

Adapting Deep Network Features to Capture Psychological Representations

Joshua C. Peterson (jpeterson@berkeley.edu)

Joshua T. Abbott (joshua.abbott@berkeley.edu)

Thomas L. Griffiths (thomas.griffiths@berkeley.edu)

Department of Psychology, University of California, Berkeley, CA 94720 USA

Abstract

Deep neural networks have become increasingly successful at solving classic perception problems such as object recognition, semantic segmentation, and scene understanding, often reaching or surpassing human-level accuracy. This success is due in part to the ability of DNNs to learn useful representations of high-dimensional inputs, a problem that humans must also solve. We examine the relationship between the representations learned by these networks and human psychological representations recovered from similarity judgments. We find that deep features learned in service of object classification account for a significant amount of the variance in human similarity judgments for a set of animal images. However, these features do not capture some qualitative distinctions that are a key part of human representations. To remedy this, we develop a method for adapting deep features to align with human similarity judgments, resulting in image representations that can potentially be used to extend the scope of psychological experiments.

Keywords: deep learning; neural networks; psychological representations; similarity

Introduction

The resurgence of neural networks in the form of *deep learning* has continued to dominate object recognition benchmarks in the field of computer vision, often attaining near or above human-level accuracy for a variety of perceptual tasks, most notably through recent advances in classifying thousands of objects within natural images (Krizhevsky, Sutskever, & Hinton, 2012; He, Zhang, Ren, & Sun, 2015). Part of the success of these models is due to their ability to learn effective feature representations of high-dimensional inputs (e.g., complex color images); a challenge that human perception must also confront (Austerweil & Griffiths, 2013). As a result, cognitive scientists have started to explore how the representations learned by these networks can be used in models of human behavior for perceptual tasks such as predicting the memorability of objects in images (Dubey, Peterson, Khosla, Yang, & Ghanem, 2015) and predicting judgments of category typicality (Lake, Zaremba, Fergus, & Gureckis, 2015).

While deep learning models continue to mimic a growing list of human-like abilities, a number of core questions remain unanswered about the relevance of these models to actual human cognition and perception. For instance, features of the input learned using these networks excel in predicting certain human judgments, but how are these feature representations related to human psychological representations? At first glance, it would seem that the ability of these representations to predict typicality judgments and stimulus memorability would constitute robust evidence of their relevance to people, however recent work has shown that neural networks

that classify images can be systematically deceived by imperceptible image transformations (Szegedy et al., 2013), casting doubt on their similarity to humans.

Understanding the relationship between the representations found by deep learning and those of humans is an important question in cognitive science, and could potentially benefit artificial intelligence. However, independent of this question, simply having a good approximation to how people represent images would allow cognitive scientists to test psychological theories using complex, realistic stimuli. Indeed, tasks such as creating stimulus sets that uniformly span psychological space are far from trivial.

In this paper, we address this question directly by examining how well features extracted from state-of-the-art deep neural networks predict human similarity judgments. An initial evaluation shows that these features account for a significant amount of variance in human judgments, but fail to capture qualitative distinctions that are key to human representations. We then develop a method for adapting deep network features to better predict human similarity judgments, and show that this approach can reproduce those qualitative distinctions. These results suggest that while raw features produced by deep learning may not be suitable for use in modeling cognition, they can be modified to bring them into close alignment with human representations.

Deep Representations

In general, deep neural networks (DNN) are neural networks that have depth in terms of their number of hidden layers between input and output (Bengio, 2009). In the past few years, training such networks to understand aspects of large, complex data sets has led to a number of advances in vision and language applications (LeCun, Bengio, & Hinton, 2015).

