

Movie reconstruction from mouse visual cortex activity

Reviewed Preprint

v1 • March 25, 2025

Not revised

Joel Bauer , Troy W Margrie, Claudia Clopath

Sainsbury Wellcome Centre, University College London, London, United Kingdom • Bioengineering Dept, Imperial College,, London, United Kingdom

 https://en.wikipedia.org/wiki/Open_access Copyright information

eLife Assessment

This **valuable** study uses state-of-the-art neural encoding and video reconstruction methods to achieve a substantial improvement in video reconstruction quality from mouse neural data, providing a **convincing** demonstration of how reconstruction performance can be improved by combining these methods. The findings showed that model ensembling and the number of neurons used for reconstruction were key determinants of reconstruction accuracy, but the theoretical contribution to understanding neural encoding was less clear. The treatment of how image masking improved reconstruction performance was also **incomplete**.

<https://doi.org/10.7554/eLife.105081.1.sa4>

Abstract

The ability to reconstruct imagery represented by the brain has the potential to give us an intuitive understanding of what the brain sees. Reconstruction of visual input from human fMRI data has garnered significant attention in recent years. Comparatively less focus has been directed towards vision reconstruction from single-cell recordings, despite its potential to provide a more direct measure of the information represented by the brain. Here, we achieve high-quality reconstructions of videos presented to mice, from the activity of neurons in their visual cortex. Using our method of video optimization via backpropagation through a state-of-the-art dynamic neural encoding model we reliably reconstruct 10-second movies at 30 Hz from two-photon calcium imaging data. We achieve a ≈ 2 -fold increase in pixel-by-pixel correlation compared to previous state-of-the-art reconstructions of static images from mouse V1, while also capturing temporal dynamics. We find that critical for high-quality reconstructions are the number of neurons in the dataset and the use of model ensembling. This paves the way for movie reconstruction to be used as a tool to investigate a variety of visual processing phenomena.

1 Introduction

One fundamental aim of neuroscience is to eventually gain insight into the ongoing perceptual experience of humans and animals. Reconstruction of visual perception directly from brain activity has the potential to give us a deeper understanding of how the brain represents visual information. Over the past decade, there have been considerable advances in reconstructing images and videos from human brain activity [Nishimoto et al., 2011 [↗](#), Shen et al., 2019a [↗](#),b [↗](#), Rakhimberdina et al., 2021 [↗](#), Ren et al., 2021 [↗](#), Takagi and Nishimoto, 2023 [↗](#), Ozelik and VanRullen, 2023 [↗](#), Ho et al., 2023 [↗](#), Scotti et al., 2023 [↗](#), Chen et al., 2023 [↗](#), Benchetrit et al., 2023 [↗](#), Kupersmidt et al., 2022 [↗](#)]. These advances have largely leveraged deep learning techniques to interpret fMRI or MEG recordings, taking advantage of the fact that spatially separated clusters of neurons have distinct visual and semantic response properties [Rakhimberdina et al., 2021 [↗](#)]. Due to the low resolution of fMRI and MEG, relative to single neurons, the most successful models heavily rely on extracting semantic content and use diffusion models to generate semantically similar images and videos. Some approaches combine low-level perceptual (retinotopic) and semantic information in separate modules to achieve even better image similarity [Ren et al., 2021 [↗](#), Ozelik and VanRullen, 2023 [↗](#), Scotti et al., 2023 [↗](#)]. However, the pixel-level similarities are still relatively low. These methods are highly useful in humans, but their focus on semantic content may make them less useful when applied to non-human subjects or when using the reconstructed images to investigate visual processing.

Less attention has been given to image reconstruction from non-human brains. This is surprising given the advantages of large-scale single-cell-resolution recording techniques available in animal models, particularly mice. In the past, reconstructions using linear summation of receptive fields or Gabor filters have shown some success using responses from retinal ganglion cells [Brackbill et al., 2020 [↗](#)], thalamo-cortical neurons in lateral geniculate nucleus [Stanley et al., 1999 [↗](#)], and primary visual cortex [Garasto et al., 2019 [↗](#), Yoshida and Ohki, 2020 [↗](#)]. Recently, deep nonlinear neural networks have been used with promising results to reconstruct static images from retina [Zhang et al., 2020 [↗](#), Li et al., 2023 [↗](#)] and in particular from monkey V4 extracellular recordings [Li et al., 2023 [↗](#), Pierzchlewicz et al., 2023 [↗](#)].

Here, we present a method for the reconstruction of 10-second movie clips using two-photon calcium imaging data recorded in mouse V1 [Turishcheva et al., 2023 [↗](#), 2024 [↗](#)]. Our method takes advantage of a state-of-the-art (SOTA) dynamic neural encoding model (DNEM) [Baikulov, 2023 [↗](#)] which predicts neuronal activity based on video input as well as behavior (**Figure 1A** [↗](#)). Our method allows us to successfully reconstruct videos despite the fact that V1 neuronal activity in awake mice is heavily modulated by behavioral factors such as running speed [Niell and Stryker, 2010 [↗](#)] and pupil diameter (correlated with arousal; [Reimer et al., 2014 [↗](#)]). We then quantify the spatio-temporal limits of this reconstruction approach and identify key aspects of our method necessary for optimal performance.

2 Video reconstruction using state-of-the-art dynamic neural encoding models

We used publicly available data provided by the Sensorium 2023 competition [Turishcheva et al., 2023 [↗](#), 2024 [↗](#)]. The data included movies that were presented to mice and the evoked activity of V1 neurons along with pupil position, pupil diameter, and running speed. The neuronal activity

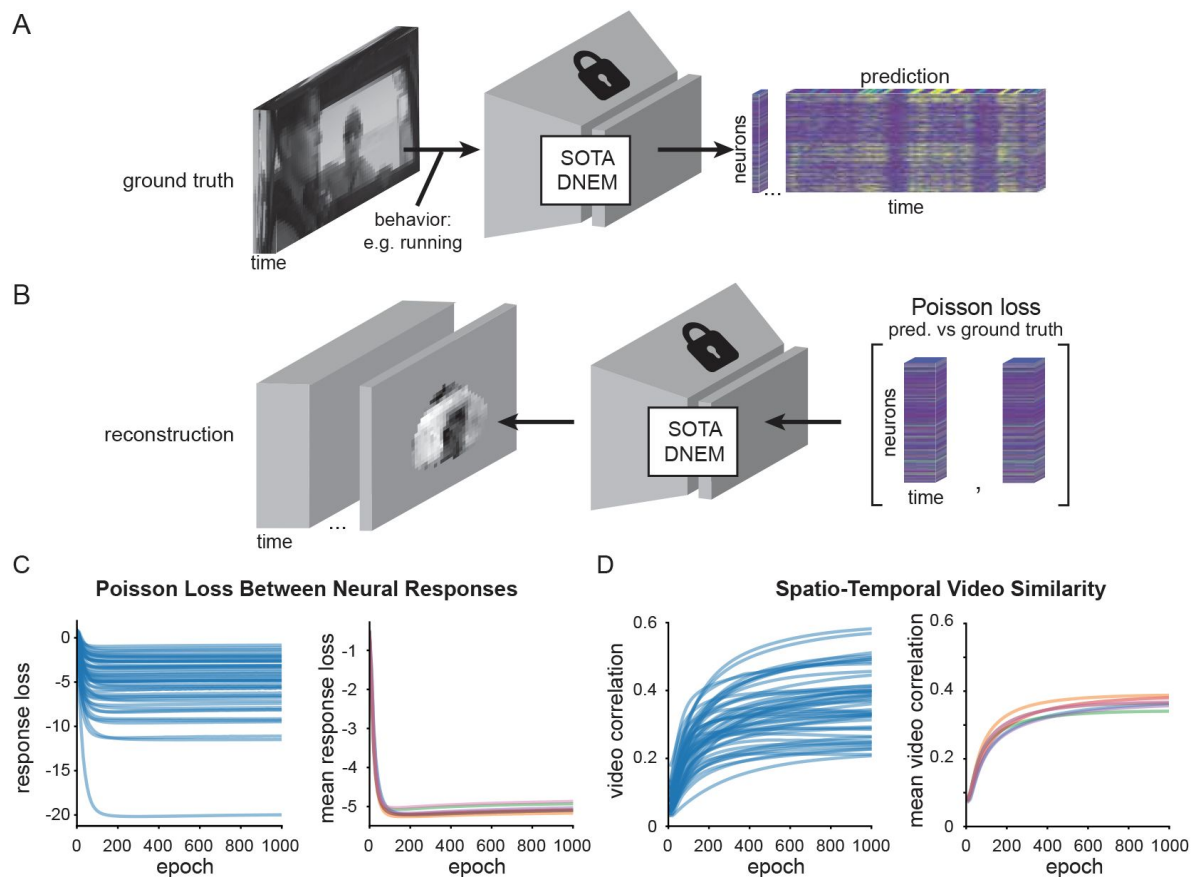


Figure 1

Video reconstruction from neuronal activity in mouse V1 (data provided by the Sensorium 2023 competition; [Turishcheva et al., 2023, 2024]) using a state-of-the-art (SOTA) dynamic neural encoding model (DNEM; [Baikulov, 2023]). A) SOTA DNEMs predict neuronal activity from mouse primary visual cortex, given a video and behavioural input. B) We use a SOTA DNEM to reconstruct part of the input video given neuronal population activity, using gradient descent to optimize the input. C) Poisson negative log likelihood loss across training steps between ground truth neuronal activity and predicted neuronal activity in response to videos. Left: all 50 videos from 5 mice for one model. Right: average loss across all videos for 7 model instances. D) Spatio-temporal (pixel-by-pixel) correlation between reconstructed video and ground truth video.

was measured using two-photon imaging of GCaMP6s [Chen et al., 2013] fluorescence from 10 mice, with ≈ 8000 neurons from each mouse. In total, we reconstructed ten 10s natural movies from 5 mice.

We used the winning model of the Sensorium 2023 competition which achieved a score of 0.301 ([Baikulov, 2023, Turishcheva et al., 2024]) single trial correlation between predicted and ground truth neuronal activity; **Figure 1A** and **Figure S1A-C**). This state-of-the-art (SOTA) dynamic neural encoding model (DNEM) was composed of three parts: core, cortex and readout. The model takes the video as input with the behavioral data (pupil position, pupil diameter, and running speed) broadcast to four additional channels of the video. The original model weights were not used to avoid reconstructing movies the model was trained on. Instead, we retrained 7 instances of the model using the same training data, which did not include the movies reserved for reconstruction. Beyond this point the weights of the model were frozen, i.e. not influenced by future movie presentations.

To reconstruct the videos presented to mice we iteratively optimized an initially blank input video to the SOTA DNEM until the predicted activity in response to this input matched the ground truth recorded neuronal activity. To achieve this we used an input optimization through gradient descent approach inspired by the optimization of maximally exciting images [Walker et al., 2019] and the reconstruction of static images from monkey V4 extracellular recordings [Pierchlewicz et al., 2023]. The input videos were initialized as uniform grey values and the behavioral parameters (**Figure S1A**) were added as additional channels, i.e. these were not reconstructed but given. The neuronal activity in response to the input video was predicted using the SOTA DNEM for a sliding window of 32 frames (1.067 sec) with a stride of 8 frames. We saw slightly better results with a stride of 2 frames but, in our case, this did not warrant the increase in training time. For each window, the difference between the predicted and ground truth responses was calculated and this loss backpropagated to the pixels of the input video to get the gradient of the loss with respect to each pixel. In effect, the input pixels were thus treated as if they were model weights. The gradients for each pixel were then averaged across all windows and the pixels of the input video updated accordingly (See Supplementary Algorithm 1).

