

意思決定理論に基づく発話プランニング

A Framework of Decision-Theoretic Utterance Planning

乾 健太郎*
Kentaro Inui

徳永 健伸*
Takenobu Tokunaga

田中 穂積*
Hozumi Tanaka

* 東京工業大学大学院情報理工学研究科
Dept. of Computer Science, Tokyo Institute of Technology, Tokyo 152, Japan.

1996年10月11日 受理

Keywords: dialogue, decision theory, utterance planning, plan recognition, Bayesian network, uncertain reasoning.

Summary

In human-computer dialogue, a diverse range of situations arise from uncertainty in reasoning regarding the user's beliefs and plans. The dialogue system is required to plan the contents of its utterances under such uncertain conditions. This paper presents a decision-theoretic framework of content planning for dialogue systems. In our framework, the uncertain domain knowledge and user's model are both declaratively described as probabilistic constraints, which represent uncertain situations as probabilistic distributions over possible worlds. We estimate this probabilistic distribution using a Bayesian network. Given a utility function to evaluate the degree of goal attainment in each possible world, the system chooses the content of its utterance based on the expected utility maximization principle. Our framework can perform content planning under uncertain conditions without any procedural planning strategy.

1. ま え が き

目的指向の対話を協調的にすすめるためには、相手の発話と文脈から相手の信念やプランを推定し、その結果をもとに自分の発話をプランニングすることが必要だとされている。相手の知識や信念が推定できれば、それと関連づけて新しい情報を提供することができる [Akiha 94]。また、相手のプランがわかれば、たとえ明示的には質問されなくても、そのプランに関して重要な情報があればそれを付加的に提供したり [Litman 90, 山田 94]、プランの誤りを教えたり [Pollack 90] することができる。本論文では、対話システムとユーザの対話を想定し、ユーザの信念やプランを推定しながら協調的に発話する対話システムを実現するための新しい機構について論じる。

ユーザの信念やプランに関する情報をユーザモデルと呼ぶ。対話プランニングに関する従来の研究の多くは、ユーザモデルを推定する作業と発話を選択する作

業を独立に扱ってきた。しかしながら、これら2つの作業の間には本来もっと緊密な相互作用が存在するはずである。たとえば、ユーザのプランをどの程度細かく推定すべきかという判断は発話選択の作業に依存する。一般に、ユーザプランを詳細に推定しようとする、プランの候補数が増大するため、個々の候補の尤度は低くなり、推定誤りの可能性が高くなる。一方、ユーザプランの詳細までわかっていた方がユーザにとって有益な情報を提供できる。したがって、ユーザのプランをどの程度細かく推定すべきかという判断は、できるだけ推定誤りを低くしようとするプラン推定側と有益な情報提供のためにできるだけ詳細なユーザプランの情報を要求する発話選択側とのトレードオフを最適化する作業であると考えられる。これを実現するためにはユーザモデルの推定と発話選択の間の緊密な相互作用が必要である。

本稿では、ユーザモデルを含む状況の不確実性を可能世界の確率分布として表現し、発話によって得られる効用の期待値にもとづいて発話候補の優先度が決ま

る意思決定理論的枠組 [乾 95] について述べる。この枠組では、個々の作業を行うための手続き的知識を必要とせず、状況に依存しない原理的なヒューリスティクスによって発話の優先度が決まるので、多様な状況に追従するシステムの実現が期待できる。また、ユーザモデルの推定と発話選択を統合的に行うため両者間の緊密な相互作用を実現することができる。

2. 発話選択の意思決定理論的枠組

2.1 信念ベースと信念モデル

対話システムは対話相手（ユーザ）の信念を含む外部の世界の情報を部分的にしか持たないので、対話システムの世界知識は一般に不確実である。我々の枠組では、この不確実性を可能世界の確率分布として表現し、可能世界の確率分布を信念ベースと信念モデルによって与える。

信念ベースは、対象領域についてシステムが持つ信念とユーザが持っているシステムが信じる信念からなり、確率的制約として記述される。信念ベースは対象領域に関する広範な制約の集合であり、我々は確率 Horn 論理 [Poole 93a] を用いてこれを記述する。一方、システムが対話のある時点で実際に参照する制約は、その時点での話題に関連するものに限られる。本論文ではこれをとくに信念モデルと呼ぶ。すなわち、信念モデルは信念ベース内の確率的制約のうち、話題に関連する制約を具現化したものである。我々はこれを Bayesian ネットワーク [Pearl 88] で表現する。以下では、信念ベースから信念モデルを動的に生成する作業を信念モデルの動的構築と呼び、伝達内容を選択する作業を信念モデルの評価と呼ぶ（図 1 参照）。

以下本章では、信念モデル上での発話選択の仕組みを示した後、信念ベースの記述方法について概説し、信念モデルの動的構築に触れる。

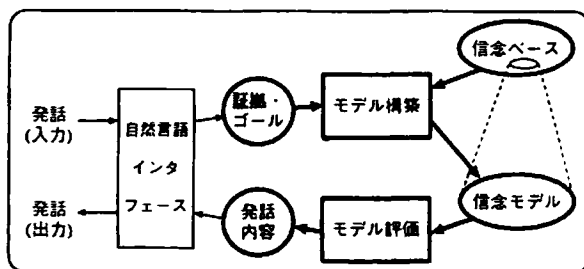


図1 信念ベースと信念モデル

2.2 信念モデル

Bayesian ネットワークは $\langle V, D, e \rangle$ の組で定義される非循環有向グラフである。 V は確率変数 V の集合である。各変数は命題を表し、それぞれネットワークの1つのノードに対応する。以下、変数 V に対する値割当てを A_V と表記する。また、 $V = \text{true}$ を $+V$ 、 $V = \text{false}$ を $-V$ と書くことがある。 D は局所的確率分布 D_V の集合である。 D_V は変数 V と V の親の変数^{*1}の集合 $\text{prt}(V)$ との局所的依存関係を規定する。これは V と $\text{prt}(V)$ のすべての値割当ての組み合わせについての条件つき確率 $P(A_V | A_{\text{prt}(V)})$ で与えられる。最後に e は観察された証拠（値割当て）の集合である。

