

当事者本人の状態像理解を深める 認知症ケアインタラクション表現モデルの構築

小野塚 優志^{†1} 中野目 あゆみ^{†2} 香山 壮太^{†2}
小俣 敦士^{†1} 石川 翔吾^{†1} 桐山伸也^{†1}

概要：認知症ケアの現場において、ケアの質をデータに基づいて客観的に評価する仕組みは未だ発展途上である。筆者らは、ケア場面における人や環境の状況を複数の方法でセンシングし、データに適切な解釈を与えてケアインタラクションの特徴を複数の観点で見える化し、よいケアとは何かを客観化するのに役立つエビデンスを創出・蓄積している。代表的な認知症ケア技法であるユマニチュードを対象に、介入行動の特徴を捉える表現モデルを設計・活用してケアの客観的な評価を進めている。本稿では、既存の行動表現モデルを拡張し、被介護者本人の状態を記述でき、介護者と被介護者の関係性を捉えられるインタラクション表現モデルを提案する。複数の事例で行動記述を実践し、ケアのエキスパートの視点で、提案モデルがケア場面における介入行動の的確な評価に有用であることを示す。

キーワード：マルチモーダル、認知症ケア、状況理解、表現モデル、エビデンスベースドケア

Multimodal Understanding the Situation Sensing for Creating Evidence for Dementia Care

MASASHI ONOZUKA^{†1} AYUMI NAKANOME^{†2} SOUTA KAYAMA^{†2}
ATSUSHI OMATA^{†1} SHOGO ISHIKAWA^{†1} SHINYA KIRIYAMA^{†1}

Keywords: multimodal, dementia care, understanding the situation, expression model, evidence based care

1. はじめに

認知症は医学的な原因で生活障害が生じる症状群であり、認知症原因疾患の回復は難しいため、関わりや環境の調節といった工夫によって穏やかな生活を持続させる取り組みが実践されている。特に、認知症の精神症状の緩和には対人コミュニケーションを中心とした非薬物療法が効果的であると明らかになってきている [1]。実際に、認知症ケアの熟練者は、認知症のある本人の状況を読み取りながら反応に基づき適切なケアを選択し実践している。しかし、認知症ケアを表現する構造の整備がされていないため、熟練者が培ってきた介護のスキルやノウハウは経験的なものに留まっている。そのため暗黙的な技術が多いこと、有効性に関するエビデンスが少ないことが課題として挙げられる。

これまで我々は、ケア従事者の行動に着目し、認知症ケアを情報学的に表現することでケアスキルを評価し、現場にフィードバックするための仕組みを開発してきた [2]。しかし、認知症のある本人の行動の意図の表現が未着手であるため、ケアのスキルと効果の関係を表現できないという課題があった。これは、ケアスキルが適切に認知症のある本人に活用されているかを評価すること

に加え、認知症のある本人の心的・身体的状態を情報学的に表現することにつながる。そこで本稿では、ケアインタラクションの観点から、認知症のある本人の行動を評価するための構造化、さらにさらにスキルと行動の関係を表現するための情報モデルとその表現能力について述べる。

2. ケアインタラクションの表現基盤

我々はこれまでに認知症ケア技法ユマニチュード® (Humanitude) に着目し、情報学的なケアスキルの評価を進めてきた [3]。ユマニチュードはマルチモーダルコミュニケーションを体系化しており、実践することで、薬物投与量の減少や行動・心理症状 (BPSD) の減少、スタッフの離職率低下などの成果が報告されている [3]。ユマニチュードにおける「見る」「話す」「触れる」という基本スキルを表現するために、下記の3層のレイヤーで表現する基本スキームを設計した [4]。

- Intra-modality : 行動の最小単位を表す。見る、話す、触れる等

^{†1} 静岡大学
Shizuoka University
^{†2} 郡山市医療介護病院
Koriyama Medical Care Hospital

- **Inter-modality :**
Intra-modality の関係性を表す. 各要素の同時性, 包括性, 連続性, 順序関係等
- **Multimodal-interaction :**
動作の行為者間の関係を表す. アイコンタクト, 対話, 非対話, 樹里, 拒否等

このようなスキームを活用することにより, 個人内の行動と個人間の行動を区別して表現することが可能となる. また, それぞれのレイヤを適切に評価するために, エキスパートの知識を基にした行動プリミティブを設計した [5]. 以下表 1 に示す. これらの構造とスキームでケアスキルを観察ベースで評価するために図 1 のマルチモーダル行動分析システムを開発し, 映像データへのアノテーションと知識ベースの設計, 統計的解析の基盤を構築している.