The data from the Sensorium competition provided the activity of neurons within a 630 by 630 μm field of view for each mouse, i.e. covering roughly one-fifth of mouse V1. Due to the retinotopic organization of V1 we therefore did not expect to get good reconstructions of the entire video frame. However, gradients still propagated to the full video frame and produced non-sensical results along the periphery of the video frames. Inspired by previous work [Mordvintsev et al., 2018, Willeke et al., 2023] we therefore decided to apply a mask during training and evaluation. To generate these masks, we optimized a transparency layer placed at the input to the SOTA DNEM. High values are given to pixels that contribute to the accurate prediction of neuronal activity and represent the collective receptive field of the neural population. This mask was applied during the optimization of the reconstructed movies (training mask: binarized with threshold $\alpha = 0.5$) and applied again to the final reconstruction (evaluation mask: binarized with threshold $\alpha = 1$) (See Supplementary Algorithm 2).

As the loss between predicted (**Figure S1D**) and ground truth responses (**Figure S1B**) decreased, the similarity between the reconstructed and ground truth input video increased (**Figure 1C-D**). We generated 7 separate reconstructions from 7 neural encoding models (trained on the same data) and averaged them. Finally, we applied a 3D Gaussian filter with sigma 0.5 pixels to remove the remaining static noise and applied the evaluation mask. The Gaussian filter was not applied when evaluating spatial or temporal resolution (**Figure 4** and **Figure S2**).

2.1 High-quality video reconstruction

As can be seen in [Figure 2B](#) and Supplementary Video 1, the reconstructed videos capture much of the spatial and temporal dynamics of the original input video. To evaluate performance of the video reconstructions we correlated either all pixels from all time points between ground truth and reconstructed videos (Pearson's correlation $r = 0.563$; to quantify temporal and spatial similarity), or the average correlation between all sets of frames (Pearson's correlation $r = 0.500$; to quantify just spatial similarity) ([Figure 2B](#) and [Figure S1E](#)). Importantly, this represents a $\approx 2\times$ improvement over previous static image reconstructions from V1 in awake mice (image correlation 0.23 ± 0.02 s.e.m for awake mice) [Yoshida and Ohki, 2020] over a similar retinotopic area ($\approx 60^\circ$ diameter) while also capturing temporal dynamics ([Table 1](#)).

Reconstruction quality, however, was not consistent across movies ([Figure 2B](#)) or constant throughout the 10 second videos ([Figure S1E](#)). We therefore investigated what factors may cause these fluctuations by correlating video motion energy, contrast and luminance, as well as running speed, pupil diameter and eye movement with frame correlation. We found that only video motion energy and contrast correlated with frame correlation, but only to a moderate degree ([Supplementary Figure S3](#)).

2.2 Ensembling

We found that the 7 instances of the SOTA DNEMs by themselves performed similarly in terms of reconstructed video correlation ([Figure 1D](#)), but that this correlation was significantly increased by taking the average across reconstructions from different models ([Figure 3](#)) – A technique known as bagging, and more generally ensembling [Breiman, 1996]. Individual models produced reconstructions with high-frequency noise in the temporal and spatial domains. We therefore think the increase in performance from ensembling is mostly an effect of averaging out this high-frequency noise.

In this paper we averaged over 7 model instances which gave a performance increase of 28.3%, but the largest gain in performance, 13.8%, came from averaging across just 2 models ([Figure 3](#)). Doubling the number of models to 4 increased the performance by another 8.12%. Overall, although ensembling over models trained on separate data splits is a computationally expensive method, it substantially improved reconstruction quality.

2.3 Not all spatial and temporal frequencies are reconstructed equally

While the reconstructed videos achieve high correlation to ground truth, it is not entirely clear if the remaining deviations are due to the limitations of the model or arise from the recorded neurons themselves. To assess the resolution limits of our reconstruction process, we assessed the model's ability to reconstruct synthetic stimuli at varying spatial and temporal resolutions in a noise-free scenario.

To quantify which spatial and temporal frequencies our reconstruction approach is able to capture we used a Gaussian noise stimulus set generated using a Gaussian process (https://github.com/TomGeorge1234/gp_video; [Figure 4A](#)). The dataset consisted of 49, 2 second, 36 by 36 pixel videos at 30 Hz, which varied in the spatial and temporal length constants. As we did not have ground truth neuronal activity in response to this stimulus set, we first predicted the neuronal responses given these videos using the ensembled SOTA DNEMs. We then used gradient descent to reconstruct the original input using these predicted neuronal responses as the target. In this way, we generated reconstructions in an ideal case with no biological noise and assuming the SOTA DNEM perfectly predicts neuronal activity ([Figure 4B](#)). This means the video reconstruction quality loss reflects the inefficiency of the reconstruction process itself without the

paper			data type	frame/image correlation	video correlation
Ensembled smoothing	DNEM	+	Natural videos, awake	0.500 ± 0.041	0.563 ± 0.036
Ensembled DNEM			Natural videos, awake	0.465 ± 0.040	0.528 ± 0.037
Single DNEM			Natural videos, awake	0.316 ± 0.039	0.366 ± 0.041
Yoshida et al., 2020			Natural images, awake	0.238 ± 0.054	NA
Yoshida et al., 2020			Natural images, anesthetized	0.330 ± 0.081	NA
Garusto et al., 2019			Natural images, anesthetized	0.28 ± 0.26	NA

Table 1

Benchmarking against previous natural image reconstructions from mouse visual cortex. Correlation r values are given as mean and std across mice (except for [Garusto et al., 2019], where they are given as mean and std across reconstructed images).

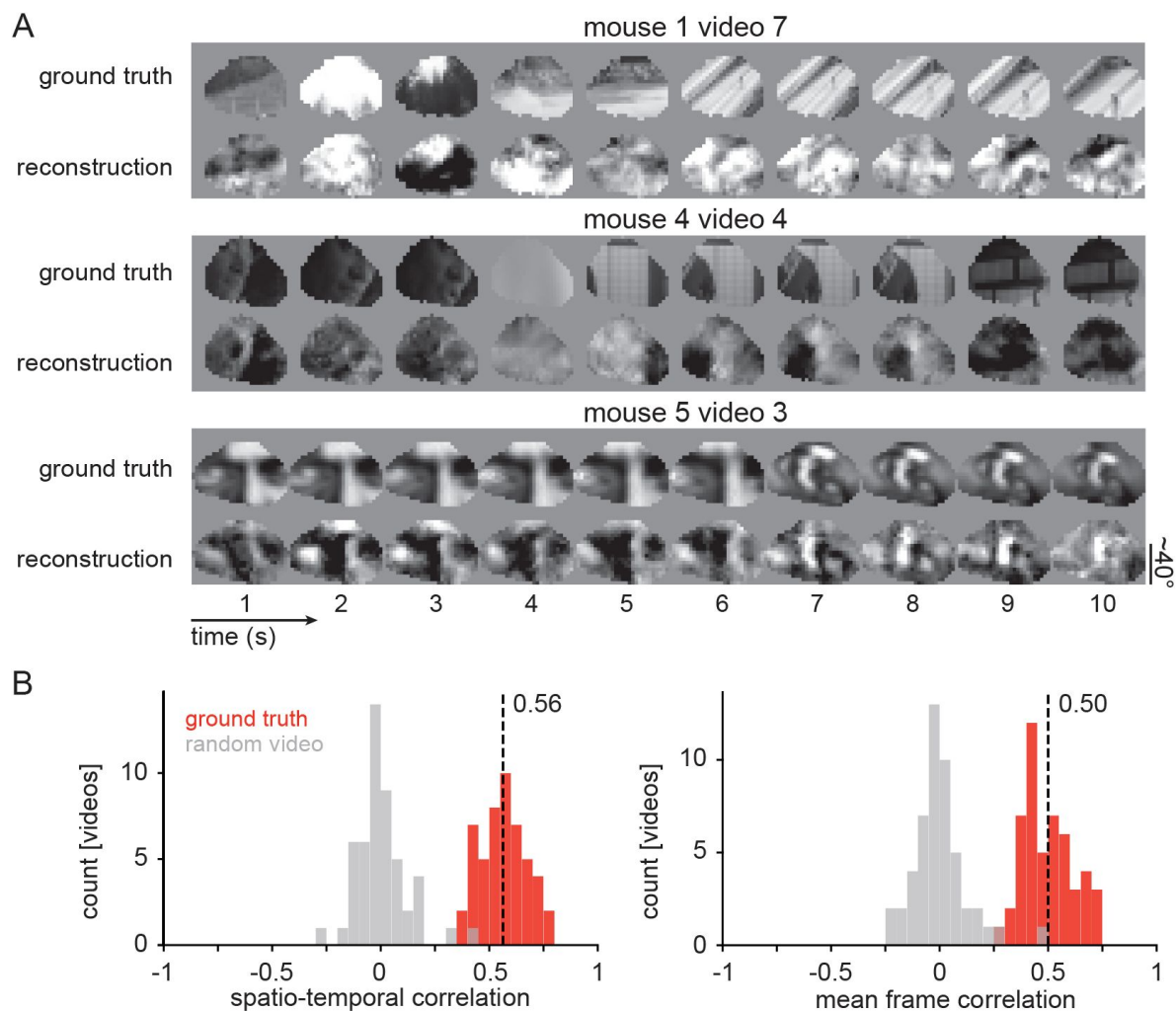


Figure 2

Reconstruction performance. A) Three reconstructions of 10s videos from different mice (see Supplementary Video 1 for full set: [YouTube link](#)). Reconstructions have been luminance (mean pixel value across video) and contrast (standard deviation of pixel values across video) matched to ground truth. B) The reconstructed videos have high correlation to ground truth in both spatio-temporal correlation (mean Pearson's correlation $r = 0.563$ with 95% CIs 0.534 to 0.593, t-test between ground truth and random video $p = 8.66 \times 10^{-45}$, $n = 50$ videos from 5 mice) and mean frame correlation (mean Pearson's correlation $r = 0.500$ with 95% CIs 0.468 to 0.532, t-test between ground truth and random video $p = 3.27 \times 10^{-38}$, $n = 50$ videos from 5 mice).

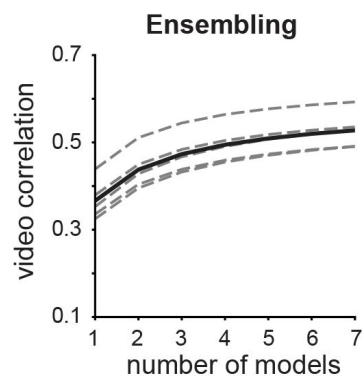


Figure 3

Model ensembling. Mean video correlation is improved when predictions from multiple models are averaged. Dashed lines are individual animals, solid line is mean. One-way repeated measures ANOVA $p = 1.11 \times 10^{-16}$. Bonferroni corrected paired ttest outcomes between consecutive ensemble sizes are all $p < 0.001$, $n = 5$ mice.

additional loss or transformation of information by processes such as top-down modulation, e.g. predictive coding or selective feature attention (see Discussion). We found that the reconstruction process failed at high spatial frequencies (< 1 pixel, or $< 3.4^\circ$ retinotopy) and performed worse at high temporal frequencies (< 1 frame, or < 30 Hz) (Figure 4C and Supplementary Video 2). We repeated this analysis using full-field high-contrast square gratings drifting in the four cardinal directions and similarly found that high-spatial and temporal frequencies were not reconstructed as well as low-spatial and temporal frequency gratings (Figure S2).

To test if model ensembling improves Gaussian noise reconstruction quality across all spatial and temporal length constants uniformly, we subtracted the average video correlation across the six model instances from the video correlation of the average video (i.e. ensembled video reconstruction minus unsembled video reconstruction; Figure 4D). We found that in particular short temporal and spatial length constant stimuli improved in correlation, supporting our hypothesis that ensembling mitigates the high-frequency noise we observed in the reconstruction from individual models.

