
System Identification of Neural Systems: If We Got It Right, Would We Know?

Yena Han¹ Tomaso Poggio¹ Brian Cheung¹

Abstract

Artificial neural networks are being proposed as models of parts of the brain. The networks are compared to recordings of biological neurons, and good performance in reproducing neural responses is considered to support the model’s validity. A key question is how much this system identification approach tells us about brain computation. Does it validate one model architecture over another? We evaluate the most commonly used comparison techniques, such as a linear encoding model and centered kernel alignment, to correctly identify a model by replacing brain recordings with known ground truth models. System identification performance is quite variable; it also depends significantly on factors independent of the ground truth architecture, such as stimuli images. In addition, we show the limitations of using functional similarity scores in identifying higher-level architectural motifs.

1. Introduction

Over the last two decades, the dominant approach for machine learning engineers in search of better performance has been to use standard benchmarks to rank networks from most relevant to least relevant. This practice has driven much of the progress in the machine learning community. A standard comparison benchmark enables the broad validation of successful ideas. Recently such benchmarks have found their way into neuroscience with the advent of experimental frameworks like Brain-Score (Schrimpf et al., 2020) and Algonauts (Cichy et al., 2021), where artificial models compete to predict recordings from real neurons in animal brains. Can engineering approaches like this be helpful in the natural sciences?

¹Center for Brains, Minds and Machines, Massachusetts Institute of Technology. Correspondence to: Yena Han <yena-han@mit.edu>.

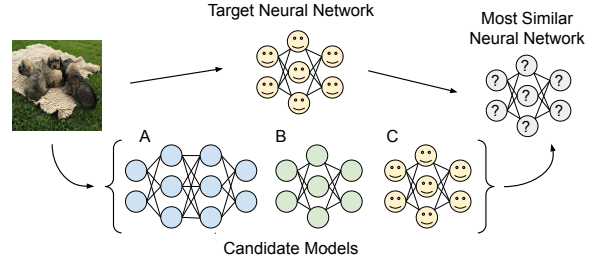


Figure 1. Illustration of the research question: If one of the candidate models matches the target neural network (brain) closely, can the current model evaluation methods accurately find that out?

The answer is clearly yes: the “engineering approach” described above ranks models that predict neural responses better as better models of animal brains. While such rankings may be a good measure of absolute performance in approximating the neural responses, which on its own is valuable for various applications (Bashivan et al., 2019), it is an open question whether they are sufficient. In neuroscience, understanding natural intelligence at the level of the underlying neural circuits requires developing model systems that reproduce the abilities of their biological analogs while respecting the constraints provided by biology, including anatomy and biophysics (Marr & Poggio, 1976; Schaeffer et al., 2022). A model that reproduces neural responses well but turns out to require connectivity or biophysical mechanisms that are different from the biological ones is thereby falsified.

Consider the conjecture that the similarity of responses between model units and brain neurons allows us to conclude that brain activity fits better, for instance, a convolutional motif rather than a dense architecture. If this were true, it would mean that functional similarity over large data sets effectively constrains architecture. Then the need for a separate test of the model at the level of anatomy would become, at least in part, less critical for model validation. Therefore, we ask the question: could functional similarity be a reliable predictor of architectural similarity?

We describe an attempt to benchmark the most popular similarity techniques by replacing the brain recordings with

data generated by various known networks with drastically different architectural motifs, such as convolution vs. attention. Such a setting provides a valuable upper bound to the identifiability of anatomical differences.

1.1. System Identification from Leaderboards

When artificial models are compared against common biological benchmarks for predictivity (Yamins & DiCarlo, 2016), models with the top score are deemed better models for neuroscience. As improvements to scores are made over time, ideally, more relevant candidates emerge. Nevertheless, if two artificial models with distinctly different architectures, trained on the same data, happen to be similar in reproducing neural activities (target model), then it would be impossible to conclude what accounts for the similarity. It can be biologically relevant *motifs from each architecture*, the properties of the *stimulus input*, or *similarity metric*. Such ambiguity is due to the many-to-one mapping of a model onto a leaderboard score. Our work shows that multiple factors play a role in representational similarities.

An interesting example is offered by Chang et al. (2021), which compares many different models with respect to their ability to reproduce neural responses in the inferotemporal (IT) cortex to face images. The study concludes that the 2D morphable model is the best, even though the operations required in the specific model, such as correspondence and vectorization, do not have an apparent biological implementation in terms of neurons and synapses. This highlights that there are multiple confounds besides the biological constraints which affect neural predictivity.

2. Related Work

While the analogy between neural network models and the brain has been well validated (Bashivan et al., 2019), the extent of this correspondence across multiple levels (Marr & Poggio, 1976) has not been fully understood. Yet, previous works often make connections between representational and architectural levels of analyses, drawing implications about computational properties directly from functional correspondence. One previous study (Kar et al., 2019) argues for recurrent connections in the ventral stream, partly based on the higher neural predictivity of recurrent neural networks. Another study (Nonaka et al., 2021) observed that purely feedforward neural networks outperform those with branching or skip connections in terms of a similarity measure on neural activations. They concluded that the results indicate the significance of feedforward connections in the brain. Additionally, their interpretation suggests that certain brain areas may be responsible for spatial integration, as fully-connected layers, which have the ability to integrate visual features spatially, better match these brain areas.

This assumed correspondence could be attributed to methodological limitations of evaluating such models simultaneously across all levels. Jonas & Kording (2017) investigated the robustness of standard analysis techniques in neuroscience with a microprocessor as a ground-truth model to determine the boundaries of what conclusions could be drawn about a known system. The presumption of correspondence could also be attributed to underappreciated variability from model hyperparameters (Schaeffer et al., 2022). In a similar spirit to Jonas & Kording (2017); Lazebnik (2002), we evaluate system identification on a known ground-truth model to establish the boundaries of what architectural motifs can be reliably uncovered. We perform our analysis under favorable experimental conditions to establish an upper bound.

As modern neural network models have grown more prominent in unison with the corresponding resources to train these models, pre-trained reference models have become more widely available in research (Wightman, 2019). Consequently, the need to compare these references along different metrics has followed suit. Kornblith et al. (2019); Morcos et al. (2018) explored using different similarity measures between the layers of artificial neural network models. Kornblith et al. (2019) propose various properties a similarity measure should be invariant such as orthogonal transformations and isotropic scaling while not invariant to invertible linear transformations. Kornblith et al. (2019) found centered kernel alignment (CKA), a method very similar to Representation Similarity Analysis (Kriegeskorte et al., 2008), to best satisfy these requirements. Ding et al. (2021) explored the sensitivity of methods like canonical correlation analysis, CKA, and orthogonal procrustes distance to changes in factors that do not impact the functional behavior of neural network models.

3. Background and Methods

The two predominant approaches to evaluating computational models of the brain are using metrics based on linear encoding analysis for neural predictivity and population-level representation similarity. The first measures how well a model can predict the activations of individual units, whereas the second metric measures how correlated the variance of internal representations is. We study the following neural predictivity scores consistent with the typical approaches: Linear Regression and Centered Kernel Alignment (CKA).

