

2 脳情報デコーディング技術

2 Brain Decoding Technology

2-1 視覚と認知をつかさどる脳機能の定量的理解とその応用に関する研究

2-1 Modeling and Decoding of Visual and Cognitive Brains

西本伸志 西田知史

Shinji NISHIMOTO and Satoshi NISHIDA

私たちの自然な日常生活は、複雑多様な情報を処理し、合目的な行動を生み出す高度な脳機能によって支えられている。近年、脳計測及び解読技術の発展に従い、動画視聴時等の自然で複雑な体験下における脳内情報を定量的に扱う研究が進展している。本稿では、脳情報の定量と解読に関する研究の枠組み及び言語モデルを利用した脳情報解読等、最近の研究成果について紹介する。

Our daily life is supported by the precise coordination of the brain functions that process massive, dynamic sensory inputs. Recent advancements in brain activity measurements and analysis techniques have enabled the quantitative modeling of representation and computation in the human brain under natural and complex experiences, such as during watching of movie clips. In this paper, we introduce our framework and recent research outcomes to model and decode human brain activity, in particular using the latent feature space derived from natural language processing techniques.

1 研究の背景と枠組み

1.1 はじめに

私たちは日常生活において視聴覚を代表とする感覚入力を介して多様な情報を受け取っており、また言語的・非言語的な手段を介して様々な情報を出力している。私たち—あるいはその情報処理を担う臓器である脳—は、このような様々な情報伝達の中核であり、情報通信の究極の受け手であり送り手である。脳がどのように情報を処理し、解釈し、また生成するかを知ることができれば、より効果的で豊かな情報伝達を実現するための基盤となり得る。

2010年代に入り、機能的磁気共鳴画像装置 (functional Magnetic Resonance Imaging: fMRI) を代表とする大規模・高精度な脳活動計測技術の発展、及び記録した脳活動を解析するための統計・機械学習手法の高度化に伴い、私たちが日常生活で触れるような多様で複雑な体験下 (例: 動画視聴時等) における脳活動を定量的に理解する研究が進展している [1][2]。従来の脳研究においては、心理学や生理学等に由来する特定の仮説を検証するためにデザインされた特殊な刺激、

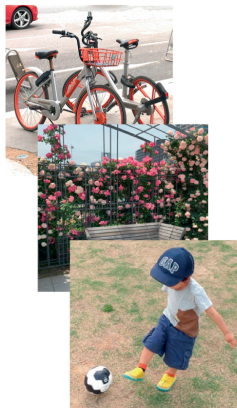
あるいは高度に統制されたタスクを用いた研究が主に進められてきた。これら従来型の研究は特定仮説の検証を行う強力な手段である一方で、得られた知見をより自然で複雑な条件下に一般化することは困難であった。上述の技術的進歩に伴って、複雑な体験下における脳活動を直接の研究対象とすることが可能になったことで、例えばCM動画を見ているときの印象内容を脳活動から定量評定する [3]、複雑な知覚内容を脳活動から文章として読み出す [4]、映像として想起した内容を解読することで意思伝達を行う [5]、等の実社会における応用、またそれらを見据えた多様な成果が生み出されつつある。

本稿では、自然で複雑な条件下における脳機能の定量理解を行うための枠組み及びその枠組みを活用した研究の具体例について紹介する。

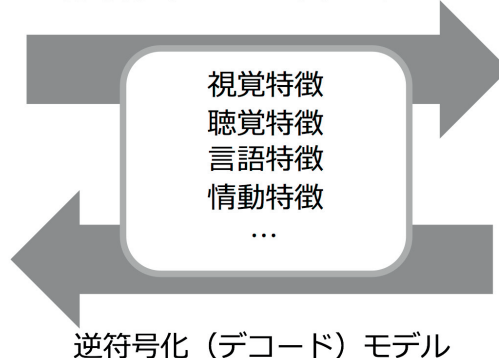
1.2 脳情報モデル構築と解読の枠組み

脳が多様な情報をどのように表現しているかを調べるため、私たちは可能な限り一般化された条件下における実験系を用いた研究を推進している。典型例としては、図1に示すような様々な動画を見ているときの

自然で多様な
知覚・認知体験
(例：動画視聴)



符号化（エンコード）モデル



逆符号化（デコード）モデル

脳神経活動
(例：fMRI計測)