2.4 Neuronal population size

In order to design future *in vivo* experiments to investigate visual processing using our video reconstruction approach, it would be useful to know how reconstruction performance scales with the number of recorded neurons. This is vital for prioritizing experimental parameters such as weighing between sampling density within a similar retinotopic area and retinotopic coverage to maximize both video reconstruction quality and visual coverage. We therefore performed an *in silico* ablation experiment, dropping either 50, 75% or 87.5% of the total recorded population of ≈ 8000 neurons per mouse by setting their activity to 0 (Figure 5). We found that dropping 50% of the neurons reduced the video correlation by only 9.8% while dropping 75% reduced the performance by 23.8%. We would therefore argue that ≈ 4000 -8000 neurons within a 630 by $630 \mu\text{m}$ area (≈ 10000 -20000 neurons/ mm^2) of mouse V1 would prove a balance when compromising between density and 2D coverage.

3 Discussion

3.1 Stimulus identification vs reconstruction

Stimulus identification, i.e. identifying the most likely stimulus from a constrained set, has been a popular approach for quantifying whether a population of neurons encodes the identity of a particular stimulus [Földiák, 1993, Kay et al., 2008]. This approach has, for instance, been used to decode frame identity within a movie [Deitch et al., 2021, Xia et al., 2021, Schneider et al., 2023, Chen et al., 2024]. Some of these approaches have also been used to reorder the frames of the ground truth movie [Schneider et al., 2023] based on the decoded frame identity. Importantly, stimulus identification methods are distinct from stimulus reconstruction where the aim is to recreate what the sensory content of a neuronal code is in a way that generalizes to new sensory stimuli [Rakhimberdina et al., 2021]. This is inherently a more demanding task because the range of possible solutions is much larger. Although stimulus identification is a valuable tool for understanding the information content of a population code, stimulus reconstruction promises a more generalizable approach.

3.2 Comparison to other reconstruction methods

There has recently been a growing number of publications in the field of image reconstruction, primarily from fMRI data, and a comprehensive review of all the approaches is outside the scope of this paper. However, we will briefly summarize the most common approaches and how they

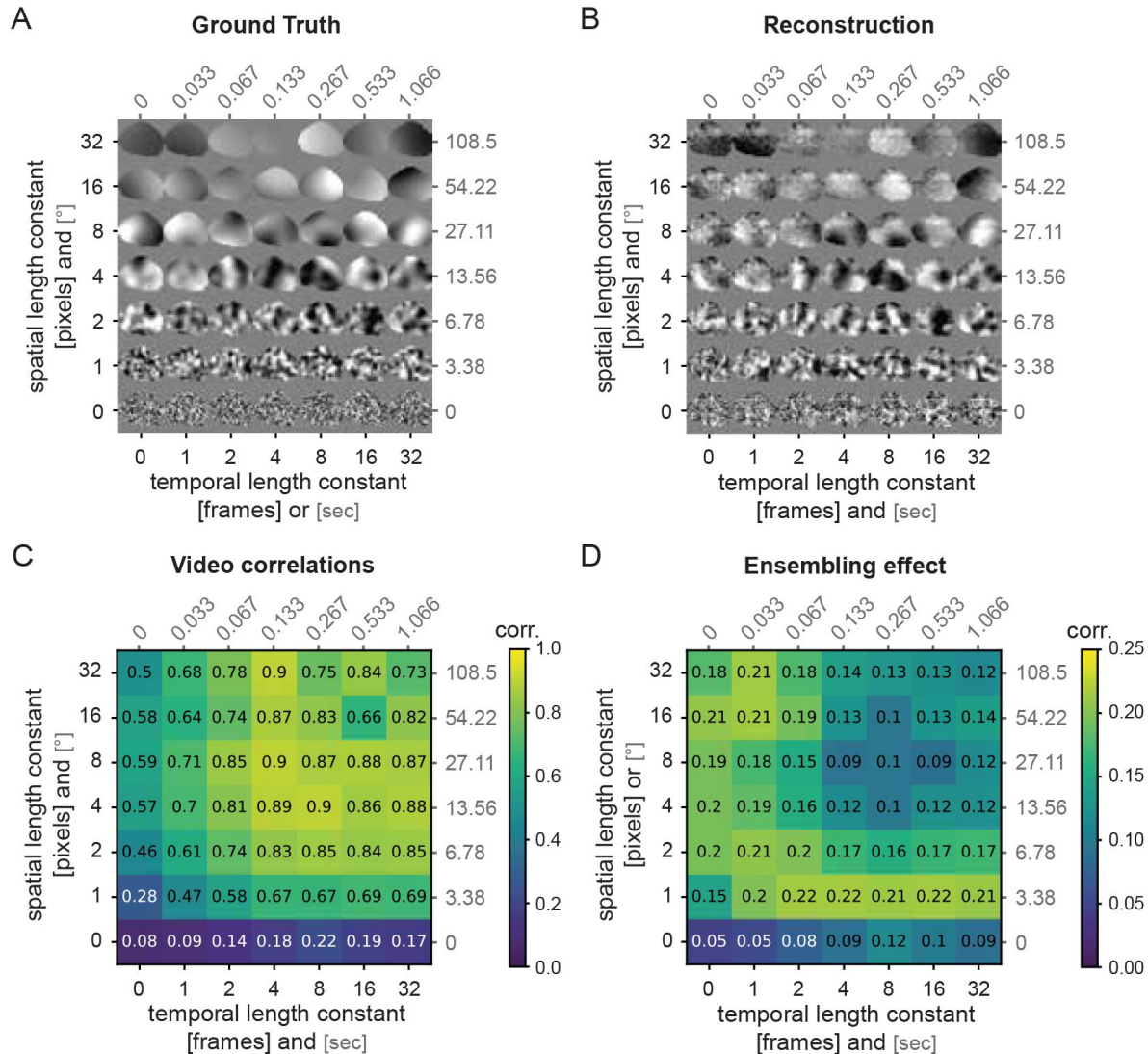


Figure 4

Reconstruction of Gaussian noise across the spatial and temporal spectrum using predicted activity. A) Example Gaussian noise stimulus set with evaluation mask for one mouse. Shown is the last frame of a 2 second video. B) Reconstructed Gaussian stimuli with SOTA DNEM predicted neuronal activity as the target (see also Supplementary Video 2: [YouTube link](#)). C) Pearson's correlation between ground truth (A) and reconstructed videos (B) across the range of spatial and temporal length constants. For each stimulus type the average correlation across 5 movies reconstructed from the SOTA DNEM of 3 mice is given. D) Ensembling effect for each stimulus. Video correlation for ensembled prediction (average videos from 7 model instances) minus the mean video correlation across the 7 individual model instances.

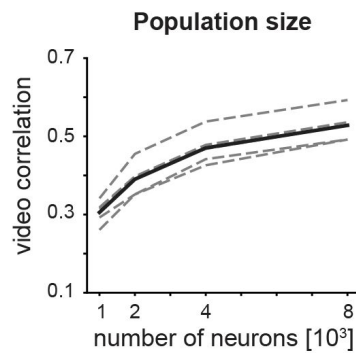


Figure 5

Video reconstruction using fewer neurons (i.e. population ablation) leads to lower reconstruction quality. Dashed lines are individual animals, solid line is mean. One-way repeated measures ANOVA $p = 1.58 \times 10^{-12}$. Bonferroni corrected paired t-test outcomes between consecutive drops in population size are all $p < 0.001$, $n = 5$ mice.

relate to our own method. In general, image reconstruction methods can be categorized into one of four groups: direct decoding models, encoder-decoder models, invertible encoding models, and encoder model input optimization.

Direct decoders, directly decode the input image/videos from neuronal activity with deep neuronal networks [Shen et al., 2019a [↗](#), Zhang et al., 2020 [↗](#), Li et al., 2023 [↗](#)]. When training direct decoders, the decoders can be pretrained [Ren et al., 2021 [↗](#)] or additional constraints can be added to the loss function to encourage the decoder to produce images that adhere to learned image statistics [Shen et al., 2019a [↗](#), Kupersmidt et al., 2022 [↗](#)]. A direct decoder approach has been used for video reconstruction in mice [Chen et al., 2024 [↗](#)], but in that case, the training and test movies were the same, meaning it is unclear if out-of-training set generalization was achieved (a key distinction between sensory reconstruction and stimulus identification, see previous section).

In encoder-decoder models the aim is to combine separately trained brain encoders (brain activity to latent space) and decoders (latent space to image/video). Recently this approach has become particularly popular because it allows the use of SOTA generative image models such as stable diffusion [Rombach et al., 2021 [↗](#), Takagi and Nishimoto, 2023 [↗](#), Scotti et al., 2023 [↗](#), Chen et al., 2023 [↗](#), Benchetrit et al., 2023 [↗](#)]. The encoder part of the models are first trained to translate brain activity into a latent space that the pretrained generative networks can interpret. Because these latent spaces are often conditioned on semantic information, this lends itself to separate processing of low-level visual and high-level semantic information from brain activity [Scotti et al., 2023 [↗](#)].

Invertible encoding models are encoding models which, once trained to predict neuronal activity, can implicitly be inverted to predict sensory input given brain activity. We would also include those models in this class which first compute the receptive field or preferred stimulus of neurons (or voxels) and reconstruct the input as the weighted sum of the receptive fields by their activity [Stanley et al., 1999 [↗](#), Thirion et al., 2006 [↗](#), Garasto et al., 2019 [↗](#), Brackbill et al., 2020 [↗](#), Yoshida and Ohki, 2020 [↗](#), Nishimoto et al., 2011 [↗](#)]. The down-side of this approach is that these encoding models generally under perform in terms of capturing the coding properties of neurons compared to more complex deep neural networks [Willeke et al., 2023 [↗](#)].

Encoder input optimization, also involves first training an encoder which predicts the activity of neurons or voxels given sensory input. Once trained the encoder is fixed and the input to the network is optimized using backpropagation until the predicted activity matches the observed activity [Pierchlewicz et al., 2023 [↗](#)]. Unlike with invertible encoding models any SOTA neuronal encoding model can be used. But like invertible models, the networks are not specifically trained to reconstruct images so they may be less likely to extrapolate information encoded by the brain by learning general image statistics.

Although outlined here as 4 distinct classes these approaches can be combined. For instance, Encoding input optimization can be combined with image diffusion [Pierchlewicz et al., 2023 [↗](#)] and in principle invertible models could also be combined in such a way.

We chose to pursue a pure encoder input optimization approach for single cell mouse visual cortex activity for two reasons. First, there have been considerable advances in the performance of neuronal encoding models for dynamic visual stimuli [Sinz et al., 2018 [↗](#), Wang et al., 2023 [↗](#), Turishcheva et al., 2024 [↗](#)] and we aimed to take advantage of these developments. Second, the addition of a generative decoder trained to producing high quality images brings with it the risk of extrapolating information based on general image statistics rather than interpreting what the brain is representing. In some cases, the brain may not be encoding coherent images and in those cases we would argue image reconstruction should fail, rather than producing an image when only the semantic information is present.

3.3 Key contributions and limitations

We demonstrate high-quality video reconstruction from mouse V1 using SOTA DNEMs to iteratively optimize the input video to match the resulting predicted activity with the recorded neuronal activity. Key to achieving high-quality reconstructions is model ensembling and using a large enough number of recorded neurons over a given retinotopic area.