Bayesian ネットワークでは、変数間にいくつかの条件つき独立性を仮定することにより、任意の完全値割当て^{*2} A_V の結合確率を D の局所的条件つき確率の積として式 (1) のように評価することができる。

$$P(A_V) = \prod_{V \in V} P(A_V \in A_V | A_{\text{prt}(V)} \in A_V) \quad (1)$$

また、証拠でない変数の集合 W に対する値割当て A_W の結合確率は

$$P(A_W) = \sum_{A_V \supset A_W} P(A_V) \quad (2)$$

で与えられる。証拠 e のもとでの A_W の条件つき確率は

$$P(A_W | e) = \frac{P(A_W, e)}{P(e)} \quad (3)$$

で与えられる。ここで、1つの A_W を1つの可能世界と考えれば、可能世界の確率分布を式 (3) で与えることができる。このように、Bayesian ネットワークをもちいると、局所的な確率分布の組み合わせによって大域的な可能世界の確率分布を表現することができる。

信念モデルが表現するのは、システムがすでに持っている証拠 e のもとでの可能世界の確率分布 $P(A_W | e)$ である。この確率分布はユーザを含む世界の状態に関する推定結果に相当する。

2.3 発話の優先度

システムが持っている証拠は、システムが対話開始

*1 ネットワーク上で変数からある変数に向かう弧があるとき、前者を後者の親、後者を前者の子と呼ぶ。

*2 ネットワークに含まれるすべての確率変数に値を割当てるとき、これを完全値割当てと呼ぶ。

前から持っている証拠と対話中に新たに得られる証拠からなる。対話中に観察するユーザの発話はユーザの信念の一部に関する証拠とみなすことができる。また、システムが発話する場合は、ユーザにとっては対話相手の発話を観察することになるので、システムの信念に関するユーザの信念に証拠を与える行為とみなすことができる。

まず、システムがある命題の値割当てに関する情報 A_Q をユーザに伝達する場合を考えよう。以下、命題の値割当てに関する情報 A_Q の伝達を情報伝達行為と呼び、 $\text{inform}(A_Q)$ と表記する。システムがユーザに情報 A_Q を伝達すると、ユーザモデルを表す信念モデルに証拠 A_Q が新たに加わり、可能世界の確率分布が $P(A_W|e)$ から $P(A_W|A_Q, e)$ に変化する。

システムがユーザとの対話をすすめるのは何らかのゴールを達成するためであると考えられる。このゴールはシステムが対話開始時から持っている場合もあれば、ユーザから与えられる場合もある。いずれにせよ、システムにとってはゴールをより満足する可能世界がより望ましい。そこで、可能世界 A_W の望ましさを程度を評価する効用関数 $U(A_W)$ を導入し、証拠 e のもとでの可能世界の確率分布の望ましさを効用の期待値 $EU(e)$ で計量することにする。

$$EU(e) \stackrel{\text{def}}{=} \sum_{A_W} U(A_W) P(A_W|e) \quad (4)$$

上で述べたように、システムが発話すると可能世界の確率分布が変化するので、期待効用も変化する。そこで、図2に示すように、期待効用をできるだけ高くするように発話を選択することにすれば、ゴールの達成をめざして合理的に対話をすすめる対話者のふるまいを部分的に実現することができると考えられる。

システムが情報伝達行為 $\text{inform}(A_Q)$ を実行すると、可能世界の確率分布が変化し、 $P(A_W|A_Q, e)$ になる。このときの期待効用は、

$$EU(A_Q, e) = \sum_{A_W} U(A_W) P(A_W|A_Q, e) \quad (5)$$

である。これを最大にする情報伝達行為 $\text{inform}(A_Q)$ が最適な行為だと考えられるので、証拠 e のもとでの $\text{inform}(A_Q)$ の優先度 $Pr(\text{inform}(A_Q)|e)$ を式(5)の期待効用で与えることにする。

$$Pr(\text{inform}(A_Q)|e) \stackrel{\text{def}}{=} EU(A_Q, e) - C \quad (6)$$

ただし、 C は情報伝達行為の実行コストとする。

ここで、 $\forall A_W: 0 \leq U(A_W) \leq 1$ が成り立つよう

に U を決めることができると仮定しよう。このとき、

$$A_Q \in \{\text{true}, \text{false}\} \quad (7)$$

$$P(+G|A_W) = U(A_W) \quad (8)$$

で与えられる確率変数 G を信念モデルに加えると、式(4)より $EU(e) = P(+G|e)$ が成り立つ。したがって、式(6)は式(9)のように変形できる。

$$\begin{aligned} Pr(\text{inform}(A_Q)|e) &\stackrel{\text{def}}{=} EU(A_Q, e) - C \\ &= P(+G|A_Q, e) - C \\ &= \frac{P(A_Q|+G, e)P(+G|e)}{P(A_Q|e)} - C \quad (9) \end{aligned}$$

式(9)より、証拠が e のモデルと $+G \wedge e$ のモデルそれぞれについて各候補 A_Q の後驗確率を計算すれば、式(9)を最大にする情報伝達行為を選択することができる。同じモデルの各変数の後驗確率は1度のモデル評価で同時に計算できるので、モデル評価を2度行えば、情報伝達行為の個々の候補の優先度が計算できるので、この最大化問題は効率的に解くことができる。

ユーザのプランがシステムにとってあいまいであり、プランによって有益な情報伝達行為が異なるような場合、そのあいまい性に関してユーザに問い合わせる方がよい場合がある。そこで、ある命題 R の値割当てを問い合わせる情報収集行為 $\text{ask}(R)$ を考えよう。システムが情報収集行為 $\text{ask}(R)$ を実行すると、ユーザは R の値割当てに関する自分の信念 A_R をシステムに伝達する ($\text{inform}(A_R)$) と期待でき、システムはその情報に依存して最適な情報伝達行為 $\text{inform}(A_Q)$ を選択することができると考えられる。このとき得られる期待効用の最大値は、

$$\max_{A_Q} EU(A_Q, A_R, e) \quad (10)$$

である。ただし、ユーザの回答が何であるかは実際にそれを観察するまで確率的にしかわからないので、システムの情報収集行為 $\text{ask}(R)$ 、ユーザの回答 $\text{inform}(A_R)$ 、システムの情報伝達行為 $\text{inform}(A_Q)$ という過程で得られる期待効用は式(10)の期待値として計算することになる。これを $\text{ask}(R)$ の優先度 $Pr(\text{ask}(R)|e)$ としよう。

$$\begin{aligned} Pr(\text{ask}(R)|e) &\stackrel{\text{def}}{=} \\ &\sum_{A_R} \max_{A_Q} EU(A_Q, A_R, e) P(A_R|e) - C' \quad (11) \end{aligned}$$

