

自己生成した説明の受容の予測因子に関する実験的検討

Empirical investigation of which factor is better to predict if self-generated explanation is accepted or not, Likeliness or Loveliness

下條 朝也[†], 三輪 和久[†], 寺井 仁[‡]

Asaya Shimojo, Kazuhisa Miwa, Hitoshi Terai

[†]名古屋大学, [‡]近畿大学

Nagoya University, Kindai University

shimojo@cog.human.nagoya-u.ac.jp

Abstract

When confronted with unfamiliar events we don't understand, we generate explanations why they occur and find the best one among them based on some criterion. In this research, we explore which criterion, *Likeliness* or *Loveliness* explainers use to evaluate self-generated explanations. The result showed that the former, *Likeliness*, is more useful for predicting if explainers accept their self-generated explanations.

Keywords — explanation, hypothesis generation, inference to the best explanation

1. 背景

従来, 科学哲学において, 科学的仮説を受け入れるか棄却するかを判断を科学者自身が行うべきであるとする立場[1]と, 他者に委任するべきであるとする立場[2]が存在し, 議論が行われてきた. 後者の主張は, 生成した説明の受容の可否を決めるには方法論レベルでの価値判断 (認識的価値判断) が存在すべきであるというものである. このことから, 科学におけるトランス・サイエンス問題に象徴されるように, 専門家同士はもちろん, 専門家・非専門家の垣根を超え, 説明者・被説明者の双方が, 不確実な説明を受容すべきかどうかの判断が求められている.

1.1. 説明の生成と評価

人間は説明的動物である[3]. 我々は日常的に, 身の回りで起こる様々な現象に対して, その原因を理解するために, 説明の生成を試みる. その際, 説明を生成する過程で用いられる推論を Abduction, その説明を評価する推論を Inference to the Best Explanation と呼ぶ[4].

しかし, どのような説明を優れていると感じるのかについては, 未だ明らかになっていない. 科学哲学において, 説明の良さを測る規準をめぐって, 対立する 2 つの立場が存在すると言われている[5]. ひとつは, なぜ事象 X_i が起こったのかに対する説明 E_i という条件の

もと, X_i が起こる確率 $P(X_i | E_i)$ (以降「Likeliness」と呼ぶ) が高いほど良い説明であるとする Bayesianism という立場である[6]. もうひとつは, 説明自体への愛着や美しさ (以降「Loveliness」と呼ぶ) が高いほど良い説明であるとする Explanationism という立場である[7]. 説明の美しさは, 単純さや一貫性, 汎用性等で構成されていると考えられている[8]. たとえば, Albert Einstein が提唱した相対性理論は, その発想の奇抜さゆえに Likeliness が高かったとは言いがたいが, その説明の美しさゆえに広く受け入れられた好例であると言える.

また, 説明の生成は, 科学的営みにおいてのみ行われる行為ではない. それは, 推理小説で犯人を予想したり, 友人が不機嫌な理由を考えたりといった日常生活まで, 様々な場面で行われている. その場合, 自らが生成した説明が精査するに足るかどうかの判断を下さなければならない. しかし, 従来の研究では, 他者によって生成された説明に対して下した評価について焦点が当てられてきた.

そこで, 本研究では, 自らが生成した説明を受け入れるかどうかを, Likeliness と Loveliness のどちらの規準がより良く予測するのかを明らかにする.

1.2. 説明の評価に関する先行研究

認知心理学の領域では, 説明の評価に関して, 大きく 2 種類の研究が展開されている.

ひとつは, 被説明者がどのような説明を高く評価するかについての研究である. 他者が生成した説明を受け入れるかどうかは, Loveliness によってより良く予測されることが明らかとなっている[9]. より厳密には, 説明の美しさはその確率判断に影響を与えており, 美しさが低い場合, 高い説明よりも 4 倍以上高い事前確率が提示されてはじめて, その影響を打ち消すことができる[8].

