

Human EEG and artificial neural networks reveal disentangled representations and processing timelines of object real-world size and depth in natural images

Zitong Lu , Julie D Golomb

Department of Psychology, The Ohio State University, Columbus, United States

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

 Copyright information

eLife Assessment

This **important** study combines EEG, neural networks and multivariate pattern analysis to show that real-world size, retinal size and real-world depth are represented at different latencies. The evidence presented is **convincing** and the work will be of broader interest to the experimental and computational vision community.

<https://doi.org/10.7554/eLife.98117.2.sa3>

Abstract

Remarkably, human brains have the ability to accurately perceive and process the real-world size of objects, despite vast differences in distance and perspective. While previous studies have delved into this phenomenon, distinguishing the processing of real-world size from other visual properties, like depth, has been challenging. Using the THINGS EEG2 dataset with human EEG recordings and more ecologically valid naturalistic stimuli, our study combines human EEG and representational similarity analysis to disentangle neural representations of object real-world size from retinal size and perceived depth, leveraging recent datasets and modeling approaches to address challenges not fully resolved in previous work. We report a representational timeline of visual object processing: object real-world depth processed first, then retinal size, and finally, real-world size. Additionally, we input both these naturalistic images and object-only images without natural background into artificial neural networks. Consistent with the human EEG findings, we also successfully disentangled representation of object real-world size from retinal size and real-world depth in all three types of artificial neural networks (visual-only ResNet, visual-language CLIP, and language-only Word2Vec). Moreover, our multi-modal representational comparison framework across human EEG and artificial neural networks reveals real-world size as a stable and higher-level dimension in object space incorporating both visual and semantic information. Our research provides a temporally resolved characterization of how certain key object properties – such as object

real-world size, depth, and retinal size – are represented in the brain, which offers further advances and insights into our understanding of object space and the construction of more brain-like visual models.

Introduction

Imagine you are viewing an apple tree while walking around an orchard: as you change your perspective and distance, the retinal size of the apple you plan to pick varies, but you still perceive the apple as having a constant real-world size. How do our brains extract object real-world size information during object recognition to allow us to understand the complex world? Behavioral studies have demonstrated that perceived real-world size is represented as an object physical property, revealing same-size priming effects (Setti et al., 2008 [DOI](#)), familiar-size stroop effects (Konkle & Oliva, 2012a [DOI](#); Long & Konkle, 2017 [DOI](#)), and canonical visual size effects (Chen et al., 2022 [DOI](#); Konkle & Oliva, 2011 [DOI](#)). Human neuroimaging studies have also found evidence of object real-world size representation (Huang et al., 2022 [DOI](#); S.-M. Khaligh-Razavi et al., 2018 [DOI](#); Konkle & Caramazza, 2013 [DOI](#); Konkle & Oliva, 2012b [DOI](#); Luo et al., 2023 [DOI](#); Quek et al., 2023 [DOI](#); R. Wang et al., 2022 [DOI](#)). These findings suggest real-world size is a fundamental dimension of object representation.

However, previous studies on object real-world size have faced several challenges. Firstly, the perception of an object's real-world size is closely related to the perception of its real-world distance in depth. For instance, imagine you are looking at photos of an apple and a basketball: if the two photos were zoomed in such that the apple and the basketball filled the same exact retinal (image) size, you could still easily perceive that the apple is the physically smaller real-world object. But, you would simultaneously infer that the apple is thus located closer to you (or the camera) than the basketball. In previous neuroimaging studies of perceived real-world size (Huang et al., 2022 [DOI](#); Konkle & Caramazza, 2013 [DOI](#); Konkle & Oliva, 2012b [DOI](#)), researchers presented images of familiar objects zoomed and cropped such that they occupied the same retinal size, finding that neural responses in ventral temporal cortex reflected the perceived real-world size (e.g. an apple smaller than a car). However, while they controlled the retinal size of objects, the intrinsic correlation between real-world size and real-world depth in these images meant that the influence of perceived real-world depth could not be entirely isolated when examining the effects of real-world size. This makes it difficult to ascertain whether their results were driven by neural representations of perceived real-world size and/or perceived real-world depth. MEG and EEG studies focused on temporal processing of object size representations (S.-M. Khaligh-Razavi et al., 2018 [DOI](#); R. Wang et al., 2022 [DOI](#)) have been similarly susceptible to this limitation. Indeed, one recent study (Quek et al., 2023 [DOI](#)) provided evidence that perceived real-world depth could influence real-world size representations, further illustrating the necessity of investigating pure real-world size representations in the brain. Secondly, the stimuli used in these studies were cropped object stimuli against a plain white or grey background, which are not particularly naturalistic. More and more studies and datasets have highlighted the important role of naturalistic context in object recognition (Allen et al., 2022 [DOI](#); Gifford et al., 2022 [DOI](#); Grootswagers et al., 2022 [DOI](#); Hebart et al., 2019 [DOI](#); Stoinski et al., 2023 [DOI](#)). In ecological contexts, inferring the real-world size/distance of an object likely relies on a combination of bottom-up visual information and top-down knowledge about canonical object sizes for familiar objects. Incorporating naturalistic background context in experimental stimuli may produce more accurate assessments of the relative influences of visual shape representations (Bracci et al., 2017 [DOI](#); Bracci & Op de Beeck, 2016 [DOI](#); Proklova et al., 2016 [DOI](#)) and higher-level semantic information (Doerig et al., 2022 [DOI](#); Huth et al., 2012 [DOI](#); A. Y. Wang et al., 2022 [DOI](#)). Furthermore, most previous studies have tended to categorize size rather broadly, such as merely differentiating between big and small objects (S.-M. Khaligh-Razavi et al., 2018 [DOI](#); Konkle & Oliva, 2012b [DOI](#); R.

Wang et al., 2022 [↗](#)) or dividing object size into seven levels from small to big. To more finely investigate the representation of object size in the brain, it may be necessary to obtain a more continuous measure of size for a more detailed characterization.

Certainly, a minority of fMRI studies have attempted to utilize natural images and also engaged in more detailed size measurements to more precisely explore the encoding of object real-world size in different brain areas (Luo et al., 2023 [↗](#); Troiani et al., 2014 [↗](#)). However, no study has yet comprehensively overcome all the challenges and unfolded a clear processing timeline for object retinal size, real-world size, and real-world depth in human visual perception.

In the current study, we overcome these challenges by combining human EEG recordings, naturalistic stimulus images, artificial neural networks, and computational modeling approaches including representational similarity analysis (RSA) and partial correlation analysis to distinguish the neural representations of object real-world size, retinal size, and real-world depth. We applied our integrated computational approach to an open EEG dataset, THINGS EEG2 (Gifford et al., 2022 [↗](#)). Firstly, the visual image stimuli used in this dataset are more naturalistic and include objects that vary in real-world size, depth, and retinal size. This allows us to employ a multi-model representational similarity analysis to investigate relatively unconfounded representations of object real-world size, partialing out – and simultaneously exploring – these confounding features. Secondly, we are able to explore the neural dynamics of object feature processing in a more ecological context based on natural images in human object recognition. Thirdly, instead of categorizing object size into discrete levels, we applied a more continuous measure based on detailed behavioral measurements from an online size rating task, allowing us to more finely decode the representation of object size in the brain.

We first focus on unfolding the neural dynamics of statistically isolated object real-world size representations. The temporal resolution of EEG allows us the opportunity to investigate the representational timecourse of visual object processing, asking whether processing of perceived object real-world size precedes or follows processing of perceived depth, if these two properties are in fact processed independently.

We then attempt to further explore the underlying mechanisms of how human brains process object size and depth in natural images by integrating artificial neural networks (ANNs). In the domain of cognitive computational neuroscience, ANNs offer a complementary tool to study visual object recognition, and an increasing number of studies support that ANNs exhibit representations similar to human visual systems (Cichy et al., 2016 [↗](#); Güçlü & van Gerven, 2015 [↗](#); Yamins et al., 2014 [↗](#); Yamins & DiCarlo, 2016 [↗](#)). Indeed, a recent study found that ANNs also represent real-world size (Huang et al., 2022 [↗](#)); however, their use of a fixed retinal size image dataset with the same cropped objects as described above makes it similarly challenging to ascertain whether the results reflected real-world size and/or depth. Additionally, some recent work indicates that artificial neural networks incorporating semantic embedding and multimodal neural components might more accurately reflect human visual representations within visual areas and even the hippocampus, compared to vision-only networks (Choksi, Mozafari, et al., 2022 [↗](#); Choksi, Vanrullen, et al., 2022 [↗](#); Conwell et al., 2022 [↗](#); Doerig et al., 2022 [↗](#); Jozwik et al., 2023 [↗](#); A. Y. Wang et al., 2022 [↗](#)). Given that perception of real-world size may incorporate both bottom-up visual and top-down semantic knowledge about familiar objects, these models offer yet another novel opportunity to investigate this question. Utilizing both visual and visual-semantic models, as well as different layers within these models, ANNs provide us the approach to extract various image features, low-level visual information from early layers and higher-level information including both visual and semantic features from late layers.

The integrated computational approach by cross-modal representational comparisons we take with the current study allows us to compare how representations of perceived real-world size and depth emerge in both human brains and artificial neural networks. Unraveling the internal

representations of object size and depth features in both human brains and ANNs enables us to investigate how distinct spatial properties—retinal size, real-world depth, and real-world size—are encoded across systems, and to uncover the representational mechanisms and temporal dynamics through which real-world size emerges as a potentially higher-level, semantically grounded feature.

Materials and methods

Experimental design, stimuli images and EEG data

We utilized the open dataset from THINGS EEG2 (Gifford et al., 2022 [↗](#)), which includes EEG data from 10 healthy human subjects (age=28.5±4, 8 female and 2 male) participating in a rapid serial visual presentation (RSVP) paradigm with an orthogonal target detection task to ensure participants paid attention to the visual stimuli. All participants were seated at a fixed distance of 0.6 meters from the screen throughout the experiment. For each trial, subjects viewed one image (sized 500 × 500 pixels) for 100ms. Each subject viewed 16740 images of objects on a natural background for 1854 object concepts from THINGS dataset (Hebart et al., 2019 [↗](#)). For the current study, we used the ‘test’ dataset portion, which includes 16000 trials per subject corresponding to 200 images (200 object concepts, one image per concept) with 80 trials per image, providing high trial counts per condition necessary for reliable EEG decoding. In contrast, images in the training set were only shown four times each. This design choice ensured higher decoding reliability and greater signal quality for RSA. Before inputting the images to the ANNs, we reshaped image sizes to 224 × 224 pixels and normalized the pixel values of images to ImageNet statistics.

EEG data were collected using a 64-channel EASYCAP and a BrainVision actiCHamp amplifier. The EEG data were originally sampled at 1000Hz and online-filtered between 0.1 Hz and 100 Hz during acquisition, with recordings referenced to the Fz electrode. For preprocessing, no additional filtering was applied. Baseline correction was performed by subtracting the mean signal during the 100 ms pre-stimulus interval from each trial and channel separately. We used already pre-processed data from 17 channels with labels beginning with “O” or “P” (O1, Oz, O2, PO7, PO3, POz, PO4, PO8, P7, P5, P3, P1, Pz, P2) ensuring full coverage of posterior regions typically involved in visual object processing. The epoched data were then down-sampled to 100 Hz. We re-epoched EEG data ranging from 100ms before stimulus onset to 300ms after onset with a sample frequency of 100Hz. Thus the shape of our EEG data matrix for each trial was 17 channels × 40 time points.

ANN models

We applied two pre-trained ANN models: one visual model (ResNet-101 (He et al., 2016 [↗](#)) pretrained on ImageNet), and one multi-modal (visual+semantic) model (CLIP with a ResNet-101 backbone (Radford et al., 2021 [↗](#)) pretrained on YFCC-15M). Our motivation for selecting ResNet-50 and CLIP ResNet-50 was not to make a definitive comparison between model classes, but rather to include two widely used representatives of their respective categories—one trained purely on visual information (ResNet-50 on ImageNet) and one trained with joint visual and linguistic supervision (CLIP ResNet-50 on image–text pairs). These models are both highly influential and commonly used in computational and cognitive neuroscience, allowing for relevant comparisons with existing work (Choksi, Mozafari, et al., 2022 [↗](#); Choksi, Vanrullen, et al., 2022 [↗](#); Conwell et al., 2024 [↗](#); Song et al., 2024 [↗](#); A. Y. Wang et al., 2022 [↗](#)). We used THINGSvision (Muttenthaler & Hebart, 2021 [↗](#)) to obtain low- and high-level feature vectors of ANN activations from early and late layers (early layer: second convolutional layer; late layer: last visual layer) for the images.

Word2Vec model

To approximate the non-visual, pure semantic space of objects, we also applied a Word2Vec model, a natural language processing model for word embedding, pretrained on Google News corpus (Mikolov et al., 2013 [↗](#)), which contains 300-dimensional vectors for 3 million words and phrases. We input the words for each image's object concept (pre-labeled in THINGS dataset: Hebart et al., 2019 [↗](#)), instead of the visual images themselves. We used Gensim (Řehůřek & Sojka, 2010 [↗](#)) to obtain Word2Vec feature vectors for the objects in images.

Representational dissimilarity matrices (RDMs)

To conduct RSA across human EEG, artificial models, and our hypotheses corresponding to different visual features, we first computed representational dissimilarity matrices (RDMs) for different modalities (Figure 2 [↗](#)). The shape of each RDM was 200×200 , corresponding to pairwise dissimilarity between the 200 images. We extracted the 19900 cells from the upper half of the diagonal of each RDM for subsequent analyses.

Neural RDMs

From the EEG signal, we constructed timepoint-by-timepoint neural RDMs for each subject with decoding accuracy as the dissimilarity index (Figure 2A [↗](#)). Since EEG has a low SNR and includes rapid transient artifacts, Pearson correlations computed over very short time windows yield unstable dissimilarity estimates (Kappenman & Luck, 2010 [↗](#); Luck, 2014 [↗](#)) and may thus fail to reliably detect differences between images. In contrast, decoding accuracy - by training classifiers to focus on task-relevant features - better mitigates noise and highlights representational differences. We first conducted timepoint-by-timepoint classification-based decoding for each subject and each pair of images (200 images, 19900 pairs in total). We applied linear Support Vector Machine (SVM) to train and test a two-class classifier, employing a 5-time 5-fold cross-validation method, to obtain an independent decoding accuracy for each image pair and each timepoint. Therefore, we ultimately acquired 40 (1 per timepoint) EEG RDMs for each subject.

Hypothesis-based (HYP) RDMs

We constructed three hypothesis-based RDMs reflecting the different types of visual object properties in the naturalistic images (Figure 2B [↗](#)): Real-World Size RDM, Retinal Size RDM, and Real-World Depth RDM. We constructed these RDMs as follows:

1. For Real-World Size RDM, we obtained human behavioral real-world size ratings of each object concept from the THINGS+ dataset (Stoinski et al., 2022). In the THINGS+ dataset, 2010 participants (different from the subjects in THINGS EEG2) did an online size rating task and completed a total of 13024 trials corresponding to 1854 object concepts using a two-step procedure. In their experiment, first, each object was rated on a 520-unit continuous slider anchored by familiar reference objects (e.g., “grain of sand,” “microwave oven,” “aircraft carrier”) representing a logarithmic size range. Participants were not shown numerical values but used semantic anchors as guides. In the second step, the scale zoomed in around the selected region to allow for finer-grained refinement of the size judgment. Final size values were derived from aggregated behavioral data and rescaled to a range of 0–519 for consistency across objects, with the actual mean ratings across subjects ranging from 100.03 (‘grain of sand’) to 423.09 (‘subway’). We used these ratings as the perceived real-world size measure of the object concept pre-labeled in THINGS dataset (Hebart et al., 2019 [↗](#)) for each image. We then constructed the representational dissimilarity matrix by calculating the absolute difference between perceived real-world size ratings for each pair of images.

2. For Retinal Size RDM, we applied Adobe Photoshop (Adobe Inc., 2019) to crop objects corresponding to object labels from images manually. All cropping and measurement were conducted by one of the authors to ensure consistency across the dataset. The cropped object images have been made publicly available at <https://github.com/ZitongLu1996/RWsize> to facilitate future reuse and reproducibility. For each image, we obtained a rectangular region that precisely contains a single object, then measured the diagonal length of the segmented object in pixels as the retinal size measure (Konkle & Oliva, 2011). Due to our calculations being at the object level, if there were more than one same objects in an image, we cropped the most complete one to get more accurate retinal size. We then constructed the RDM by calculating the absolute difference between measured retinal size for each pair of images.
3. For Real-World Depth RDM, we calculated the perceived depth based on the measured retinal size index and behavioral real-world size ratings, such that $\text{real-world depth} / \text{visual image depth} = \text{real-world size} / \text{retinal size}$. Since visual image depth (viewing distance) is held constant across images in the task, inferred real-world depth is proportional to $\text{real-world size} / \text{retinal size}$. We then constructed the RDM by calculating the absolute difference between inferred real-world depth index for each pair of images.

ANN (and Word2Vec) model RDMs

We constructed a total of five model-based RDMs (Figure 2C). Our primary analyses used four ANN RDMs, corresponding to the early and late layers for both ResNet and CLIP (Figure S1). The early layer in ResNet refers to ResNet.maxpool layer, and the late layer in ResNet refers to ResNet.avgpool layer. The early layer in CLIP refers to CLIP.visual.avgpool layer, and the late layer in CLIP refers to CLIP.visual.attnpool layer. We also calculated a single Word2Vec RDM for the pure semantic analysis (Figure S2). For each RDM, we got the dissimilarities by calculating $1 - \text{Pearson correlation coefficient}$ between each pair of two vectors of the model features corresponding to two input images.

Representational similarity analyses (RSA) and statistical analyses

We conducted cross-modal representational similarity analyses between the three types of RDMs (Figure 2). All decoding and RSA analyses were implemented using NeuroRA (Lu & Ku, 2020).

EEG × ANN (or W2V) RSA

To measure the representational similarity between human brains and ANNs and confirm that ANNs have significantly similar representations to human brains, we calculated the Spearman correlation between the 40 timepoint-by-timepoint EEG neural RDMs and the 4 ANN RDMs corresponding to the representations of ResNet early layer, ResNet late layer, CLIP early layer, CLIP late layer, respectively. We also calculated temporal representational similarity between human brains (EEG RDMs) and the Word2Vec model RDM. Cluster-based permutation tests were conducted to determine the time windows of significant representational similarity. First, we performed one-sample t-tests (one-tailed testing) against zero to get the t-value for each timepoint, and extracted significant clusters. We computed the clustering statistic as the sum of t-values in each cluster. Then we conducted 1000 permutations of each subject's timepoint-by-timepoint similarities to calculate a null distribution of maximum clustering statistics. Finally, we assigned cluster-level p-values to each cluster of the actual representational timecourse by comparing its cluster statistic with the null distribution. Time-windows were determined to be significant if the p-value of the corresponding cluster was <0.05 .

