

暗黙的提示に基づくアイロニーの解釈モデル

内海 彰

東京工業大学 大学院総合理工学研究科知能システム科学専攻

utsumi@utm.dis.titech.ac.jp

本稿では、アイロニーは現在の発話状況がアイロニー環境であることを暗黙的に提示する表現であるという観点から、アイロニー解釈の計算モデルを提案する。提案するモデルは従来のアイロニー論に比べてより多くのアイロニーを説明することが可能であり、さらにアイロニーに関して心理学や言語学の分野で得られている多くの知見と合致している。このモデルに基づくアルゴリズムは、現在の発話状況がアイロニー環境であるかどうかを調べた後に、アイロニーの表現上の特徴がどの程度満たされているかを評価することによって、与えられた言語表現・発話がアイロニーであるかどうかを決定する。そしてアイロニーである可能性が高い場合には、話し手のアイロニカルな意図を推測する。

A Computational Model of Irony Interpretation Using Implicit Display

Akira Utsumi

Department of Computational Intelligence and Systems Science
Tokyo Institute of Technology

This paper presents a computational model for interpreting verbal irony on the basis of the notion of implicit display of ironic environment. Its main features are an underlying comprehensive theory that covers a wider range of irony than previous theories, and the consistency with various aspects of irony. The algorithm that embodies the model, after checking a discourse context for situational settings for irony, decides whether a given utterance is ironic by measuring to what degree it satisfies linguistic properties of irony, and then infers the speaker's ironic intention.

1 はじめに

アイロニーは一見現実と反することを言いながら真意をはのめかす言語表現であり、比喩とともによく観察される非字義的な表現である。例えば待ち合わせの時間になかなか遅れて来た恋人に対しての発話「ずいぶんと早かったね」はその典型的な例である。アイロニカルな発話によって時間に遅れたことへの非難やいやみを伝えながらも直接的な表現に比べて相手に与えるダメージが弱まるなど対人関係の維持に効果的であり、それゆえにアイロニーは知的で機知に富んだ会話や文章には欠かせない言語現象である。よってアイロニーの解釈機構を計算機上で実現することは人工知能の研究課題になり得るし、さらに (a) アイロニーの認識や理解には文脈・状況の理解が不可欠である、(b) より自然な人間-機械系への貢献が期待できる (Hulstijn & Nijholt 1996)、(c) アイロニーはさまざまな伝達目標を達成することができる (Roberts & Kreuz 1994) などの理由から自然言語処理や計算言語学の分野でも重要なテーマである。

しかしアイロニーの計算モデルや計算機処理を指向した研究はいくつかの例外 (伊藤 & 滝澤 1994; 滝澤 & 伊藤 1994) を除いて皆無である。これは文字通りの意味の反対を意図する単純な修辞法であると考えられていたアイロニーが実はかなり複雑な現象であり、アイロニーを特徴づけるのは非常に困難であることに起因する。この複雑さのため従来の言語学や心理学におけるアイロニー研究においても、アイロニーとは何であり、聞き手がどのようにしてある発話をアイロニーと認識・解釈するのか、という問いに対する十分に満足な解答を与えていない。アイロニーの計算モデルを構築するためには、少なくともこれらの問いに答えることのできるアイロニー論が不可欠である。

筆者はこのようなアイロニー論として暗黙的提示理論 (Utsumi 1996b; 内海 1996) を提案しており、本稿ではこの暗黙的提示の概念に基づいたアイロニーの計算モデルを提案する。また本稿のモデルがアイロニーに関して今までに得られている知見と合致することを示す。

2 アイロニーに関する知見

- (A1) アイロニーは単に文字通りの意味の反対を伝達する表現ではない。

例えば以下の発話(1a)には「文字通りの反対の意味」という概念が適用可能に見えるが、他の発話(1b)~(1e)は平叙文でないのでこの概念を適用することはできない。

状況1 散らっている部屋をきれいにしなさいと注意したが

- が何もしない息子に対して母親が
- a. ものすごくきれいな部屋だわ。
 - b. お母さんは部屋をきれいにしている子が好きよ。
 - c. もうしわけありませんが、どうか部屋を片付けて頂けませんかでしょうか？
 - d. ちょっと散らかっているみたいね。
 - e. こんな快適なところにいて何か感じない？