もうひとつは, 説明の生成から評価までの一連の認知プロセスをモデル化する研究である. HyGene と呼ばれるモデルでは, 説明の生成は評価に影響を与えると

考えており、それを踏まえて2段階で構成されている。第1段階では、ある事象 X_1 を観測したとき、一般的に X_1 に関連すると思われる説明 E_1 を、エピソード記憶から取り出す[10]。第2段階では、第1段階において活性化された説明 E_1 を、意味記憶における既知の情報と照合することで評価する。その際、説明 E_1 の確率判断は、 E_1 以外に思いついた説明（以降「代替説明」と呼ぶ）の存在の影響を受けたり、ワーキングメモリー容量の低下あるいは時間の制約によって高く見積もったりすることが明らかとなっている。しかし、この一連の研究では、確率判断のみを評価の対象として扱っており、説明の美しさは考慮されていない。

実際、科学者を含め、我々は他者が提案した説明よりも、自らが生成した説明を信じる可能性が高いことが指摘されている[11]。これは、ある対象を評価する際に、それを「自分のもの」と認識するだけで、より肯定的な評価をするようになる The Mere Ownership Effect という現象が関与していると考えられる。このことから、自らが生成した説明に対しては、他者が生成した説明に対する評価と比べて、説明への愛着や美しさを示す Loveliness の値が高まると予想する。実験では、自らが生成した説明に対して評価を行った「自己生成群」と、他者が生成した説明に対して評価を行った「他者生成群」を設け、Loveliness と Likelihood の比較を行うことで、どちらの基準が説明者の説明受容の可否をより良く予測するのかを検討する。

また、説明評価に関する従来の研究では、選択形式の問題を用いて検討が試みられてきた。しかし、日常生活で説明を生成する際、他者から選択肢が与えられる状況は限られており、現実的であるとは言えない。そこで、本研究では、より現実的な状況を想定するため、複数の回答が考えられる問題を用い、自由記述で説明を生成してもらうこととする。

2. 実験1

本実験では、説明を生成するという行為が、その説明の評価に用いる基準にどのような影響を及ぼすのかについて検討するため、説明生成群に対して実験を実施した。

2.1. 参加者

名古屋大学の学部生 36 名を対象に実験を行った。

2.2. 課題

物語文を3題用いた。1つ目は、洞察問題 ([12]を改

変)である。2つ目は、[13]を改変したもので、一般的に洞察問題とされているが、解に至るまでに論理の飛躍があり、公式に解とされているものを正解と見做さない人が存在することが予備調査で判明した（正答者4人中3人（75%）が、他の回答を最終的な解とした）ため、本研究では準洞察課題と定義する（例1）。最後は、[9]で使用されたものを改変したもので、非洞察問題である。これら全ての問題に共通する設定として、一人の登場人物が何かの結論を出せずに悩んでおり、参加者には、なぜその人物が悩んでいるのか、その理由を説明してもらった。

太郎君はワゴン車を借りて、3人の友達を乗せてラスベガスまでドライブしていましたが、車が途中の小さな町で故障してしまいました。そこで、車を修理している間、彼は床屋に行くことに決めました。

その町には、アルフの店とバリーの店の、2軒の床屋があります。太郎君は、どちらの床屋に入ろうか悩んでいます。

アルフの店は、町の西の方の、ビルの一階にあります。同じビルには、事務用品を取り扱うオフィスが入居しています。バリーは、町の東側を走る通り沿いに店を構えています。その2軒隣は、食品を売るスーパーマーケットです。

店主であるアルフの髪は、整髪料で整えられてはいるものの、切り方がいい加減なせいで、お世じにも綺麗だとは言えません。うなじもほとんど手入れされていないようです。一方、バリーの髪型は丁寧に切りそろえられており、とても好感の持てる髪型です。うなじも綺麗に整えられていました。

アルフとバリーは親友同士で、頻繁に二人で遊びに出かけたり、交互に家を訪ねて家事を手伝ったりしているそうです。

また、アルフは夜遅くまで店を開けています。晩御飯を知り合いの店で食べることがよくあるそうです。バリーの店は、朝早くから営業しています。近くを散歩することが、彼の日課になっているそうです。

友達たちは皆、口を揃えて「絶対バリーの店に行った方がいいよ。自分の髪もきちんと手入れしていない人に切ってもらうべきではない」と太郎君に言いました。しかし、太郎君はアルフの店に行こうかどうか悩んでいました。