EEG × HYP RSA

To evaluate how human brains temporally represent different visual features, we calculated the timecourse of representational similarity between the timepoint-by-timepoint EEG neural RDMs and the three hypothesis-based RDMs. To avoid correlations between hypothesis-based RDMs (**Figure 3A**) influencing comparison results, we calculated partial correlations and used a one-tailed test against the null hypothesis that the partial correlation was less than or equal to zero, testing whether the partial correlation was significantly greater than zero. In EEG × HYP partial correlation (**Figure 3D**), we correlated EEG RDMs with one hypothesis-based RDM (e.g., real-world size), while controlling for the other two (e.g., retinal size and real-world depth). Cluster-based permutation tests were performed as described above to determine the time windows of significant representational similarity. In addition, we conducted peak latency analysis to determine the latency of peak representational similarity for each type of visual information with the EEG signal. To assess the stability of peak latency estimates for each subject, we performed a bootstrap procedure across stimulus pairs. At each time point, the EEG RDM was vectorized by extracting the lower triangle (excluding the diagonal), resulting in 19,900 unique pairwise values. For each bootstrap sample, we resampled these 19,900 pairwise entries with replacement to generate a new pseudo-RDM of the same size. We then computed the partial correlation between the EEG pseudo-RDM and a given hypothesis RDM (e.g., real-world size), controlling for other feature RDMs, and obtained a time course of partial correlations. Repeating this procedure 1000 times and extracting the peak latency within the significant time window yielded a distribution of bootstrapped latencies, from which we got the bootstrapped mean latencies per subject. Paired t-tests (two-tailed) were conducted to assess the statistical differences in peak latencies between different visual features.

ANN (or W2V) × HYP RSA

To evaluate how different visual information is represented in ANNs, we calculated representational similarity between the ANN RDMs and hypothesis-based RDMs. As in the EEG × HYP RSA, we calculated partial correlations to avoid correlations between hypothesis-based RDMs. In ANN (or W2V) × HYP partial correlation (**Figure 3E** and **Figure 5A**), we correlated ANN (or W2V) RDMs with one hypothesis-based RDM (e.g., real-world size), while partialling out the other two. We also calculated the partial correlations between hypothesis-based RDMs and the Word2Vec RDM. To determine statistical significance, we conducted a bootstrap test. We shuffled the order of the cells above the diagonal in each ANN (or Word2Vec) RDM 1000 times. For each iteration, we calculated partial correlations corresponding to the three hypothesis-based RDMs. This produced a 1000-sample null distribution for each HYP × ANN (or W2V) RSA. We hypothesized that if the real similarity was higher than the 95% confidence interval of the null distribution, it indicated that ANN (or W2V) features validly encoded the corresponding visual feature.

Additionally, to explore how the naturalistic background present in the images might influence object real-world size, retinal size, and real-world depth representations, we conducted another version of the analysis by inputting cropped object images without background into ANN models to obtain object-only ANN RDMs (**Figure S3**). Then we performed the same ANN × HYP similarity analysis to calculate partial correlations between the hypothesis-based RDMs and object-only ANN RDM. (We didn't conduct the similarity analysis between timepoint-by-timepoint EEG neural RDMs with subjects viewing natural images and object-only ANN RDMs due to the input differences.)

Data and code accessibility

EEG data and images from THINGS EEG2 data are publicly available on OSF (<https://osf.io/3jk45/>). All Python analysis scripts will be available post-publication on GitHub (<https://github.com/ZitongLu1996/RWsize>).

Results

We conducted a cross-modal representational similarity analysis (**Figures 1**–**2**, see Method section for details) comparing the patterns of human brain activation while participants viewed naturalistic object images (timepoint-by-timepoint decoding of EEG data), the output of different layers of artificial neural networks and semantic language models fed the same stimuli (ANN and Word2Vec models), and hypothetical patterns of representational similarity based on behavioral and mathematical measurements of different visual image properties (perceived real-world object size, displayed retinal object size, and inferred real-world object depth).

Dynamic representations of object size and depth in human brains

To explore if and when human brains contain distinct representations of perceived real-world size, retinal size, and real-world depth, we constructed timepoint-by-timepoint EEG neural RDMs (**Figure 2A**), and compared these to three hypothesis-based RDMs corresponding to different visual image properties (**Figure 2B**). Firstly, we confirmed that the hypothesis-based RDMs were indeed correlated with each other (**Figure 3A**), and without accounting for the confounding variables, Spearman correlations between the EEG and each hypothesis-based RDM revealed overlapping periods of representational similarity (**Figure 3B**). In particular, representational similarity with real-world size (from 90 to 120ms and from 170 to 240ms) overlapped with the significant time-windows of other features, including retinal size from 70 to 210ms, and real-world depth from 60 to 130ms and from 180 to 230ms. But critically, with the partial correlations, we isolated their independent representations. The partial correlation results reveal a relatively unconfounded representation of object real-world size in the human brain from 170 to 240ms after stimulus onset, independent from retinal size and real-world depth, which showed significant representational similarity at different time windows (retinal size from 90 to 200ms, and real-world depth from 60 to 130ms and 270 to 300ms) (**Figure 3D**).

Peak latency results showed that neural representations of real-world size, retinal size and real-world depth reached their peaks at different latencies after stimulus onset (real-world depth: ~87ms, retinal size: ~138ms, real-world size: ~206ms, **Figure 3C**). The representation of real-world size had a significantly later peak latency than that of both retinal size, $t(9)=4.30$, $p=.002$, and real-world depth, $t(9)=18.58$, $p<.001$. And retinal size representation had a significantly later peak latency than real-world depth, $t(9)=3.72$, $p=.005$. These varying peak latencies imply an encoding order for distinct visual features, transitioning from real-world depth through retinal size, and then to real-world size.

Artificial neural networks also reflect distinct representations of object size and depth

To test how ANNs process these visual properties, we input the same stimulus images into ANN models and got their latent features from early and late layers (**Figure 2C**), and then conducted comparisons between the ANN RDMs and hypothesis-based RDMs. Parallel to our findings of dissociable representations of real-world size, retinal size, and real-world depth in the human brain signal, we also found dissociable representations of these visual features in ANNs (**Figure 3E**). Our partial correlation RSA analysis showed that early layers of both ResNet and CLIP had significant real-world depth and retinal size representations, whereas the late layers of both ANNs were dominated by real-world size representations, though there was also weaker retinal size representation in the late layer of ResNet and real-world depth representation in the late layer of CLIP (additional results of the extended analysis of multiple layers in ResNet and CLIP are shown in **Figure S4**). The detailed statistical results are shown in **Table S1**.

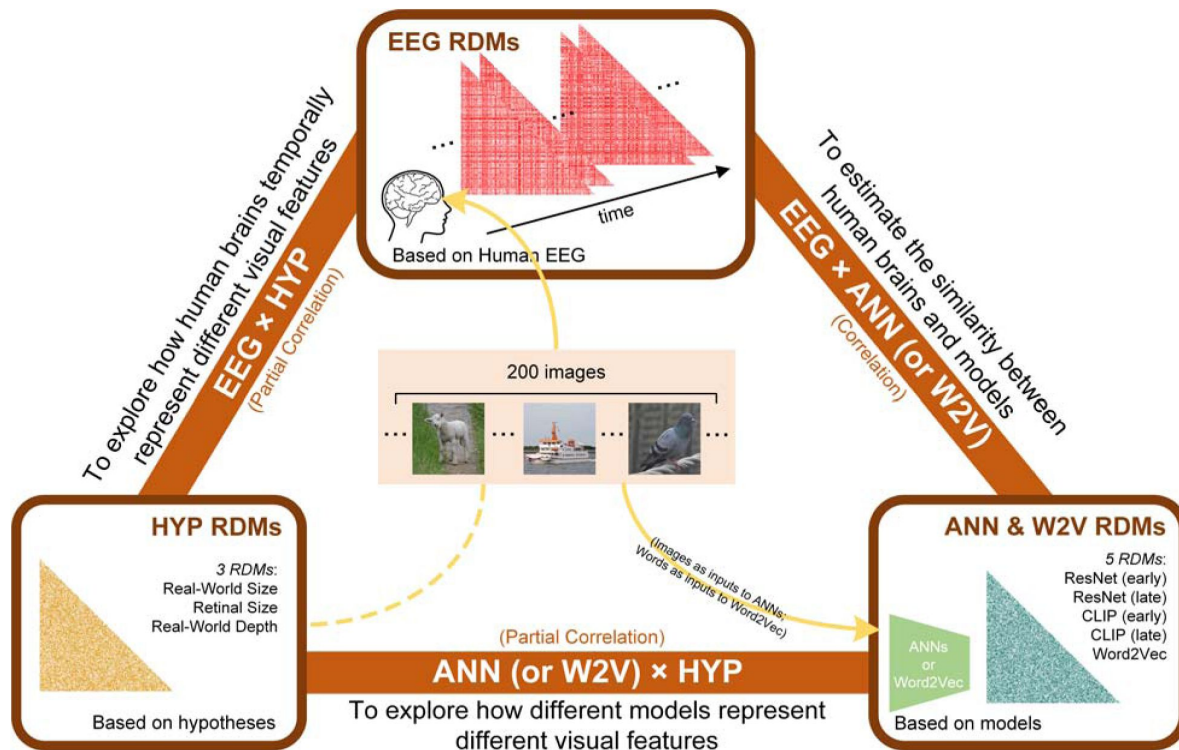


Figure 1

Overview of our analysis pipeline including constructing three types of RDMs and conducting comparisons between them.

We computed RDMs from three sources: neural data (EEG), hypothesized object features (real-world size, retinal size, and real-world depth), and artificial models (ResNet, CLIP, and Word2Vec). Then we conducted cross-modal representational similarity analyses between: EEG x HYP (partial correlation, controlling for other two HYP features), ANN (or W2V) x HYP (partial correlation, controlling for other two HYP features), and EEG x ANN (correlation).

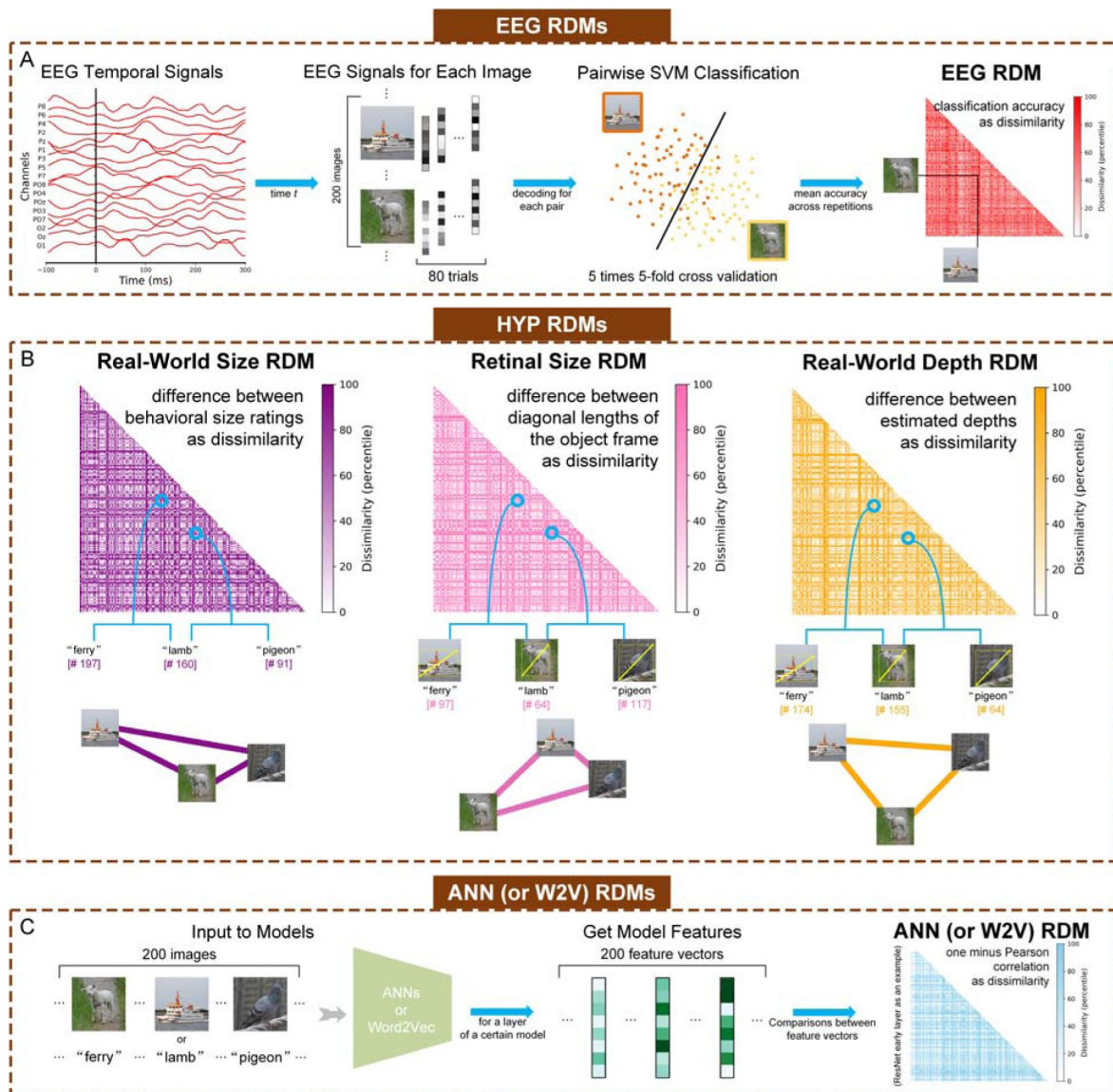


Figure 2

Methods for calculating neural (EEG), hypothesis-based (HYP), and artificial neural network (ANN) & semantic language processing (Word2Vec, W2V) model-based representational dissimilarity matrices (RDMs).

(A) Steps of computing the neural RDMs from EEG data. EEG analyses were performed in a time-resolved manner on 17 channels as features. For each time t , we conducted pairwise cross-validated SVM classification. The classification accuracy values across different image pairs resulted in each 200×200 RDM for each time point. (B) Calculating the three hypothesis-based RDMs: Real-World Size RDM, Retinal Size RDM, and Real-World Depth RDM. Real-world size, retinal size, and real-world depth were calculated for the object in each of the 200 stimulus images. The number in the bracket represents the rank (out of 200, in ascending order) based on each feature corresponding to the object in each stimulus image (e.g. “ferry” ranks 197th in real-world size from small to big out of 200 objects). The connection graph to the right of each RDM represents the relative representational distance of three stimuli in the corresponding feature space. (C) Steps of computing the ANN and Word2Vec RDMs. For ANNs, the inputs were the resized images, and for Word2Vec, the inputs were the words of object concepts. For clearer visualization, the shown RDMs were separately histogram-equalized (percentile units).

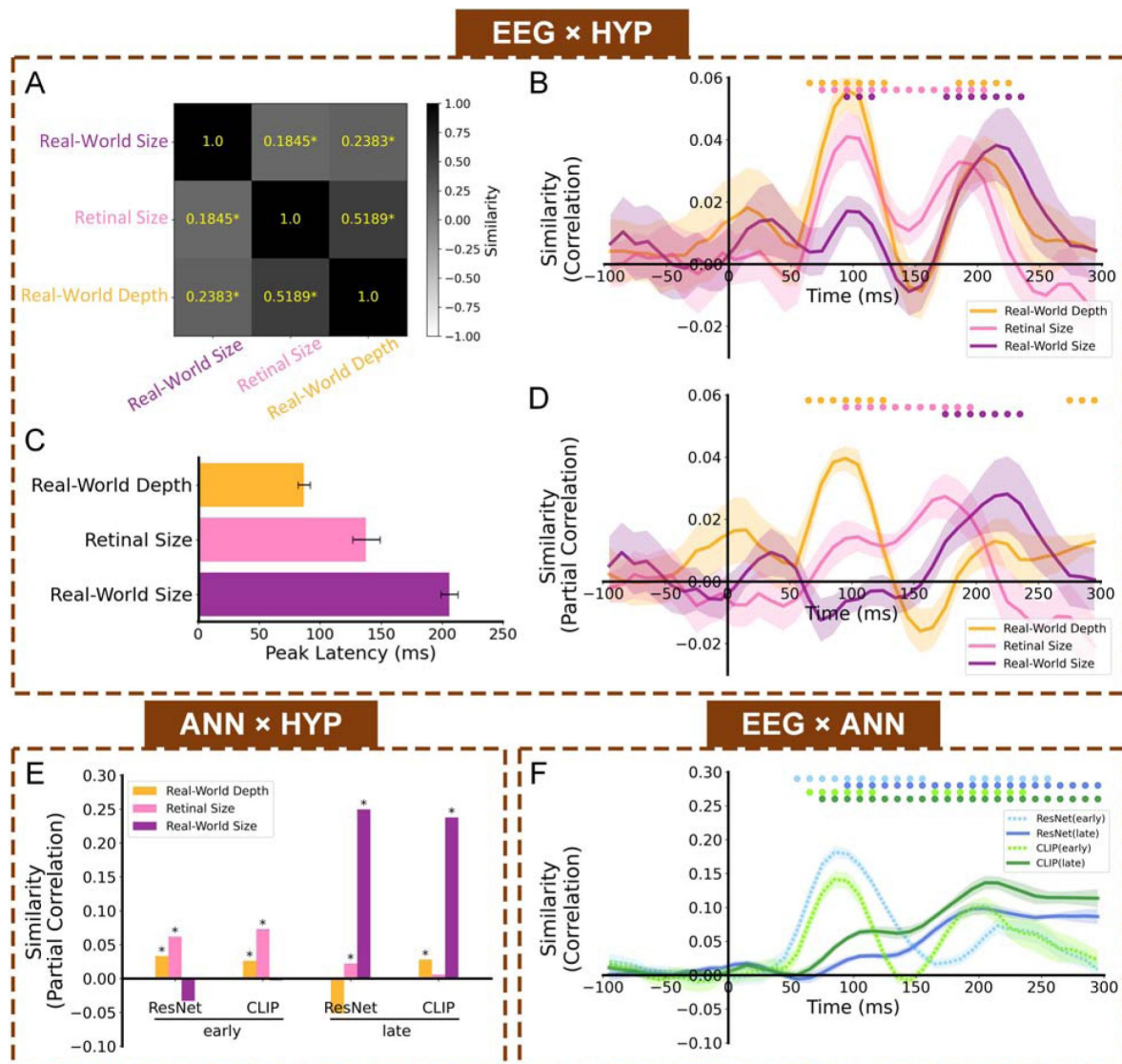


Figure 3

Cross-modal RSA results.