- (A2) アイロニーは必ずしも選択制限違反や語用論的な違反を含むとは限らない。

語用論的アプローチ(e.g., Grice 1989)では質の公準違反などの語用論的な違反に気付くことによりアイロニーが認識されるとしている。しかし上記の発話—一般化表現(1b)、過度に丁寧な表現(1c)、控え目な表現(1d)—には質の公準違反が認められないにもかかわらず我々はアイロニーとして容易に理解することができる。

- (A3) アイロニーはある人の期待や社会的・文化的慣習を言及(mention)する。

言及理論(Sperber & Wilson 1986; Wilson & Sperber 1992)では、アイロニーとはある人の見解や発話に言及することによって話し手の心的態度を表明するものであると考える。「言及」という概念はアイロニーを説明する上で不可欠であるが、彼らの「言及」の概念は主におうむ返し的なアイロニーの説明が中心であり、上記や以下のアイロニーが何を言及しているのを説明できない。

状況2 妻がおなかが減って眠れないので食べようとしたビザをまるごと平らげてしまった夫に向かって

- (2) a. おなか为本当にいっぱいだわ。
 - b. ああ、ビザが食べれて満足だわ。
 - c. ビザはここにあったのを見なかった？
 - d. これでぐっすり眠れるわ。
 - e. ビザをのこらず食べてくれてありがとう。
- (A4) 字づらの意味が肯定的な(positive)内容のものほどアイロニーとして解釈されやすいが、否定的な(negative)表現によるアイロニーも存在する。

この非対称性は「言及」の概念で以下のように説明可能である(Kreuz & Glucksberg 1989)。一般に言及の対象となる期待や規律は肯定的な内容であるのでそれらを言及しているアイロニー表現も肯定的な内容となるが、以下の発話(3)のように否定的な内容であっても言及の対象となる期待が話し手と聞き手双方にとって明らかな状況下では、期待に言及していることが容易に認識できるのでアイロニーとして解釈される。

状況3 いくら注意しても息子が部屋を散らかしたままなので、妻が夫に息子を叱ってほしいと頼んだ。そこで2人は彼の部屋に行ってみると、部屋はきれいに片付いていた。夫が妻に向かって

- (3) おお、なんと汚い部屋だ。

- (A5) アイロニーは通常さまざまなアイロニー標識(ironic cues)を伴って言語化されるが、このような標識がなくてもアイロニーと解釈される。

アイロニー標識には形容詞や副詞による誇張表現(例: (1a)の「ものすごく」)やアイロニカルな意味を含む語彙(例:「おめでたい」、「えらい」)、感嘆詞(例: (3)の「おお」)、韻律上の特徴(イントネーション・音調やアクセントなど)などの言語的情報の他に表情や身振りなどの非言語的な標識も含まれる。このような標識がアイロニーを認識しやすくすることに疑いの余地はないが、このような標識がなくてもアイロニカルに解釈されることが心理学実験を通じて明らかになっているし(Gibbs 1994)、逆にこのような標識はアイロニー以外の発話にも使われる(Barbe 1995)。

- (A6) 聞き手が語用論的な違反や言及に気付いていなくてもアイロニーを理解することができる。

例えば状況4において、恋人は自分が料理が下手だと思っていなくても彼の発話(4)のアイロニカルな意図を汲み取ることが可能である。

状況4 料理に自信があると自負している恋人の手料理を初めて食べたときに、その料理のまずさを皮肉って

- (4) 君って、本当に料理が上手だね。

- (A7) 話し手が意図していないのに聞き手がアイロニーと解釈するという意図的でない(unintentional)アイロニーが存在する(Gibbs 1994; Barbe 1995)。

3 アイロニー解釈の計算モデル

3.1 アイロニー環境の暗黙的提示

上記の知見を説明するためには、アイロニーであるために必要な状況設定とアイロニーの表現上の特徴を明確に区別することが必要である。そこで暗黙的提示理論では、アイロニーの成立に必要な状況設定としてアイロニー環境(ironic environment)、言語表現上の特徴として暗黙的提示(implicit display)という概念を提案し、「アイロニーは現在の発話状況がアイロニー環境によって囲まれた状況であることを聞き手に暗黙的に提示する言語表現である」と主張する。なお従来のアイロニー研究の問題点をどのように解決しているかなどの暗黙的提示理論の詳細については(Utsumi 1996b, 1996a)を参照されたい。