In computational neuroscience, we usually have a neural system (brain) that we are interested in modeling. We call this network a *target* and the proposed candidate model a *source*. Formally, for a layer with p_x units in a source model, let $X \in \mathbb{R}^{n \times p_x}$ be the matrix of representations with p_x features over n stimulus images. Similarly, let $Y \in \mathbb{R}^{n \times p_y}$ be a matrix of representations with p_y features of the target

model (or layer) on the same n stimulus images. Unless otherwise noted, we subsample 3000 target units to test an analogous condition as in biology, where recordings are far from exhaustive (more discussion in Appendix A.4). Our analyses partially depend upon the target coverage, and we later examine the effect of increasing the number of target units.

3.1. Encoding Model: Linear Regression

Following the procedure developed by previous works (Schrumpf et al., 2020; Yamins et al., 2014; Conwell et al., 2021; Kar et al., 2019; Mitchell et al., 2008), we linearly project the feature space of a single layer in a source model to map onto a single unit in a target model (a column of Y). The linear regression score is the Pearson’s correlation $r(\cdot, \cdot)$ coefficient between the predicted responses of a source model and the ground-truth target responses to a set of stimulus images.

$$\hat{\beta} = \operatorname{argmin}_{\beta} \|Y - XS\beta\|_F^2 + \lambda \|\beta\|_F^2 \quad (1)$$

$$LR(X, Y) = r(XS\hat{\beta}, Y) \quad (2)$$

We first extract activations on the same set of stimulus images for source and target models. To reduce computational costs without sacrificing predictivity, we apply sparse random projection $S \in \mathbb{R}^{p_x \times q_x}$ for $q_x \ll p_x$, on the activations of the source model (Conwell et al., 2021). This projection reduces the dimensionality of the features to q_x while still preserving relative distances between points (Li et al., 2006). We determine the dimensionality of random projection based on the Johnson–Lindenstrauss lemma (Johnson, 1984) with the distortion factor $\epsilon = 0.1$, which leads to 6862 random projections for 3000 stimulus images. Unlike principal component analysis, sparse random projection is a dataset-independent dimensionality reduction method. This removes any data-dependent confounds from our processing pipeline for linear regression and isolates dataset dependence into our variables of interest: linear regression and candidate model.

We apply ridge regression on every layer of a source model to predict a target unit using these features. We use nested cross-validations in which the regularization parameter λ is chosen in the inner loop and a linear model is fitted in the outer loop. The list of tested λ is $[0.01, 0.1, 1.0, 10.0, 100]$. We use 4-fold cross-validation for inner loops and 5-fold cross-validation for outer loops. As there are multiple target units, the median of Pearson’s correlation coefficients between predicted and true responses is the aggregate score for layer-wise comparison between source and target models. Note that a layer of a target model is usually assumed to correspond to a visual area, e.g., V1 or IT, in the visual cor-

tex. For a layer-mapped model, we report maximum linear regression scores across source layers for target layers.

3.2. Centered Kernel Alignment

Another widely used type of metric builds upon the idea of measuring the representational similarity between the activations of two neural networks for each pair of images. While variants of this metric abound, including RSA or re-weighted RSA (Kriegeskorte et al., 2008; Khaligh-Razavi et al., 2017), we use CKA (Cortes et al., 2012) as Kornblith et al. (2019) showed strong correspondence between layers of models trained with different initializations, which we will further discuss as a validity test we perform. We consider linear CKA in this work:

$$\text{CKA}(X, Y) = \frac{\|Y^T X\|_F^2}{\|X^T X\|_F \|Y^T Y\|_F} \quad (3)$$

Kornblith et al. (2019) showed that the variance explained by (unregularized) linear regression accounts for the singular values of the source representation. In contrast, linear CKA depends on the singular values of both target and source representations. Recent work (Diedrichsen et al., 2020) notes that linear CKA is equivalent to a whitened representational dissimilarity matrix (RDM) in RSA under certain conditions. We also call CKA a neural predictivity score because a target network is observable, whereas a source network gives predicted responses.

3.3. Identifiability Index

To quantify how selective predictivity scores are when a source matches the target architecture compared to when the architecture differs between source and target networks, we define an identifiability index as:

$$\text{Identifiability Index} = \frac{\text{score}(s = t) - \overline{\text{score}}(s \neq t)}{\text{score}(s = t) + \overline{\text{score}}(s \neq t)} \quad (4)$$

where s is the source or candidate model, and t is the target model. In brief, it is a normalized difference between the score for the true positive and the mean score ($\overline{\text{score}}$) for the true negatives. Previous works (Dobs et al., 2022; Freiwald & Tsao, 2010) defined selectivity indices in the same way in similar contexts, such as the selectivity of a neuron to specific tasks.

3.4. Simulated Environment

If a target network is a brain, it is essentially a black box, making it challenging to understand the properties or limitations of the comparison metrics. Therefore, we instead

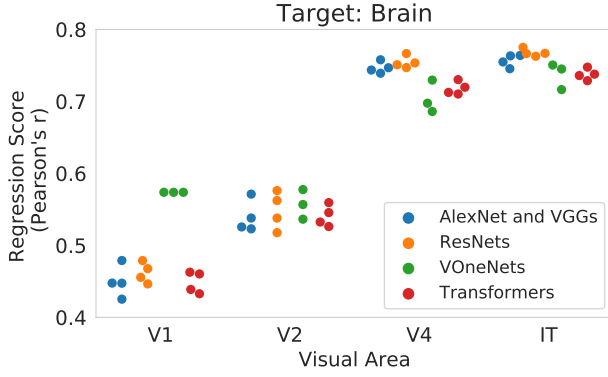


Figure 2. Linear regression scores of deep neural networks for brain activations in the macaque visual cortex. Architecture list in Appendix A.1. For V1, the top performing three models are in the VOneNet family (Dapello et al., 2020), which are explicitly designed to mimic the known properties of V1.

use artificial neural networks of our choice as targets for our experiments.

We investigate the reliability of a metric to compare models, mainly to discriminate the underlying computations specified by the model’s architecture. We intentionally create favorable conditions for identifiability in a simulated environment where the ground truth model is a candidate among the source models. Taking these ideal conditions further, our target and source models are deterministic and do not include adverse conditions typically encountered in biological recordings, such as noise and temporal processing. We consider the following architectures:

Convolutional: AlexNet (Krizhevsky et al., 2012), VGG11 (Simonyan & Zisserman, 2014), ResNet18 (He et al., 2016)

Recurrent: CORnet-S (Kubilius et al., 2019)

Transformer: ViT-B/32 (Dosovitskiy et al., 2020)

Mixer: MLP-Mixer-B/16 (Tolstikhin et al., 2021)

These architectures are emblematic of the vision-based models used today. Each architecture has a distinct motif, making it unique from other models. For example, transformer networks use the soft-attention operation as a core motif, whereas convolutional networks use convolution. Recurrent networks implement feedback connections which may be critical in the visual cortex for object recognition (Kubilius et al., 2019; Kar et al., 2019). Moreover, mixer networks (Tolstikhin et al., 2021; Touvron et al., 2021; Melas-Kyriazi, 2021) uniquely perform fully-connected operations over image patches, alternating between the feature and patch dimensions.