While we averaged the video reconstructions from several models, an alternative method would be to average the gradients calculated by multiple models at each epoch, as has been done for the generation of maximally exciting images in the past [Walker et al., 2019]. However, this requires a large amount of GPU memory when using video models and is likely not practical with most hardware limitations. However, there might be situations in which averaging gradients yields better reconstructions. For instance, there may be multiple solutions for the activation pattern of a neural population, e.g. if their responses are translation/phase invariant [Ito et al., 1995, Tacchetti et al., 2018]. In such a case, averaging ‘misaligned’ reconstructions from multiple models might degrade overall quality.

The SOTA DNEM we used takes video data at an angular resolution of $3.4^\circ/\text{pixels}$ at the center of the screen which is about 3x worse than the visual acuity of mice ($\approx 0.5 \text{ cycles}^\circ$ [Prusky and Douglas, 2004]). As our model can reconstruct Gaussian noise stimuli down to a spatial length constant of 1 pixel, and drifting gratings up to a spatial frequency of $0.071 \text{ cycles}^\circ$, there is still some potential for improving spatial resolution. To close this gap and achieve reconstructions equivalent to the limit of mouse visual acuity, a different dataset and model would likely need to be developed. However, the frame rate of the videos the SOTA DNEM takes as input (30 Hz) is faster than the flicker fusion frequency of mice (14 Hz [Nomura et al., 2019]) and our tests with Gaussian noise and drifting grating stimuli show that the temporal resolution of reconstruction is close to this expected limit. Future efforts should therefore focus on the spatial resolution of video reconstruction rather than the temporal resolution.

It is, however, unclear how closely the representation of vision by the brain is expected to match the actual input. There are a number of visual processing phenomena that have previously been identified which leads us to suspect that some deviations between video reconstructions and ground truth input are to be expected. One such phenomenon is predictive coding [Rao and Ballard, 1999, Fiser et al., 2016]. It is possible that the unexpected parts of visual stimuli are sharper and have higher contrast compared to the expected parts when reconstructed from neuronal activity. Alternatively, perceptual learning is a phenomenon where visual stimulus detection or discriminability is enhanced through prolonged training [Li, 2015] and is associated with changes in the tuning distribution of neurons in the visual system [Goltstein et al., 2013, Poort et al., 2015, Jurjut et al., 2017, Schumacher et al., 2022]. Similarly, selective feature attention can modulate the response amplitude of neurons that have a preference for the features that are currently being attended to [Kanamori and Mrsic-Flogel, 2022]. Visual task engagement and training could therefore alter the accuracy and biases of what features of a video can accurately be reconstructed from the neuronal activity.

Such visual processing phenomena are likely difficult to investigate using existing fMRI-based reconstruction approaches, due to the low spatial and temporal resolution of the data. Additionally, many of these fMRI-based reconstruction approaches rely on the use of pretrained generative diffusion models to achieve more naturalistic and semantically interpretable images [Takagi and Nishimoto, 2023, Ozcelik and VanRullen, 2023, Scotti et al., 2023, Chen et al., 2023], but very likely at the cost of introducing information that may not be present in the actual neuronal representation. In contrast, our video reconstruction approach using single-cell resolution recordings, without a pretrained generative model, provides a more accurate method to investigate visual processing phenomena such as predictive coding, perceptual learning, and selective feature attention.

In conclusion, we reconstruct videos presented to mice based on the activity of neurons in the mouse visual cortex, with a ≈ 2 -fold improvement in pixel-by-pixel correlation compared to previous static image reconstruction methods. This paves the way to using movie reconstruction as a tool to investigate a variety of visual processing phenomena.

4 Methods

4.1 Source data

The data was provided by the Sensorium 2023 competition [Turishcheva et al., 2023 [↗](#), 2024 [↗](#)] and downloaded from <https://gin.g-node.org/pollytur/Sensorium2023Data> [↗](#) and https://gin.g-node.org/pollytur/sensorium_2023_dataset [↗](#). The data included grayscale movies presented to the mice at 30 Hz on a 31.8 by 56.5 cm monitor 15 cm from and perpendicular to the left eye. The movies were provided as spatially downsampled versions of the original screen resolution to 36 by 64 pixels, corresponding to an angular resolution of 3.4° /pixel at the center of the screen. The pupil position and diameter were recorded at 20 Hz and the running at 100 Hz. The neuronal activity was measured using two-photon imaging [Denk et al., 1990 [↗](#)] of GCaMP6s [Chen et al., 2013 [↗](#)] fluorescence at 8 Hz, extracted and deconvolved using the CAIMAN pipeline [Giovannucci et al., 2019 [↗](#)]. For each of the 10 mice, the activity of ≈ 8000 neurons was provided. The different data types were resampled to 30 Hz.

4.2 State-of-the-art dynamic neural encoding model

We used the winning model of the Sensorium 2024 competition, DwiseNeuro [Turishcheva et al., 2023 [↗](#), 2024 [↗](#)]. The code for the SOTA DNEM was downloaded from <https://github.com/IRomul/sensorium> [↗](#). The full model consists of 3 main components: core, cortex, and readout. The core largely consisted of factorized 3D convolution blocks with residual connections, positional encoding [Vaswani et al., 2017 [↗](#)] and SiLU activations [Elfwing et al., 2017 [↗](#)] followed by spatial average pooling. The cortex consisted of three fully connected layers. The readout consisted of a 1D convolution for each mouse with a final Softplus nonlinearity, that gives activity predictions for all neurons of each mouse. The kernel of the input layer had size 16 with a dilation of 2 in the time dimension, so spanned 32 video frames.

The original ensemble of models consisted of 7 model instances trained on a 7-fold cross-validation split of all available Sensorium 2023 competition data (≈ 1 hour of training data and ≈ 8 min of cross-validation data per fold from each mouse). Each model instance was trained on 6 of 7 data folds, with different validation data excluded from training for each model. To allow ensembled reconstructions of videos without test set contamination we instead retrained the models with a shared validation fold, i.e. we retrained the models leaving out the same validation data for all 7 model instances. The only other difference in the training procedure was that we retrained the models using a batch size of 24 instead of 32, this did not change the performance of neuronal response prediction on the withheld data folds (mean validation fold predicted vs ground truth response correlation for original weights: 0.293; and retrained weights: 0.291). We also did not use model distillation, while the original model did (see <https://github.com/IRomul/sensorium> [↗](#)).

4.3 Additional visual stimuli

The Gaussian noise stimuli were downloaded from https://github.com/TomGeorge1234/gp_video [↗](#) and spanned a range of 0 to 32 pixels in spatial length constant and 0 to 32 frames in temporal length constant used in the Gaussian process. The drifting grating stimuli were produced using PsychoPy [Peirce et al., 2019 [↗](#)] and ranged from 0.5 to 0.062 cycles/degree and 0.5 to 0

cycles/second, with 2 seconds of movie for each cardinal direction. These ranges were chosen to avoid aliasing effects in the 36 by 64 pixel videos. The highest temporal frequency corresponds to a flicker stimulus.

4.4 Mask training

To generate the transparency masks we used an alpha blending approach inspired by [Mordvintsev et al., 2018, Willeke et al., 2023]. A transparency layer was placed at the input to the SOTA DNEM. This transparency layer was used to alpha blend the true video V with another randomly selected background video BG from the data:

$$V_{BG} = V * \alpha + V_{BG} * (1 - \alpha) \quad (1)$$

where α is the 2D transparency mask and V_{BG} is the blended input video. This mask was optimized using stochastic gradient descent (for 1000 epochs with learning rate 10) with mean squared error (MSE) loss between the true responses y and the predicted responses $\hat{y}$ scaled by the average weight of the transparency mask $\bar{\alpha}$:

$$MSE(y, \hat{y}) = \sum_{i=1}^n (y - \hat{y})^2 \quad (2)$$

$$Loss = MSE(y, \hat{y} * (1 - \bar{\alpha})) \quad (3)$$

where n is the total number of neurons. The mask was initialized as uniform noise between 0 and 0.05. At each epoch the neuronal activity in response to a randomly selected 32 frame video segment from the training set was predicted and the gradients of the loss (Equation 3) with respect to the pixels in the transparency mask α were calculated for each video frame. The gradients were normalized by their matrix norm, clipped to between -1 and 1 and averaged across frames. The gradients were smoothed with a 2D Gaussian kernel of $\sigma = 5$ and subtracted from the transparency mask. The transparency mask was only calculated using one SOTA DNEM instance using its validation fold. See Supplementary Algorithm 2.

The transparency mask was thresholded and binarized at 0.5 for the masked gradients ∇_{masked} or 1 for the masked videos for evaluation V_{eval} :

$$\nabla_{masked} = \nabla * (\alpha > 0.5) \quad (4)$$

$$V_{eval} = V * (\alpha > 1) \quad (5)$$

where ∇ is the gradients of the loss with respect to each pixel in the video and V is the reconstructed video before masking. These masks were trained independently for each mouse using one model instance with the original weights of the model <https://github.com/IRomul/sensorium>, not the retrained models used in the rest of this paper to reconstruct the videos.

4.5 Video reconstruction

To reconstruct the input video we initialized the video as uniform gray values and concatenated the ground truth behavioral parameters. The SOTA DNEM took 32 frames at a time and we shifted this window by 8 frames until all frames of the whole 10s video were covered. For each 32-frame window, the Poisson negative log-likelihood loss between the predicted and true neuronal responses was calculated:

$$Loss(y, \hat{y}) = \sum \hat{y}_{neuron} - y_{neuron} \log(\hat{y}_{neuron} + 10^{-8}) \quad (6)$$

where $\hat{y}$ are the predicted responses and y are the ground truth responses. The gradients of the loss with respect to each pixel of the input video were calculated for each window of frames and averaged across all windows. The gradients for each pixel were normalized by the matrix norm across all gradients and clipped to between -1 and 1. The gradients were masked (Equation 4) and applied to the input video using Adam without second order momentum [Kingma and Ba, 2014] ($\beta_1 = 0.9$) for 1000 epochs and a learning rate of 1000, with a learning rate warm-up for the first 10 epochs. After each epoch, the video was clipped to between 0 and 255. The optimization was run for 1000 epochs. 7 reconstructions from 7 model instances were averaged, denoised with a 3D Gaussian filter $\sigma = 0.5$ (unless specified otherwise), and masked with the evaluation mask. See Supplementary Algorithm 1. Optimizing each 10-second video with one model instance for 1000 epochs took ≈ 60 min using a desktop with an RTX4070 GPU.

