

動物の視知覚を探る

牛谷智一

千葉大学 大学院人文科学研究院

目に見えない内的な情報処理過程を客観的に探究するため、認知科学は、その過程についてのモデルを構成し、そのモデルの予測する反応が得られるかどうかでモデルの妥当性を評価する。複数のモデルが考えられる場合は、妥当性の高さを比較し、より高い方を採用することになる。この方法では内観報告が必要ないため、ヒト以外の動物の視覚情報処理過程、すなわち、彼らにとってこの世界がどのように見えているか、をヒトを対象とする研究同様の妥当性で解明できる。本発表では、オペラント条件づけを利用し、行動指標を用いてハトがどのように視覚情報を処理しているか調べた発表者自身の実験を紹介し、比較認知科学と呼ばれる領域の方法論と意義について考察する。

視覚情報処理過程の代表的なモデルに、Treisman & Gelade (1980) による特徴統合モデルがある。このモデルでは、視界にあるオブジェクトの様々な属性がまず同時並行でマッピングされ、必要に応じてそれら複数のマッピングを比較照合して標的となるオブジェクトが探索されると想定している。背景とは異なるテクスチャ領域の探索をハトに訓練した Cook (1992) の実験では、探索正答率のパターンがこの特徴統合モデルに合致していた。しかし、発表者が Treisman & Gelade (1980) により近い場面でハトをテストした実験で得られた反応時間のデータは、特徴統合モデルでは完全には説明できなかった。ヒトの視覚研究では、特徴統合モデルを修正した Wolfe (1989) の誘導探索モデルが提唱されている。3つの属性次元で標的刺激とディストラクタを構成した実験におけるハトの反応時間データは、ハトにおいても特徴統合モデルより誘導探索モデルの妥当性が高いことを示していた。

より高次の視覚情報処理過程として、欠損した輪郭や表面の情報を、実際に見えている情報を使って補う知覚的補間（補完）が知られている。知覚的補間は、ある物体が別の物体に一部隠蔽されているとき、その隠蔽された部分の輪郭や表面を補間して認識するアモーダル補間と、手前にある物体が背景と類似しているため見えにくくなっているとき、その物体と背景との間にある別の物体の輪郭情報を利用して、手前にある物体が実際にあるように知覚されるモーダル補間（主観的輪郭）とに大別される。モーダル補間には、アモーダル補間が伴うため、アモーダル補間を実現するメカニズムがモーダル補間を実現しているという説（同一説）がある。発表者は、ハトにおけるアモーダル補間とモーダル補間を調べた。ハトに一本につながった長方形と中央が分断された長方形の弁別を訓練し、テストでは中央が別の図形によって隠された長方形を呈示したところ、ハトはこのテスト長方形を、中央が分断された長方形であると報告した。この結果は、異なる結果となったチンパンジーやオマキザ

ルとは対照的に、ハトが別の物体に隠蔽された長方形の中央部分をアモーダルに補間していないことを示している。一方、ヒトではモーダル補間によって三角形や四角形だと知覚できる図形の弁別をハトに訓練したところ、ハトはこれを学習し、また新たなパターンへも弁別が転移したことから、ハトはモーダル補間することが示唆された。両実験の結果は、モーダル補間の実現に、アモーダル補間の過程は必ずしも必要ではないことを示しており、同一説を支持しない。この一連の研究は、動物実験による知見が、ヒトの認知モデル研究に貢献する例を示している。