

Factorized visual representations in the primate visual system and deep neural networks

Reviewed Preprint

v2 • June 5, 2024

Revised by authors

Reviewed Preprint

v1 • February 2, 2024

Jack W. Lindsey, Elias B. Issa 

Zuckerman Mind Brain Behavior Institute, Columbia University • Department of Neuroscience, Columbia University

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

Abstract

Object classification has been proposed as a principal objective of the primate ventral visual stream and has been used as an optimization target for deep neural network models (DNNs) of the visual system. However, visual brain areas represent many different types of information, and optimizing for classification of object identity alone does not constrain how other information may be encoded in visual representations. Information about different scene parameters may be discarded altogether (“invariance”), represented in non-interfering subspaces of population activity (“factorization”) or encoded in an entangled fashion. In this work, we provide evidence that factorization is a normative principle of biological visual representations. In the monkey ventral visual hierarchy, we found that factorization of object pose and background information from object identity increased in higher-level regions and strongly contributed to improving object identity decoding performance. We then conducted a large-scale analysis of factorization of individual scene parameters – lighting, background, camera viewpoint, and object pose – in a diverse library of DNN models of the visual system. Models which best matched neural, fMRI and behavioral data from both monkeys and humans across 12 datasets tended to be those which factorized scene parameters most strongly. Notably, invariance to these parameters was not as consistently associated with matches to neural and behavioral data, suggesting that maintaining non-class information in factorized activity subspaces is often preferred to dropping it altogether. Thus, we propose that factorization of visual scene information is a widely used strategy in brains and DNN models thereof.

eLife assessment

The study makes a **valuable** empirical contribution to our understanding of visual processing in primates and deep neural networks, with a specific focus on the concept of factorization. The analyses provide **convincing** evidence that high factorization scores are correlated with neural predictivity. This work will be of interest to systems neuroscientists studying vision and could inspire further research that ultimately may lead to better models of or a better understanding of the brain.

<https://doi.org/10.7554/eLife.91685.2.sa2>

Introduction

Artificial deep neural networks (DNNs) are the most predictive models of neural responses to images in the primate high-level visual cortex^{1,2}. Many studies have reported that DNNs trained to perform image classification produce internal feature representations broadly similar to those in areas V4 and IT of the primate cortex, and that this similarity tends to be greater in models with better classification performance³. However, it remains opaque what aspects of the representations of these more performant models drive them to better match neural data. Moreover, beyond a certain threshold level of object classification performance, further improvement fails to produce a concomitant improvement in predicting primate neural responses^{2,4,5}. This weakening trend motivates finding new normative principles, besides object classification ability, that push models to better match primate visual representations.

One strategy for achieving high object classification performance is to form neural representations that discard some (are tolerant to) or all (are invariant to) information besides object class. Invariance in neural representations is in some sense a zero-sum strategy: building invariance to some parameters improves the ability to decode others. We also note that our use of “invariance” in this context refers to invariance in neural representations, rather than behavioral or perceptual invariance⁶. However, high-level cortical neurons in the primate ventral visual stream are known to simultaneously encode many forms of information about visual input besides object identity, such as object pose^{7–10}. In this work, we seek to characterize how the brain simultaneously represents different forms of information.

In particular, we introduce methods to quantify the relationships between different types of visual information in a population code (e.g., object pose vs. camera viewpoint), and specifically the degree to which different forms of information are “factorized.” Intuitively, if the variance driven by one parameter is encoded along orthogonal dimensions of population activity space compared to the variance driven by other scene parameters, we say that this representation is factorized. We note that our definition of factorization is closely related to the existing concept of manifold disentanglement^{6,11} and can be seen as a generalization of disentanglement to high-dimensional visual scene parameters like object pose. Factorization can enable simultaneous decoding of many parameters at once, supporting diverse visually guided behaviors (e.g., spatial navigation, object manipulation or object classification)¹².

Using existing neural datasets, we found that both factorization of and invariance to object category and position information increase across the macaque ventral visual cortical hierarchy. Next, we leveraged the flexibility afforded by *in silico* models of visual representations to probe different forms of factorization and invariance in more detail, focusing on several scene parameters of interest: background content, lighting conditions, object pose, and camera viewpoint. Across a broad library of DNN models that varied in their architecture and training objectives, we found that factorization of all of the above scene parameters in DNN feature representations was positively correlated with models’ matches to neural and behavioral data. Interestingly, while neural invariance to some scene parameters (background scene and lighting conditions) predicted neural fits, invariance to others (object pose and camera viewpoint) did not. Our results generalized across both monkey and human datasets using different measures (neural spiking, fMRI, and behavior; 12 datasets total) and could not be accounted for by models’ classification performance. Thus, we suggest that factorized encoding of multiple behaviorally-relevant scene variables is an important consideration, alongside other desiderata such as classification performance, in building more brain-like models of vision.

Results

Disentangling object identity manifolds in neural population responses can be achieved by qualitatively different strategies. These include building invariance of responses to non-identity scene parameters (or, more realistically, partial invariance⁶) and/or factorizing non-identity-driven response variance into isolated (factorized) subspaces (**Figure 1A**, left vs. center panels, cylindrical/spherical shaded regions represent object manifolds). Both strategies maintain an “identity subspace” in which object manifolds are linearly separable. In a non-invariant, non-factorized representation, other variables like camera viewpoint also drive variance within the identity subspace, “entangling” the representations of the two variables (**Figure 1A**, right; viewpoint driven variance is mainly in identity subspace, orange flat shaded region).

To formalize these different representational strategies, we introduced measures of factorization and invariance to scene parameters in neural population responses (**Figure 1B**; see **Equations 2–4** in **Methods**). Concretely, invariance to a scene variable (e.g. object motion) is computed by measuring the degree to which varying that parameter alone changes neural responses, relative to the changes induced by varying other parameters (lower relative influence on neural activity corresponds to higher invariance, or tolerance, to that parameter). Factorization is computed by identifying the axes in neural population activity space that are influenced by varying the parameter of interest and assessing how much it overlaps the axes influenced by other parameters (“ a ” in **Figure 1B,C**; lower overlap corresponds to higher factorization). We quantified this overlap in two different ways (“PCA based” and “covariance based” factorization, corresponding to **Equations 2** and **4** in **Methods**) which produced similar results when compared in subsequent analyses (unless otherwise noted, factorization scores will generally refer to the PCA based method, and the covariance method is shown in **Figures 5–7** for comparison). Intuitively, a neural population in which one neural subpopulation encodes object identity and another separate subpopulation encodes object position exhibits a high degree of factorization of those two parameters (however, note that factorization may also be achieved by neural populations with mixed selectivity in which the “subpopulations” correspond to subspaces, or independent orthogonal linear projections, of neural activity space rather than physical subpopulations). Though the example presented in **Figure 1** focused on factorization of and invariance to object identity versus non-identity variables, we stress that our definitions can be applied to any scene variables of interest. Furthermore, we presented a simplified visual depiction of the geometry within each scene variable subspace in **Figure 1**. We emphasize that our factorization metric does not require a particular geometry within a variable’s subspace, whether parallel linearly ordered coding of viewpoint as in the cylindrical class manifolds shown in **Figure 1A** and **1B**, or a more complex geometry where there is a lack of parallelism and/or a more nonlinear layout.

While factorization and invariance are not mutually exclusive representational strategies, they are qualitatively different. Factorization, unlike invariance, has the potential to enable the simultaneous representation of multiple scene parameters in a decodable fashion. Intuitively, factorization increases with higher dimensionality as this decreases overlap, all other things being equal (in the limit, the angle between points will approach 90° or a fully orthogonal code in high dimensions), and for a given finite, fixed dimension, factorization is mainly driven by the angle between this dimension and the other variable subspaces which measures the degree of contamination (**Figure 1C**; square vs. parallelogram). In a simulation, we found that the extent to which the variables of interest were represented in a factorized way (i.e., along orthogonal axes, rather than correlated axes) influenced the ability of a linear discriminator to successfully decode both variables in a generalizable fashion from a few training samples (**Figure 1C**).

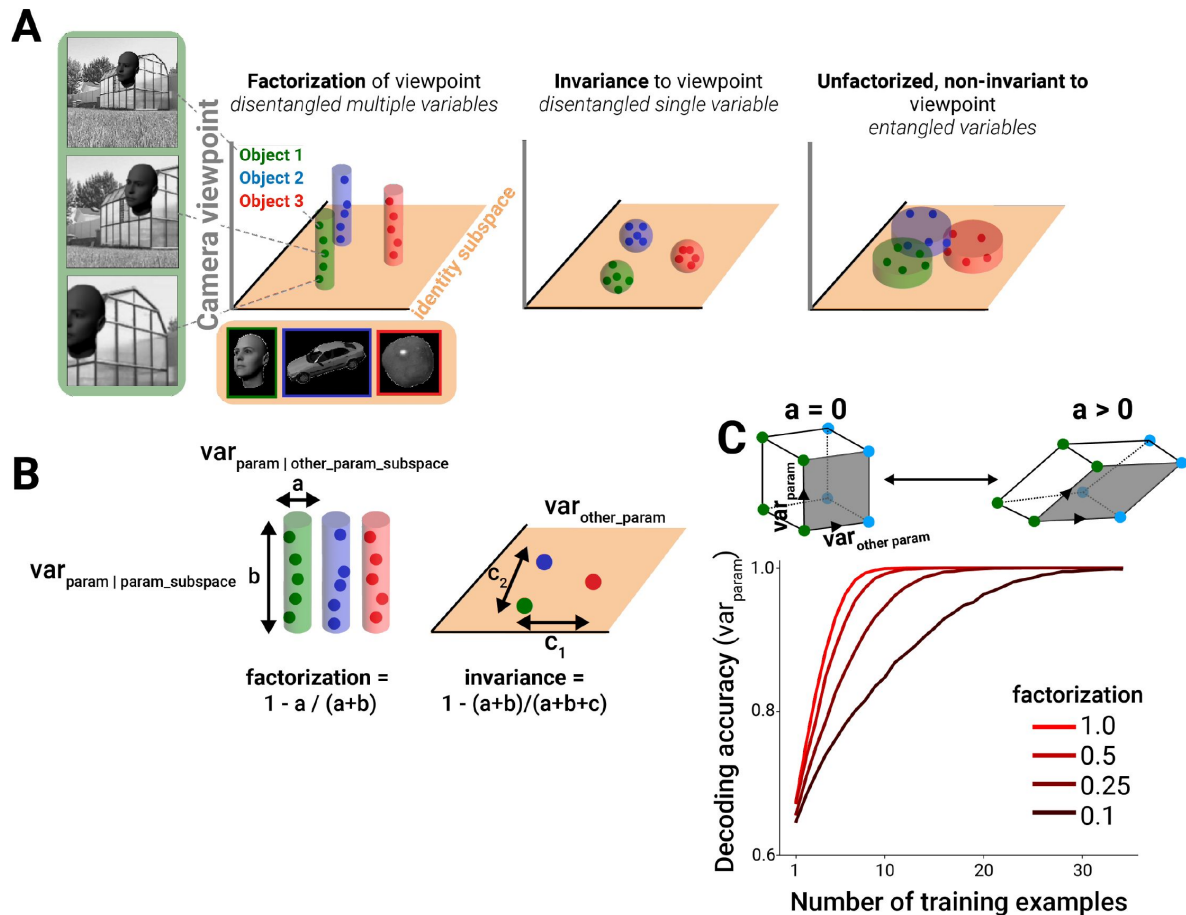


Figure 1.

Framework for quantifying factorization in neural and model representations.

(A) A subspace for encoding a variable, for example object identity, in a linearly separable manner can be achieved by becoming invariant to non-class variables (compact spheres, middle column, where the volume of the sphere corresponds to the degree of neural invariance, or tolerance, for non-class variables; colored dots represent example images within each class) and/or by encoding variance induced by non-identity variables in orthogonal neural axes to the identity subspace (extended cylinders, left column). Only the factorization strategy simultaneously represents multiple variables in a disentangled fashion. A code that is sensitive to non-identity parameters within the identity subspace corrupts the ability to decode identity (right column) (identity subspace denoted by orange plane). (B) Variance across images within a class can be measured in two different linear subspaces: that containing the majority of variance for all other parameters (a , "other_param_subspace") and that containing the majority of the variance for that parameter (b , "param_subspace"). Factorization is defined as the fraction of parameter-induced variance that avoids the other-parameter subspace (left). By contrast, invariance to the parameter of interest is computed by comparing the overall parameter-induced variance to the variance in response to other parameters (c , "var_other_param")(right). (C) In a simulation of coding strategies for two binary variables out of 10 total dimensions that are varying (see **Methods**), a decrease in orthogonality of the relationship between the encoding of the two variables (alignment $a > 0$, or going from a square to a parallelogram geometry), despite maintaining linear separability of variables, results in poor classifier performance in the few training-samples regime when i.i.d. Gaussian noise is present in the data samples (only 3 of 10 dimensions used in simulation are shown).

Given the theoretically desirable properties of factorized representations, we next asked whether such representations are observed in neural data, and how much factorization contributes empirically to downstream decoding performance in real data. Specifically, we took advantage of an existing dataset in which the tested images independently varied object identity versus object pose plus background context¹³. We found that both V4 and IT responses exhibited more significant factorization of object identity information from non-identity information than a shuffle control (which accounts for effects on factorization due to dimensionality of these regions) (**Figure S1**; see **Methods**). Furthermore, the degree of factorization increased from V4 to IT (**Figure 2A**). Consistent with prior studies, we also found that invariance to non-identity information increased from V4 to IT in our analysis (**Figure 2A**, right, solid lines)¹⁴. Invariance to non-identity information was even more pronounced when measured in the subspace of population activity capturing the bulk (90%) of identity-driven variance, as a consequence of increased factorization of identity from non-identity information (**Figure 2A**, right, dashed lines).

To illustrate the beneficial effect of factorization on decoding performance, we performed a statistical lesion experiment that precisely targeted this aspect of representational geometry. Specifically, we analyzed a transformed neural representation obtained by rotating the population data so that inter-class variance more strongly overlapped with the principal components of the intra-class variance in the data (see **Equation 1** in **Methods**). Note that this transformation, designed to decrease factorization, acts on the angle between latent variable subspaces. The applied linear basis rotation leaves all other activity statistics completely intact (such as mean neural firing rates, covariance structure of the population, and its invariance to non-class variables) yet has the effect of strongly reducing object identity decoding performance in both V4 and IT (**Figure 2B**). Our analysis shows that maintaining invariance alone in the neural population code was insufficient to account for a large fraction of decoding performance in high-level visual cortex; factorization of non-identity variables is key to the decoding performance achieved by V4 and IT representations.

We next asked whether factorization is found in deep neural network (DNN) model representations and whether this novel, heretofore unconsidered metric is a strong indicator of more brainlike models. When working with computational models, we have the liberty to test an arbitrary number of stimuli; therefore, we could independently vary multiple scene parameters at sufficient scale to enable computing factorization and invariance for each, and we explored factorization in DNN model representations in more depth than previously measured in existing neural experiments. To gain insight back into neural representations, we also assessed the ability of each model to predict separately collected neural and behavioral data. In this fashion, we may indirectly assess the relative significance of geometric properties like factorization and invariance to biological visual representations – if, for instance, models with more factorized representations consistently match neural data more closely, we may infer that those neural representations likely exhibit factorization themselves (**Figure 3**). To measure factorization, invariance, and decoding properties of DNN models, we generated an augmented image set, based on the images used in the previous dataset (**Figure 2**), in which we independently varied the foreground object identity, foreground object pose, background identity, scene lighting, and 2D scene viewpoint. Specifically for each base image from the original dataset, we generated sets of images that varied exactly one of the above scene parameters while keeping the others constant, allowing us to measure the variance induced by each parameter relative to the variance across all scene parameters (**Figure 3**, top left; 100 base scenes and 10 transformed images for each source of variation). We presented this large image dataset to models (4000 images total) to assess the relative degree of representational factorization of and invariance to each scene parameter. We conducted this analysis across a broad range of DNNs varying in architecture and objective as well as other implementational choices to obtain the widest possible range of DNN representations for testing

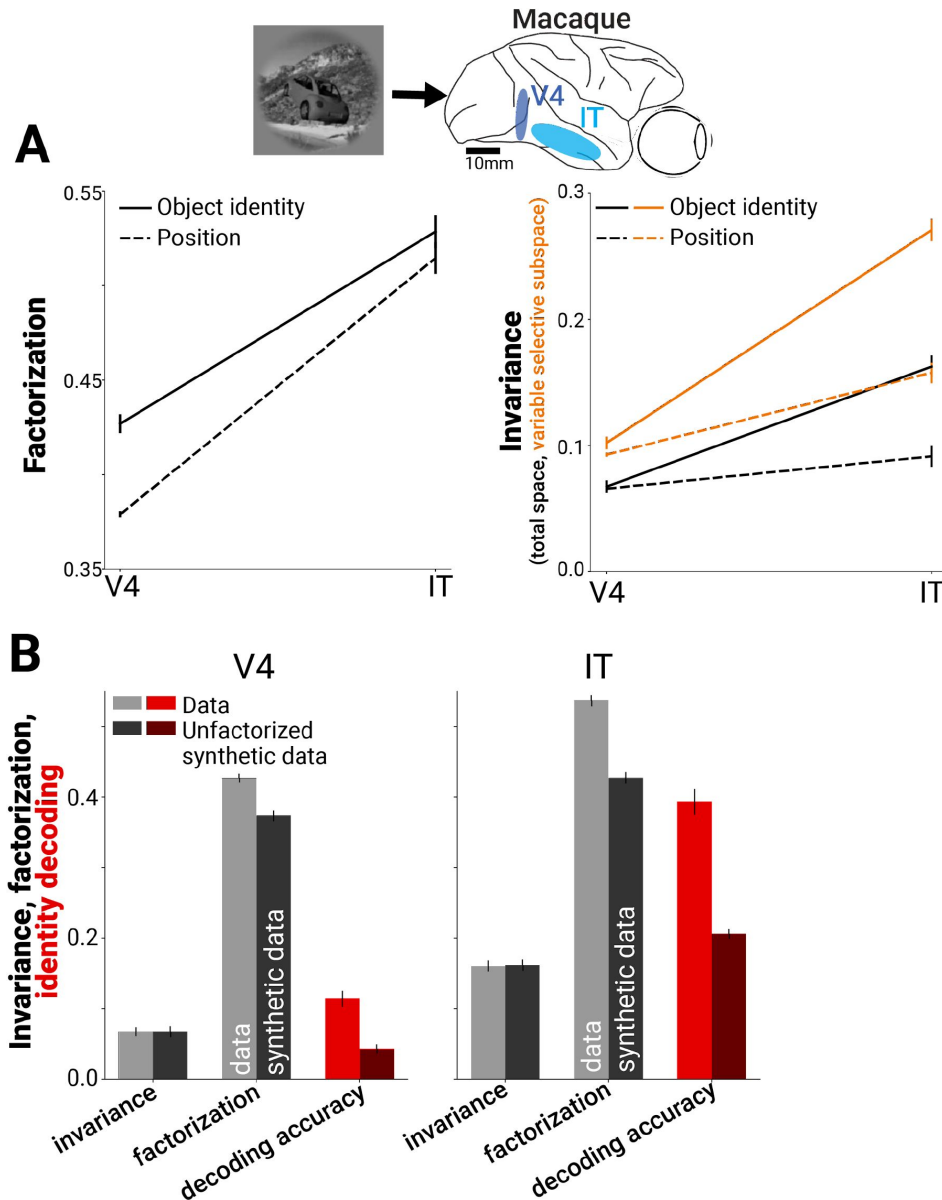


Figure 2.