4.6 Reconstruction quality assessment

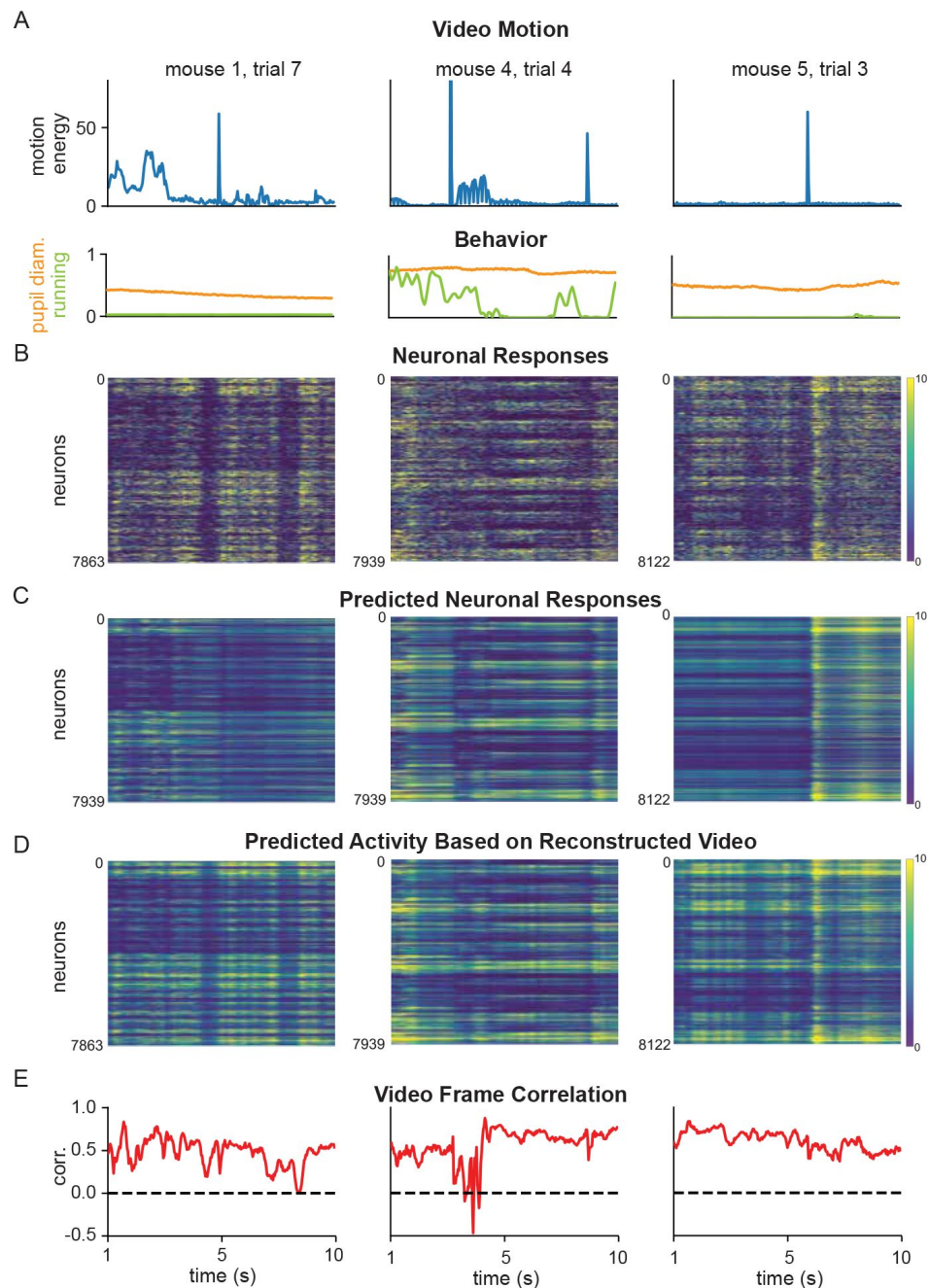
To evaluate the similarity between reconstructed and ground truth videos, we used the mean Pearson's correlation between pixels of corresponding frames to evaluate spatial similarity:

$$\text{mean frame correlation} = \frac{1}{f} \sum_{i=1}^f \frac{\text{cov}(x_i, \hat{x}_i)}{\sigma_{x_i} \sigma_{\hat{x}_i}} \quad (7)$$

where f is the number of frames, and x_i and $\hat{x}_i$ are the ground truth and reconstructed frames. To evaluate temporal and spatial similarity between ground truth and reconstructed videos we used the Pearson's correlation between all pixels of the whole movie:

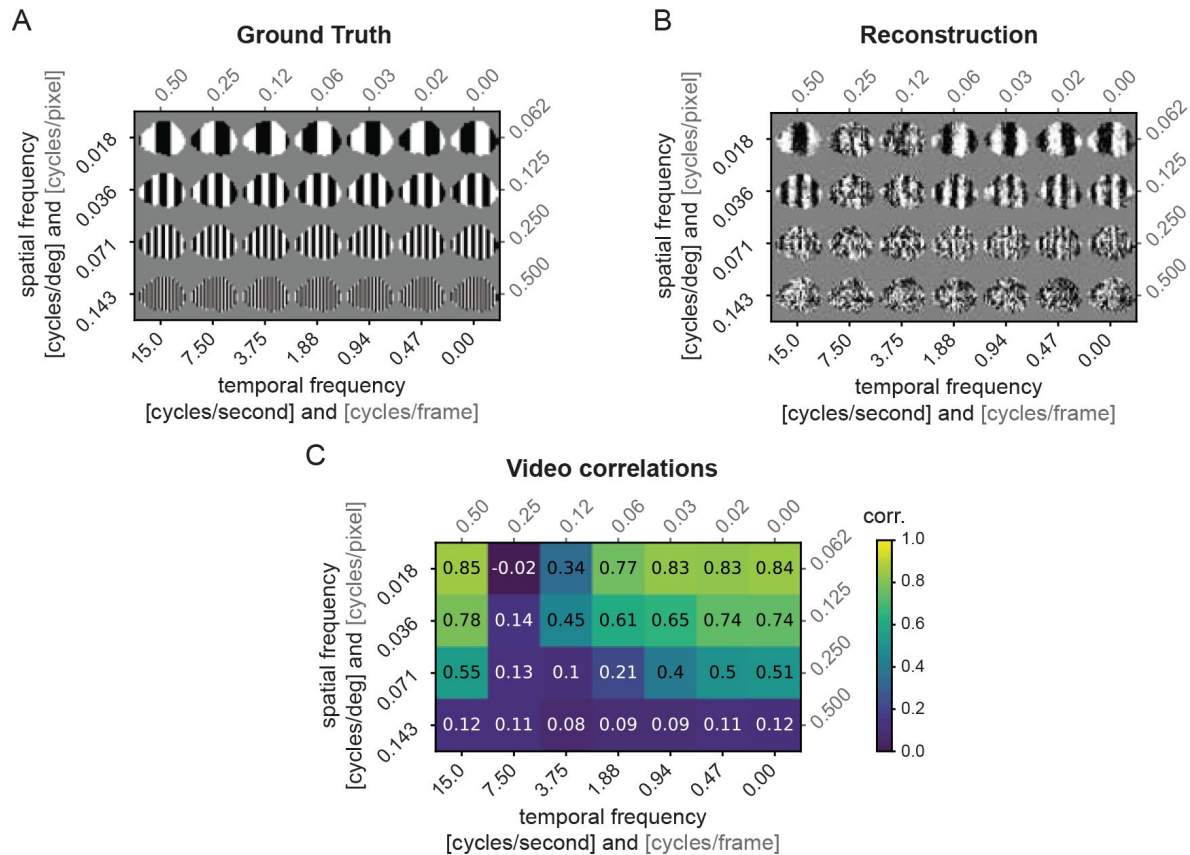
$$\text{video correlation} = \frac{\text{cov}(x, \hat{x})}{\sigma_x \sigma_{\hat{x}}} \quad (8)$$

A Appendix / supplemental material



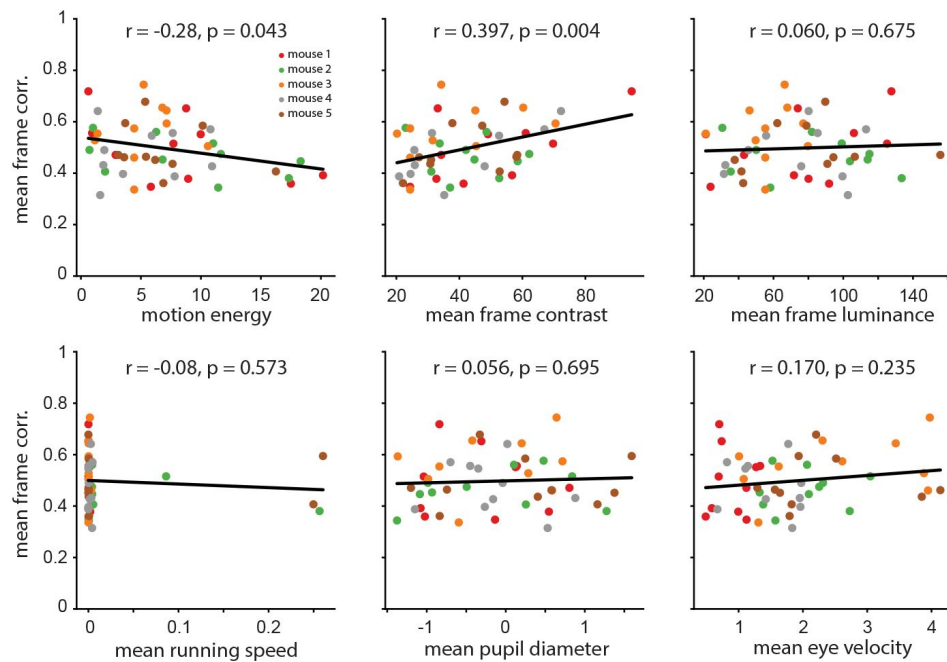
Supplementary Figure S1

Summary ethogram of SOTA DNEM inputs, output predictions, and video reconstruction over time for three videos from three mice (same as **Figure 2A**). A) Top: motion energy of the input video. Bottom: pupil diameter and running speed of the mouse during the video. B) Ground truth neuronal activity. C) Predicted neuronal activity in response to input video and behavioural parameters. D) predicted neuronal activity given reconstructed video and ground truth behaviour as input. E) Frame by frame correlation between reconstructed and ground truth video.



Supplementary Figure S2

Reconstruction of drifting grating stimuli with different spatial and temporal frequencies using predicted activity. A) Example drifting grating stimuli (rightwards moving) masked with the evaluation mask for one mouse. Shown is the 31st frame of a 2 second video. B) Reconstructed drifting grating stimuli with SOTA DNEM predicted neuronal activity as the target (see also Supplementary Video 3: [YouTube link](#)). C) Pearson's correlation between ground truth (A) and reconstructed videos (B) across the range of spatial and temporal frequencies. For each stimulus type the average correlation across 4 directions (up, down, left, right) reconstructed from the SOTA DNEM of 1 mouse is given. Interestingly, video correlation at 15 cycles/second (half the video frame rate of 30 Hz) is much higher than 7.5 cycles/second. This is an artifact of using predicted responses rather than true neural responses. The DNEM input layer convolution has dilation 2. The predicted activity is therefore based on every second frame, with the effect that the activity is predicted as the response of 2 static images which are then interleaved.



Supplementary Figure S3

Reconstruction performance correlates with video motion energy and frame contrast but not with behavioral parameters. Pearson's correlation between mean frame correlation per movie, and 3 movie parameters and 3 behavioral parameters. Linear fit as black line.