表 1. 行動プリミティブ

モダリティ	内容	詳細
見る	見る対象	介護者, 被介護者, 物, 空間など
	見ている部位	目, 顔, 体など
	対象との距離	20cm, 50cm, 50cm以上
	水平方向の位置	真ん中, 右, 左
	垂直方向の位置	真ん中, 上, 下
話す	話す対象	介護者, 被介護者, 独り言, 複数人など
	発話内容	発話の書き起こし
	発話内容分類	依頼, 命令, 約束, 質問, 否定, 警告, 感謝など
	音の高さ	高い, 低い, 普通
	声の大きさ	大きい, 小さい, 普通
	声のスピード	速い, 遅い, 普通
触れる	触る対象	介護者, 被介護者, 別介護者など
	触れる部位	指, 掌, 腕, 肩, 頭, 口腔内, 唇, 背中など
	使用する手	右, 左
	接触部位	指, 手全体, 腕から先全体など
	親指の使用	あり, なし
	媒介物	手袋, 歯ブラシ, 衣類, なし
	着陸	できている, できていない, わからない
	離陸	できている, できていない, わからない
	ストローク	速い, ゆっくり, 普通, なし

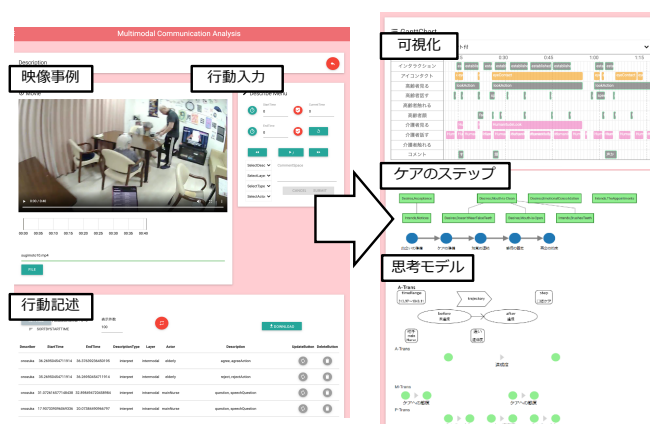


図 1. マルチモーダル認知症ケアインタラクション分析ツール

で収集された映像事例に対して, 知識構造を用いて表現を行う. 表現された映像事例を用いてエキスパートにヒアリングすることによりエキスパートの主観的な知見を抽出する. 取り出された知見をもとに知識構造のアップデートを行い, エキスパートの考えとズレがないか再度ヒアリングを行う. このようなサイクルを形成することで, センシング基盤を継続的に発展させる仕組みを構築する. 具体的には, 認知症のある本人の状態を表現するための基本行動プリミティブの設計を 3 節で, インタラクションを評価するためのインタラクション表現モデルを 4 節で述べる.

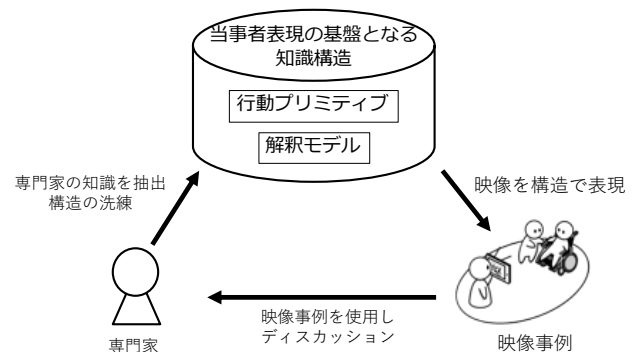


図 2. 知識構造設計の流れ

3. 認知症のある本人の行動表現モデルの設計と評価

3.1 認知症のある本人の行動表現モデルの設計

本基盤で設計する知識構造ではインタラクションを表現することを想定しているが, 先行研究における 3 層のレイヤーにおいても, Multimodal-Interaction レイヤーにおいてインタラクションの表現を行うことができる. よって今回設計する知識構造では先行研究における表現の構造を利用し, 新たなインタラクションのモデルを追加することで発展させることを目指す. 同様に, 3 層のレイヤーの Intra-modality を表現するために設計された, 表 1 の行動プリミティブを使用することとする. しかし, 認知症当事者は言語能力が低下し, 非言語的な情報の重要度が高い [6]. 実際に, 今回使用する郡山医療介護病院で撮影されてきた映像事例は, ベッドで寝ている認知症当事者に対して行われるケアが多く, 学習顔で得られたケアのインストラクターからのコメントには表情に関するものも得られていた [2]. そこで, 表 1 の行動プリミティブに対して, 被介護者の顔の情報を捉える要素を追加する. 改良した行動プリミティブを表 2 に示す.