Benefit of factorization to neural decoding in macaque V4 and IT.

(A) Factorization of object identity and position increased from macaque V4 to IT (PCA based factorization, see **Methods**; dataset E1 – multiunit activity in macaque visual cortex)(left). Like factorization, invariance also increased from V4 to IT (note, “identity” refers to invariance to all non-identity position factors, solid black line)(right). Combined with increased factorization of the remaining variance, this led to higher invariance within the variable’s subspace (orange lines), representing a neural subspace for identity information with invariance to nuisance parameters which decoders can target for read-out. (B) An experiment to test the importance of factorization for supporting object class decoding performance in neural responses. We applied a transformation to the neural data (linear basis rotation) that rotated the relative positions of mean responses to object classes without changing the relative proportion of within- vs. between-class variance (**Equation 1** in **Methods**). This transformation preserved invariance to non-class factors (leftmost pair of bars in each plot), while decreasing factorization of class information from non-class factors (center pair of bars in each plot). Concurrently, it had the effect of significantly reducing object class decoding performance (light vs. dark red bars in each plot, chance = 1/64; n=128 multi-unit sites in V4 and 128 in IT).

our hypothesis. These included models using supervised training for object classification^{15,16}, contrastive self-supervised training^{17,18}, and self-supervised models trained using auxiliary objective functions^{19–22} (see **Methods** and **Table S2**).

First, we asked whether, in the course of training, DNN models develop factorized representations at all. We found that the final layers of trained networks exhibited consistent increases in factorization of all tested scene parameters relative to a randomly initialized (untrained) baseline with the same architecture (**Figure 4A**, top row, rightward shift relative to black cross, a randomly initialized ResNet-50). By contrast, training DNNs produced mixed effects on invariance, typically increasing it for background and lighting but reducing it for object pose and camera viewpoint (**Figure 4A**, bottom row, leftward shift relative to black cross for left two panels). Moreover, we found that the degree of factorization in models correlated with the degree to which they predicted neural activity for single-unit IT data (**Figure 4A**, top row), which can be seen as correlative evidence that neural representations in IT exhibit factorization of all scene variables tested. Interestingly, we saw a different pattern for representational invariance to a scene parameter. Invariance showed mixed correlations with neural predictivity (**Figure 4A**, bottom row), suggesting that IT neural representations build invariance to some scene information (background and lighting) but not to others (object pose and observer viewpoint). Similar effects were observed when we assessed correlations between these metrics and fits to human behavioral data rather than macaque neural data (**Figure 4B**).

To assess the robustness of these findings to choice of images and brain regions used in an experiment, we conducted the same analyses across a large and diverse set of previously collected neural and behavioral datasets, from different primate species and visual regions (6 macaque datasets^{13,23,24}; two V4, two IT, and two behavior; 6 human datasets^{24–26}; two V4, two HVC, and two behavior; **Table S1**). Consistently, increased factorization of scene parameters in model representations correlated with models being more predictive of neural spiking responses, voxel BOLD signal, and behavioral responses to images (**Figure 5A**, black bars; see **Figure S2** for scatter plots across all datasets). Although invariance to appearance factors (background identity and scene lighting) correlated with more brainlike models, invariance for spatial transforms (object pose and camera viewpoint) consistently did not (zero or negative correlation values; **Figure 5C**, red and green open circles). Our results were preserved when we re-ran the analyses using only the subset of models with the identical ResNet-50 architecture (**Figure S3**) or when we evaluated model predictivity using representational dissimilarity matrices of the population (RDMs) instead of linear regression (encoding) fits of individual neurons or voxels (**Figure S4**). Furthermore, the main finding of a positive correlation between factorization and neural predictivity was robust to the particular choice of PCA threshold we used to quantify factorization (**Figure S5**). We found similar results using a covariance based method for computing factorization that does not have any free parameters (**Figure 5C** faded filled circles; see Equations 4 in **Methods**).

Finally, we tested whether our results generalized across the particular image set used for computing the model factorization scores in the first place. Here, instead of relying on our synthetically generated images, where each scene parameter was directly controlled, we re-computed factorization from two types of relatively unconstrained natural movies, one where the observer moves in an urban environment (approximates camera viewpoint changes)²⁷ and another where objects move in front of a fairly stationary observer (approximates object pose changes)²⁸. Similar to the result found for factorization measured using augmentations of synthetic images, factorization of frame-by-frame variance (local in time, presumably dominated by either observer or camera motion; see **Methods**) from other sources of variance across natural movies (non-local in time) was correlated with improved neural predictivity in both macaque and human data while invariance to local frame-by-frame differences was not (**Figure 5B**; black versus gray bars). Thus, we have shown that a main finding – the importance of object pose and camera viewpoint factorization for achieving brainlike representations – holds across types of

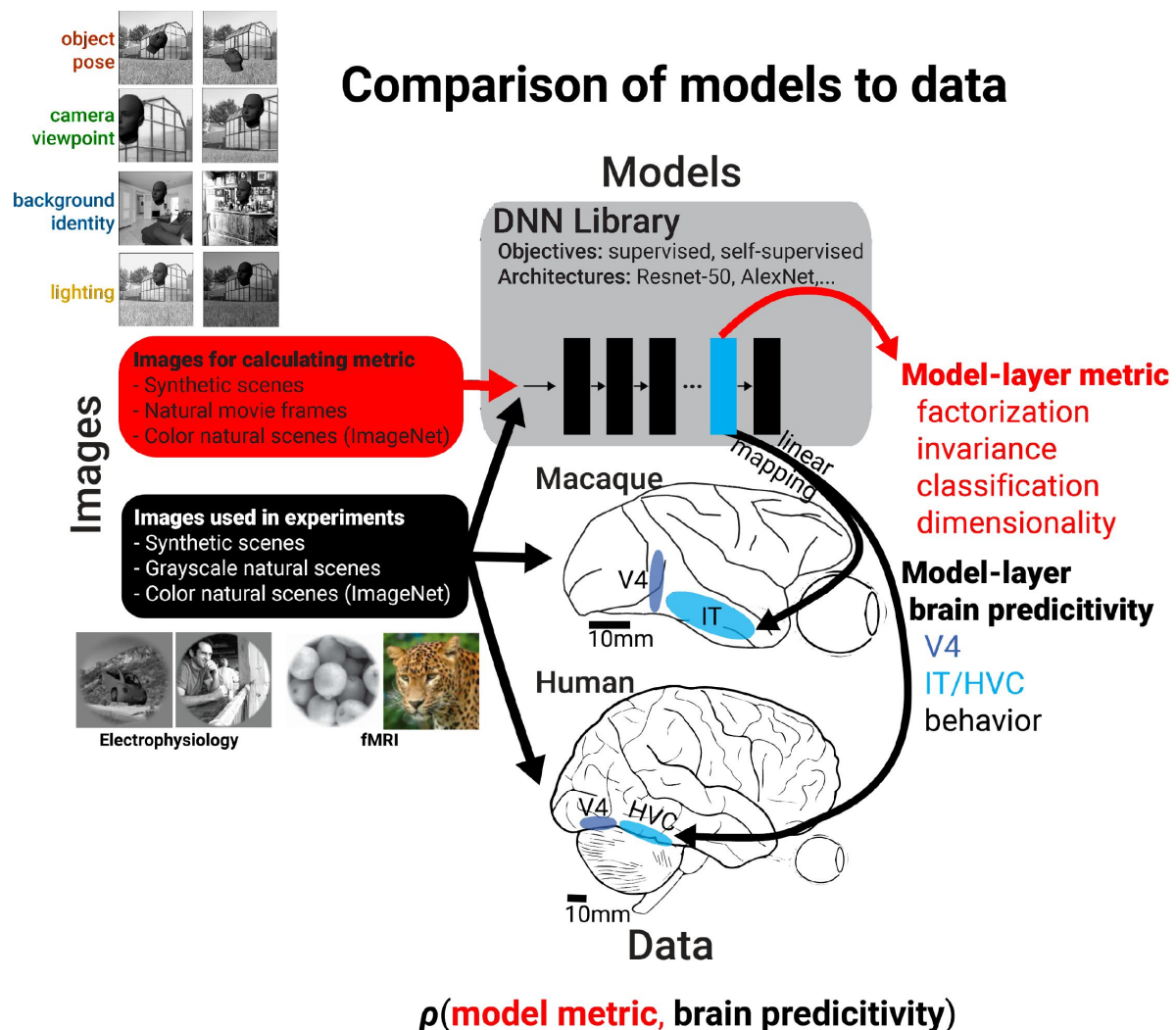


Figure 3.

Measurement of factorization in DNN models and comparison to brain data.

Schematic showing how metaanalysis on models and brain data was conducted by first computing various representational metrics on models and then measuring a model's predictive power across a variety of datasets. For computing the representational metrics of factorization and invariance to a scene parameter, variance in model responses was induced by individually varying each of four scene parameters ($n=10$ parameter levels) for each base scene ($n=100$ base scenes) (see images at top left). The combination of model-layer metric and model-layer dataset predictivity for a choice of model, layer, metric, and dataset specifies the coordinates of a single dot on the scatter plots in **Figures 4** and **7**, and the across model correlation coefficient between a particular representational metric and neural predictivity for a dataset summarizes the potential importance of the metric in producing more brainlike models (see **Figures 5** & **6**).

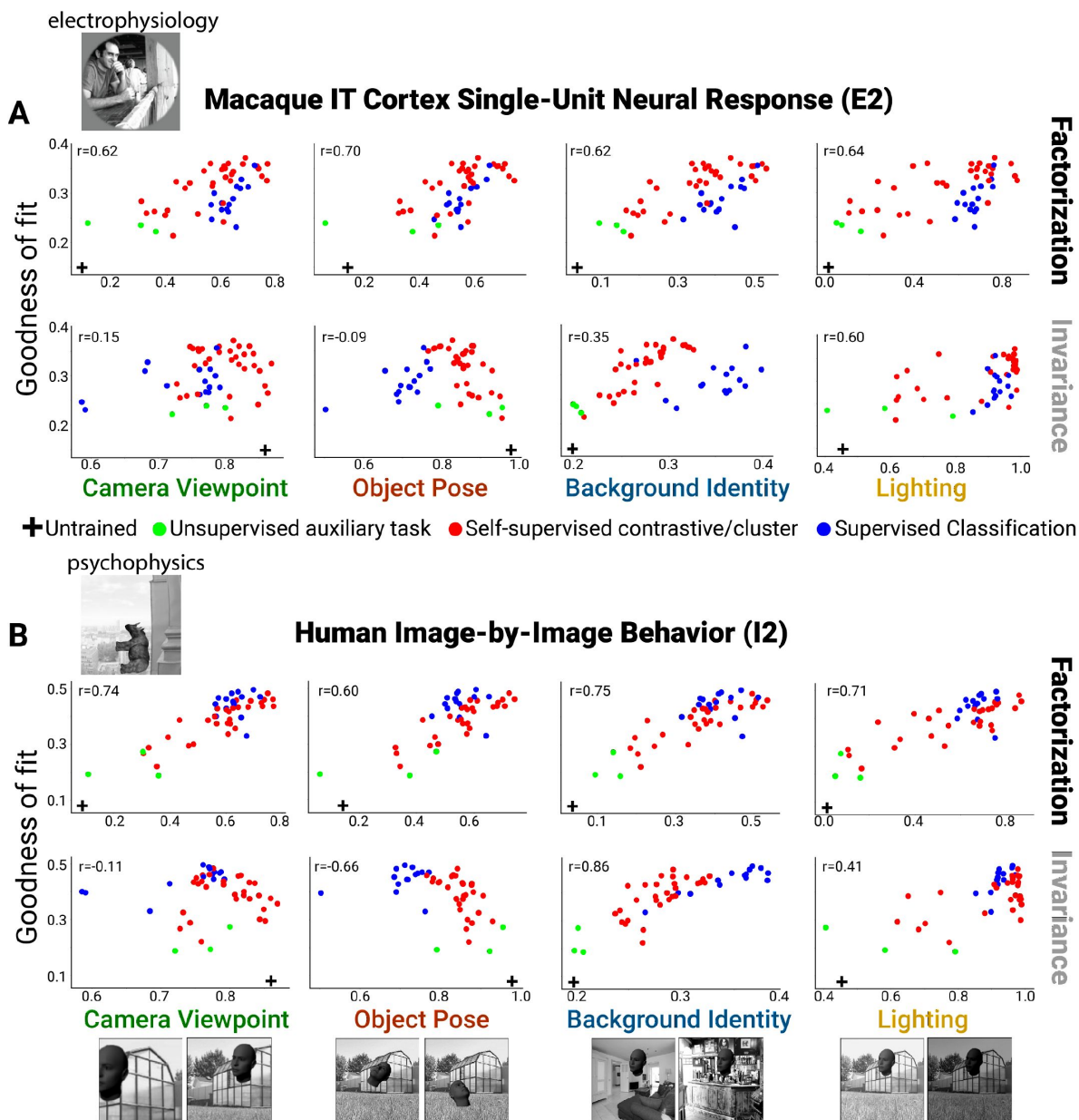


Figure 4.

Neural and behavioral predictivity of models versus their factorization and invariance properties.

(A) Scatter plots for example neural dataset (IT single-units, macaque E2 dataset) showing the correlation between a model's predictive power as an encoding model for IT neural data versus a model's ability to factorize or become invariant to different scene parameters (each dot is a different model, using each model's penultimate layer). Note that factorization (PCA based, see **Methods**) in trained models is consistently higher than that for an untrained, randomly initialized Resnet-50 DNN architecture (rightward shift relative to black cross). Invariance to background and lighting but not to object pose and viewpoint increased in trained models relative to the untrained control (rightward versus leftward shift relative to black cross). (B) Same as (A) except for human behavior performance patterns across images (human I2 dataset). Increasing scene parameter factorization in models generally correlated with better neural predictivity (top row). A noticeable drop in neural predictivity was seen for high levels of invariance to object pose (bottom row, second panel).

