

An image-computable model of speeded decision-making

Paul I. Jaffe , Gustavo X. Santiago-Reyes, Robert J. Schafer, Patrick G. Bissett, Russell A. Poldrack

Department of Psychology, Stanford University • Department of Bioengineering, Stanford University • Lumos Labs

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

 Copyright information

Reviewed Preprint

v1 • June 13, 2024

Not revised

Abstract

Evidence accumulation models (EAMs) are the dominant framework for modeling response time (RT) data from speeded decision-making tasks. While providing a good quantitative description of RT data in terms of abstract perceptual representations, EAMs do not explain how the visual system extracts these representations in the first place. To address this limitation, we introduce the visual accumulator model (VAM), in which convolutional neural network models of visual processing and traditional EAMs are jointly fitted to trial-level RTs and raw (pixel-space) visual stimuli from individual subjects. Models fitted to largescale cognitive training data from a stylized flanker task captured individual differences in congruency effects, RTs, and accuracy. We find evidence that the selection of task-relevant information occurs through the orthogonalization of relevant and irrelevant representations, demonstrating how our framework can be used to relate visual representations to behavioral outputs. Together, our work provides a probabilistic framework for both constraining neural network models of vision with behavioral data and studying how the visual system extracts representations that guide decisions.

eLife assessment

This **important** study presents an original and promising approach to combine convolutional neural networks of visual processing with evidence accumulation models of decision-making. While the methodological approach itself is strong and technically sophisticated, the evidence supporting the conclusions is currently **incomplete**. The study will be of interest to researchers working in the fields of machine learning and cognitive modeling.

<https://doi.org/10.7554/eLife.98351.1.sa3>

Introduction

Decision-making under time pressure is deeply embedded within the activities of daily life. To study the cognitive and neural processes underlying decision-making, psychologists fit computational models to response time (RT) data gathered from relatively simple cognitive tasks. Evidence accumulation models (EAMs), such as the diffusion-decision model [58 , 59 ] and linear ballistic accumulator model (LBA) [8 , are the most successful and widely-used computational models of decision-making. In the EAM framework, sensory evidence is

accumulated (possibly with noise) until reaching a threshold, at which point an overt response is generated. As such, decision-making is described in terms of a small number of interpretable parameters that capture basic cognitive processes. Empirically, EAMs have been shown to capture RT distributions from a variety of cognitive tasks at the individual subject level [8], [58], [60], [80].

Despite this empirical success and theoretical merit, the simple characterization of decision-making encapsulated by EAMs results in somewhat restrictive applications [17]. For example, EAMs do not typically provide a detailed specification of how stimuli are transformed into evidence or of the motor processes that physically generate a response. On the sensory side, the modeler instead converts the raw high-dimensional stimulus into a small number of abstract perceptual features that drive the rate of evidence accumulation. In the context of visual tasks, as we study here, this implies that EAMs are not amenable to studying the visual processing steps underlying decision-making, despite decades of experimental and theoretical vision research. It is also difficult to model tasks with complex (e.g. naturalistic) visual stimuli with EAMs, since the relevant visual features must be specified by hand.

Here, we integrate convolutional neural network (CNN) models for visual feature extraction with traditional accumulator models of speeded decision-making in a framework we call the visual accumulator model (VAM). We use CNNs since they are *image-computable*: they accept arbitrary images as input, such as visual stimuli from cognitive tasks. CNNs also capture key features of biological vision [37], [41], [85]. Each layer of the CNN applies a convolution, nonlinearity, and pooling operation to the output of the previous layer with a bank of trainable filters. In early network layers, this results in spatially-organized units with local receptive fields, analogous to the retinotopic organization of the early mammalian visual pathway. The concatenation of multiple layers forms a hierarchy in which progressively more abstract features are extracted in deeper layers, in a way that is globally similar to the primate visual cortical hierarchy.