(A) Similarities (Spearman correlations) between three hypothesis-based RDMs. Asterisks indicate a significant similarity, $p < .05$. (B) Representational similarity time courses (full Spearman correlations) between EEG neural RDMs and hypothesis-based RDMs. (C) Temporal latencies for peak similarity (partial Spearman correlations) between EEG and the 3 types of object information. Error bars indicate $\pm$ SEM. Asterisks indicate significant differences across conditions ($p < .05$); (D) Representational similarity time courses (partial Spearman correlations) between EEG neural RDMs and hypothesis-based RDMs. (E) Representational similarities (partial Spearman correlations) between the four ANN RDMs and hypothesis-based RDMs of real-world depth, retinal size, and real-world size. Asterisks indicate significant partial correlations (bootstrap test, $p < .05$). (F) Representational similarity time courses (Spearman correlations) between EEG neural RDMs and ANN RDMs. Color-coded small dots at the top indicate significant timepoints (cluster-based permutation test, $p < .05$). Shaded area reflects $\pm$ SEM across the 10 subjects.

Thus, ANNs provide another approach to understand the formation of different visual features, offering convergent results with the EEG representational analysis, where retinal size was reflected most in the early layers of ANNs, while object real-world size representations didn't emerge until late layers of ANNs, consistent with a potential role of higher-level visual information, such as the semantic information of object concepts.

Representational similarity between human EEG and artificial neural networks

To directly examine the representational similarity between ANNs and human EEG signals, we compared the timepoint-by-timepoint EEG neural RDMs and the ANN RDMs. This analysis allowed us to assess how different stages of visual processing in the human brain align temporally with hierarchical representations in ANNs. As shown in **Figure 3F**, the early layer representations of both ResNet and CLIP (ResNet.maxpool layer and CLIP.visual.avgpool) showed significant correlations with early EEG time windows (early layer of ResNet: 40-280ms, early layer of CLIP: 50-130ms and 160-260ms), while the late layers (ResNet.avgpool layer and CLIP.visual.attnpool layer) showed correlations extending into later time windows (late layer of ResNet: 80-300ms, late layer of CLIP: 70-300ms). Although there is substantial temporal overlap between early and late model layers, the overall pattern suggests a rough correspondence between model hierarchy and neural processing stages.

We further extended this analysis across intermediate layers of both ResNet and CLIP models (from early to late, ResNet: ResNet.maxpool, ResNet.layer1, ResNet.layer2, ResNet.layer3, ResNet.layer4, ResNet.avgpool; from early to late, CLIP: CLIP.visual.avgpool, CLIP.visual.layer1, CLIP.visual.layer2, CLIP.visual.layer3, CLIP.visual.layer4, CLIP.visual.attnpool). The results, now included in **Figure S5**, show a consistent trend: early layers exhibit higher similarity to early EEG time points, and deeper layers show increased similarity to later EEG stages. This pattern of early-to-late correspondence aligns with previous findings that convolutional neural networks exhibit similar hierarchical representations to those in the brain visual cortex (Cichy et al., 2016; Güçlü & van Gerven, 2015; Kietzmann et al., 2019; Yamins & DiCarlo, 2016): that both the early stage of brain processing and the early layer of the ANN encode lower-level visual information, while the late stage of the brain and the late layer of the ANN encode higher-level visual information. Notably, early brain responses showed stronger similarity to early ResNet layers than to CLIP layers, consistent with prior work suggesting that early visual processing is more closely aligned with purely visual models (Greene & Hansen, 2020). In contrast, at later time windows, brain activity more closely resembled late CLIP layers, possibly reflecting the integration of visual and semantic information. However, it is also possible that these differences between ResNet and CLIP reflect differences in training data scale and domain.

To contextualize how much of the shared variance between EEG and ANN representations is driven by the specific visual object features we tested above, we conducted a partial correlation analysis between EEG RDMs and ANN RDMs controlling for the three hypothesis-based RDMs (**Figure S6**). The EEG×ANN similarity results remained largely unchanged, suggesting that ANN representations capture much more additional rich representational structure beyond these features. Similarly, a supplemental EEG noise ceiling analysis (**Figure S7**) shows that the observed EEG–HYP similarity values are substantially below a theoretical upper bound of explainable variance. This outcome is fully expected: Each of our HYP RDMs captures only a specific aspect of the neural representational structure, rather than attempting to account for the totality of the EEG or ANN signal. Our goal is not to optimize model performance or maximize fit, but to probe if these different components of object information are reflected – and dissociated – in the temporal dynamics and processing stages of the brain's responses. In parallel, these supplemental findings help to situate our hypothesis-driven analysis within the broader representational capacity of the brain and the models.

Real-world size as a stable and higher-level dimension in object space

An important aspect of the current study is the use of naturalistic visual images as stimuli, in which objects were presented in their natural contexts, as opposed to cropped images of objects without backgrounds. In natural images, background can play an important role in object perception. How dependent are the above results on the presence of naturalistic background context? To investigate how image context influences object size and depth representations, we next applied a reverse engineering method, feeding the ANNs with modified versions of the stimulus images containing cropped objects without background, and evaluating the ensuing ANN representations compared to the same original hypothesis-based RDMs. If the background significantly contributes to the formation of certain feature representations, we may see some encoding patterns in ANNs disappear when the input only includes the pure object but no background.

Compared to results based on images with background, the ANNs based on cropped-object modified images showed weaker overall representational similarity for all features (**Figure 4**). In the early layers of both ANNs, we now only observed significantly preserved retinal size representations (which is a nice validity check, since retinal size measurements were based purely on the physical object dimensions in the image, independent of the background). Real-world depth representations were almost totally eliminated, with only a small effect in the late layer of ResNet. However, we still observed a preserved pattern of real-world size representations, with significant representational similarity in the late layers of both ResNet and CLIP, and not in the early layers. The detailed statistical results are shown in **Table S2**. Even though the magnitude of representational similarity for object real-world size decreased when we removed the background, this high-level representation was not entirely eliminated. This finding suggests that background information does indeed influence object processing, but the representation of real-world size seems to be a relatively stable higher-level feature. On the other hand, representational formats of real-world depth changed when the input lacked background information. The deficiency of real-world depth representations in early layers, compared to when using full-background images, might suggest that the human brain typically uses background information to estimate object depth, though the significant effect in the late layer of ResNet in background-absent condition might also suggest that the brain (or at least ANN) has additional ability to integrate size information to infer depth when there is no background. These results show that real-world size emerges in the later layers of ANNs regardless of image background information, and – based on our prior EEG results – although we could not test object-only images in the EEG data, we hypothesize that a similar temporal profile would be observed in the brain, even for object-only images.

The relative robustness of real-world size representations – even in the absence of background – raises an important question: might real-world size be encoded as a more abstract, conceptual-level dimension in the brain's object representation space? If so, we might expect it to be driven not only by higher-level visual information, but also potentially by purely semantic information about familiar objects. To test this, we extracted object names from each image and input the object names into a Word2Vec model to obtain a Word2Vec RDM (**Figure S2**), and then conducted a partial correlation RSA comparing the Word2Vec representations with the hypothesis-based RDMs (**Figure 5A**). The results showed a significant real-world size representation ($r=0.1871$, $p<0.001$) but no representation of retinal size ($r=-0.0064$, $p=0.8148$) or real-world depth ($r=-0.0040$, $p=.7151$) from Word2Vec. Also, the significant time-window (90-300ms) of similarity between Word2Vec RDM and EEG RDMs (**Figure 5B**) contained the significant time-window of EEG x real-world size representational similarity (**Figure 3B**).

Figure 4

Contribution of image backgrounds to object size and depth representations.

Representational similarity results (partial Spearman correlations) between ANNs fed inputs of cropped object images without backgrounds and the hypothesis-based RDMs. Stars above bars indicate significant partial correlations (bootstrap test, $p < .05$).

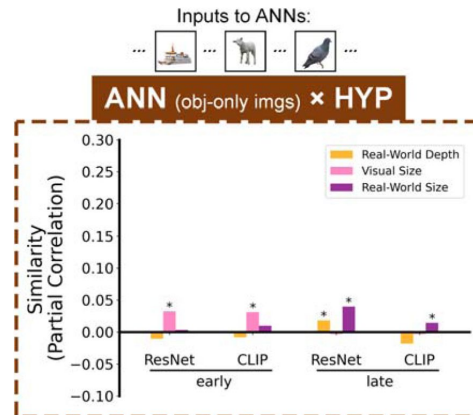
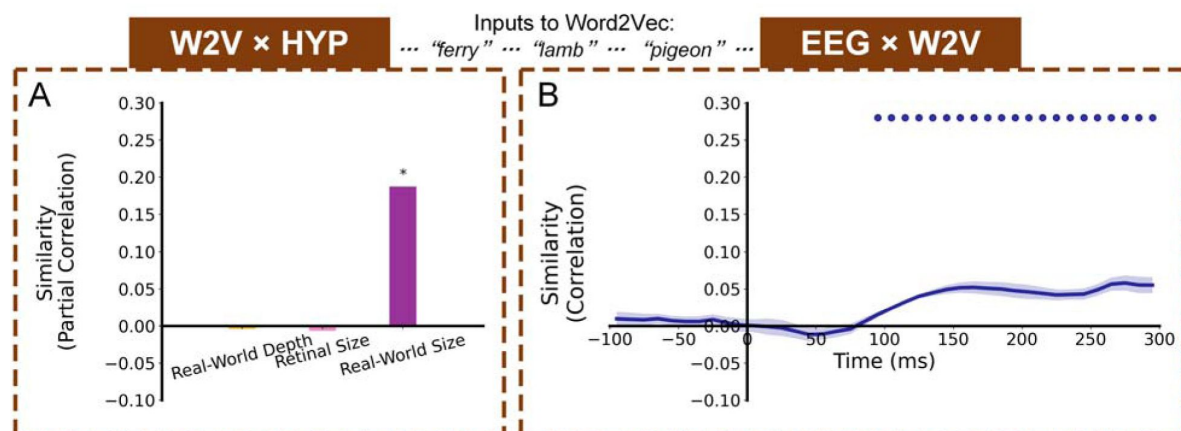


Figure 5

Representation similarity with a non-visual semantic language processing model (Word2Vec) fed word inputs corresponding to the images' object concepts.

(A) Representational similarity results (partial Spearman correlations) between Word2Vec RDM and hypothesis-based RDMs. Stars above bars indicate significant partial correlations (bootstrap test, $p < .05$). (B) Representational similarity time course (Spearman correlations) between EEG RDMs (neural activity while viewing images) and Word2Vec RDM (fed corresponding word inputs). Color-coded small dots at the top indicate significant timepoints (cluster-based permutation test, $p < .05$). Line width reflects $\pm$ SEM across the 10 subjects.



To further probe the relationship between real-world size and semantic information, and to examine whether Word2Vec captures variances in EEG signals beyond that explained by visual models, we conducted two additional analyses. First, we performed a partial correlation between EEG RDMs and the Word2Vec RDM, while regressing out four ANN RDMs (early and late layers of both ResNet and CLIP) (**Figure S8** [↗](#)). We found that semantic similarity remained significantly correlated with EEG signals across sustained time windows (100-190ms and 250-300ms), indicating that Word2Vec captures neural variance not fully explained by visual or visual-language models. Second, we conducted a variance partitioning analysis, in which we decomposed the variance in EEG RDMs explained by three hypothesis-based RDMs and the semantic RDM (Word2Vec RDM), and we still found that real-world size explained unique variance in EEG even after accounting for semantic similarity (**Figure S9** [↗](#)). And we also observed a substantial shared variance jointly explained by real-world size and semantic similarity and a unique variance of semantic information. These results suggest that real-world size is indeed partially semantic in nature, but also has independent neural representation not fully explained by general semantic similarity.

Both the reverse engineering manipulation and Word2Vec findings corroborate that object real-world size representation, unlike retinal size and real-world depth, emerges in both image- and semantic-level in object space.

Discussion

Our study applied computational methods to distinguish the representations of objects' perceived real-world size, retinal size, and inferred real-world depth features in both human brains and ANNs. Consistent with prior studies reporting real-world size representations ([Huang et al., 2022](#) [↗](#); [S.-M. Khaligh-Razavi et al., 2018](#) [↗](#); [Konkle & Caramazza, 2013](#) [↗](#); [Konkle & Oliva, 2012b](#) [↗](#); [Luo et al., 2023](#) [↗](#); [Quek et al., 2023](#) [↗](#); [R. Wang et al., 2022](#) [↗](#)), we found that both human brains and ANNs contain significant information about real-world size. Critically, our study goes beyond the contributions of prior studies in several key ways, offering both theoretical and methodological advances: (a) we eliminated the confounding impact of perceived real-world depth (in addition to retinal size) on the real-world size representation; (b) we conducted analyses based on more ecologically valid naturalistic images; (c) we obtained precise feature values for each object in every image, instead of simply dividing objects into 2 or 7 coarse categories; and (d) we utilized a multi-modal, partial correlation RSA that combines EEG, hypothesis-based models, and ANNs. Our approach allowed us to investigate representational time courses and reverse engineering manipulations in unparalleled detail. By integrating EEG data with hypothesis-based models and ANNs, this method offers a powerful tool for dissecting the neural underpinnings of object size and depth perception in more ecological contexts, which enriches our comprehension of the brain's representational mechanisms.

Using EEG we uncovered a representational timeline for visual object processing, with object real-world depth information represented first, followed by retinal size, and finally real-world size. While size and depth are highly correlated to each other, our results suggest that the human brain indeed has dissociated time courses and mechanisms to process them. The later representation time-window for object real-world size may suggest that the brain requires more sophisticated, higher-level information to form this representation, perhaps incorporating semantic and/or memory information about familiar objects, which was corroborated by our ANN and Word2Vec analyses. These findings also align with a recent fMRI study ([Luo et al., 2023](#) [↗](#)) using natural images to explore the neural selectivity for real-world size, finding that low-level visual information could hardly account for neural size preferences, although that study did not consider covariables like retinal size and real-world depth.

In contrast to the later-emerging real-world size representations, it makes sense that retinal size representations could be processed more quickly based on more fundamental, lower-level information such as shape and edge discrimination. Interestingly, real-world depth representations emerged even earlier than retinal size, suggesting that depth feature may precede real-world size processing. This early emergence might reflect the role of depth cues in facilitating object segmentation, as proposed in classical theories of vision (Marr, 1982 [↗](#)), where depth helps delineate object boundaries from complex backgrounds. Additionally, there was a secondary, albeit substantially later, significant depth representation time-window, which might indicate that our brains also have the ability to integrate object retinal size and higher-level real-size information to form the final representation of real-world depth. Our comparisons between human brains and artificial models and explorations on ANNs and Word2Vec offer further insights and suggest that although real-world object size and depth are closely related, object real-world size appears to be a more stable and higher-level dimension.

The concept of ‘object space’ in cognitive neuroscience research is crucial for understanding how various visual features of objects are represented. Historically, various visual features have been considered important dimensions in constructing object space, including animate-inanimate (Kriegeskorte et al., 2008 [↗](#); Naselaris et al., 2012 [↗](#)), spikiness (Bao et al., 2020 [↗](#); Coggan & Tong, 2023 [↗](#)), and physical appearance (Edelman et al., 1998 [↗](#)). In this study, we focus on one particular dimension, real-world size (Huang et al., 2022 [↗](#); Konkle & Caramazza, 2013 [↗](#); Konkle & Oliva, 2012b [↗](#)). How we generate neural distinctions of different object real-world size and where this ability comes from remain uncertain. Some previous studies found that object shape rather than texture information could trigger neural size representations (Huang et al., 2022 [↗](#); Long et al., 2016 [↗](#), 2018 [↗](#); R. Wang et al., 2022 [↗](#)). Our results attempt to further advance their findings, that object real-world size is a stable and higher-level dimension substantially driven by object semantics in object space.

Increasingly, research has begun to use ANNs to study the mechanisms of object recognition (Ayzenberg et al., 2023 [↗](#); Cichy & Kaiser, 2019 [↗](#); Doerig et al., 2023 [↗](#); Kanwisher et al., 2023 [↗](#)). We can explore how the human brain processes information at different levels by comparing brain activity with models (Cichy et al., 2016 [↗](#); S. M. Khaligh-Razavi & Kriegeskorte, 2014 [↗](#); Kuzovkin et al., 2018 [↗](#); Xie et al., 2020 [↗](#)), and we can also analyze the representation patterns of the models with some specific manipulations and infer potential processing mechanisms in the brain (Golan et al., 2020 [↗](#); Huang et al., 2022 [↗](#); Lu & Ku, 2023 [↗](#); Xu et al., 2021 [↗](#)). In the current study, our comparisons result between EEG signals and different ANNs showed that the visual model’s early layer had a higher similarity to the brain in the early stage, while the visual-semantic model’s late layer had a higher similarity to the brain in the late stage. In addition, to assess whether this EEG-ANN similarity pattern generalizes to more biologically inspired architectures, we conducted the same analyses and found that CORnet show similar patterns to those observed for ResNet and CLIP, providing converging evidence across models (**Figure S10** [↗](#)). However, for the representation of objects, partial correlation results for different ANNs didn’t demonstrate the superiority of the multi-modal model at late layers. This might be due to models like CLIP, which contain semantic information, learning more complex image descriptive information (like the relationship between object and the background in the image). Real-world size might be a semantic dimension of the object itself, and its representation does not require overall semantic descriptive information of the image. In contrast, retinal size and real-world depth could rely on image background information for estimation, thus their representations in the CLIP late layer disappeared when input images had only pure object but no background.

Although our study does not directly test specific models of visual object processing, the observed temporal dynamics provide important constraints for theoretical interpretations. In particular, we find that real-world size representations emerge significantly later than low-level visual features such as retinal size and depth. This temporal profile is difficult to reconcile with a purely feedforward account of visual processing (e.g., DiCarlo et al., 2012 [↗](#)), which posits that object

properties are rapidly computed in a sequential hierarchy of increasingly complex visual features. Instead, our results are more consistent with frameworks that emphasize recurrent or top-down processing, such as the reverse hierarchy theory (Hochstein & Ahissar, 2002 [DOI](#)), which suggests that high-level conceptual information may emerge later and involve feedback to earlier visual areas. This interpretation is further supported by representational similarities with late-stage artificial neural network layers and with a semantic word embedding model (Word2Vec), both of which reflect learned, abstract knowledge rather than low-level visual features. Taken together, these findings suggest that real-world size is not merely a perceptual attribute, but one that draws on conceptual or semantic-level representations acquired through experience. While our EEG analyses focused on posterior electrodes and thus cannot definitively localize cortical sources, we see this study as a step toward linking low-level visual input with higher-level semantic knowledge. Future work incorporating broader spatial coverage (e.g., anterior sensors), source localization, or complementary modalities such as MEG and fMRI will be critical to adjudicate between alternative models of object representation and to more precisely trace the origin and flow of real-world size information in the brain. Moreover, building on the promising findings of our study, future work may further delve into the detailed processes of object processing and object space. One important problem to solve is how real-world size interacts with other object dimensions in object space. In addition, our approach could be used with future studies investigating other influences on object processing, such as how different task conditions impact and modulate the processing of various visual features.

