Williams and Steve Young. Partially observable markov decision processes for spoken dialog systems. *Computer Speech & Language*, Vol. 21, No. 2, pp. 393–422, 2007.
- [40] Jonathan Allen. Synthesis of speech from unrestricted text. *Proceedings of the IEEE*, Vol. 64, No. 4, pp. 433–442, 1976.

- [41] 目黒豊美, 東中竜一郎, 堂坂浩二, 南泰浩. 聞き役対話の分析および分析に基づいた対話制御部の構築. 情報処理学会論文誌, Vol. 53, No. 12, pp. 2787–2801, 2012.
- [42] Norman M Fraser and G Nigel Gilbert. Simulating speech systems. *Computer Speech & Language*, Vol. 5, No. 1, pp. 81–99, 1991.
- [43] 杉山昂太郎, 戸田智基, Graham Neubig, Sakriani Sakti, 中村哲. エントレインメント現象を用いた音声認識に適した発話様式誘導に関する調査. 日本音響学会 2014 年春季研究発表会 (ASJ), 東京, 3 2014.

発表リスト

国際会議

1. Yuiko Tsunomori, Graham Neubig, Sakriani Sakti, Tomoki Toda and Satoshi Nakamura. An Analysis Towards Dialogue-based Deception Detection. International Workshop Series on Spoken Dialogue System Technology (IWSDS), Korea, Jan. 2015.

研究会

1. 角森唯子, Graham Neubig, Sakriani Sakti, 戸田智基, 中村哲. 対話中の音声言語特徴量に着目した嘘の検出法と日英間比較. NLP 若手の会 (YANS) 第9回シンポジウム, 神奈川, Sep. 2014.
2. 角森唯子, Graham Neubig, Sakriani Sakti, 戸田智基, 中村哲. 対話中における嘘検出に有効な特徴量比較と質問分析. 第72回人工知能学会 音声・言語理解と対話処理研究会 (SIG-SLUD), 東京, Dec. 2014.

大会発表

1. 小田悠介, 笹倉隆史, 角森唯子, 赤部晃一, 大島悠司, 鶴田さくら, 西垣友理, 田代 光輝, 中村 哲. オンラインコミュニティの盛衰と挨拶表現の関連性の調査. 情報社会学会 2014 年度研究発表大会, 神奈川, May. 2014.
2. 角森唯子, Graham Neubig, Sakriani Sakti, 戸田智基, 中村哲. 対話中の音声言語特徴量に着目した嘘の検出法と日英間比較. 日本音響学会 2013 年秋季研究発表会 (ASJ), 北海道, Sep. 2013.