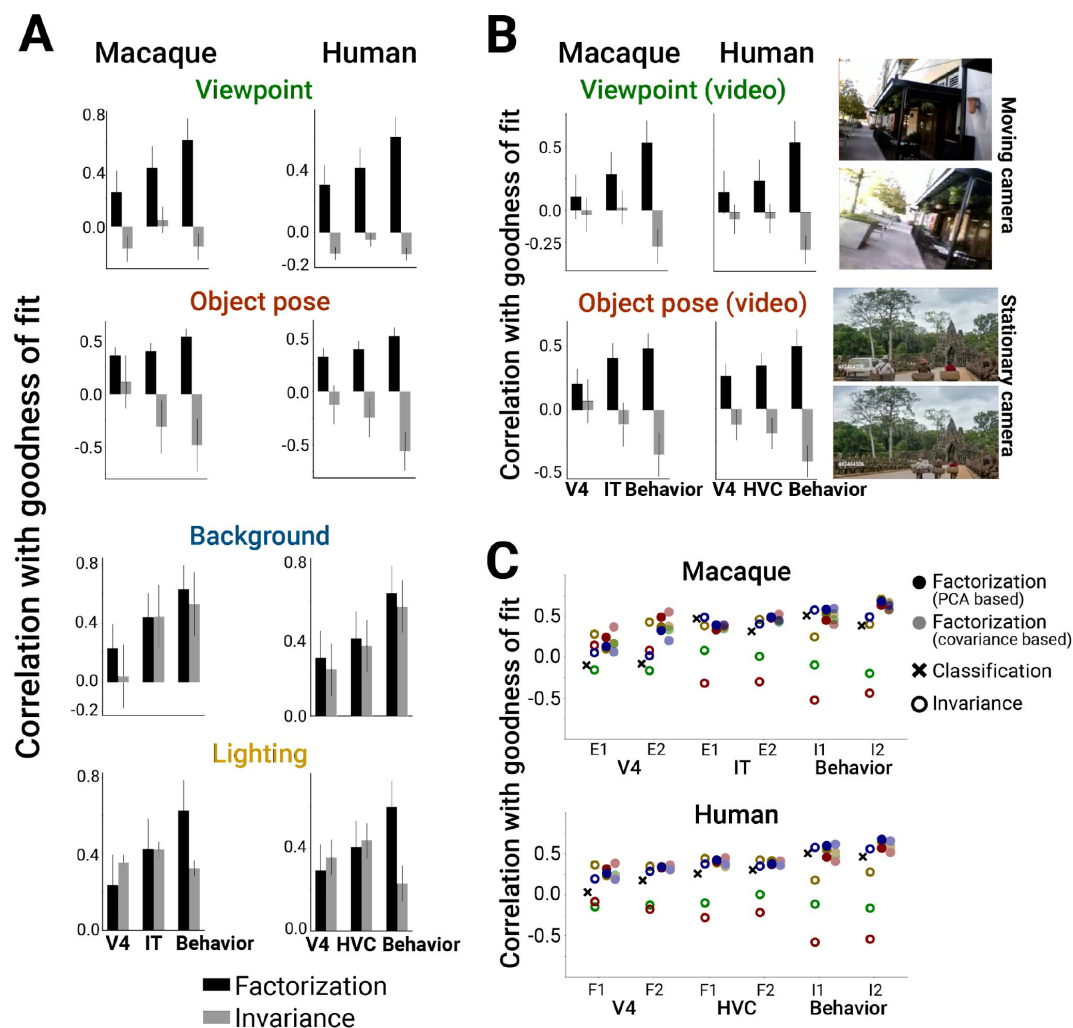


Figure 5.

Scene parameter factorization correlates with more brainlike DNN models.

(A) Factorization of scene parameters in model representations computed using the PCA based method consistently correlated with a model being more brainlike across multiple independent datasets measuring monkey neurons, human fMRI voxels, or behavioral performance in both macaques and humans (left vs. right column)(black bars). By contrast, increased invariance to camera viewpoint or object pose was not indicative of brainlike models (gray bars). In all cases, model representational metric and neural predictivity score were computed by averaging scores across the last 5 model layers. (B) Instead of computing factorization scores using our synthetic images (Figure 3, top left), recomputing camera viewpoint or object pose factorization from natural movie datasets that primarily contained camera or object motion, respectively, gave similar results for predicting which model representations would be more brainlike (right: example movie frames; also see **Methods**). Error bars in (A) & (B) are standard deviations over bootstrapped resampling of the models. (C) Summary of the results from (A) across datasets (x-axis) for invariance (open symbols) versus factorization (closed symbols)(for reference, 'x' symbols indicate predictive power when using model classification performance). Results using a comparable, alternative method for computing factorization (covariance based, Equation 4 in **Methods**; light closed symbols) are shown adjacent to the original factorization metric (PCA based, Equation 2 in **Methods**; dark closed symbols).

brain signal (spiking vs. BOLD), species (monkey vs. human), cortical brain areas (V4 vs. IT), images for testing in experiments (synthetic, grayscale vs. natural, color), and image sets for computing the metric (synthetic images vs. natural movies).

Our analysis of DNN models provides strong evidence that greater factorization of a variety of scene variables is consistently associated with a stronger match to neural and behavioral data. Prior work had identified a similar correlation between object classification performance (measured fitting a decoder for object class using model representations) and fidelity to neural data^{3,29,30}. A priori, it is possible that the correlations we have demonstrated between scene parameter factorization and neural fit can be entirely captured by the known correlation between classification performance and neural fits^{2,3,30}, as factorization and classification may themselves be correlated. However, we found that factorization scores significantly boosted cross-validated predictive power of neural/behavioral fit performance compared to simply using object classification alone, and factorization boosted predictive power as much if not slightly more when using RDMs instead of linear regression fits to quantify the match to the brain/behavior (**Figure 6**). Thus, considering factorization in addition to object classification performance improves upon our prior understanding of the properties of more brainlike models (**Figure 7**).

Discussion

Object classification, which has been proposed as a normative principle for the function of the ventral visual stream, can be supported by qualitatively different representational geometries^{3,29,30}. These include representations that are completely invariant to non-class information^{30,31} and representations that retain a high-dimensional but factorized encoding of non-class information, which disentangles the representation of multiple variables (**Figure 1A**). Here, we presented evidence that factorization of non-class information is an important strategy used, alongside invariance, by the high-level visual cortex (**Figure 2**) and by DNNs that are predictive of primate neural and behavioral data (**Figures 4 & 5**).

Prior work has indicated that building representations that support object classification performance and representations that preserve high-dimensional information about natural images are both important principles of the primate visual system^{1,32} (though see Conwell et al.³³). Critically, our results cannot be accounted for by classification performance or dimensionality alone (**Figure 6**, gray and pink bars); that is, the relationship between factorization and matches to neural data was not entirely mediated by classification or dimensionality. That said, we do not regard factorization and dimensionality, or factorization and object classification performance, as mutually exclusive hypotheses for useful principles of visual representations. Indeed, high-dimensional representations could be regarded as a means to facilitate factorization, and likewise factorized representations can better support classification (**Figure 1C**).

Our notion of factorization is related to, but distinct from, several other concepts in the literature. Many prior studies in machine learning have considered the notion of disentanglement, often defined as the problem of inferring independent factors responsible for generating the observed data^{34–36}. One prior study notably found that machine learning models designed to infer disentangled representations of visual data displayed single-unit responses that resembled those of individual neurons in macaque IT³⁷. Our definition of factorization is more flexible, requiring only that independent factors be encoded in orthogonal subspaces, rather than by distinct individual neurons. Moreover, our definition applies to generative factors, such as camera viewpoint or object pose, that are multidimensional and context dependent. Factorization is also related to a measure of “abstraction” in representational geometry introduced in a recent line of work^{38,39}, which is observed to emerge in trained neural networks^{12,40}. In these studies, an abstract representation is defined as one in which variables are encoded and can be decoded in

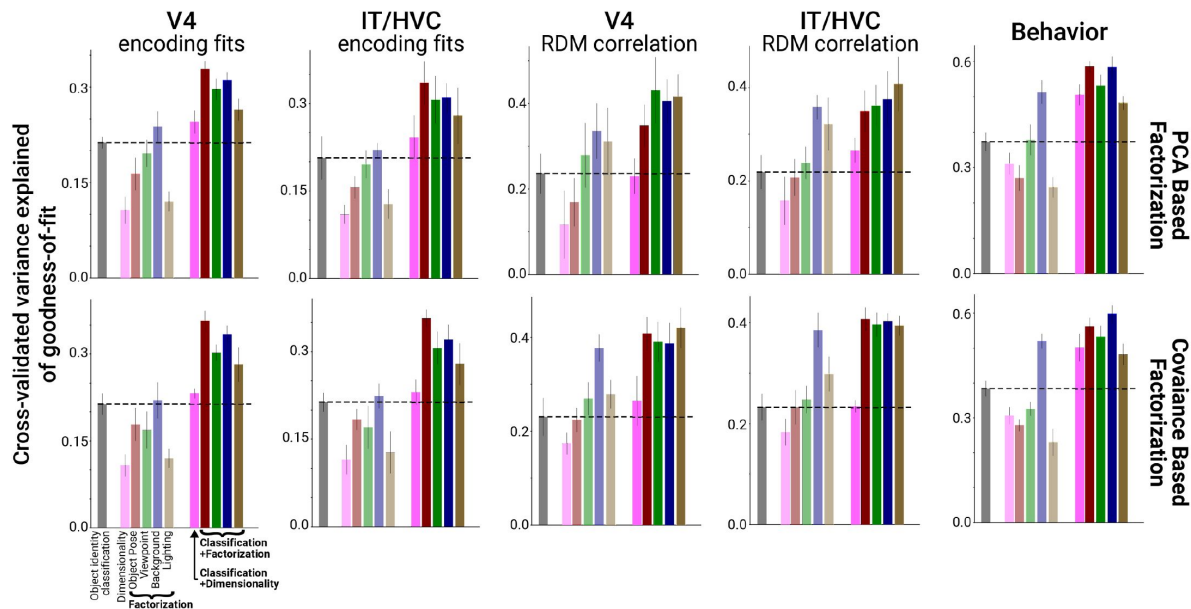


Figure 6.

Scene parameter factorization combined with object identity classification improves correlations with neural predictivity.

Average across datasets of brain predictivity of classification (faded black bar), dimensionality (faded pink bar), and factorization (remaining faded colored bars) in a model representation. Linearly combining factorization with classification in a regression model (unfaded bars at right) produced significant improvements in predicting the most brainlike models (performance cross-validated across models and averaged across datasets, $n=4$ datasets for each of V4, IT/HVC and behavior). The boost from factorization in predicting the most brainlike models was not observed for neural and fMRI data when combining classification with a model's overall dimensionality (solid pink bars; compare to black dashed line for brain predictivity when using classification alone). Results are shown for both the PCA based and covariance based factorization metric (top versus bottom row). Error bars are standard deviations over bootstrapped resampling of the models.

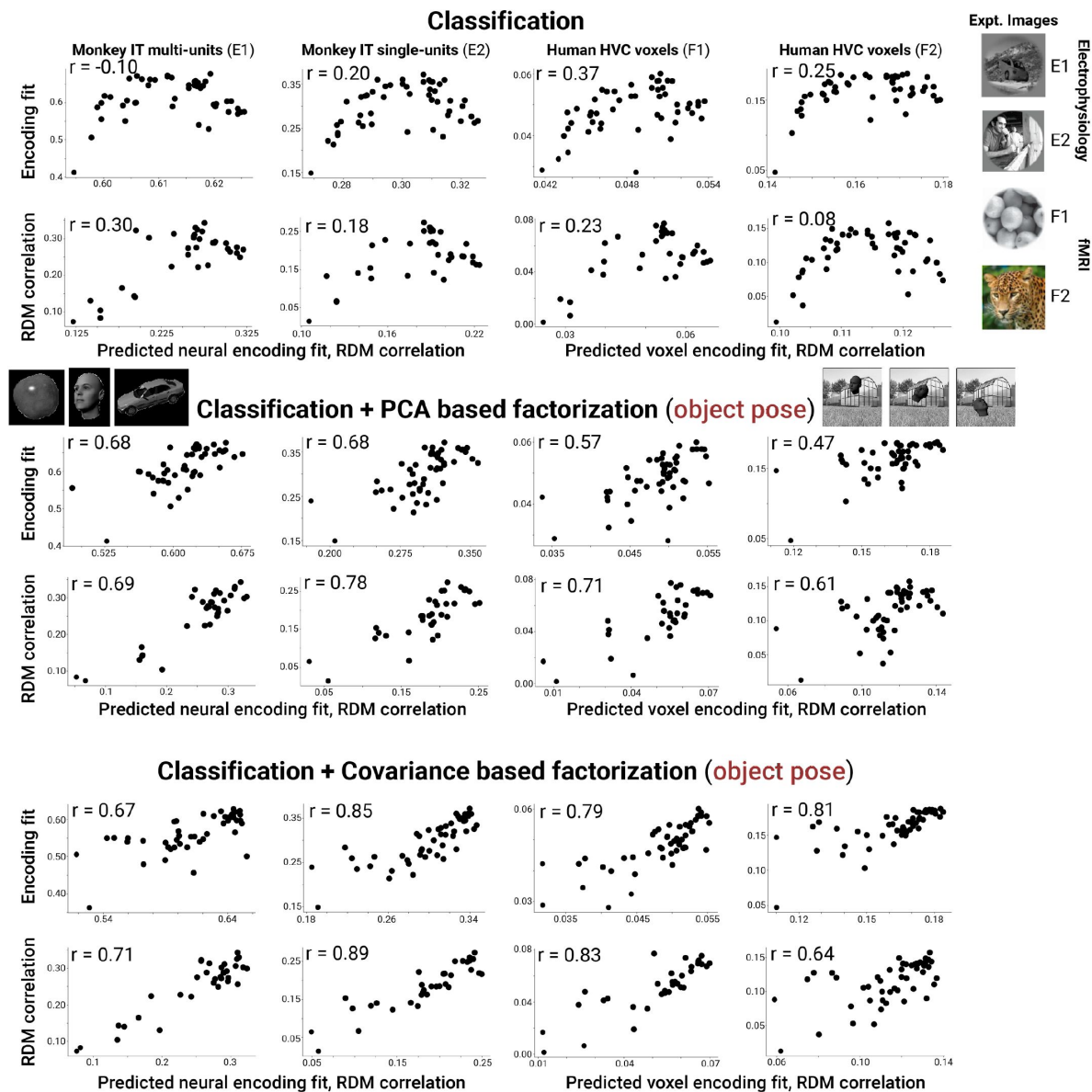


Figure 7.

Combining classification performance with object pose factorization improves predictions of the most brainlike models on IT/HVC data.

Example scatter plots for neural and fMRI datasets (macaque E1 & E2, IT multi-units & single-units; human F1 & F2, fMRI voxels) showing a saturating and sometimes reversing trend in neural (voxel) predictivity for models that are increasingly good at classification (top row). This saturating/reversing trend is no longer present when adding object pose factorization to classification as a combined, predictive metric for brainlikeness of a model (middle and bottom rows). The x-axis of each plot indicates the predicted encoding fit or RDM correlation after fitting a linear regression model with the indicated metrics as input (either classification or classification + factorization).

a consistent fashion regardless of the values of other variables. A fully factorized representation should be highly abstract according to this definition, though factorization emphasizes the geometric properties of the population representation while these studies emphasize the consequences for decoding performance in training downstream linear read-outs. Relatedly, another recent study found that orthogonal encoding of class and non-class information is one of several factors that determines few-shot classification performance⁴¹. Our work can be seen as complementary to work on representational straightening of natural movie trajectories in the population space⁴². This work suggested that visual representations maintain a locally linear code of latent variables like camera viewpoint, while our work focused on the global arrangement of the linear subspaces affected by different variables (e.g., overall coding of camera viewpoint driven variance versus sources of variance from other scene variables in a movie). Local straightening of natural movies was found to be important for early visual cortex neural responses but not necessarily for high-level visual cortex⁴³ – where the present work suggests factorization may play a role.

Our work has several limitations. First, our analysis is primarily correlative. Going forward, we suggest that factorization could prove to be a useful objective function for optimizing neural network models that better resemble primate visual systems, or that factorization of latent variables should at least be a byproduct of other objectives that lead to more brain-like models. An important direction for future work is finding ways to directly incentivize factorization in model objective functions so as to test its causal impact on the fidelity of learned representations to neural data. Second, our choice of scene variables to analyze in this study was heuristic and somewhat arbitrary. Future work could consider unsupervised methods (in the vein of independent components analysis) for uncovering the latent sources of variance that generate visual data, and assessing to what extent these latent factors are encoded in factorized form. Third, in our work we do not specify the details of how a particular scene parameter is encoded within its factorized subspace, including whether the code is linear (“straightened”) or nonlinear^{42,44}. Neural codes could adopt different strategies resulting in similar factorization scores at the population level, each with some support in visual cortex literature: (1) Each neuron encodes a single latent variable^{45,46}, (2) Separate brain subregions encode qualitatively different latent variables but using distributed representations within each region^{47–49}, (3) Each neuron encodes multiple variables in a distributed population code, such that the factorization of different variables is only apparent as independent directions when assessed in high-dimensional population activity space^{45,50}. Future work can disambiguate among these possibilities by systematically examining ventral visual stream subregions^{9,49,51} and the single neuron tuning curves within them^{52,53}.

Methods

Monkey datasets

Macaque monkey datasets were of single-unit neural recordings²³, multi-unit neural recordings¹³, and object recognition behavior²⁴. Single-unit spiking responses to natural images were measured in V4 and anterior ventral IT²³. The advantages of this dataset are that it contains well-isolated single neurons, the gold standard for electrophysiology. Furthermore, the IT recordings were obtained from penetrating electrodes targeting the anterior ventral portion of IT near the base of skull, reflecting the highest level of the IT hierarchy. On the other hand, the multi-unit dataset was obtained from across IT with a bias toward where multi-unit arrays are more easily placed such as CIT and PIT¹³, complementing the recording locations of the single-unit dataset. An advantage of the multi-unit dataset using chronic recording arrays is that an order of magnitude more images were tested per recording site (see dataset comparisons in **Table S1**). Finally, the monkey behavioral dataset came from a third study examining the image-by-image object classification performance of macaques and humans²⁴.

Human datasets

Three datasets from humans were used, two fMRI datasets and one object recognition behavior dataset^{4,24,25}. The fMRI datasets used different images (color versus grayscale) but otherwise used fairly similar number of images and voxel resolution in MR imaging. Human fMRI studies have found that different DNN layers tend to map to V4 and HVC human fMRI voxels⁴. The human behavioral dataset measured image-by-image classification performance and was collected in the same study as the monkey behavioral signatures²⁴.