本研究では, これらの基盤技術を活用し, 図 2 に示すケア現場

表 2. 改良した行動プリミティブ

モダリティ	内容	詳細
見る	見る対象	介護者, 被介護者, 副介護者, 物, 空間
	見ている部位	目, 顔, 体, 例外
	対象との距離	20cm, 50cm, 50cm以上
	水平方向の位置	真ん中, 右, 左, 例外
	垂直方向の位置	真ん中, 上, 下, 例外
話す	話す対象	介護者, 被介護者, 独り言, 複数人, 対象者不明
	発話内容	発話の書き起こし
	発話内容分類	依頼, 命令, 嘆願, 約束, 陳述, オートフィードバック, 質問, 許可, 否定, 笑い声, 忠告, 警告, 祝福, 感嘆, 感謝, 陳謝, 歌, その他, 例外
	音の高さ	高い, 低い, 普通
	声の大きさ	大きい, 小さい, 普通
触れる	触る対象	介護者, 被介護者, 副介護者
	触れる部位	指, 掌, 腕, 肩, 頭, 目蓋, 鼻, 口腔内, 唇, 背中, 首, あご, 胸, 腹部, 脚, 足
	使用する手	右, 左
	接触部位	指, 手全体, 腕から先全体など
	親指の使用	あり, なし
顔	媒介物	手袋, 歯ブラシ, 衣類, なし
	着陸	できている, できていない, わからない
	離陸	できている, できていない, わからない
	ストローク	速い, ゆっくり, 普通, なし
	部位	全体, 目, 眉間, 口
	種類	うなずく, 横にふる, 動く, シワがよる

次に, 3層のレイヤーに対して, 行動プリミティブの組み合わせにより付与するための解釈モデルの設計を行う。

まず Intra-modality について設計を行う。Intra-modality では行動の最小単位を行動プリミティブから表現する。先行研究では, 介護者の見る, 話す, 触れる, のモダリティに関してユマニチュードにおいて良いとされる Intra-modality の解釈である HumanitudeLook, HumanitudeSpeech, HumanitudeTouch がモデル化されている。被介護者の表現に関しては, 重度の認知症当事者とのインタラクションでは非言語的コミュニケーションの重要度が増す。改良を行った行動プリミティブでは介護者のスキルだけでなく, 被介護者の状態を捉える目的も含んでいる。介護者のアクションに対してリアクションがあったのか, を捉えるためには, 各モダリティの対象が被介護者から介護者に向けられて行われているものを取り出す必要があると考える。これにより, 上位レイヤーでのモデル設計の簡略化につなげることができる。なお顔のモダリティに関しては対象を必要としないため, 行為自体を抽象化する。以上を一般化した Intra-modality に含めることとする。先行研究と今回設計した Intra-modality のモデルを表 3 に示す。

表 3. 改良した Intra-modality の解釈モデル

対象とする行為者	解釈名	付与ルール
介護者	HumanitudeLook	行為者が対象者の目を正面から30cm程度の距離で見ている
	HumanitudeSpeech	行為者が対象者に発話している
	HumanitudeTouch	行為者が相手を親指を使わず媒介物を通して指より大きい面積で触れる
被介護者	lookAction	行為者が対象者の顔か目を見ている
	speechAction	行為者が対象者に話しかけている
	touchAction	行為者が対象者に触れている
	faceAction	行為者の顔に何らかの動作が起こっている

次に Inter-modality について検討を行う。Inter-modality では Intra-modality の関係性について記述する。ここでは, Intra-modality で抽象化した行動をさらに抽象化することで, 被介護者の行動の有無を表現する Action を追加する。Inter-modality の解釈モデルを表 4 に示す。

表 4. Inter-modality の解釈モデル

対象とする行為者	解釈名	付与ルール
被介護者	Action	lookAction, speechAction, touchAction, faceActionのいずれかが行われている

最後に Multimodal-interaction について検討を行う。Multimodal-interaction は介護者と被介護者の行動の関係性について記述する。先行研究においては, 介護者と被介護者の目が合っている際にアイコンタクトの解釈を付与している。ユマニチュードでは, 介護者は被介護者に対して反応を求める話しかけを行なった際, 3秒待つ, というものがある。そこで, ここでは介護者の話しかけた行為に対して被介護者が反応を示しているかを示す interaction を追加する。Multimodal-interaction のモデルを表 5 に示す。

表 5. Multimodal-interaction の解釈モデル

解釈名	付与ルール: 介護者	付与ルール: 被介護者
EyeContact	HumanitudeLookをしている	介護者の目を見ている
interaction	HumanitudeSpeechをしている	介護者のHumanitudeSpeech後3秒以内にActionを行う

3.2 認知症のある本人の行動表現モデルの再現性に関する評価実験

作成した知識構造が現場の解釈を表現できているか検証を行うため, 映像事例中の認知症当事者の状態に対する解釈のデータ収集実験を行った。実験手順は次の通りである。

1. 5名の看護師の, 歯磨きケアを行う 10本の映像事例の認知症当事者の状態に関して, 郡山医療介護病院のインストラクターのコメントデータを収集
2. 映像事例に対し, 設計した知識構造とコメントデータを付与し, 可視化する
3. オンラインにてディスカッションを行い, 可視化した結果とコメントを照らし合わせながら認知症当事者の状態についてヒアリングを行い, その内容が表現できているか検証する

本実験のディスカッションには, 郡山医療介護病院の看護師とインストラクター各1名が参加した。映像事例に対し知識構造とコメントを付与し可視化したものを図 3 に示す。図の左にある項目が各モダリティであり, 右がタイムライン表示となっている。各モダリティに要素が存在する場合, タイムライン上に四角く表示される。