Two features of the VAM distinguish our framework from other models in which neural networks were used to generate RTs or RT proxies [22], [39], [55], [71], [77]. First, the VAM generates RTs via a traditional EAM, a component that prior models lack. The VAM therefore inherits the advantages in interpretability afforded by EAMs and allows us to connect to the rich decision-making literature from psychology. Second, the CNN and EAM components of the VAM are jointly fitted to the RT, choice, and raw (pixel-space) visual stimuli from a single participant in a unified Bayesian framework. Thus, in contrast to previous models, the representations learned by the CNN and EAM parameters are directly constrained by behavioral data. We leverage these advantages to explore how abstract, task-relevant information is extracted from raw sensory inputs, and investigate how the behavioral phenomenon of congruency effects arises as a consequence of the representation geometry learned by the CNN. In doing so, our framework also addresses one of the main criticisms of deep neural network models of vision that are optimized to perform particular tasks (e.g. object identification): these models account for few results from psychology [2], [5], [18], [42], [44]. Indeed, as we will show, fitting the CNN with human data—rather than optimizing the model to perform a task—has significant consequences for the representations learned by the model.

Modeling framework

We first describe the task—Lost in Migration (LIM), available as part of the Lumosity cognitive training platform—which will serve to ground the discussion of the model. LIM is a stylized version of the well-known arrows flanker task used in the study of cognitive control and visual attention (Fig. 1A; [73]). In LIM, participants are shown stimuli composed of several arrow-shaped birds. The task is to indicate the direction of the central bird (the target) using the arrow keys on the keyboard, ignoring the direction of the surrounding birds (the flankers). Participants engage with the task at home and are rewarded via a composite score that takes into account both speed and accuracy.

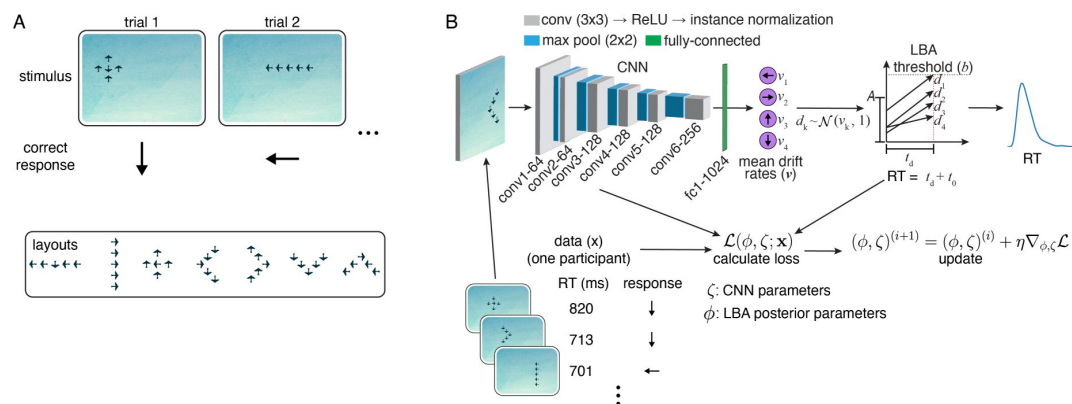


Figure 1

Task and model.

A) Top, Lost in Migration task. Bottom, the seven stimulus layouts (random target/flanker directions). **B)** VAM schematic. The numbers after the CNN layer names correspond to the number of channels used in that layer. See Methods for additional details.

The stimuli used in LIM vary along several dimensions that are not present in the standard flanker task, implying potentially complicated stimulus-behavior dependencies that are well-suited to the VAM. First, both targets and flankers can be independently oriented left, right, up, or down, such that there are four possible response directions (the standard arrow flanker task allows for only left and right responses). Second, the layout of the targets and the flankers can appear in one of seven configurations: left/right/up/down ‘V’, horizontal/vertical line, and cross (**Fig. 1A**). Last, the target can be centered anywhere within the 640x480 pixel game window, subject to edge constraints.

As with other flanker tasks, a given trial is said to be congruent if targets and flankers are oriented in the same direction, and incongruent otherwise. A consistent observation from studies employing the flanker task and related conflict tasks are *congruency effects*: participants are slower and less accurate on incongruent trials [16], [69], [74]. Congruency effects are considered to index a specific aspect of cognitive control: they indicate the extent to which an individual can selectively attend to task-relevant information and ignore irrelevant information.