4. Results

4.1. Different Models with Equivalent Neural Predictivity

We compare various artificial neural networks with publicly shared neural recordings in primates (Majaj et al., 2015; Freeman et al., 2013) via the Brain-Score framework (Schrimpf et al., 2020). Our experiments show that the differences between markedly different neural network architectures are minimal after training (Figure 2), consistent with the previous work (Schrimpf et al., 2020; Kubilius et al., 2019; Conwell et al., 2021). Previous works focused on the relative ranking of models and investigated which model yields the highest score. However, if we take a closer look at the result, the performance difference is minimal, with the range of scores having a standard deviation < 0.03 (for V2=0.021, V4=0.023, IT=0.016) except for V1. For V1, VOneNets (Dapello et al., 2020), which explicitly build in properties observed from experimental works in neuroscience, significantly outperform other models. Notably, the models we consider have quite different architectures based on combinations of various components, such as convolutional layers, attention layers, and skip connections. This suggests that architectures with different computational operations reach almost equivalent performance after training on the same large-scale dataset, i.e., ImageNet.

4.2. Identification of Architectures in an Ideal Setting

One potential interpretation of the above result would be that different neural network architectures are equally good (or bad) models of the visual cortex. An alternative explanation would be that the method we use to compare models with the brain has limitations in identifying precise computational operations. To test the hypothesis, we focus on where underlying target neural networks are known instead of being a black box, as with biological brains. Specifically, we replace the target neural recordings with artificial neural network activations. By examining whether the candidate source model with the highest predictivity is identical to the target model, we can evaluate to what extent we can identify architectures with the current approach.

4.2.1. LINEAR REGRESSION

We first compare various source models (AlexNet, VGG11, ResNet18, CORnet-S, ViT-B/32, MLP-Mixer-B/16) with a target network, the same architecture as one of the source models and is trained on the same dataset but initialized with a different seed. We test a dataset composed of 3200 images of synthetic objects studied in (Majaj et al., 2015) to be consistent with the evaluation pipeline of Brain-Score. The ground-truth source model will yield a higher score than other models if the model comparison pipeline is reliable.

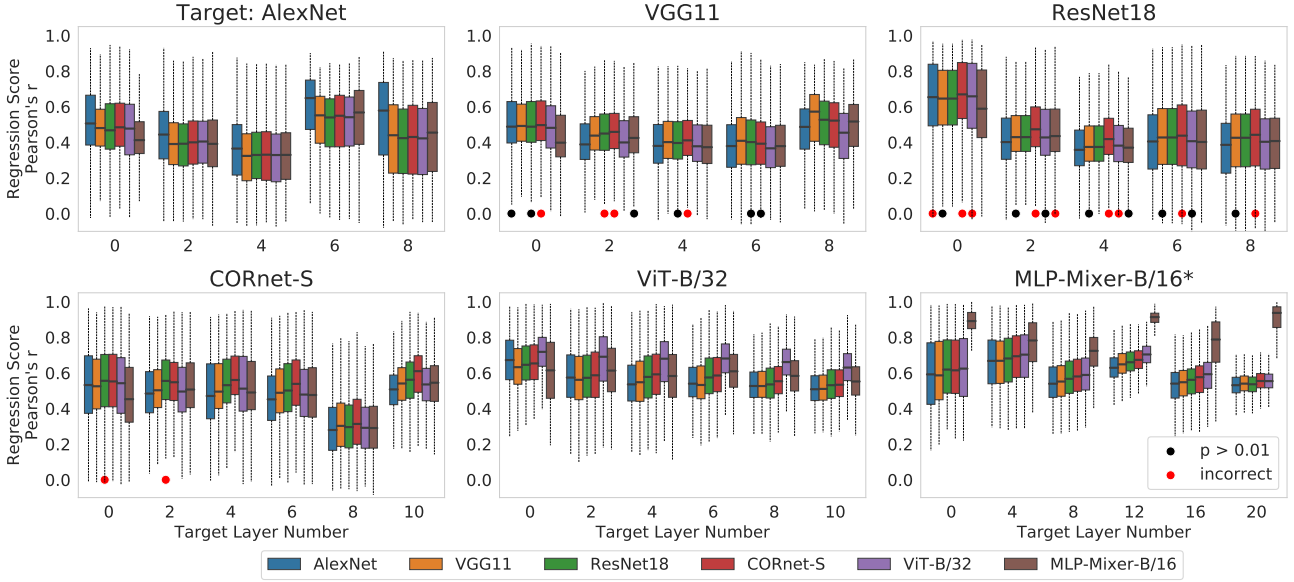


Figure 3. Linear regression scores as a function of target network layer. Each unit in a target layer yields an independent score. Boxplots show the distribution of scores for subsampled 3000 units / target layer. The statistics displayed by boxplots are: min, max, median, and the first and third quartiles. When the architecture of source and target networks matches, different initialization seeds are used for source networks, except for one (MLP-Mixer-B/16*), which uses identical weights. Black dots ● indicate that the correct architecture does not outperform others with statistically significant differences ($p > 0.01$). Red dots ● indicate that the median scores for those networks are higher than the median for the correct architecture. The ranking of source models is typically based on median scores (Schrimpf et al., 2020).

For most target layers, source networks with the highest median score are the correct network (Figure 3). However, strikingly, for several layers in VGG11, ResNet18, and CORnet-S the best-matched layers belong to a source model that is not the correct architecture. In other words, given the activations of ResNet18, for instance, and based on linear regression scores, we would make an incorrect prediction that the system’s underlying architecture is closest to a recurrent neural network CORnet-S. While it has been argued that recurrent neural networks essentially correspond to ResNets with weight shared across layers (Liao & Poggio, 2016), the ResNet18 architecture we consider here does not have such weight-sharing. The confusion between ResNet18 and CORnet-S is especially noteworthy in that the prediction leads to an incorrect inference about the presence of recurrent connections in the target network.

In addition, because of our ideal setting, where an identical network is one of the source models, we expect to see a significant difference between matching and non-matching models. However, for multiple target layers, linear regression scores for the non-identical architectures, when compared with those for the identical one, do not show a significant decrease in predictivity based on Welch’s t-test with $p < 0.01$ applied as a threshold (Figure 3). This result suggests that the identification of the underlying architectures

of unknown neural systems is far from perfect.

4.2.2. CENTERED KERNEL ALIGNMENT

Next, we examine another widely used metric, CKA, for comparing representations. Again, we compare different source models to a target model, also an artificial neural network. For the target models we tested, the ground-truth source models achieve the highest score (Figure 4). Still, some unmatched source networks lead to scores close to the matched networks, even for the target MLP-Mixer-B/16, where the source network of the same architecture type also has identical weights.