発話状況が以下の3つの条件を満たしているとき、その状況はアイロニー環境によって囲まれていると言う。

1. 時間 T_0 において話し手がある期待 E を持っている。

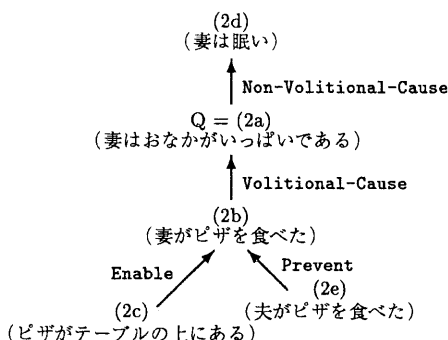


図 1: 状況 2 における婉曲的言及

2. T_0 と同じかまたは T_0 より後の時間 T_1 において話し手の期待 E が満たされていない。
3. 上記の期待と現実との間の不一致に対して、話し手が否定的な心的態度／感情（失望、怒り、非難、ねたみ）を持っている。

例えば状況 2 は話し手である妻の期待「空腹をみたく」が夫の行為（ピザを食べる）によって満たすことができず、妻がそれに対して不満や怒りを抱いているので、アイロニー環境によって囲まれた状況である。

そして言語表現や発話 U が以下の 3 条件を満たすとき、 U はアイロニー環境を暗黙的提示するという。

1. U は話し手の期待 E に婉曲的に言及 (allude) する。
2. U は何らかの語用論の原則に表面上違反することによって語用論的不誠実性を含む。
3. U は話し手の心的態度・感情を暗示する。

本理論では (A3) の言及の概念として、接続関係 (coherence relation) (Mann & Thompson 1987) によって定義される婉曲的言及 (allusion) という概念を用いる。 P を言語表現 U の命題内容を表す命題または行為、 Q を期待 E の内容を表す命題とすると、 P または P の構成要素（これらをまとめて P_i と表す）と Q を関係づける接続関係の連鎖が存在するとき、言語表現 U は期待 E に婉曲的に言及するということ。特に P_i と Q が同じであるとき、 U と E は言及理論でいうところの言及の関係にある。例えば発話 (1a) や (1b) は母の期待「部屋がきれいである」を言及しており、また状況 2 のそれぞれの発話は図 1 に示すように妻の期待に婉曲的に言及している。なお図 1 において Volitional-Cause, Non-Volitional-Cause, Enable, Prevent はそれぞれ意志的な原因、無意志的な原因、可能化、妨害の接続関係を表している。

語用論的不誠実性は (A2) の選択制限違反や語用論的な違反を拡張した概念であり、少なくとも何らかの語用論的な原則に表面上違反しているとき、その発話は語用論的に不誠実になる。これらの語用論的な原則には Grice の質の公準や言語行為の適切性条件（誠実性、

準備、命題内容条件）だけでなく、関連性の原理、丁寧さの原理や量の公準なども含まれている。例えば、(2c) は Ask-if という言語行為の準備条件「話し手はその質問の答を知らない」を違反していることで語用論的な不誠実性を含んでいる。同様に、一般化表現 (1b) は、自分の部屋をきれいにしている子供が存在する文脈において (1b) が「母はその人が好きである」を含意する点で関連性理論 (Sperber & Wilson 1986) の関連性の原理を満たすことができるが、状況 1 の場合にはそのような子供が現在の文脈にいないことから関連性の原理に違反しており、過度に丁寧な表現 (1c) は適切な丁寧さで発話を行なうという丁寧さの原理に違反しているし、控え目な表現 (1d) は量の公準に違反している。

話し手の心的態度・感情の間接的な伝達は (A5) で述べたさまざまな言語的・非言語的標識の使用によって達成される。さらに発話 (2e) のように感謝や賛辞などの実際とは相反する感情を表出することによっても本当の心的態度が間接的に伝達されると考えられる。

アイロニー環境と暗黙的提示の定義により、発話状況がアイロニー環境でないか、またはアイロニー環境の条件を直接的に表現しているときには、他の条件が満たされていてもアイロニーにはなり得ないことがわかる。例えば以下の発話 (2f) の下線部は (2a) と同じ表現であるが、明らかにアイロニカルな意図を含んでいない。