The visual accumulator is an image-computable EAM, composed of a CNN and EAM chained together (**Fig. 1B**). Minimally processed (pixel-space) visual stimuli are provided as inputs to the CNN. The outputs of the CNN correspond to the mean rates at which evidence is accumulated (the drift rates), one for each possible response (four in the case of LIM). The EAM then generates choices and RTs through a noisy evidence accumulation process. For each participant in the dataset, we fit one such model (CNN + EAM) jointly using that participant’s visual stimuli, RTs, and choices. Building on prior work [11], [38], we developed an automatic differentiation variational inference (ADVI) algorithm that simultaneously optimizes the parameters of the CNN and learns the posterior distribution over the LBA parameters (Methods).

The particular EAM we adopt is the linear ballistic accumulator model (LBA [8]; **Fig. 1B**) since it can be applied to tasks with more than two possible responses, though the general VAM framework is compatible with other EAMs that have a closed-form or easily approximated likelihood. The LBA parameters fitted in the VAM are the decision threshold b , the non-decision time t_0 , and a parameter A that controls the dispersion of the initial accumulator values (the drift rate means are fitted implicitly via the CNN). We used a seven layer CNN architecture (6 convolutional layers, 1 fully-connected layer) in the VAM (**Fig. 1B**). To speed up the training process, the first two convolutional layers were initialized with the parameters from a 16-layer VGG CNN trained to classify the ImageNet dataset [12], [70].

All of the code (<https://github.com/pauljaffe/vam>) and data (<https://doi.org/10.5281/zenodo.10775514>) used to train the VAMs and reproduce our results are publicly-available.

Results

Models capture human behavioral data

We fitted a separate VAM to LIM data (visual stimuli, RTs, choices) from each of 75 Lumosity users (participants). We selected participants who had practiced Lost in Migration extensively ($\geq 25,000$ trials) and used data from a later stage of practice to minimize learning effects ($\geq$ each participant’s 50th gameplay). These participants varied in age (23–87 years) and in their behavior (mean RT, accuracy, congruency effects), allowing us to examine how well our framework captured individual differences.

To assess the model fits, we compared the mean RT, accuracy, RT congruency effect (incongruent trial mean RT minus congruent trial mean RT), and accuracy congruency effect (congruent trial accuracy minus incongruent trial accuracy) of each model/participant pair on a set of holdout

stimuli, separate from those used to train the models (Figs. 2 and S1). For each behavioral summary statistic, the responses of the fitted models were highly correlated with those of the participants (Pearson's $r > 0.75$), with slopes close to unity (Fig. 2B-E, statistics in figure legend). We found that the RT congruency effect could be attributed to a reduction in the mean target drift rate parameter on incongruent vs. congruent trials, while the accuracy congruency effect could be attributed to a higher mean flanker drift rate on incongruent trials relative to the non-target (other) mean drift rates on congruent trials (Fig. 2F). The latter observation follows from the fact that the drift rates on trial i are sampled from $N(v^{(i)}, 1)$, where the $v^{(i)}$ are the mean drift rates shown in Fig. 2F. Since the flanker drift rates on incongruent trials are higher (less negative) than the non-target drift rates on congruent trials, more errors result on incongruent trials, giving rise to the accuracy congruency effect.

We also examined whether the fitted models captured demographic effects, focusing on the well-established slowing of RTs that occurs with age [23], [50], [57]. Mean RTs began increasing around age 50, effects captured by the fitted models (Fig. 2G). Consistent with prior work in a variety of decision-making tasks [20], [57], [67], [72], we found that an age-dependent increase in the non-decision time (t_0) component of the response contributed to longer RTs in the models from older adults (Fig. S2A). In contrast to these prior studies, we did not observe increased response caution (measured as $b - A$) in the models from older adults (Fig. S2B). This discrepancy could be explained by differences in the task design or the fact that our participants had practiced the task substantially more than in typical studies [72]. We also observed an age-dependent reduction in the target drift rates and no age-dependence of the flanker drift rates (Fig. S2C-D), findings that have received mixed support in the literature [3], [20], [57], [67], [72].

One virtue of the VAM is that the model implicitly learns which stimulus properties influence behavior. As a simple demonstration of this, we investigated how two high-level visual features influenced participant behavior—stimulus layout and both horizontal/vertical stimulus position—and whether the fitted models captured these effects. We focused on RT effects, since no participants exhibited a significant effect of layout or horizontal/vertical position on accuracy ($p > 0.05$ for all stimulus features, chi-squared test), though we note that ceiling effects resulting from the high accuracy of the participants may have hindered our ability to detect these accuracy effects.

