

確率的制約にもとづく発話プランニング

乾健太郎, 徳永健伸, 田中穂積

東京工業大学 工学部

{inui,take,tanaka}@cs.titech.ac.jp

協調的で効率的な発話のプランニングを実現する上で起こる2つの問題, すなわち, 相手のプランがあいまいなときにどう対応するか, 相手の推論の予測をどのようにプランニングに埋め込むかに焦点をあて, 状況に依存しないヒューリスティクスをつかって発話を選択する枠組について論じる. ここでつかうヒューリスティクスとは, 最大の効用をもつ発話を選択するというものである. 我々の枠組では, すべての知識を確率的制約として記述し, その制約の上で発話の効用を計算する. 本稿では確率的制約の表現と発話の効用の計算についてのべる.

Utterance Planning based on Probabilistic Constraints

INUI Kentaro, TOKUNAGA Takenobu and TANAKA Hozumi

Department of Computer Science, Tokyo Institute of Technology

(2-12-1 Ōokayama Meguro Tokyo 152 Japan)

This paper presents a framework for utterance planning where content planning is guided by a situation-independent heuristics. Here we address two issues on planning cooperative and efficient utterance: how the system could react when ambiguity of the user's plan can not be completely resolved, and how the system could take into consideration user's expected reaction to utterance of itself when planning. In our framework, every piece of knowledge is represented as a probabilistic constraint, on which the utility of each utterance plan is evaluated. Contents of utterance are chosen according to their utility. This paper focuses on the representation of probabilistic constraints and evaluation of utterance's utility.

1 はじめに

目的指向の対話を協調的にすすめるためには、まず相手の発話と文脈から相手のプランや信念状態を的確に推論し、推論した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。これまでに相手の発話意図やプランを認識するタスク (e.g. [8, 14]), また知識を推測するタスク (e.g. [1]) について多くの研究がなされているが、これらのタスクの主要な目的は次の発話 (一般には行為) をより協調的なものにするることである。相手のプランがわかれば、明示的には聞かれていないが重要な情報を提供したり [8], 相手のプランがまちがっていること教えたり [14] することができる。また、相手の知識が推測できれば、それと関連させながら新しい情報を説明することができる [1, 9]。本稿では、このような背景のもとで生じる以下の2つの問題をとりあげる。

1. 相手のプランや信念状態が不確実にはかわからないときにどのように行為すべきか。
2. 自分の発話によって引き起こされる相手の推論 (の予測) をどのようにプランニング中に考慮するか。

言語理解の多くの研究が示すように、相手のプランや信念状態を一意に特定するのは一般に容易なことではない。このため、これまでに開発された対話システムの多くは、相手のプランや質問の内容を一意に特定するまでユーザにとって無益な質問をくりかえすという問題をかかえていた。相手のプランや要求があいまいな状況で発話をプランニングする機構が欠けていたためである。これに対し、van Beek ら [18] と Nagao [12] らはプランが一意に特定できなくても相手にとって有益な内容を発話できる場合があることを指摘している。

ここで1つの例題を考えてみよう。以下の例は van Beek, Nagao らがともにつけた例をもとにしている。

例 (1) 料理人 *c* が菜食主義の客 *g* にだす料理をつくっているとき専門家 *s* に「いまマリナラソースをつくるところですが、ワインは赤がいいですか?」とたずねたとする。*s* は、*c* がマリナラソースをつくっていることから、料理が

- スパゲティーマリナラ
- フェットチーニマリナラ
- 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だろうと (*s* が) 思っているとしよう。すると *s* は、料理が肉料理でなく、スパゲティーマリナラかフェットチーニマリナラのいずれかだろうと推測できる。いずれの場合も白ワインが適しているので、*s* は白ワインをだすように *c* に要求する。

この例で *s* は、*c* のプランにあいまい性があるにもかかわらず、*c* にとって有益な答えをかえすことができる。

このような発話プランニングを実現するためには、プランライブラリをつかってシステムのゴールからプランを展開する従来の方法 (e.g. [10, 4]) では不十分である。ユーザのプランや信念状態が不確実な場合、どれとどれが不確実かについての状況の記述は組み合わせ的になり、状況ごとにプランニングの規則を書いていたのでは、規則の数が組み合わせ的に増えると予想されるからである。