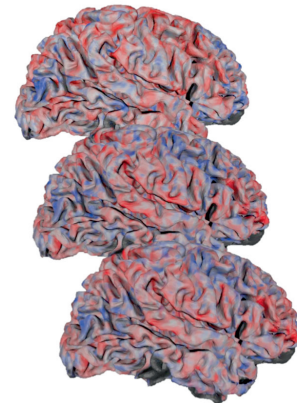


図1 脳情報モデル構築の概念図

被験者の脳活動を fMRI 等によって全脳にわたって連続的に撮像（記録）する。近年開発された Simultaneous multi-slice (SMS) 法等の高速撮像手法を用いると、例えば 2 mm 角の解像度で毎秒 1 ボリュームの全脳活動を 3 次元的に撮像することができる。このような記録を（休憩を入れながら）3 時間程度にわたって行うとすると、約 6～8 万の空間点（ヒト大人の全脳の場合）について約 1 万の時間サンプルといった大規模な時空間データを得ることができる。こうして得られた脳活動データは、そのまま観察する限りではほとんどノイズにしか見えない（図 1 右：赤青の色が活動強度を示す）ものであり、体験内容（動画そのもの、あるいはそれが誘起する知覚・認知内容）が何らかの形で反映された暗号のようなものと考えることができる。私たちの研究は、このようなデータからヒトの体験内容と脳活動の間の隠れた関係性を解読する（両者の関係を説明する定量的な予測モデルを構築する）ことにある。

体験内容と脳活動というペアとなった時空間データの関係性を調べるには、2 つの方向性を持ったモデルを考えることができる [6]。1 つは、体験内容のどのような特徴が脳活動として符号化されているかを調べる符号化（エンコード）モデルである（図 1）。これは、脳が情報をどのように表現しているかを直接的に検証する手段の 1 つである [1][2]。また、符号化モデルは脳と同じ働きをする（脳の振る舞いを予測する）言わば人工脳モデルであることから、同モデルの構築は脳に近い振る舞いを示すシステムを実現するための基盤技術としても注目を集めている [7]。もう 1 つは、特定の脳活動パターンが計測されたときにそれがどのような体験内容を意味するのかを推定する、逆符号化（デコード）モデルである。これは、脳活動がどのような情報を表現しているかを調べるもう 1 つの手法で

あると同時に、脳活動から被験者が何を感じて／考えていたかを推定する、言わば脳解読機ととらえることができる。このことから、逆符号化モデルはいわゆる脳・機械インターフェース (Brain Machine Interface: BMI) を実現するための数理的基盤としても注目を集めている。

上記の符号化・逆符号化モデルを構成するに当たり、体験内容のどのような側面に着目するかによって、多様な情報表現に関する検証を行うことが可能になる。これは、例えば同じ視覚入力があったとしても、私たちがそれを様々なレベル（映像、物体、印象、…）で記述することができることに対応する。例えば映像特徴としては、色や動き、テキスト等の記述があり得るが、これらは各種の画像特徴フィルタや畳み込みニューラルネットワーク (Convolutional Neural Network: CNN) における中間層の活動パターン等として定量化することができる。体験内容（映像）をこのような特徴が張る空間（特徴空間）上の点ととらえると、上記のモデル化は脳活動パターンが張る空間と特徴空間の間の投射関係を求めることに帰着できる。同様にして、「建物」「話す」「かわいい」といった言葉で表現できる意味内容や印象等については、自然言語処理由来の特徴空間（次項で詳述）を用いることで定量的に扱うことが可能となる。また、本稿では詳しくは扱わないが、同様の特徴空間を用いた枠組みは、視覚や聴覚だけでなく味覚や嗅覚、運動、また情動や言語、想起といったより高次の認知情報にも同様に適用することが可能である [1][2][6]。これらの様々な特徴空間やそれらの組合せを利用したモデルが任意の新規体験内容と脳活動の関係を説明（予測）することができるなら、そこで用いた特徴空間は脳内における情報表現を何らかの形で反映したものと推定することができる。ここ