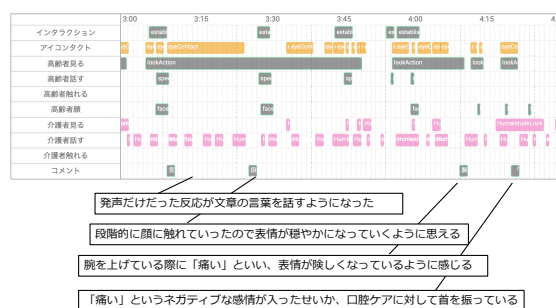


図 3. 映像事例の可視化一例

本実験の結果、77 個のコメントデータを収集した。
得られたコメントの例を表 6 に示す。

表 6. コメントと知識構造による表現の可否

番号	表現出来たコメント	表現出来なかったコメント
1	あつてすぐ返事をし目が合っている	「はい」と言った後に笑顔が見られた
2	顔いて発声で返事をしている	笑顔で接してくれ、いい関係ができている
3	声かけに対してうなずきと発声で返事をしている	覚えてますか？という言葉で困惑の表情のように見える
4	お願いしたことに対して協力的に口を開けてくれている	笑顔で会いにきたことでこやかな表情になっている
5	たんの絡みのせいか、うなずきや発声が減少している	職員に対して手を合わせ感謝の気持ちを示している
6	アイコンタクトをとっている	
7	質問に対してうんうんと顔いてくれる	
8	暖かいタオルが心地良かったのか、発声・うなずきが見られる	
9	協力的に口を開けてくれている	

結果として、表情に関するものと、手を合わせるなどの行為については表現することができていなかったが、介護者の問いかけに対しての反応に関するコメントは、作成したモデルで表現することができており、設計したモデルが現場の人の解釈を表現することに使用できると考えられる。

4. ケアインタラクション表現モデルの設計と評価

4.1 インタラクション表現モデルの設計

インストラクターのコメントと先行研究における知見を踏まえ、知識構造のアップデートを行う。

・intra-modality の解釈モデルの検討

コメントデータの中に、「質問に対して顔いている」「うなずきと発声で返事をしている」とある。これらは、うなずきや発声といった行為が同意を意味する行為であると考えられる。また、「首を振っている」というコメントがあることから、拒否を示す行為であると考えられる。こうした同意拒否に関しての行動が発話として表出する可能性もある。また「声かけ」や「質問」というワードがいくつかあることから、発話を質問的な意味で捉えられるものとそれ以外で分類することにも意味があると考えられる。以上のことから、表 7 のようにモデルを更新した。

・Inter-modality の解釈モデルの検討

Intra-modality で検討した「同意」「拒否」の行動をこのレイヤーでまとめて表現することで、認知症当事者の反応をまとめて表現することが出来、上位レイヤーのモデルの設計や可視化の際に役に立つ。以上のことから表 8 のようにモデルを更新した。

表 7. 更新後の Intra-modality の解釈モデル

対象とする行為者	解釈名	付与ルール
介護者	HumanitudeLook	行為者が対象者の目を正面から30cm程度の距離で見ている
	HumanitudeSpeech	行為者が対象者に発話している
	HumanitudeTouch	行為者が相手を親指を使わず媒介物を通さず指より大きい面積で触れる
被介護者	lookAction	行為者が対象者の顔か目を見ている
	speechAction	行為者が対象者に話しかけている
	speechAgree	行為者の発話内容分類が"許可"
	speechReject	行為者の発話内容分類が"否定"
	touchAction	行為者が対象者に触れている
	faceAction	行為者の顔の動きが、口が動いている、眉間にシワがよっている、目が動いている、のいずれか
介護者/被介護者	faceAgree	行為者が顔いている
	faceAgree	行為者が顔を横に振っている
介護者/被介護者	speechQuestion	行為者の発話分類が"質問"、"依頼"、"命令"、"懇願"のいずれか

表 8. 更新後の Inter-modality の解釈モデル

対象とする行為者	解釈名	付与ルール
被介護者	Action	lookAction, speechAction, touchAction, faceActionのいずれかが行われている
	Agree	speechAgree, faceAgreeのいずれか
	Reject	speechReject, faceRejectのいずれか

・Multimodal-interaction の解釈モデルの検討

まず、Inter-modality レイヤーで同意・拒否の表現ができていることから、介護者の問いかけに対して被介護者の同意・拒否が行われたインタラクションを表現することができるようになった。これにより表 5 に示した interaction を分解し、同意しているインタラクションと拒否しているインタラクションと分けることができる。ディスカッションにおいて、「うなずきで反応している」のようなコメントが付与された場面が、被介護者に質問的な発話をしている場面であったことから、介護者の発話は speechQuestion に限定する。また、「お願いしたことに対して協力的に口を開けてくれている」というコメントがあるが、使用した映像が歯磨きケアであることから表出したコメントであると考えられる。歯磨きケアの際、口を開けるよう依頼したことに対しては口を開ける行為が同意した、と解釈できると考えられる。また否定に関して、介護者の声かけの際、被介護者が目線を逸らす行為を行っていた場面が映像事例の中で見受けられた。その事例の可視化ビューを図 4 に示す。ここから、声かけに対して目線を逸らした場面は拒否と解釈できると考えられる。