When applying CKA to compare representations, we subsample a set (3000) of target units to mimic the limited coverage of single-unit recordings. Assuming we can increase the coverage for future experiments with more advanced measurement techniques, we test whether the identifiability improves if we include all target units. Additionally, methods similar to CKA, such as RSA, are often applied to interpret neural recordings, including fMRI, MEG, and EEG (Cichy & Oliva, 2020; Cichy & Pantazis, 2017), which can have full coverage of the entire brain or a specific region of interest. Therefore, we simulate such analyses by having all units in the targets. Overall, the ground-truth source models

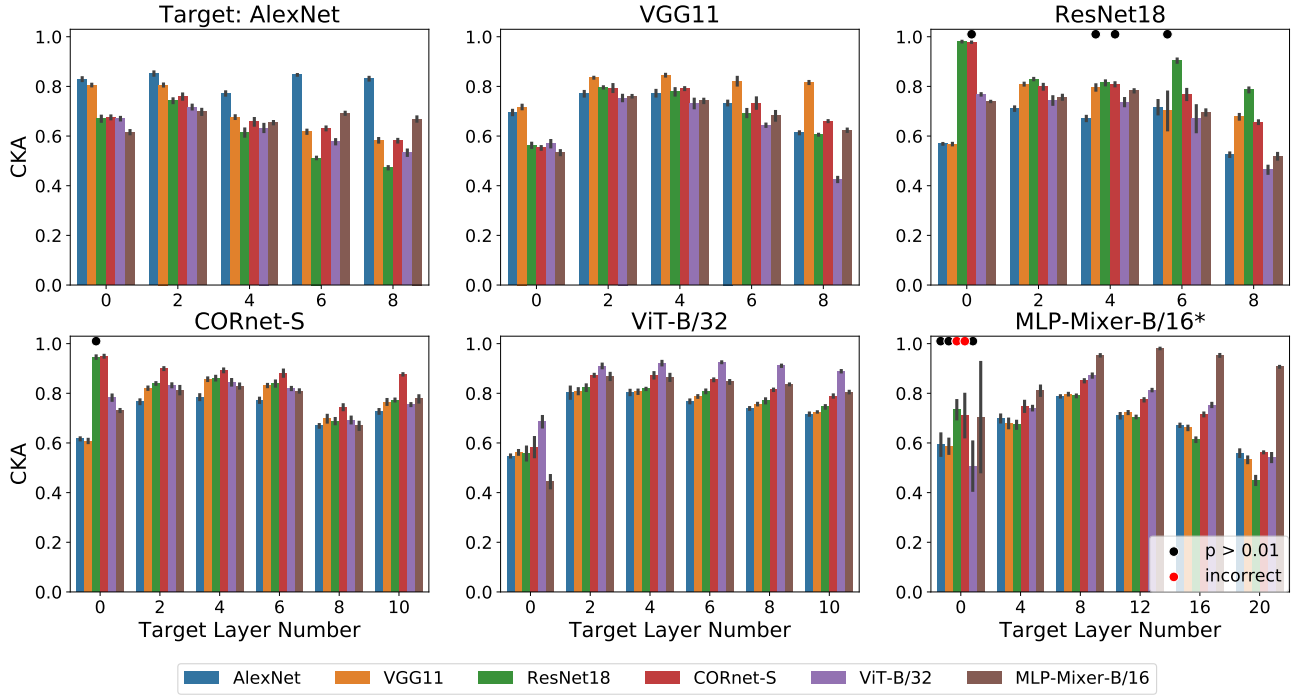


Figure 4. CKA for different source and target networks. The number of subsampled units in a target layer is 3000. Bar plots display the mean and error bars indicate the standard deviation for subsampling a set of target units multiple times ($N=4$). The experimental setup is identical as in Figure 3 except that CKA is used as the metric instead of linear regression. Note that, unlike linear regression scores, CKA yields an aggregated single score for each target layer.

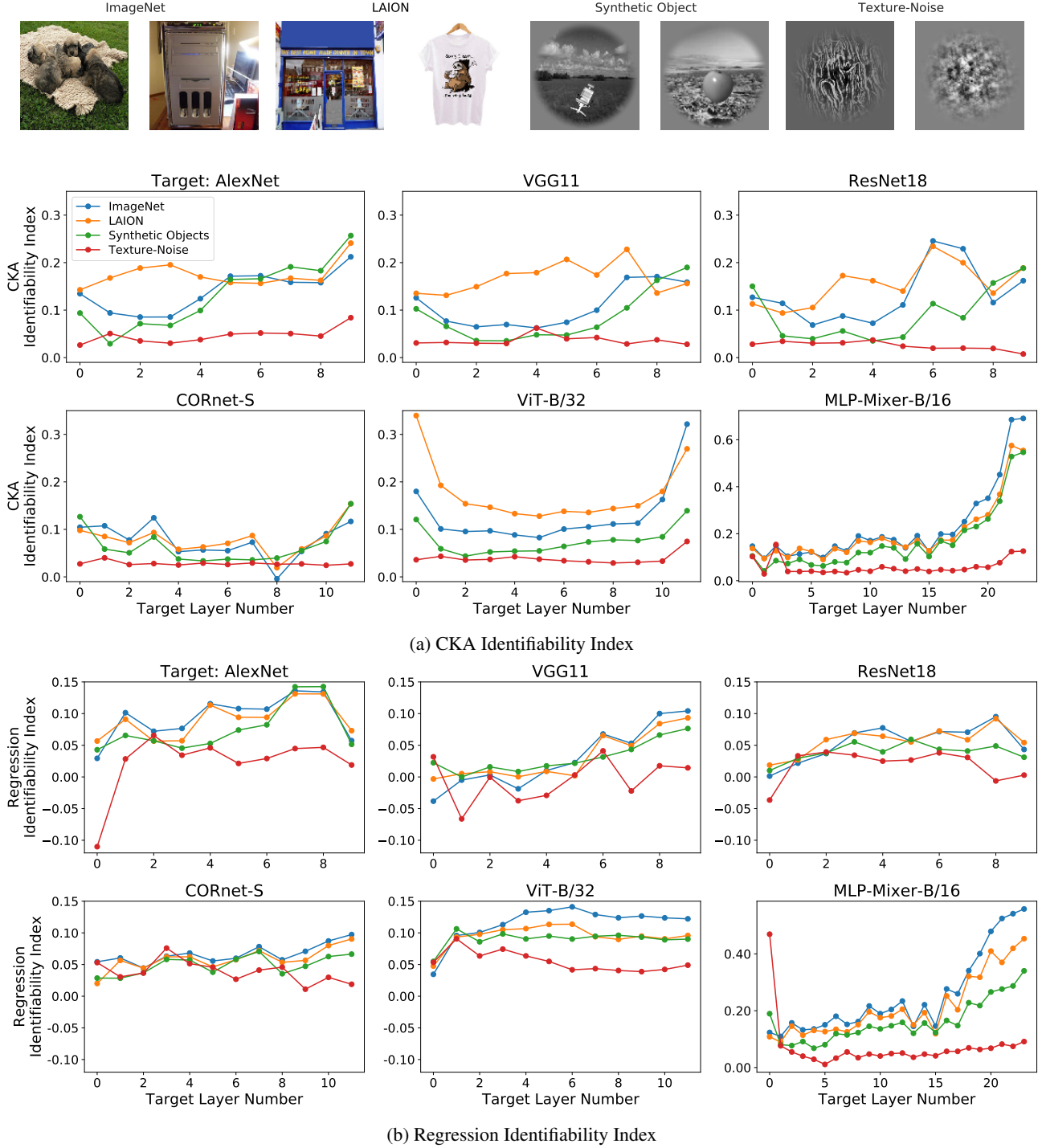


Figure 5. **Top** Sample images of each stimulus image type. **(a)** Identifiability index using CKA and **(b)** linear regression for various types of stimulus images and target networks.