Computational models

In recent years, a variety of approaches to training DNN vision models have been developed that learn representations that can be used for downstream classification (and other) tasks. Models differ in a variety of implementational choices including in their architecture, objective function, and training dataset. In the models we sampled, objectives included supervised learning of object classification (AlexNet, ResNet), self-supervised contrastive learning (MoCo, SimCLR), and other unsupervised learning algorithms based on auxiliary tasks (e.g., reconstruction or colorization). A majority of the models that we considered relied on the widely used, performant ResNet-50 architecture, though some in our library utilized different architectures. The randomly initialized network control utilized ResNet-50 (see **Figure 4A,B**). The set of models we used is listed in **Table S2**.

Simulation of factorized versus non-factorized representational geometries

For the simulation in **Figure 1C**, we generated data in the following way. First we randomly sampled the values of $N=10$ binary features. Feature values corresponded to positions in an N -dimensional vector space as follows: each feature was assigned an axis in N -dimensional space, and the value of each feature (+1 or -1) was treated as a coefficient indicating the position along that axis. All but two of the feature axes were orthogonal to the rest. The last two features, which served as targets for the trained linear decoders, were assigned axes whose alignment ranged from 0 (orthogonal) to 1 (identical). In the noiseless case, factorization of these two variables with respect to one another is given by subtracting the square of the cosine of the angle between the axes from 1. We added Gaussian noise to the positions of each data point and randomly sampled K positive and negative examples for each variable of interest to use as training data for the linear classifier (a support vector machine).

Macaque neural data analyses

For the shuffle control used as a null model for factorization, we shuffled the object identity labels of the images (**Figure S1**). For the transformation used in **Figure 2B**, we computed the principal components of the mean neural activity response to each object class (“class centers,” $\mathbf{x}^c$), referred to as the inter-class PCs, $\mathbf{v}_1^{inter}, \mathbf{v}_2^{inter}, \dots, \mathbf{v}_N^{inter}$. We also computed the principal components of the data with corresponding class centers subtracted (i.e., $\mathbf{x} - \mathbf{x}^c$) from each activity pattern, referred to as the intra-class $\mathbf{v}_1^{intra}, \mathbf{v}_2^{intra}, \dots, \mathbf{v}_N^{intra}$. We transformed the data by applying to the class centers a change of basis matrix $W_{inter \rightarrow intra}$ that rotated each inter-class PC into the corresponding intra-class PC: $W_{inter \rightarrow intra} = \mathbf{v}_1^{intra} (\mathbf{v}_1^{inter})^T + \dots + \mathbf{v}_N^{intra} (\mathbf{v}_N^{inter})^T$. That is, the class centers were transformed by this matrix, but the relative positions of activity patterns within a given class were fixed. For an activation vector $\mathbf{x}$ belonging to a class c for which the average activity vector over all images of class c is $\mathbf{x}^c$, the transformed vector was:

$$\mathbf{x}_{transformed} = W_{inter \rightarrow intra} \mathbf{x}^c + (\mathbf{x} - \mathbf{x}^c) \quad (1)$$

This transformation has the effect of preserving intra-class variance statistics exactly from the original data and preserving everything about the statistics of inter-class variance except its orientation relative to intra-class variance. That is, the transformation is designed to affect (specifically decrease) factorization while controlling for all other statistics of the activity data that may be relevant to object classification performance (considering the simulation in [Figure 1C](#) of two binary variables, this basis change of the neural data in [Figure 2B](#) is equivalent to turning a square into the maximally flat parallelogram, the degenerate one where all the points are collinear).

Scene parameter variation

Our generated scenes consisted of foreground objects imposed upon natural backgrounds. To measure variance associated with a particular parameter like the background identity, we randomly sampled ten different backgrounds while holding the other variables (e.g., foreground object identity and pose constant). To measure variance associated with foreground object pose, we randomly varied object angle from $[-90, 90]$ along all three axes independently, object position on the two in-plane axes, horizontal $[-30\%, 30\%]$ and vertical $[-60\%, 60\%]$, and object size $[\times 1/1.6, \times 1.6]$. To measure variance associated with camera position, we took crops of the image with scale uniformly varying from 20% to 100% of the image size, and position uniformly distributed across the image. To measure variance associated with lighting conditions we applied random jitters to the brightness, contrast, saturation, and hue of an image, with jitter value bounds of $[-0.4, 0.4]$ for brightness, contrast, and saturation and $[-0.1, 0.1]$ for hue. These parameter choices follow standard data augmentation practices for self-supervised neural network training, as used, for example, in the SimCLR and MoCo models tested here^{17,18}.

Factorization and invariance metrics

Factorization and invariance were measured according to the following equations:

$$\text{factorization}_{\text{param}} = 1 - \text{var}_{\text{param}|\text{other_param_subspace}} / \text{var}_{\text{param}} \quad (2)$$

$$\text{invariance}_{\text{param}} = 1 - \text{var}_{\text{param}} / \text{var}_{\text{all param}} \quad (3)$$

Variance induced by a parameter ($\text{var}_{\text{param}}$) is computed by measuring the variance (summed across all dimensions of neural activity space) of neural responses to the 10 augmented versions of a base image where the augmentations are those obtained by varying the parameter of interest. This quantity is then averaged across the 100 base images. The variance induced by all parameters is simply the sum of the variances across all images and augmentations. To define the “other-parameter subspace,” we averaged neural responses for a given base image over all augmentations using the parameter of interest, and ran PCA on the resulting set of averaged responses. The subspace was defined as the space spanned by top PCA components containing 90% of the variance of these responses. Intuitively, this space captures the bulk of the variance driven by all parameters other than the parameter of interest (due to the averaging step). The variance of the parameter of interest *within* this “other-parameter subspace,” $\text{var}_{\text{param}|\text{other_param_subspace}}$, was computed the same way as $\text{var}_{\text{param}}$ but using the projections of neural activity responses onto the other-parameter subspace. In the main text, we refer to this method of computing factorization as PCA based factorization.

We also considered an alternative definition of factorization referred to as covariance based factorization. In this alternative definition, we measured the covariance matrices $\text{cov}_{\text{param}}$ and $\text{cov}_{\text{other_param}}$ induced by varying (in the same fashion as above) the parameter of interest, and all other parameters. Factorization was measured by the following equation:

$$\text{factorization}_{\text{param}} = \frac{1 - \text{Trace}[(\text{cov}_{\text{param}})^T \text{cov}_{\text{other_param}}]}{(\text{Trace}[(\text{cov}_{\text{param}})^T \text{cov}_{\text{param}}] \text{Trace}[(\text{cov}_{\text{other_param}})^T \text{cov}_{\text{other_param}}])^{1/2}} \quad (4)$$

This is equal to 1 minus the dot product between the normalized, flattened covariance matrices, and thus covariance based factorization is a measure of the discrepancy of the covariance structure induced by the parameter of interest and other parameters. The main findings were unaffected by our choice of method for computing the factorization metric, whether PCA or covariance based (**Figures 5** [↗](#) **7** [↗](#)). An advantage of the PCA based method is that as an intermediate one recovers the linear subspaces containing parameter variance, but in so doing requires an arbitrary choice of the explained variance threshold used to choose the number of PCs. By contrast, the covariance based method is more straightforward to compute and has no free parameters. Thus, these two metrics are complementary, and somewhat analogous in methodology to two metrics commonly used for measuring dimensionality (the number of components needed to explain a certain fraction of the variance, analogous to our original PCA based definition, and the participation ratio, analogous to our covariance based definition)⁵⁴ [↗](#), ⁵⁵ [↗](#).

Natural movie factorization metrics

For natural movies, variance is not induced by explicit control of a parameter as in our synthetic scenes but implicitly, by considering contiguous frames (separated by 200ms in real time) as reflective of changes in one of two motion parameters (object versus observer motion) depending on how stationary the observer is (MIT Moments in Time movie set: stationary observer; UT-Austin Egocentric movie set: nonstationary)²⁷ [↗](#), ²⁸ [↗](#). Here, the *all parameters* condition is simply the variance across all movie frames which in the case of MIT Moments in Time dataset includes variance across thousands of video clips taken in many different settings and in the case of the UT-Austin Egocentric movie dataset includes variance across only 4 movies but over long durations of time during which an observer translates extensively in an environment (3-5 hours).

Thus, movie clips in the MIT Moments in Time movie set contained new scenes with different object identities, backgrounds, and lightings and thus effectively captured variance induced by these non-spatial parameters²⁸ [↗](#). In the UT Austin Egocentric movie set, new objects and backgrounds are encountered as the subject navigates around the urban landscape²⁷ [↗](#).

Model neural encoding fits

Linear mappings between model features and neuron (or voxel) responses were computed using ridge regression (with regularization coefficient selected by cross validation) on a low-dimensional linear projection of model features (top 300 PCA components computed using images in each dataset). We also tested an alternative approach to measuring representational similarity between models and experimental data based on representational similarity analysis (RSA)⁵⁶ [↗](#), computing dot product similarities of the representations of all pairs of images and measuring the Spearman correlation coefficient between these pairwise similarity matrices obtained from a given model and neural dataset, respectively.

Model behavioral signatures

We followed the approach of Rajalingham, Issa et al.²⁴ [↗](#). We took human and macaque behavioral data from the object classification task and used it to create signatures of image-level difficulty (the “I1” vector) and image-by-distractor-object confusion rates (the “I2” matrix). We did the same for the DNN models, extracting model “behavior” by training logistic regression classifiers to classify object identity in the same image dataset used in the experiments of Rajalingham, Issa et al.²⁴ [↗](#), using model layer activations as inputs. Model behavioral accuracy rates on image by distractor object pairs were assessed using the classification probabilities output by the logistic regression model, and these were used to compute I1 and I2 metrics as was done for the true behavioral data. Behavioral similarity between models and data was assessed by measuring the correlation between the entries of the I1 vectors and I2 matrices (both I1 and I2 results are reported).

Model layer choices

The scatter plots in **Figure 4A,B** and **Figure S2** use metrics (factorization, invariance, and goodness of neural fit) taken from the final representational layer of the network (the layer prior to the logits layer used for classification in supervised network, prior to the embedding head in contrastive learning models, or prior to any auxiliary task-specific layers in unsupervised models trained using auxiliary tasks). However, representational geometries of model activations, and their match to neural activity and behavior, vary across layers. This variability arises because different model layers correspond to different stages of processing in the model (convolutional layers in some cases, and pooling operations in others), and may even have different dimensionalities. To ensure that our results do not depend on idiosyncrasies of representations in one particular model layer and the particular network operations that precede it, summary correlation statistics in all other figures (**Figures 5-7** & **S3-S5**) show the results of the analysis in question averaged over the five final representational layers of the model. That is, the metrics of interest (factorization, invariance, neural encoding fits, RDM correlation, behavioral similarity scores) were computed independently for each of the five final representational layers of each model, and these five values were averaged prior to computing correlations between different metrics.

Correlation of model predictions and experimental data

A Spearman's linear correlation coefficient was calculated for each model layer by biological dataset combination (6 monkey datasets and 6 human datasets). Here, we do not correct for noise in the biological data when computing the correlation coefficient, as this would require trial repeats (for computing intertrial variability) that were limited or not available in the fMRI data used. In any event, normalizing by the data noise ceiling applies a uniform scaling to all model prediction scores and does not affect model comparison, which only depends on ranking models as being relatively better or worse in predicting brain data.

Finally, we estimated the effectiveness of model factorization, invariance, or dimensionality in combination with model object classification performance for predicting model neural and behavioral fit by performing a linear regression on the particular dual metric combination (e.g., classification plus object pose factorization) and reporting the Spearman correlation coefficient of the linearly weighted metric combination. The correlation was assessed on held-out models (80% used for training, 20% for testing), and the results were averaged over 100 randomly sampled train/test splits.

Acknowledgements

This work was performed on the Columbia Zuckerman Institute Axon GPU cluster and via generous access to Cloud TPUs from Google's TPU Research Cloud (TRC). JWL was supported by the DOE CSGF (DE-SC0020347). EBI was supported by a Klingenstein-Simons fellowship, Sloan Foundation fellowship, and Grossman-Kavli Scholar Award. We thank Erica Shook for comments on a previous version of the manuscript. The authors declare no competing interests.

Author contributions

JWL and EBI designed the research. JWL performed the computational modeling and data analysis. JWL and EBI wrote the manuscript. EBI supervised the research.

Supplement

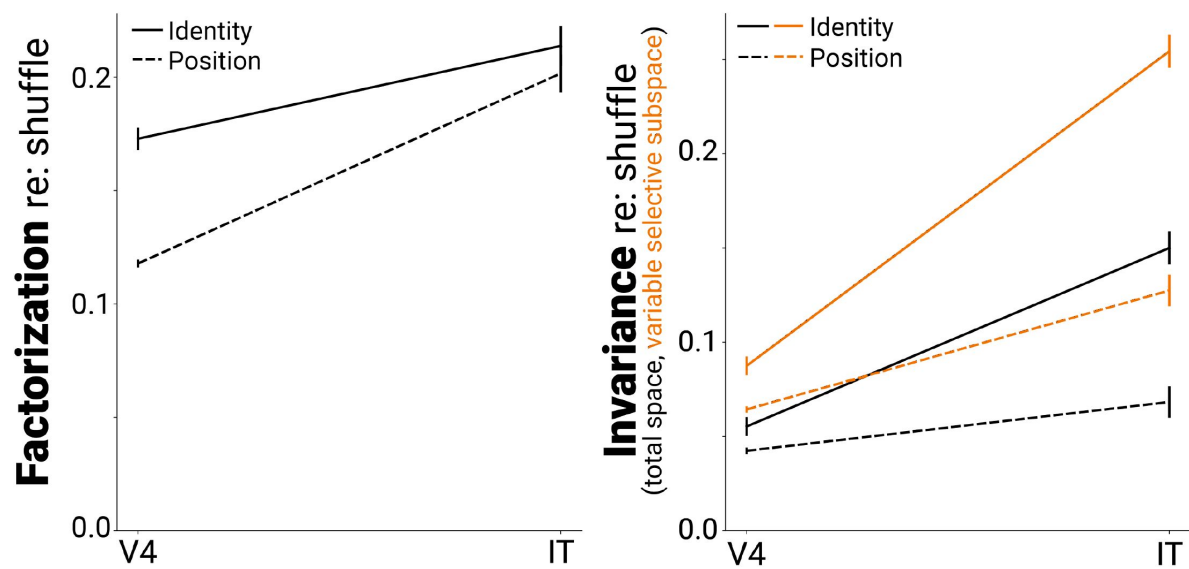


Figure S1.

Factorization and invariance in V4 & IT neural data.

Normalized factorization and invariance as in [Figure 2A](#) but after subtracting shuffle control for V4 and IT neural datasets. Shuffling the image identities of each population vector accounts for increases in factorization driven purely by changes in the covariance statistics of population responses between V4 and IT. However, the normalized factorization scores remained significantly above zero for both brain areas.

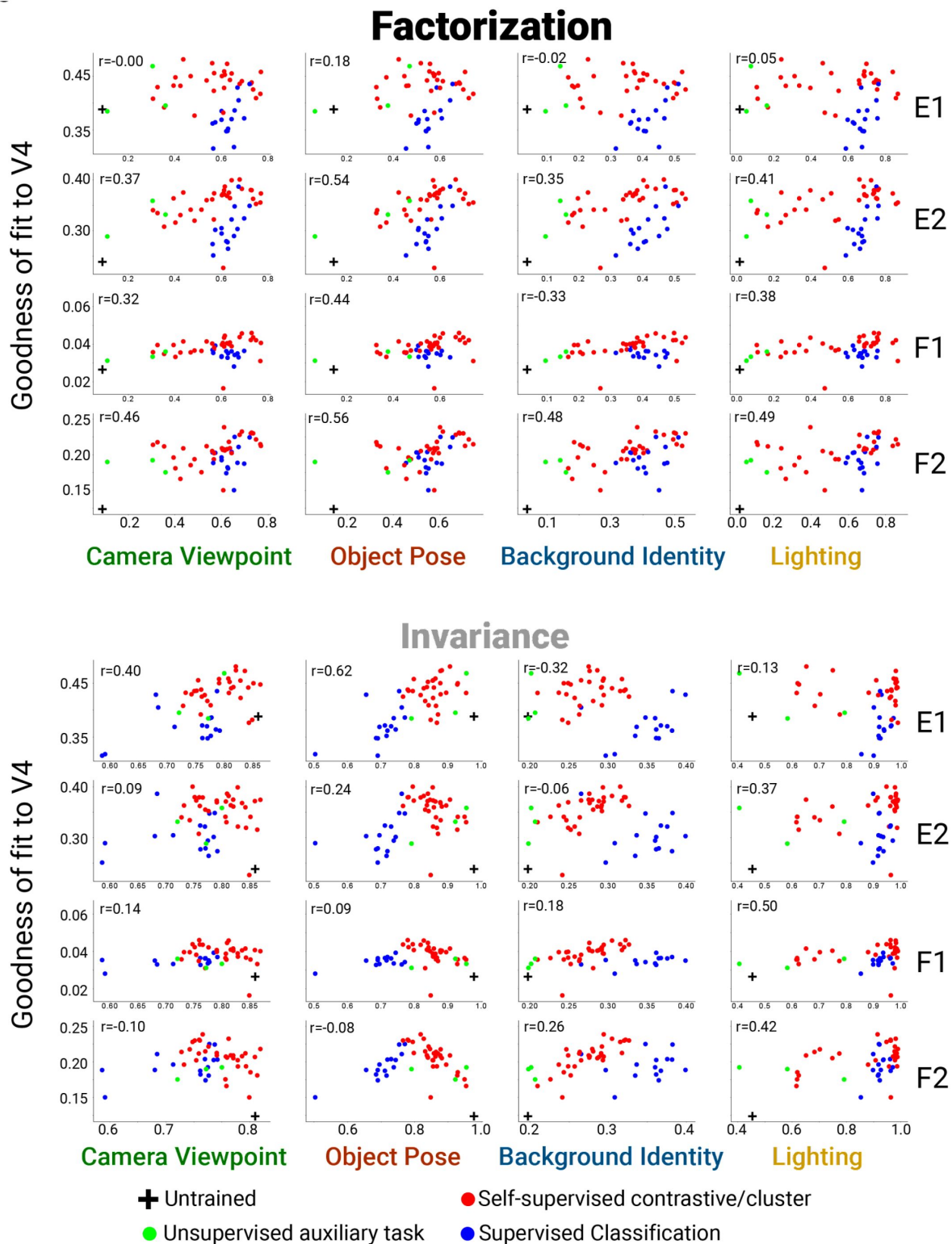


Figure S2.