状況 2' 妻はおなかが減って眠れないのでピザを食べようとしたが、夫がまるごと平らげてしまった。それに気付いて「ごめんね」と謝った夫に対して

- (2) f. 気にしないでよ。 おなかが本当にいっぱいだわ。

同様に以下の発話 (1f) ~ (1h) はそれぞれ話し手の期待、期待と現実のずれ、心的態度を直接的に表現しているので、たとえアイロニー環境の 3 条件を満たしている状況 1 で発せられたとしてもアイロニーではない。

- (1) f. きれいな部屋を期待していたのよ。
g. この部屋、本当に散らかってるわ。
h. 散らかったままで、もうがっかり。

3.2 アイロニーの認識と解釈

3.1 節の議論から、聞き手がアイロニーを認識するためにはそれ以前にアイロニー環境の成立を知っていて、かつそれをもとにして暗黙的提示の 3 条件の成立を同定することが必要であるように思われるかもしれない。しかしこれでは 2 章で述べた (A6) や (A7) のような事実を説明することはできない。これはアイロニカルな表現がアイロニー環境を提示する仕方がまさに暗黙的だからであり、本稿では、このような性質を捉えるために、以下の条件を満たすとき聞き手は言語表現 U がアイロニーである（可能性が高い）と認識すると考える。

1. 聞き手は発話状況がアイロニー環境を満たす可能性を否定できない。
2. U は発話状況がアイロニー環境に囲まれていることを直接的に表現していない。

3. 聞き手は U が暗黙的提示の3つの条件のうち少なくとも2つを満たしていることを認識できる。

4. 話し手の期待が既知である場合には、聞き手は U がその期待を婉曲的に言及していると認識できる。

条件1と2は3.1節で述べた明らかにアイロニーない場合を排除するための条件である。また条件1は解釈前にアイロニー環境の成立が既知である必要はないことを述べており、条件3とともに (A6) のアイロニー環境であることを認識していない (話し手の期待や期待と現実の不一致に気付かない) けれどもアイロニーであると判断する場合を説明できる。さらに条件3によって (A7) の意図的ではないアイロニーだけでなく、(A5) で述べたようにアイロニー標識が聞き手に認識されなくてもアイロニーを理解することを許容する。なお条件4は、例えば以下の発話 (1i) が状況1で発せられた場合を考えるとわかるように、婉曲的言及以外の2つの暗黙的提示の条件を満たしているが明らかに文脈と関連を持たないような発話を排除するためのものである。

(1) i. へー、太陽は西から上るんだ。

本稿の解釈モデルでは、条件1と2の成立を検証したあとに条件3と4を考慮した言語表現 U のアイロニー度 $d(U)$ の算出を行う。アイロニー度は0から3までの実数値を取り、5つの尺度 — 言及の婉曲度 d_a 、不誠実性度 d_i 、感情の暗示度 d_e 、話し手の期待の明白度 d_o 、 U の内容の評価値 d_d — を用いて以下の式で定義される。

$$d(U) = d_o \cdot d_a + (1 - d_o) \cdot d_d + f(d_a, d_o) \cdot (d_i + d_e)$$

上式の $f(d_a, d_o)$ は $d_a = 0$ のときに $1 - d_o$ を、その他の場合には1を出力する関数であり、アイロニー度の値が1.5より大きい場合には U がアイロニーであると判断する。この式は上記の条件3と4を正確に反映している。つまり、条件3の暗黙的提示の3つの条件のうち少なくとも2つが成立 (i.e., d_a, d_i, d_e の少なくとも2つが $\gg 0$) している場合には $d(U)$ の値は1.5を越すが、既知である話し手の期待への婉曲的言及がない (i.e., $d_o > 0$ か $d_a = 0$) 場合には1.5を越えることはない。さらに d_d の値が大きい (つまり肯定的内容を持つ) 言語表現ほどアイロニー度が高くなるが、 $d_o = 1$ (つまり期待が共有されている) 場合には d_d の値はアイロニー度に影響を与えないので否定的な内容の言語表現もアイロニーと判断され、(A4) とも整合性がある。