In computer vision, the majority of this progress has been driven by a particular DNN called a convolutional neural network (CNN) (LeCun et al., 1989). CNNs get their name from the use of convolutional layers, which learn a set of image filters that produce feature maps of spatially-organized inputs like images. This allows for a drastic decrease in the number of parameters the network must learn, which would otherwise explode exponentially in a fully connected network with high-dimensional inputs. The typical CNN architecture includes a series of hidden convolutional layers, followed by a smaller number of fully connected layers, and finally a layer that generates the final output or classification. While CNNs were initially developed over two decades ago, they came to mainstream popularity in 2012 when a 7-layer architecture named AlexNet (Krizhevsky, Sutskever, & Hin-

ton, 2012) won the ImageNet Large Scale Visual Recognition Challenge (ILSVRC), reducing the previous winner's error rate by an uncommonly large margin. Since then, a deeper CNN has won the contest every year, currently dominated by Microsoft's 150-layer network which obtained a best-of-top-5 error rate of 4.94%, surpassing the accuracy of non-expert humans at 5.1% (He, Zhang, Ren, & Sun, 2015).

Interestingly, CNNs produce much more than just their outputs (e.g., a category label for an image); they can also return feature representations at each layer of the network. The "deep representations" learned by these networks have proven useful in predicting human behavior. Dubey, Peterson, Khosla, Yang, & Ghanem (2015) used representations extracted from the last fully-connected layer of a CNN to predict the intrinsic memorability of objects. That is, the objects that humans are jointly likely to remember or forget in a large complex natural scene database. The correlation between estimates of memorability and the original memorability scores for each object matched human consistency (i.e. the correlation between memorability scores of random splits of the full sample of subjects). Similarly, Lake, Zaremba, Fergus, & Gureckis (2015) were able to reliably predict human typicality ratings of eight object categories using the same network and features, and called for cognitive scientists to pay attention to deep learning since categorization is a foundational problem in the field.

Deep representations are also beginning to interest the neuroscience community. For example, CNN activations have been used to predict monkey IT cortex activity (Yamins et al., 2014), as well as both low- and high-level activity in human visual areas (Agrawal, Stansbury, Malik, & Gallant, 2014). Delving deeper, Khaligh-Razavi & Kriegeskorte (2014) found that a CNN best explained IT cortex representations out of a set of 37 well-known models from both the computer vision and neuroscience fields, although no model completely explained all of the variance, unsupervised models being the worst of all of them.

Although CNN representations currently do the best job of predicting neural activity as measured by Blood Oxygenation Level Dependent (BOLD) response, this does not guarantee that we can explain psychological representations as a result. In fact, Mur et al. (2013) was partly successful in predicting human similarity judgments (a classic index of psychological representations) from IT cortex representations. However, the key categorical distinctions in the human representations were not well predicted: human IT cortex representations were more similar to monkey IT cortex representations than they were to human psychological representations. In the remainder of the paper, we use a similar approach to evaluate how well deep network features align with human psychological representations, and to explore how the correspondence between the two can be increased.

Evaluating Representations

Our first step is to evaluate the potential correspondence between deep network features and psychological representations. Unlike neural representations, psychological representations cannot be measured directly. However, both spatial and hierarchical psychological representations for N objects can be recovered given an $N \times N$ matrix of similarity judgments using methods such as multidimensional scaling and hierarchical clustering (Shepard, 1980). We thus reduce the problem to one of capturing human similarity judgments, subjecting both human judgments and model predictions to these different methods of extracting representations. We approach this problem by taking the inner-product of the deep feature representations of each pair of images (a measure of similarity between two vectors). We then compute the correlation between these pairwise vector similarities and human similarity judgments for the same stimulus pairs, which gives us a measure of the correspondence we want to evaluate.

Stimuli. Our stimulus set consisted of 120 color photographs of animals (sample images are shown in Figure 1). Images were cropped to 300×300 pixels, resulting in close-ups of either the animal's face or body. The set was constructed to include both inter- and intraspecies variation.

Behavioral Experiment. We collected pairwise similarity ratings for our animal stimulus set through Amazon Mechanical Turk. Participants were instructed to rate the similarity of four pairs of animal images on a scale from 0 (not similar at all) to 10 (very similar). We paid workers \$0.02 per set of four comparisons. Before each task, eight examples were shown to help prevent bias in early judgments. Amazon workers could repeat the task with new animal pairs as many times as they wanted. There were 7,140 possible image comparisons, each of which was rated by 10 unique participants, for a total of 71,400 ratings from 209 different participants. The result was a 120×120 similarity matrix after averaging over judgments.