The majority of participants exhibited layout-dependent RT biases ($p < 0.05$ for 60/75 participants, one-way ANOVA; see examples in Fig. 2H and Fig. S1B). We quantified how well the models captured these effects by calculating the Pearson's r between model/participant mean RTs across each layout for the participants that exhibited significant layout-dependent RT modulation (Fig. 2J). The median Pearson's r was 0.67, demonstrating good correspondence between model/participant behavior. There was considerable heterogeneity in the particular layout RT biases exhibited by both the participants and models, though responses to trials with the vertical line layout were on average somewhat faster than the other layouts for both participants and models (Figs. S1B and S3A).

The majority of participants also exhibited both horizontal and vertical position-dependent RT biases (horizontal position: $p < 0.05$ for 72/75 participants, vertical position: $p < 0.05$ for 69/75 participants, one-way ANOVA; see examples in Fig. 2I and Fig. S1C-D). The fitted models captured these effects adequately: the median Pearson's r between model/participant mean RTs across stimulus position bins was 0.78 for horizontal position and 0.55 for vertical position (Fig. 2J). While there was substantial variability across both participants and models in the particular position-dependent RT biases they exhibited, most participants/models responded more slowly when the stimuli were positioned close to the horizontal edges and, to a lesser extent, the vertical edges of the task window (Figs. S1C-D and S3B-C).

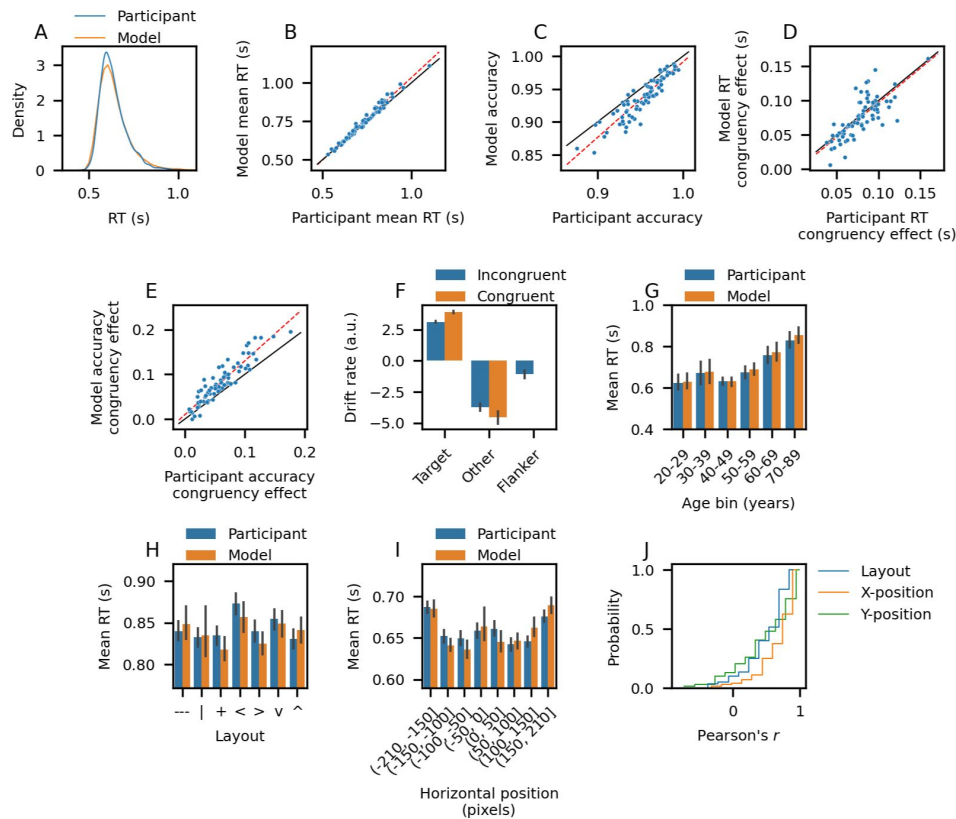


Figure 2

Comparison of model/participant behavior.