Moreover, we must also emphasize that in this study, we were concerned with perceived real-world size and depth reflecting a perceptual estimation of our world, which are slightly different from absolute physical size and depth. The differences in brain encoding between perceived and absolute physical size and depth require more comprehensive measurements of an object's physical attributes for further exploration. Also, we focused on perceiving depth and size from 2D images in this study, which might have some differences in brain mechanism compared to physically exploring the 3D world. Nevertheless, we believe our study offers a valuable contribution to object recognition, especially the encoding process of object real-world size in natural images. Additionally, we acknowledge that our metric for real-world depth was derived indirectly as the ratio of perceived real-world size to retinal size. While this formulation is grounded in geometric principles of perspective projection and served the purpose of analytically dissociating depth from size in our RSA framework, it remains a proxy rather than a direct measure of perceived egocentric distance. Future work incorporating behavioral or psychophysical depth ratings would be valuable for validating and refining this metric.

In conclusion, we used computational methods to distinguish the representations of real-world size, retinal size, and real-world depth features of objects in ecologically natural images in both human brains and ANNs. We found an unconfounded representation of object real-world size, which emerged at later time windows in the human EEG signal and at later layers of artificial neural networks compared to real-world depth, and which also appeared to be preserved as a stable dimension in object space. Thus, although size and depth properties are closely correlated, the processing of perceived object size and depth may arise through dissociated time courses and mechanisms. Our research provides a temporally resolved characterization of how certain key object properties – such as object real-world size, depth, and retinal size – are represented in the brain, which advances our understanding of real-world size encoding as a stable and semantically grounded dimension in object space, and contributes to the development of more brain-like visual models.

Acknowledgements

This work was supported by research grants from the National Institutes of Health (R01-EY025648) and from the National Science Foundation (NSF 1848939) to JDG. The authors declare no competing financial interests.

Supplementary materials

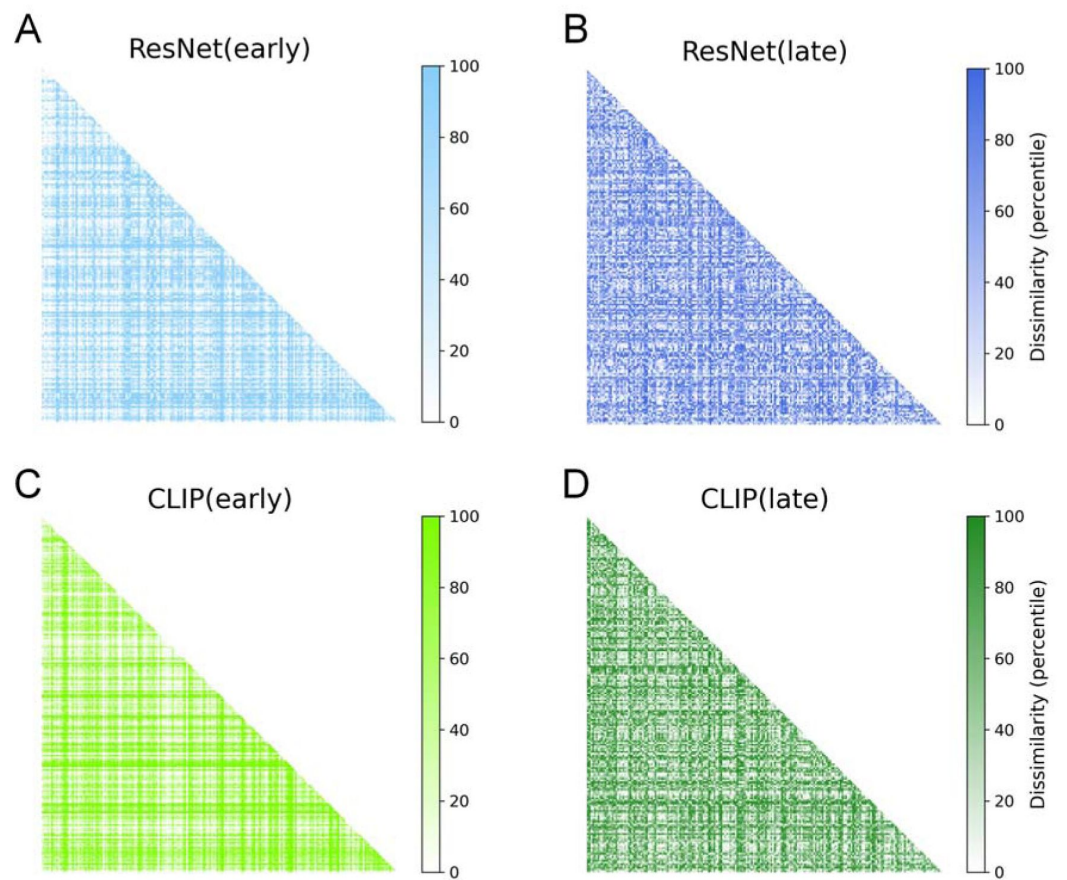


Figure S1.

Four ANN RDMs of ResNet early layer, ResNet late layer, CLIP early layer, and CLIP late later.

Figure S2.

Word2Vec RDMs.

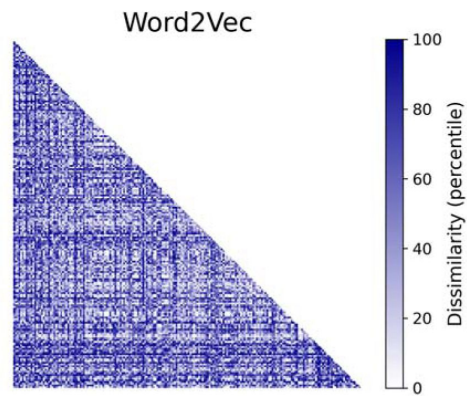


Figure S3.

Four ANN RDMs with inputs of cropped object images without background of ResNet early layer, ResNet late layer, CLIP early layer, and CLIP late later.

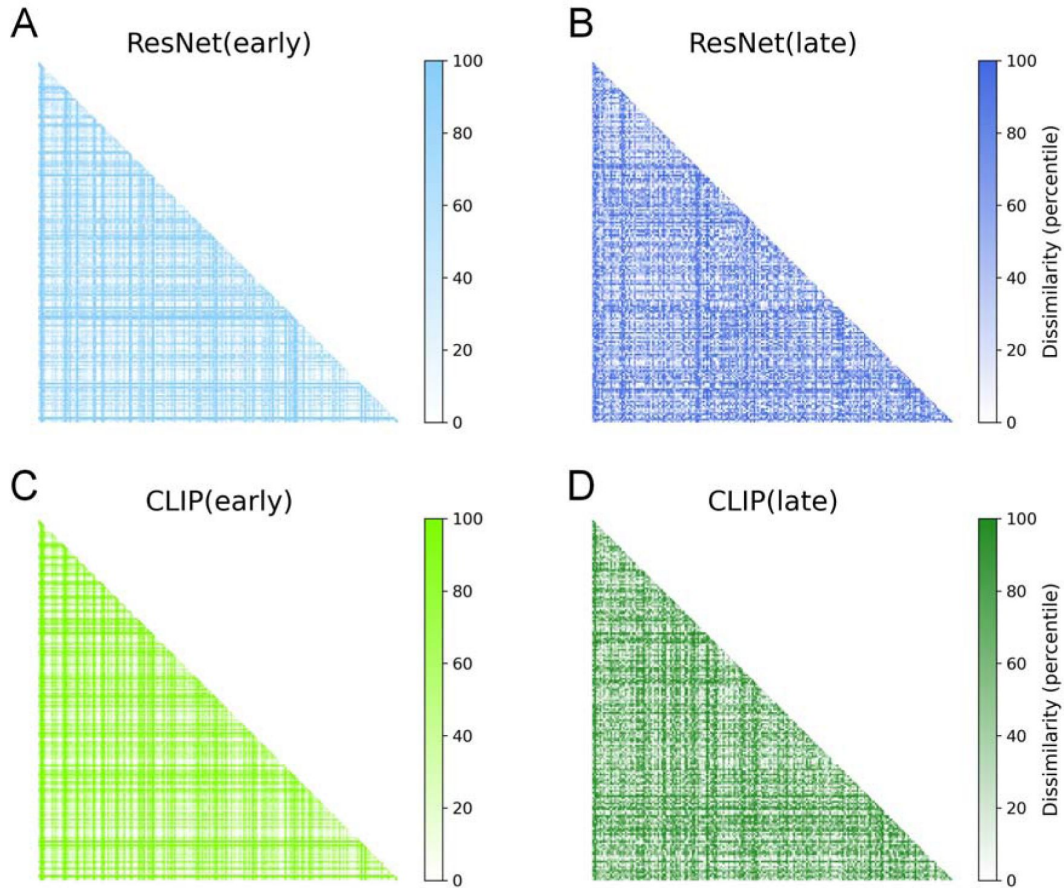


Figure S4.

Representational similarities (partial Spearman correlations) between ANN RDMs from multiple layers in ResNet and CLIP and hypothesis-based RDMs of real-world depth, retinal size, and real-world size (as extended results of Figure 3E).

Asterisks indicate significant partial correlations (bootstrap test, $p < .05$).

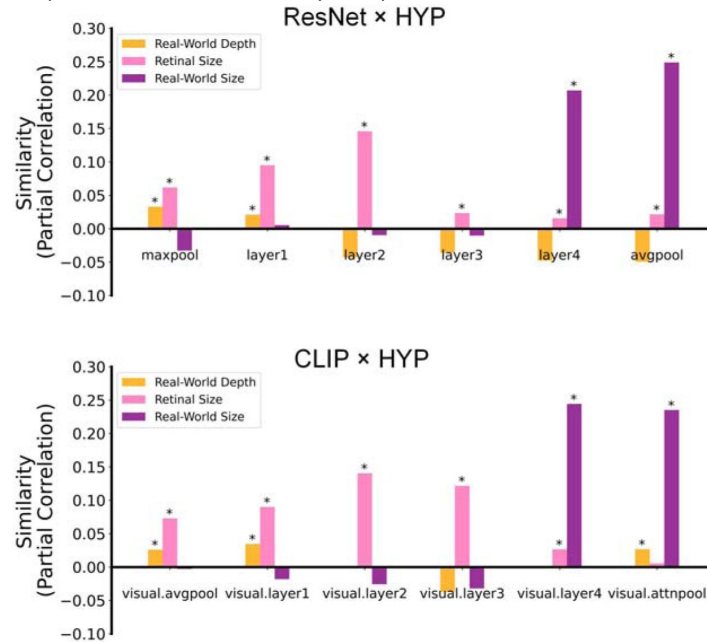


Figure S5.

Representational similarity time courses (Spearman correlations) between EEG neural RDMs and ANN RDMs from multiple layers (as extended results of Figure 3F).

(A) Similarity between EEG and ResNet's multiple layers (from early layer to late layer: ResNet.maxpool layer, ResNet.layer1 layer, ResNet.layer2 layer, ResNet.layer3 layer, ResNet.layer4 layer, ResNet.avgpool). (B) Similarity between EEG and CLIP's multiple layers (from early layer to late layer: CLIP.visual.avgpool layer, CLIP.visual.layer1 layer, CLIP.visual.layer2 layer, CLIP.visual.layer3 layer, CLIP.visual.layer4 layer, CLIP.visual.attnpool). Color-coded small dots at the top indicate significant timepoints (cluster-based permutation test, $p < .05$). Shaded area reflects $\pm$ SEM across the 10 subjects.

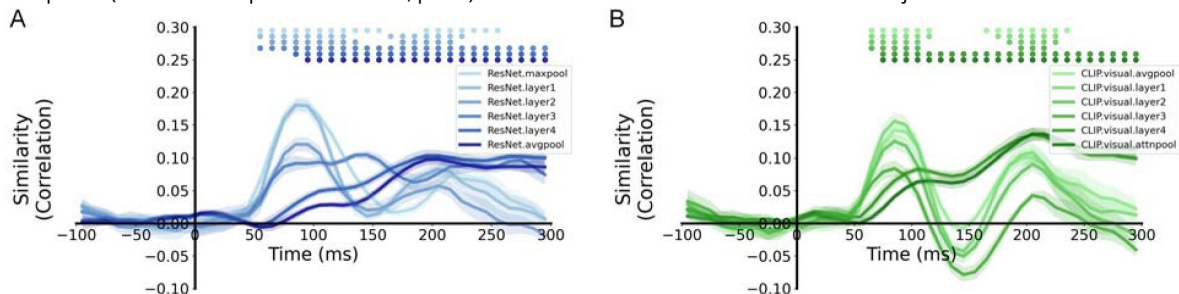


Figure S6.

Representational similarity time courses (partial Spearman correlations) between EEG neural RDMs and ANN RDMs controlling for the three hypothesis-based RDMs.

Color-coded small dots at the top indicate significant timepoints (cluster-based permutation test, $p < .05$). Shaded area reflects $\pm$ SEM across the 10 subjects.

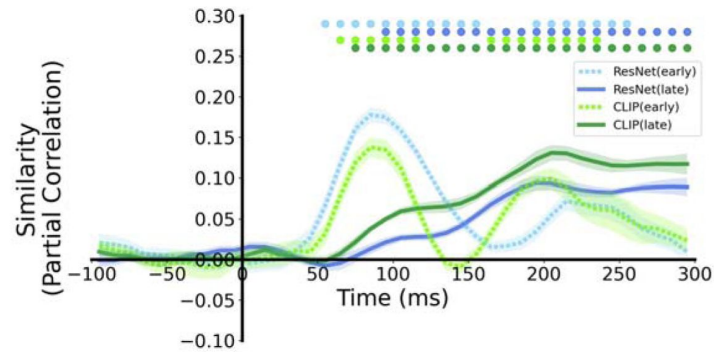


Figure S7.

Noise ceiling of representational similarity analysis based on EEG neural RDMs.

Shaded area reflects the lower and upper bound.

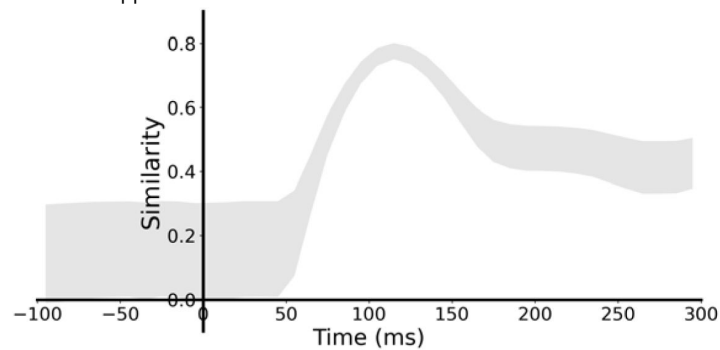


Figure S8.

Representational similarity time courses (partial Spearman correlations) between EEG neural RDMs and the Word2Vec RDM controlling for four ANN RDMs (ResNet early/late and CLIP early/late layers).

Color-coded small dots at the top indicate significant timepoints (cluster-based permutation test, $p < .05$). Shaded area reflects $\pm$ SEM across the 10 subjects.

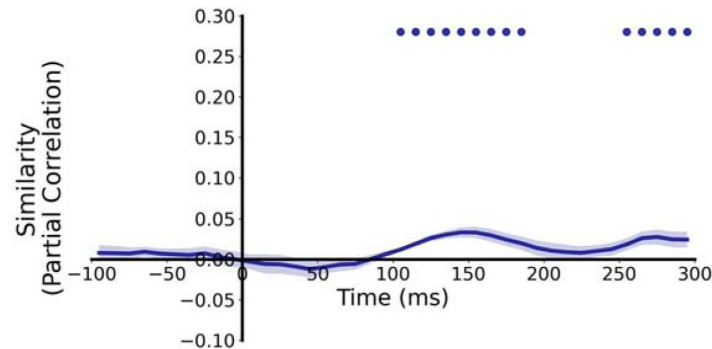


Figure S9.

Unique and shared variance of real-world size and semantic information (from Word2Vec) in EEG neural representations using variance partitioning analysis while controlling retinal size and real-world depth.

Color-coded small dots at the top indicate significant timepoints (cluster-based permutation test, $p < .05$). Shaded area reflects $\pm$ SEM across the 10 subjects.

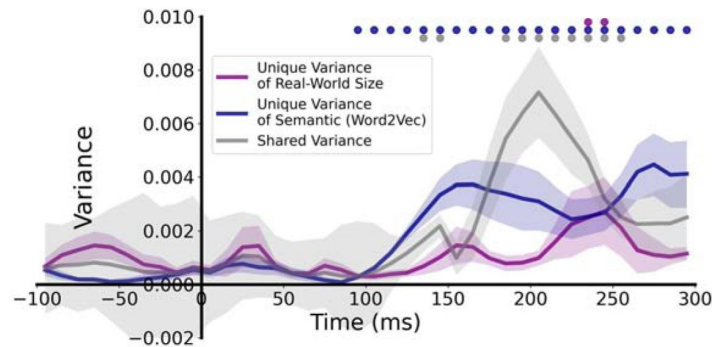


Figure S10.

RSA results based on CORnet.

(A) Representational similarities (partial Spearman correlations) between the five CORnet RDMs from multiple layers (from early layer to late layer: CORnet.V1, CORnet.V2, CORnet.V4, CORnet.IT, and CORnet.decoder.avgpool) and hypothesis-based RDMs of real-world depth, retinal size, and real-world size. Asterisks indicate significant partial correlations (bootstrap test, $p < .05$). (B) Representational similarity time courses (Spearman correlations) between EEG neural RDMs and the five CORnet RDMs. Color-coded small dots at the top indicate significant timepoints (cluster-based permutation test, $p < .05$). Shaded area reflects $\pm$ SEM across the 10 subjects.

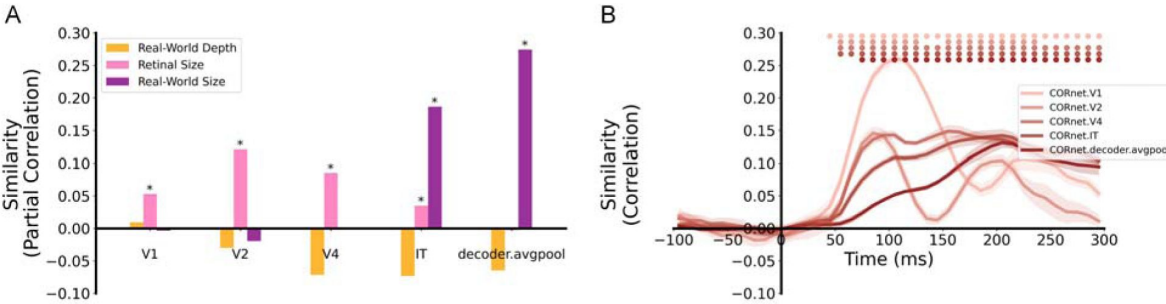


Table S1

Statistical results of similarities (partial Spearman correlations) between four ANN RDMs and three hypothesis-based RDMs.