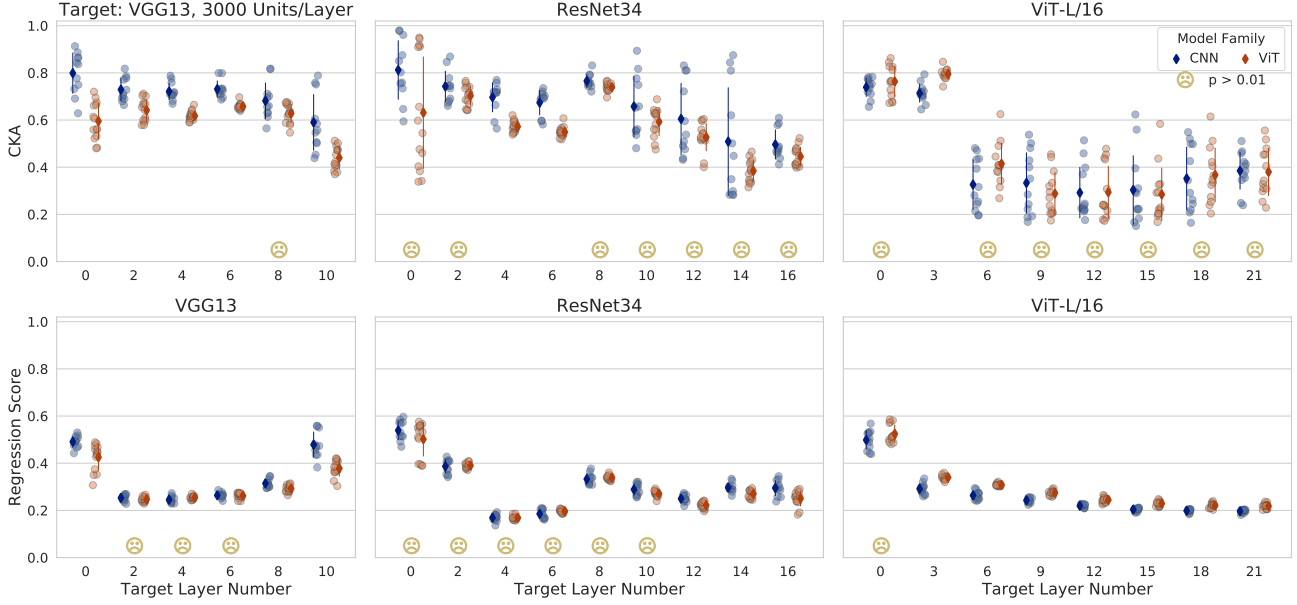


Figure 6. Different architectural variants (12 CNNs and 14 ViTs) are compared with two CNNs and one ViT target network. Each data point represents the maximum score of an architecture for the corresponding target layers. Solid markers indicate the mean score of the corresponding model class, and error bars display standard deviations. ☺ indicates that the corresponding layer does not show a statistically significant difference between model classes.

outperform the other source models by a significant margin (Figure 7 in the Appendix). This suggests system identification can be more reliable with more recording units and complete target coverage.

4.3. Effects of the Stimulus Distribution on Identifiability

A potentially significant variable overlooked in comparing computational models of the brain is the type of stimulus images. What types of stimulus images are suited for evaluating competing models? In Brain-Score, stimulus images for comparing models of the high-level visual areas, V4 and IT, are images of synthetic objects (Majaj et al., 2015). In contrast, those for the lower visual areas, V1 and V2, are images of texture and noise (Freeman et al., 2013). To examine the effect of using different stimulus images, we test images of synthetic objects (3200 images) and texture and noise (135 images), and additionally, ImageNet (3000 images) and web-crawled images LAION (3000 images) (Schuhmann et al., 2021) which are more natural images than the first two datasets.

In Figure 5, we analyze Identifiability Index for different stimulus images. More realistic stimulus images, i.e., synthetic objects, ImageNet, and LAION, show higher identifiability than texture and noise images for all target models. We observe identifiability increases with layer depth. Notably, even for early layers in target models, which would

correspond to V1 and V2 in the visual cortex, texture and noise images fail to give higher identifiability.

Additionally, testing LAION images leads to Identifiability Indices as high as, or often higher than, those obtained using ImageNet images, on which all target systems are trained. This result indicates that the stimuli do not need to be exactly in-distribution with the training data for better identification and suggests that natural image statistics may be sufficient.

It is important to note that the images of texture and noise used in the experiment help characterize certain aspects of V1 and V2, as shown in previous work (Freeman et al., 2013). Specifically, the original study investigated the functional role of V2 in comparison to V1 and demonstrated that naturalistic structure modulates V2. While the image set plays an influential role as a variable in a carefully designed experiment to address specific hypotheses, it does not serve as a sufficient test set for any general hypothesis, such as evaluating different neural networks.

4.4. Challenges of Identifying Key Architectural Motifs

Interesting hypotheses for a more biologically plausible design principle of brain-like models often involve key high-level architectural motifs. For instance, potential questions are whether recurrent connections are crucial in visual processing or, with the recent success of transformer models in deep learning, whether the brain similarly implements computations like attention layers in transformers. The de-

tails beyond the key motif, such as the number of layers or the exact type of activation functions, may vary and be underdetermined within the scope of such research questions. Likewise, it is unlikely that candidate models proposed by scientists align with the brain at every level, from low-level specifics to high-level computation. Therefore, an ideal methodology for comparing models should help separate the key properties of interest while being invariant to other confounds.

Considering it is a timely question, with the increased interest in transformers as models of the brain in different domains (Schrimpf et al., 2021; Berrios & Deza, 2022; Whittington et al., 2021), we focus on the problem of identifying convolution vs. attention. We test 12 Convolutional Networks and 14 Vision Transformers of different architectures (list in Appendix A.2), and to maximize identifiability, we use ImageNet stimulus images. Note that an identical architecture with the target network is not included as a source network.

Figure 6 shows that for both CKA and regression, there is high inter-class variance for many target layers. For CKA, one layer in VGG13, 7 layers in ResNet34, and 7 layers in ViT-L/16, and for regression, three layers in VGG13, 6 layers in ResNet34, and one layer in ViT-L/16 do not show a statistically significant difference between the two model classes based on Welch’s t-test with $p < 0.01$ used as a threshold. The significant variance among source models suggests that model class identification can be incorrect depending on the precise variation we choose, especially if we rely on a limited set of models.

5. Discussion

Under idealized settings, we tested the identifiability of various artificial neural networks with differing architectures. We present two contrasting interpretations of model identifiability based on our results, one optimistic (Glass half full) and one pessimistic (Glass half empty).

Glass half full: Despite the many factors that can lead to variable scores, linear regression and CKA give reasonable identification capability under unrealistically ideal conditions. Across all the architectures tested, identifiability improves as a function of depth.

Glass half empty: However, system identification is highly variable and dependent on the properties of the target architecture and the stimulus data used to probe the candidate models. For architecture-wide motifs, like convolution vs. attention, scores overlap significantly across almost all layers. This indicates that such distinct motifs do not play a significant role in the score.

Our results suggest two future directions for improving

system identification with current approaches: 1) Using stimuli images that are more natural. 2) With more neurons recorded in the brain, neural predictivity scores can be more reliable in finding the underlying architecture.