For panels B-E, each point is one model/participant ($n = 75$), black line: unity, red line: linear best-fit. **A)** Example model/participant RT distributions. **B)** Mean RT (Pearson's $r = 0.99$, bootstrap 95% CI = (0.99, 0.99), best-fit slope = 1.07). **C)** Accuracy ($r = 0.91$, 95% CI = (0.87, 0.94), slope = 1.15). **D)** RT congruency effect ($r = 0.77$, 95% CI = (0.67, 0.86), slope = 1.01). **E)** Accuracy congruency effect ($r = 0.92$, 95% CI = (0.88, 0.94), slope = 1.20). **F)** Drift rates averaged across all trials and models. **G)** Mean RT vs. age averaged across models. **H)** Example model/participant mean RT vs. stimulus layout (Pearson's $r = 0.67$). **I)** Example model/participant mean RT vs. horizontal stimulus position (negative values: left of center; Pearson's $r = 0.79$). **J)** Empirical CDF of Pearson's r between model/participant mean RTs across stimulus feature bins (only participants with significant RT modulation are shown; layout: $n = 60$ models/participants, x-position: $n = 72$, y-position: $n = 69$). Error bars in panels F-I are bootstrap 95% confidence intervals.

While the analyses above demonstrate that the mean congruency effects from the participants were captured well by the VAMs, we observed a mismatch between model and participant behavior for other commonly used behavioral indices that take into account additional information from the RT distribution: RT delta plots and conditional accuracy functions [30], [63], [78], [82], [83] (Fig. S4A-B). However, when we fitted a separate model to trials from each of $n=5$ RT quantiles from a given participant's data, quantitative fits improved and qualitative trends were captured for both delta plots and conditional accuracy functions (Fig. S4C-D, Methods). Thus, a small modification of the VAM training paradigm can be used to investigate how congruency effects change as a function of RT. However, a thorough investigation of these effects is outside the scope of this manuscript, and we focus on the models fit to all trials in the remainder of our study.

In summary, the VAM captured individual differences in RTs, accuracy, and congruency effects, and the dependence of RTs on stimulus layout and position. To lay the groundwork for understanding how the learned visual representations of the models relate to these behavioral effects, we sought to characterize the general properties of these representations that enable proficient execution of the task.

Representations of task-relevant stimulus information

To perform LIM well, the models must learn representations that select the task-relevant information (target direction) and diminish the influence of irrelevant information (flanker direction, stimulus layout, and stimulus position). To investigate these representations, we analyzed a set of stimulus *features* in each CNN layer, derived from the models' responses to a holdout set of LIM stimuli (separate from the training stimuli; Fig. 3A, Methods). The features in each layer form an $N \times K_l$ matrix, where N is the number of stimuli in the holdout set (5000) and K_l is the number of active units in layer l (a variable fraction of units in each layer did not respond to any stimuli and were excluded from the feature matrix).

We characterized the emergence of target selectivity in the model by quantifying how well the target direction could be decoded from the feature matrix of a given layer with a linear support vector machine [66] (SVM; Fig. 3B). Note that only incongruent trials were used in these analyses, since the classifier could use the flanker direction to classify the target on congruent trials, artificially inflating performance. Decoding performance on holdout data for target direction was near chance (27%) in the first network layer and increased to nearly perfect decoding ($\geq 97\%$) at layer 4, with a sharp increase between the second and third convolutional layers (Fig. 3B). The increase in target decoding accuracy from shallower to deeper network layers is generally consistent with neural recording studies in mammals and other neural networks studies that document more accurate decoding of abstract variables (e.g. object or category identity) in higher visual/cortical regions and deeper neural network layers [7], [47], [76].

We also investigated target selectivity at the single-unit level by quantifying the mutual information between each unit's activity and target direction [47], [76] (Fig. 3C). In contrast to the population-level decoding results, information for target direction exhibited a non-monotonic "hunchback" profile across the convolutional layers, then increased sharply in the final fully-connected layer. The observation that single-unit information for target direction decreased between the fourth and final convolutional layers while population-level decoding remained high is especially noteworthy in that it implies a transition from representing target direction with specialized "target neurons" to a more distributed, ensemble-level code. Notably, a similar transition in coding properties takes place along the cortical hierarchy, with higher-order cortical regions exhibiting a reduction in units with "pure" selectivity and a corresponding increase in units with mixed selectivity for multiple task or stimulus features [45], [46], [64]. The high proportion of units with mixed selectivity results in a high-dimensional representation in which all task-relevant information can easily be extracted with simple linear decoders [10], [51].