ANN $\times$ HYP	<i>ResNet (early)</i>	<i>CLIP (early)</i>	<i>ResNet (late)</i>	<i>CLIP (late)</i>
<i>Real-World depth</i>	$r=.0330, p<.001$	$r=.0262, p<.001$	$r=-.0513, p=1$	$r=.0278, p<.001$
<i>Retinal Size</i>	$r=.0618, p<.001$	$r=.0730, p<.001$	$r=.0221, p<.001$	$r=.0058, p=.1788$
<i>Real-World Size</i>	$r=-.0330, p=1$	$r=-.0027, p=.8710$	$r=.02497, p<.001$	$r=.02378, p<.001$

Table S2

Statistical results of similarities (partial Spearman correlations) between four ANN RDMs with inputs of cropped object images without background and three hypothesis-based RDMs.

ANN (obj-only imgs) $\times$ HYP	<i>ResNet (early)</i>	<i>CLIP (early)</i>	<i>ResNet (late)</i>	<i>CLIP (late)</i>
<i>Real-World depth</i>	$r=-.0109, p=.9382$	$r=-.0086, p=.8862$	$r=.0183, p=.0049$	$r=-.0176, p=.9934$
<i>Retinal Size</i>	$r=.0323, p<.001$	$r=.0315, p<.001$	$r=-.0032, p=.6725$	$r=-.0031, p=.6705$
<i>Real-World Size</i>	$r=.0022, p=.3781$	$r=.0084, p=.1193$	$r=.0402, p<.001$	$r=.0154, p=.0149$

Additional information

Funding

National Institutes of Health (R01-EY025648)

National Science Foundation (NSF 1848939)

References

Allen E. J., St-Yves G., Wu Y., Breedlove J. L., Prince J. S., Dowdle L. T., Nau M., Caron B., Pestilli F., Charest I., Hutchinson J. B., Naselaris T., Kay K (2022) **A massive 7T fMRI dataset to bridge cognitive neuroscience and artificial intelligence** *Nature Neuroscience* **25**:116–126 [Google Scholar](#)

Ayzenberg V., Blauch N., Behrmann M. (2023) **Using deep neural networks to address the how of object recognition** *PsyArXiv* [Google Scholar](#)

Bao P., She L., McGill M., Tsao D. Y (2020) **A map of object space in primate inferotemporal cortex** *Nature* **583**:103–108 [Google Scholar](#)

Bracci S., Daniels N., Op de Beeck H. (2017) **Task Context Overrides Object- and Category-Related Representational Content in the Human Parietal Cortex** *Cerebral Cortex* **27**:310–321 [Google Scholar](#)

Bracci S., Op de Beeck H. (2016) **Dissociations and associations between shape and category representations in the two visual pathways** *Journal of Neuroscience* **36**:432–444 [Google Scholar](#)

Chen Y. C., Deza A., Konkle T (2022) **How big should this object be? Perceptual influences on viewing-size preferences** *Cognition* **225**:105114 [Google Scholar](#)

Choksi B., Mozafari M., VanRullen R., Reddy L (2022) **Multimodal neural networks better explain multivoxel patterns in the hippocampus** *Neural Networks* **154**:538–542 [Google Scholar](#)

Choksi B., Vanrullen R., Reddy L (2022) **Do multimodal neural networks better explain human visual representations than vision-only networks?** In: *Conference on Cognitive Computational Neuroscience 2022* [Google Scholar](#)

Cichy R. M., Kaiser D (2019) **Deep Neural Networks as Scientific Models** *Trends in Cognitive Sciences* **23**:305–317 [Google Scholar](#)

Cichy R. M., Khosla A., Pantazis D., Torralba A., Oliva A (2016) **Comparison of deep neural networks to spatio-temporal cortical dynamics of human visual object recognition reveals hierarchical correspondence** *Scientific Reports* **6**:1–13 [Google Scholar](#)

Coggan D. D., Tong F (2023) **Spikiness and animacy as potential organizing principles of human ventral visual cortex** *Cerebral Cortex* **33**:8194–8217 [Google Scholar](#)

- Conwell C., Prince J. S., Alvarez G. A., Konkle T. (2022) **Large-Scale Benchmarking of Diverse Artificial Vision Models in Prediction of 7T Human Neuroimaging Data** *BioRxiv* [Google Scholar](#)
- Conwell C., Prince J. S., Kay K. N., Alvarez G. A., Konkle T. (2024) **A large-scale examination of inductive biases shaping high-level visual representation in brains and machines** *Nature Communications* **15**:1–18 [Google Scholar](#)
- DiCarlo J. J., Zoccolan D., Rust N. C. (2012) **How does the brain solve visual object recognition?** *Neuron* **73**:415–434 [Google Scholar](#)
- Doerig A., Kietzmann T. C., Allen E., Wu Y., Naselaris T., Kay K., Charest I. (2022) **Semantic scene descriptions as an objective of human vision** *ArXiv* [Google Scholar](#)
- Doerig A., Sommers R., Seeliger K., Richards B., Ismael J., Lindsay G., Kording K., Konkle T., Van Gerven M. A. J., Kriegeskorte N., Kietzmann T. C. (2023) **The neuroconnectionist research programme** *Nature Reviews Neuroscience* **24**:431–450 [Google Scholar](#)
- Edelman S., Grill-Spector K., Kushnir T., Malach R. (1998) **Toward direct visualization of the internal shape representation space by fMRI** *Psychobiology* **26**:309–321 [Google Scholar](#)
- Gifford A. T., Dwivedi K., Roig G., Cichy R. M. (2022) **A large and rich EEG dataset for modeling human visual object recognition** *NeuroImage* **264**:119754 [Google Scholar](#)
- Golan T., Raju P. C., Kriegeskorte N. (2020) **Controversial stimuli: Pitting neural networks against each other as models of human cognition** *Proceedings of the National Academy of Sciences of the United States of America* **117**:29330–29337 [Google Scholar](#)
- Greene M. R., Hansen B. C. (2020) **Disentangling the Independent Contributions of Visual and Conceptual Features to the Spatiotemporal Dynamics of Scene Categorization** *Journal of Neuroscience* **40**:5283–5299 [Google Scholar](#)
- Grootswagers T., Zhou I., Robinson A. K., Hebart M. N., Carlson T. A. (2022) **Human EEG recordings for 1,854 concepts presented in rapid serial visual presentation streams** *Scientific Data* **9**:1–7 [Google Scholar](#)
- Güçlü U., van Gerven M. A. J. (2015) **Deep Neural Networks Reveal a Gradient in the Complexity of Neural Representations across the Ventral Stream** *Journal of Neuroscience* **35**:10005–10014 [Google Scholar](#)
- He K., Zhang X., Ren S., Sun J. (2016) **Deep Residual Learning for Image Recognition** In: *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR)* pp. 770–778 [Google Scholar](#)
- Hebart M. N., Dickter A. H., Kidder A., Kwok W. Y., Corriveau A., Van Wicklin C., Baker C. I. (2019) **THINGS: A database of 1,854 object concepts and more than 26,000 naturalistic object images** *PLoS ONE* **14**:1–24 [Google Scholar](#)
- Hochstein S., Ahissar M. (2002) **View from the top: Hierarchies and reverse hierarchies in the visual system** *Neuron* **36**:791–804 [Google Scholar](#)
- Huang T., Song Y., Liu J. (2022) **Real-world size of objects serves as an axis of object space** *Communications Biology* **5**:1–12 [Google Scholar](#)

Huth A. G., Nishimoto S., Vu A. T., Gallant J. L. (2012) **A Continuous Semantic Space Describes the Representation of Thousands of Object and Action Categories across the Human Brain** *Neuron* **76**:1210–1224 [Google Scholar](#)

Jozwik K. M., Kietzmann T. C., Cichy R. M., Kriegeskorte N., Mur M (2023) **Deep Neural Networks and Visuo-Semantic Models Explain Complementary Components of Human Ventral-Stream Representational Dynamics** *Journal of Neuroscience* **43**:1731–1741 [Google Scholar](#)

Kanwisher N., Khosla M., Dobs K (2023) **Using artificial neural networks to ask ‘why’ questions of minds and brains** *Trends in Neurosciences* **46**:240–254 [Google Scholar](#)

Kappenman E. S., Luck S. J. (2010) **The effects of electrode impedance on data quality and statistical significance in ERP recordings** *Psychophysiology* **47**:888–904 [Google Scholar](#)

Khaligh-Razavi S.-M., Cichy R. M., Pantazis D., Oliva A (2018) **Tracking the Spatiotemporal Neural Dynamics of Real-world Object Size and Animacy in the Human Brain** *Journal of Cognitive Neuroscience* **30**:1559–1576 [Google Scholar](#)

Khaligh-Razavi S. M., Kriegeskorte N (2014) **Deep Supervised, but Not Unsupervised, Models May Explain IT Cortical Representation** *PLoS Computational Biology* **10**:e1003915 [Google Scholar](#)

Kietzmann T. C., Spoerer C. J., Sörensen L. K. A., Cichy R. M., Hauk O., Kriegeskorte N (2019) **Recurrence is required to capture the representational dynamics of the human visual system** *Proceedings of the National Academy of Sciences of the United States of America* **116**:21854–21863 [Google Scholar](#)

Konkle T., Caramazza A (2013) **Tripartite Organization of the Ventral Stream by Animacy and Object Size** *Journal of Neuroscience* **33**:10235–10242 [Google Scholar](#)

Konkle T., Oliva A (2011) **Canonical Visual Size for Real-World Objects** *Journal of Experimental Psychology: Human Perception and Performance* **37**:23–37 [Google Scholar](#)

Konkle T., Oliva A (2012a) **A familiar-size Stroop effect: Real-world size is an automatic property of object representation** *Journal of Experimental Psychology: Human Perception and Performance* **38**:561 [Google Scholar](#)

Konkle T., Oliva A (2012b) **A Real-World Size Organization of Object Responses in Occipitotemporal Cortex** *Neuron* **74**:1114–1124 [Google Scholar](#)

Kriegeskorte N., Mur M., Ruff D. A., Kiani R., Bodurka J., Esteky H., Tanaka K., Bandettini P. A (2008) **Matching Categorical Object Representations in Inferior Temporal Cortex of Man and Monkey** *Neuron* **60**:1126–1141 [Google Scholar](#)

Kubilius J., Schrimpf M., Kar K., Rajalingham R., Hong H., Majaj N. J., Issa E. B., Bashivan P., Prescott-Roy J., Schmidt K., Nayeibi A., Bear D., Yamins D. L. K., Dicarlo J. J (2019) **Brain-Like Object Recognition with High-Performing Shallow Recurrent ANNs** *Advances in Neural Information Processing Systems (NeurIPS)* **32** [Google Scholar](#)

Kuzovkin I., Vicente R., Petton M., Lachaux J. P., Baciú M., Kahane P., Rheims S., Vidal J. R., Aru J (2018) **Activations of deep convolutional neural networks are aligned with gamma band activity of human visual cortex** *Communications Biology* **1**:1–12 [Google Scholar](#)

- Long B., Konkle T (2017) **A familiar-size Stroop effect in the absence of basic-level recognition** *Cognition* **168**:234–242 [Google Scholar](#)
- Long B., Konkle T., Cohen M. A., Alvarez G. A (2016) **Mid-level perceptual features distinguish objects of different real-world sizes** *Journal of Experimental Psychology: General* **145**:95–109 [Google Scholar](#)
- Long B., Yu C. P., Konkle T (2018) **Mid-level visual features underlie the high-level categorical organization of the ventral stream** *Proceedings of the National Academy of Sciences of the United States of America* **115**:E9015–E9024 [Google Scholar](#)
- Lu Z., Ku Y (2020) **NeuroRA: A Python Toolbox of Representational Analysis From Multi-Modal Neural Data** *Frontiers in Neuroinformatics* **14**:61 [Google Scholar](#)
- Lu Z., Ku Y (2023) **Bridging the gap between EEG and DCNNs reveals a fatigue mechanism of facial repetition suppression** *IScience* **26**:108501 [Google Scholar](#)
- Luck S. J (2014) **An Introduction to the Event-Related Potential Technique** In: *An introduction to the event-related potential technique* MIT Press [Google Scholar](#)
- Luo A. F., Wehbe L., Tarr M. J., Henderson M. M. (2023) **Neural Selectivity for Real-World Object Size In Natural Images** **Abbreviated title : Neural Selectivity for Real-World Size** *BioRxiv* [Google Scholar](#)
- Marr D (1982) **Vision : a computational investigation into the human representation and processing of visual information** The MIT Press [Google Scholar](#)
- Mikolov T., Chen K., Corrado G., Dean J (2013) **Efficient Estimation of Word Representations in Vector Space** In: *Proceedings of the International Conference on Learning Representations (ICLR)* [Google Scholar](#)
- Muttenthaler L., Hebart M. N (2021) **THINGSvision: A Python Toolbox for Streamlining the Extraction of Activations From Deep Neural Networks** *Frontiers in Neuroinformatics* **15**:45 [Google Scholar](#)
- Naselaris T., Stansbury D. E., Gallant J. L (2012) **Cortical representation of animate and inanimate objects in complex natural scenes** *Journal of Physiology Paris* **106**:239–249 [Google Scholar](#)
- Proklova D., Kaiser D., Peelen M. V (2016) **Disentangling Representations of Object Shape and Object Category in Human Visual Cortex: The Animate–Inanimate Distinction** *Journal of Cognitive Neuroscience* **28**:680–692 [Google Scholar](#)
- Quek G., Theodorou A., Peelen M. V. (2023) **Better together : Objects in familiar constellations evoke high-level representations of real-world size** *BioRxiv* [Google Scholar](#)
- Radford A., Kim J. W., Hallacy C., Ramesh A., Goh G., Agarwal S., Sastry G., Askell A., Mishkin P., Clark J., Krueger G., Sutskever I (2021) **Learning Transferable Visual Models From Natural Language Supervision** In: *Proceedings of the International Conference on Machine Learning (ICML)* [Google Scholar](#)
- Řehůřek R., Sojka P. (2010) **Software framework for topic modelling with large corpora** In: *Proceedings of the LREC 2010 Workshop on New Challenges for NLP Frameworks* pp. 45–50 [Google Scholar](#)

- Setti A., Caramelli N., Borghi A. M (2008) **Conceptual information about size of objects in nouns** *European Journal of Cognitive Psychology* **21**:1022–1044 [Google Scholar](#)
- Song Y., Liu B., Li X., Shi N., Wang Y., Gao X. (2024) **Decoding Natural Images from EEG for Object Recognition** In: *International Conference on Learning Representations (ICLR)* [Google Scholar](#)
- Stoinski L. M., Perkuhn J., Hebart M. N (2023) **THINGSplus: New norms and metadata for the THINGS database of 1854 object concepts and 26,107 natural object images** *Behavior Research Methods* :1–21 [Google Scholar](#)
- Troiani V., Stigliani A., Smith M. E., Epstein R. A (2014) **Multiple object properties drive scene-selective regions** *Cerebral Cortex* **24**:883–897 [Google Scholar](#)
- Wang A. Y., Kay K., Naselaris T., Tarr M. J., Wehbe L. (2022) **Incorporating natural language into vision models improves prediction and understanding of higher visual cortex** *BioRxiv* [Google Scholar](#)
- Wang R., Janini D., Konkle T (2022) **Mid-level feature differences support early animacy and object size distinctions: Evidence from electroencephalography decoding** *Journal of Cognitive Neuroscience* **34**:1670–1680 [Google Scholar](#)
- Xie S., Kaiser D., Cichy R. M (2020) **Visual Imagery and Perception Share Neural Representations in the Alpha Frequency Band** *Current Biology* **30**:2621–2627 [Google Scholar](#)
- Xu S., Zhang Y., Zhen Z., Liu J (2021) **The Face Module Emerged in a Deep Convolutional Neural Network Selectively Deprived of Face Experience** *Frontiers in Computational Neuroscience* **15**:1–12 [Google Scholar](#)
- Yamins D. L. K., DiCarlo J. J (2016) **Using goal-driven deep learning models to understand sensory cortex** *Nature Neuroscience* **19**:356–365 [Google Scholar](#)
- Yamins D. L. K., 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 of the United States of America* **111**:8619–8624 [Google Scholar](#)

Author information

Zitong Lu

Department of Psychology, The Ohio State University, Columbus, United States
ORCID iD: [0000-0002-7953-6742](#)

For correspondence: lu.2637@osu.edu

Julie D Golomb

Department of Psychology, The Ohio State University, Columbus, United States
ORCID iD: [0000-0003-3489-0702](#)

Editors

Reviewing Editor

Arun SP

Indian Institute of Science Bangalore, Bangalore, India

Senior Editor

Yanchao Bi

Peking University, Beijing, China

Reviewer #1 (Public review):

Lu & Golomb combined EEG, artificial neural networks, and multivariate pattern analyses to examine how different visual variables are processed in the brain. The conclusions of the paper are mostly well supported.

The authors find that not only real-world size is represented in the brain (which was known), but both retinal size and real-world depth is represented, at different time points or latencies, which may reflect different stages of processing. Prior work has not been able to answer the question of real-world depth due to stimuli used. The authors made this possible by assess real-world depth and testing it with appropriate methodology, accounting for retinal and real-world size. The methodological approach combining behavior, RSA, and ANNs is creative and well thought out to appropriately assess the research questions, and the findings may be very compelling if backed up with some clarifications and further analyses.

The work will be of interest to experimental and computational vision scientists, as well as the broader computational cognitive neuroscience community as the methodology is of interest and the code is or will be made available. The work is important as it is currently not clear what the correspondence between many deep neural network models are and the brain are, and this work pushes our knowledge forward on this front. Furthermore, the availability of methods and data will be useful for the scientific community.

<https://doi.org/10.7554/eLife.98117.2.sa2>

Reviewer #3 (Public review):

The authors used an open EEG dataset of observers viewing real-world objects. Each object had a real-world size value (from human rankings), a retinal size value (measured from each image), and a scene depth value (inferred from the above). The authors combined the EEG and object measurements with extant, pre-trained models (a deep convolutional neural network, a multimodal ANN, and Word2vec) to assess the time course of processing object size (retinal and real-world) and depth. They found that depth was processed first, followed by retinal size, and then real-world size. The depth time course roughly corresponded to the visual ANNs, while the real-world size time course roughly corresponded to the more semantic models.