ただし、 C' はシステムの情報収集行為、ユーザの回答、システムの情報伝達行為に要するコストの合計である。

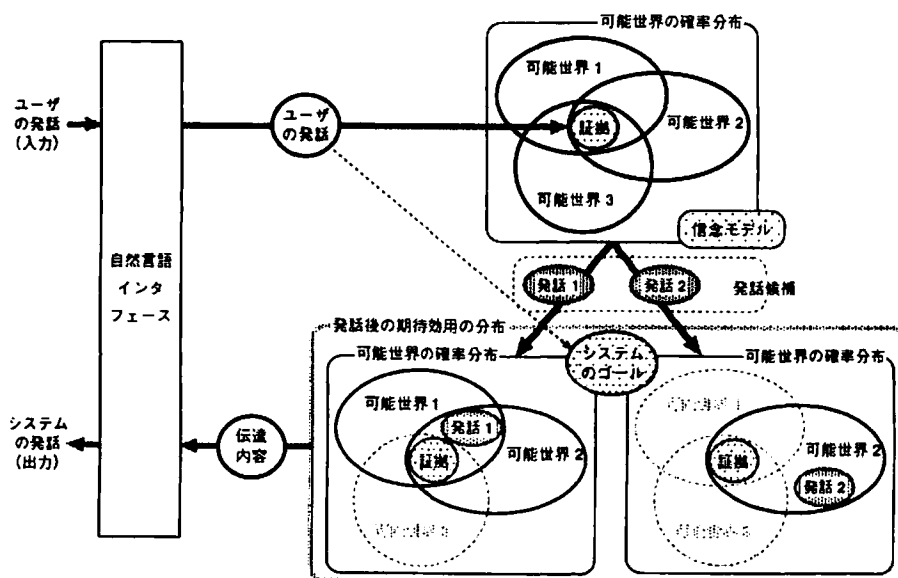


図2 発話選択の仕組み

式 (11) は Bayes 則により次のように書ける。

$$Pr(ask(R)|e) = \sum_{A_R} \left(\max_{A_Q} \frac{P(A_Q, A_R | +G, e) P(+G|e)}{P(A_Q, A_R | e)} \right) \cdot P(A_R|e) - C' \quad (12)$$

したがって、式 (9) と同様、ask(R) の優先度も証拠が e の場合と $+G \wedge e$ の場合の 2 つのモデルを評価すれば計算できる。

システムは対話の各時点で、式 (9)、式 (12) によって決まる優先度のもっとも高い発話を選択する。たとえば、現在の証拠集合のもとで最適な情報伝達行為を実行して得られる期待効用に比べ、情報収集後にユーザからの回答に依存して最適な情報伝達行為を実行する場合の期待効用の方が、情報収集のコストを支払っても高い場合、システムは情報収集行為を選択する。発話の優先度は、可能世界の確率分布を与える信念モデルの上で直接評価される。我々の枠組では、ユーザモデルの推定を発話選択と独立に行うのではなく、両者の作業を完全に統合しているため、それらの間の緊密な相互作用を実現することができる。

2.4 信念ベースの記述形式

信念ベースは確率的制約の集合によって与える。確率的制約の記述形式は、信念モデルの構築に必要な情報を与えるものであればどんなものでもよい。我々は現在 Poole の確率 Horn 論理 [Poole 93a] を使っている。Poole の表現形式は一階述語 Horn 論理に確率付きの仮説集合を導入したものである。確率 Horn 論

理は、それによって記述された理論の任意のインスタンスに対し等価な Bayesian ネットワークを作ることができることが証明されている [Poole 93a]。以下に確率 Horn 論理の概要を述べる。

- (1) 項は、論理変数 (大文字ではじまる文字列)、定数 (小文字ではじまる文字列)、または $f(t_1, \dots, t_n)$ (ただし、 f は関数記号、 t_i は項) である。要素式は、 $p(t_1, \dots, t_n)$ (ただし、 p は述語記号、 t_i は項) である。
- (2) 確定節の形は、 $a, a_1, \dots, a_n$ をそれぞれ要素式とするとき、 a 、または $a \leftarrow a_1, \dots, a_n$ である。
- (3) 排他宣言の形は、 $disjoint([h_1 : p_1, \dots, h_n : p_n])$ である。ただし、 h_i は要素式、 p_i は $0 \leq p_i \leq 1$ 、 $p_1 + \dots + p_n = 1$ をみたす実数である。このとき、 h_i を仮説と呼ぶ。ある排他宣言に含まれる h_i が変数 X を持つとき、同じ排他宣言に含まれる他の仮説も変数 X を持たなければならない。
- (4) 確率 Horn 論理で記述される理論 T は、確定節の集合と排他宣言の集合からなる。ただし、ある基底要素式 h がある排他宣言の仮説のインスタンスであるとき、 h は T の他の排他宣言に現れるあらゆる仮説のインスタンスでない。
- (5) 理論 T が与えられたとき、 T の事実集合は T の確定節と次の形の節の集合である。 $false \leftarrow h_i, h_j$ ただし、 h_i, h_j ($i \neq j$) は T 中の同じ排他宣言に現れる任意の仮説の組である。
- (6) T の仮説集合は T の排他宣言に現れる仮説の集合である。

- (7) 基底要素式 h' がある仮説 h のインスタンスであり, $h:p$ が T のある排他宣言に現れるとき, $P(h') = p$ である. ただし, $P(Q)$ は Q が成り立つ確率を表す.