Scatter plots for all datasets.

Scatter plots as in [Figure 4A,B](#) for all datasets. Brain metrics (y-axes) by panel are: (A) macaque neuron/human voxel fits in V4 cortex, (B) macaque neuron/human voxel fits in ITC/HVC, and (C) macaque/human per-image classification performance (I1) and image-by-distractor class performance (I2). In all panels, the plots in the top half use DNN factorization scores (PCA based method) on the x-axis while the bottom half use DNN invariance scores.

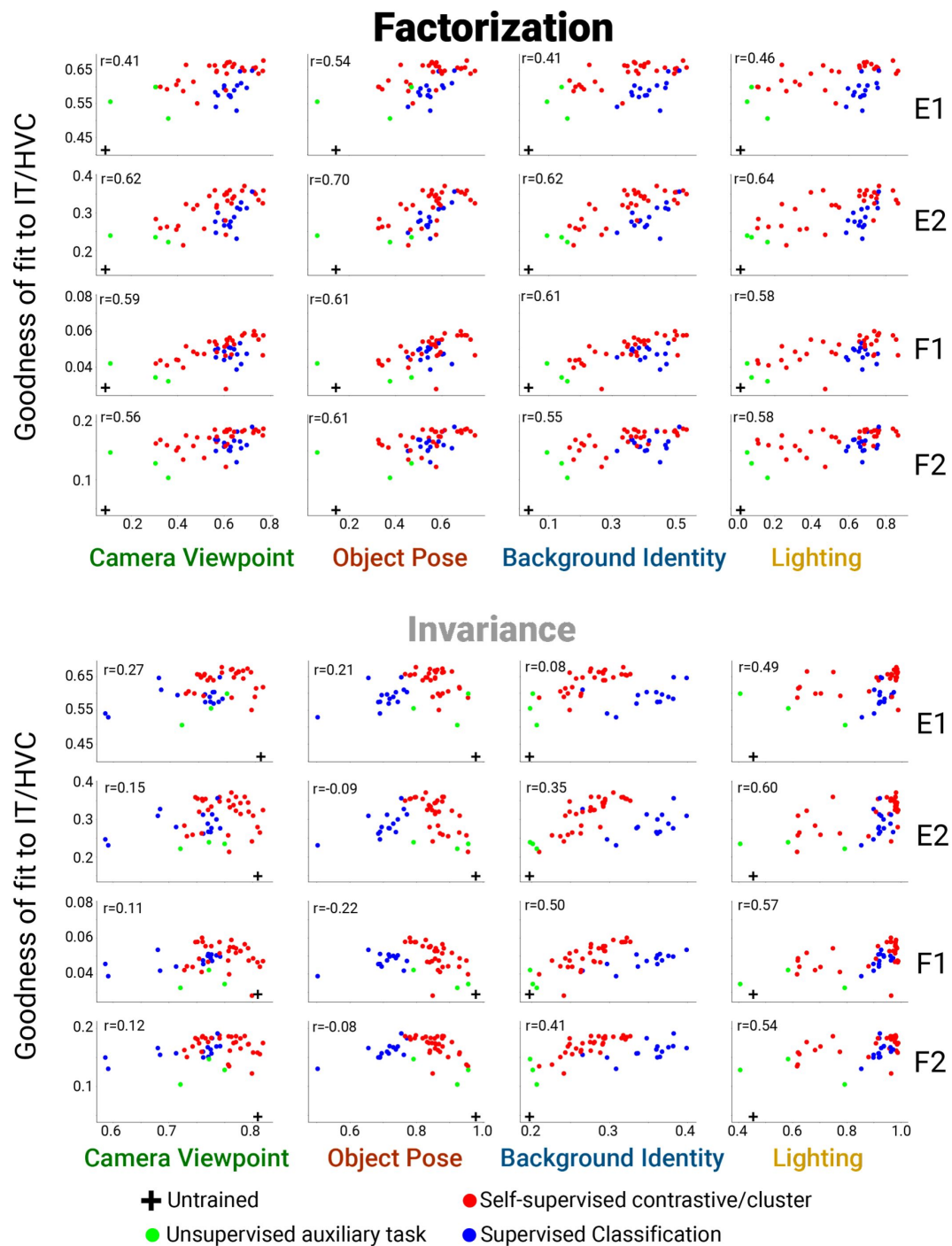


Figure S2. (continued)

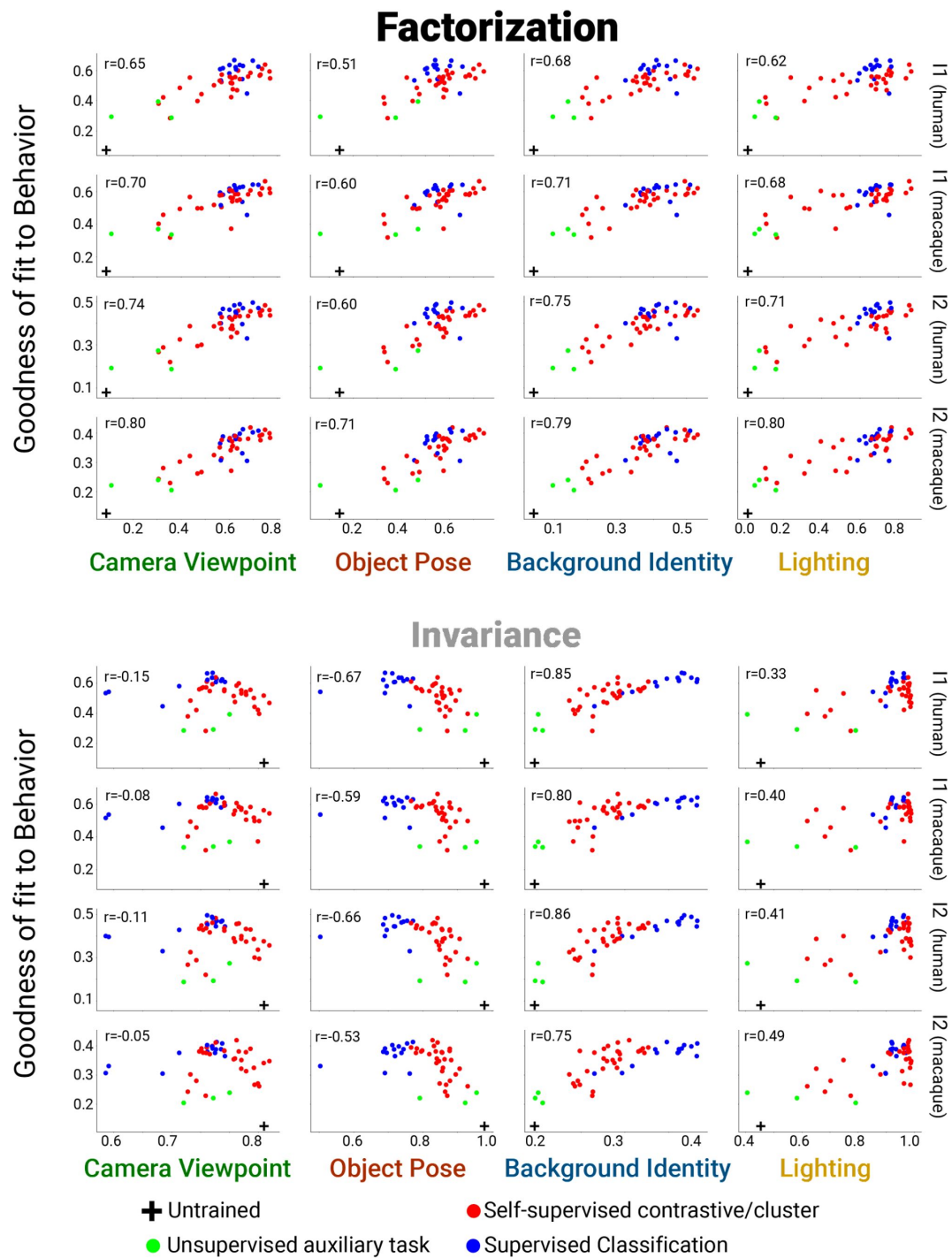


Figure S2. (continued)

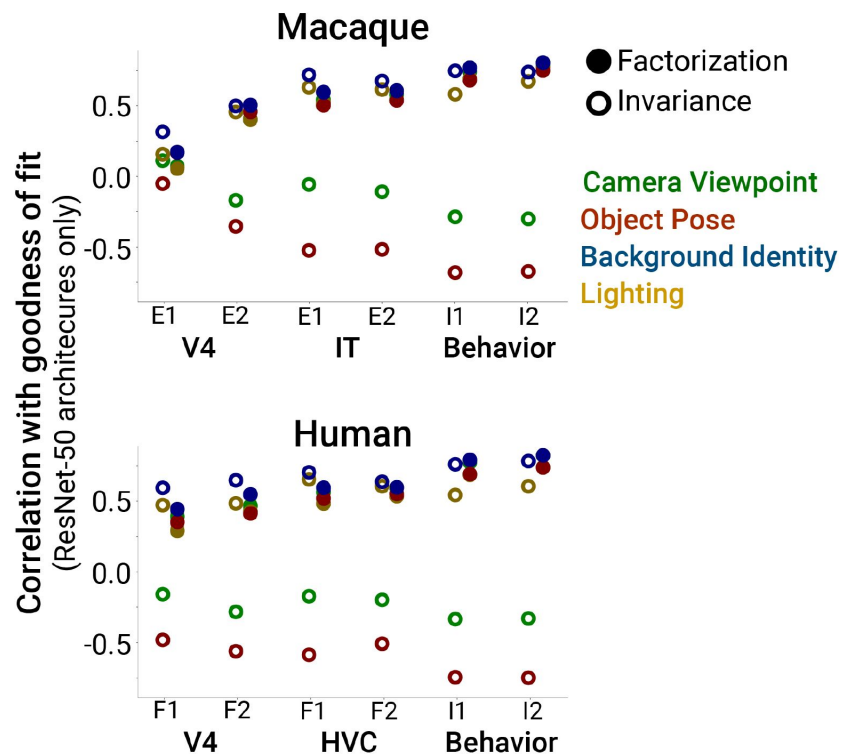


Figure S3.

Predictivity of factorization and invariance restricting to ResNet-50 model architectures.

Same format as [Figure 5C](#) except with the analyses restricted to using only models with the Resnet-50 architecture. The main finding of factorization of scene parameters in DNNs being generally positively correlated with better predictions of brain data is replicated using this architecture-matched subset of models, controlling for potential confounds from model architecture.

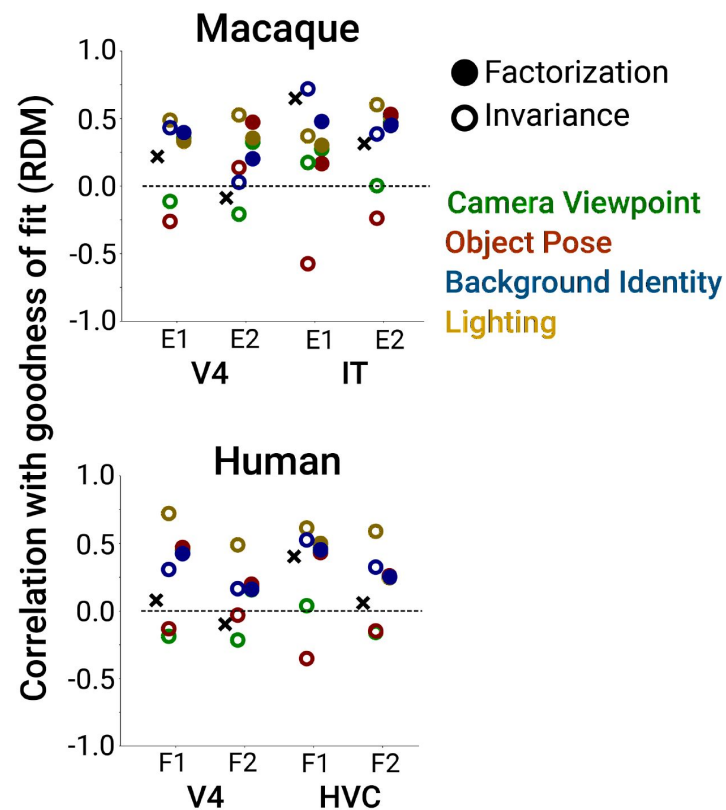


Figure S4.

Predictivity of factorization and invariance for RDMs.

Same format as [Figure 5C](#) except for predicting population representational dissimilarity matrices (RDMs) of macaque neurophysiological and human fMRI data (in [Figure 5C](#) linear encoding fits of each single neuron/voxel were used to measure brain predictivity of a model). The main finding of factorization of scene parameters in DNNs being positively correlated with better predictions of brain data is replicated using RDMs instead of neural/voxel goodness of encoding fit.

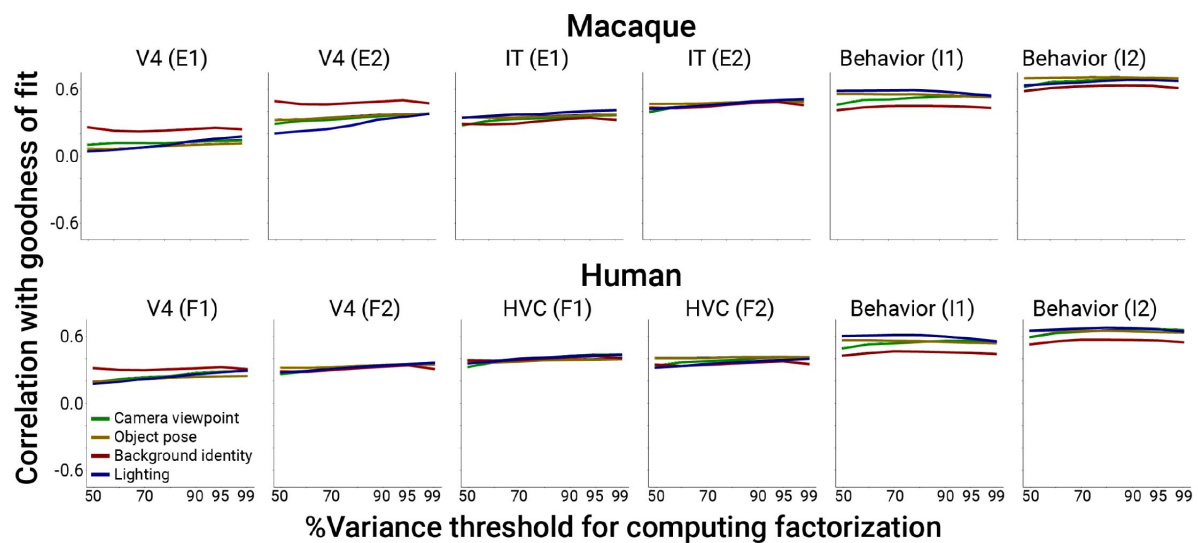


Figure S5.

Effect on neural and behavioral predictivity of PCA threshold for computing PCA based factorization, Related to Figure 5.

The % variance threshold used in the main text for estimating a PCA linear subspace capturing the bulk of the variance induced by all other parameters besides the parameter of interest is somewhat arbitrary. Here we show that results of our main analysis change little if we vary this parameter from 50-99%. In the main text, a PCA threshold of 90% was used for computing PCA based factorization scores.

Dataset	Key	#neurons, voxels	#subj	Image Stimuli	#images
DiCarlo-Majaj-Hong 2015¹ Macaque V4 multi-unit activity	E1	128	2	6°, grayscale, synthetic	5760
DiCarlo-Majaj-Hong 2015¹ Macaque IT multi-unit activity	E1	168	2	6°, grayscale, synthetic	5760
DiCarlo-Rust 2012² Macaque V4 single neuron	E2	143	2	5°, grayscale, natural	300
DiCarlo-Rust 2012² Macaque IT single neuron	E2	142	2	5°, grayscale, natural	300
DiCarlo-Rajalingham-Issa 2018³ Macaque behavior Image-level classification	I1	N/A	5	6-8°, grayscale, synthetic	240
DiCarlo-Rajalingham-Issa 2018³ Macaque behavior Image x class confusion matrix	I2	N/A	5	6-8°, grayscale, synthetic	240
Gallant-Kay 2008⁴ Human V4 fMRI (dataset)	F1	2,557	2	20°, grayscale, natural	1870
Gallant-Kay 2008⁴ Human HVC fMRI (dataset)	F1	1,286	2	20°, grayscale, natural	1870
Horikawa-Kamitani 2019⁵ Human V4 fMRI (dataset)	F2	3,377	3	12°, color, natural	1250
Horikawa-Kamitani 2019⁵ Human HVC fMRI (dataset)	F2	14,465	3	12°, color, natural	1250
DiCarlo-Rajalingham-Issa 2018³ Human behavior Image-level classification	I1	N/A	1472	6-8°, grayscale, synthetic	240
DiCarlo-Rajalingham-Issa 2018³ Human behavior Image x class confusion matrix	I2	N/A	1472	6-8°, grayscale, synthetic	240

Table S1.