つぎに上述の問題の2点目について考えるために、もう1つの例を見てみよう。

例 (2) 例 (1) と同様、*s* は *c* がマリナラソースをつくっていることを *c* の発話から知っている。ここで、*g* が菜食主義であることを *c* は知らないと *s* が思っているとしよう。このとき *c* の可能なプランは前述の3通りである。さらに、*c* が肉料理を好んでつくることを *s* が知っていたとすると、チキンマリナラの可能性が高いことになる。そこで *s* は「客は菜食主義者だよ」と *c* に教える。

ここで *s* はつぎのように推論したと考えられる。まず、*c* のプランは失敗する可能性が比較的高いことに気づく。また、菜食主義者は肉料理を食べないことを *c* も知っているのも、もし *c* が *g* が菜食主義であることを知っていれば肉料理は作らないだろうと推論する。したがって、「客は菜食主義者だ」教えれば、*c* は肉料理をつくらず、失敗を回避できる。

ここで重要なのは、*s* の発話に対する *c* の推論を予測した上で *s* がその発話を実行したことである。

このことは、発話プランニングというタスクのなかに、聞き手の推論の予測というタスクが埋め込まれていることを示唆する。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したいときに特に重要である。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こと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と考えられる。このことは、例(2)のように「客が肉食主義である」ことを教えるだけでcのプランの誤りが解消されるということからもわかる。

このようなプランニングの実現を考える場合も、例(1)での議論と同様、発話の戦略を規則として記述する方法は現実的でない。聞き手の推論は、話し手の発話の内容ごとに、また文脈ごとに非常に多様な方向にすすむことが予想されるが、各々の場合についてプランニングの規則を用意する方法ではその多様性をあつかいきれない。

これらの問題に対するアプローチとして、我々は確率的制約にもとづく発話プランニングの枠組の構築をすすめている。この枠組では、対話者の推論を引き起こす信念状態間の関係や信念状態の尤度をすべて確率的制約として宣言的に記述し、それらの制約のもとでの発話の確率的な効用を計算することによって発話の選択をおこなう。

発話プランニングは、モデル構築・モデル評価の2つの処理からなる(図1)。

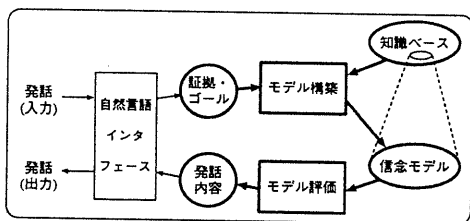


図1 全体の枠組

システムは、ユーザの発話から発話内容を抽出すると、まず知識ベースに照らして信念モデルを動的に構築(すでにある場合は更新)する。発話内容は後述のように信念モデルにおける証拠またはゴールとしてあつかわれる。知識ベースはさまざまな命題間の局所的な確率的制約の集合であると仮定する。現在我々は Bayes 法による確率の定式化 [13] をもとにしているの、局所的な確率的制約は局所的な

確率的依存関係として記述される。また、知識ベースは領域全体を包含する仮想的な Bayesian ネットワークとみなせる。Bayes の方法は、大域的な信念状態間の関係や信念状態の尤度を局所的な確率的制約の組み合わせとして表現し、計算することができるというよい性質をもっている。ただし、現実的な領域では知識ベースが現実的な速度で計算できる大きさをはるかに越えるため、ユーザから得られた発話やシステムのゴールに関係のある部分だけを切り出す処理が必要になる。切り出した部分は Bayesian ネットワークとして明示的に表現される。本稿ではこれを特に信念モデルとよぶ。このような処理は一般に動的モデル構築とよばれ、これまでにさまざまな研究がおこなわれてきた [6, 19]。信念モデルの動的構築については、現在 Carroll と Charniak が提案している確率的マーカパッシング [3] の拡張を検討している。

システムは、つぎに信念モデルを評価することによってシステムの発話の候補の効用をそれぞれ計算する。発話の効用とは発話がゴールの達成に貢献する度合である。システムは、効用のもっとも大きい候補を選択するというヒューリスティクスにしたがって発話内容を決定する。我々の方法では、相手のプランや信念状態をあいまいなまに保ちながら効用を計算するので、ユーザのプランがあいまいな場合のプランニングとあいまいでない場合のプランニングを統一的にあつかうことができる。また Bayesian ネットワークをつかうことによってさまざまな方向に進む推論の複雑な組み合わせを容易にあつかうことができる。