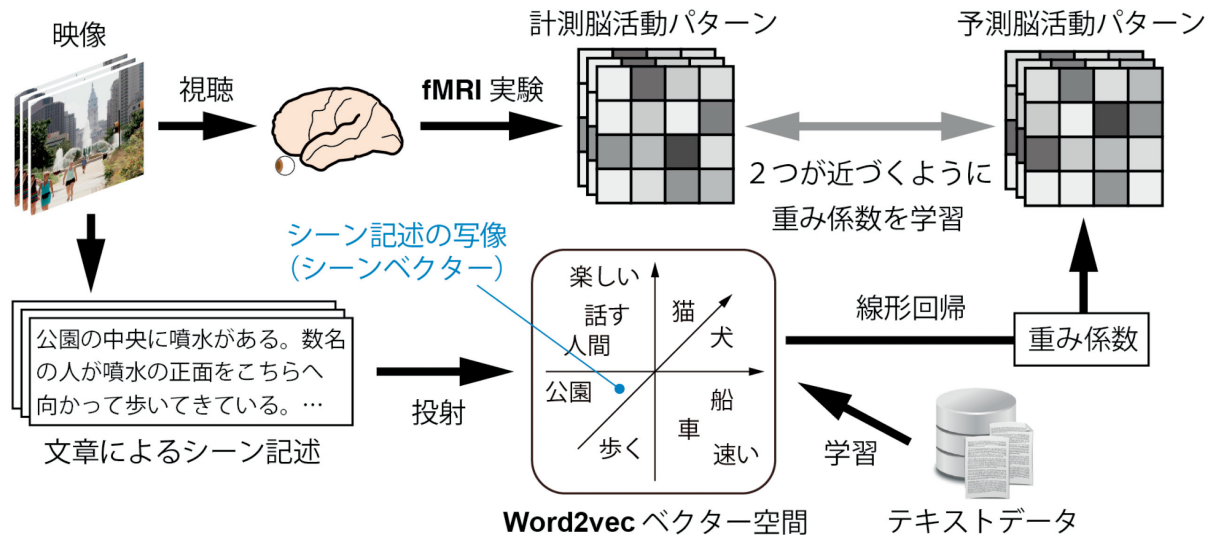


図2 Word2 vec ベクター空間を用いた符号化モデル

で、各々のモデルやモデル内の個別特徴が脳のどの部位における活動をどれだけ説明するかを解析することで、脳のどこでどのような情報が表現されているかを定量的に評価することができる。また脳内における特徴空間の構造を調べることで、脳が多様な体験内容をどのように構造化して表現しているかを推定することができる。さらには、逆符号化モデルを用いることで、脳活動から該当する特徴空間に関連した体験内容を推定することも可能になる。

次項では、これらの脳情報のモデル構築と解釈の枠組みを用いた具体的な研究成果について紹介する。

2 意味認知の定量的理解とその応用

2.1 脳内意味情報のモデル化

この世界にあふれる様々な情報に対し、ヒトは意味という単位でまとまりを見だし、言語を割り当てて理解する。意味及び言語により情報を表現することで、世界に存在する情報を体系的に解釈及び操作することが可能となる。ヒト脳内における意味情報の表現を理解することは、言語を扱う地球上唯一の生物であるヒトの知性を理解するためにひととき重要であろう。しかし、ヒト脳内の意味情報表現は、心理学や神経科学において長く研究されてきたが、いまだ不明な点が多い [8]。上述の符号化・逆符号化モデリングは、そのような脳内の意味情報表現を理解するうえで、非常に強力な方法論を提供する。

私たちは、脳内意味情報を定量化するため、言語特徴を用いた符号化・逆符号化モデリングに取り組んだ。言語特徴は、自然言語処理アルゴリズムの1つである word2 vec [9] により生成した。Word2 vec は Wikipedia コーパスのような大規模テキストデータから、単

語と単語の意味的関係性を適切に表現する多次元のベクター空間を学習する。この空間では数万から数十万個の単語がベクターとして表現され、単語同士の意味的な近さが、対応する単語ベクターの空間内の距離として表現される。また、単語間のアナロジー（類似）関係が適切に表現され、単語ベクターを用いて「王－男性＋女性＝女王」のような仮想的な加算・減算も行える。このような意味認知における私たちの直観と合致するベクター空間を言語特徴としてモデル化に用いることで、効果的に脳内意味情報を定量化できると考えた。

私たちは脳内意味表現を定量的に理解するために、word2 vec ベクター空間を用いた符号化モデルを構築した (図2)。まず、自然動画により誘発される大脳皮質の fMRI 応答を計測した。さらに、動画から言語特徴を抽出するため、各動画シーンに対して、複数名の記述者から文章によるシーン記述を取得した。シーン記述に含まれる単語を学習済みの word2 vec ベクター空間に投射して、各シーンに対応する word2 vec ベクター（以降、シーンベクターと呼ぶ）を算出し、それらの時系列を説明変数として fMRI 応答の時系列を予測するモデルを線形回帰により学習した。こうして学習したモデルを用いて、新規の動画に対応するシーンベクターから予測した fMRI 応答と、同じ動画が誘発した計測 fMRI 応答のピアソン相関により、各ボクセルの予測精度を評価した。その結果、視覚皮質をはじめとする広い皮質領域で高い予測精度を示すことが分かった [10]。また、同じく脳内意味表現の定量化を行った先行研究の符号化モデル [11] と比較しても、私たちのモデルの方がより高い予測精度を示したことから、word2 vec ベクター空間を用いた符号化モデルが脳内意味表現を上手く定量化することが分かった。