暗黙的提示理論より明らかなように、アイロニーはアイロニー環境の3条件が現在の発話状況で成立していることを聞き手に知らせるという発話内行為を遂行する。この伝達内容を知ることがアイロニーの解釈であり、聞き手は既知である情報を用いて既知でない条件の成立を推論することによって発話内行為を認識し、それらの情報を自分の信念および話し手との共有信念として登録することになる。例えば前述した状況4の例の場合には、発話 (4) をアイロニーと認識することによって彼の心的態度とともに料理が下手であると彼が思っていると知ることになる。

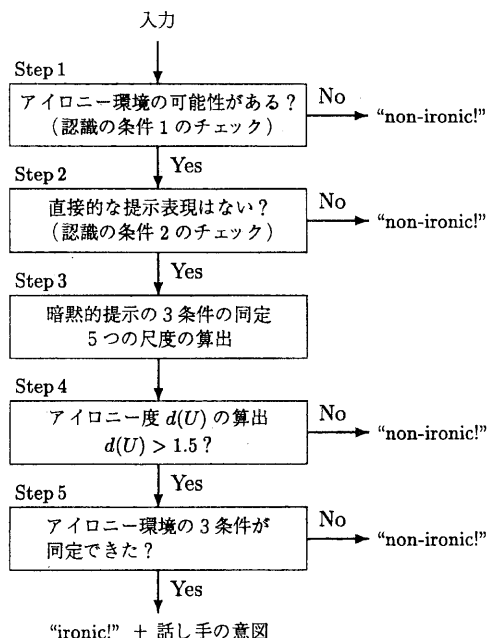


図2: アイロニー解釈アルゴリズム

3.3 アイロニーの解釈アルゴリズム

アイロニー解釈モデルのアルゴリズムの概要を図2に示す。このアルゴリズムは現在のところ英語の表現を対象に Common Lisp プログラムとして計算機に実装されている。アルゴリズムは1) 言語表現 U の命題内容 P 、2) U の字義的な発話内行為 F 、3) U の素性構造による意味表現 M 、4) 発話状況 W と共有知識 K から成る聞き手の信念の集合 C を入力として受け取り¹、入力された言語表現 U がアイロニーであるかどうかを出力する。さらにアイロニーであると判断された場合にはアイロニーの発話内行為を出力する。入力のうち1) から3) までは適当なパーサによって得られると仮定する。本モデルでは言語表現の命題内容 P や聞き手の信念を図3に示すような論理式で表現する。論理式 $F = (! = (B \ T) \ I)$ はできごとや状態を表し、その情報内容 (infor) I と状況 $(B \ T)$ の支持関係として規定される。状況はその信念の主体を表す信念空間 B とそのできごとの時間 T から構成される。本モデルでは信念空間として、“H” (聞き手の信念)、“SH” (話し手の信念に関する聞き手の信念) や “MBH” (聞き手側の共有信念) などを用いる。なおあくまでも聞き手の側からの表現なので、話し手の側からの話し手と聞き手の信念空間はそれぞれ S と HS となる。また論理式 $(\text{neg } F)$ は $(! = (B \ T) \ (\text{not } I))$ を表し、 $\neg F$ は $(B \ T)$ が I で不成立であることを意味する。変数は “?” を用いて表現し、特に “:” の付い

¹ 話し言葉の場合には韻律上の特徴や非言語的特徴も考慮しなければいけないが、本研究ではこれらの定式化は行わない。

```

(!= (MBH ?T) (wife x y))
(!= (MBH ?T) (husband y x))
(!= (MBH ?T) (pizza a))
(!= (MBH ?T) (eatable a))
(!= (MBH ?T) (accessible x Lt))
(!= (MBH ?T) (accessible y Lt))
(!= (MBH ?T) T1>T0)
(!= (MBH T0) (located a Lt))
(!= (MBH T1) (did eat(y,a)))
(!= (MBH T1) (not (did eat(x,a))))
(!= (MBH T1) (blameworthy eat(y,a) x))
(!= (H T0) (hungry x))
(!= (H T0) (hope x (!= (* ?T: ?T>T0)
                        (not (hungry x)))))