ただし以下では, 記述を簡略化するため, 上の (2) の形式を以下のように拡張する.

- (2') 「 $a_1, \dots, a_n$ が成り立つとき a が確率 p で成り立つ」という制約を

$$a \leftarrow_p a_1, \dots, a_n.$$

と書く. これは, 各節に固有の仮説 r を新たに作り,

$$a \leftarrow a_1, \dots, a_n, r.$$

$$\text{disjoint}([r:p, \neg r:1-p]).$$

と書くことと等価である.

2.5 信念モデルの動的構築

上で述べたように, 確率 Horn 論理で記述した信念ベースの任意のインスタンスは等価な Bayesian ネットワークに変換することができる. つまり, 確率的制約のインスタンスは, 無限の大きさを持つ仮想的な Bayesian ネットワークと見なせる. 信念モデルの構築は, この仮想的な Bayesian ネットワークから伝達内容のプランニングに関連する部分だけを切り出す作業と見なせる.

このような動的モデル構築についてはすでに様々な先行研究がある [Breese 94]. たとえば Breese は, Poole の確率 Horn 論理と類似した記述形式の信念ベースから Bayesian ネットワークを動的構築する手続きを提案している [Breese 92]. Breese の手続きは, 注目したい変数を起点に後向き推論と前向き推論をそれぞれ行い, 根ノード, 葉ノード, 証拠ノードが見つかった時点で推論を終了する. Bayesian ネットワークはこの推論手続きによって具現化された要素式と証拠から構築される. この手続きで Bayesian ネットワークを構築すると, 注目したい変数とネットワークに含まれない任意の変数が, 与えられた証拠のもとで条件つき独立になる. 変数 X と Y が証拠 e のもとで条件つき独立であるとは, $P(X|Y, e) = P(X|e)$ が成り立つということである. すなわち, Breese の手続きをつかえば, 注目する変数の後驗確率を計算するのに必要十分な大きさのモデルを構築することができる.

この手続きは Poole の確率 Horn 論理にもそのまま適用することができる. 信念モデルは, まず Breese の手続きによって確率的制約を具現化し, 次にそれを Poole の変換手続き [Poole 93a] で Bayesian ネット

ワークに変換することによって得られる. 確率的制約を具現化する際にシステムのゴールを起点に Breese の手続きを実行することにより, ゴールに対応する確率変数と条件つき独立な変数を含まない信念モデルを作ることができる. Poole の変換手続きでは, 個々の排他宣言のインスタンスをそれぞれ 1 つの確率変数で表現し, Bayesian ネットワーク上の 1 つの根ノードに対応させる. このとき, 確率変数の値域は対応する (具現化された) 排他宣言の要素式の集合である. その他, 排他宣言に含まれない要素式のインスタンスはそれぞれ 1 つの Boolean 型確率変数に変換する. また, 確率つき確定節で定義される要素式間の依存関係は, Bayesian ネットワークを構成する局所的条件つき確率に変換する.

3. 例 題

3.1 対 話 例

次の例 (文献 [Nagao 93, van Beek 91] の例題を修正) を考えよう.

例 (1) 登場人物は, 専門家 s , 料理人 c , 菜食主義の客 g の 3 人である. c は g に出す料理を作ろうとしており, 料理について s に相談することができる.

c : いまマリナラソースを作るところです. ワインは赤がいいですか?

s : パスタ料理には白ワインが合います.

s は, c の質問から c がマリナラソースを作っていることを知り, c のプランがスパゲティーマリナラ, フェットチーマリナラ, チキンマリナラのどれかであることがわかる. ここでさらに, s は g が菜食主義であることを知っており, c もそれを知っていると (s が) 信じているとする.すると s は, プランが肉料理 (チキンマリナラ) でなく, スパゲティーマリナラかフェットチーマリナラのいずれかだろうと推定できるが, いずれの料理もパスタ料理であり, パスタ料理には白ワインが適しているので, 上の例のように答えることができる. c のプランにあいまい性があるにもかかわらず, s が c にとって有益な答えをかせしている点に注意したい. 従来の枠組では, 2 つのプランの候補 (スパゲティーマリナラとフェットチーマリナラ) に顕著な尤度の差がなければ, 「(s) あなたが作ろうとしている料理はスパゲティーマリナラですか, フェットチーマリナラですか?」のように聞き返さなければならない. しかしながら, s のこの質問はあきらかに不必要である.

例 (2) 登場人物は例 (1) と同じである. ただし,

g が肉食主義者であることを s 自身は知っているが、c は知らない (と s が信じている) とする。この状況で次の対話を考える。

c: いまマリナラソースを作るところです。

s: 客はベジタリアンだよ。

例 (1) と同様、c のプランの候補はスパゲティーマリナラ、フェトチーニマリナラ、チキンマリナラの3つである。しかしながら、いま s は c には肉料理を避ける理由がないと思っているので、s の信念の中ではチキンマリナラのプランは棄却されない。さらに、c が肉料理を好んで作ると s が信じているとすると、チキンマリナラの可能性が高くなる。一方、s は g が肉食主義者であることを知っているため、チキンマリナラのプランが適当でないことに気づく。そこで g が肉食主義者であることを c に教えれば、c もチキンマリナラのプランが適当でないことに気づき、失敗を回避できる。

ここで、s が c に客が肉食主義者であるという情報を与えるだけで、c 自身がチキンマリナラのプランが適当でない気がつくことが期待できることに注意したい。s は自分の発話に対する c の推論を予測した上でその発話を実行したと考えられる。このことは、伝達内容をプランニングする作業のなかに聞き手 (ユーザ) の推論を予測する作業が埋め込まれていることを意味する。聞き手の推論の予測は、簡潔な発話で効率的な情報伝達を実現するためには特に重要である。多くの場合、聞き手は、話し手の発話に明示的に含まれる情報より多くの情報を推論により得ることができる。この性質を利用すれば、より簡潔な発話でより多くの情報を伝えることができる。