さらに私たちはこの符号化モデルを用いて、word2 vec ベクター空間自体から定量化した言語意味表現と、それをモデルにより変換した脳内意味表現を対比し、脳内意味表現の特性について分析を行った[12]。この分析では、word2 vec ベクター空間に含まれる 10 万単語を意味的な類似性に基づき 47 個の意味グループに分け、各グループの言語表現及び脳内表現を算出した。そして、特定の表現系における要素間の類似性・非類似性を定量化する階層クラスタリングと呼ばれる手法を適用して、言語表現及び脳内表現における意味グループ間の類似関係の構造を可視化した。図 3 は意味グループの脳内表現の構造を示す。図中のラベルが意味グループを表す。木構造では、低い枝で連結した近接する意味グループほど表現が近く、高い枝で連結した離れた意味グループほど表現が遠いことを表す。図中上部から下部に向けて、人間的・個人的・主観的な意味グループから、非人間的・社会的・客観的な意味グループへの連続的な変化が見られる。このような傾向は意味グループの言語表現では見られなかった。すなわち、言語意味表現と対比して、脳内意味表現では人間あるいは自己に関連する意味を基準とした、連続性をもつ意味表現の構造を有することが示唆された。このような自己を基準とする意味表現構造は、自己から外界・他者へと世界認識が広がっていくヒトの発達過程との対照を考えるうえでも興味深い。

私たちが各個人で定量化した脳内意味表現には、個人間の差異が存在する。この個人差は、世界を意味的に解釈するうえでの個性を反映すると考える。私たちは、今後このような脳内意味表現の個人差に着目し、個性の形成に関わる神経基盤を明らかにしたいと考えている。また、京都大学医学部精神科と共同で、意味認知の変容をもたらす精神疾患において、脳内意味表

現を定量化することによる疾患の理解と治療を目的とした取組も行っている。

2.2 意味認知内容の解読

世界からヒトが受容する情報は絶えず変化しており、ヒトが認知する意味情報は刻一刻と移り変わる。そのような意味認知内容を脳情報から解読できれば、ヒト-機械インターフェースや非言語コミュニケーションといった次世代の情報通信技術に応用できる。言語特徴を用いた逆符号化モデリングは、そのような意味認知内容の解読に利用可能な優れた技術基盤を提供する。

私たちは、脳活動から意味認知内容を言葉の形で解読することを目的として、word2 vec ベクター空間を用いた逆符号化モデルの構築を行った(図 4)[3]。このモデルでも同様に、自然動画が誘発する大脳皮質の fMRI 応答と、同じ動画に対するシーン記述を word2 vec ベクター空間に投射したシーンベクターを利用した。ただし、符号化モデルとは逆に、全皮質ボクセルの fMRI 応答の時系列を説明変数として、シーンベクターの時系列を予測するモデルを線形回帰により学習した。さらに、word2 vec ベクター空間において、予測したシーンベクターに対する 1 万語分の各単語ベクターの距離を計算し、距離が近いほど意味認知内容を強く反映する単語とみなして、単語による意味認知内容の推定を行った。

このモデルを用いて、新規の動画を視聴中の fMRI 応答から、単一のシーンに対して単語の形で意味認知内容の推定を行った結果を図 5 に示す。図中では、左のシーンを観察中の脳活動を解読し、1 万語の単語の中から意味認知内容を反映すると推定された上位 7 単語を、名詞・動詞・形容詞に分けて表示している。名詞は物体、動詞は行動、形容詞は印象の意味認知内容