Feature Extraction. We extracted features for each image in our data set using three different popular *off-the-shelf* CNNs of varying complexity that were pretrained in Caffe (Jia et al., 2014). Specifically, we used CaffeNet (based on original AlexNet), VGG16 (Simonyan & Zisserman, 2014), and GoogLeNet (Szegedy et al., 2014), the layer depths of which were 7, 16, and 22 respectively. GoogLeNet and VGG16 achieve roughly half the error rates of AlexNet. Each network had already been trained to classify 1000 object categories from previous ILSVRC competitions. A feedforward pass of each flattened image vector into each network yields feature responses at each layer. For our analysis, we extracted the last layer of each network before the classification layer. For CaffeNet and VGG16, this is a 4096-dimensional fully-connected layer, while the last



Figure 1: Samples from the set of 120 animal photographs.

Table 1: Correlations between human and deep similarities.

	CaffeNet	Google	VGG	HOG+SIFT
R^2	.32	.35	.43	.008

layer in GoogleNet is a 1000-dimensional average pooling layer. Lastly, we also extracted Histograms of Oriented Gradients (HOG) and Scale-Invariant Feature Transform (SIFT) representations for comparison since such features represent the generic representations of choice for tasks in computer vision prior to the popularity of deep learning.

Results

Table 1 gives performance (R^2) for each model. Raw representations from all three networks show medium to high correlations with the human data. In general, deeper networks with better ImageNet classification accuracy like GoogLeNet and VGG16 did better than CaffeNet, which is considerably more shallow. The HOG+SIFT baseline did surprisingly poorly, explaining very little variance as compared to the deep representations, suggesting that while these features are useful for many computer vision tasks, they differ in large part from the representations humans employ when judging animal similarity.

Although the VGG representation explained a fair amount of variance, further analyses revealed that the most crucial structural aspects of the human representations were not preserved. The first and second panels of Figure 2 show multi-dimensional scaling (MDS) solutions for the original human data and the predictions from the unaltered deep representations. While the structure of the MDS solutions for the predicted judgments looks reasonable (e.g., zebras are next to other zebras), major categorical divisions are not preserved. Hierarchical clusterings of the actual and predicted human judgments (the first and second panels of Figure 3) show a similar pattern of results: human judgments exhibit several major categorical divisions, whereas much of this structure is lost in the predicted data.

Adapting Representations

After quantifying the discrepancy between deep and human representations, we can attempt to bring them into closer alignment. First, consider that the final hidden layer feature representation in a neural network can be thought of as the input to a final linear classification layer, such that the problem solved by the final weight matrix is a linear transformation (which is then often scaled by a softmax function to convert to class probabilities). This can be thought of as a rescaling of the final stimulus representation to solve the categorization problem. This suggests that we should not think about the features extracted by the network as a static representation, but as the ingredients for a transformation that solves a problem. Thinking in these terms, we show that we can easily solve for a linear transformation that better captures human similarity judgments.

Similarity Model. Any similarity matrix $\mathbf{S}$ can be decomposed into the matrix product of a feature-by-object matrix $\mathbf{F}$, its transpose, and a diagonal weight matrix $\mathbf{W}$,

$$\mathbf{S} = \mathbf{F}\mathbf{W}\mathbf{F}^T \quad (1)$$

This formulation is similar to that employed by additive clustering models (Shepard & Arabie, 1979), wherein $\mathbf{F}$ represents a binary feature identity matrix (and is similar to Tversky’s (1977) model of similarity). When used with continuous features, this approach is akin to factor analysis. Given an existing feature-by-object matrix $\mathbf{F}$, the diagonal of $\mathbf{W}$ can be solved for using linear regression where the predictors for each similarity s_{ij} are the product of the values of each feature for the objects i and j . When $\mathbf{W}$ is the identity matrix, this reduces to the model evaluated in the previous section.

$$s_{ij} = \sum_{k=1}^{N_f} w_k f_{ik} f_{jk}. \quad (2)$$

The result is a convex optimization problem that can be solved straightforwardly, allowing us to find a transformation of the deep features with a closer correspondence to human similarity judgments.