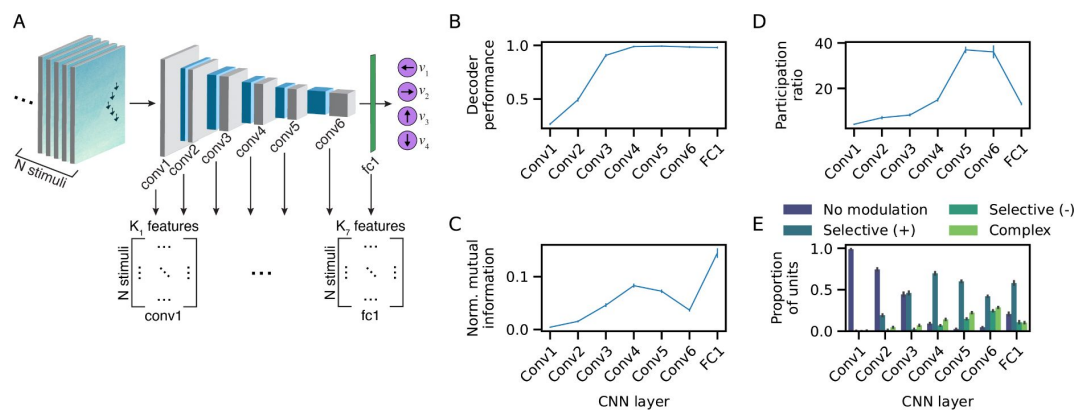


Figure 3

Neural representations of target direction.

A) Schematic of the CNN features extracted from each network layer. **B)** Decoding accuracy of stimulus target direction. **C)** Normalized mutual information for target direction conveyed by single units, averaged across units. Mutual information was normalized by the entropy of the target direction distribution. **D)** Dimensionality of target representations as measured by the participation ratio of the target-centered feature covariance matrix. **E)** Proportion of units exhibiting selectivity for target direction. Panels B-E show the average across $n = 75$ models; error bars correspond to bootstrap 95% confidence intervals.

Consistent with these ideas, we found that the dimensionality of target representations as measured by the participation ratio (Methods) increased sharply in the last two convolutional layers (Conv5-Conv6), paralleling the reduction in single-unit information for target direction in these layers (**Fig. 3D**). The highdimensional representation observed in these layers may enable the VAM to capture the rich stimulusbehavior dependencies present in the participant data. The reduction in dimensionality in the final fully-connected layer, and concomitant increase in single-unit information for target direction, may reflect a strong constraint to select the correct target direction imposed by the task. Notably, the hunchback-shaped profile we observed in the dimensionality of representations has also been observed along visual cortical regions of the rat ventral stream and in other neural network studies [1], [47].

To more explicitly characterize the degree of directional selectivity in each layer, we quantified the proportion of units that responded preferentially to one of the four target directions. We separated units according to the sign of modulation and degree of directional selectivity: “selective (+)” and “selective (-)” units were more (or less) active for one target direction relative to the other three; “complex” units exhibited significant modulation by target direction without a clear directional preference (Methods).

The proportion of units that were selective for a particular direction increased steadily from the second to the fourth convolutional layer, with most units exhibiting positive modulation (**Figs. 3E & S5**). Between the fourth and final convolutional layers, the proportion of units with positive target direction modulation decreased, while the proportion of units with negative and complex modulation increased.

In summary, a strong representation for target direction emerged in the middle convolutional layers, and was initially supported by a simple low-dimensional code, with information for each target direction concentrated in separate populations of directionally-tuned units. In later convolutional layers, target direction was supported by a more distributed, complex, and high-dimensional code, with weaker directional selectivity at the single-unit level.

Extraction and suppression of distracting stimulus information

How do the models’ visual representations enable target selectivity for stimuli that vary along several irrelevant dimensions? The mammalian visual system solves the analogous problem of visual object recognition through representations that become increasingly invariant to different object positions, sizes, and lighting conditions [14], [66]. To determine whether the models learned representations that share these characteristics, we initially assessed whether the model representations for target direction are indeed invariant (or more generally, tolerant), to variation in flanker direction, stimulus layout, and horizontal/vertical position. Alternatively, the models could have learned a coding scheme in which different representations are responsible for selecting the target in each distracter context, e.g. with units that respond to the conjunction of particular target direction and layout combinations.