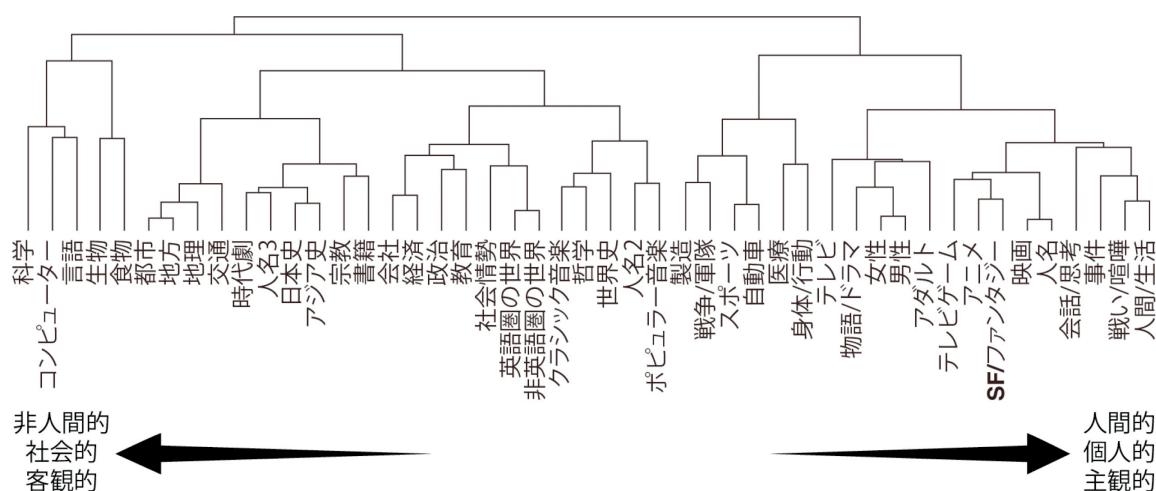


図 3 意味グループの脳内表現を表す木構造

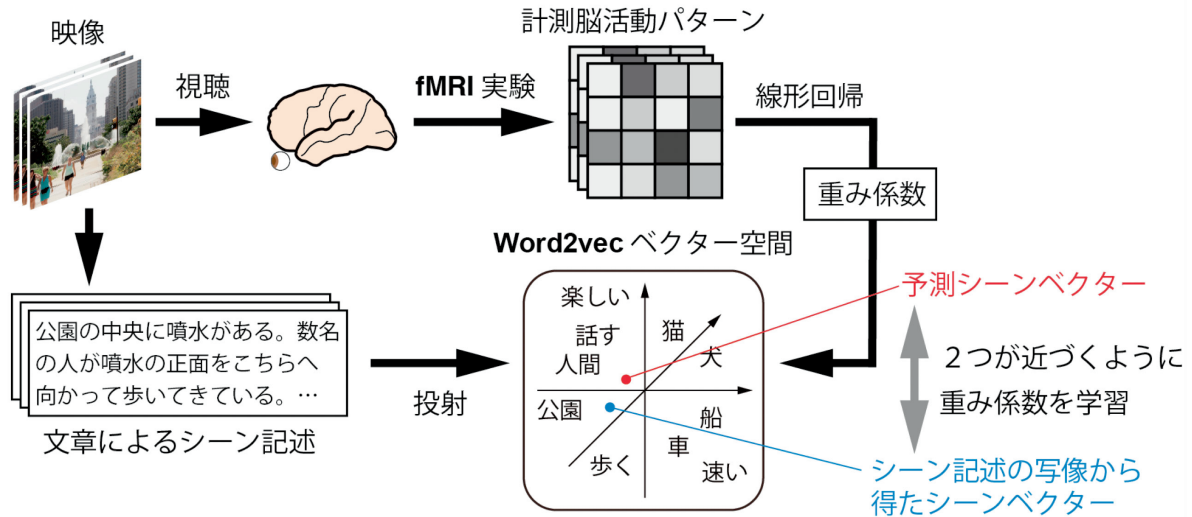


図4 word2vec ベクター空間を用いた逆符号化モデル

	名詞	動詞	形容詞
1	髪	着る	可愛い
2	髪型	喋る	親しい
3	金髪	気に入る	優しい
4	容姿	明かす	幼い
5	顔	慕う	怖い
6	女性	付き合う	欲しい
7	服	懂れる	面白い

もっと詳しい

図5 意味認知内容の推定結果例

にそれぞれ対応すると考えると、推定された単語がシーンを適切に説明していることが見て取れる。私たちは、このモデルが様々なシーンにおいて、シーン記述に合致する形で意味認知内容の推定ができることを統計的に示した。また、シーン記述の個人差が大きなシーンほど、脳活動からの推定結果も個人差が大きくなることを統計的に確認し、このモデルが意味認知の個人差を反映可能なことも示した [3]。

このモデルを用いた解読技術の利点は、1万単語という大きな語彙データを用いて意味認知内容を推定できるとともに、形容詞を含むことで印象内容も推定可能な点にある。類似した先行技術も存在するが、ここでは形容詞を含まない500単語未満の語彙データを用いているにすぎない [13]。実社会においてヒトが受容する意味情報は多種多様であり、私たちの解読技術の利点は、実社会における技術応用の可能性を広げている。実際、私たちはこの解読技術の実社会応用の1つとして、株式会社NTTデータと共同で、脳情報解読に基づく映像コンテンツの感性評価の事業化を2016年度に開始した（参考：NeM sweets DONUTs, <http://www.nem-sweets.com>）。この事業は、脳情報を用いた次世代ビジネスとして各界から注目されており、今後、更なる技術発展を目指し、研究開発を継続していきたい。