図 4. 視線を逸らした認知症当事者

次に、ユマニチュードの理念や先行研究により蓄積してきたインストラクターの知見を用いて、解釈のモデル化を試みる。

まず今回のコメントでも出ていた「協力的」について検討する。コメントの中で「協力的」とコメントが出したのは「口を開けてくれる」という動作に対して付与されたコメントであるが、時間窓を広げケア全体を通じて協力的であったことを表現するために必要な項目について検討する。協力的な動作の一例としては、「うなずく」「(歯磨きケアにおいて)口を開ける」が挙げられる。これらは Inter-modality において Agree と表現したものである。しかし Agree の頻発により協力的と解釈すると、介護者との相互作用の中で協力している、という意味にはならない。よって介護者の発話に対して Agree を示す必要があると考えられる。すでに定義した agreeInteraction は介護者の質問発話に限定して表現しているため、HumanitudeSpeech に対して Agree が表出したインタラクションを新たに定義し、このインタラクションの連続発生により「協力的」という解釈の付与が可能なのではないかと考える。よってこのようなインタラクションを positiveInteraction と定義する。同様に、「協力的でない」を表現できるモデルを設計することで被介護者が望まないケアが表現できると考え、rejectInteraction とは異なるインタラクションとして negativeInteraction を定義する。

次に、ユマニチュードの本質である認知症当事者に対する気遣いの逆となる、「一方的なケア」について検討する [1]。一方的なケアとは、介護者が目的の遂行を優先するあまり、被介護者の状態を考慮せずに行われるケアを意味している。重度の認知症当事者は、「顔」や「見る」のモダリティで自身の状態を伝えようとすることが多い [6]。介護者が、被介護者のこれらのサインの表出に気づくことが出来ないインタラクションや、相手のリアクションを待てずに畳みかけて話しかけるようなインタラクションが多く出た場合に「一方的なケア」という解釈ができると考える。よって、相手のリアクションを待てないインタラクションを actionMiss として定義する。当事者が発するサインの見逃しとして、同意していることを見逃しているようなインタラクションは、被介護者のことに意識を向けられていないインタラクションであると考え、agreeMiss と定義する。サインが捉えられない場面とは逆に、サインを捉えている事例について検討を行ったところ、図 4 の事例において視線を逸らした被介護者に対して介護者が気

づき、質問発話を行っている場面でもあることがわかった。この場面より、被介護者のサインを捉える signCatch を定義する。

最後に、人と人との関わりとして沈黙が続く時間を考慮することには意味があると考え、ユマニチュードにおいては、例えば話すのが困難な認知症当事者であっても話しかけ続けることが重要であると説いている [1]。また先行研究では、ユマニチュード技法を用いたケアのなかではいけないことをまとめたワーニングアラートが開発されてきた。その中で沈黙は 30 秒以上介護者・被介護者の発話がない場合と定義されているため [7]、今回はこの指標を利用し silence と定義する。

以上を踏まえた Multimodal-interaction のモデルを表 9 に示す。

表 9. 更新した Multimodal-interaction の解釈モデル

解釈名	付与ルール：介護者	付与ルール：被介護者
EyeContact	HumanitudeLookをしている	介護者の目を見ている
agreeInteraction	介護者のspeechQuestion後3秒以内にspeechQuestionをしている	介護者のspeechQuestion後3秒以内にAgreeか、介護者のspeechQuestion後3秒以内に口を動かす動作を行う
rejectInteraction	介護者のspeechQuestionをしている	介護者のHumanitudeSpeech後3秒以内にrejectか、介護者のHumanitudeSpeech後3秒以内に被介護者のlookActionが外れEyeContactが外れる
positiveInteraction	HumanitudeSpeechをしている	介護者のspeechQuestion後3秒以内にAgree
negativeInteraction	HumanitudeSpeechをしている	介護者のHumanitudeSpeech後3秒以内にreject
signCatch	HumanitudeSpeechの3秒後に被介護者の指定リアクションが行われている、かつその3秒後にspeechQuestion	介護者のspeechQuestion後3秒以内にfaceAction
agreeMiss	介護者のspeechQuestionの3秒以内に被介護者の指定リアクションが行われている、かつ介護者が1度目のspeechQuestionから秒以内にspeechQuestion	介護者のspeechQuestion後3秒以内にAgree
actionMiss	介護者のspeechQuestionの3秒以内に被介護者のリアクションがなく、かつ介護者が1度目のspeechQuestionから3秒以内にspeechQuestion	介護者のspeechQuestionの3秒以内にActionが発生しない
silence	HumanitudeSpeechがなく被介護者の指定したリアクションがある状態が30秒続く	speechActionがない