3.2 確率的制約の記述例

[1] 信念スペース

対話者を含む個々の登場人物はそれぞれ異なる信念を持ち得るので、この違いを表現するために登場人物ごとに1つの信念スペース [Ballim 91] を割り当てる。ある登場人物 a がある命題 p を信じているとき、p が a の信念スペースに属する、または信念スペースで成り立つと言う。各登場人物は他の登場人物の信念に関する信念を持っているので、信念スペースは入れ子構造を形成する。システム (専門家 s) の信念スペースを [s]、システムが信じるユーザ (料理人 c) の信念スペースを [c,s] のように書く。以下では、ある要素式 $e(t_1, \dots, t_n)$ で表される命題がある信念スペース Bel で成り立つとき、これを $e(t_1, \dots, t_n, Bel)$ と書く。

異なる層の信念スペースに属する命題間に成り立つ

制約に信念の浸透 (belief percolation) [Ballim 91] の制約がある。信念の浸透の制約は次のように記述できる。

$$\begin{aligned} \text{mk}(A, X, [A, B | \text{Bel}]) &\leftarrow_{p_{13}} \\ \text{mk}(A, X, [B | \text{Bel}]) & \end{aligned} \quad (13)$$

$$\begin{aligned} \text{srv}(A, G, \text{wine}(X), [A, B | \text{Bel}]) &\leftarrow_{p_{14}} \\ \text{srv}(A, G, \text{wine}(X), [B | \text{Bel}]) & \end{aligned} \quad (14)$$

たとえば、式 (13) で $A = c$, $B = s$, $\text{Bel} = []$ とすると、式 (13) は

c が X を作るプランを立てていると s が信じているならば、c 自身も自分が X を作るプランを立てていることを信じていると s は信じている

という制約が確率 p_{13} で成り立つことを表す。

[2] タスクプランのライブラリ

ユーザ (c) が遂行するタスクプランのうちプリミティブなものは排他宣言として記述できる。

$$\begin{aligned} \text{disjoint}([\text{mk}(c, \text{spa}, [s]) : p_{15.1}, \\ \text{mk}(c, \text{fet}, [s]) : p_{15.2}, \\ \text{mk}(c, \text{chi}, [s]) : p_{15.3}, \\ \neg \text{mk}(c, \text{main}, [s]) : p_{15.4}]) & \quad (15) \\ \text{disjoint}([\text{mk}(c, \text{mari}, [s]) : p_{16.1}, \\ \text{mk}(c, \text{pesto_sause}, [s]) : p_{16.2}, \\ \neg \text{mk}(c, \text{sause}, [s]) : p_{16.3}]) & \quad (16) \end{aligned}$$

式 (15) は、メインとなる食材を料理する3種類のプランに関する記述で、スパゲティーマリナラ・フェトチーニ・チキンを同じプランのサブプランとして同時に料理することはないという制約である。 $p_{15.1}$ は、証拠が何もない状況で c が「スパゲティーマリナラをゆでる」プランを立てる ($\text{mk}(c, \text{mari}, [s])$) 確率である。 $\neg \text{mk}(c, \text{main}, [s])$ は、排他宣言内の確率の和を1にするために便宜的に加えた命題で、c が上の3種類のいずれのプランも立てていないことを表す。

プランとサブプランの関係は次のように記述できる。

$$\begin{aligned} \text{srv}(A, G, \text{spaMari}, [s]) &\leftarrow_{p_{17}} \\ \text{mk}(A, \text{spa}, [s]), \text{mk}(A, \text{mari}, [s]) & \end{aligned} \quad (17)$$

式 (17) は、A が「スパゲティーマリナラをゆでる ($\text{mk}(A, \text{spa})$)」プランと「マリナラソースを作る ($\text{mk}(A, \text{mari})$)」プランを同時に立てているならば、2つのプランは確率 p_{17} で「スパゲティーマリナラを作る人 G に出す ($\text{srv}(A, G, \text{spaMari})$)」という同一プランのサブプラン

であるという制約である。

プランの効果は次のように記述できる。

$$\begin{aligned} \text{happy}(G, \text{Bel}) &\leftarrow_{p_{18}} \\ &\text{srv}(A, G, \text{dish}(\text{meat}), \text{Bel}), \\ &\text{srv}(A, G, \text{wine}(\text{Wine}), \text{Bel}), \\ &\text{for}(\text{Wine}, \text{meat}, \text{Bel}), \\ &\text{vegi}(G, \text{false}, \text{Bel}). \end{aligned} \quad (18)$$

式(18)は、「AがGに肉料理とワイン wine(Wine)を出す (srv(A,G,wine(Wine)))」とき、「wine(Wine)が肉料理に合い (for(Wine,pasta))」かつ「Gが菜食主義者でない (vegi(G,false))」ことが信念スペース Bel で信じられていれば、Bel では確率 p_{18} で「Gがしあわせになれる (happy(G))」ことを表す。

〔3〕 概念階層

概念間の上位下位関係は次のように記述できる。

$$\begin{aligned} \text{srv}(A, G, \text{dish}(\text{pasta}), \text{Bel}) &\leftarrow_{p_{19}} \\ &\text{srv}(A, G, \text{spaMari}, \text{Bel}). \end{aligned} \quad (19)$$

式(19)は「スパゲティーマリナラを出すプラン」は「パスタ料理を出すプラン」の下位概念であるという記述である。

〔4〕 システムの信念とユーザモデル

システムの信念とシステムが信じるユーザの信念の違いは信念スペースをつかえば表現できる。たとえば、節が成り立つ確からしさに関する信念の違いは、信念スペースごとに節を記述し、各節に異なる確率を割当てることによって表現できる。

また、信念スペースごとに異なる証拠が存在するという状態も信念の違いを表現する。たとえば、専門家(s)は、

- パスタ料理には白ワインが合う
- 客 g が菜食主義者である

などの信念を持っている。これはそれぞれ証拠

$$\text{for}(\text{white}, \text{pasta}, [s]) \quad (20)$$

$$\text{vegi}(g, \text{true}, [s]) \quad (21)$$

として表現できる。一方、料理人(c)も同じ知識を持っているかどうかは(s)にはわからないので、その不確実性を特定のユーザcに依存した確率分布によって表現する。確率分布の初期状態の決め方については、ユーザをステレオタイプに分類する方法[Chin 88]が考えられる。

$$\begin{aligned} \text{disjoint}([\text{for}(\text{white}, \text{pasta}, [c, s]) : p_{22.1}, \\ \text{for}(\text{red}, \text{pasta}, [c, s]) : p_{22.2}]). \end{aligned} \quad (22)$$

$$\begin{aligned} \text{disjoint}([\text{vegi}(G, \text{true}, [c, s]) : p_{23.1}, \\ \text{vegi}(G, \text{false}, [c, s]) : p_{23.2}]). \end{aligned} \quad (23)$$