Datasets used for measuring similarity of models to the brain.

Datasets from both macaque and human high-level visual cortex as well as high-level visual behavior were collated for testing the brainlikeness of computational models. For neural and fMRI datasets, the features in the model were used to predict the image-by-image response pattern of each neuron or voxel. For behavior datasets, the performance of linear decoders built atop model representations were compared to performance per image of macaques and humans.

Model	Architecture	Loss Function	Customization
SimCLR ⁶	ResNet-50	Self-supervised (contrastive)	-----
SimCLR ⁶	ResNet-50	Self-supervised (contrastive)	2x wide
SimCLR ⁶	ResNet-152	Self-supervised (contrastive)	2x wide
SimCLR ⁶	ResNet-50	Self-supervised (contrastive)	no projection head
SimCLR ⁶	ResNet-50	Self-supervised (contrastive)	only crop augmentations
SimCLR ⁶	ResNet-50	Self-supervised (contrastive)	only crop augmentations, temperature 0.2
SimCLR ⁶	ResNet-50	Self-supervised (contrastive)	only crop augmentations, temperature 0.05
SimCLR ⁶	ResNet-50	Self-supervised (contrastive)	only crop and blur augmentations
SimCLR ⁶	ResNet-50	Self-supervised (contrastive)	only crop and non-hue color jitter augmentations
SimCLR ⁶	ResNet-50	Self-supervised (contrastive)	only crop, blur, and non-hue color jitter augmentations
MOCO ⁷	ResNet-50	Self-supervised (contrastive)	-----
MOCO v2 ⁸	ResNet-50	Self-supervised (contrastive)	-----
MOCO v2 ⁸	ResNet-50	Self-supervised (contrastive)	only crop augmentations
MOCO v2 ⁸	ResNet-50	Self-supervised (contrastive)	only crop, color jitter, and grayscale augmentations
MOCO v2 ⁸	ResNet-50	Self-supervised (contrastive)	only crop augmentations, all image inputs preprocessed to grayscale
MOCO v2 ⁸	ResNext-50	Self-supervised (contrastive)	-----
MOCO v2 ⁸	ResNet-18	Self-supervised (contrastive)	-----
Instance discrimination ⁹	ResNet-50	Self-supervised (image discrimination)	-----
InfoMin ¹⁰	ResNet-50	Self-supervised (contrastive)	-----
InfoMin ¹⁰	ResNext-101	Self-supervised (contrastive)	-----

Table S2.

Models tested.

For each model, we measured representational factorization and invariance in each of the final five layers of the model as well as evaluating their brainlikeness using the datasets in **Table S1**.

InfoMin ¹⁰	ResNext-152	Self-supervised (contrastive)	-----
SwAV ¹¹	ResNet-50	Self-supervised (cluster)	-----
Deep clustering v2 ¹²	ResNet-50	Self-supervised (cluster)	-----
BYOL ¹³	ResNet-50	Self-supervised (no negative examples)	-----
BYOL ¹³	ResNet-50	Self-supervised (no negative examples)	only crop augmentations during training
BYOL ¹³	ResNet-50	Self-supervised (no negative examples)	only crop and blur augmentations
BYOL ¹³	ResNet-50	Self-supervised (no negative examples)	without color jitter augmentation
BYOL ¹³	ResNet-50	Self-supervised (no negative examples)	without grayscale augmentation
BYOL ¹³	ResNet-50	Self-supervised (no negative examples)	batch size 64
BYOL ¹³	ResNet-50	Self-supervised (no negative examples)	batch size 512
Relative patch location ¹⁴	ResNet-50	Auxiliary task (determine relative positions of image patches)	-----
Rotation prediction ¹⁴	ResNet-50	Auxiliary task (infer rotations that were applied given a set of images)	-----
Colorization ¹⁵	ResNet-50	Auxiliary task: (colorize grayscale images)	-----
Jigsaw puzzle ¹⁶	ResNet-50	Auxiliary task: (determine relative positions of image patches)	-----
Big BiGAN ¹⁷	ResNet-50	Auxiliary task (autoencoder objective with	-----

Table S2. (continued)

		reconstruction error parameterized using a neural network discriminator)	
ResNet ¹⁸	ResNet-50	Supervised (classification)	-----
ResNet ¹⁸	ResNet-50	Supervised (classification)	MOCO data augmentations used during training
ResNet ¹⁸	ResNet-50	Supervised (classification)	no data augmentation used during training
ResNet ¹⁸	ResNet-18	Supervised (classification)	-----
ResNet ¹⁸	ResNet-101	Supervised (classification)	-----
Wide ResNet ¹⁹	ResNet-50	Supervised (classification)	-----
AlexNet ²⁰	AlexNet	Supervised (classification)	-----
GoogLeNet ²¹	GoogLeNet	Supervised (classification)	-----
Inception-v3 ²²	Inception-v3	Supervised (classification)	-----
DenseNet ²³	DenseNet-169	Supervised (classification)	-----
DenseNet ²³	DenseNet-121	Supervised (classification)	-----
VGG ²⁴	VGG-11	Supervised (classification)	-----
VGG ²⁴	VGG-13	Supervised (classification)	-----
VGG ²⁴	VGG-16	Supervised (classification)	-----
VGG ²⁴	VGG-19	Supervised (classification)	-----
MobileNet ²⁵	MobileNet	Supervised (classification)	-----
SqueezeNet ²⁶	SqueezeNet-10	Supervised (classification)	-----
SqueezeNet ²⁶	SqueezeNet-11	Supervised (classification)	-----
ResNet ²⁷	ResNext-50	Supervised (classification)	-----

Table S2. (continued)

ResNet ²⁷	ResNet-101	Supervised (classification)	-----
MnasNet ²⁸	MnasNet_05	Supervised (classification)	-----
MnasNet ²⁸	MnasNet_10	Supervised (classification)	-----
ShuffleNet ²⁹	ShuffleNet_05	Supervised (classification)	-----
ShuffleNet ²⁹	ShuffleNet_10	Supervised (classification)	-----

Table S2. (continued)

References

1. Cadieu C. F., et al. (2014) **Deep Neural Networks Rival the Representation of Primate IT Cortex for Core Visual Object Recognition** *PLOS Comput. Biol* **10**
2. Schrimpf M., et al. (2020) **Brain-Score: Which Artificial Neural Network for Object Recognition is most Brain-Like?** *bioRxiv* **407007** <https://doi.org/10.1101/407007>
3. Yamins D. L. K., et al. (2014) **Performance-optimized hierarchical models predict neural responses in higher visual cortex** *Proc. Natl. Acad. Sci* **201403112** <https://doi.org/10.1073/pnas.1403112111>
4. Nonaka S., Majima K., Aoki S. C., Kamitani Y (2021) **Brain hierarchy score: Which deep neural networks are hierarchically brain-like?** *iScience* **24**
5. Linsley D., et al. (2023) **Linsley, D., et al. Performance-optimized deep neural networks are evolving into worse models of inferotemporal visual cortex. Preprint at 10.48550/arXiv.2306.03779 (2023).** <https://doi.org/10.48550/arXiv.2306.03779>
6. DiCarlo J. J., Cox D. D (2007) **Untangling invariant object recognition** *Trends Cogn. Sci* **11**:333–341
7. Freiwald W. A., Tsao D. Y (2010) **Functional Compartmentalization and Viewpoint Generalization Within the Macaque Face-Processing System** *Science* **330**:845–851
8. Hong H., Yamins D. L. K., Majaj N. J., DiCarlo J. J (2016) **Explicit information for category-orthogonal object properties increases along the ventral stream** *Nat. Neurosci* **19**:613–622
9. Kravitz D. J., Saleem K. S., Baker C. I., Ungerleider L. G., Mishkin M (2013) **The ventral visual pathway: an expanded neural framework for the processing of object quality** *Trends Cogn. Sci* **17**:26–49
10. Peters B., Kriegeskorte N (2021) **Capturing the objects of vision with neural networks.** *Nat Hum. Behav* **5**:1127–1144
11. Chung S., Lee D. D., Sompolinsky H (2017) **Classification and Geometry of General Perceptual Manifolds** *ArXiv171006487 Cond-Mat Q-Bio Stat*
12. Johnston W. J., Fusi S (2023) **Abstract representations emerge naturally in neural networks trained to perform multiple tasks** *Nat. Commun* **14**
13. Majaj N. J., Hong H., Solomon E. A., DiCarlo J. J (2015) **Simple Learned Weighted Sums of Inferior Temporal Neuronal Firing Rates Accurately Predict Human Core Object Recognition Performance** *J. Neurosci* **35**:13402–13418
14. Rust N. C., DiCarlo J. J (2010) **Selectivity and Tolerance (“Invariance”) Both Increase as Visual Information Propagates from Cortical Area V4 to IT** *J. Neurosci* **30**:12978–12995
15. Krizhevsky A., Sutskever I., Hinton G. E. (2012) **ImageNet Classification with Deep Convolutional Neural Networks** *Advances in Neural Information Processing Systems* **25**:1097–1105

16. He K., Zhang X., Ren S., Sun J. (2015) **Deep Residual Learning for Image Recognition** *ArXiv151203385 Cs*
17. He K., Fan H., Wu Y., Xie S., Girshick R (2020) **Momentum Contrast for Unsupervised Visual Representation Learning** *ArXiv191105722 Cs*
18. Chen T., Kornblith S., Norouzi M., Hinton G (2020) **A Simple Framework for Contrastive Learning of Visual Representations** *ArXiv200205709 Cs Stat*
19. Tian Y., Krishnan D., Isola P. (2019) **Contrastive Multiview Coding** *ArXiv190605849 Cs*
20. Doersch C., Gupta A., Efros A. A. (2016) **Unsupervised Visual Representation Learning by Context Prediction** *ArXiv150505192 Cs*
21. He K., Gkioxari G., Dollar P., Girshick R. (2017) **Mask R-CNN** *Proceedings of the IEEE International Conference on Computer Vision* :2961–2969
22. Donahue J., Simonyan K (2019) **Large Scale Adversarial Representation Learning** *Advances in Neural Information Processing Systems* **32**
23. Rust N. C., DiCarlo J. J (2012) **Balanced Increases in Selectivity and Tolerance Produce Constant Sparseness along the Ventral Visual Stream** *J. Neurosci* **32**:10170–10182
24. Rajalingham R., et al. (2018) **Large-Scale, High-Resolution Comparison of the Core Visual Object Recognition Behavior of Humans, Monkeys, and State-of-the-Art Deep Artificial Neural Networks** *J. Neurosci* **38**:7255–7269
25. Kay K. N., Naselaris T., Prenger R. J., Gallant J. L (2008) **Identifying natural images from human brain activity** *Nature* **452**:352–355
26. Shen G., Horikawa T., Majima K., Kamitani Y (2019) **Deep image reconstruction from human brain activity** *PLOS Comput. Biol* **15**
27. Lee Y. J., Ghosh J., Grauman K (2012) **Discovering important people and objects for egocentric video summarization** *IEEE Conference on Computer Vision and Pattern Recognition* <https://doi.org/10.1109/CVPR.2012.6247820>
28. Monfort M., et al. (2019) **Monfort, M., et al. Moments in Time Dataset: one million videos for event understanding. Preprint at 10.48550/arXiv.1801.03150 (2019).** <https://doi.org/10.48550/arXiv.1801.03150>
29. Nayebe A., et al. (2021) **Goal-Driven Recurrent Neural Network Models of the Ventral Visual Stream** *bioRxiv* <https://doi.org/10.1101/2021.02.17.431717>
30. Caron M., Bojanowski P., Joulin A., Douze M. (2019) **Deep Clustering for Unsupervised Learning of Visual Features** *ArXiv180705520 Cs*
31. Caron M., et al. (2020) **Unsupervised Learning of Visual Features by Contrasting Cluster Assignments** *ArXiv200609882 Cs*
32. Elmoznino E., Bonner M. F. (2022) **High-performing neural network models of visual cortex benefit from high latent dimensionality** *bioRxiv* <https://doi.org/10.1101/2022.07.13.499969>

33. Conwell C., Prince J. S., Kay K. N., Alvarez G. A., Konkle T. (2023) **What can 1.8 billion regressions tell us about the pressures shaping high-level visual representation in brains and machines?** *bioRxiv* <https://doi.org/10.1101/2022.03.28.485868>
34. Kim H., Mnih A. (2018) **Disentangling by Factorising** *Proceedings of the 35th International Conference on Machine Learning* :2649–2658
35. Eastwood C., Williams C. K. I. (2018) **A Framework for the Quantitative Evaluation of Disentangled Representations** . in *International conference on learning representations*
36. Higgins I., et al. (2018) **Higgins, I., et al. Towards a Definition of Disentangled Representations. Preprint at 10.48550/arXiv.1812.02230 (2018).** <https://doi.org/10.48550/arXiv.1812.02230>
37. Higgins I., et al. (2020) **Unsupervised deep learning identifies semantic disentanglement in single inferotemporal neurons** *ArXiv200614304 Q-Bio*
38. Bernardi S., et al. (2020) **The Geometry of Abstraction in the Hippocampus and Prefrontal Cortex** *Cell* **183**:954–967
39. Boyle L. M., Posani L., Irfan S., Siegelbaum S. A., Fusi S (2024) **Tuned geometries of hippocampal representations meet the computational demands of social memory** *Neuron* <https://doi.org/10.1016/j.neuron.2024.01.021>
40. Alleman M., Lindsey J. W., Fusi S. (2024) **Alleman, M., Lindsey, J. W. & Fusi, S. Task structure and nonlinearity jointly determine learned representational geometry. Preprint at 10.48550/arXiv.2401.13558 (2024).** <https://doi.org/10.48550/arXiv.2401.13558>
41. Sorscher B., Ganguli S., Sompolsky H (2022) **Neural representational geometry underlies few-shot concept learning** *Proc. Natl. Acad. Sci* **119**
42. Hénaff O. J., et al. (2021) **Primary visual cortex straightens natural video trajectories** *Nat. Commun* **12**
43. Toosi T., Issa E (2022) **Brain-like representational straightening of natural movies in robust feedforward neural networks** *The Eleventh International Conference on Learning Representations* **11**
44. Hénaff O. J., Goris R. L. T., Simoncelli E. P (2019) **Perceptual straightening of natural videos** *Nat. Neurosci* **22**
45. Field D. J (1994) **What Is the Goal of Sensory Coding?** *Neural Comput* **6**:559–601
46. Chang L., Tsao D. Y (2017) **The Code for Facial Identity in the Primate Brain** *Cell* **169**:1013–1028
47. Tsao D. Y., Freiwald W. A., Tootell R. B. H., Livingstone M. S (2006) **A Cortical Region Consisting Entirely of Face-Selective Cells** *Science* **311**:670–674
48. Lafer-Sousa R., Conway B. R (2013) **Parallel, multi-stage processing of colors, faces and shapes in macaque inferior temporal cortex** *Nat. Neurosci* **16**:1870–1878
49. Vaziri S., Carlson E. T., Wang Z., Connor C. E (2014) **A Channel for 3D Environmental Shape in Anterior Inferotemporal Cortex** *Neuron* **84**:55–62

50. Rigotti M., et al. (2013) **The importance of mixed selectivity in complex cognitive tasks** *Nature* **497**:585–590
51. Kravitz D. J., Saleem K. S., Baker C. I., Mishkin M (2011) **A new neural framework for visuospatial processing** *Nat. Rev. Neurosci* **12**:217–230
52. Leopold D. A., Bondar I. V., Giese M. A (2006) **Norm-based face encoding by single neurons in the monkey inferotemporal cortex** *Nature* **442**:572–575
53. Freiwald W. A., Tsao D. Y., Livingstone M. S (2009) **A face feature space in the macaque temporal lobe** *Nat. Neurosci* **12**:1187–1196
54. Abbott L. F., Rajan K., Sompolinsky H. (2010) **Abbott, L. F., Rajan, K. & Sompolinsky, H. Interactions between Intrinsic and Stimulus-Evoked Activity in Recurrent Neural Networks. Preprint at 10.48550/arXiv.0912.3832 (2010).** <https://doi.org/10.48550/arXiv.0912.3832>
55. Litwin-Kumar A., Harris K. D., Axel R., Sompolinsky H., Abbott L. F (2017) **Optimal Degrees of Synaptic Connectivity** *Neuron* **93**:1153–1164
56. Kriegeskorte N., Kievit R. A (2013) **Representational geometry: integrating cognition, computation, and the brain** *Trends Cogn. Sci* **17**:401–412
1. Majaj N. J., Hong H., Solomon E. A., DiCarlo J. J (2015) **Simple Learned Weighted Sums of Inferior Temporal Neuronal Firing Rates Accurately Predict Human Core Object Recognition Performance** *J. Neurosci* **35**:13402–13418
2. Rust N. C., DiCarlo J. J (2012) **Balanced Increases in Selectivity and Tolerance Produce Constant Sparseness along the Ventral Visual Stream** *J. Neurosci* **32**:10170–10182
3. Rajalingham R., et al. (2018) **Large-Scale, High-Resolution Comparison of the Core Visual Object Recognition Behavior of Humans, Monkeys, and State-of-the-Art Deep Artificial Neural Networks** *J. Neurosci* **38**:7255–7269
4. Kay K. N., Naselaris T., Prenger R. J., Gallant J. L (2008) **Identifying natural images from human brain activity** *Nature* **452**:352–355
5. Shen G., Horikawa T., Majima K., Kamitani Y (2019) **Deep image reconstruction from human brain activity** *PLOS Comput. Biol* **15**
6. Chen T., Kornblith S., Norouzi M., Hinton G (2020) **A Simple Framework for Contrastive Learning of Visual Representations** *ArXiv200205709 Cs Stat*
7. He K., Fan H., Wu Y., Xie S., Girshick R (2020) **Momentum Contrast for Unsupervised Visual Representation Learning** *ArXiv191105722 Cs*
8. Chen X., Fan H., Girshick R., He K. (2020) **Improved Baselines with Momentum Contrastive Learning** *ArXiv200304297 Cs*
9. Wu Z., Xiong Y., Yu S., Lin D. (2018) **Unsupervised Feature Learning via Non-Parametric Instance-level Discrimination** *ArXiv180501978 Cs*
10. Tian Y., et al. (2020) **What makes for good views for contrastive learning** *ArXiv200510243 Cs*