4.2 インタラクション表現演算子

設計した解釈モデルのルールを映像事例に付与するため、図 1 のツールに実装されているクエリの発行機能を利用する。これは、登録したクエリを映像事例に対して発行し、クエリの条件に当てはまるラベルを付与することができる機能である。クエリの設計では、付与されるラベルの指定と、条件の指定を行うことができる。条件として、1つのモダリティだけでなく、複数のモダリティが関わり合うような条件も指定できるような演算子が用意されている。登録されている演算子を次に示す。

- **Equal**
条件に合致する記述と同じ時間区間を登録する。
条件は1つのみ指定可能。
- **Not**
条件に合致する記述が存在しない時間区間を結果として登録する。条件は一つのみ指定可能。
- **And**
各条件に合致する記述が同時に存在する時間区間を結果として登録する。条件は複数指定可能。

- **Or**
各条件に合致する記述のいずれかが存在する時間区間を結果として登録する。条件は複数以上指定可能。
- **Continue**
条件に合致している記述の中でも、指定した時間続いている、または、続いていない記述の時間区間を結果として登録する。
- **Index**
条件に合致する記述を、時間順(昇順)に並べた時、指定した位置にある記述の時間区間を結果として登録する。
- **Interval**
条件に合致する記述を、時間順(昇順)に並べた時、記述と記述の時間区間を登録する。
- **After**
2つの条件を x, y とする。条件 x の後に条件 y が指定時間内に行われた時、条件 x の開始点と条件 y の終点を時間区間として登録する。
- **Before**
2つの条件を x, y とする。条件 x の前に条件 y が指定時間内に行われた時、条件 y の開始点と条件 x の終点を時間区間として登録する。

今回設計した解釈モデルを表現するにあたり、演算子 After を活用できる。演算子 After は、2つの条件が指定時間間隔で行われている場合を取りだす仕様となっている。条件 x の終了地点から P 秒以内に条件 y を満たしたときに新たな解釈を付与するクエリを作成することができる。この演算子により、介護者と被介護者のインタラクションを表現することができる。演算子 After のイメージを図 5 に示す。これらの演算子を駆使して、モデルを表現するクエリを作成し実装した。

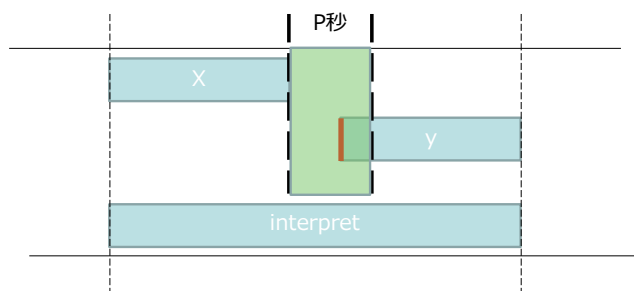


図 5. 演算子 After の仕様

4.3 ケアインタラクション表現モデルの再現性に関する評価実験

設計した解釈モデルが、映像事例にエキスパートが付与する解釈と一致するか検証を行うため、付与されたラベルに対するエキスパートの評価実験を行った。実験手順は次の通りである。

1. 6名の看護師の、歯磨きケアを行う 11本の映像事例に対し、設計した解釈モデルを付与し可視化したビューを用意
2. 各映像事例に対して付与された解釈 1つ 1つに対して、映像事例を見ながらどう解釈できるのか意見を収集
3. 収集した意見と付与した解釈モデルを比較し、一致率を検証

本実験は遠隔で行われ、郡山医療介護病院の看護師 1 名が参加した。意見収集の際に使用した可視化ビューの例を図 6 に示す。タイムライン上に、付与された解釈とその説明を入れ、対応する動画名を記載している。なお、倫理的な観点から、映像事例のオンライン上での共有ができないため、事前に動画をメモリーディスクに入れ送付している。

本実験では、3つの映像事例に付与された解釈について意見を収集することができた。結果を表 10 に示す。

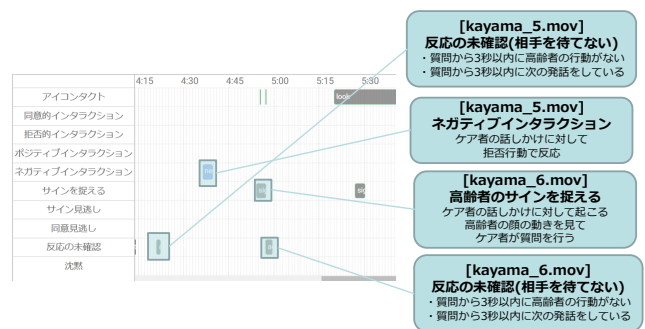


図 6. 映像事例に付与した解釈モデルの可視化例

表 10. 解釈モデルの評価実験の結果

	事例1		事例2		事例3		合計		正答率 (%)
	適当とされた解釈	不適当とされた解釈	適当とされた解釈	不適当とされた解釈	適当とされた解釈	不適当とされた解釈	適当とされた解釈	不適当とされた解釈	
agree			1		2		3	0	100
interaction					1		1	0	100
reject									
positive	2		1	1	1	1	4	2	66.67
interaction									
negative	1				1		2	0	100
interaction									
signCatch	2	1					2	1	66.67
agreeMiss					1		1	0	100
actionMiss	2	2	1				3	2	60