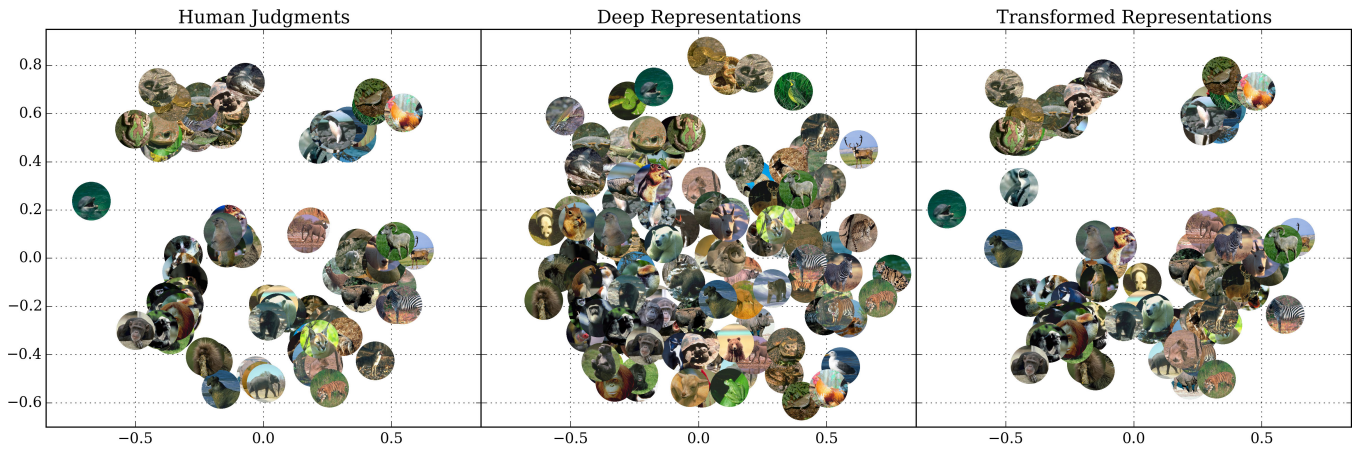


Figure 2: Multidimensional scaling solutions for similarity matrices obtained from human judgements (left), non-transformed deep representations (center), and transformed deep representations (right).