本稿では主として、信念モデルの評価による発話選択について述べる。

2 確率的制約

2.1 信念モデル

信念モデル M は $\langle V, C, e \rangle$ の組で定義される非循環有効グラフである。 V は確率変数 V_i の集合である ($i = 1, \dots, n$)。各変数はそれぞれネットワークの1つのノードに対応する。本稿では簡単のためすべての変数の値域を $\{true, false\}$ とする¹。以

¹我々はすべての確率的計算を Bayesian ネットワークの定式化にしたがっておこなうので、任意の値域をゆるすように一

下, 変数 V に対する値割当てを v , $V = \text{true}$ を $+V$, $V = \text{false}$ を $-V$ と書くことがある. C は確率的制約 C_V の集合である. C_V は変数 V と V の親の変数²の集合 X_V との局所的依存関係を規定する. これは V と X_V のすべての値割当ての組み合わせ (v, x_V) についての条件つき確率 $P(v|x_V)$ であたえられる. 変数 V が根ノードに対応する場合, C_V は $P(v)$ であたえられる. 最後に e は観察された証拠 $e_j \in \{+E_j, -E_j\}$ の集合である.

V が n 個の親をもつとき C_V のパラメタの数は 2^n 個になり, すべてに適当な数値を割当てるのは一般に容易でない. そこで, 確率的制約の単位をもう少し小さくし, さらに条件つき確率の与え方の自由度をわずかにきびしくする. 例として, $+Q$ なら確率 p_1 で $+S$ が起こり, $+R_1$ かつ $+R_2$ なら確率 p_2 で $+S$ が起こるという知識が別々に得られたとする. 閉世界仮説を仮定し, S の親が上の 2 つで数えつくされているとすると, この知識を次の 3 つの確率的制約として表現する.

- (1) $P(+S_1 | +Q) = p_1$
 $P(+S_1 | \text{otherwise}) = 0$
- (2) $P(+S_2 | +R_1 \wedge +R_2) = p_2$
 $P(+S_2 | \text{otherwise}) = 0$
- (3) $P(+S | +S_1 \vee +S_2) = 1$
 $P(+S | \text{otherwise}) = 0$

ここでは, S とその親との依存関係を 2 つの制約 (1), (2) をつかって表現している. 各制約のパラメタが 1 つであることに注意したい. (3) は S に対する 2 通りの説明 (1), (2) の noisy-or [13] を自然に表現している. このため, Q と R_1, R_2 が独立でない場合も適切にあつかうことができる. また, 上の 2 通りの説明には排他性が要求されない.

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. これにより, 変数の集合 V について, その値割当て v の結合確率は, C があたえる局所的な条件つき確率の積として計算することができる.

$$\forall v P(v) = \prod_{v_i \in v} P(v_i | x_{V_i})$$

したがって, Bayesian ネットワークをもちいると,

般化することは容易である.

² 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, 後者を前者の子とよぶ.

局所的な状態間の制約の組み合わせによって大域的な制約を表現することができる.

変数の値割当て v_i について, その証拠のもとでの条件つき確率 $P(v_i | e)$ を v_i の後驗確率という. 後驗確率の計算は一般に NP hard であることが知られており [13], 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 [16] があるが, 上でのべた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 [17]. そこで我々は, Poole が提案している可能世界の足し合わせによる近似法 [15] 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た. 現在その性能について実験中であり, これについては別の機会に報告する.