The time course result for the three object attributes is very clear and a novel contribution to the literature. The authors have revised the ANN motivations to increase clarity. Additionally, the authors have appropriately toned down some of the language about novelty, and the addition of a noise ceiling has helped the robustness of the work.

While I appreciate the addition of Cornet in the Supplement, I am less compelled by the authors' argument for Word2Vec over LLMs for "pure" semantic embeddings. While I'm not digging in on this point, this choice may prematurely age this work.

Author response:

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

Reviewer #1 (Public Review):

Lu & Golomb combined EEG, artificial neural networks, and multivariate pattern analyses to examine how different visual variables are processed in the brain. The conclusions of the paper are mostly well supported, but some aspects of methods and data analysis would benefit from clarification and potential extensions.

The authors find that not only real-world size is represented in the brain (which was known), but both retinal size and real-world depth are represented, at different time points or latencies, which may reflect different stages of processing. Prior work has not been able to answer the question of real-world depth due to the stimuli used. The authors made this possible by assessing real-world depth and testing it with appropriate methodology, accounting for retinal and real-world size. The methodological approach combining behavior, RSA, and ANNs is creative and well thought out to appropriately assess the research questions, and the findings may be very compelling if backed up with some clarifications and further analyses.

The work will be of interest to experimental and computational vision scientists, as well as the broader computational cognitive neuroscience community as the methodology is of interest and the code is or will be made available. The work is important as it is currently not clear what the correspondence between many deep neural network models and the brain is, and this work pushes our knowledge forward on this front. Furthermore, the availability of methods and data will be useful for the scientific community.

Reviewer #2 (Public Review):

Summary:

This paper aims to test if neural representations of images of objects in the human brain contain a 'pure' dimension of real-world size that is independent of retinal size or perceived depth. To this end, they apply representational similarity analysis on EEG responses in 10 human subjects to a set of 200 images from a publicly available database (THINGS-EEG2), correlating pairwise distinctions in evoked activity between images with pairwise differences in human ratings of real-world size (from THINGS+). By partialling out correlations with metrics of retinal size and perceived depth from the resulting EEG correlation time courses, the paper claims to identify an independent representation of real-world size starting at 170 ms in the EEG signal. Further comparisons with artificial neural networks and language embeddings lead the authors to claim this correlation reflects a relatively 'high-level' and 'stable' neural representation.

Strengths:

The paper features insightful figures/illustrations and clear figures.

The limitations of prior work motivating the current study are clearly explained and seem reasonable (although the rationale for why using 'ecological' stimuli with backgrounds matters when studying real-world size could be made clearer; one could also argue the

opposite, that to get a 'pure' representation of the real-world size of an 'object concept', one should actually show objects in isolation).

The partial correlation analysis convincingly demonstrates how correlations between feature spaces can affect their correlations with EEG responses (and how taking into account these correlations can disentangle them better).

The RSA analysis and associated statistical methods appear solid.

Weaknesses:

The claim of methodological novelty is overblown. Comparing image metrics, behavioral measurements, and ANN activations against EEG using RSA is a commonly used approach to study neural object representations. The dataset size (200 test images from THINGS) is not particularly large, and neither is comparing pre-trained DNNs and language models, or using partial correlations.

Thanks for your feedback. We agree that the methods used in our study – such as RSA, partial correlations, and the use of pretrained ANN and language models – are indeed well-established in the literature. We therefore revised the manuscript to more carefully frame our contribution: rather than emphasizing methodological novelty in isolation, we now highlight the combination of techniques, the application to human EEG data with naturalistic images, and the explicit dissociation of real-world size, retinal size, and depth representations as the primary strengths of our approach. Corresponding language in the Abstract, Introduction, and Discussion has been adjusted to reflect this more precise positioning:

(Abstract, line 34 to 37) “our study combines human EEG and representational similarity analysis to disentangle neural representations of object real-world size from retinal size and perceived depth, leveraging recent datasets and modeling approaches to address challenges not fully resolved in previous work.”

(Introduction, line 104 to 106) “we overcome these challenges by combining human EEG recordings, naturalistic stimulus images, artificial neural networks, and computational modeling approaches including representational similarity analysis (RSA) and partial correlation analysis ...”

(Introduction, line 108) “We applied our integrated computational approach to an open EEG dataset...”

(Introduction, line 142 to 143) “The integrated computational approach by cross-modal representational comparisons we take with the current study...”

(Discussion, line 550 to 552) “our study goes beyond the contributions of prior studies in several key ways, offering both theoretical and methodological advances: ...”

The claims also seem too broad given the fairly small set of RDMs that are used here (3 size metrics, 4 ANN layers, 1 Word2Vec RDM): there are many aspects of object processing not studied here, so it's not correct to say this study provides a 'detailed and clear characterization of the object processing process'.

Thanks for pointing this out. We softened language in our manuscript to reflect that our findings provide a temporally resolved characterization of selected object features, rather than a comprehensive account of object processing:

(line 34 to 37) “our study combines human EEG and representational similarity analysis to disentangle neural representations of object real-world size from retinal size and perceived

depth, leveraging recent datasets and modeling approaches to address challenges not fully resolved in previous work.”

(line 46 to 48) “Our research provides a temporally resolved characterization of how certain key object properties – such as object real-world size, depth, and retinal size – are represented in the brain, ...”

The paper lacks an analysis demonstrating the validity of the real-world depth measure, which is here computed from the other two metrics by simply dividing them. The rationale and logic of this metric is not clearly explained. Is it intended to reflect the hypothesized egocentric distance to the object in the image if the person had in fact been 'inside' the image? How do we know this is valid? It would be helpful if the authors provided a validation of this metric.

We appreciate the comment regarding the real-world depth metric. Specifically, this metric was computed as the ratio of real-world size (obtained via behavioral ratings) to measured retinal size. The rationale behind this computation is grounded in the basic principles of perspective projection: for two objects subtending the same retinal size, the physically larger object is presumed to be farther away. This ratio thus serves as a proxy for perceived egocentric depth under the simplifying assumption of consistent viewing geometry across images.

We acknowledge that this is a derived estimate and not a direct measurement of perceived depth. While it provides a useful approximation that allows us to analytically dissociate the contributions of real-world size and depth in our RSA framework, we agree that future work would benefit from independent perceptual depth ratings to validate or refine this metric. We added more discussions about this to our revised manuscript:

(line 652 to 657) “Additionally, we acknowledge that our metric for real-world depth was derived indirectly as the ratio of perceived real-world size to retinal size. While this formulation is grounded in geometric principles of perspective projection and served the purpose of analytically dissociating depth from size in our RSA framework, it remains a proxy rather than a direct measure of perceived egocentric distance. Future work incorporating behavioral or psychophysical depth ratings would be valuable for validating and refining this metric.”

Given that there is only 1 image/concept here, the factor of real-world size may be confounded with other things, such as semantic category (e.g. buildings vs. tools). While the comparison of the real-world size metric appears to be effectively disentangled from retinal size and (the author's metric of) depth here, there are still many other object properties that are likely correlated with real-world size and therefore will confound identifying a 'pure' representation of real-world size in EEG. This could be addressed by adding more hypothesis RDMs reflecting different aspects of the images that may correlate with real-world size.

We thank the reviewer for this thoughtful and important point. We agree that semantic category and real-world size may be correlated, and that semantic structure is one of the plausible sources of variance contributing to real-world size representations. However, we would like to clarify that our original goal was to isolate real-world size from two key physical image features — retinal size and inferred real-world depth — which have been major confounds in prior work on this topic. We acknowledge that although our analysis disentangled real-world size from depth and retinal size, this does not imply a fully “pure” representation; therefore, we now refer to the real-world size representations as “partially disentangled” throughout the manuscript to reflect this nuance.

Interestingly, after controlling for these physical features, we still found a robust and statistically isolated representation of real-world size in the EEG signal. This motivated the idea that realworld size may be more than a purely perceptual or image-based property — it may be at least partially semantic. Supporting this interpretation, both the late layers of ANN models and the non-visual semantic model (Word2Vec) also captured real-world size structure. Rather than treating semantic information as an unwanted confound, we propose that semantic structure may be an inherent component of how the brain encodes real-world size.

To directly address the your concern, we conducted an additional variance partitioning analysis, in which we decomposed the variance in EEG RDMs explained by four RDMs: real-world depth, retinal size, real-world size, and semantic information (from Word2Vec). Specifically, for each EEG timepoint, we quantified (1) the unique variance of real-world size, after controlling for semantic similarity, depth, and retinal size; (2) the unique variance of semantic information, after controlling for real-world size, depth, and retinal size; (3) the shared variance jointly explained by real-world size and semantic similarity, controlling for depth and retinal size. This analysis revealed that real-world size explained unique variance in EEG even after accounting for semantic similarity. And there was also a substantial shared variance, indicating partial overlap between semantic structure and size. Semantic information also contributed unique explanatory power, as expected. These results suggest that real-world size is indeed partially semantic in nature, but also has independent neural representation not fully explained by general semantic similarity. This strengthens our conclusion that real-world size functions as a meaningful, higher-level dimension in object representation space.

We now include this new analysis and a corresponding figure (Figure S8) in the revised manuscript:

(line 532 to 539) “Second, we conducted a variance partitioning analysis, in which we decomposed the variance in EEG RDMs explained by three hypothesis-based RDMs and the semantic RDM (Word2Vec RDM), and we still found that real-world size explained unique variance in EEG even after accounting for semantic similarity (Figure S9). And we also observed a substantial shared variance jointly explained by real-world size and semantic similarity and a unique variance of semantic information. These results suggest that real-world size is indeed partially semantic in nature, but also has independent neural representation not fully explained by general semantic similarity.”

The choice of ANNs lacks a clear motivation. Why these two particular networks? Why pick only 2 somewhat arbitrary layers? If the goal is to identify more semantic representations using CLIP, the comparison between CLIP and vision-only ResNet should be done with models trained on the same training datasets (to exclude the effect of training dataset size & quality; cf Wang et al., 2023). This is necessary to substantiate the claims on page 19 which attributed the differences between models in terms of their EEG correlations to one of them being a 'visual model' vs. 'visual-semantic model'.

We agree that the choice and comparison of models should be better contextualized.

First, our motivation for selecting ResNet-50 and CLIP ResNet-50 was not to make a definitive comparison between model classes, but rather to include two widely used representatives of their respective categories—one trained purely on visual information (ResNet-50 on ImageNet) and one trained with joint visual and linguistic supervision (CLIP ResNet-50 on image–text pairs). These models are both highly influential and commonly used in computational and cognitive neuroscience, allowing for relevant comparisons with existing work (line 181-187).

Second, we recognize that limiting the EEG $\times$ ANN correlation analyses to only early and late layers may be viewed as insufficiently comprehensive. To address this point, we have computed the EEG correlations with multiple layers in both ResNet and CLIP models (ResNet: ResNet.maxpool, ResNet.layer1, ResNet.layer2, ResNet.layer3, ResNet.layer4, ResNet.avgpool; CLIP: CLIP.visual.avgpool, CLIP.visual.layer1, CLIP.visual.layer2, CLIP.visual.layer3, CLIP.visual.layer4, CLIP.visual.attnpool). The results, now included in Figure S4, show a consistent trend: early layers exhibit higher similarity to early EEG time points, and deeper layers show increased similarity to later EEG stages. We chose to highlight early and late layers in the main text to simplify interpretation.

Third, we appreciate the reviewer's point that differences in training datasets (ImageNet vs. CLIP's dataset) may confound any attribution of differences in brain alignment to the models' architectural or learning differences. We agree that the comparisons between models trained on matched datasets (e.g., vision-only vs. multimodal models trained on the same image-text corpus) would allow for more rigorous conclusions. Thus, we explicitly acknowledged this limitation in the text:

(line 443 to 445) "However, it is also possible that these differences between ResNet and CLIP reflect differences in training data scale and domain."

The first part of the claim on page 22 based on Figure 4 'The above results reveal that realworld size emerges with later peak neural latencies and in the later layers of ANNs, regardless of image background information' is not valid since no EEG results for images without backgrounds are shown (only ANNs).

We revised the sentence to clarify that this is a hypothesis based on the ANN results, not an empirical EEG finding:

(line 491 to 495) "These results show that real-world size emerges in the later layers of ANNs regardless of image background information, and – based on our prior EEG results – although we could not test object-only images in the EEG data, we hypothesize that a similar temporal profile would be observed in the brain, even for object-only images."

While we only had the EEG data of human subjects viewing naturalistic images, the ANN results suggest that real-world size representations may still emerge at later processing stages even in the absence of background, consistent with what we observed in EEG under with-background conditions.

The paper is likely to impact the field by showcasing how using partial correlations in RSA is useful, rather than providing conclusive evidence regarding neural representations of objects and their sizes.

Additional context important to consider when interpreting this work:

Page 20, the authors point out similarities of peak correlations between models ('Interestingly, the peaks of significant time windows for the EEG $\times$ HYP RSA also correspond with the peaks of the EEG $\times$ ANN RSA timecourse (Figure 3D,F)'. Although not explicitly stated, this seems to imply that they infer from this that the ANN-EEG correlation might be driven by their representation of the hypothesized feature spaces. However this does not follow: in EEG-image metric model comparisons it is very typical to see multiple peaks, for any type of model, this simply reflects specific time points in EEG at which visual inputs (images) yield distinctive EEG amplitudes (perhaps due to stereotypical waves of neural processing?), but one cannot infer the information being processed is the same. To investigate this, one could for example conduct variance partitioning or commonality analysis to see if there is variance at these specific

timepoints that is shared by a specific combination of the hypothesis and ANN feature spaces.

Thanks for your thoughtful observation! Upon reflection, we agree that the sentence – "Interestingly, the peaks of significant time windows for the EEG × HYP RSA also correspond with the peaks of the EEG × ANN RSA timecourse" – was speculative and risked implying a causal link that our data do not warrant. As you rightly points out, observing coincident peak latencies across different models does not necessarily imply shared representational content, given the stereotypical dynamics of evoked EEG responses. And we think even variance partitioning analysis would still not suffice to infer that ANN-EEG correlations are driven specifically by hypothesized feature spaces. Accordingly, we have removed this sentence from the manuscript to avoid overinterpretation.

Page 22 mentions 'The significant time-window (90-300ms) of similarity between Word2Vec RDM and EEG RDMs (Figure 5B) contained the significant time-window of EEG x real-world size representational similarity (Figure 3B)'. This is not particularly meaningful given that the Word2Vec correlation is significant for the entire EEG epoch (from the time-point of the signal 'arriving' in visual cortex around ~90 ms) and is thus much less temporally specific than the realworld size EEG correlation. Again a stronger test of whether Word2Vec indeed captures neural representations of real-world size could be to identify EEG time-points at which there are unique Word2Vec correlations that are not explained by either ResNet or CLIP, and see if those timepoints share variance with the real-world size hypothesized RDM.

We appreciate your insightful comment. Upon reflection, we agree that the sentence – "The significant time-window (90-300ms) of similarity between Word2Vec RDM and EEG RDMs (Figure 5B) contained the significant time-window of EEG x real-world size representational similarity (Figure 3B)" – was speculative. And we have removed this sentence from the manuscript to avoid overinterpretation.

Additionally, we conducted two analyses as you suggested in the supplement. First, we calculated the partial correlation between EEG RDMs and the Word2Vec RDM while controlling for four ANN RDMs (ResNet early/late and CLIP early/late) (Figure S8). Even after regressing out these ANN-derived features, we observed significant correlations between Word2Vec and EEG RDMs in the 100–190 ms and 250–300 ms time windows. This result suggests that

Word2Vec captures semantic structure in the neural signal that is not accounted for by ResNet or CLIP. Second, we conducted an additional variance partitioning analysis, in which we decomposed the variance in EEG RDMs explained by four RDMs: real-world depth, retinal size, real-world size, and semantic information (from Word2Vec) (Figure S9). And we found significant shared variance between Word2Vec and real-world size at 130–150 ms and 180–250 ms. These results indicate a partially overlapping representational structure between semantic content and real-world size in the brain.

We also added these in our revised manuscript:

(line 525 to 539) "To further probe the relationship between real-world size and semantic information, and to examine whether Word2Vec captures variances in EEG signals beyond that explained by visual models, we conducted two additional analyses. First, we performed a partial correlation between EEG RDMs and the Word2Vec RDM, while regressing out four ANN RDMs (early and late layers of both ResNet and CLIP) (Figure S8). We found that semantic similarity remained significantly correlated with EEG signals across sustained time windows (100-190ms and 250-300ms), indicating that Word2Vec captures neural variance not fully explained by visual or visual-language models. Second, we conducted a variance

partitioning analysis, in which we decomposed the variance in EEG RDMs explained by three hypothesis-based RDMs and the semantic RDM (Word2Vec RDM), and we still found that real-world size explained unique variance in EEG even after accounting for semantic similarity (Figure S9). And we also observed a substantial shared variance jointly explained by realworld size and semantic similarity and a unique variance of semantic information. These results suggest that real-world size is indeed partially semantic in nature, but also has independent neural representation not fully explained by general semantic similarity.”

Reviewer #3 (Public Review):

The authors used an open EEG dataset of observers viewing real-world objects. Each object had a real-world size value (from human rankings), a retinal size value (measured from each image), and a scene depth value (inferred from the above). The authors combined the EEG and object measurements with extant, pre-trained models (a deep convolutional neural network, a multimodal ANN, and Word2vec) to assess the time course of processing object size (retinal and real-world) and depth. They found that depth was processed first, followed by retinal size, and then real-world size. The depth time course roughly corresponded to the visual ANNs, while the real-world size time course roughly corresponded to the more semantic models.

The time course result for the three object attributes is very clear and a novel contribution to the literature. However, the motivations for the ANNs could be better developed, the manuscript could better link to existing theories and literature, and the ANN analysis could be modernized. I have some suggestions for improving specific methods.

(1) Manuscript motivations

The authors motivate the paper in several places by asking "whether biological and artificial systems represent object real-world size". This seems odd for a couple of reasons. Firstly, the brain must represent real-world size somehow, given that we can reason about this question. Second, given the large behavioral and fMRI literature on the topic, combined with the growing ANN literature, this seems like a foregone conclusion and undermines the novelty of this contribution.

Thanks for your helpful comment. We agree that asking whether the brain represents real-world size is not a novel question, given the existing behavioral and neuroimaging evidence supporting this. Our intended focus was not on the existence of real-world size representations per se, but the nature of these representations, particularly the relationship between the temporal dynamics and potential mechanisms of representations of real-world size versus other related perceptual properties (e.g., retinal size and real-world depth). We revised the relevant sentence to better reflect our focus, shifting from a binary framing (“whether or not size is represented”) to a more mechanistic and time-resolved inquiry (“how and when such representations emerge”):

(line 144 to 149) “Unraveling the internal representations of object size and depth features in both human brains and ANNs enables us to investigate how distinct spatial properties—retinal size, realworld depth, and real-world size—are encoded across systems, and to uncover the representational mechanisms and temporal dynamics through which real-world size emerges as a potentially higherlevel, semantically grounded feature.”