Algorithm 1

Movie reconstruction

```

1: Parameters:
2: Dynamic Neural Encoding Model (DNEM), Ground truth neuronal response  $y$ , predicted neuronal
   responses  $\hat{y}$ , predicted input video  $\hat{x}$ , video width  $w = 64$ , video height  $h = 36$ , number of
   frames  $n = 300$ , transparency mask  $\alpha$ , sliding window size  $k = 32$ , sliding window stride  $s = 8$ ,
   total number of windows  $N \leftarrow \lfloor 2 + (n - k) / s \rfloor$ , learning rate  $lr = 1000$ ,  $\beta_1 = 0.9$ , loss function
    $Loss(y, \hat{y}) = \sum \hat{y}_{neuron} - y_{neuron} \log(\hat{y}_{neuron} + 10^{-8})$ , number of epochs  $epochs = 1000$ ,
   number of model instances  $q = 7$ 

3: Objective:  $F(y, \hat{x}) = Loss(y, DNEM(\hat{x}))$ 

4: Initialize variables:
5:  $\hat{x}_{f,g,w} \leftarrow 125$ 
6:  $G_{i,f,h,w} \leftarrow 0$ 
7:  $g_{f,h,w} \leftarrow 0$ 
8:  $m_{f,h,w} \leftarrow 0$ 

9: for iteration  $t = 1, \dots, epochs$  do

10: Gradients across all windows:
11: for iteration  $i = 1, \dots, N$  do
12:   if  $i < N$  then
13:      $f \leftarrow [s(i - 1), \dots, k + s(i - 1)]$ 
14:   else if  $i = N$  then
15:      $f \leftarrow [n - k, \dots, n]$ 
16:   end if

17:    $G_{i,f} \leftarrow \nabla F(y_f, \hat{x}_f)$ 
18:    $G_{i,f} \leftarrow \frac{G_{i,f}}{\|G_{i,f}\| + 10^{-6}}$ 
19: end for

20: Average gradients across windows:  $g \leftarrow \frac{1}{N} \sum_{i=1}^N G_i$ 

21: Clip and mask gradients:
22:  $g \leftarrow \text{clip}(g, -1, 1)$ 
23:  $g_{f,h,w} \leftarrow g_{f,h,w} \odot (\alpha_{h,w} > 0.5)$ 

24: Update movie:
25: if  $t < 10$  then
26:    $lr_{\text{current}} \leftarrow lr \cdot \frac{t}{10}$ 
27: else if  $t > 10$  then
28:    $lr_{\text{current}} \leftarrow lr$ 
29: end if

30:  $lr_{\text{current}} \leftarrow lr_{\text{current}} \sqrt{1 - 0.999^t} / (1 - \beta_1^t)$ 
31:  $m \leftarrow \beta_1 \cdot m + (1 - \beta_1) \cdot g$ 
32:  $\hat{m} \leftarrow \frac{m}{1 - \beta_1^t}$ 
33:  $\hat{x} \leftarrow \hat{x} - lr_{\text{current}} \cdot \hat{m}$ 

34: Clip movie:  $\hat{x} \leftarrow \text{clip}(\hat{x}, 0, 255)$ 

35: end for

36: Ensembling:  $\hat{x} \leftarrow \frac{1}{q} \sum_{i=1}^q \hat{x}$ 

37: Denoise (optional):  $\hat{x} \leftarrow \text{GaussianBlur3D}(\hat{x}, \sigma = (0.5, 0.5, 0.5))$ 

38: Mask movie:  $\hat{x}_{f,h,w} \leftarrow \hat{x}_{f,h,w} \odot (\alpha_{h,w} > 1)$ 

```

Algorithm 2

Mask training

```

1: Parameters:
2: Dynamic Neural Encoding Model (DNEM), Ground truth neuronal response  $y$ , predicted neuronal
   responses  $\hat{y}$ , ground truth input video  $x$ , background video  $background$ , video width  $w = 64$ ,
   video height  $h = 36$ , frames  $f$ , number of frames  $n = 32$ , transparency mask  $\alpha$ , learning rate
    $lr = 10$ , loss function  $Loss(y, \hat{y}, \alpha) = MSE(y, \hat{y}(1 - \bar{\alpha}))$ , number of epochs  $epochs = 1000$ ,
   alpha blending function  $Blend(x, background, \alpha) = x \odot \alpha + background \odot (1 - \alpha)$ 

3: Objective:
4:  $F(y, x, background, \alpha) = Loss(y, DNEM(Blend(x, background, clip(\alpha, -1, 1))), \alpha)$ 

5: Initialize variables:
6:  $\alpha_{h,w} \leftarrow U(0, 0.5)$ 

7: for iteration  $t = 1, \dots, epochs$  do

8:    $x \leftarrow$  get random video
9:    $background \leftarrow$  get different random video
10:   $background \leftarrow Flip(background, dim = 1)$ 
11:   $background \leftarrow Flip(background, dim = 2)$ 
12:   $background \leftarrow Flip(background, dim = 3)$ 

13:  Gradients with respect to transparency mask:
14:   $G \leftarrow \frac{\partial F(y, x, background, \alpha)}{\partial \alpha}$ 
15:   $G \leftarrow \frac{G}{\|G\| + 10^{-6}}$ 

16:  Average gradients across frames:
17:   $g \leftarrow \frac{1}{n} \sum_{f=1}^n G_f$ 

18:  Clip and mask gradients:
19:   $g \leftarrow clip(g, -1, 1)$ 
20:   $g \leftarrow GaussianBlur2D(g, \sigma = (5, 5))$ 

21:  Update mask:
22:   $\alpha \leftarrow \alpha - lr \cdot g$ 

23: end for

```

Acknowledgements

We would like to thank Emmanuel Bauer, Sandra Reinert and the anonymous reviewers for useful input and discussions, and Tom George for the Gaussian noise stimulus set. T.W.M. is funded by The Wellcome Trust (219627/Z/19/Z; 214333/Z/18/Z) and Gatsby Charitable Foundation (GAT3755) and J.B. is funded by EMBO (ALTF 415-2024).

Additional information

5 Code

The code is available at https://github.com/Joel-Bauer/movie_reconstruction_code .

Additional files

Supplementary Videos. *Supplementary Video 1: Reconstructed natural videos from mouse brain activity. Odd rows are ground truth (GT) movie clips presented to mice. Even rows are the reconstructed movies from the activity of ≈ 8000 V1 neurons. Reconstructed movies are smoothed ($\sigma = 0.5$ pixels), masked, and contrast (std) and luminance (mean) matched to ground truth movies. Supplementary Video 2: Gaussian noise stimuli and reconstructions. Odd rows are ground truth (GT) video inputs to the model. Even rows are the reconstructed videos from the predicted neuronal activity for 1 mouse. Reconstructed movies are masked, and contrast (std) and luminance (mean) matched to ground truth videos. Supplementary Video 3: Drifting grating stimuli and reconstructions. Odd rows are ground truth (GT) video inputs to the model. Even rows are the reconstructed videos from the predicted neuronal activity for 1 mouse. Reconstructed movies are masked, and contrast (std) and luminance (mean) matched to ground truth videos.* [🔗](#)

References

- Baikulov Ruslan (2023) **Solution for Sensorium 2023 Competition (v23.11.22)** *Zenodo* <https://doi.org/10.5281/zenodo.10155151>
- Benchetrit Yohann, Banville Hubert, King Jean-Rémi (2023) **Brain decoding: toward real-time reconstruction of visual perception** *arXiv* <https://doi.org/10.48550/arxiv.2310.19812>
- Brackbill Nora, Rhoades Colleen, Kling Alexandra, Shah Nishal P, Sher Alexander, Litke Alan M, Chichilnisky EJ (2020) **Reconstruction of natural images from responses of primate retinal ganglion cells** *eLife* **9**:e58516 <https://doi.org/10.7554/elife.58516>
- Breiman Leo (1996) **Stacked regressions** *Machine Learning* **24**:49–64 <https://doi.org/10.1007/bf00117832>
- Chen Tsai-Wen, Wardill Trevor J., Sun Yi, Pulver Stefan R., Renninger Sabine L., Baohan Amy, Schreiter Eric R., Kerr Rex A., Orger Michael B., Jayaraman Vivek, Looger Loren L., Svoboda Karel, Kim Douglas S. (2013) **Ultrasensitive fluorescent proteins for imaging neuronal activity** *Nature* **499**:295–300 <https://doi.org/10.1038/nature12354>
- Chen Ye, Beech Peter, Yin Ziwei, Jia Shanshan, Zhang Jiayi, Yu Zhaofoei, Liu Jian K. (2024) **Decoding dynamic visual scenes across the brain hierarchy** *PLOS Computational Biology* **20**:e1012297 <https://doi.org/10.1371/journal.pcbi.1012297>
- Chen Zijiao, Qing Jiaxin, Zhou Juan Helen (2023) **Cinematic Mindscapes: High-quality Video Reconstruction from Brain Activity** *arXiv* <https://doi.org/10.48550/arxiv.2305.11675>
- Deitch Daniel, Rubin Alon, Ziv Yaniv (2021) **Representational drift in the mouse visual cortex** *Current Biology* **31**:4327–4339 <https://doi.org/10.1016/j.cub.2021.07.062>
- Denk Winifried, Strickler James H., Webb Watt W. (1990) **Two-Photon Laser Scanning Fluorescence Microscopy** *Science* **248**:73–76 <https://doi.org/10.1126/science.2321027>
- Elfwing Stefan, Uchibe Eiji, Doya Kenji (2017) **Sigmoid-Weighted Linear Units for Neural Network Function Approximation in Reinforcement Learning** *arXiv* <https://doi.org/10.48550/arxiv.1702.03118>
- Fiser Aris, Mahringer David, Oyibo Hassana K, Petersen Anders V, Leinweber Marcus, Keller Georg B (2016) **Experience-dependent spatial expectations in mouse visual cortex** *Nature Neuroscience* **19**:1658–1664 <https://doi.org/10.1038/nn.4385>
- Földiák Peter (1993) **The ‘Ideal Homunculus’: Statistical Inference from Neural Population Responses** In: *Computation and Neural Systems* pp. 55–60 https://doi.org/10.1007/978-1-4615-3254-5_9
- Garasto Stef, Nicola Wilten, Bharath Anil A., Schultz Simon R. (2019) **Neural Sampling Strategies for Visual Stimulus Reconstruction from Two-photon Imaging of Mouse Primary Visual Cortex** In: *2019 9th International IEEE/EMBS Conference on Neural Engineering (NER)* pp. 566–570 <https://doi.org/10.1109/ner.2019.8716934>

Giovannucci Andrea, Friedrich Johannes, Gunn Pat, Kalfon Jérémie, Brown Brandon L, Koay Sue Ann, Taxisdis Jiannis, Najafi Farzaneh, Gauthier Jeffrey L, Zhou Pengcheng, Khakh Baljit S, Tank David W, Chklovskii Dmitri B, Pnevmatikakis Eftychios A (2019) **CaImAn an open source tool for scalable calcium imaging data analysis** *eLife* **8**:e38173 <https://doi.org/10.7554/elife.38173>

Goltstein Pieter M., Coffey Emily B. J., Roelfsema Pieter R., Pennartz Cyriel M. A. (2013) **In Vivo Two-Photon Ca²⁺ Imaging Reveals Selective Reward Effects on Stimulus-Specific Assemblies in Mouse Visual Cortex** *The Journal of Neuroscience* **33**:11540–11555 <https://doi.org/10.1523/jneurosci.1341-12.2013>

Ho Jun Kai, Horikawa Tomoyasu, Majima Kei, Cheng Fan, Kamitani Yukiyasu (2023) **Inter-individual deep image reconstruction via hierarchical neural code conversion** *NeuroImage* **271**:120007 <https://doi.org/10.1016/j.neuroimage.2023.120007>

Ito M., Tamura H., Fujita I., Tanaka K. (1995) **Size and position invariance of neuronal responses in monkey inferotemporal cortex** *Journal of Neurophysiology* **73**:218–226 <https://doi.org/10.1152/jn.1995.73.1.218>

Jurjut Ovidiu, Georgieva Petya, Busse Laura, Katzner Steffen (2017) **Learning Enhances Sensory Processing in Mouse V1 before Improving Behavior** *The Journal of Neuroscience* **37**:6460–6474 <https://doi.org/10.1523/jneurosci.3485-16.2017>

Kanamori Takahiro, Mscic-Flogel Thomas D. (2022) **Independent response modulation of visual cortical neurons by attentional and behavioral states** *Neuron* **110**:3907–3918 <https://doi.org/10.1016/j.neuron.2022.08.028>

Kay Kendrick N., Naselaris Thomas, Prenger Ryan J., Gallant Jack L. (2008) **Identifying natural images from human brain activity** *Nature* **452**:352–355 <https://doi.org/10.1038/nature06713>

Kingma Diederik P, Ba Jimmy (2014) **Adam: A Method for Stochastic Optimization** *arXiv* <https://doi.org/10.48550/arxiv.1412.6980>

Kupersmidt Ganit, Bely Roman, Gaziv Guy, Irani Michal (2022) **A Penny for Your (visual) Thoughts: Self-Supervised Reconstruction of Natural Movies from Brain Activity** *arXiv* <https://doi.org/10.48550/arxiv.2206.03544>

