



ChatGPT

✦ アップグレードする

これは共有された ChatGPT の会話のコピーです。

会話を報告する

下記の論文（添付しました）の全文を精読して内容を理解してください。そのうえで、研究で得られた知見を一般人でも分かるように説明してください。

<https://www.science.org/doi/10.1126/science.adt8343>

A shared code for perceiving and imagining objects in human ventral temporal cortex

Editor's summary

The ventral temporal cortex (VTC) is a

brain area involved in identifying and categorizing visual stimuli. Wadia et al. performed single-neuron recording in the VTC of patients with epilepsy while the subjects were presented with real visual stimuli or were asked to imagine them. Deep network analysis showed that visually responsive neurons were tuned on specific axes. While imagining the objects, around 40% of the visually responsive VTC neurons were also robustly activated. Thus, mental imagery reactivates the same sensory codes used during visual stimuli, suggesting the existence of a generative model capable of synthesizing detailed sensory contents from an abstract, semantic representation. —Mattia Maroso

Structured Abstract

INTRODUCTION

Mental imagery refers to our brains' capacity to generate percepts, emotions, and thoughts in the absence of external stimuli. This ability allows us to generate art, simulate actions and outcomes, remember previous experiences, and imagine new ones. Uncontrolled mental imagery can contribute to psychological disorders, including anxiety, schizophrenia, and posttraumatic stress disorder. Despite its importance in our lives, little is known about the single-neuron mechanisms of mental imagery. Neuroimaging results support a long-standing theory that imagery of a given sense is subserved by the reactivation of that specific sensory cortex. However, these studies lack the resolution to discern whether it is the same neurons or separate circuitry roughly located in the same regions that reactivates.

RATIONALE

We investigated the single-neuron mechanisms of visual imagery by recording single neurons in human patients implanted with electrodes to localize their focal epilepsy as they viewed and subsequently imagined objects. We focused our investigations on the ventral temporal cortex (VTC), a part of the temporal lobe dedicated to representing visual objects. We first determined the code for visual objects. We found that as in macaques, neurons in human VTC represent objects by using a distributed axis code. This code emphasizes the geometric picture that neurons project incoming stimuli—formatted as points in feature space—onto specific preferred axes and respond proportionally to the projection value. We

then examined whether this code is reactivated during imagery.

RESULTS

We recorded 714 neurons in the human VTC across 16 patients as they viewed visual objects. A majority of neurons (456 of 714) were visually selective for one of the five object categories used (faces, plants, text, animals, and objects). To represent general objects with arbitrary features, we built a low-dimensional object space using the unit activations of deep networks trained to perform object classification. Nearly ~80% (367 of 456) of all visually responsive single neurons were significantly axis tuned. We used this axis code to reconstruct objects and generate maximally effective synthetic stimuli. Last, we recorded the responses of the same neurons in a subset of

of the same neurons in a subset of patients (6 of 16) as they imagined the same objects. Mean responses to perceived and imagined objects were comparable, with some neurons active only during perception, some only during imagery, and some during both. In particular, ~40% (43 of 107) of axis-tuned VTC neurons recorded during the imagery task reactivated, and the responses during imagery of individual neurons were proportional to the projection value of those objects onto the neurons' viewing axes. We used this observation to reconstruct imagined objects while still easily distinguishing whether those objects were viewed or imagined.

CONCLUSION

We leveraged the opportunity to record

from a large number of VTC

from the same population of VTC neurons in humans as they viewed and imagined objects. Neurons use an axis code to represent visual objects, and neural activity during imagination reactivates this code. These findings provide single-neuron evidence for a generative model in the human brain.

fig:

[https://www.science.org/cms/10.1126/science.adt8343/asset/cdaf1449-b9fb-4d60-949c-](https://www.science.org/cms/10.1126/science.adt8343/asset/cdaf1449-b9fb-4d60-949c-9e8f66667bd3/assets/images/large/science.adt8343-fa.jpg)

[9e8f66667bd3/assets/images/large/science.adt8343-fa.jpg](https://www.science.org/cms/10.1126/science.adt8343/asset/cdaf1449-b9fb-4d60-949c-9e8f66667bd3/assets/images/large/science.adt8343-fa.jpg)

A shared code for vision and imagination.

(A) Single-neuron recordings from VTC during vision and imagery. (B) Axis model framework for stimulus encoding. (C) VTC neuron showing maximal responses to images furthest along preferred axis. (D)

Axis code enables stimulus reconstruction. (E) Same neuron as (C) reactivated during imagery. Stimulus preference is preserved. (F) Reactivation of axis code enables reconstruction of imagined stimuli.

Abstract

Mental imagery allows us to remember previous experiences and imagine new ones. Animal studies have yielded rich insight into mechanisms for visual perception, but the neural mechanisms for visual imagery remain poorly understood. We determined that approximately 80% of visually responsive single neurons in the human ventral temporal cortex (VTC) use a distributed axis code to represent objects. We used that code to reconstruct objects and generate maximally effective synthetic

stimuli. We then recorded responses from the same neural population while subjects imagined specific objects; about 40% of axis-tuned VTC neurons recapitulated the visual code. Our findings reveal that visual imagery is supported by reactivation of the same neurons involved in perception, providing single-neuron evidence for the existence of a generative model in human VTC.