While the introduction further promises to "also investigate possible mechanisms of object realworld size representations.", I was left wishing for more in this department. The authors report correlations between neural activity and object attributes, as well as between neural activity and ANNs. It would be nice to link the results to theories of object processing (e.g., a feedforward sweep, such as DiCarlo and colleagues have suggested,

versus a reverse hierarchy, such as suggested by Hochstein, among others). What is semantic about real-world size, and where might this information come from? (Although you may have to expand beyond the posterior electrodes to do this analysis).

We thank the reviewer for this insightful comment. We agree that understanding the mechanisms underlying real-world size representations is a critical question. While our current study does not directly test specific theoretical frameworks such as the feedforward sweep model or the reverse hierarchy theory, our results do offer several relevant insights: The temporal dynamics revealed by EEG—where real-world size emerges later than retinal size and depth—suggest that such representations likely arise beyond early visual feedforward stages, potentially involving higher-level semantic processing. This interpretation is further supported by the fact that real-world size is strongly captured by late layers of ANNs and by a purely semantic model (Word2Vec), suggesting its dependence on learned conceptual knowledge.

While we acknowledge that our analyses were limited to posterior electrodes and thus cannot directly localize the cortical sources of these effects, we view this work as a first step toward bridging low-level perceptual features and higher-level semantic representations. We hope future work combining broader spatial sampling (e.g., anterior EEG sensors or source localization) and multimodal recordings (e.g., MEG, fMRI) can build on these findings to directly test competing models of object processing and representation hierarchy.

We also added these to the Discussion section:

(line 619 to 638) “Although our study does not directly test specific models of visual object processing, the observed temporal dynamics provide important constraints for theoretical interpretations. In particular, we find that real-world size representations emerge significantly later than low-level visual features such as retinal size and depth. This temporal profile is difficult to reconcile with a purely feedforward account of visual processing (e.g., DiCarlo et al., 2012), which posits that object properties are rapidly computed in a sequential hierarchy of increasingly complex visual features. Instead, our results are more consistent with frameworks that emphasize recurrent or top-down processing, such as the reverse hierarchy theory (Hochstein & Ahissar, 2002), which suggests that high-level conceptual information may emerge later and involve feedback to earlier visual areas. This interpretation is further supported by representational similarities with late-stage artificial neural network layers and with a semantic word embedding model (Word2Vec), both of which reflect learned, abstract knowledge rather than low-level visual features. Taken together, these findings suggest that real-world size is not merely a perceptual attribute, but one that draws on conceptual or semantic-level representations acquired through experience. While our EEG analyses focused on posterior electrodes and thus cannot definitively localize cortical sources, we see this study as a step toward linking low-level visual input with higher-level semantic knowledge. Future work incorporating broader spatial coverage (e.g., anterior sensors), source localization, or complementary modalities such as MEG and fMRI will be critical to adjudicate between alternative models of object representation and to more precisely trace the origin and flow of real-world size information in the brain.”

Finally, several places in the manuscript tout the “novel computational approach”. This seems odd because the computational framework and pipeline have been the most common approach in cognitive computational neuroscience in the past 5-10 years.

We have revised relevant statements throughout the manuscript to avoid overstating novelty and to better reflect the contribution of our study.

(2) Suggestion: modernize the approach

I was surprised that the computational models used in this manuscript were all 8-10 years old. Specifically, because there are now deep nets that more explicitly model the human brain (e.g., Cornet) as well as more sophisticated models of semantics (e.g., LLMs), I was left hoping that the authors had used more state-of-the-art models in the work. Moreover, the use of a single dCNN, a single multi-modal model, and a single word embedding model makes it difficult to generalize about visual, multimodal, and semantic features in general.

Thanks for your suggestion. Indeed, our choice of ResNet and CLIP was motivated by their widespread use in the cognitive and computational neuroscience area. These models have served as standard benchmarks in many studies exploring correspondence between ANNs and human brain activity. To address your concern, we have now added additional results from the more biologically inspired model, CORnet, in the supplementary (Figure S10). The results for CORnet show similar patterns to those observed for ResNet and CLIP, providing converging evidence across models.

Regarding semantic modeling, we intentionally chose Word2Vec rather than large language models (LLMs), because our goal was to examine concept-level, context-free semantic representations. Word2Vec remains the most widely adopted approach for obtaining noncontextualized embeddings that reflect core conceptual similarity, as opposed to the contextdependent embeddings produced by LLMs, which are less directly suited for capturing stable concept-level structure across stimuli.

(3) Methodological considerations

(a) Validity of the real-world size measurement

I was concerned about a few aspects of the real-world size rankings. First, I am trying to understand why the scale goes from 100-519. This seems very arbitrary; please clarify. Second, are we to assume that this scale is linear? Is this appropriate when real-world object size is best expressed on a log scale? Third, the authors provide "sand" as an example of the smallest realworld object. This is tricky because sand is more "stuff" than "thing", so I imagine it leaves observers wondering whether the experimenter intends a grain of sand or a sandy scene region. What is the variability in real-world size ratings? Might the variability also provide additional insights in this experiment?

We now clarify the origin, scaling, and interpretation of the real-world size values obtained from the THINGS+ dataset.

In their experiment, participants first rated the size of a single object concept (word shown on the screen) by clicking on a continuous slider of 520 units, which was anchored by nine familiar real-world reference objects (e.g., “grain of sand,” “microwave oven,” “aircraft carrier”) that spanned the full expected size range on a logarithmic scale. Importantly, participants were not shown any numerical values on the scale—they were guided purely by the semantic meaning and relative size of the anchor objects. After the initial response, the scale zoomed in around the selected region (covering 160 units of the 520-point scale) and presented finer anchor points between the previous reference objects. Participants then refined their rating by dragging from the lower to upper end of the typical size range for that object. If the object was standardized in size (e.g., “soccer ball”), a single click sufficed. These size judgments were collected across at least 50 participants per object, and final scores were derived from the central tendency of these responses. Although the final size values numerically range from 0 to 519 (after scaling), this range is not known to participants and is only applied post hoc to construct the size RDMs.

Regarding the term “sand”: the THINGS+ dataset distinguished between object meanings when ambiguity was present. For “sand,” participants were instructed to treat it as “a grain of sand”—consistent with the intended meaning of a discrete, minimal-size reference object.

Finally, we acknowledge that real-world size ratings may carry some degree of variability across individuals. However, the dataset includes ratings from 2010 participants across 1854 object concepts, with each object receiving at least 50 independent ratings. Given this large and diverse sample, the mean size estimates are expected to be stable and robust across subjects. While we did not include variability metrics in our main analysis, we believe the aggregated ratings provide a reliable estimate of perceived real-world size.

We added these details in the Materials and Method section:

(line 219 to 230) “In the THINGS+ dataset, 2010 participants (different from the subjects in THINGS EEG2) did an online size rating task and completed a total of 13024 trials corresponding to 1854 object concepts using a two-step procedure. In their experiment, first, each object was rated on a 520-unit continuous slider anchored by familiar reference objects (e.g., “grain of sand,” “microwave oven,” “aircraft carrier”) representing a logarithmic size range. Participants were not shown numerical values but used semantic anchors as guides. In the second step, the scale zoomed in around the selected region to allow for finer-grained refinement of the size judgment. Final size values were derived from aggregated behavioral data and rescaled to a range of 0–519 for consistency across objects, with the actual mean ratings across subjects ranging from 100.03 (‘grain of sand’) to 423.09 (‘subway’).”

(b) This work has no noise ceiling to establish how strong the model fits are, relative to the intrinsic noise of the data. I strongly suggest that these are included.

We have now computed noise ceiling estimates for the EEG RDMs across time. The noise ceiling was calculated by correlating each participant’s EEG RDM with the average EEG RDM across the remaining participants (leave-one-subject-out), at each time point. This provides an upper-bound estimate of the explainable variance, reflecting the maximum similarity that any model—no matter how complex—could potentially achieve, given the intrinsic variability in the EEG data.

Importantly, the observed EEG–model similarity values are substantially below this upper bound. This outcome is fully expected: Each of our model RDMs (e.g., real-world size, ANN layers) captures only a specific aspect of the neural representational structure, rather than attempting to account for the totality of the EEG signal. Our goal is not to optimize model performance or maximize fit, but to probe which components of object information are reflected in the spatiotemporal dynamics of the brain’s responses.

For clarity and accessibility of the main findings, we present the noise ceiling time courses separately in the supplementary materials (Figure S7). Including them directly in the EEG × HYP or EEG × ANN plots would conflate distinct interpretive goals: the model RDMs are hypothesis-driven probes of specific representational content, whereas the noise ceiling offers a normative upper bound for total explainable variance. Keeping these separate ensures each visualization remains focused and interpretable.

Reviewer #1 (Recommendations For The Authors)::

Some analyses are incomplete, which would be improved if the authors showed analyses with other layers of the networks and various additional partial correlation analyses.

Clarity

(1) Partial correlations methods incomplete - it is not clear what is being partialled out in each analysis. It is possible to guess sometimes, but it is not entirely clear for each analysis. This is important as it is difficult to assess if the partial correlations are sensible/correct in each case. Also, the Figure 1 caption is short and unclear.

For example, ANN-EEG partial correlations - "Finally, we directly compared the timepoint-bytimepoint EEG neural RDMs and the ANN RDMs (Figure 3F). The early layer representations of both ResNet and CLIP were significantly correlated with early representations in the human brain" What is being partialled out? Figure 3F says partial correlation

We apologize for the confusion. We made several key clarifications and corrections in the revised version.

First, we identified and corrected a labeling error in both Figure 1 and Figure 3F. Specifically, our EEG $\times$ ANN analysis used Spearman correlation, not partial correlation as mistakenly indicated in the original figure label and text. We conducted partial correlations for EEG $\times$ HYP and ANN $\times$ HYP. But for EEG $\times$ ANN, we directly calculated the correlation between EEG RDMs and ANN RDM corresponding to different layers respectively. We corrected these errors: (1) In Figure 1, we removed the erroneous “partial” label from the EEG $\times$ ANN path and updated the caption to clearly outline which comparisons used partial correlation. (2) In Figure 3F, we corrected the Y-axis label to “(correlation)”.

Second, to improve clarity, we have now revised the Materials and Methods section to explicitly describe what is partialled out in each partial correlation analysis:

(line 284 to 286) “In EEG $\times$ HYP partial correlation (Figure 3D), we correlated EEG RDMs with one hypothesis-based RDM (e.g., real-world size), while controlling for the other two (retinal size and real-world depth).”

(line 303 to 305) “In ANN (or W2V) $\times$ HYP partial correlation (Figure 3E and Figure 5A), we correlated ANN (or W2V) RDMs with one hypothesis-based RDM (e.g., real-world size), while partialling out the other two.”

Finally, the caption of Figure 1 has been expanded to clarify the full analysis pipeline and explicitly specify the partial correlation or correlation in each comparison.

(line 327 to 332) “Figure 1 Overview of our analysis pipeline including constructing three types of RDMs and conducting comparisons between them. We computed RDMs from three sources: neural data (EEG), hypothesized object features (real-world size, retinal size, and real-world depth), and artificial models (ResNet, CLIP, and Word2Vec). Then we conducted cross-modal representational similarity analyses between: EEG $\times$ HYP (partial correlation, controlling for other two HYP features), ANN (or W2V) $\times$ HYP (partial correlation, controlling for other two HYP features), and EEG $\times$ ANN (correlation).”

We believe these revisions now make all analytic comparisons and correlation types full clear and interpretable.

Issues / open questions

(2) Semantic representations vs hypothesized (hyp) RDMs (real-world size, etc) - are the representations explained by variables in hyp RDMs or are there semantic representations over and above these? E.g., For ANN correlation with the brain, you could partial out hyp RDMs - and assess whether there is still semantic information left over, or is the variance explained by the hyp RDMs?

Thank for this suggestion. As you suggested, we conducted the partial correlation analysis between EEG RDMs and ANN RDMs, controlling for the three hypothesis-based RDMs. The results (Figure S6) revealed that the EEG×ANN representational similarity remained largely unchanged, indicating that ANN representations capture much more additional representational structure not accounted for by the current hypothesized features. This is also consistent with the observation that EEG×HYP partial correlations were themselves small, but EEG×ANN correlations were much greater.

We also added this statement to the main text:

(line 446 to 451) “To contextualize how much of the shared variance between EEG and ANN representations is driven by the specific visual object features we tested above, we conducted a partial correlation analysis between EEG RDMs and ANN RDMs controlling for the three hypothesis-based RDMs (Figure S6). The EEG×ANN similarity results remained largely unchanged, suggesting that ANN representations capture much more additional rich representational structure beyond these features.”

(3) Why only early and late layers? I can see how it's clearer to present the EEG results. However, the many layers in these networks are an opportunity - we can see how simple/complex linear/non-linear the transformation is over layers in these models. It would be very interesting and informative to see if the correlations do in fact linearly increase from early to later layers, or if the story is a bit more complex. If not in the main text, then at least in the supplement.

Thank you for the thoughtful suggestion. To address this point, we have computed the EEG correlations with multiple layers in both ResNet and CLIP models (ResNet: ResNet.maxpool, ResNet.layer1, ResNet.layer2, ResNet.layer3, ResNet.layer4, ResNet.avgpool; CLIP: CLIP.visual.avgpool, CLIP.visual.layer1, CLIP.visual.layer2, CLIP.visual.layer3, CLIP.visual.layer4, CLIP.visual.attnpool). The results, now included in Figure S4 and S5, show a consistent trend: early layers exhibit higher similarity to early EEG time points, and deeper layers show increased similarity to later EEG stages. We chose to highlight early and late layers in the main text to simplify interpretation, but now provide the full layerwise profile for completeness.

(4) Peak latency analysis - Estimating peaks per ppt is presumably noisy, so it seems important to show how reliable this is. One option is to find the bootstrapped mean latencies per subject.

Thanks for your suggestion. To estimate the robustness of peak latency values, we implemented a bootstrap procedure by resampling the pairwise entries of the EEG RDM with replacement. For each bootstrap sample, we computed a new EEG RDM and recalculated the partial correlation time course with the hypothesis RDMs. We then extracted the peak latency within the predefined significant time window. Repeating this process 1000 times allowed us to get the bootstrapped mean latencies per subject as the more stable peak latency result. Notably, the bootstrapped results showed minimal deviation from the original latency estimates, confirming the robustness of our findings. Accordingly, we updated the Figure 3D and added these in the Materials and Methods section:

(line 289 to 298) “To assess the stability of peak latency estimates for each subject, we performed a bootstrap procedure across stimulus pairs. At each time point, the EEG RDM was vectorized by extracting the lower triangle (excluding the diagonal), resulting in 19,900 unique pairwise values. For each bootstrap sample, we resampled these 19,900 pairwise entries with replacement to generate a new pseudo-RDM of the same size. We then computed the partial correlation between the EEG pseudo-RDM and a given hypothesis RDM (e.g., real-

world size), controlling for other feature RDMs, and obtained a time course of partial correlations. Repeating this procedure 1000 times and extracting the peak latency within the significant time window yielded a distribution of bootstrapped latencies, from which we got the bootstrapped mean latencies per subject.”

(5) *"Due to our calculations being at the object level, if there were more than one of the same objects in an image, we cropped the most complete one to get a more accurate retinal size. " Did EEG experimenters make sure everyone sat the same distance from the screen? and remain the same distance? This would also affect real-world depth measures.*

Yes, the EEG dataset we used (THINGS EEG2; Gifford et al., 2022) was collected under carefully controlled experimental conditions. We have confirmed that all participants were seated at a fixed distance of 0.6 meters from the screen throughout the experiment. We also added this information in the method (line 156 to 157).

Minor issues/questions - note that these are not raised in the Public Review

(6) *Title - less about rigor/quality of the work but I feel like the title could be improved/extended. The work tells us not only about real object size, but also retinal size and depth. In fact, isn't the most novel part of this the real-world depth aspect? Furthermore, it feels like the current title restricts its relevance and impact... Also doesn't touch on the temporal aspect, or processing stages, which is also very interesting. There may be something better, but simply adding something like"...disentangled features of real-world size, depth, and retinal size over time OR processing stages".*

Thanks for your suggestion! We changed our title – “Human EEG and artificial neural networks reveal disentangled representations and processing timelines of object real-world size and depth in natural images”.

(7) *"Each subject viewed 16740 images of objects on a natural background for 1854 object concepts from the THINGS dataset (Hebart et al., 2019). For the current study, we used the 'test' dataset portion, which includes 16000 trials per subject corresponding to 200 images." Why test images? Worth explaining.*

We chose to use the “test set” of the THINGS EEG2 dataset for the following two reasons:

(1) Higher trial count per condition: In the test set, each of the 200 object images was presented 80 times per subject, whereas in the training set, each image was shown only 4 times. This much higher trial count per condition in the test set allows for substantially higher signal-to-noise ratio in the EEG data.

(2) Improved decoding reliability: Our analysis relies on constructing EEG RDMs based on pairwise decoding accuracy using linear SVM classifiers. Reliable decoding estimates require a sufficient number of trials per condition. The test set design is thus better suited to support high-fidelity decoding and robust representational similarity analysis.

We also added these explanations to our revised manuscript (line 161 to 164).

(8) *"For Real-World Size RDM, we obtained human behavioral real-world size ratings of each object concept from the THINGS+ dataset (Stoinski et al., 2022).... The range of possible size ratings was from 0 to 519 in their online size rating task..." How were the ratings made? What is this scale - do people know the numbers? Was it on a continuous slider?*

We should clarify how the real-world size values were obtained from the THINGS+ dataset.

In their experiment, participants first rated the size of a single object concept (word shown on the screen) by clicking on a continuous slider of 520 units, which was anchored by nine familiar real-world reference objects (e.g., “grain of sand,” “microwave oven,” “aircraft carrier”) that spanned the full expected size range on a logarithmic scale. Importantly, participants were not shown any numerical values on the scale—they were guided purely by the semantic meaning and relative size of the anchor objects. After the initial response, the scale zoomed in around the selected region (covering 160 units of the 520-point scale) and presented finer anchor points between the previous reference objects. Participants then refined their rating by dragging from the lower to upper end of the typical size range for that object. If the object was standardized in size (e.g., “soccer ball”), a single click sufficed. These size judgments were collected across at least 50 participants per object, and final scores were derived from the central tendency of these responses. Although the final size values numerically range from 0 to 519 (after scaling), this range is not known to participants and is only applied post hoc to construct the size RDMs.