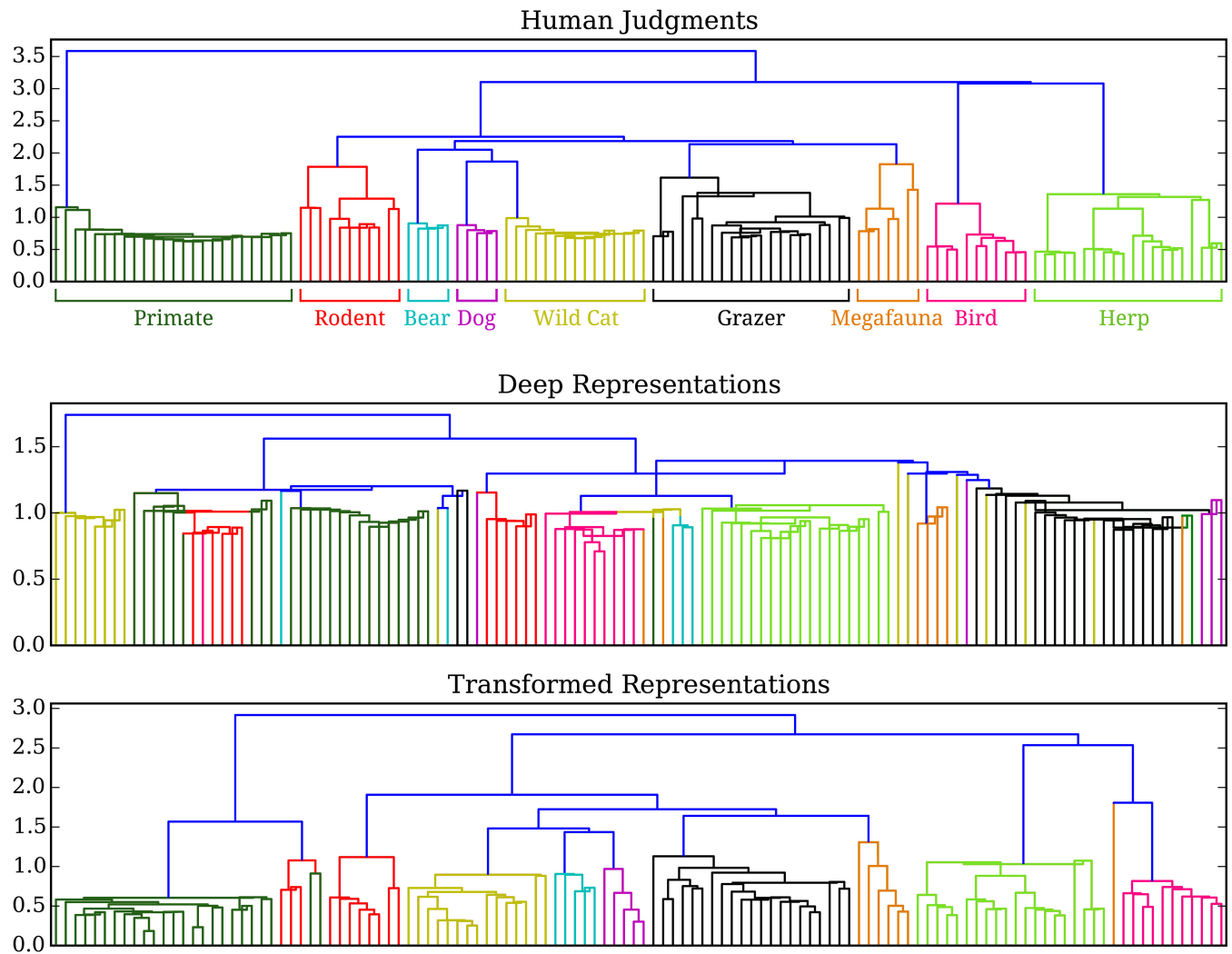


Figure 3: Hierarchical clustering of human judgements (top), deep representations (middle), and transformed representations (bottom). Human judgments resulted in nine interpretable clusters, grouped by color and semantic category label in the top panel. The leaves of the deep and transformed representation clusterings are color-coded relative to the human judgements.

Table 2: Model performance using adjusted CNN features.

	CaffeNet	Google	VGG	SIFT
R^2	.69	.72	.84	.09

Analysis. With such a large number of predictors, regularization is critical to avoid overfitting. We used ridge regression ($L2$ regularization) and performed grid search on cross-validated generalization performance to find the best regularization parameter. We predicted only the upper triangle of the similarity matrix since the matrix is symmetric. Each model was evaluated via its generalization performance in 6-fold cross-validation. We did this for the feature vectors extracted at each layer of the network.

As an additional control against overfitting, we compared model performance with several baselines. In Baseline 1, we shuffled the rows of the feature matrix such that the feature representation of one image was replaced with that of a different randomly chosen image. In Baseline 2, the columns of the feature matrix were randomly permuted for each row separately. Lastly, Baseline 3 simply combined the shuffling schemes from the first two baselines. In all three cases, the randomized feature matrices were subjected to the same set of analyses as the true features, allowing us to check for spurious correlations.

Results

Table 2 shows performance for each network using our adjustment of the representations. The R^2 values reported are the average values across all six folds of the crossvalidation. All five models performed considerably well, each showing improvement over the original non-weighted models. Most notably, VGG16 performed best, accounting for 84% of the variance. Training using the estimated regularization parameter on the entire dataset yielded an R^2 of 91%. In contrast, all three baseline models explained essentially no variance ($R^2 < 0.01$), suggesting that our results were not spurious correlations resulting from our large sets of predictors.

Crucially, the MDS solution for the improved predictions is almost identical to the original human spatial representation. The same improvements were found in hierarchical clusterings of actual and predicted similarity matrices (1st and 3rd panels of Figure 3), this time largely in the form of top-level parent nodes.