```

図 3: 状況 2 における発話状況

た変数は条件づけされた変数である。命題内容 P は論理式だけでなく行為を表す述語によっても表現される。例えば、状態 $(!= (* T1) (not (hungry x)))$ と行為 $eat(y,a)$ はそれぞれ (2a) と (2e) の命題内容である。なお (2a) の命題内容は特定の主体の信念ではないので、信念空間は変数 “ x ” になっている。また字義的な発話内行為 F としては、*Inform*, *Ask-if*, *Ask-ref*, *Request*, *Offer*, *Thank*, *Apologize* の 7 つが用意されている。

聞き手と話し手が共有していると仮定する知識 K は図 4 に示すように領域知識、言語行為に関する知識、感情導出規則から構成される。すべての知識はできごとや状態間の因果関係 $(=> (!= Sit1 I1) A (!= Sit2 I2))$ で表現され、 $(!= Sit1 I1)$ が成立するという前提のもとで行為 A を遂行すると $(!= Sit2 I2)$ が成立するという関係を意味する。意図的な行為 A が関係しない場合は規則 $(=> (!= Sit1 I1) (!= Sit2 I2))$ となり、これは無意志的な因果関係と考える。これらの因果関係を用いて、婉曲的言及に必要な接続関係は以下のように定義できる。例えば $(=> F1 A F2)$ と $(=> F3 F4)$ を仮定すると、 $(vol-cause A F2)$, $(enable F1 A)$, $(non-vol-cause F3 F4)$, $(prevent (neg F1) A)$ の接続関係が成り立つ。また言語行為に関する知識において、 $(ascribe ?P ?X)$ は信念 $?P$ の主体を $?X$ に帰した結果得られる論理式を表す表記である。例えば、 $(ascribe (!= (* T1) (not (hungry x))) x) = (!= (SH T1) (not (hungry x)))$ である。感情導出規則は O’Rourke & Ortony (1994) が提案した規則を修正したものを用いている。本モデルでは “expect”, “want”, “hope” の感情要素で「期待」を表現し、“disappointment”, “angry”, “reproach” の 3 つの感情要素で話し手の否定的な態度を表現する。

図 2 のアルゴリズムはまず Step 1 で話し手の期待 E を W から選択することを試み、 E が見つかった場合には「アイロニー環境の他の 2 つの条件のいずれかが不成立である」ことが C から帰結されるかを調べる。帰結されれば “non-ironic” となり、帰結されなければ Step 2 に進む。例えば図 3 で表現された状況 2 では、聞き手の信念である $(!= (H T0) (hope x ...))$ が E とし

```

Domain Knowledge:
(=> (!= (?B ?T0) (and (accessible ?X ?L)
                      (located ?A ?L) (eatable ?A))))
eat(?X, ?A)
(!= (?B ?T1) (and (not (hungry ?X))
                  (not (located ?A ?L)))))

Speech Act Schemes:
(=> (ascribe ?P ?S) Inform(?S, ?H, ?P)
    (!= (H ?T1) (intend ?S Convince(?S, ?H, ?P))))
(=> (and (!= (SH ?T0) (want ?S (KnowIf ?S ?P)))
        (NotKnowIf ?S ?P)) Ask-if(?S, ?H, ?P)
    (!= (H ?T1) (intend ?S Inform-if(?H, ?S, ?P))))

Emotion-Eliciting Rules:
(<=> (!= (?B ?T) (hope ?X ?I))
    (!= (?B ?T) (and (want ?X ?I) (expect ?X ?I))))
(=> (and (!= (?B ?T0) (hope ?X (!= (?B1 ?T: ?T>?T0)
                                     (not ?I))))
    (!= (?B ?T1: ?T1>?T0) ?I))
(!= (?B ?T1) (disappointed ?X (!= (?B1 ?T) ?I))))
(=> (and (!= (?B ?T0) (hope ?X (!= (?B1 ?T: ?T>?T0)
                                     (not ?I))))
    (!= (?B1 ?T1: ?T1>?T0) (and ?I (did ?A)
                                (blameworthy ?A ?X))))
(vol-cause ?A (!= (?B1 ?T1) ?I))
(!= (?B ?T1) (angry-at ?X Agent(?A) ?A)))
(=> (!= (?B ?T) (and (did ?A) (blameworthy ?A ?X)))
    (!= (?B ?T) (reproach ?X Agent(?A) ?A)))