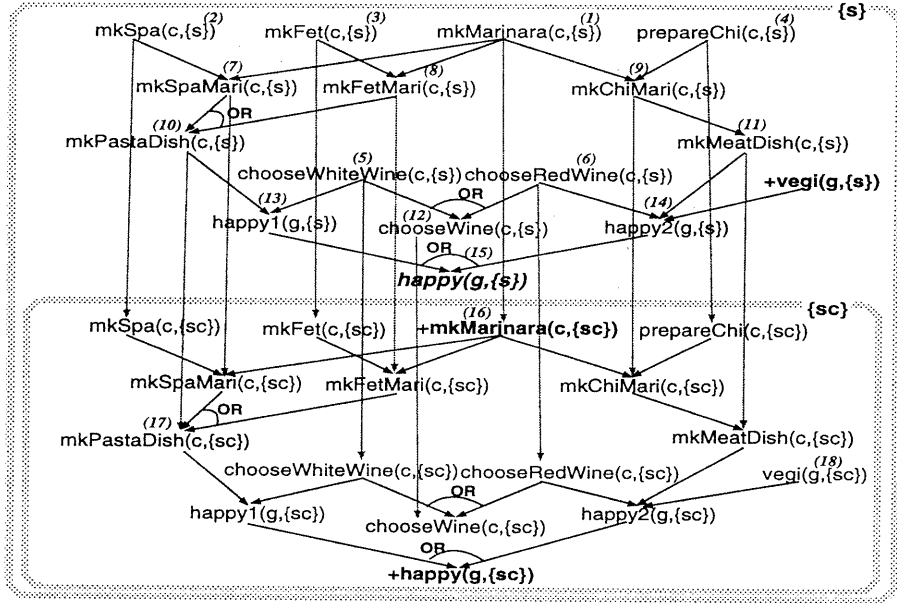
2.2 信念状態

1 節でのべたように, システムはユーザの発話が観察されると知識ベースから顕現性の高い部分だけを切り出し, 信念モデルを構築 (または更新) する. このときシステムがおこなう推論は非決定的なので, 信念モデルは一般にいくつかの可能な推論 (可能世界) を含むことになる. たとえば, プランの認識が重要な状況では, 信念モデルは複数のプランの候補を含んでいるかもしれない. このように, モデル構築はあたえられた状況のもとで比較的尤度の高い可能世界をいくつか選択する処理である.

これに対し, 一度信念モデルが構築されると, たとえそこにあいまいな可能世界が複数存在しても, その解消は明示的にはおこなわれない. 個々の可能世界はそれぞれ異なる尤度をもつが, その差は可能世界の確率分布として信念モデルのなかに潜在的に存在するだけである. すなわち, 個々の確率は明示的に計算されるのではなく, 信念モデルの確率的制約によって宣言的にあたえられる. このような, ある時点での信念モデルの状態を以下では信念状態とよぶ. システムが計算するのは可能世界の確率分布ではなく, それをあたえる信念状態のもとでの発話の効用であることに注意してほしい.

ここで, 1 節で紹介した例題に関する信念モデルを図 2 にしめそう.

まず, s と c の信念の違いを表現するために信念スペース [2] を用意する. s の信念スペースを $\{s\}$,



- (1) $P(+mkMarinara(c, \{s\})) = .00001$
- (2) $P(+mkSpa(c, \{s\})) = .0001$
- (3) $P(+mkFet(c, \{s\})) = .0001$
- (4) $P(+mkChi(c, \{s\})) = .0003$
- (5) $P(+chooseWhiteWine(c, \{s\})) = .0001$
- (6) $P(+chooseRedWine(c, \{s\})) = .0001$
- (7) $P(+mkSpaMari(c, \{s\}) | +mkMarinara(c, \{s\}) \wedge +mkSpa(c, \{s\})) = .001$
- (8) $P(+mkFetMari(c, \{s\}) | +mkMarinara(c, \{s\}) \wedge +mkFet(c, \{s\})) = .001$
- (9) $P(+mkChiMari(c, \{s\}) | +mkMarinara(c, \{s\}) \wedge +mkChi(c, \{s\})) = .001$
- (10) $P(+mkPastaDish(c, \{s\}) | +mkSpaMari(c, \{s\}) \vee +mkFetMari(c, \{s\})) = 1$
- (11) $P(+mkMeatDish(c, \{s\}) | +mkChiMari(c, \{s\})) = 1$
- (12) $P(+chooseWine(c, \{s\}) | +chooseWhiteWine(c, \{s\}) \vee +chooseRedWine(c, \{s\})) = 1$
- (13) $P(+happy1(g, \{s\}) | +mkPastaDish(c, \{s\}) \wedge +chooseWhiteWine(c, \{s\})) = .9$
- (14) $P(+happy2(g, \{s\}) | +mkMeatDish(c, \{s\}) \wedge +chooseRedWine(c, \{s\}) \wedge -vegi(g, \{s\})) = .9$
- (15) $P(+happy(g, \{s\}) | +happy1(g, \{s\}) \vee +happy2(g, \{s\})) = 1$
- (16) $P(+mkMarinara(c, \{sc\}) | +mkMarinara(c, \{s\})) = 1$
- (17) $P(+mkPastaDish(c, \{sc\}) | [+mkSpaMari(c, \{s\}) \vee +mkFetMari(c, \{s\})] \wedge +mkPastaDish(c, \{s\})) = 1$
- (18) $P(+vegi(g, \{sc\})) = .95$