11. Caron M., et al. (2020) **Unsupervised Learning of Visual Features by Contrasting Cluster Assignments** *ArXiv200609882 Cs*
12. Caron M., Bojanowski P., Joulin A., Douze M. (2019) **Deep Clustering for Unsupervised Learning of Visual Features** *ArXiv180705520 Cs*
13. Grill J.-B., et al. (2020) **Bootstrap Your Own Latent: A New Approach to Self-Supervised Learning** *ArXiv200607733 Cs Stat*
14. Gidaris S., Bursuc A., Komodakis N., Pérez P., Cord M. (2019) **Gidaris, S., Bursuc, A., Komodakis, N., Pérez, P. & Cord, M. Boosting Few-Shot Visual Learning with Self-Supervision. Preprint at 10.48550/arXiv.1906.05186 (2019).** <https://doi.org/10.48550/arXiv.1906.05186>
15. Zhang R., Isola P., Efros A. A. (2016) **Zhang, R., Isola, P. & Efros, A. A. Colorful Image Colorization. Preprint at 10.48550/arXiv.1603.08511 (2016).** <https://doi.org/10.48550/arXiv.1603.08511>
16. Noroozi M., Favaro P. (2017) **Noroozi, M. & Favaro, P. Unsupervised Learning of Visual Representations by Solving Jigsaw Puzzles. Preprint at 10.48550/arXiv.1603.09246 (2017).** <https://doi.org/10.48550/arXiv.1603.09246>
17. Donahue J., Simonyan K. (2019) **Large Scale Adversarial Representation Learning** *Advances in Neural Information Processing Systems* **32**
18. He K., Zhang X., Ren S., Sun J. (2015) **Deep Residual Learning for Image Recognition** *ArXiv151203385 Cs*
19. Zagoruyko S., Komodakis N. (2017) **Zagoruyko, S. & Komodakis, N. Wide Residual Networks. Preprint at 10.48550/arXiv.1605.07146 (2017).** <https://doi.org/10.48550/arXiv.1605.07146>
20. Krizhevsky A., Sutskever I., Hinton G. E. (2012) **ImageNet Classification with Deep Convolutional Neural Networks** *Advances in Neural Information Processing Systems* **25**:1097–1105
21. Szegedy C., et al. (2014) **Going Deeper with Convolutions** *ArXiv14094842 Cs*
22. Szegedy C., Vanhoucke V., Ioffe S., Shlens J., Wojna Z. (2015) **Szegedy, C., Vanhoucke, V., Ioffe, S., Shlens, J. & Wojna, Z. Rethinking the Inception Architecture for Computer Vision. Preprint at 10.48550/arXiv.1512.00567 (2015).** *Rethinking the Inception Architecture for Computer Vision* <https://doi.org/10.48550/arXiv.1512.00567>
23. Huang G., Liu Z., van der Maaten L., Weinberger K. Q. (2018) **Densely Connected Convolutional Networks** <https://doi.org/10.48550/arXiv.1608.06993>
24. Simonyan K., Zisserman A. (2015) **Simonyan, K. & Zisserman, A. Very Deep Convolutional Networks for Large-Scale Image Recognition. Preprint at 10.48550/arXiv.1409.1556 (2015).** <https://doi.org/10.48550/arXiv.1409.1556>
25. Howard A. G., et al. (2017) **Howard, A. G., et al. MobileNets: Efficient Convolutional Neural Networks for Mobile Vision Applications. Preprint at 10.48550/arXiv.1704.04861 (2017).** <https://doi.org/10.48550/arXiv.1704.04861>

26. Iandola F. N., et al. (2016) **Iandola, F. N., et al. SqueezeNet: AlexNet-level accuracy with 50x fewer parameters and <0.5MB model size. Preprint at 10.48550/arXiv.1602.07360 (2016).** <https://doi.org/10.48550/arXiv.1602.07360>
27. Xie S., Girshick R., Dollár P., Tu Z., He K. (2017) **Aggregated Residual Transformations for Deep Neural Networks** *ArXiv161105431 Cs*
28. Tan M., et al. (2019) **Tan, M., et al. MnasNet: Platform-Aware Neural Architecture Search for Mobile. Preprint at 10.48550/arXiv.1807.11626 (2019).** <https://doi.org/10.48550/arXiv.1807.11626>
29. Zhang X., Zhou X., Lin M., Sun J. (2017) **Zhang, X., Zhou, X., Lin, M. & Sun, J. ShuffeNet: An Extremely Efficient Convolutional Neural Network for Mobile Devices. Preprint at 10.48550/arXiv.1707.01083 (2017).** <https://doi.org/10.48550/arXiv.1707.01083>

Editors

Reviewing Editor

Srdjan Ostojic

École Normale Supérieure - PSL, Paris, France

Senior Editor

Joshua Gold

University of Pennsylvania, Philadelphia, United States of America

Reviewer #2 (Public Review):

Summary:

The dominant paradigm in the past decade for modeling the ventral visual stream's response to images has been to train deep neural networks on object classification tasks and regress neural responses from units of these networks. While object classification performance is correlated to variance explained in the neural data, this approach has recently hit a plateau of variance explained, beyond which increases in classification performance do not yield improvements in neural predictivity. This suggests that classification performance may not be a sufficient objective for building better models of the ventral stream. Lindsey & Issa study the role of factorization in predicting neural responses to images, where factorization is the degree to which variables such as object pose and lighting are represented independently in orthogonal subspaces. They propose factorization as a candidate objective for breaking through the plateau suffered by models trained only on object classification. They show the degree of factorization in a model captures aspects of neural variance that classification accuracy alone does not capture, hence factorization may be an objective that could lead to better models of ventral stream. I think the most important figure for a reader to see is Fig. 6.

Strengths:

This paper challenges the dominant approach to modeling neural responses in the ventral stream, which itself is valuable for diversifying the space of ideas.

This paper uses a wide variety of datasets, spanning multiple brain areas and species. The results are consistent across the datasets, which is a great sign of robustness.

The paper uses a large set of models from many prior works. This is impressively thorough and rigorous.

The authors are very transparent, particularly in the supplementary material, showing results on all datasets. This is excellent practice.

Weaknesses:

The authors have addressed many of the weaknesses in the original review. The weaknesses that remain are limitations of the work that cannot be easily addressed. In addition to the limitations stated at the end of the discussion, I'll add two:

(1) This work shows that factorization is correlated with neural similarity, and notably explains some variance in neural similarity that classification accuracy does not explain. This suggests that factorization could be used as an objective (along with classification accuracy) to build better models of the brain. However, this paper does not do that - using factorization to build better models of the brain is left to future work.

<https://doi.org/10.7554/eLife.91685.2.sa1>

Reviewer #3 (Public Review):

Summary:

Object classification serves as a vital normative principle in both the study of the primate ventral visual stream and deep learning. Different models exhibit varying classification performances and organize information differently. Consequently, a thriving research area in computational neuroscience involves identifying meaningful properties of neural representations that act as bridges connecting performance and neural implementation. In the work of Lindsey and Issa, the concept of factorization is explored, which has strong connections with emerging concepts like disentanglement [1,2,3] and abstraction [4,5]. Their primary contributions encompass two facets: (1) The proposition of a straightforward method for quantifying the degree of factorization in visual representations. (2) A comprehensive examination of this quantification through correlation analysis across deep learning models.

To elaborate, their methodology, inspired by prior studies [6], employs visual inputs featuring a foreground object superimposed onto natural backgrounds. Four types of scene variables, such as object pose, are manipulated to induce variations. To assess the level of factorization within a model, they systematically alter one of the scene variables of interest and estimate the proportion of encoding variances attributable to the parameter under consideration.

The central assertion of this research is that factorization represents a normative principle governing biological visual representation. The authors substantiate this claim by demonstrating an increase in factorization from macaque V4 to IT, supported by evidence from correlated analyses revealing a positive correlation between factorization and decoding performance. Furthermore, they advocate for the inclusion of factorization as part of the objective function for training artificial neural networks. To validate this proposal, the authors systematically conduct correlation analyses across a wide spectrum of deep neural networks and datasets sourced from human and monkey subjects. Specifically, their findings indicate that the degree of factorization in a deep model positively correlates with its predictability concerning neural data (i.e., goodness of fit).

Strengths:

The primary strength of this paper is the authors' efforts in systematically conducting analysis across different organisms and recording methods. Also, the definition of factorization is simple and intuitive to understand.

Weaknesses:

Comments on revised version:

I thank the authors for addressing the weaknesses I brought up regarding the manuscript.

<https://doi.org/10.7554/eLife.91685.2.sa0>

Author response:

The following is the authors' response to the original reviews.

eLife assessment

The study makes a valuable empirical contribution to our understanding of visual processing in primates and deep neural networks, with a specific focus on the concept of factorization. The analyses provide solid evidence that high factorization scores are correlated with neural predictivity, yet more evidence would be needed to show that neural responses show factorization. Consequently, while several aspects require further clarification, in its current form this work is interesting to systems neuroscientists studying vision and could inspire further research that ultimately may lead to better models of or a better understanding of the brain.

Public Reviews:

Reviewer #1 (Public Review):

Summary:

The paper investigates visual processing in primates and deep neural networks (DNNs), focusing on factorization in the encoding of scene parameters. It challenges the conventional view that object classification is the primary function of the ventral visual stream, suggesting instead that the visual system employs a nuanced strategy involving both factorization and invariance. The study also presents empirical findings suggesting a correlation between high factorization scores and good neural predictivity.

Strengths:

(1) Novel Perspective: The paper introduces a fresh viewpoint on visual processing by emphasizing the factorization of non-class information.

(2) Methodology: The use of diverse datasets from primates and humans, alongside various computational models, strengthens the validity of the findings.

(3) Detailed Analysis: The paper suggests metrics for factorization and invariance, contributing to a future understanding & measurements of these concepts.

Weaknesses:

(1) Vagueness (Perceptual or Neural Invariance?): The paper uses the term 'invariance', typically referring to perceptual stability despite stimulus variability [1], as the complete discarding of nuisance information in neural activity. This oversimplification overlooks the nuanced distinction between perceptual invariance (e.g., invariant object recognition) and neural invariance (e.g., no change in neural activity). It seems that by 'invariance' the authors mean 'neural' invariance (rather than 'perceptual' invariance) in this paper, which is vague. The paper could benefit from changing what is called 'invariance' in the paper to 'neural invariance' and distinguish it from 'perceptual invariance,' to avoid

potential confusion for future readers. The assignment of 'compact' representation to 'invariance' in Figure 1A is misleading (although it can be addressed by the clarification on the term invariance). [1] DiCarlo JJ, Cox DD. Untangling invariant object recognition. Trends in cognitive sciences. 2007 Aug 1;11(8):333-41.

Thanks for pointing out this ambiguity. In our Introduction we now explicitly clarify that we use “invariance” to refer to neural, rather than perceptual invariance, and we point out that both factorization and (neural) invariance may be useful for obtaining behavioral/perceptual invariance.

(2) Details on Metrics: The paper's explanation of factorization as encoding variance independently or uncorrelatedly needs more justification and elaboration. The definition of 'factorization' in Figure 1B seems to be potentially misleading, as the metric for factorization in the paper seems to be defined regardless of class information (can be defined within a single class). Does the factorization metric as defined in the paper (orthogonality of different sources of variation) warrant that responses for different object classes are aligned/parallel like in 1B (middle)? More clarification around this point could make the paper much richer and more interesting.

Our factorization metric measures the degree to which two sets of scene variables are factorized from one another. In the example of Fig. 1B, we apply this definition to the case of factorization of class vs. non-class information. Elsewhere in the paper we measure factorization of several other quantities unrelated to class, specifically camera viewpoint, lighting conditions, background content, and object pose. In our revised manuscript we have clarified the exposition surrounding Fig. 1B to make it clear that factorization, as we define it, can be applied to other quantities as well and that responses do not need to be aligned/parallel but simply live in a different set of dimensions whether linearly or nonlinearly arranged. Thanks for raising the need to clarify this point.

(3) Factorization vs. Invariance: Is it fair to present invariance vs. factorization as mutually exclusive options in representational hypothesis space? Perhaps a more fair comparison would be factorization vs. object recognition, as it is possible to have different levels of neural variability (or neural invariance) underlying both factorization and object recognition tasks.

We do not mean to imply that factorization and invariance are mutually exclusive, or that they fully characterize the space of possible representations. However, they are qualitatively distinct strategies for achieving behavioral capabilities like object recognition. In the revised manuscript we also include a comparison to object classification performance (Figures 5C & S4, black x's) as a predictor of brain-like representations, alongside the results for factorization and invariance.

In our revised Introduction and beginning of the Results section, we make it more clear that factorization and invariance are not mutually exclusive – indeed, our results show that both factorization and invariance for some scene variables like lighting and background identity are signatures of brain-like representations. Our study focuses on factorization because we believe its importance has not been studied or highlighted to the degree that invariance to “nuisance” parameters has in concert with selectivity to object identity in individual neuron tuning functions. Moreover, the loss functions used for supervised training functions of neural networks for image classification would seem to encourage invariance as a representational strategy. Thus, the finding that factorization of scene parameters is an equally good if not better predictor of brain-like representations may motivate new objective functions for neural network training.

(4) Potential Confounding Factors in Empirical Findings: The correlation observed in Figure 3 between factorization and neural predictivity might be influenced by data dimensionality, rather than factorization per se [2]. Incorporating discussions around this recent finding could strengthen the paper.

[2] Elmoznino E, Bonner MF. High-performing neural network models of the visual cortex benefit from high latent dimensionality. bioRxiv. 2022 Jul 13:2022-07.

We thank the Reviewer for pointing out this important, potential confound and the need for a direct quantification. We have now included an analysis computing how well dimensionality (measured using the participation ratio metric for natural images, as was done in [2] Elmoznino& Bonner bioRxiv. 2022) can account for model goodness-of-fit (additional pink bars in Figure 6). Factorization of scene parameters appears to add more predictive power than dimensionality on average (Figure 6, light shaded bars), and critically, factorization+classification jointly predict goodness-of-fit significantly better than dimensionality+classification for V4 and IT/HVC brain areas (Figure 6, dark shaded bars). Indeed, dimensionality+classification is only slightly more predictive than classification alone for V4 and IT/HVC indicating some redundancy in those measures with respect to neural predictivity of models (Figure 6, compare dark shaded pink bar to dashed line).

That said, high-dimensional representations can, in principle, better support factorization, and thus we do not regard these two representational strategies necessarily in competition. Rather, our results suggest (consistent with [2]) that dimensionality is predictive of brain-like representation to some degree, such that some (but not all) of factorization's predictive power may indeed owe to a partial correlation with dimensionality. We elaborate in the Discussion where this point comes up and now refer to the updated Figure 6 that shows the control for dimensionality.

Conclusion:

The paper offers insightful empirical research with useful implications for understanding visual processing in primates and DNNs. The paper would benefit from a more nuanced discussion of perceptual and neural invariance, as well as a deeper discussion of the coexistence of factorization, recognition, and invariance in neural representation geometry. Additionally, addressing the potential confounding factors in the empirical findings on the correlation between factorization and neural predictivity would strengthen the paper's conclusions.

Taken together, we hope that the changes described above address the distinction between neural and perceptual invariance, provide a more balanced understanding of the contributions of factorization, invariance, and local representational geometry, and rule against dimensionality for natural images as contributing to the main finding of the benefits from factorization of scene parameters.

Reviewer #2 (Public Review):

Summary:

The dominant paradigm in the past decade for modeling the ventral visual stream's response to images has been to train deep neural networks on object classification tasks and regress neural responses from units of these networks. While object classification performance is correlated to the variance explained in the neural data, this approach has recently hit a plateau of variance explained, beyond which increases in classification performance do not yield improvements in neural predictivity. This suggests that classification performance may not be a sufficient objective for building better models of

the ventral stream. Lindsey & Issa study the role of factorization in predicting neural responses to images, where factorization is the degree to which variables such as object pose and lighting are represented independently in orthogonal subspaces. They propose factorization as a candidate objective for breaking through the plateau suffered by models trained only on object classification.

They claim that (i) maintaining these non-class variables in a factorized manner yields better neural predictivity than ignoring non-class information entirely, and (ii) factorization may be a representational strategy used by the brain.

The first of these claims is supported by their data. The second claim does not seem well-supported, and the usefulness of their observations is not entirely clear.

Strengths:

This paper challenges the dominant approach to modeling neural responses in the ventral stream, which itself is valuable for diversifying the space of ideas.