〔5〕 発話の観察

ユーザモデルはユーザからの発話を観察することによって更新される。ユーザモデルの更新はユーザモデルを規定する確率的制約の更新によって実現される。

たとえば、例(1)の料理人cの発話「いま、マリナラソースを作るところです」を観察すると、システムは「cは自分がマリナラソースを作るプランを立てていると信じている」という証拠を得たことになり、

$$\text{mk}(c, \text{mari}, [c, s]) \quad (24)$$

という証拠を追加する。

また、「赤ワインがいいですか」という料理人の質問からは、「cはパスタ料理または肉料理に合うワインの種類がわからない」という情報が得られる。この情報は式(22)で与えられた先験確率を変化させる(肉料理の場合も同様)。たとえば、

$$\begin{aligned} \text{disjoint}([\text{for}(\text{white}, \text{pasta}, [c, s]) : 0.5, \\ \text{for}(\text{red}, \text{pasta}, [c, s]) : 0.5]). \end{aligned} \quad (25)$$

一方、システムの発話をユーザが観察する場合もユーザモデルが変化する。たとえば、システムが「パスタ料理には赤ワインが合います」と発話すると、ユーザは「システムが for(white, pasta) を信じている」ことを知ることになる。これは信念スペース [s, c, s] に次の証拠を追加することに相当する。

$$\text{for}(\text{white}, \text{pasta}, [s, c, s]) \quad (26)$$

〔6〕 ユーザのゴール

例(1)では、ユーザが happy(g) というゴールを持っていることをシステムが対話の開始時点で知っていたものとしよう。ユーザは現在自分が立てているプランを遂行することによって自分のゴールが達成されると信じていると仮定できる。これは証拠

$$\text{happy}(g, [c, s]) \quad (27)$$

として表現される。

3.3 発話選択

いま s と c が協調的に対話をすすめると仮定しているので、s のゴールは c のゴールが達成できることを s 自身が信じていることができることである。すなわち、s の効用関数は $P(+\text{happy}(g, [s]) | e)$ で与えられる。

例(1)で c の発話「いまマリナラソースを作るところです。赤ワインがいいですか?」が観察された時点での信念モデルは図3のような Bayesian ネットワー

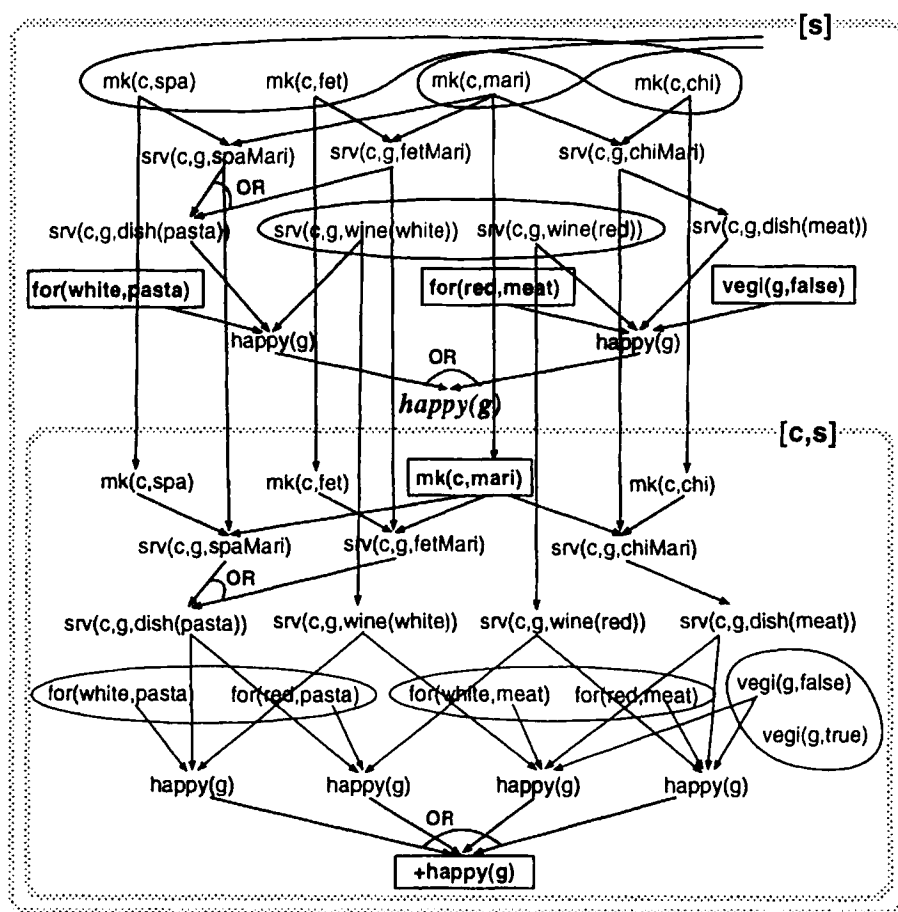


図3 例(1)の信念モデル

クで表現される。ただし、2.5節で述べたように排他宣言のインスタンスに対応する確率変数は Boolean 型でないため、図ではそれらを特別なノードとして扱い、実曲線によって表現している。すなわち、実曲線で囲まれたノードの集合は、それぞれある排他宣言のインスタンスに対応する1つの確率変数を表し、個々のノードはそれぞれ確率変数のとりうる値を表すものとする。たとえば、図中、信念スペース [s] に属する $mk(c,spa)$, $mk(c,fet)$, $mk(c,chi)$ を囲む曲線は排他宣言 (15) に対応する確率変数を表している。曲線で囲まれていない一般のノードは、それぞれ単独で Boolean 型確率変数を表す。