図 2 信念モデル

s が信じる c の信念スペース $\{sc\}$ のように書くことにすると、信念スペースは図のような入れ子構造をつくる。

(1)～(6) は証拠が何もない状態で c がそれぞれのプランをおこなう確率である³。確率の意味論に

³以下、2.1節でのべた $P(V|otherwise) = 0$ の記述を一貫して省略する。

については Charniak らの議論を参考にした [5]。 (7)～(9) はプランの部分全体関係に関する知識である。たとえば、(7) はスパゲティーマリナラをつくるプランは2つのサブプラン(マリナラソースを準備するプランとスパゲティを準備するプラン)からなるという知識である。(10)～(12) はプランの概念の上下位関係をあわわしている。たとえば、(11)

はチキンマリナラをつくるプランの上位概念が肉料理をつくるプランであることをあらわしている。(13), (14) はそれぞれプランとその効果の因果関係を記述している。(13) はパスタ料理のときに白ワインをだすと客 g がよろこぶという知識であり, (14) は肉料理のときは赤ワインがいいが, g が菜食主義であってはいけないという知識である。(15) は (13) と (14) の noisy-or である。信念スペース $\{sc\}$ でもだいたい上と同じような制約がはたらくと仮定する。

一方, $\{s\}$ と $\{sc\}$ の間には信念の浸透 (belief percolation) [2] に関する制約がある。たとえば (16) は, c がマリナラソースをつくっていると s が信じているなら, そのことを c 自身も信じているはずだと s は考えるという知識である。ここでは簡単のため, c の行為および知識についての信念のみ $\{s\}$ と $\{sc\}$ の間で浸透がおこるとする。たとえば, $vegi(g)$ については浸透がおこらない。 $vegi(g)$ は g についての命題であり, s がそれを信じているからといって, c も同じことを信じているとはかぎらないからである。ここでは, s 自身は $vegi(g)$ を過去の観察によって知っているが, c が $vegi(g)$ を信じているかどうかは確率的にしかわからないと仮定する (18)。

上の確率的制約のもとで c の発話が観察されると, c がマリナラソースをつくっており, 現在のプランで g を楽しませることができると信じていることがわかる。したがって, $+mkMarinara(b, \{sc\})$ と $+happy(g, \{sc\})$ が証拠としてあたえられる。この新しい証拠をふくむ確率的制約の状態が s の現在の信念状態である。各可能世界は信念状態中に潜在的に存在する。たとえば上の例では, c がスパゲティーをつくって g を楽しませる可能世界と, フェトチーニをつくって g を楽しませる可能世界の確率は同程度である。しかしながら, s と同様 c も $vegi(g)$ を信じている確率が高いので, c がチキンマリナラをつくって g を楽しませることに失敗するという可能世界の確率は相対的に低い。

3 発話プランニング

3.1 ゴール関数

対話者 s のゴール関数 G を信念モデル $M = (V, C, e)$ 中のある変数 $V_g \in V$ に対するある値割当

で v_g の後驗確率 $G = P(v_g|e)$ としよう。 $v_g = +V_g$ のとき, V_g は s にとって成り立ってほしいと思っている命題である。一方, $v_g = -V_g$ のとき, V_g は s にとって成り立ってほしくない命題である。 G は, M において s のゴール v_g が成り立つ確率であり, 現在の確率分布のもとでのゴールの達成度とみなすこともできる。

s は複数のゴール v_{g1}, v_{g2} を同時に達成したいと思うかもしれない。その場合, 一般にゴール間には重要度の相対的な差があると予想される。たとえば, Nagao らの対話プランニングでは, 自分のゴール関数と対話相手のゴール関数の重みつき和を各対話者のゴール関数としている [11]。そこで, 個々のゴールの相対的な重要度 (重み) を w_1, w_2 ($w_1 + w_2 \leq 1$) とし, 仮想的な変数 G とその値割当て g を考える。このとき, G は

$$G = P(g|e)$$

ただし, v^- をある値割当て v の否定とすると,

$$P(g|v_{g1}, v_{g2}) = w_1 + w_2$$

$$P(g|v_{g1}, v_{g2}^-) = w_1$$

$$P(g|v_{g1}^-, v_{g2}) = w_2$$

$$P(g|v_{g1}^-, v_{g2}^-) = 0$$

3.2 発話の選択

2節でのべたように, 信念モデルでは対話者の発話内容 (命題) を証拠としてあつかう。発話の観察は信念モデルに新たな証拠をあたえることに対応する。新たにあたえられた証拠は信念状態を変化させ, 可能世界の確率分布を変化させる。この変化を発話の効果と考える。そこで発話の効用をつぎのように定義する。