On the other hand, it is worthwhile to note that we may have reached close to the ceiling using neural predictivity scores for system identification. As an example, when our source network is AlexNet, its regression scores against the brain (Figure 2) are on par with, or slightly higher than, the scores against another AlexNet (Figure 3). In other words, based on the current methods, AlexNet predicts the brain as well as, if not better than, predicting itself. This observation is not limited to AlexNet but applies to other target networks. This fundamental limitation of present evaluation techniques, such as the linear encoding analysis used in isolation, emphasizes the need to develop new approaches beyond comparing functional similarities.

As we argued earlier, ranking models in terms of their agreement with neural recordings is the first step in verifying or falsifying a neuroscience model. Since several different models are very close in ranking, the next step – architectural validation – is the key. Furthermore, it may have to be done independently of functional validation and with little guidance from it, using standard experimental tools in neuroscience. A parallel direction is, however, to try to develop specially designed, critical stimuli to distinguish between different architectures instead of measuring the overall fit to data. As a simple example, it may be possible to discriminate between dense and local (e.g., CNN) network architectures by measuring the presence or absence of interactions between parts of a visual stimulus that are spatially separated.

Acknowledgements

This material is based upon work supported by the Center for Brains, Minds and Machines (CBMM), funded by NSF STC award CCF-1231216. Y. Han is a recipient of the Lore Harp McGovern Fellowship.

References

- Bashivan, P., Kar, K., and DiCarlo, J. J. Neural population control via deep image synthesis. *Science*, 364(6439): eaav9436, 2019.
- Berrios, W. and Deza, A. Joint rotational invariance and adversarial training of a dual-stream transformer yields state of the art brain-score for area v4. *arXiv preprint arXiv:2203.06649*, 2022.
- Chang, L., Egger, B., Vetter, T., and Tsao, D. Y. Explaining face representation in the primate brain using different computational models. *Current Bi-*

- ology, 31(13):2785–2795.e4, 2021. ISSN 0960-9822. doi: <https://doi.org/10.1016/j.cub.2021.04.014>. URL <https://www.sciencedirect.com/science/article/pii/S0960982221005273>.
- Cichy, R. M. and Oliva, A. Am/eeg-fmri fusion primer: resolving human brain responses in space and time. *Neuron*, 107(5):772–781, 2020.
- Cichy, R. M. and Pantazis, D. Multivariate pattern analysis of meg and eeg: A comparison of representational structure in time and space. *NeuroImage*, 158:441–454, 2017.
- Cichy, R. M., Dwivedi, K., Lahner, B., Lascelles, A., Iamshchinina, P., Graumann, M., Andonian, A., Murty, N., Kay, K., Roig, G., et al. The algonauts project 2021 challenge: How the human brain makes sense of a world in motion. *arXiv preprint arXiv:2104.13714*, 2021.
- Collins, C. E., Airey, D. C., Young, N. A., Leitch, D. B., and Kaas, J. H. Neuron densities vary across and within cortical areas in primates. *Proceedings of the National Academy of Sciences*, 107(36):15927–15932, 2010.
- Conwell, C., Mayo, D., Barbu, A., Buice, M., Alvarez, G., and Katz, B. Neural regression, representational similarity, model zoology & neural taskonomy at scale in rodent visual cortex. *Advances in Neural Information Processing Systems*, 34, 2021.
- Cortes, C., Mohri, M., and Rostamizadeh, A. Algorithms for learning kernels based on centered alignment. *The Journal of Machine Learning Research*, 13:795–828, 2012.
- Dapello, J., Marques, T., Schrimpf, M., Geiger, F., Cox, D., and DiCarlo, J. J. Simulating a primary visual cortex at the front of cnns improves robustness to image perturbations. *Advances in Neural Information Processing Systems*, 33:13073–13087, 2020.
- Diedrichsen, J., Berlot, E., Mur, M., Schütt, H. H., Shahbazi, M., and Kriegeskorte, N. Comparing representational geometries using whitened unbiased-distance-matrix similarity. *arXiv preprint arXiv:2007.02789*, 2020.
- Ding, F., Denain, J.-S., and Steinhart, J. Grounding representation similarity with statistical testing. *arXiv preprint arXiv:2108.01661*, 2021.
- Dobs, K., Martinez, J., Kell, A. J., and Kanwisher, N. Brain-like functional specialization emerges spontaneously in deep neural networks. *Science advances*, 8(11):eabl8913, 2022.
- Dosovitskiy, A., Beyer, L., Kolesnikov, A., Weissenborn, D., Zhai, X., Unterthiner, T., Dehghani, M., Minderer, M., Heigold, G., Gelly, S., et al. An image is worth 16x16 words: Transformers for image recognition at scale. *arXiv preprint arXiv:2010.11929*, 2020.
- Freeman, J., Ziemba, C. M., Heeger, D. J., Simoncelli, E. P., and Movshon, J. A. A functional and perceptual signature of the second visual area in primates. *Nature neuroscience*, 16(7):974–981, 2013.
- Freiwald, W. A. and Tsao, D. Y. Functional compartmentalization and viewpoint generalization within the macaque face-processing system. *Science*, 330(6005):845–851, 2010.
- He, K., Zhang, X., Ren, S., and Sun, J. Deep residual learning for image recognition. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pp. 770–778, 2016.
- Johnson, W. B. Extensions of lipschitz mappings into a hilbert space. *Contemp. Math.*, 26:189–206, 1984.
- Jonas, E. and Kording, K. P. Could a neuroscientist understand a microprocessor? *PLoS computational biology*, 13(1):e1005268, 2017.
- Kar, K., Kubilius, J., Schmidt, K., Issa, E. B., and DiCarlo, J. J. Evidence that recurrent circuits are critical to the ventral stream’s execution of core object recognition behavior. *Nature neuroscience*, 22(6):974–983, 2019.
- Khaligh-Razavi, S.-M., Henriksson, L., Kay, K., and Kriegeskorte, N. Fixed versus mixed rsa: Explaining visual representations by fixed and mixed feature sets from shallow and deep computational models. *Journal of Mathematical Psychology*, 76:184–197, 2017.
- Kornblith, S., Norouzi, M., Lee, H., and Hinton, G. Similarity of neural network representations revisited. In *International Conference on Machine Learning*, pp. 3519–3529. PMLR, 2019.
- Kriegeskorte, N., Mur, M., and Bandettini, P. A. Representational similarity analysis-connecting the branches of systems neuroscience. *Frontiers in systems neuroscience*, 2:4, 2008.
- Krizhevsky, A., Sutskever, I., and Hinton, G. E. Imagenet classification with deep convolutional neural networks. *Advances in neural information processing systems*, 25, 2012.
- Kubilius, J., Schrimpf, M., Kar, K., Rajalingham, R., Hong, H., Majaj, N., Issa, E., Bashivan, P., Prescott-Roy, J., Schmidt, K., et al. Brain-like object recognition with high-performing shallow recurrent anns. *Advances in neural information processing systems*, 32, 2019.