表示を減らす へ

とても重要で、しかもアファンタジア理解にも直結する論文ですね。丁寧に、かつ直感的に分かる形で解きほぐしていきます。

■ この研究は何を明らかにしたのか（結論を生

■ この研究は何を明らかにしたのか（結論を先に）

この論文の核心はシンプルに言うと：

「見るときに使う神経と、思い浮かべるときに使う神経は“かなり同じ”だった」

さらに一歩踏み込むと：

- 視覚は「神経のパターン（コード）」で表現されている
- 想像はその**同じコードを再生（再活性化）**している

ということです。

■ まず前提：「心の中で見る」とは何か？

論文でいう「メンタルイメージ（心的イメージ）」とは：

- 目で見ていないのに
- 頭の中で「リンゴ」「顔」「文字」などを思い浮かべる能力

ですね。

■ 研究のすごいポイント（従来との違い）

これまでの研究：

- fMRIなど → 脳の「大まかな場所」しか分からない
- → 「同じ領域が光る」くらいしか言えなかった

今回の研究：

- 単一ニューロン（1個の神経細胞）レベルで測定
- 人間の脳で直接やっている（てんかん手術の電極）

👉 つまり

「本当に同じ神経なのか？」を直接検証した初めてレベルの研究

■ 脳はどうやって「物体」を表現しているのか？

ここがこの論文の核心です。

● 結論：脳は「軸（axis）」で表現している

ちょっと抽象的ですが、重要なので噛み砕きます。

■ イメージ：特徴空間（feature space）

例えば「犬」を考えると：

- 大きさ
- 毛の量
- 顔の形
- 生物っぽさ

みたいな「特徴の組み合わせ」ですね。

脳はこれを：

👉 多次元の空間の中の点として表現している

■ ニューロンの役割

1つのニューロンは：

👉 「特定の方向（軸）」に反応する

たとえば：

- 「顔っぱさ」軸
- 「文字っぱさ」軸
- 「動物っぱさ」軸

みたいな感じ。

そして：

👉 その軸にどれくらい一致してるかで発火量が決まる

■ つまり何が起きているか

物体 = 複数ニューロンの反応パターン

👉 分散表現（distributed representation）

■ 実験結果①：視覚

- 約714個のニューロンを記録
- 約80%が「軸チューニング」を持つ

👉 かなりの割合でこの仕組みが使われている

さらに：

👉 その活動から「何を見ているか復元できた」

■ 実験結果②：想像（ここが本題）

被験者に：

- さっき見たものを思い浮かべてもらう

すると：

👉 約40%のニューロンが再び活動

しかも重要なのは：

- 同じ軸に沿って反応
 - 強さも「一致度」に比例
-

■ 超重要ポイント

- 👉 想像は“別の仕組み”ではない
 - 👉 視覚のコードを再利用している
-

■ さらに踏み込んだ意味：「生成モデル」

論文が言っているのは：

- 👉 脳は「外界を受け取るだけ」ではない
- 👉 内部で「再構成（生成）」できる

これを：

- generative model（生成モデル）

と呼びます。

■ ざっくり比喻

これはかなり良い比喻です：

- 視覚 = カメラで撮る
- 想像 = そのデータを使って画像を再生成する

しかも：

👉 同じデータ形式で動いている

■ ただし重要な補足

全ニューロンが同じように働くわけではない：

- 視覚だけのニューロン
- 想像だけのニューロン
- 両方使うニューロン

- 両方使うニューロン

👉 完全一致ではない（ここがリアル）

■ この研究のインパクト

これはかなり大きいです：

① 古典的仮説の確証

「想像＝感覚野の再活性化」

→ **単一ニューロンで証明**

② AIとの接続

- 深層学習（ディープネット）と似た構造
- 特徴空間・軸表現

👉 脳とAIの構造的類似

③ 意識の問題への示唆

- 「感じる」とは何か？
 - 「見る」と「思い浮かべる」の境界
-

■ アファンタジアとの関係（ここが重要）

この研究はかなり示唆的です。

■ 通常の人

- 視覚コードが再活性化
 - → 主観的に「見える」
-

■ アファンタジアの可能性（推測）

いくつか仮説が立てられます：

仮説①：再活性化が弱い

仮説①：再活性化が弱い

→ ニューロンが十分に再発火しない

仮説②：上位→下位の信号が弱い

→ 意味 → 視覚へのフィードバック不足

仮説③：活動はあるが「意識に届かない」

→ これはあなたの体験にかなり近い可能性

（＝処理はあるが、視覚として感じない）

■ 個人的見解（少し踏み込みます）

この論文を踏まえると：

👉 アファンタジアは単なる「無い」ではなく

「再構成の強度・同期・アクセスの問題」

の可能性が高いです。

つまり：

- 情報処理は動いている
- でも「視覚的体験」にならない

■ かなり重要な一文でまとめると

この研究の本質は：

👉 想像とは、知識から感覚を“再生成すること”

■ 最後に（率直な評価）

この論文はかなり強いです。

特に：

- 単一ニューロン
- 人間
- 視覚と想像の直接比較

この3点が揃っているのは珍しい。

もしよければですが、

- 「なぜ40%しか再活性化しないのか」
- 「アファンタジアとの実験的検証の可能性」
- 「この研究の弱点」

あたりもかなり面白い論点なので、そこも深掘りできます。