This paper uses a wide variety of datasets, spanning multiple brain areas and species. The results are consistent across the datasets, which is a great sign of robustness.

The paper uses a large set of models from many prior works. This is impressively thorough and rigorous.

The authors are very transparent, particularly in the supplementary material, showing results on all datasets. This is excellent practice.

Weaknesses:

(1) The primary weakness of this paper is a lack of clarity about what exactly is the contribution. I see two main interpretations: (1-A) As introducing a heuristic for predicting neural responses that improve over-classification accuracy, and (1-B) as a model of the brain's representational strategy. These two interpretations are distinct goals, each of which is valuable. However, I don't think the paper in its current form supports either of them very well:

*(1-A) Heuristic for neural predictivity. The claim here is that by optimizing for factorization, we could improve models' neural predictivity to break through the current predictivity plateau. To frame the paper in this way, the key contribution should be a new heuristic that correlates with neural predictivity better than classification accuracy. The paper currently does not do this. The main piece of evidence that factorization may yield a more useful heuristic than classification accuracy alone comes from Figure 5. However, in Figure 5 it seems that factorization along some factors is more useful than others, and different linear combinations of factorization and classification may be best for different data. There is no single heuristic presented and defended. If the authors want to frame this paper as a new heuristic for neural predictivity, I recommend the authors present and defend a specific heuristic that others can use, e.g. $[K * \text{factorization_of_pose} + \text{classification}]$ for some constant K , and show that (i) this correlates with neural predictivity better than classification alone, and (ii) this can be used to build models with higher neural predictivity. For (ii), they could fine-tune a state-of-the-art model to improve this heuristic and show that doing so achieves a new state-of-the-art neural predictivity. That would be convincing evidence that their contribution is useful.*

Our paper does not make any strong claim regarding the Reviewer's point 1-A (on heuristics for neural predictivity). In the Discussion, last paragraph, we better specify that our work is merely suggestive of claim 1-A about heuristics for more neurally predictive, more brainlike

models. We believe that our paper supports the Reviewer's point 1-B (on brain representation) as we discuss below.

We leave it to future work to determine if factorization could help optimize models to be more brainlike. This treatment may require exploration of novel model architectures and loss functions, and potentially also more thorough neural datasets that systematically vary many different forms of visual information for validating any new models.

(1-B) Model of representation in the brain. The claim here is that factorization is a general principle of representation in the brain. However, neural predictivity is not a suitable metric for this, because (i) neural predictivity allows arbitrary linear decoders, hence is invariant to the orthogonality requirement of factorization, and (ii) neural predictivity does not match the network representation to the brain representation. A better metric is representational dissimilarity matrices. However, the RDM results in Figure S4 actually seem to show that factorization does not do a very good job of predicting neural similarity (though the comparison to classification accuracy is not shown), which suggests that factorization may not be a general principle of the brain. If the authors want to frame the paper in terms of discovering a general principle of the brain, I suggest they use a metric (or suite of metrics) of brain similarity that is sensitive to the desiderata of factorization, e.g. doesn't apply arbitrary linear transformations, and compare to classification accuracy in addition to invariance.

We agree with the Reviewer about the shortcomings of neural predictivity for comparing representational geometries, and in our revised manuscript we have provided a more comprehensive set of results that includes RDM predictivity in new Figures 6 & 7, alongside the results for neural fit predictivity. In addition, as suggested we added classification accuracy predictivity in Figures 5C & S4 (black x's) for visual comparison to factorization/invariance. In Figure S4 on RDMs, it is apparent how factorization is at least as good a predictor as classification on all V4 & IT datasets from both monkeys and humans (compared x's to filled circles in Figure S4; note that some of the points from the original Figure S4 changed as we discovered a bug in the code that specifically affected the RDM analysis for a few of the datasets).

We find that the newly included RDM analyses in Figures 6 & 7 are consistent with the conclusions of the neural fit regression analyses: that the correlation of factorization metrics with RDM matches are strong, comparable in magnitude to that of classification accuracy (Figure 6, 3rd & 4th columns, compare black dashed line to faded colored bars) and are not fully accounted for by the model's classification accuracy alone (Figure 6, 3rd & 4th columns, higher unfaded bars for classification combined with factorization, and see corresponding example scatters in Figure 7 middle/bottom rows).

It is encouraging that the added benefit of factorization for RDM predictivity accounting for classification performance is at least as good as the improvement seen for neural fit predictivity (Figure 6, 1st & 2nd columns for encoding fits versus 3rd & 4th columns for RDM correlations).

(2) I think the comparison to invariance, which is pervasive throughout the paper, is not very informative. First, it is not surprising that invariance is more weakly correlated with neural predictivity than factorization, because invariant representations lose information compared to factorized representations. Second, there has long been extensive evidence that responses throughout the ventral stream are not invariant to the factors the authors consider, so we already knew that invariance is not a good characterization of ventral stream data.

While we appreciate the Reviewer's intuition that highly invariant representations are not strongly supported in the high-level visual cortex, we nevertheless thought it was valuable to put this intuition to a quantitative, detailed test. As a result, we uncovered effects that were

not obvious a priori, at least to us – for example, that invariance for some scene parameters (camera view, object pose) is negatively correlated with neural predictions while invariance to others (background, lighting) is positively correlated. Thus, our work exercises the details of invariance for different types of information.

(3) The formalization of the factorization metric is not particularly elegant, because it relies on computing top K principal components for the other-parameter space, where K is arbitrarily chosen as 10. While the authors do show that in their datasets the results are not very sensitive to K (Figure S5), that is not guaranteed to be the case in general. I suggest the authors try to come up with a formalization that doesn't have arbitrary constants. For example, one possibility that comes to mind is $E[\delta_a \times \delta_b]$, where 'x' is the normalized cross product, δ_a , and δ_b are deltas in representation space induced by perturbations of factors a and b, and the expectation is taken over all base points and deltas. This is just the first thing that comes to mind, and I'm sure the authors can come up with something better. The literature on disentangling metrics in machine learning may be useful for ideas on measuring factorization.

Thanks to the Reviewer for raising this point. First, we wish to clarify a potential misunderstanding of the factorization metric: the number K of principal components we choose is not an arbitrary constant, but rather calibrated to capture a certain fraction of variance, set to 90% by default in our analyses. While this variance threshold is indeed an arbitrary hyperparameter, it has a more intuitive interpretation than the number of principal components.

Nonetheless, the Reviewer's comment did inspire us to consider another metric for factorization that does not depend on any arbitrary parameters. In the revised version, we now include a covariance matrix based metric which simply measures the elementwise correlation of the covariance matrices induced by varying the scene parameter of interest and the covariance matrix induced by varying the other parameters (and then subtracts this quantity from 1).

Correspondingly, we now present results for both the new covariance based measure and the original PCA based one in Figures 5C, 6, and 7. The main findings remain largely the same when using the covariance based metric, and the covariance based metric (Figure 5C, compare light shaded to dark shaded filled circles; Figure 6, compare top row to bottom row; Figure 7, compare middle rows to bottom rows).

Ultimately, we believe these two metrics are complementary and somewhat analogous to two metrics commonly used for measuring dimensionality (the number of components needed to explain a certain fraction of the variance, analogous to our original PCA based definition; the participation ratio, analogous to our covariance based definition). We have added the formula for the covariance based factorization metric along with a brief description to the Methods.

(4) The authors defined the term "factorization" according to their metric. I think introducing this new term is not necessary and can be confusing because the term "factorization" is vague and used by different researchers in different ways. Perhaps a better term is "orthogonality", because that is clear and seems to be what the authors' metric is measuring.

We agree with the Reviewer that factorization has become an overloaded term. At the same time, we think that in this context, the connotation of the term factorization effectively conveys the notion of separating out different latent sources of variance (factors) such that they can be encoded in orthogonal subspaces.

To aid clarity, we now mention in the Introduction that factorization defined here is meant to measure orthogonalization of scene factors. Additionally, in the Discussion section, we now go into more detail comparing our metric to others previously used in the literature, including orthogonality, to help put it in context.

(5) One general weakness of the factorization paradigm is the reliance on a choice of factors. This is a subjective choice and becomes an issue as you scale to more complex images where the choice of factors is not obvious. While this choice of factors cannot be avoided, I suggest the authors add two things: First, an analysis of how sensitive the results are to the choice of factors (e.g. transform the basis set of factors and re-run the metric); second, include some discussion about how factors may be chosen in general (e.g. based on temporal statistics of the world, independent components analysis, or something else).

The Reviewer raises a very reasonable point about the limitation of this work. While we limited our analysis to generative scene factors that we know about and that could be manipulated, there are many potential factors to consider. It is not clear to us exactly how to implement the Reviewer's suggestion of transforming the basis set of factors, as the factors we consider are highly nonlinear in the input space. Ultimately, we believe that finding unsupervised methods to characterize the "true" set of factors that is most useful for understanding visual representations is an important subject for future work, but outside the scope of this particular study. We have added a comment to this effect in the Discussion.

Reviewer #3 (Public Review):

Summary:

Object classification serves as a vital normative principle in both the study of the primate ventral visual stream and deep learning. Different models exhibit varying classification performances and organize information differently. Consequently, a thriving research area in computational neuroscience involves identifying meaningful properties of neural representations that act as bridges connecting performance and neural implementation. In the work of Lindsey and Issa, the concept of factorization is explored, which has strong connections with emerging concepts like disentanglement [1,2,3] and abstraction [4,5]. Their primary contributions encompass two facets: (1) The proposition of a straightforward method for quantifying the degree of factorization in visual representations. (2) A comprehensive examination of this quantification through correlation analysis across deep learning models.

To elaborate, their methodology, inspired by prior studies [6], employs visual inputs featuring a foreground object superimposed onto natural backgrounds. Four types of scene variables, such as object pose, are manipulated to induce variations. To assess the level of factorization within a model, they systematically alter one of the scene variables of interest and estimate the proportion of encoding variances attributable to the parameter under consideration.

The central assertion of this research is that factorization represents a normative principle governing biological visual representation. The authors substantiate this claim by demonstrating an increase in factorization from macaque V4 to IT, supported by evidence from correlated analyses revealing a positive correlation between factorization and decoding performance. Furthermore, they advocate for the inclusion of factorization as part of the objective function for training artificial neural networks. To validate this proposal, the authors systematically conduct correlation analyses across a wide spectrum of deep neural networks and datasets sourced from human and monkey subjects. Specifically, their findings indicate that the degree of factorization in a deep

model positively correlates with its predictability concerning neural data (i.e., goodness of fit).

Strengths:

The primary strength of this paper is the authors' efforts in systematically conducting analysis across different organisms and recording methods. Also, the definition of factorization is simple and intuitive to understand.

Weaknesses:

This work exhibits two primary weaknesses that warrant attention: (i) the definition of factorization and its comparison to previous, relevant definitions, and (ii) the chosen analysis method.

Firstly, the definition of factorization presented in this paper is founded upon the variances of representations under different stimuli variations. However, this definition can be seen as a structural assumption rather than capturing the effective geometric properties pertinent to computation. More precisely, the definition here is primarily statistical in nature, whereas previous methodologies incorporate computational aspects such as deviation from ideal regressors [1], symmetry transformations [3], generalization [5], among others. It would greatly enhance the paper's depth and clarity if the authors devoted a section to comparing their approach with previous methodologies [1,2,3,4,5], elucidating any novel insights and advantages stemming from this new definition.

[1] Eastwood, Cian, and Christopher KI Williams. "A framework for the quantitative evaluation of disentangled representations." International conference on learning representations. 2018.

[2] Kim, Hyunjik, and Andriy Mnih. "Disentangling by factorising." International Conference on Machine Learning. PMLR, 2018.

[3] Higgins, Irina, et al. "Towards a definition of disentangled representations." arXiv preprint arXiv:1812.02230 (2018).

[4] Bernardi, Silvia, et al. "The geometry of abstraction in the hippocampus and prefrontal cortex." Cell 183.4 (2020): 954-967.

[5] Johnston, W. Jeffrey, and Stefano Fusi. "Abstract representations emerge naturally in neural networks trained to perform multiple tasks." Nature Communications 14.1 (2023): 1040.

Thanks to the Reviewer for this suggestion. We agree that our initial submission did not sufficiently contextualize our definition of factorization with respect to other related notions in the literature. We have added additional discussion of these points to the Discussion section in the revised manuscript and have included therein the citations provided by the Reviewer (please see the third paragraph of Discussion).

Secondly, in order to establish a meaningful connection between factorization and computation, the authors rely on a straightforward synthetic model (Figure 1c) and employ multiple correlation analyses to investigate relationships between the degree of factorization, decoding performance, and goodness of fit. Nevertheless, the results derived from the synthetic model are limited to the low training-sample regime. It remains unclear whether the biological datasets under consideration fall within this low training-sample regime or not.

We agree that our model in Figure 1C is very simple and does not fully capture the complex interactions between task performance and features of representational geometry, like factorization. We intend it only as a proof of concept to illustrate how factorized representations can be beneficial for some downstream task use cases. While the benefits of factorized representations disappear for large numbers of samples in this simulation, we believe this is primarily a consequence of the simplicity and low dimensionality of the simulation. Real-world visual information is complex and high-dimensional, and as such the relevant sample size regime in which factorization offers tasks benefits may be much greater. As a first step toward this real-world setting, Figure 2 shows how decreasing the amount of factorization in neural population data in macaque V4/IT can have an effect on object identity decoding.

Recommendations for the authors

Reviewer #1 (Recommendations For The Authors):

Missing citations: The paper could benefit from discussions & references to related papers, such as:

Higgins I, Chang L, Langston V, Hassabis D, Summerfield C, Tsao D, Botvinick M. Unsupervised deep learning identifies semantic disentanglement in single inferotemporal face patch neurons. Nature communications. 2021 Nov 9;12(1):6456.

We have added additional discussion of related work, including the suggested reference and others on disentanglement, to the Discussion section in the revised manuscript.

Reviewer #2 (Recommendations For The Authors):

Here are several small recommendations for the authors, all much more minor than those in the public review:

I suggest more use of equations in methods sections about Figure 1C and macaque neural data analysis.

Thanks for this suggestion. We have added new Equation 1 for the method transforming neural data to reduce factorization of a variable while preserving other firing rate statistics.

In Figure 1-C, the methods indicate that Gaussian noise was added. This is a very important detail, and complexifies the interpretation of the figure because it adds an assumption about the structure of noise. In other words, if I understand correctly, the correct interpretation of Figure 1C is "assuming i.i.d. noise, decoding accuracy improves with factorization." The i.i.d. noise is a big assumption, and it is debated how well the brain satisfies this assumption. I suggest you either omit noise for this figure or clearly state in the main text (e.g. caption) that the figure must be interpreted under an i.i.d. noise assumption.

We have added an explicit statement of the i.i.d. noise assumption to the Figure 1C legend.

For Figure 2B, I suggest labeling the x-axis clearly below the axis on both panels. Currently, it is difficult to read, particularly in print.

We have made the x-axis labels more clear and included on both panels.

Figure 3A is difficult to read because of the very small task. I suggest avoiding such small fonts.

We agree that Figure 3A is difficult to read. We have broken out Figure 3 into two new Figures 3 & 4 to increase clarity and sizing of text in Figure 3A.

Reviewer #3 (Recommendations For The Authors):

To strengthen this work, it is advisable to incorporate more comprehensive comparisons with previous research, particularly within the machine learning (ML) community. For instance, it would be beneficial to explore and reference works focusing on disentanglement [1,2,3]. This would provide valuable context and facilitate a more robust understanding of the contributions and novel insights presented in the current study.

We have added additional discussion of related work and other notions similar to factorization to the Discussion section in the revised manuscript.

Additionally, improving the quality of the figures is crucial to enhance the clarity of the findings:

- *Figure 2: The caption of subfigure B could be revised for greater clarity.*

Thank you, we have substantially clarified this figure caption.

- *Figure 3: Consider a more equitable approach for computing the correlation coefficient, such as calculating it separately for different types of models. In the case of supervised models, it appears that the correlation between invariance and goodness of fit may not be negligible across various scene parameters.*

We appreciate the suggestion, but we are not confident in our ability to conclude much from analyses restricted to particular model classes, given the relatively small N and the fact that the different model classes themselves are an important source of variance in our data.

- *Figure 4: To enhance the interpretability of subfigures A and B, it may be beneficial to include p-values (indicating confidence levels).*

As we supply bootstrapped confidence intervals for our results, which provide at least as much information as p-values, and most of the effects of interest are fairly stark when comparing invariance to factorization, p-values were not needed to support our points. We added a sentence to the legend of new Figure 5 (previously Figure 4) indicating that error bars reflect standard deviations over bootstrap resampling of the models.

- *Figure 5: For subfigure B, it could be advantageous to plot the results solely for factorization, allowing for a clear assessment of whether the high correlation observed in Classification+Factorization arises from the combined effects of both factors or predominantly from factorization alone.*

First, we clarify/note that the scatters solely for factorization that the Reviewer seeks are already presented earlier in the manuscript across all conditions in Figures 4A,B and Figure S2.

While we could also include these in new Figure 7 (previously Figure 5B) as the Reviewer suggests, we believe it would distract from the message of that figure at the end of the manuscript – which is that factorization is useful as a supplement to classification in predictive matches to neural data. Nonetheless, new Figure 6 (old Figure 5A) provides a

summary quantification of the information that the reviewer requests (Fig. 6, faded colored bars reflect the contribution of factorization alone).

<https://doi.org/10.7554/eLife.91685.2.sa3>