- Lazebnik, Y. Can a biologist fix a radio?—or, what i learned while studying apoptosis. *Cancer cell*, 2(3):179–182, 2002.
- Li, P., Hastie, T. J., and Church, K. W. Very sparse random projections. In *Proceedings of the 12th ACM SIGKDD international conference on Knowledge discovery and data mining*, pp. 287–296, 2006.
- Liao, Q. and Poggio, T. Bridging the gaps between residual learning, recurrent neural networks and visual cortex. *arXiv preprint arXiv:1604.03640*, 2016.
- Majaj, N. J., Hong, H., Solomon, E. A., and DiCarlo, J. J. Simple learned weighted sums of inferior temporal neuronal firing rates accurately predict human core object recognition performance. *Journal of Neuroscience*, 35 (39):13402–13418, 2015.
- Marr, D. and Poggio, T. From understanding computation to understanding neural circuitry. 1976.
- Melas-Kyriazi, L. Do you even need attention? a stack of feed-forward layers does surprisingly well on imagenet. *arXiv preprint arXiv:2105.02723*, 2021.
- Mitchell, T. M., Shinkareva, S. V., Carlson, A., Chang, K.-M., Malave, V. L., Mason, R. A., and Just, M. A. Predicting human brain activity associated with the meanings of nouns. *science*, 320(5880):1191–1195, 2008.
- Morcos, A., Raghu, M., and Bengio, S. Insights on representational similarity in neural networks with canonical correlation. *Advances in Neural Information Processing Systems*, 31, 2018.
- Nonaka, S., Majima, K., Aoki, S. C., and Kamitani, Y. Brain hierarchy score: Which deep neural networks are hierarchically brain-like? *IScience*, 24(9):103013, 2021.
- Schaeffer, R., Khona, M., and Fiete, I. No free lunch from deep learning in neuroscience: A case study through models of the entorhinal-hippocampal circuit. *bioRxiv*, 2022.
- Schrimpf, M., Kubilius, J., Hong, H., Majaj, N. J., Rajalingham, R., Issa, E. B., Kar, K., Bashivan, P., Prescott-Roy, J., Geiger, F., et al. Brain-score: Which artificial neural network for object recognition is most brain-like? *BioRxiv*, pp. 407007, 2020.
- Schrimpf, M., Blank, I. A., Tuckute, G., Kauf, C., Hosseini, E. A., Kanwisher, N., Tenenbaum, J. B., and Fedorenko, E. The neural architecture of language: Integrative modeling converges on predictive processing. *Proceedings of the National Academy of Sciences*, 118(45), 2021.
- Schuhmann, C., Vencu, R., Beaumont, R., Kaczmarczyk, R., Mullis, C., Katta, A., Coombes, T., Jitsev, J., and Komatsuzaki, A. Laion-400m: Open dataset of clip-filtered 400 million image-text pairs. *arXiv preprint arXiv:2111.02114*, 2021.
- Simonyan, K. and Zisserman, A. Very deep convolutional networks for large-scale image recognition. *arXiv preprint arXiv:1409.1556*, 2014.
- Tolstikhin, I. O., Houlsby, N., Kolesnikov, A., Beyer, L., Zhai, X., Unterthiner, T., Yung, J., Steiner, A., Keysers, D., Uszkoreit, J., et al. Mlp-mixer: An all-mlp architecture for vision. *Advances in Neural Information Processing Systems*, 34, 2021.
- Touvron, H., Bojanowski, P., Caron, M., Cord, M., El-Nouby, A., Grave, E., Izacard, G., Joulin, A., Synnaeve, G., Verbeek, J., et al. Resmlp: Feedforward networks for image classification with data-efficient training. *arXiv preprint arXiv:2105.03404*, 2021.
- Whittington, J. C., Warren, J., and Behrens, T. E. Relating transformers to models and neural representations of the hippocampal formation. *arXiv preprint arXiv:2112.04035*, 2021.
- Wightman, R. Pytorch image models. <https://github.com/rwightman/pytorch-image-models>, 2019.
- Yamins, D. L. and DiCarlo, J. J. Using goal-driven deep learning models to understand sensory cortex. *Nature neuroscience*, 19(3):356–365, 2016.
- Yamins, D. L., Hong, H., Cadieu, C. F., Solomon, E. A., Seibert, D., and DiCarlo, J. J. Performance-optimized hierarchical models predict neural responses in higher visual cortex. *Proceedings of the national academy of sciences*, 111(23):8619–8624, 2014.
- Yuan, L., Chen, Y., Wang, T., Yu, W., Shi, Y., Jiang, Z.-H., Tay, F. E., Feng, J., and Yan, S. Tokens-to-token vit: Training vision transformers from scratch on imagenet. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*, pp. 558–567, 2021.

A. Appendix

A.1. Model details for Section 4.1: Brain-Score

Below is the full list of models tested on the benchmarks of Brain-Score as reported in Section 4.1. In addition to testing vision models pre-trained on ImageNet available from PyTorch’s torchvision model package version 0.12, we test VOneNets that are pre-trained on ImageNet and made publicly available by the authors (Dapello et al., 2020). VOneNets are also a family of CNNs.

Convolutional Networks: AlexNet, VGG11, VGG13, VGG19, ResNet18, ResNet34, ResNet50, ResNet101, VOneAlexNet, VOneResNet50, VOneCORnet-S

Transformer Networks: ViT-B/16, ViT-B/32, ViT-L/16, ViT-L/32

A.2. Model details for Section 4.5: finding the key architectural motif

For each target network reported in Section 4.5, namely VGG13, ResNet34, and ViT-L/16, below is the full list of source models tested to compare two model classes, CNN and transformer. For Tokens-to-token ViTs (T2T) (Yuan et al., 2021), we use models released by the authors. We use Twins Vision Transformers and a Visformer from timm library (Wightman, 2019). All other models are available from PyTorch’s torchvision model package version 0.12. All models are pre-trained on ImageNet.

Convolutional Networks: AlexNet, VGG11, VGG13, VGG16, VGG13_bn, ResNet18, ResNet34, ResNet50, Wide-ResNet50_2, SqueezeNet1.0, Densenet121, MobileNet_v2

Transformer Networks: ViT-B/16, ViT-B/32, ViT-L/16, ViT-L/32, T2T-ViT_t-14, T2T-ViT_t-19, T2T-ViT-7, T2T-ViT-10, Swin-B, Swin-S, Swin-T, Twins-PCPVT-Small, Twins-SVT-Small, Visformer-Small