発話の効用 対話者のゴール関数を $G = P(g|e)$, ある発話 j の内容をあらわす命題を Q_j とするとき, 発話 j の効用 U_j は,

$$U_j = \frac{P(g|+Q_j, e)}{P(g|e)}.$$

すなわち, 発話の効用は対応する命題を証拠としたときにゴール関数の評価値が大きくなる度合をはかる尺度である。したがって, いつも効用の高い発話を選択するという戦略は, 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を実現するためのよいヒューリスティクスになっている。

可能な発話の候補の集合を j とすると、最適な発話 j^* はつぎのように計算できる。

$$\begin{aligned} j^* &= \arg \max_{j \in J} U_j \\ &= \arg \max_{j \in J} \frac{P(g|Q_j, e)}{P(g|e)} \\ &= \arg \max_{j \in J} \frac{P(+Q_j|g, e)}{P(+Q_j|e)} \end{aligned}$$

この式から、証拠が e の場合と $+Q \wedge e$ の場合の各変数の後驗確率を計算すれば j^* がもとまることがわかる。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度おこなえば、それぞれの行為の効用が計算でき、 j^* を見つけることができる。発話の候補ごとに別々に効用を計算する必要がない点が重要である。

3.3 例題

ここで再び1節の例にもどろう(図2)。いま s がシステム、 c がユーザであるとする。 s は c と協調的に対話をすすめるので、現在の s のゴールは c のプランを成功させることである。すなわち、 s のゴール関数は

$$P(+happy(g, \{s\})|e)$$

であたえられる。

例(1)では $P(+vegi(g, \{sc\}))$ が高いため(18), $+mkMarinara(b, \{sc\})$ と $+happy(g, \{sc\})$ が観察された時点では、 $+mkChiMari(b, \{s\})$ にくらべ $+mkSpaMari(b, \{s\})$, $+mkFetMari(b, \{s\})$ の確率が十分に高い。また、パスタ料理に白ワインがあることを c も知っていると s は信じているので、 $+chooseWhiteWine(b, \{s\})$ の確率も十分に高い。したがって、 $+happy(a, \{s\})$ の確率はすでに高く、 s は発話する動機をもたない。つぎに質問「ワインは赤がいいですか」が観察されると、パスタ料理には白があるという知識を c がもっていないことがわかり、信念スペース $\{sc\}$ の状態が図3のように変化する⁴。更新後の信念モデルでは $+chooseWhiteWine(b, \{sc\})$ の確率が下がり、結果としてゴール関数の評価値も下がる。したがって、 $+chooseWhiteWine(b, \{sc\})$ の確率が再び上がるよ

⁴この処理は今のところモデル構築(更新)の過程でおこなわれると仮定しており、どのように実現すべきかは今後の課題である。

うに発話すれば(たとえば、「白ワインに下さい」と発話する)、ゴール関数の評価値も上がることがわかる。実際、発話の効用を計算すると約3である。効用が3であるということは、その行為を実行すればゴール関数の評価値が3倍になることを意味する⁵。

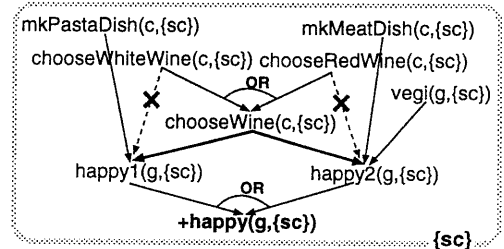


図3 信念モデルの更新

このような推論は状況によって非常に多様な方向に進むので、推論の方向をあらかじめ規定する方法では本質的にあつかいきれない。我々の枠組では、このような推論を1段ごとに明示的におこなうのではなく、効用の評価過程のなかで暗黙に考慮することができる。また、この性質によって、我々は1つの特定な可能世界に明示的にコミットすることなく発話の効用を計算することができる。例(1)では、 s にとって c のプランが依然あいまいであるにもかかわらず発話の効用を計算している。