4.4 解釈モデルの考察

評価実験の結果として、不適当と解釈されたものについて、各事例についてのコメントを踏まえて考察する。

• positiveInteraction

解釈が不適当とされた 1 つ目の映像事例を図 7 に示す。看護師からは以下のような意見を得た。

- 被介護者の反応に対する評価をするべきだった。
「返事してくれましたね」などのポジティブな評価を

すべき。



図 7. positiveInteraction の解釈不適当な場面 1

解釈が不適当とされた 2 つ目の映像例を図 8 に示す。看護師からは以下のような意見を得た。

- 被介護者の同意を待たずに「準備しますね」と言っている気がする
- アイコンタクトがないなど、介護者の意思がうまく伝わらずにケアが進んでいる印象

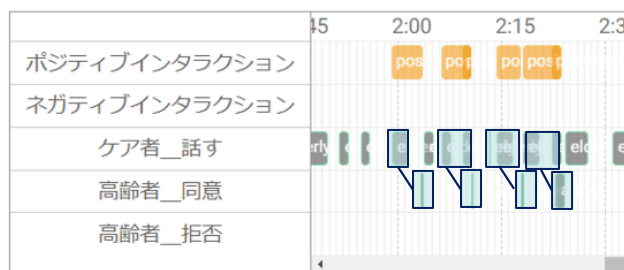


図 8. positiveInteraction の解釈不適当な場面 2

どちらの場面も、positiveInteraction は、介護者の話しかけに対して被介護者が同意を示していることから付与されている。しかし、得られたコメントとしてはポジティブではない、という意味合いの強いコメントであった。

1 つ目の事例では、被介護者が反応したことに対するポジティブな評価が行われていないことが指摘された。被介護者が介護者の問いかけに反応するだけでなく、反応に対して評価するまでがポジティブなインタラクション、と評価することが必要だと考えられる。また 2 つ目の事例では、同意を待っていない点とアイコンタクト等が行われないことから意思疎通が図られていない点を指摘された。この場面では、positiveInteraction の他に、同意見逃しや、反応の未確認も付与されていた。可視化したビューを図 9 に示す。ユマニチュードのケアにおいて負を示す解釈と競合した際には positiveInteraction を付与しないようなルール設計を行うことが必要であると考えられる。

[1]

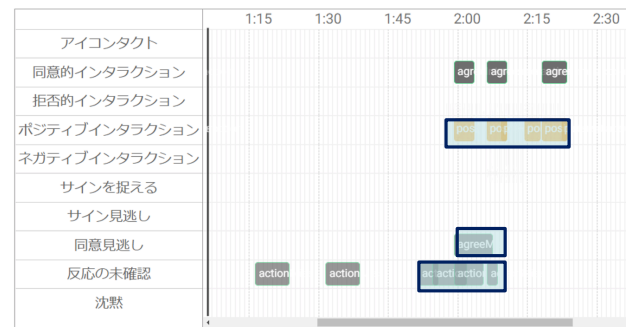


図 9. positiveInteraction の解釈不適当な場面 2 において付与された解釈一覧

- signCatch (サインを捉える)

解釈が不適当とされた 1 つ目の映像事例を図 10 に示す。看護師からは以下のような意見を得た。

- この場面で表出している口の動きは、歯磨きの刺激で動いているもの、この被介護者は歯ブラシを噛んでしまう癖がある

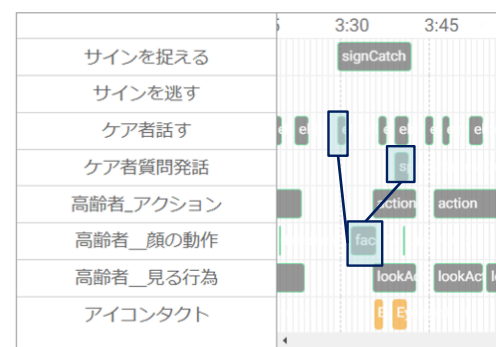


図 10. signMiss の解釈不適当な場面

この事例では、signCatch は、介護者の話しかけに対して被介護者が口を動かし、介護者が質問する流れを検知し付与されている。しかし、ここで表出した口の動きは歯ブラシが口に入った際に反射的に行われた動きであった。signCatch の設計の際、被介護者が出すサインとして非言語コミュニケーションである顔の動きとしたが、中でもさらに細分化して分けていく必要があると考えられる。

- actionMiss (反応の未確認・相手を待てない)