We added these details in the Materials and Method section:

(line 219 to 230) “In the THINGS+ dataset, 2010 participants (different from the subjects in THINGS EEG2) did an online size rating task and completed a total of 13024 trials corresponding to 1854 object concepts using a two-step procedure. In their experiment, first, each object was rated on a 520unit continuous slider anchored by familiar reference objects (e.g., “grain of sand,” “microwave oven,” “aircraft carrier”) representing a logarithmic size range. Participants were not shown numerical values but used semantic anchors as guides. In the second step, the scale zoomed in around the selected region to allow for finer-grained refinement of the size judgment. Final size values were derived from aggregated behavioral data and rescaled to a range of 0–519 for consistency across objects, with the actual mean ratings across subjects ranging from 100.03 (‘grain of sand’) to 423.09 (‘subway’).”

(9) *"For Retinal Size RDM, we applied Adobe Photoshop (Adobe Inc., 2019) to crop objects corresponding to object labels from images manually... " Was this by one person? Worth noting, and worth sharing these values per image if not already for other researchers as it could be a valuable resource (and increase citations).*

Yes, all object cropping were performed consistently by one of the authors to ensure uniformity across images. We agree that this dataset could be a useful resource to the community. We have now made the cropped object images publicly available <https://github.com/ZitongLu1996/RWsize>.

We also updated the manuscript accordingly to note this (line 236 to 239).

(10) *"Neural RDMs. From the EEG signal, we constructed timepoint-by-timepoint neural RDMs for each subject with decoding accuracy as the dissimilarity index " Decoding accuracy is presumably a similarity index. Maybe 1-accuracy (proportion correct) for dissimilarity?*

Decoding accuracy is a dissimilarity index instead of a similarity index, as higher decoding accuracy between two conditions indicates that they are more distinguishable – i.e., less similar – in the neural response space. This approach aligns with prior work using classification-based representational dissimilarity measures (Grootswagers et al., 2017; Xie et al., 2020), where better decoding implies greater dissimilarity between conditions. Therefore, there is no need to invert the decoding accuracy values (e.g., using 1 - accuracy).

Grootswagers, T., Wardle, S. G., & Carlson, T. A. (2017). Decoding dynamic brain patterns from evoked responses: A tutorial on multivariate pattern analysis applied to time series neuroimaging data. *Journal of Cognitive Neuroscience*, 29(4), 677-697.

Xie, S., Kaiser, D., & Cichy, R. M. (2020). Visual imagery and perception share neural representations in the alpha frequency band. *Current Biology*, 30(13), 2621-2627.

(11) *Figure 1 caption is very short - Could do with a more complete caption. Unclear what the partial correlations are (what is being partialled out in each case), what are the comparisons "between them" - both in the figure and the caption. Details should at least be in the main text.*

Related to your comment (1). We revised the caption and the corresponding text.

Reviewer #2 (Recommendations For The Authors):

(1) *Intro:*

Quek et al., (2023) is referred to as a behavioral study, but it has EEG analyses.

We corrected this – “..., one recent study (Quek et al., 2023) ...”

The phrase 'high temporal resolution EEG' is a bit strange - isn't all EEG high temporal resolution? Especially when down-sampling to 100 Hz (40 time points/epoch) this does not qualify as particularly high-res.

We removed this phrasing in our manuscript.

(2) *Methods:*

It would be good to provide more details on the EEG preprocessing. Were the data low-pass filtered, for example?

We added more details to the manuscript:

(line 167 to 174) “The EEG data were originally sampled at 1000Hz and online-filtered between 0.1 Hz and 100 Hz during acquisition, with recordings referenced to the Fz electrode. For preprocessing, no additional filtering was applied. Baseline correction was performed by subtracting the mean signal during the 100 ms pre-stimulus interval from each trial and channel separately. We used already preprocessed data from 17 channels with labels beginning with “O” or “P” (O1, Oz, O2, PO7, PO3, POz, PO4, PO8, P7, P5, P3, P1, Pz, P2) ensuring full coverage of posterior regions typically involved in visual object processing. The epoched data were then down-sampled to 100 Hz.”

It is important to provide more motivation about the specific ANN layers chosen. Were these layers cherry-picked, or did they truly represent a gradual shift over the course of layers?

We appreciate the reviewer’s concern and fully agree that it is important to ensure transparency in how ANN layers were selected. The early and late layers reported in the main text were not cherry-picked to maximize effects, but rather intended to serve as illustrative examples representing the lower and higher ends of the network hierarchy. To address this point directly, we have computed the EEG correlations with multiple layers in both ResNet and CLIP models (ResNet: ResNet.maxpool, ResNet.layer1, ResNet.layer2, ResNet.layer3, ResNet.layer4, ResNet.avgpool; CLIP: CLIP.visual.avgpool, CLIP.visual.layer1,

CLIP.visual.layer2, CLIP.visual.layer3, CLIP.visual.layer4, CLIP.visual.attnpool). The results, now included in Figure S4, show a consistent trend: early layers exhibit higher similarity to early EEG time points, and deeper layers show increased similarity to later EEG stages.

It is important to provide more specific information about the specific ANN layers chosen. 'Second convolutional layer': is this block 2, the ReLu layer, the maxpool layer? What is the 'last visual layer'?

Apologize for the confusing! We added more details about the layer chosen:

(line 255 to 257) “The early layer in ResNet refers to ResNet.maxpool layer, and the late layer in ResNet refers to ResNet.avgpool layer. The early layer in CLIP refers to CLIP.visual.avgpool layer, and the late layer in CLIP refers to CLIP.visual.attnpool layer.”

Again the claim 'novel' is a bit overblown here since the real-world size ratings were also already collected as part of THINGS+, so all data used here is available.

We removed this phrasing in our manuscript.

Real-world size ratings ranged 'from 0 - 519'; it seems unlikely this was the actual scale presented to subjects, I assume it was some sort of slider?

You are correct. We should clarify how the real-world size values were obtained from the THINGS+ dataset.

In their experiment, participants first rated the size of a single object concept (word shown on the screen) by clicking on a continuous slider of 520 units, which was anchored by nine familiar real-world reference objects (e.g., “grain of sand,” “microwave oven,” “aircraft carrier”) that spanned the full expected size range on a logarithmic scale. Importantly, participants were not shown any numerical values on the scale—they were guided purely by the semantic meaning and relative size of the anchor objects. After the initial response, the scale zoomed in around the selected region (covering 160 units of the 520-point scale) and presented finer anchor points between the previous reference objects. Participants then refined their rating by dragging from the lower to upper end of the typical size range for that object. If the object was standardized in size (e.g., “soccer ball”), a single click sufficed. These size judgments were collected across at least 50 participants per object, and final scores were derived from the central tendency of these responses. Although the final size values numerically range from 0 to 519 (after scaling), this range is not known to participants and is only applied post hoc to construct the size RDMS.

We added these details in the Materials and Method section:

(line 219 to 230) “In the THINGS+ dataset, 2010 participants (different from the subjects in THINGS EEG2) did an online size rating task and completed a total of 13024 trials corresponding to 1854 object concepts using a two-step procedure. In their experiment, first, each object was rated on a 520unit continuous slider anchored by familiar reference objects (e.g., “grain of sand,” “microwave oven,” “aircraft carrier”) representing a logarithmic size range. Participants were not shown numerical values but used semantic anchors as guides. In the second step, the scale zoomed in around the selected region to allow for finer-grained refinement of the size judgment. Final size values were derived from aggregated behavioral data and rescaled to a range of 0–519 for consistency across objects, with the actual mean ratings across subjects ranging from 100.03 (‘grain of sand’) to 423.09 (‘subway’).”

Why is conducting a one-tailed ($p < 0.05$) test valid for EEG-ANN comparisons? Shouldn't this be two-tailed?

Our use of one-tailed tests was based on the directional hypothesis that representational similarity between EEG and ANN RDMs would be positive, as supported by prior literature showing correspondence between hierarchical neural networks and human brain representations (e.g., Cichy et al., 2016; Kuzovkin et al., 2014). This is consistent with a large number of RSA studies which conduct one-tailed tests (i.e., testing the hypothesis that coefficients were greater than zero: e.g., Kuzovkin et al., 2018; Nili et al., 2014; Hebart et al., 2018; Kaiser et al., 2019; Kaiser et al., 2020; Kaiser et al., 2022). Thus, we specifically tested whether the similarity was significantly greater than zero.

Cichy, R. M., Khosla, A., Pantazis, D., Torralba, A., & Oliva, A. (2016). Comparison of deep neural networks to spatio-temporal cortical dynamics of human visual object recognition reveals hierarchical correspondence. *Scientific reports*, 6(1), 27755.

Kuzovkin, I., Vicente, R., Petton, M., Lachaux, J. P., Baci, M., Kahane, P., ... & Aru, J. (2018). Activations of deep convolutional neural networks are aligned with gamma band activity of human visual cortex. *Communications biology*, 1(1), 107.

Nili, H., Wingfield, C., Walther, A., Su, L., Marslen-Wilson, W., & Kriegeskorte, N. (2014). A toolbox for representational similarity analysis. *PLoS computational biology*, 10(4), e1003553.

Hebart, M. N., Bankson, B. B., Harel, A., Baker, C. I., & Cichy, R. M. (2018). The representational dynamics of task and object processing in humans. *Elife*, 7, e32816.

Kaiser, D., Turini, J., & Cichy, R. M. (2019). A neural mechanism for contextualizing fragmented inputs during naturalistic vision. *elife*, 8, e48182.

Kaiser, D., Inciuraite, G., & Cichy, R. M. (2020). Rapid contextualization of fragmented scene information in the human visual system. *Neuroimage*, 219, 117045.

Kaiser, D., Jacobs, A. M., & Cichy, R. M. (2022). Modelling brain representations of abstract concepts. *PLoS Computational Biology*, 18(2), e1009837.

Importantly, we note that using a two-tailed test instead would not change the significance of our results. However, we believe the one-tailed test remains more appropriate given our theoretical prediction of positive similarity between ANN and brain representations.

The sentence on the partial correlation description (page 11 'we calculated partial correlations with one-tailed test against the alternative hypothesis that the partial correlation was positive (greater than zero)') didn't make sense to me; are you referring to the null hypothesis here?

We revised this sentence to clarify that we tested against the null hypothesis that the partial correlation was less than or equal to zero, using a one-tailed test to assess whether the correlation was significantly greater than zero.

(line 281 to 284) "..., we calculated partial correlations and used a one-tailed test against the null hypothesis that the partial correlation was less than or equal to zero, testing whether the partial correlation was significantly greater than zero."

(3) Results:

I would prevent the use of the word 'pure', your measurement is one specific operationalization of this concept of real-world size that is not guaranteed to result in unconfounded representations. This is in fact impossible whenever one is using a finite set of natural stimuli and calculating metrics on those - there can always be a factor or metric that was not considered that could explain some of the variance in your

measurement. It is overconfident to claim to have achieved some form of Platonic ideal here and to have taken into account all confounds.

Your point is well taken. Our original use of the term “pure” was intended to reflect statistical control for known confounding factors, but we recognize that this wording may imply a stronger claim than warranted. In response, we revised all relevant language in the manuscript to instead describe the statistically isolated or relatively unconfounded representation of real-world size, clarifying that our findings pertain to the unique contribution of real-world size after accounting for retinal size and real-world depth.

Figure 2C: It's not clear why peak latencies are computed on the 'full' correlations rather than the partial ones.

No. The peak latency results in Figure 2C were computed on the partial correlation results – we mentioned this in the figure caption – “Temporal latencies for peak similarity (partial Spearman correlations) between EEG and the 3 types of object information.”

SEM = SEM across the 10 subjects?

Yes. We added this in the figure caption.

Figure 3F y-axis says it's partial correlations but not clear what is partialled out here.

We identified and corrected a labeling error in both Figure 1 and Figure 3F. Specifically, our EEG × ANN analysis used Spearman correlation, not partial correlation as mistakenly indicated in the original figure label and text. We conducted partial correlations for EEG × HYP and ANN × HYP. But for EEG × ANN, we directly calculated the correlation between EEG RDMs and ANN RDM corresponding to different layers respectively. We corrected these errors: (1) In Figure 1, we removed the erroneous “partial” label from the EEG × ANN path and updated the caption to clearly outline which comparisons used partial correlation. (2) In Figure 3F, we corrected the Y-axis label to “(correlation)”.

Reviewer #3 (Recommendations For The Authors):

(1) Several methodologies should be clarified:

(a) It's stated that EEG was sampled at 100 Hz. I assume this was downsampled? From what original frequency?

Yes. We added more detailed about EEG data:

(line 167 to 174) “The EEG data were originally sampled at 1000Hz and online-filtered between 0.1 Hz and 100 Hz during acquisition, with recordings referenced to the Fz electrode. For preprocessing, no additional filtering was applied. Baseline correction was performed by subtracting the mean signal during the 100 ms pre-stimulus interval from each trial and channel separately. We used already preprocessed data from 17 channels with labels beginning with “O” or “P” (O1, Oz, O2, PO7, PO3, POz, PO4, PO8, P7, P5, P3, P1, Pz, P2) ensuring full coverage of posterior regions typically involved in visual object processing. The epoched data were then down-sampled to 100 Hz.”

(b) Why was decoding accuracy used as the human RDM method rather than the EEG data themselves?

Thanks for your question! We would like to address why we used decoding accuracy for EEG RDMs rather than correlation. While fMRI RDMs are typically calculated using 1 minus

correlation coefficient, decoding accuracy is more commonly used for EEG RDMs (Grootswager et al., 2017; Xie et al., 2020). The primary reason is that EEG signals are more susceptible to noise than fMRI data. Correlation-based methods are particularly sensitive to noise and may not reliably capture the functional differences between EEG patterns for different conditions. Decoding accuracy, by training classifiers to focus on task-relevant features, can effectively mitigate the impact of noisy signals and capture the representational difference between two conditions.

Grootswagers, T., Wardle, S. G., & Carlson, T. A. (2017). Decoding dynamic brain patterns from evoked responses: A tutorial on multivariate pattern analysis applied to time series neuroimaging data. *Journal of Cognitive Neuroscience*, 29(4), 677-697.

Xie, S., Kaiser, D., & Cichy, R. M. (2020). Visual imagery and perception share neural representations in the alpha frequency band. *Current Biology*, 30(13), 2621-2627.

We added this explanation to the manuscript:

(line 204 to 209) “Since EEG has a low SNR and includes rapid transient artifacts, Pearson correlations computed over very short time windows yield unstable dissimilarity estimates (Kappenman & Luck, 2010; Luck, 2014) and may thus fail to reliably detect differences between images. In contrast, decoding accuracy - by training classifiers to focus on task-relevant features - better mitigates noise and highlights representational differences.”

| (c) *How were the specific posterior electrodes selected?*

The 17 posterior electrodes used in our analyses were pre-selected and provided in the THINGS EEG2 dataset, and corresponding to standard occipital and parietal sites based on the 10-10 EEG system. Specifically, we included all 17 electrodes with labels beginning with “O” or “P”, ensuring full coverage of posterior regions typically involved in visual object processing (Page 7).

| (d) *The specific layers should be named rather than the vague (“last visual”)*

Apologize for the confusing! We added more details about the layer information:

(line 255 to 257) “The early layer in ResNet refers to ResNet.maxpool layer, and the late layer in ResNet refers to ResNet.avgpool layer. The early layer in CLIP refers to CLIP.visual.avgpool layer, and the late layer in CLIP refers to CLIP.visual.attnpool layer.”

(line 420 to 434) “As shown in Figure 3F, the early layer representations of both ResNet and CLIP (ResNet.maxpool layer and CLIP.visual.avgpool) showed significant correlations with early EEG time windows (early layer of ResNet: 40-280ms, early layer of CLIP: 50-130ms and 160-260ms), while the late layers (ResNet.avgpool layer and CLIP.visual.attnpool layer) showed correlations extending into later time windows (late layer of ResNet: 80-300ms, late layer of CLIP: 70-300ms). Although there is substantial temporal overlap between early and late model layers, the overall pattern suggests a rough correspondence between model hierarchy and neural processing stages.

We further extended this analysis across intermediate layers of both ResNet and CLIP models (from early to late, ResNet: ResNet.maxpool, ResNet.layer1, ResNet.layer2, ResNet.layer3, ResNet.layer4, ResNet.avgpool; from early to late, CLIP: CLIP.visual.avgpool, CLIP.visual.layer1, CLIP.visual.layer2, CLIP.visual.layer3, CLIP.visual.layer4, CLIP.visual.attnpool).”

| (e) *p19: please change the reporting of t-statistics to standard APA format.*

Thanks for the suggestion. We changed the reporting format accordingly:

(line 392 to 394) “The representation of real-word size had a significantly later peak latency than that of both retinal size, $t(9)=4.30$, $p=.002$, and real-world depth, $t(9)=18.58$, $p<.001$. And retinal size representation had a significantly later peak latency than real-world depth, $t(9)=3.72$, $p=.005$.”

(2) "early layer of CLIP: 50-130ms and 160-260ms), while the late layer representations of two ANNs were significantly correlated with later representations in the human brain (late layer of ResNet: 80-300ms, late layer of CLIP: 70-300ms)."

This seems a little strong, given the large amount of overlap between these models.

We agree that our original wording may have overstated the distinction between early and late layers, given the substantial temporal overlap in their EEG correlations. We revised this sentence to soften the language to reflect the graded nature of the correspondence, and now describe the pattern as a general trend rather than a strict dissociation:

(line 420 to 427) “As shown in Figure 3F, the early layer representations of both ResNet and CLIP (ResNet.maxpool layer and CLIP.visual.avgpool) showed significant correlations with early EEG time windows (early layer of ResNet: 40-280ms, early layer of CLIP: 50-130ms and 160-260ms), while the late layers (ResNet.avgpool layer and CLIP.visual.attnpool layer) showed correlations extending into later time windows (late layer of ResNet: 80-300ms, late layer of CLIP: 70-300ms). Although there is substantial temporal overlap between early and late model layers, the overall pattern suggests a rough correspondence between model hierarchy and neural processing stages.”

(3) "Also, human brain representations showed a higher similarity to the early layer representation of the visual model (ResNet) than to the visual-semantic model (CLIP) at an early stage. "

This has been previously reported by Greene & Hansen, 2020 J Neuro.

Thanks! We added this reference.

(4) "ANN (and Word2Vec) model RDMs"

Why not just "model RDMs"? Might provide more clarity.

We chose to use the phrasing “ANN (and Word2Vec) model RDMs” to maintain clarity and avoid ambiguity. In the literature, the term “model RDMs” is sometimes used more broadly to include hypothesis-based feature spaces or conceptual models, and we wanted to clearly distinguish our use of RDMs derived from artificial neural networks and language models. Additionally, explicitly referring to ANN or Word2Vec RDMs improves clarity by specifying the model source of each RDM. We hope this clarification justifies our choice to retain the original phrasing for clarity.

<https://doi.org/10.7554/eLife.98117.2.sa0>