Li Wenyi, Zheng Shengjie, Liao Yufan, Hong Rongqi, He Chenggang, Chen Weiliang, Deng Chunshan, Li Xiaojian (2023) **The brain-inspired decoder for natural visual image reconstruction** *Frontiers in Neuroscience* **17** <https://doi.org/10.3389/fnins.2023.1130606>

Li Wu (2015) **Perceptual Learning: Use-Dependent Cortical Plasticity** *Annual Review of Vision Science* **2**:1–22 <https://doi.org/10.1146/annurev-vision-111815-114351>

Mordvintsev Alexander, Pezzotti Nicola, Schubert Ludwig, Olah Chris (2018) **Differentiable image parameterizations** *Distill* **3** <https://doi.org/10.23915/distill.00012>

Niell Christopher M., Stryker Michael P. (2010) **Modulation of Visual Responses by Behavioral State in Mouse Visual Cortex** *Neuron* **65**:472–479 <https://doi.org/10.1016/j.neuron.2010.01.033>

Nishimoto Shinji, Vu An T., Naselaris Thomas, Benjamini Yuval, Yu Bin, Gallant Jack L. (2011) **Reconstructing Visual Experiences from Brain Activity Evoked by Natural Movies** *Current Biology* **21**:1641–1646 <https://doi.org/10.1016/j.cub.2011.08.031>

Nomura Yuichiro, Ikuta Shohei, Yokota Satoshi, Mita Junpei, Oikawa Mami, Matsushima Hiroki, Amano Akira, Shimonomura Kazuhiro, Seya Yasuhiro, Koike Chieko (2019) **Evaluation of critical flicker-fusion frequency measurement methods using a touchscreen-based visual temporal discrimination task in the behaving mouse** *Neuroscience Research* **148**:28–33 <https://doi.org/10.1016/j.neures.2018.12.001>

Ozcelik Furkan, VanRullen Rufin (2023) **Natural scene reconstruction from fMRI signals using generative latent diffusion** *Scientific Reports* **13**:15666 <https://doi.org/10.1038/s41598-023-42891-8>

Peirce Jonathan, Gray Jeremy R, Simpson Sol, MacAskill Michael, Sogo Hiroyuki, Kastman Erik, Kristoffer Lindeløv Jonas (2019) **Psychopy2: Experiments in behavior made easy** *Behavior research methods* **51**:195–203

Pierzchlewicz Pawel A., Willeke Konstantin F., Nix Arne F., Elumalai Pavithra, Restivo Kelli, Shinn Tori, Nealley Cate, Rodriguez Gabrielle, Patel Saumil, Franke Katrin, Tolias Andreas S., Sinz Fabian H. (2023) **Energy Guided Diffusion for Generating Neurally Exciting Images** *bioRxiv* <https://doi.org/10.1101/2023.05.18.541176>

Poort Jasper, Khan Adil G., Pachitariu Marius, Nemri Abdellatif, Orsolich Ivana, Krupic Julija, Bauza Marius, Sahani Maneesh, Keller Georg B., Mrsic-Flogel Thomas D., Hofer Sonja B. (2015) **Learning Enhances Sensory and Multiple Non-sensory Representations in Primary Visual Cortex** *Neuron* **86**:1478–1490 <https://doi.org/10.1016/j.neuron.2015.05.037>

Prusky G.T., Douglas R.M. (2004) **Characterization of mouse cortical spatial vision** *Vision Research* **44**:3411–3418 <https://doi.org/10.1016/j.visres.2004.09.001>

Rakhimberdina Zarina, Jodelet Quentin, Liu Xin, Murata Tsuyoshi (2021) **Natural Image Reconstruction From fMRI Using Deep Learning: A Survey** *Frontiers in Neuroscience* **15**:0795488 <https://doi.org/10.3389/fnins.2021.795488>

Rao Rajesh P. N., Ballard Dana H. (1999) **Predictive coding in the visual cortex: a functional interpretation of some extra-classical receptive-field effects** *Nature Neuroscience* **2**:79–87 <https://doi.org/10.1038/4580>

Reimer Jacob, Froudarakis Emmanouil, Cadwell Cathryn R., Yatsenko Dimitri, Denfield George H., Tolias Andreas S. (2014) **Pupil Fluctuations Track Fast Switching of Cortical States during Quiet Wakefulness** *Neuron* **84**:355–362 <https://doi.org/10.1016/j.neuron.2014.09.033>

Ren Ziqi, Li Jie, Xue Xuetong, Li Xin, Yang Fan, Jiao Zhicheng, Gao Xinbo (2021) **Reconstructing seen image from brain activity by visually-guided cognitive representation and adversarial learning** *NeuroImage* **228**:117602 <https://doi.org/10.1016/j.neuroimage.2020.117602>

Rombach Robin, Blattmann Andreas, Lorenz Dominik, Esser Patrick, Ommer Björn (2021) **High-Resolution Image Synthesis with Latent Diffusion Models** *arXiv* <https://doi.org/10.48550/arxiv.2112.10752>

Schneider Steffen, Lee Jin Hwa, Mathis Mackenzie Weygandt (2023) **Learnable latent embeddings for joint behavioural and neural analysis** *Nature* **12**:5170 <https://doi.org/10.1038/s41586-023-06031-6>

Schumacher Joseph W., McCann Matthew K., Maximov Katherine J., Fitzpatrick David (2022) **Selective enhancement of neural coding in V1 underlies fine-discrimination learning in**

tree shrew *Current Biology* <https://doi.org/10.1016/j.cub.2022.06.009>

Scotti Paul S, Banerjee Atmadeep, Goode Jimmie, Shabalin Stepan, Nguyen Alex, Cohen Ethan, Dempster Aidan J, Verlinde Nathalie, Yundler Elad, Weisberg David, Norman Kenneth A, Abraham Tanishq Mathew (2023) **Reconstructing the Mind's Eye: fMRI-to-Image with Contrastive Learning and Diffusion Priors** *arXiv* <https://doi.org/10.48550/arxiv.2305.18274>

Shen Guohua, Dwivedi Kshitij, Majima Kei, Horikawa Tomoyasu, Kamitani Yukiyasu (2019a) **End-to-End Deep Image Reconstruction From Human Brain Activity** *Frontiers in Computational Neuroscience* **13**:21 <https://doi.org/10.3389/fncom.2019.00021>

Shen Guohua, Horikawa Tomoyasu, Majima Kei, Kamitani Yukiyasu (2019b) **Deep image reconstruction from human brain activity** *PLoS Computational Biology* **15**:e1006633 <https://doi.org/10.1371/journal.pcbi.1006633>

Sinz Fabian H., Ecker Alexander S., Fahey Paul G., Walker Edgar Y., Cobos Erick, Froudarakis Emmanouil, Yatsenko Dimitri, Pitkow Xaq, Reimer Jacob, Tolias Andreas S. (2018) **Stimulus domain transfer in recurrent models for large scale cortical population prediction on video** *bioRxiv* :452672 <https://doi.org/10.1101/452672>

Stanley Garrett B, Li Fei F, Dan Yang (1999) **Reconstruction of Natural Scenes from Ensemble Responses in the Lateral Geniculate Nucleus** *The Journal of Neuroscience* **19**:8036–8042 <https://doi.org/10.1523/jneurosci.19-18-08036.1999>

Tacchetti Andrea, Isik Leyla, Poggio Tomaso A. (2018) **Invariant Recognition Shapes Neural Representations of Visual Input** *Annual Review of Vision Science* **4**:403–422 <https://doi.org/10.1146/annurev-vision-091517-034103>

Takagi Yu, Nishimoto Shinji (2023) **High-resolution image reconstruction with latent diffusion models from human brain activity** *bioRxiv* <https://doi.org/10.1101/2022.11.18.517004>

Thirion Bertrand, Duchesnay Edouard, Hubbard Edward, Dubois Jessica, Poline Jean-Baptiste, Lebihan Denis, Dehaene Stanislas (2006) **Inverse retinotopy: Inferring the visual content of images from brain activation patterns** *NeuroImage* **33**:1104–1116 <https://doi.org/10.1016/j.neuroimage.2006.06.062>

Turishcheva Polina, Fahey Paul G, Hansel Laura, Froebe Rachel, Ponder Kayla, Vystrcilová Michaela, Willeke Konstantin F, Bashiri Mohammad, Wang Eric, Ding Zhiwei, Tolias Andreas S, Sinz Fabian H, Ecker Alexander S (2023) **The Dynamic Sensorium competition for predicting large-scale mouse visual cortex activity from videos** *arXiv* <https://doi.org/10.48550/arxiv.2305.19654>

Turishcheva Polina, Fahey Paul G, Vystrcilová Michaela, Hansel Laura, Froebe Rachel, Ponder Kayla, Qiu Yongrong, Willeke Konstantin F, Bashiri Mohammad, Baikulov Ruslan, Zhu Yu, Ma Lei, Yu Shan, Huang Tiejun, Li Bryan M, De Wulf Wolf, Kudryashova Nina, Hennig Matthias H, Rochefort Nathalie L, Onken Arno, Wang Eric, Ding Zhiwei, Tolias Andreas S, Sinz Fabian H, Ecker Alexander S (2024) **Retrospective for the Dynamic Sensorium Competition for predicting large-scale mouse primary visual cortex activity from videos** *arXiv*

Vaswani Ashish, Shazeer Noam, Parmar Niki, Uszkoreit Jakob, Jones Llion, Gomez Aidan N, Kaiser Lukasz, Polosukhin Illia (2017) **Attention Is All You Need** *arXiv* <https://doi.org/10.48550/arxiv.1706.03762>

Walker Edgar Y., Sinz Fabian H., Cobos Erick, Muhammad Taliah, Froudarakis Emmanouil, Fahey Paul G., Ecker Alexander S., Reimer Jacob, Pitkow Xaq, Tolias Andreas S. (2019) **Inception loops discover what excites neurons most using deep predictive models** *Nature Neuroscience* **22**:2060–2065 <https://doi.org/10.1038/s41593-019-0517-x>

Wang Eric Y, Fahey Paul G, Ponder Kayla, Ding Zhuokun, Chang Andersen, Muhammad Taliah, Patel Saumil, Ding Zhiwei, Tran Dat, Fu Jiakun, Papadopoulos Stelios, Franke Katrin, Ecker Alexander S, Reimer Jacob, Pitkow Xaq, Sinz Fabian H, Tolias Andreas S (2023) **Towards a Foundation Model of the Mouse Visual Cortex** *bioRxiv* <https://doi.org/10.1101/2023.03.21.533548>

Willeke Konstantin F., Restivo Kelli, Franke Katrin, Nix Arne F., Cadena Santiago A., Shinn Tori, Nealley Cate, Rodriguez Gabrielle, Patel Saumil, Ecker Alexander S., Sinz Fabian H., Tolias Andreas S. (2023) **Deep learning-driven characterization of single cell tuning in primate visual area V4 unveils topological organization** *bioRxiv* <https://doi.org/10.1101/2023.05.12.540591>

Xia Ji, Marks Tyler D., Goard Michael J., Wessel Ralf (2021) **Stable representation of a naturalistic movie emerges from episodic activity with gain variability** *Nature Communications* **12**:5170 <https://doi.org/10.1038/s41467-021-25437-2>

Yoshida Takashi, Ohki Kenichi (2020) **Natural images are reliably represented by sparse and variable populations of neurons in visual cortex** *Nature Communications* **11**:872 <https://doi.org/10.1038/s41467-020-14645-x>

Zhang Yichen, Jia Shanshan, Zheng Yajing, Yu Zhaofei, Tian Yonghong, Ma Siwei, Huang Tiejun, Liu Jian K. (2020) **Reconstruction of natural visual scenes from neural spikes with deep neural networks** *Neural Networks* **125**:19–30 <https://doi.org/10.1016/j.neunet.2020.01.033>

Author information

Joel Bauer

Sainsbury Wellcome Centre, University College London, London, United Kingdom,
Bioengineering Dept, Imperial College, London, United Kingdom
ORCID iD: [0000-0001-5858-166X](https://orcid.org/0000-0001-5858-166X)

For correspondence: joel.bauer@ucl.ac.uk

Troy W Margrie[†]

Sainsbury Wellcome Centre, University College London, London, United Kingdom
ORCID iD: [0000-0002-5526-4578](https://orcid.org/0000-0002-5526-4578)

[†]These authors contributed equally as last authors

Claudia Clopath[†]

Sainsbury Wellcome Centre, University College London, London, United Kingdom,
Bioengineering Dept, Imperial College, London, United Kingdom
ORCID iD: [0000-0003-4507-8648](https://orcid.org/0000-0003-4507-8648)

[†]These authors contributed equally as last authors

Editors

Reviewing Editor

Rachel Denison

Boston University, Boston, United States of America

Senior Editor

Yanchao Bi

Beijing Normal University, Beijing, China

Reviewer #1 (Public review):**Summary:**

This paper presents a method for reconstructing videos from mouse visual cortex neuronal activity using a state-of-the-art dynamic neural encoding model. The authors achieve high-quality reconstructions of 10-second movies at 30 Hz from two-photon calcium imaging data, reporting a 2-fold increase in pixel-by-pixel correlation compared to previous methods. They identify key factors for successful reconstruction including the number of recorded neurons and model ensembling techniques.