例1 実験問題文例（準洞察問題）

2.3. 実験の流れ

参加者は、物語文を読み、なぜ物語がそのような結末になったのかを説明し、その説明について評価することが求められた（図1参照）。実験は、すべて質問紙上で行った。

評価は、3つの質問によって行われた。ひとつ目は、その説明を受け入れるかどうかについての質問である（以降「Accept」と呼ぶ）。参加者は、「はい」また

は「いいえ」での回答が求められた。ふたつ目は、その説明が有り得る可能性 (Likeliness) はどの程度だと思うかについての質問である (以降「Probability」と呼ぶ)。参加者は、0% (絶対に X でない) から 100% (絶対に X である) までの間の数値で回答が求められた。最後は、その説明の質 (Loveliness) はどの程度だと思うかについての質問である (以降「Quality」と呼ぶ)。

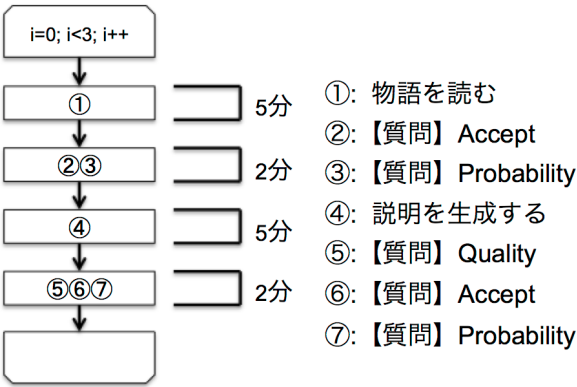


図 1 実験フローについて

参加者は、1 (とても低い) から 7 (とても高い) までの間の数値で回答が求められた。以上は、[9]で用いられていた質問と同内容であり、同様に、Probability を Likeliness の、Quality を Loveliness の測定指標とする。

また、説明の生成以前にも、Accept と Probability についての質問を行うが (図 1 を参照)、それは説明の生成前後で Accept と Probability の評価を比較するためである。これは、ある事象 X_1 についての説明 E_I を採択するのは、説明 E_I のもとで現象 X_1 が起こる確率 $P(X_1 | E_I)$ が $P(X_1)$ より大きくなる時に限るという理論である Bayesian Confirmation Theory を満たしているかを考慮するためである。

3. 結果

3. 1. 洞察問題 ([12]を改変)

36 名を分析した。この問題における説明生成後の Probability の値の平均値は 75.3 ($SD=19.67$) で、最大値が 100, 最小値が 30 であった (単位は「%」)。また、Quality の値の平均値は 4.26 ($SD=1.57$) で、最大値は 7, 最小値は 1 であった。

各説明における Probability・Quality の値が、対応する説明の Accept の可否をどの程度予測するかについて検討を行う。そのために、説明変数を Probability あるいは Quality, 被説明変数を Accept の可否とし、単ロジ

スティック回帰分析を行った。また、その際の Probability・Quality 間の相関は、0.603 ($p<0.01$) だった。

Probability の値を Accept の可否に回帰させたモデル ($p>0.05$, 図 3) と、Quality の値を Accept の可否に回帰させたモデル ($p<0.05$, 図 4) を、AIC と BIC を用いて比較した結果、後者の方がより良く Accept の可否を予測することができた (表 1)。

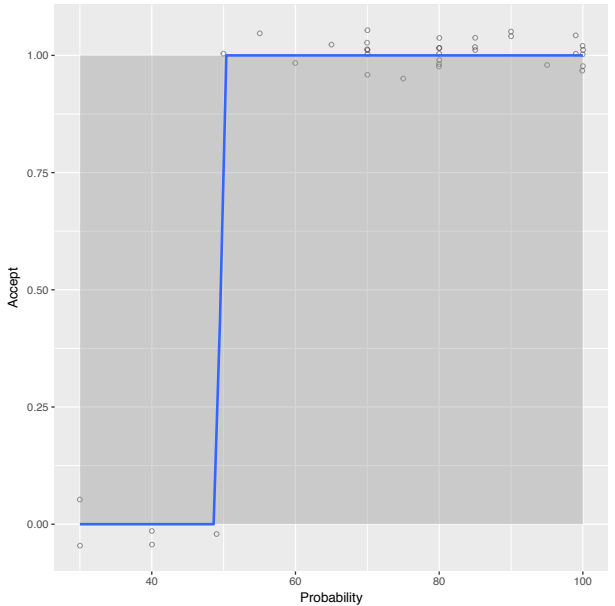


図 3 Probability による Accept の可否の予測 (洞察問題)