To assess the degree of representation tolerance, we quantified how well a linear SVM trained to classify target direction in one distracter context generalized to a different distracter context, where the distracter context is defined by a particular type of task-irrelevant information [4], [66]. For example, to assess the degree of tolerance to flanker direction, we split trials into a training set where the flanker direction was fixed to a particular value (e.g. down), and a generalization set with all other flanker directions (e.g. flanker direction = left/right/up). The performance of the classifier on the generalization set measures the extent to which the target representations are tolerant (or invariant) to variability in a given type of distracting information.

For each type of distracting information, we found that generalization performance was initially near chance (25%) and increased steadily through the network layers (**Fig. 4A**). The lower generalization performance for flanker direction observed in the deepest network layers (~60%)

can be attributed to the robust congruency effects exhibited by the models (**Fig. 2D-E**): for the vast majority of incongruent error trials, the model chose the flanker direction, implying that flanker direction has a strong impact on model representations.

The tolerance of target direction representations to variability in irrelevant stimulus features suggests that the irrelevant information was progressively suppressed from shallower to deeper network layers. To examine this explicitly, we quantified how well each irrelevant stimulus feature (flanker direction, stimulus layout, horizontal/vertical position) could be decoded from the activity in each layer using the same SVM classifier methodology as we did for target direction decoding. We found that decoding accuracy for flanker direction and stimulus layout exhibited a hunchback profile in which decoding performance started low in early layers, increased in intermediate layers, and decreased in later layers (**Fig. 4B**). The decoding accuracy for horizontal and vertical position followed a similar but shifted pattern: there was a steady increase in accuracy until the second-to-last layer, followed by a slight drop in the last layer. Partial (rather than complete) suppression of irrelevant stimulus features is expected given that these features all impact behavior (**Fig. 2C**). Note that the increase in receptive field size that occurs from shallower to deeper layers is necessary but not *sufficient* for accurate decoding of the irrelevant stimulus features, and does not explain the reduction in decoding accuracy in later layers.

We also examined suppression at the single-unit level by quantifying the mutual information between each unit's activations and the values of each task-irrelevant stimulus feature [76] (**Fig. 4C**). The mutual information for a given stimulus feature was normalized by the entropy of the feature distribution to facilitate comparisons between the features [47]. In agreement with the population-level decoding analyses, information for the irrelevant stimulus features exhibited a pronounced hunchback profile in the progression from shallower to deeper network layers. It is noteworthy that the suppression of irrelevant information is more pronounced at the single-unit level in that decoding accuracy remains relatively high in later layers, particularly for flanker direction. This parallels the findings discussed above for target direction, and again suggests a transition from a simple code with populations of units that are selective for particular stimulus features to a more distributed, ensemble-level code.

Orthogonality of task-relevant and irrelevant information predicts behavior

A noteworthy feature of visual attention is that the selectivity and tolerance identified above is not absolute: irrelevant information cannot be filtered out completely, as illustrated by the congruency effects observed in the flanker task and related conflict tasks. A common framework for understanding these behavioral phenomena posits that task-relevant and irrelevant information compete for control over response execution, and that congruency effects arise from incomplete suppression of irrelevant information [78], [83]. Motivated by these theories, we investigated whether the degree of suppression of irrelevant information in the trained models was correlated with congruency effects across the models we analyzed.

To this end, we operationalized suppression with two metrics of the model representations that we investigated above: the accuracy of decoders trained to classify flanker direction from the model representations and the mutual information for flanker direction conveyed by single units. We expected to observe a positive correlation between both of these metrics and both RT and accuracy congruency effects: higher decoder accuracy or mutual information for flanker direction corresponds to less suppression and therefore higher congruency effects. However, we did not observe a significant positive correlation between either decoding accuracy or mutual information and RT or accuracy congruency effects in any of the model layers, with the exception of a single significant positive correlation between flanker mutual information and accuracy congruency effects in layer Conv3 (**Fig. S6**).