図3の状態では、 $P(vegi(g,true,[c,s]))$ が高いことから、 $P(srv(c,g,dish(meat),[c,s])|e)$ が抑制され、 $P(srv(c,g,dish(pasta),[c,s])|e)$ が高くなる。ところが、 c の発話から $P(for(white,pasta,[c,s]))$ が十分に高くないという証拠が得られ、効用関数 $P(+happy(g,[s])|e)$ の値も十分に高くない。ここで効用関数の値を高くするには、情報伝達行為 $inform(for(white,pasta))$ (「パスタ料理には白ワインが合います」) を実行するのが

もっとも効果的である。先に述べたように、 c のプランにあいまい性があるにもかかわらず、 c にとって有益な発話の選択ができる点に注意したい。

この例では、 $inform(srv(c,g,wine(white)))$ の優先度も同様に高くなる。これは s が c にプラン $srv(c,g,wine(white))$ を実行するよう要求すること (たとえば「白ワインを出しなさい」) に相当する。 $inform(for(white,pasta))$ のように事実関係の情報を伝達する行為と $inform(srv(c,g,wine(white)))$ のように行為を要求する行為のどちらを優先すべきかは、対話者間の対人的関係や対話の話題などに依存すると考えられる。この問題についての考察は今後の課題である。

次に例(2)の場合を考えよう。この例では、「パスタ料理に白ワインが合い、肉料理に赤ワインが合う」ことを c が信じていると s が信じていると仮定する。これは

$$\begin{aligned} & disjoint([for(white,pasta,[c,s]) : 0.95, \\ & \quad for(red,pasta,[c,s]) : 0.05]), \\ & disjoint([for(white,meat,[c,s]) : 0.05, \\ & \quad for(red,meat,[c,s]) : 0.95]). \end{aligned}$$

という制約として表現できる。一方、 g が肉食主義者であることを c は信じていないと s が信じているという制約は

$$\begin{aligned} & disjoint([vegi(g, true, [c, s]) : 0.001, \\ & \quad vegi(g, false, [c, s]) : 0.999]) \end{aligned}$$

のように表現できる。

現在の信念モデルでは、 $vegi(g, true, [c, s])$ の確率が低いので、 $srv(c, g, dish(meat), [c, s])$ の確率はとくに抑制されず、相対的に $srv(c, g, dish(pasta), [c, s])$ の確率は低くなる。信念スペース s では $vegi(g, true, [s])$ が証拠として与えられているので、 $srv(c, g, dish(meat))$ のプランは成功せず、 $P(+happy(g, [s])|e)$ は高くない。ここで、情報伝達行為 $inform(vegi(g, true))$ (「客はベジタリアンだよ」) をシステムが実行したとすると、証拠 $+vegi(g, true, [c, s])$ が信念モデルに追加される。これによって、 $srv(c, g, chiMari)$ のプランが成功しないことに c が気づくと予測されるため、このプランの確率が下がる。この予測は相対的に他の2つのプランの確率を上げるため、ゴール関数の評価値も上がる。したがって、この行為の優先度は高いと評価される。このように、発話の効用は、発話が相手に観察されたときに相手が行いそうな推論の先読みを含めて評価される。

4. む す び

本論文では、協調的な発話を意思決定理論的に選択する枠組について述べた。発話選択に手続き的知識をつかわない点、ユーザモデルの推定と発話選択を統合的に行う点、ユーザが行う推論を予測し発話選択に利用する点、がこの枠組の特徴である。確率モデルの表現をさらに工夫すれば、相手プランの誤りの指摘、自主的な情報伝達、問い合わせなどを状況依存的に選択するメカニズムの説明が可能になると期待できる。

同様の特徴を持つアプローチに Nagao らの力学的制約にもとづく行為選択の研究 [Nagao 93] がある。Nagao らの枠組では、システムやユーザの推論に関する制約を力学的制約として宣言的に記述し、その制約のもとでシステムとユーザのゴールの達成の度合を表す効用関数を大きく増加させることが期待できる伝達内容を優先する。我々の枠組と同様に発話プランニングのための手続き的規則を必要とせず、ユーザの多様な推論のシミュレートと状況追従性の高いプランニングの実現が期待できる。しかしながら、この枠組では力学的制約を構成するパラメタの意味論が明らかでない。これに対し、我々の枠組では確率論的基礎にもとづい

て制約を記述することができる。

Bayesian ネットワークの評価は一般に NP hard であることが知られており [Pearl 88]、高速な近似アルゴリズムがいくつか提案されている。代表的なものに統計シミュレーションをつかって近似する方法 [Shachter 90] があるが、確率 Horn 論理で規定されるような論理ゲートを多用するモデルの評価には向かないことが証明されている [Shavez 90]。これに対し我々は、Poole が提案している可能世界の足し合わせによる近似法 [Poole 93b] を改良し、論理ゲートを含む信念モデルを高速に評価するアルゴリズムを開発し実装している [乾 95]。

発話の優先度の計量は発話内容が複数の命題の値割当てを含む場合にも適用できるが、その場合計算量が組み合わせ的に増大するため、何らかの工夫が必要である。ただし、対話では、相手の信念に対する不確実性が高いため、完全なプランを立ててから実行するよりも、局所的な優先度による部分的なプランの立案と実行をインタリーブさせる方が望ましい場合が多い。このことは、発話の優先度計算において複数の命題の組み合わせを考える必要が必ずしもないことを示唆している。伝達する命題の数を1つに限定して優先度計算する場合にどの程度妥当な選択ができるかをさまざまな対象領域について実験的に調べる必要があるだろう。これについては今後の課題である。

信念モデルの動的構築において、条件つき独立性だけでモデルの大きさが十分に抑えられない場合は、よりゴールとの関連性が高い部分だけを具現化する必要がある。ゴールとの関連性を見積るにはモデルをある程度評価する必要があるため、モデルの構築と評価とのインタリーブによって漸進的に洗練されたモデルを得る枠組が最近注目されている [Goldman 92, Poh 93, Provan 94]。また、コストの安い確率的マーカパッシングで確率の低い仮説を枝刈りする方法も提案されている [Carroll 91]。これらの研究成果をどのように我々の枠組に応用するかは今後の課題である。