Feature Analysis. While higher layers in CNNs tend to produce the most generic high-level features for domain transfer across image applications, the choice of feature depth is ultimately dependent on the task (Sainath, Kingsbury, Mohamed, & Ramabhadran, 2013). This implies that layer responses at different depths may explain different types of human similarity judgments (e.g. tasks that involve comparing visual features versus conceptual information). We examined

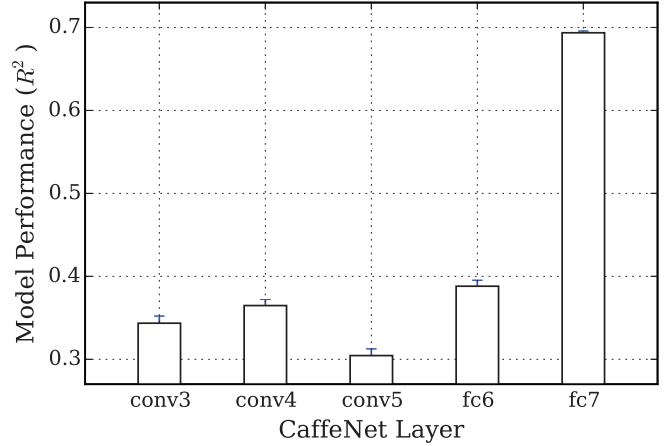


Figure 4: Model performance as a function of feature layer depth in CaffeNet.

our model’s performance in predicting similarity judgments as a function of feature depth using CaffeNet, given its more straightforward architecture and manageably-sized layers. Specifically, we compared performance across the last three convolutional layers and the last two fully-connected layers. The results are shown in Figure 4. Performance does appear to correspond strongly to layer depth, although fully connected layers perform much better than convolutional layers, suggesting that human similarity judgments may not be explained well from simpler image features.

Reweighted Classification. We investigated the effect of our fine-tuned representations on a separate animal classifier, using a new animal data set consisting of 1,740 images from 19 animal classes (*bear, cougar, cow, coyote, deer, elephant, giraffe, goat, gorilla, horse, kangaroo, leopard, lion, panda, penguin, sheep, skunk, tiger, zebra*) (Afkhani, Targhi, Eklundh, & Pronobis, 2008). We used multinomial logistic regression with 6-fold cross-validation to classify animals using fine-tuned representations as predictors. We fine-tuned these representations by pairwise multiplying the original VGG16 representations with the square-root of the weights obtained through prediction of the human similarity data. However, because some of the weights of the original solution are negative, we used Elastic Net regression to solve for weights constrained to be positive. We ran the same model using the original unaltered VGG16 representation to serve as baseline performance. The original model performed very well (average $R^2 = .94$), whereas the fine-tuned model performed consistently worse ($R^2 = .89$).

Discussion

This analysis constitutes the first formal comparison of deep representations to human psychological representations. Initial results using currently high-performing convolutional neural networks show that the two representations were moderately correlated, but diverge in terms of crucial structural

characteristics, a problem exhibited by similar experiments using neural representations as opposed to deep features (Mur et al., 2013).

Our method of overcoming this problem, by a parsimonious adjustment of the feature representation inspired by a classic model of similarity, appears to have been largely successful. Indeed, the human representations were almost completely reconstructed by our adjusted CNN features. Using features extracted from deep convolutional neural networks provides an opportunity to estimate psychological representations from real, raw sensory inputs (e.g. pixels). However, one potential limitation of this work is the generalizability of the transformation acquired to broader stimulus contexts. Testing this question will require replication and transfer across several domains. To the extent that this can be established, we envision our method as a standard tool for studying cognitive science using natural stimulus sets, on par with modern artificial intelligence.

Beyond this, we see potential for such an interface between cognitive science and artificial intelligence to be exploited for the benefit of each. While our attempt to improve a common categorization objective in computer vision (i.e. one-versus-all classification) using human-tuned representations was not successful, it does raise interesting distinctions between the computational problems solved by humans and CNNs. After all, the full breadth of human categorization behavior exhibits complex patterns such as overlapping class assignments, which are not likely to be well-represented when the learning objective is defined through images and objects characterized by a single label. Further, one might ask if poor categorization performance of the one-versus-all kind is the price paid for a more flexible system of categorization with respect to a set of complex objects that can be partitioned using several “good” configurations, depending on the context and task at hand. Given this possibility, one should be careful not to equate CNN classification performance with human categorization abilities in general.

Acknowledgments. This work was supported by grant number FA9550-13-1-0170 from the Air Force Office of Scientific Research. We thank Alex Huth for help with image selection.