```

この図において、 $?S$, $?H$, $?P$ はそれぞれ話し手、聞き手、発話の命題内容を表している。また $(KnowIf ?S ?P) = (or (ascribe ?P ?S) (ascribe (not ?P) ?S))$, $(NotKnowIf ?S ?P) = \neg(KnowIf ?S ?P)$ である。

図 4: 本モデルで用いられる共有知識の例

て選択され、期待が現実と一致していることを示す $(!= (* ?T: ?T>T0) (not (hungry x)))$ が C から帰結されるかを求めることになる。なお推論の際には、 $(!= (?B ?T0) ?I)$ と $\neg(!= (?B ?T: ?T>T0) (not ?I))$ が成立していれば $(!= (?B ?T) ?I)$ も成立するという仮定を用いている。また E が見つからない（つまり期待が未知である）場合にはそのまま Step 2 に進む。Step 2 も同様にして直接的な表現であるかどうかを調べる。

次に Step 3 では言語表現が暗黙的提示の 3 条件を満たしているかどうかそれぞれ調べられる。婉曲的言及を調べるには探索手法を用いるが、探索木の深さを 5 に制限しており、すべての P_i に対してこの制限内で解が見つからない場合には、 U は話し手の期待に婉曲的に言及していないと判断する。なお E が未知の場合には探索は行なわない。語用論的不誠実性についてはまず F のすべての適切性条件を調べ、少なくとも一つの条件 Fi に対して (i) 話し手の信念 SH から $(not Fi)$ が帰結されるか、(ii) 話し手の信念からは Fi と $(not Fi)$ のいずれも帰結されないが聞き手自身の信念 H から $(not Fi)$ が帰結される場合には、 P が適切性条件を違反していると判断する。これらの違反が見つからない場合には関連性の原理などの他の語用論の原則の違反を調べる。なお関連性の原理の違反は筆者の関連性の定式化 (内海

& 菅野 1996) に基づいているが、丁寧さの原理や量の公準については、現在の時点では経験的に違反と思われるパターン（例えば「申し訳ありませんが〜」）との照合しか行っていない。

また Step3 ではアイロニー度の算出に必要な5つの尺度の値を以下のように計算する。

- 接続関係の連鎖の長さ n により $d_a = 1 - 0.1n$ とする。婉曲的言及がないときは、 $d_a = 0$ とする。

- U が上記の条件 (i) により誠実性条件違反と判断されたときは $d_i = 1$ 、条件 (ii) の場合には $d_i = 0.8$ とする。 U が他の語用論的な原則を違反しているときには $d_i = 2/3$ 、違反が検出されないときは $d_i = 0$ とする。

- 検出された標識の個数 n より $d_e = 1.0$ ($n > 1$)、 $d_e = 2/3$ ($n = 1$)、 $d_e = 0$ ($n = 0$) とする。

- E が Step1 で見つかったときには、それが共有信念であれば $d_o = 1$ とし、それ以外は $d_o = 0.8$ とする。 E が見つからないときには $d_o = 0$ とする。

- 意味表現 $M = (\text{Rel}:\text{sem1}, \text{Theme}:\text{M1})$ に対して、 $d_d(M) = e(\text{sem1}) + 0.5 \times d_d(M1)$ とする。ただし $e(\text{sem1})$ は (sem1) の評価値)/3 を表し、評価値は3から-3までの整数である。また $d_d(\text{Rel}:\text{not}, \text{Theme}:\text{M}) = -d_d(M)$ とする。そして $d_d(M) > 1$ のとき $d_d = 1$ 、 $d_d(M) < 0$ のとき $d_d = 0$ 、それ以外は $d_d = d_d(M)$ とする。

例えば発話 (2a) を入力文とすると、婉曲的言及をしていると判断され、図1より $d_a = 1 - 0.1 \cdot 0 = 1$ となる。また C は $(\text{!} = (\text{SH T1}) (\text{not} (\text{hungry } x)))$ と $(\text{!} = (\text{SH T1}) (\text{hungry } x))$ のどちらも帰結しないが $(\text{!} = (\text{H T1}) (\text{hungry } x))$ を帰結するので、条件 (ii) より $F = \text{Inform}$ の誠実性条件の違反があると判断され、 $d_i = 0.8$ となる。さらに U は標識を含まないので $d_e = 0$ 、Step1 で見つかった期待は共有信念ではないので $d_o = 0.8$ となる。尺度 d_d は $- \{e(\text{hungry}) + 0.5 \times M(\text{Pro } ?x (\text{I}, ?x))\} = -(-2/3 + 0) = 0.67$ と算出される。