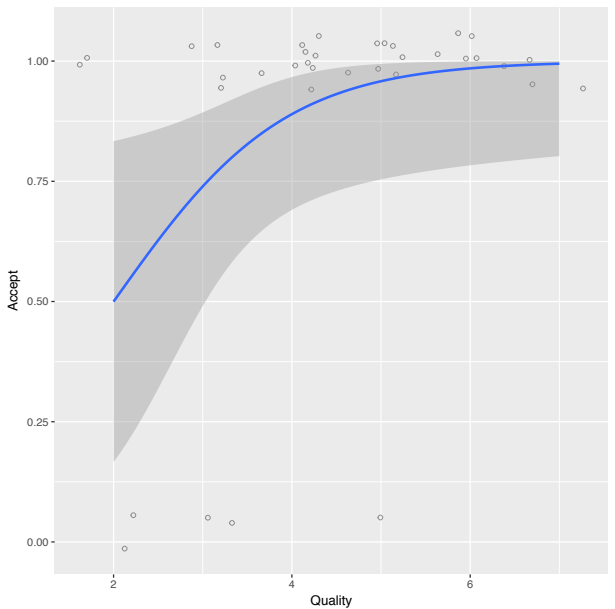


図 4 Quality による Accept の可否の予測 (洞察問題)

	AIC	BIC
Probability	26.60618	29.773221
Quality	4.00000	7.167038

表 1 AIC, BIC を用いた両モデルの比較 (洞察問題)

3.2. 準洞察問題（[13]を改変）

36名を分析した。この問題における説明生成後の Probability の値の平均値は 67.11($SD=21.19$)で、最大値が 100、最小値が 10 であった（単位は「%」）。また、Quality の値の平均値は 4.11($SD=1.81$)で、最大値は 7、最小値は 1 であった。

2.4.1.と同様、説明変数を Probability あるいは Quality、被説明変数を Accept の可否とし、単ロジスティック回帰分析を行った。また、その際の Probability・Quality 間の相関は、0.589 ($p<0.01$)だった。

Probability の値を Accept の可否に回帰させたモデル ($p<0.01$, 図 5) と、Quality の値を Accept の可否に回帰させたモデル ($p<0.05$, 図 6) を、AIC と BIC を用いて比較した結果、前者の方がより良く Accept の可否を予測することができた（表 2）

3.3. 非洞察問題（[9]を改変）

36名を分析した。この問題における説明生成後の Probability の値の平均値は 46.12($SD=21.28$)で、最大値が 90、最小値が 0 であった（単位は「%」）。また、Quality の値の平均値は 3.32($SD=1.60$)で、最大値は 6、最小値は 1 であった。

2.4.1.と同様、説明変数を Probability あるいは Quality、被説明変数を Accept の可否とし、単ロジスティック回帰分析を行った。また、その際の Probability・Quality 間の相関は、0.277 ($p>0.1$)だった。

Probability の値を Accept の可否に回帰させたモデル ($p<0.01$, 図 5) と、Quality の値を Accept の可否に回帰させたモデル ($p<0.05$, 図 6) を、AIC と BIC を用いて比較した結果、前者の方がより良く Accept の可否を予測することができた（表 2）

4. 考察と展望

本実験では、説明を生成するという行為が、その説明の評価に用いる基準にどのような影響を及ぼすのかについて検討した。

結果として、準洞察問題と非洞察問題に対する説明は Probability を用いて、その受容の可否を決定していることが明らかとなった。一方、洞察問題に関しては、Probability を従属変数とした場合、ロジスティック関数に基づく回帰を行うことができなかった。そこで、シグモイド関数ではなく、

$$\text{Accept (Probability)} = \begin{cases} 0 & (\text{Probability} < 50) \\ 1 & (\text{Probability} \geq 50) \end{cases}$$

のステップ関数を仮定すると、Accept の可否をより良く予測できるのは Probability のモデルであると言える。このことから、説明者は、課題のタイプに依存せず、自らが生成した説明に対して Probability を用いて評価していることが示唆される。