Strengths:

- (1) A comprehensive technical approach combining state-of-the-art neural encoding models with gradient-based optimization for video reconstruction.
- (2) Thorough evaluation of reconstruction quality across different spatial and temporal frequencies using both natural videos and synthetic stimuli.
- (3) Detailed analysis of factors affecting reconstruction quality, including population size and model ensembling effects.
- (4) Clear methodology presentation with well-documented algorithms and reproducible code.
- (5) Potential applications for investigating visual processing phenomena like predictive coding and perceptual learning.

Weaknesses:

The main metric of success (pixel correlation) may not be the most meaningful measure of reconstruction quality:

High correlation may not capture perceptually relevant features.

Different stimuli producing similar neural responses could have low pixel correlations The paper doesn't fully justify why high pixel correlation is a valuable goal

Comparison to previous work (Yoshida et al.) has methodological concerns: Direct comparison of correlation values across different datasets may be misleading; Large

differences in the number of recorded neurons (10x more in the current study); Different stimulus types (dynamic vs static) make comparison difficult; No implementation of previous methods on the current dataset or vice versa.

Limited exploration of how the reconstruction method could provide insights into neural coding principles beyond demonstrating technical capability.

The claim that "stimulus reconstruction promises a more generalizable approach" (line 180) is not well supported with concrete examples or evidence.

The paper would benefit from addressing how the method handles cases where different stimuli produce similar neural responses, particularly for high-speed moving stimuli where phase differences might be lost in calcium imaging temporal resolution.

<https://doi.org/10.7554/eLife.105081.1.sa3>

Reviewer #2 (Public review):

This is an interesting study exploring methods for reconstructing visual stimuli from neural activity in the mouse visual cortex. Specifically, it uses a competition dataset (published in the Dynamic Sensorium benchmark study) and a recent winning model architecture (DNEM, dynamic neural encoding model) to recover visual information stored in ensembles of the mouse visual cortex.

This is a great project - the physiological data were measured at a single-cell resolution, the movies were reasonably naturalistic and representative of the real world, the study did not ignore important correlates such as eye position and pupil diameter, and of course, the reconstruction quality exceeded anything achieved by previous studies. Overall, it is great that teams are working towards exploring image reconstruction. Arguably, reconstruction may serve as an endgame method for examining the information content within neuronal ensembles - an alternative to training interminable numbers of supervised classifiers, as has been done in other studies. Put differently, if a reconstruction recovers a lot of visual features (maybe most of them), then it tells us a lot about what the visual brain is trying to do: to keep as much information as possible about the natural world in which its internal motor circuits may act consequently.

While we enjoyed reading the manuscript, we admit that the overall advance was in the range of those that one finds in a great machine learning conference proceedings paper. More specifically, we found no major technical flaws in the study, only a few potential major confounds (which should be addressable with new analyses), and the manuscript did not make claims that were not supported by its findings, yet the specific conceptual advance and significance seemed modest. Below, we will go through some of the claims, and ask about their potential significance.

(1) The study showed that it could achieve high-quality video reconstructions from mouse visual cortex activity using a neural encoding model (DNEM), recovering 10-second video sequences and approaching a two-fold improvement in pixel-by-pixel correlation over attempts. As a reader, I am left with the question: okay, does this mean that we should all switch to DNEM for our investigations of the mouse visual cortex? What makes this encoding model special? It is introduced as "a winning model of the Sensorium 2023 competition which achieved a score of 0.301... single-trial correlation between predicted and ground truth neuronal activity," but as someone who does not follow this competition (most eLife readers are not likely to do so, either), I do not know how to gauge my response. Is this impressive? What is the best achievable score, in theory, given data noise? Is the model inspired by the mouse brain in terms of mechanisms or architecture, or was it optimized to win the

competition by overfitting it to the nuances of the data set? Of course, I know that as a reader, I am invited to read the references, but the study would stand better on its own if clarified how its findings depended on this model.

(2) Along those lines, two major conclusions were that "critical for high-quality reconstructions are the number of neurons in the dataset and the use of model ensembling." If true, then these principles should be applicable to networks with different architectures. How well can they do with other network types?

(3) One major claim was that the quality of the reconstructions depended on the number of neurons in the dataset. There were approximately 8000 neurons recorded per mouse. The correlation difference between the reconstruction achieved by 1 neuron and 8000 neurons was ~ 0.2 . Is that a lot or a little? One might hypothesize that $\sim 7,999$ additional neurons could contribute more information, but perhaps, those neurons were redundant if their receptive fields were too close together or if they had the same orientation or spatiotemporal tuning. How correlated were these neurons in response to a given movie? Why did so many neurons offer such a limited increase in correlation?

(4) On a related note, the authors address the confound of RF location and extent. The study resorted to the use of a mask on the image during reconstruction, applied during training and evaluation (Line 87). The mask depends on pixels that contribute to the accurate prediction of neuronal activity. The problem for me is that it reads as if the RF/mask estimate was obtained during the very same process of reconstruction optimization, which could be considered a form of double-dipping (see the "Dead salmon" article, [https://doi.org/10.1016/S1053-8119\(09\)71202-9](https://doi.org/10.1016/S1053-8119(09)71202-9)). This could inflate the reconstruction estimate. My concern would be ameliorated if the mask was obtained using a held-out set of movies or image presentations; further, the mask should shift with eye position, if it indeed corresponded to the "collective receptive field of the neural population." Ideally, the team would also provide the characteristics of these putative RFs, such as their weight and spatial distribution, and whether they matched the biological receptive fields of the neurons (if measured independently).

(5) We appreciated the experiments testing the capacity of the reconstruction process, by using synthetic stimuli created under a Gaussian process in a noise-free way. But this further raised questions: what is the theoretical capability for the reconstruction of this processing pipeline, as a whole? Is 0.563 the best that one could achieve given the noisiness and/or neuron count of the Sensorium project? What if the team applied the pipeline to reconstruct the activity of a given artificial neural network's layer (e.g., some ResNet convolutional layer), using hidden units as proxies for neuronal calcium activity?

(6) As the authors mentioned, this reconstruction method provided a more accurate way to investigate how neurons process visual information. However, this method consisted of two parts: one was the state-of-the-art (SOTA) dynamic neural encoding model (DNEM), which predicts neuronal activity from the input video, and the other part reconstructed the video to produce a response similar to the predicted neuronal activity. Therefore, the reconstructed video was related to neuronal activity through an intermediate model (i.e., SOTA DNEM). If one observes a failure in reconstructing certain visual features of the video (for example, high-spatial frequency details), the reader does not know whether this failure was due to a lack of information in the neural code itself or a failure of the neuronal model to capture this information from the neural code (assuming a perfect reconstruction process). Could the authors address this by outlining the limitations of the SOTA DNEM encoding model and disentangling failures in the reconstruction from failures in the encoding model?

(7) The authors mentioned that a key factor in achieving high-quality reconstructions was model assembling. However, this averaging acts as a form of smoothing, which reduces the reconstruction's acuity and may limit the high-frequency content of the videos (as mentioned

in the manuscript). This averaging constrains the tool's capacity to assess how visual neurons process the low-frequency content of visual input. Perhaps the authors could elaborate on potential approaches to address this limitation, given the critical importance of high-frequency visual features for our visual perception.

<https://doi.org/10.7554/eLife.105081.1.sa2>

Reviewer #3 (Public review):

Summary:

This paper presents a method for reconstructing input videos shown to a mouse from the simultaneously recorded visual cortex activity (two-photon calcium imaging data). The publicly available experimental dataset is taken from a recent brain-encoding challenge, and the (publicly available) neural network model that serves to reconstruct the videos is the winning model from that challenge (by distinct authors). The present study applies gradient-based input optimization by backpropagating the brain-encoding error through this selected model (a method that has been proposed in the past, with other datasets). The main contribution of the paper is, therefore, the choice of applying this existing method to this specific dataset with this specific neural network model. The quantitative results appear to go beyond previous attempts at video input reconstruction (although measured with distinct datasets). The conclusions have potential practical interest for the field of brain decoding, and theoretical interest for possible future uses in functional brain exploration.

Strengths:

The authors use a validated optimization method on a recent large-scale dataset, with a state-of-the-art brain encoding model. The use of an ensemble of 7 distinct model instances (trained on distinct subsets of the dataset, with distinct random initializations) significantly improves the reconstructions. The exploration of the relation between reconstruction quality and the number of recorded neurons will be useful to those planning future experiments.

Weaknesses:

The main contribution is methodological, and the methodology combines pre-existing components without any new original components. The movie reconstructions include a learned "transparency mask" to concentrate on the most informative area of the frame; it is not clear how this choice impacts the comparison with prior experiments. Did they all employ this same strategy? If not, shouldn't the quantitative results also be reported without masking, for a fair comparison?

<https://doi.org/10.7554/eLife.105081.1.sa1>

Author response:

We thank the reviewers for their thorough review of our manuscript and their constructive feedback. We will address their comments and concerns in a point-by-point response at a later stage but would like to clarify some minor misunderstanding to not confuse any readers in the meantime.

- In regard to population ablation: When investigating the contribution of population size to reconstruction quality, we used 12.5, 25, 50 or 100% of the recorded neuronal population, which corresponds to ~1000/2000/4000/8000 neurons per animal. We did not produce reconstructions from only 1 neuron.

- In regard to the training of the transparency masks: The transparency masks were not produced using the same movies we reconstructed. We apologize for the lack of clarity on this point in the manuscript. We calculated the masks using an original model instance rather than a retrained instances used in the rest of the paper. Specifically, the masks were calculated using the original model instance ‘fold 1’ and data fold 1, which is it’s validation fold. In contrast, the model instances used in the paper for movie reconstruction were retrained while omitting the same validation fold across all instances (fold 0) and all the reconstructed movies in the paper are from data fold 0.
- In regard to reconstruction based on predicted activity: We always reconstructed the videos based on the true neural responses not the predicted neural response, with the exception of the Gaussian noise and drifting grating stimuli in Figure 4 and Supplementary Figure S2 where no recorded neural activity was available).

<https://doi.org/10.7554/eLife.105081.1.sa0>