2.3 解読技術の拡張と脳融合型人工知能

私たちの解読技術の実社会応用を進める過程で、様々な障壁も明らかになってきた。その1つが、fMRIの計測コストの大きさである。MRI装置は非常に高価であり、維持費及び利用料も高額である。また、多種多様な映像の誘発する意味認知内容を評価するためには、その分の被験者実験を行う必要があり、人的及び時間的コストも甚大である。そこで私たちは、脳情報解読においてfMRIの計測コストを大幅に削減する新技術の開発を行った [14]。この技術が従来の解読技術と異なる点は、計測したfMRI応答の代わりに、符号化モデルにより感覚入力から予測したfMRI応答を逆符号化モデルに入力して、知覚・認知内容を解読する点である。これにより、最初に数時間分のfMRI応答データを用いて符号化及び逆符号化モデルの学習を行った後は、追加のfMRI計測を一切要さずに、任意の感覚入力から知覚・認知内容を推定できる。こうして構築したシステムは、感覚信号を入力、知覚・認知内容を出力とする一種の人工知能システムとして機能するため、脳情報を融合した新しい形の人工知能とみなせる。

図6は、この提案技術の実証実験で私たちが設計したシステムの概要である。このシステムでは、任意の映像入力からfMRI応答の予測を行う符号化モデルとして、最新の人工知能技術の一種である畳み込みニューラルネットワーク [15] の内部表現を用いたモデルを利用した。畳み込みニューラルネットワークは、視覚物体のカテゴリ判別において高い精度を示す深層学習モデルの1つである。その内部表現は、fMRI応答予測に応用して高い精度を示すことが知られている [16]。符号化モデルの学習では、映像のフレーム画像を畳み込みニューラルネットワークに入力し、その際に得られる内部表

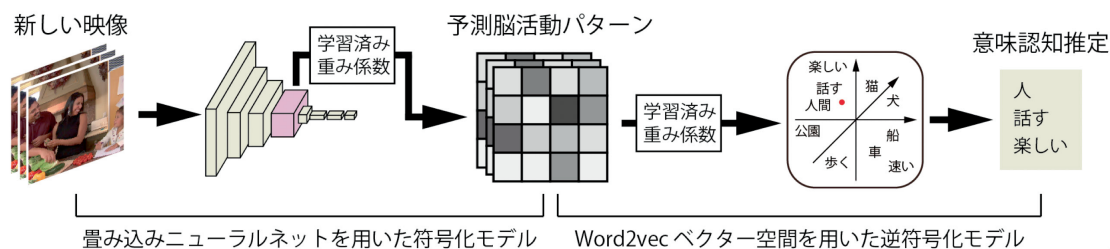


図6 追加のfMRI計測を必要としない意味認知解読システム

従来の解読技術による推定				提案システムによる推定			
	名詞	動詞	形容詞		名詞	動詞	形容詞
1	池	囲む	広い	1	池	咲く	広い
2	水路	広がる	古い	2	湖	囲む	古い
3	湖	造る	近い	3	山	広がる	狭い
4	海岸	面す	狭い	4	岩	降る	近い
5	岩	沿う	新しい	5	溪谷	造る	遠い
6	庭園	採る	遠い	6	滝	流れる	細長い
7	山	建てる	細長い	7	山頂	降りる	美しい

図7 意味認知推定における従来技術と提案システムの比較

現の時系列を説明変数として、fMRI 応答の時系列を予測するモデルを線形回帰により学習した。そして、その符号化モデルが新規の映像に対して予測したfMRI 応答を、word2 vec ベクター空間を用いた逆符号化モデルに入力し、意味認知内容の推定を行った。

図7は、提案システムによる意味認知内容の推定結果例である。比較のため、同じ映像シーンにおいて、計測fMRI 応答を直接入力したときの、word2 vec ベクター空間を用いた逆符号化モデルの推定結果も示している。図5と同様に、意味認知内容を反映する最上位の単語7個を名詞・動詞・形容詞に分けて表示している。いずれの方法を用いた場合も、シーンを説明する適切な単語が推定できていることが見て取れ、提案システムの推定が従来の解読技術と比べても遜色ないことが分かる。これを定量的に評価するために、私たちは10分間の新規な映像において、それに対応するシーン記述への合致度合いに基づき、提案システムと従来技術の推定精度の比較を行った。その結果、提案システムはむしろ従来技術と比べて統計的に有意に高い推定精度を示すことが分かった[14]。これは、符号化モデルによる予測を介することで、fMRI 応答に元々多く含まれる計測ノイズや、映像視聴に無関係な被験者の認知状態に基づくノイズが削減されたためだと考える。また、提案システムの重要な利点は、符号化・逆符号化モデルの学習を被験者個人ごとに行える点である。興味深いことに、従来の解読技術と同様に、映像に対するシーン記述の個人差が大きなシーンほど、