また、先行研究[9]と類似の課題（非洞察問題）においては、[9]において示唆された基準（Quality）とは異なる基準（Probability）を用いて評価していることが明らかとなった。この相違は、説明評価者が説明者本人か被説明者かによるものであると考えられる。そのため、続く実験では、本実験で生成された説明を物語文に組み込み、参加者に他者が生成した説明を与えるという、先行研究[9]と同様の状況を設けて実験を行う。そうすることで、非洞察問題において、被説明者による説明受容の判断が Quality によってより良く予測できるかどうかを追試する。また、洞察問題・準洞察問題において、説明評価者の立場（説明者であるか被説明者であるか）によって、評価に用いる基準に違いが見られるかどうかについても検討する。

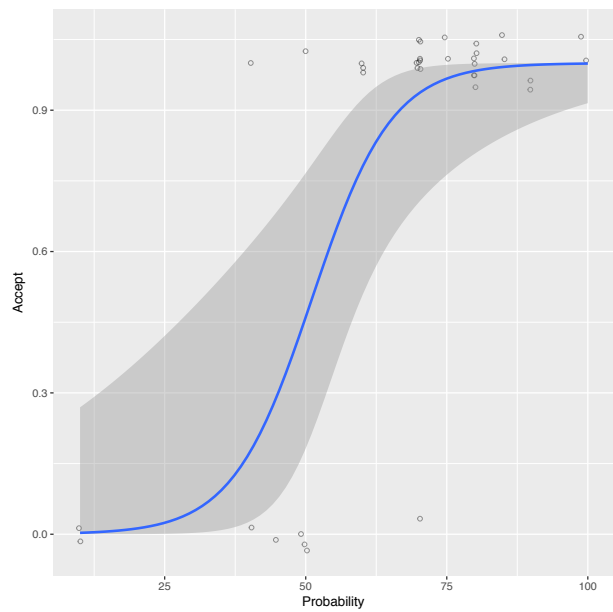


図5 Probability による Accept の可否の予測(準洞察問題)

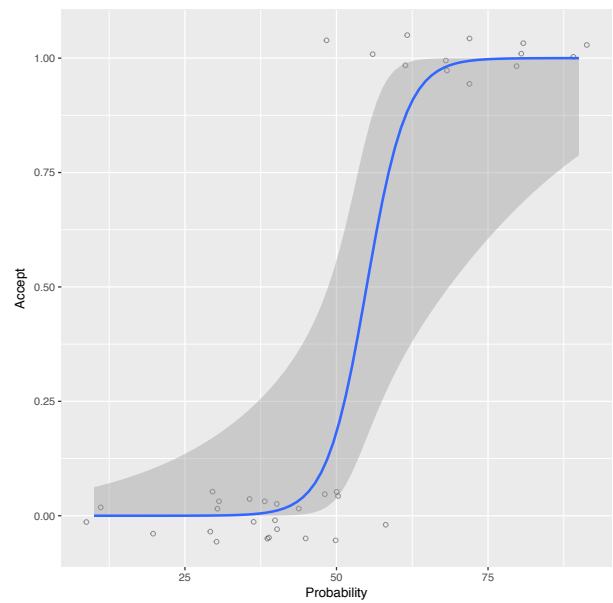


図7 Probability による Accept の可否の予測 (非洞察問題)

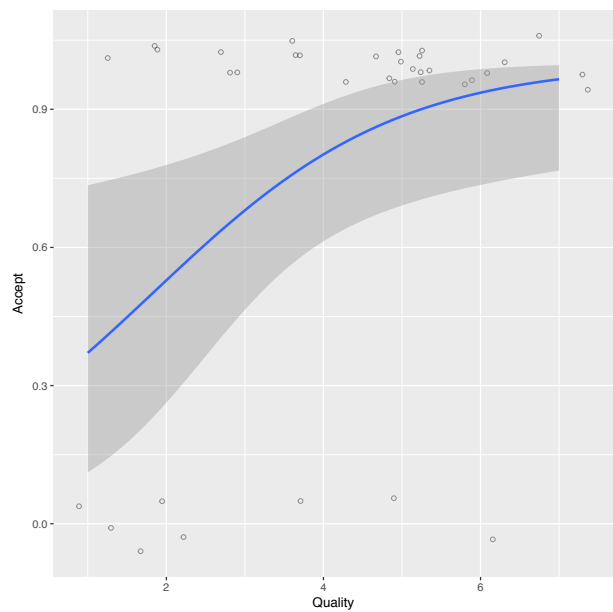


図6 Quality による Accept の可否の予測(準洞察問題)