解釈が不適当とされた 1 つ目の映像事例を図 11 に示す。看護師からは以下のような意見を得た。

- 被介護者が頷いたことに対して介護者はしっかり答えている

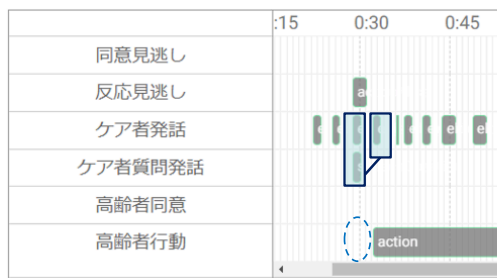


図 11. actionMiss の解釈不適当な場面 1

解釈が不適当とされた 2 つ目の映像事例を図 12 に示す。看護師からは以下のような意見を得た。

- 被介護者は、ケアに対して同意はしていないが、介護者は反応を待つことはしている

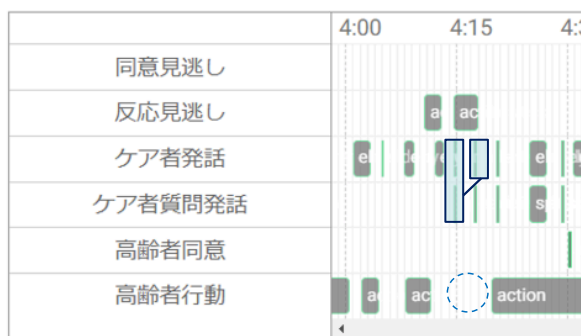


図 12. actionMiss の解釈不適当な場面 2

どちらの場面も、介護者の質問に対して被介護者の行動が見られず、次の発話が行われているため解釈が付与されている。

1 つ目の事例は、被介護者側は介護者に質問に対して頷いており、介護者も反応できている、という指摘を得た。この事例の冒頭においては、撮影を行う画角の関係上被介護者の様子が見て取れず、観察による顔の行為は確認できていなかった。そのため本事例のように間違った解釈が付与されている。2 つ目の事例では、介護者が待つことができているという指摘を得た。この場面において、介護者側は 3 度質問をしており、被介護者は 2 回目の質問のあとに行動を行っている。つまり得られたコメントとしては、3 回質問し返答を待っているか、という視点での指摘であり、今回付与されたものは 1 回のやりとりのみから判断し付与された解釈であることから、このような場面では時間軸を広げた解釈付与が必要であったと考えられる。

今回行った評価実験により、作成した解釈モデルに対するエキスパートの意見を得ることができ、作成した解釈モデルはエキスパートの意見と一致している部分があった一方で、一致しない部分もあった。今回の評価実験では病院の都合上 1 名の看護師のみの参加であったが、人により考え方は多様であり正解は存在しない。解釈が一致しない場合はデータとして蓄積し、何度も議論を重ねることで、解釈モデルをアップデートし続ける必要があると考えられる。

5. おわりに

認知症ケアにおける介入スキルを事例データに基づいて客観的に評価できる行動表現モデルを構築した。既存のモデルを被介護者本人の状態を記述できるものに拡張した結果、被介護者の非言語的な動作が客観化でき、介護者の観察スキルの詳細評価に活用できることが示唆された。また、介護者と被介護者のインタラクションの特徴をルールベースで抽出できる機能を設計開発し、高次の介入スキル分析に活用できる見通しを得た。

参考文献

- [1] 本田美和子, イヴ・ジネスト, ロゼット・マレスコッティ, ユマニチュード入門, 医学書院, 2014.
- [2] 石川翔吾, 竹林洋一, “エビデンスを生み出す認知症情報学-情動理解基盤技術とコミュニケーション支援,” 人工知能学会誌, 第 32, 第 1, pp. 103-110, 2017.
- [3] H. Miwako, I. Shogo, T. Yoichi and T. J. Tierney, "Reduction of Behavioral Symptoms of Dementia by Multimodal Comprehensive Care for Vulnerable Geriatric Patients in an Acute Care Hospital: A Case Series," *Case Reports in Medicine*, vol. 2016, no. Article ID 4813196, 2016.
- [4] I. Shogo, I. Mio, H. Miwako, T. Yoichi, “The skill representation of a multimodal communication care method for people with dementia,” *JJAP Conference Proceeding*, 第 4, 2016.
- [5] 石川翔吾, 菊池拓也, 本田美和子, G. Yves, 竹林洋一, “認知症ケア技法ユマニチュードにおけるコミュニケーションスキルの分析,” 第 4 回コモンセンス 知識と情動研究会, 第 11, 第 22, pp. 19-22, 2014.
- [6] 小川敬之, 竹田徳則, 認知症の作業療法 ソーシャルインクルージョンを目指して, 医歯薬出版株式会社, 2016.
- [7] A. Aung, S. Ishikawa, Y. Sakane, M. Ito, M. Honda, Y. Takebayashi, “A Visualization of Dementia Care Skills Based on Multimodal Communication Features,” *n 2016 AAAI Spring Symposium Series*, 2016.