次に Step4 で3.2節のアイロニー度を算出しアイロニーかどうかを判断してから、Step5 でアイロニー環境の3条件を感情導出規則などを用いて推論することによってアイロニーの発語内行為を得る。特に話し手の期待が未知である場合には、Step3 と同じ探索手法により最も近い肯定的な内容をもつできごと・状態を話し手の期待とする。この際に、あるできごと・状態が一般的に望ましいかどうかは d_d の値を計算することによって判断する。なお Step5 でも話し手の期待を推測できなかった場合や、推測はできたけれども現実との不一致がなかったりそれに対して否定的な態度を持っていない場合には、“non-ironic” を出力する。そして3つの条件がすべて同定されれば話し手の意図を出力する。例えば (2a) の場合には、 $d(U) = 0.8 \cdot 1 + (1 - 0.8) \cdot 0.67 + 1 \cdot (0.8 + 0) = 1.73$ から Step5 に進み、以下の発語内行為を出力する。

```
Inform(x,y,(and (!= (SH T0) (hope x (!=
    (?* ?T:?T>T0) (not (hungry x))))))
    (!= (SH T1) (hungry x))
    (!= (SH T1) (angry-at x y eat(y,a))))))
```

4 おわりに

本稿ではアイロニー解釈の計算モデルを提案し、本モデルが今までに得られているアイロニーに関する知見を矛盾なく説明できることを示した。しかし現時点ではまだ改良の余地が多く残されている。主な今後の課題としては、アイロニーによって伝達される発語媒介行為的な話し手の意図・含意の解釈、韻律的な特徴や非言語的標識の考慮、人間の判断との比較による本モデルの評価、などが挙げられる。今後はこれらの課題をふまえて解釈モデルを改良していくつもりである。

参考文献

- Barbe, K. (1995). *Irony in Context*. John Benjamins Publishing Company.
- Gibbs, R. (1994). *The Poetics of Mind*. Cambridge University Press.
- Grice, H. (1989). *Studies in the Way of Words*. Harvard University Press.
- Hulstijn, J. & Nijholt, A. (Eds.). (1996). *Proceedings of the International Workshop on Computational Humor*.
- 伊藤 昭, 滝澤 修 (1994). 対話のモデルに基づくアイロニーの一定式化. 人工知能学会誌, 9(2), 283-289.
- Kreuz, R. & Glucksberg, S. (1989). How to be sarcastic: The echoic reminder theory of verbal irony. *Journal of Experimental Psychology: General*, 118(4), 374-386.
- Mann, W. & Thompson, S. (1987). Rhetorical structure theory: toward a functional theory of text organization. *Text*, 8(3), 167-182.
- O'Rourke, P. & Ortony, A. (1994). Explaining emotions. *Cognitive Science*, 18, 283-323.
- Roberts, R. & Kreuz, R. (1994). Why do people use figurative language?. *Psychological Science*, 5(3), 159-163.
- Sperber, D. & Wilson, D. (1986). *Relevance: Communication and Cognition*. Oxford, Basil Blackwell.
- 滝澤 修, 伊藤 昭 (1994). アイロニー表現検出の一手法. 人工知能学会誌, 9(6), 875-881.
- Utsumi, A. (1996a). Implicit display theory of verbal irony: Towards a computational model of irony. In *Proceedings of the International Workshop on Computational Humor (IWCH'96)*, pp. 29-38.
- Utsumi, A. (1996b). A unified theory of irony and its computational formalization. In *Proceedings of the 16th International Conference on Computational Linguistics (COLING 96)*, pp. 962-967.
- 内海 彰 (1996). アイロニーとは何か? — 言語現象としてのアイロニーのモデル化の試み —. 言語処理学会第2回年次大会発表論文集, pp. 289-292.
- 内海 彰, 菅野 道夫 (1996). 関連性理論を用いた文脈のなかの隠喩解釈の計算モデル. 情報処理学会論文誌, 37(6), 1017-1029.
- Wilson, D. & Sperber, D. (1992). On verbal irony. *Lingua*, 87, 53-76.