謝 辞

本研究を進めるに当たり貴重なコメントをいただきました電子技術総合研究所の橋田浩一氏、ソニーコンピュータサイエンス研究所の長尾確氏、日立製作所基礎研究所の岩山真氏に感謝いたします。

◇ 参 考 文 献 ◇

- [Akiha 94] T. Akiha and H. Tanaka. A Bayesian approach for user modelling in dialog systems. In *Proceedings of the 14th International Conference on Computational*

- Linguistics*. COLING '94, pp.1212-1218, 1994.
- [Ballim 91] A. Ballim and Y. Wilks. *Artificial Believers: The Ascription of Belief*. Lawrence Erlbaum Associates, 1991.
- [Breese 92] J.S. Breese. Construction of belief and decision networks. *Computational Intelligence*, Vol.8, No.4, pp.624-647, 1992.
- [Breese 94] J.S. Breese, R.P. Goldman, and M.P. Wellman. Introduction to the special section on knowledge-based construction of probabilistic and decision models. *IEEE Transactions on Systems, Man, and Cybernetics*, Vol.24, No.11, 1994.
- [Carroll 91] G. Carroll and E. Charniak. A probabilistic analysis of marker-passing techniques for plan-recognition. In *Proceedings of the 7th International Conference on Uncertainty in Artificial Intelligence*, pp.69-76, 1991.
- [Chin 88] D.N. Chin. Exploiting user expertise in answer expression. In *Proceedings of the 7th National Conference on Artificial Intelligence*, Vol.2, pp.756-760. AAAI '88, 1988.
- [Cohen 90] P.R. Cohen, J. Morgan, and M.E. Pollack, editors. *Intentions in Communication*. The MIT Press, 1990.
- [Goldman 92] R.P. Goldman and J.S. Breese. Integrating model construction and evaluation. In *Proceedings of the 8th International Conference on Uncertainty in Artificial Intelligence*, pp.104-111, 1992.
- [乾 95] 乾健太郎. 自然言語生成における相互依存的制約の扱いに関する研究. 博士論文, 東京工業大学, 1995.
- [Litman 90] D.J. Litman and J.F. Allen. Discourse processing and commonsense plans. In Cohen, et al. [Cohen 90], chapter 17, pp.365-388.
- [Nagao 93] K. Nagao, K. Hasida, and T. Miyata. Emergent planning: A computational architecture for situated planning. In *Proceedings of the 5th European Workshop on Modelling Autonomous Agents in a Multi-Agent World*. MAAMAW '93, 1993.
- [Pearl 88] J. Pearl. *Probabilistic Reasoning in Intelligent Systems*. Morgan Kaufmann, 1988.
- [Poh 93] K.L. Poh and E.J. Horvitz. Reasoning about the value of decision-model refinement: Methods and applications. In *Proceedings of the 10th International Conference on Uncertainty in Artificial Intelligence*, pp.174-182, 1993.
- [Pollack 90] M.E. Pollack. Plans as complex mental attitudes. In Cohen, et al. [Cohen 90], chapter 5, pp.77-105.
- [Poole 93a] D. Poole. Probabilistic Horn abduction and Bayesian networks. *Artificial Intelligence*, Vol.64, No.1, pp.81-129, 11 1993.
- [Poole 93b] D. Poole. Average-case analysis of a search algorithm for estimating prior and posterior probabilities in Bayesian networks with extreme probabilities. In *Proceedings of the 13th International Joint Conference on Artificial Intelligence*, Vol.1, pp.606-612, 1993.
- [Provan 94] G.M. Provan. Tradeoffs in knowledge-based construction of probabilistic models. *IEEE Transactions on Systems, Man, and Cybernetics*, Vol.24, No.11, pp.1580-1592, 1994.
- [Shachter 90] R.D. Shachter. An ordered examination of influence diagrams. *Networks*, Vol.20, pp.535-562, 1990.
- [Shavez 90] M.R. Shavez and G.F. Cooper. A randomized approximation algorithm for probabilistic inference on Bayesian belief networks. *Networks*, Vol.20, pp.661-685, 1990.
- [van Beek 91] P. van Beek and R. Cohen. Resolving plan

ambiguity for cooperative response generation. In *Proceedings of the 12th International Joint Conference on Artificial Intelligence*, Vol.2, pp.938-944. IJCAI '91, 8 1991.

[山田 94] 山田耕一, 溝口理一郎, 原田直樹. 質問応答システムにおけるユーザ発話モデルと協調的応答の生成. 情報処理学会論文誌, Vol.35, No.11, pp.2265-2275, 1994.

[担当編集委員・査読者: 山本秀樹]

著者紹介

乾 健太郎(正会員)



1990年東京工業大学工学部情報工学科卒業. 1995年同大学大学院理工学研究科博士課程修了. 同年同大学助手, 現在にいたる. 博士(工学). 自然言語処理, 言語学習支援などの研究に従事. 言語処理学会, 機能言語学会各会員.

<inui@cs.titech.ac.jp>

徳永 健伸(正会員)



1983年東京工業大学工学部情報工学科卒業. 1985年同大学院理工学研究科修士課程修了. 同年(株)三菱総合研究所入社. 1986年東京工業大学大学院博士課程入学. 現在, 同大学大学院情報理工学研究科助教授. 博士(工学). 自然言語処理, 情報検索などの研究に従事. 情報処理学会, 計量国語学会, 言語処理学会, Association for Computational Linguistics, ACM SIGIR 各会員.

<take@cs.titech.ac.jp>

田中 穂積(正会員)



1964年東京工業大学工学部情報工学科卒業. 1966年同大学院理工学研究科修士課程修了. 同年, 電気試験所(現電子総合研究所)入所. 1980年東京工業大学助教授. 1983年同大学教授. 現在, 同大学大学院理工学研究科計算工学専攻教授. 工学博士. 人工知能, 自然言語処理に関する研究に従事. 情報処理学会, 電子情報通信学会, 日本認知科学会, 言語処理学会, 計量国語学会各会員.

<tanaka@cs.titech.ac.jp>