つぎに例(2)の場合を考えてみよう。 c が $vegi(g)$ を信じていないと s が信じているという状況は $P(+vegi(g, \{sc\}))$ が十分小さいことに対応する。

$$(18') P(+vegi(g, \{sc\})) = .001$$

1節でのべたように、この信念状態のもとでは $+mkChiMari(b, \{s\})$ の確率が抑制されない。逆に、図2の制約(2), (3), (4)にしたがい、 $+mkChiMari(b)$ の確率は他のプランの確率より高くなる。したがって、 s の信念スペースのなかで c のプランが失敗する確率が高くなる。ここで s が「 g は菜食主義者だよ」と教えると、 $+vegi(g, \{sc\})$ という証拠が c の信念スペースに加わる。すると、チキンマリナラ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、そ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

⁵対話者が行為の実行にふみきるかどうかを決める要因として、効用の閾値を仮定することもできるかもしれない。

るため、ゴール関数の評価値も上がる。このことは $+vegi(g, \{sc\})$ の効用が高いことに反映される。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おこないそうな推論の先読みをふくめて評価される。

4 おわりに

本稿では、協調的で効率的な発話のプランニングを実現する上でおこる2つの問題、すなわち、相手のプランがあいまいなときにどう対応するか、相手の推論の予測をどのようにプランニングに埋め込むかに焦点をあて、確率的制約にもとづく状況に依存しないヒューリスティクスをつかって発話を選択する枠組について論じた。我々の枠組は以下のように多くの興味深い研究課題をふくんでいる。

- 発話行為と物理的行為の統一的なあつかい、行為のコスト、信念の浸透や時間的な概念のあつかい、無知を表現するための3値論理の導入の有効性など、種々の知識をどのように確率的制約として表現するか?
- 効率的なモデル構築アルゴリズムの洗練。Nagao, Hasida らが指摘しているように [11], 状況に依存しないモデル構築の制御方法の開発は本質的な問題である。
- 確率の意味論

参考文献

- [1] T. Akiha and H. Tanaka. A Bayesian approach for user modelling in dialog systems. In *Proceedings of the 15th International Conference on Computational Linguistics*, 1994.
- [2] A. Ballim and Y. Wilks. *Artificial Believers: The Ascription of Belief*. 1991.
- [3] G. Carroll and E. Charniak. A probabilistic analysis of marker-passing techniques for plan-recognition. In *Proceedings of the 7th Conference on Uncertainty in Artificial Intelligence*, pp. 69–76, 1991.
- [4] A. Cawsey. Generating interactive explanations. In *Proceedings of the 9th National Conference on Artificial Intelligence*, Vol. 1, pp. 86–91, 1991.
- [5] E. Charniak and R. Goldman. A semantics for probabilistic quantifier-free first-order languages, with particular application to story understanding. In *Proceedings of the 11th International Joint Conference on Artificial Intelligence*, pp. 1074–1079, 1989.
- [6] E. Charniak and R. P. Goldman. A Bayesian model of plan recognition. *Artificial Intelligence*, Vol. 64, No. 1, pp. 53–79, 11 1993.
- [7] P. R. Cohen, J. Morgan, and M. E. Pollack, editors. *Intentions in Communication*. The MIT Press, 1990.
- [8] D. J. Litman and J. F. Allen. Discourse processing and commonsense plans. In Cohen, et al. [7], chapter 17, pp. 365–388.
- [9] 榎本英治, 垣内隆志, 上原邦昭, 豊田順一. 対話型システムにおける文脈情報を利用した文章生成について. 情報処理学会自然言語処理研究会, Vol. NL59-8, 1986.
- [10] J. D. Moore. *A Reactive Approach to Explanation in Expert and Advice-Giving Systems*. 博士論文, UCLA, 1989.
- [11] 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.
- [12] K. Nagao and A. Takeuchi. Social interaction: Multimodal conversation with social agents. In *National Conference on Artificial Intelligence*, 1994.
- [13] J. Pearl. *Probabilistic Reasoning in Intelligent Systems*. Morgan Kaufmann, 1988.
- [14] M. E. Pollack. Plans as complex mental attitudes. In Cohen, et al. [7], chapter 5, pp. 77–105.
- [15] 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, 9 1993.
- [16] R. D. Shachter. An ordered examination of influence diagrams. *Networks*, Vol. 20, pp. 535–562, 1990.
- [17] 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.
- [18] 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, 8 1991.
- [19] M. P. Wellman, J. S. Breese, and R. P. Goldman. From knowledge bases to decision models. *The Knowledge Engineering Review*, Vol. 7, No. 1, pp. 35–53, 1992.