A.3. Model details: number of layers included for each model

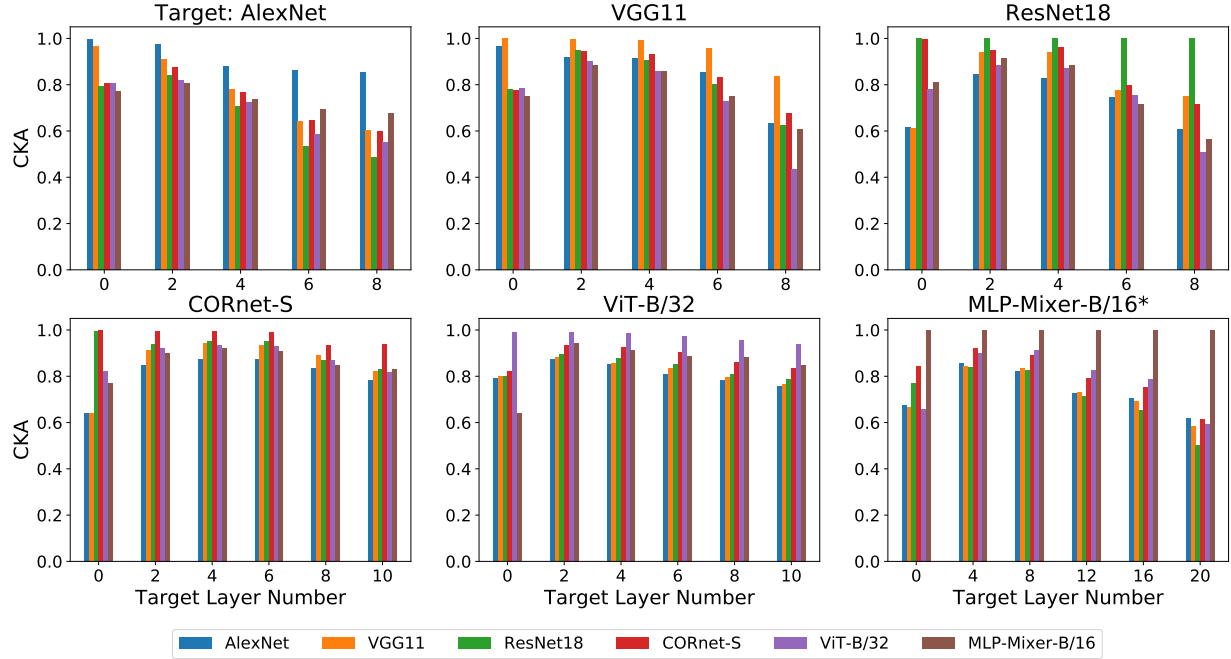
Table 1.

Model	Number of Layers
AlexNet	10
CORnet-S	12
Densenet121	30
MLP-Mixer-B16-224	24
Mobilenet_v2	14
ResNet18	10
ResNet34	18
ResNet50	18
Squeezenet1.0	13
Swin-B	24
Swin-S	24
Swin-T	12
T2T-ViT-10	13
T2T-ViT-7	10
T2T-ViT_t-14	17
T2T-ViT_t-19	22
Twins-PCPVT-Small	16
Twins-SVT-Small	18
VGG11	10
VGG13	12
VGG13-BN	12
VGG16	15
Visformer-Small	16
ViT-B-16	12
ViT-B-32	12
ViT-L-16	24
ViT-L-32	24
Wide-ResNet502	18

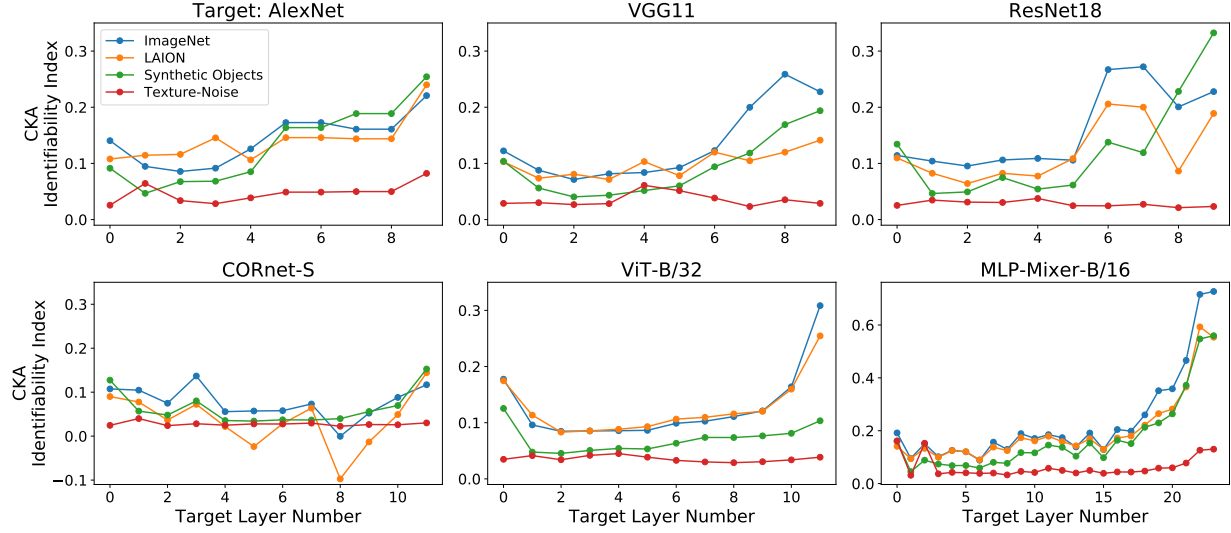
A.4. Coverage of neural recordings in target systems

For single-unit neural recordings, we have a limited number of recording sites. In our experimental settings, this is analogous to subsampling units in a target system. [Collins et al. \(2010\)](#) estimates that there are approximately 35 million neurons in V1 in the primate visual cortex, with a decreasing number of neurons in higher visual areas, such as MT, which has 1.6 million neurons. When different visual areas are combined, single-neuron recording sites generally amount to a few hundred. To maximize coverage, assuming approximately 100 sites for each of V1, V2, V4, and IT combined in one visual area, we can approximate the number of neural recording sites as 400 neurons for a back-of-the-envelope calculation. If we apply a similar coverage ratio from the brain to neural network models, the number of subsampled units would range from 1 to 200 for a single layer, considering that the number of units in models roughly ranges from 800K (early layers of ResNet) to 4096 (later layer of AlexNet). This range is smaller than our experimental condition of subsampling 3000 neurons per target layer.

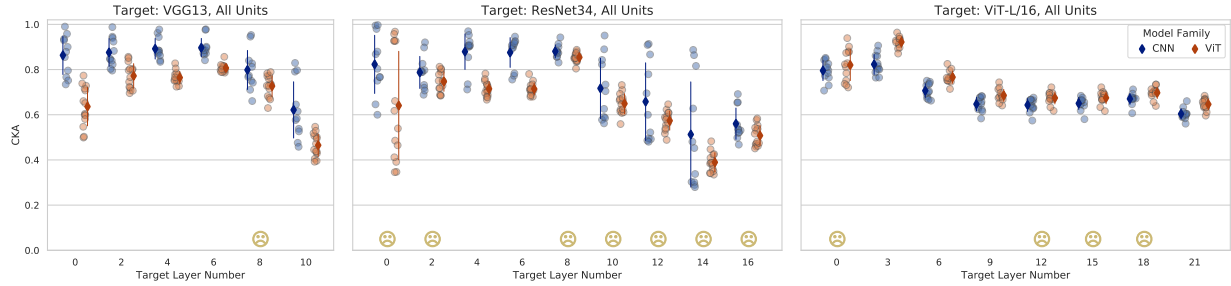
While increasing the number of recording sites for better system identification may not be easily feasible currently, we cannot rule out the possibility that recording techniques will be substantially improved in the future. Furthermore, applying CKA as a comparison metric may be more analogous to using RSA on neuroimaging data, such as fMRI or MEG, which captures signals from the whole brain. Therefore, we consider the condition of measuring all target units (Figure 7).



(a) CKA scores for top-performing layer in each source network



(b) Identifiability index for different types of stimuli images



(c) Comparing CKA scores for the model families CNN and ViT

Figure 7. CKA scores when all target units are tested. Experimental setups are identical to Figures 4, 5, and 6 otherwise. See Section 4.2.2 for discussion. For the target MLP-Mixer-B/16 in (a), the source MLP-Mixer-B/16 has identical weights with the target; thus CKA is trivially 1.