各個人で作成した提案システムの推定結果における個人差も大きくなる傾向が統計的に確認された[14]。このことは、提案システムを用いることで、映像に対する意味認知の個人差を、任意の映像に対して評価できる可能性を示唆している。

以上の実証実験では、符号化・逆符号化モデルに特定のものをを用いたが、私たちの提案技術は任意の符号化・逆符号化モデルに適用できる。これにより、視覚だけでなく聴覚などの別の感覚種を入力として扱える。また、意味認知だけでなく、多様な知覚・認知推定にも適用でき、個人の主観が強く反映する美醜判断や好き嫌いのような感性の推定を行う情報処理システムとして利用できる可能性がある。さらに、知覚・認知の個人差も反映し得るとすれば、既存の人工知能技術の弱点を大幅に補う、画期的な情報基盤技術として実社会に大きな影響を与えるだろう。今後は、様々な符号化・逆符号化モデル組み入れて提案技術の拡張を行いながら、産業応用にも取り組み、技術の開発と応用を進めていく予定である。

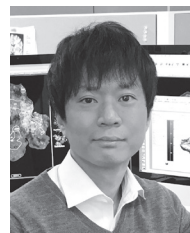
3 脳情報解読の展開と展望

本稿では、脳内情報について符号化・逆符号化モデルを用いて解明する研究の枠組み、またその具体例としての言語情報表現を介した脳情報のモデル化と解読及びその社会実装に関する取組等について紹介した。本研究の枠組みは、動画視聴時等の一般的な知覚・認

知条件下における脳機能の解明を目指すものである。このため、基礎から応用までのステップが比較的少なく、学術及び社会実装の両面において多様な展開を見据えた研究が可能である。例として、本稿では主にfMRIを用いた脳活動計測の成果について述べたが、同様の枠組みは侵襲的計測手法(手術等を必要とするが、より詳細な活動計測や刺激が行える手法)と組み合わせることが可能である。これにより、脳情報伝達の基本単位である神経細胞単位の情報表現の解明[17]や、頭蓋内埋め込み電極を用いたBMIの実現等の臨床応用を目指す研究[18]を並行して推進している。また眼球運動等の各種行動・生理指標と脳情報の関係性を調べる[19][20]ことで、物理的な感覚入力とヒトが感知する知覚内容の乖離に関する脳内メカニズムの解明等に関する研究も進めている。これらの試みは、将来的には簡易的な生理指標から脳情報を推定する技術等の応用につながる可能性がある。また多様な個性や趣味嗜好を脳機能面から理解するため、音楽や絵画等の芸術面に着目した脳情報表現を解明する研究も進めている[21]。さらには、本稿で紹介した符号化・逆符号化モデルは、深層学習や再帰的ネットワークといった近年進展が著しい人工知能技術との融合的活用が可能である[4][22]。今後も関連分野の進展を逐次取り入れることで、より高次の脳機能の定量理解、またその知見を利用した応用等、多様な展開を見据えた研究を推進する。