References

- Afkham, H. M., Targhi, A. A. T., Eklundh, J.-O., & Pronobis, J.-O. E. A. (2008). Joint visual vocabulary for animal classification. In *19th International Conference on Pattern Recognition (ICPR)* (pp. 1–4).
- Agrawal, P., Stansbury, D., Malik, J., & Gallant, J. L. (2014). Pixels to Voxels: Modeling Visual Representation in the Human Brain. *arXiv preprint arXiv:1407.5104*.
- Austerweil, J. L., & Griffiths, T. L. (2013). A nonparametric Bayesian framework for constructing flexible feature representations. *Psychological Review*, 120(4), 817.
- Bengio, Y. (2009). Learning deep architectures for AI. *Foundations and Trends in Machine Learning*, 2(1), 1–127.
- Dubey, R., Peterson, J., Khosla, A., Yang, M.-H., & Ghanem, B. (2015). What makes an object memorable? In *International Conference on Computer Vision (ICCV)*.
- He, K., Zhang, X., Ren, S., & Sun, J. (2015). Delving deep into rectifiers: Surpassing human-level performance on imagenet classification. *arXiv preprint arXiv:1502.01852*.
- Jia, Y., Shelhamer, E., Donahue, J., Karayev, S., Long, J., Girshick, R., ... Darrell, T. (2014). Caffe: Convolutional architecture for fast feature embedding. *arXiv preprint arXiv:1408.5093*.
- Khaligh-Razavi, S.-M., & Kriegeskorte, N. (2014). Deep supervised, but not unsupervised, models may explain it cortical representation. *PLoS Comput Biol*, 10(11).
- Krizhevsky, A., Sutskever, I., & Hinton, G. E. (2012). Imagenet classification with deep convolutional neural networks. In *Advances in Neural Information Processing Systems* (pp. 1097–1105).
- Lake, B. M., Zaremba, W., Fergus, R., & Gureckis, T. M. (2015). Deep neural networks predict category typicality ratings for images. In *Proceedings of the 37th Annual Cognitive Science Society*.
- LeCun, Y., Bengio, Y., & Hinton, G. (2015). Deep learning. *Nature*, 521(7553), 436–444.
- LeCun, Y., Boser, B., Denker, J. S., Henderson, D., Howard, R. E., Hubbard, W., & Jackel, L. D. (1989). Backpropagation applied to handwritten zip code recognition. *Neural computation*, 1(4), 541–551.
- Mur, M., Meys, M., Bodurka, J., Goebel, R., Bandettini, P. A., & Kriegeskorte, N. (2013). Human object-similarity judgments reflect and transcend the primate-IT object representation. *Frontiers in Psychology*, 4.
- Sainath, T. N., Kingsbury, B., Mohamed, A.-r., & Ramabhadran, B. (2013). Learning filter banks within a deep neural network framework. In *Ieee workshop on Automatic Speech Recognition and Understanding* (pp. 297–302).
- Shepard, R. N. (1980). Multidimensional scaling, tree-fitting, and clustering. *Science*, 210(4468), 390–398.
- Shepard, R. N., & Arabie, P. (1979). Additive clustering: Representation of similarities as combinations of discrete overlapping properties. *Psychological Review*, 86(2), 87.
- Simonyan, K., & Zisserman, A. (2014). Very deep convolutional networks for large-scale image recognition. *arXiv preprint arXiv:1409.1556*.
- Szegedy, C., Liu, W., Jia, Y., Sermanet, P., Reed, S., Anguelov, D., ... Rabinovich, A. (2014). Going deeper with convolutions. *arXiv preprint arXiv:1409.4842*.
- Szegedy, C., Zaremba, W., Sutskever, I., Bruna, J., Erhan, D., Goodfellow, I., & Fergus, R. (2013). Intriguing properties of neural networks. *arXiv preprint arXiv:1312.6199*.
- Yamins, D. L., Hong, H., Cadieu, C. F., Solomon, E. A., Seibert, D., & DiCarlo, J. J. (2014). Performance-optimized hierarchical models predict neural responses in higher visual cortex. *Proceedings of the National Academy of Sciences*, 111(23), 8619–8624.