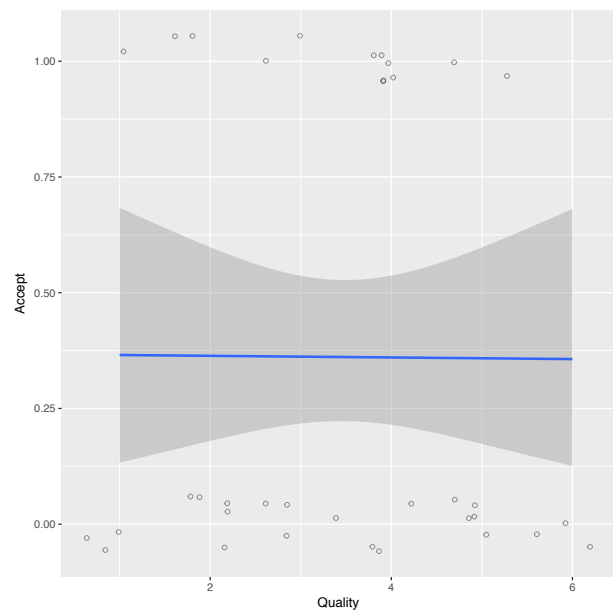


図8 Quality による Accept の可否の予測 (非洞察問題)

	AIC	BIC
Probability	22.21283	25.37987
Quality	35.46508	38.63212

表2 AIC, BIC を用いた両モデルの比較(準洞察問題)

	AIC	BIC
Probability	15.04278	18.20982
Quality	51.09089	54.25793

表3 AIC, BIC を用いた両モデルの比較(非洞察問題)

参考文献

- [1] Richard C. J., (1956), "Valuation and Acceptance of Scientific Hypothesis", *Philosophy of Science*, 23, 237-246.
- [2] Rudner R., (1953), "The Scientist Qua Scientist Makes Value Judgments", *Philosophy of Science*, 20, 1-6.
- [3] Norman. D. A., (1988), "*The Psychology Of Everyday Things*", Basic Books
- [4] Campos, G., (2011), "On the distinction between Peirce's abduction and Lipton's inference to the best explanation". *Synthese*. 180 (3): 419-442.
- [5] McGrew, T., (2003) "Confirmation, heuristics, and explanatory reasoning", *British Journal for the Philosophy of Science*, 54, 553-67.
- [6] Johnston, A. M., Johnson, S. G. B., Koven, M. L., & Keil, F. C., (2016), "Little Bayesians or little Einsteins? Probability and explanatory virtue in children's inferences", *Developmental Science*, 1-14.
- [7] Lipton, P., (2004) "*Inference to the Best Explanation*", Second edition. London and New York: Routledge.
- [8] Lombrozo, T., (2007), "Simplicity and probability in causal explanation", *Cognitive Psychology*, 55, 232-257.
- [9] Douven, I. and Mirabile, P., (2018), "Best, second-best, and good-enough explanations: How they matter to reasoning", *Journal of Experimental Psychology: Learning, Memory, and Cognition*, Feb 01, No Pagination Specified
- [10] Dougherty, M. R. P., & Hunter, J. E. (2003), "Hypothesis generation, probability judgment, and individual differences in working memory capacity". *Acta Psychologica*, 113, 263-282.
- [11] Schunn, C. D., & Klahr, D., (1993), "Self vs. other-generated hypothesis in scientific discovery", *Proceedings of the 15th Annual Conference of the Cognitive Science Society*, 900-905.
- [12] Schooler, J. W., Ohlsson, S., & Brooks, K. (1993), "Thoughts beyond words: When language overshadows insight", *Journal of Experimental Psychology: General*, 122, 166-183.
- [13] Gardner, M., (1978), "*Aha! Insight*", Mathematical Association of America, (邦訳, 「Aha! insight ひらめき思考」, 島田一男 訳, 日本経済新聞出版社, (2009))