【参考文献】

- 1 J.L. Gallant, S. Nishimoto, T. Naselaris, and M.C.K. Wu, "System identification, encoding models and decoding models: a powerful new approach to fMRI research," In Kriegeskorte N and Kreiman G (eds.), Visual Population Codes (pp.163-188), Cambridge: MIT press., 2011.
- 2 S. Nishimoto, "エンコーディングモデルを用いた視覚情報処理研究: 情報表現, 予測, デコーディング" 日本神経回路学会誌, vol.19, pp.39-49, 2012.
- 3 S. Nishida, and S. Nishimoto, "Decoding naturalistic experiences from human brain activity via distributed representations of words," NeuroImage, 2017, <http://dx.doi.org/10.1016/j.neuroimage.2017.08.017>.
- 4 E. Matsuo, I. Kobayashi, S. Nishimoto, S. Nishida, and H. Asoh, "Generating natural language descriptions for semantic representations of human brain activity," Proceedings of the ACL 2016 Student Research Workshop, pp.22-29, 2016.
- 5 T. Naselaris, C.A. Olman, D.E. Stansbury, K. Ugurbil, and J.L. Gallant, "A voxel-wise encoding model for early visual areas decodes mental images of remembered scenes," Neuroimage, vol.105, pp.215-228, 2015.
- 6 T. Naselaris, K.N. Kay, S. Nishimoto, and J.L. Gallant, "Encoding and decoding in fMRI," Neuroimage, vol.56, no.2, pp.400-410, 2011.
- 7 R.C. Fong, W.J. Scheirer, and D.D. Cox, "Using human brain activity to guide machine learning," Scientific Reports, vol.8, no.1, Article 5397, 2018.
- 8 K. Patterson, P. J. Nestor, and T. T. Rogers, "Where do you know what you know? The representation of semantic knowledge in the human brain," Nature Reviews Neuroscience, vol.8, no.12, pp.976-987, 2007.
- 9 T. Mikolov, K. Chen, G. Corrado, and J. Dean, "Distributed representations of words and phrases and their compositionality," Advances in Neural Information Processing Systems, vol.26, pp.3111-3119, 2013.
- 10 S. Nishida, A. G. Huth, J. L. Gallant, and S. Nishimoto, "Word statistics in large-scale texts explain the human cortical semantic representation of objects, actions, and impressions," 45th Annual Meeting of the Society for Neuroscience, 333.13, 2015.
- 11 A. G. Huth, S. Nishimoto, A. T. Vu, and J. L. Gallant, "A continuous semantic space describes the representation of thousands of object and action categories across the human brain," Neuron, vol.76, no.6, pp.1210-1224, 2012.
- 12 S. Nishida, A. G. Huth, J. L. Gallant, and S. Nishimoto, "Voxelwise modeling with distributed word representations reveals the similarity and dissimilarity of semantic structures between a large-scale text corpus and the human brain," 第39回日本神経科学大会, I-139, 2016.
- 13 A. G. Huth, T. Lee, S. Nishimoto, N. Y. Bilenko, A. T. Vu, and J. L. Gallant, "Decoding the semantic content of natural movies from human brain activity," Frontiers in Systems Neuroscience, vol.10, Article 81, 2016.
- 14 S. Nishida and S. Nishimoto, "A brain-mediated computational model to estimate perceptual experiences evoked by arbitrary naturalistic visual scenes," Vision Science Society Annual Meeting 2018, 33.347, 2018.
- 15 K. Simonyan and A. Zisserman, "Very deep convolutional networks for large-scale image recognition," ICLR, 2015.
- 16 U. Güçl and M. A. J. van Gerven, "Deep neural networks reveal a gradient in the complexity of neural representations across the ventral stream," Journal of Neuroscience, vol.35, no.27, pp.10005-10014, 2015.
- 17 K. Ikezoe, M. Amano, S. Nishimoto, and I. Fujita, "Mapping stimulus feature selectivity in macaque V1 by two-photon Ca^{2+} imaging: Encoding-model analysis of fluorescence responses to natural movies," NeuroImage, 2018, <http://dx.doi.org/10.1016/j.neuroimage.2018.01.009>.
- 18 R. Fukuma et al., "Decoding visual stimulus in semantic space from electrocorticography signals," IEEE International Conference on Systems, Man, and Cybernetics (SMC), accepted, 2018.
- 19 S. Nishimoto, A.G. Huth, N.Y. Bilenko, and J.L. Gallant, "Eye movement-invariant representations in the human visual system," Journal of Vision, vol.17, no.1, Article 11, 2017.
- 20 H. Yamaguchi, S. Nishida, and S. Nishimoto, "An encoding model reveals spatiotemporal specificity of eye-related effects during natural vision," Cosyne 2018, III-35, 2018.
- 21 T. Nakai, N. Koide-Majima, and S. Nishimoto, "Encoding and decoding of music-genre representations in the human brain," IEEE International Conference on Systems, Man, and Cybernetics (SMC), accepted, 2018.
- 22 C. Fujiyama, I. Kobayashi, S. Nishimoto, S. Nishida, and H. Asoh, "A study on a correlation between a predictive model of motion pictures imitating the predictive coding of the cerebral cortex and brain activity," 第31回人工知能学会全国大会, 2 K3-OS-33 a-1 in2, 2017.



西本伸志 (にしもと しんじ)

脳情報通信融合研究センター
脳情報通信融合研究室
主任研究員
博士(理学)
神経生理学、認知神経科学、脳情報デコーディング



西田知史 (にしだ さとし)

脳情報通信融合研究センター
脳情報通信融合研究室
研究員
博士(医学)
認知神経科学、脳情報デコーディング、ニューロマーケティング