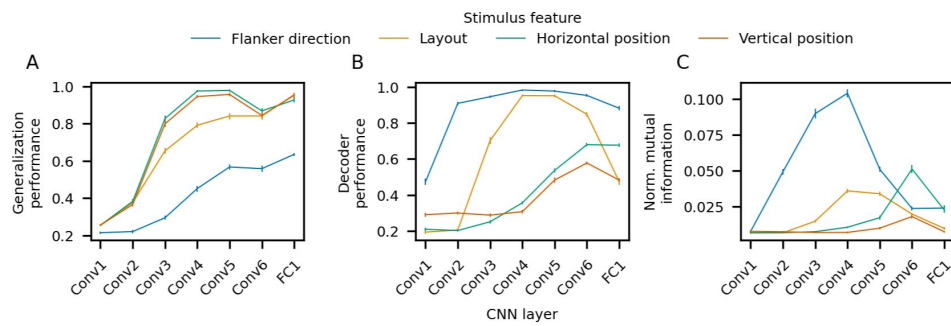


Figure 4

Suppression of task-irrelevant information and tolerance in task-relevant representations.

A) Decoding accuracy of stimulus target direction in a new distracter context (generalization performance). Context was defined by the values of a given stimulus feature (flanker direction, layout, horizontal/vertical position). **B)** Decoding accuracy of irrelevant stimulus features. **C)** Normalized mutual information for irrelevant stimulus features conveyed by single units, averaged across units. For each stimulus feature, the mutual information was normalized by the entropy of the stimulus feature distribution. All panels show the average across $n = 75$ models; error bars correspond to bootstrap 95% confidence intervals.

Given this somewhat surprising negative result, we were motivated to consider an alternative account of congruency effects, one that takes into account the relative geometry of the task-irrelevant (flanker) *and* taskrelevant (target) information. In particular, we considered the possibility that task-relevant and irrelevant information could be orthogonalized in the high-dimensional space of neural activity, such that task-relevant information is shielded from distracter interference. The general idea that neural representations for different types of information can be orthogonalized to prevent interference has received support from a number of neural recording studies [19], [32], [40], [52], [65].

To examine whether target and flanker representations are orthogonalized, we first defined target and flanker subspaces from the target direction and flanker direction classifiers used in the decoding analyses described above (Methods). The target direction classifier for a given network layer implicitly defines four decoding vectors, one for each target direction [4], [40]. We define the subspace of the K -dimensional feature space spanned by these four vectors as the target subspace; the flanker subspace is defined analogously.

To measure the orthogonality between target and flanker subspaces, we calculated the average of the cosine of the principal angles between the target and flanker subspaces, a metric we refer to as subspace alignment (Methods). Principal angles generalizes the idea of angles between lines or planes to arbitrary dimensions [31], [88]. The subspace alignment metric has a simple interpretation: it is equal to one if the subspaces are completely parallel and zero if they are completely orthogonal.

We first characterized how target/flanker subspace alignment develops across the layers of the trained models (Fig. 5A). Given that target direction decoding is poor in the first two layers (< 50%; Fig. 3B), the decoding vectors used to define the target subspace are not particularly meaningful, and we do not attempt to interpret the subspace alignment metric in these layers. Beginning at the third convolutional layer, when decoding accuracy for both targets and flankers is high (> 90%), we found that target and flanker subspaces are well-aligned. We interpret the high alignment as evidence that the model has learned a common representation for direction that is shared for both targets and flankers. In later layers, we found that target and flanker representations become increasingly orthogonal, consistent with the view that the processing in later layers acts to reduce interference between task-relevant and irrelevant information.

If orthogonalizing target and flanker representations reduces interference from the irrelevant (flanker) information, we should observe a positive correlation between subspace alignment and congruency effects across models, since greater alignment results in more interference. Consistent with this idea, we observed a significant positive correlation between subspace alignment and accuracy congruency effects across models in each layer beginning with the fourth convolutional layer (Fig. 5B-C; adjusted p -value < 0.05 for layers 3–7, permutation test, Bonferroni correction for 7 comparisons). In contrast, we did not observe a significant correlation between target/flanker subspace alignment and RT congruency effects in any network layer, suggesting a mechanistic dissociation between RT and accuracy congruency effects (Fig. S7).

Representation geometry of task-optimized models

Researchers who use neural network models to study neural representations typically optimize the model to perform a task, rather than fit the model to behavioral data from the task as we do here [45], [81], [86], [87], cf. [13], [29], [75]. This raises the possibility that these two training paradigms induce different representations in the models. To investigate this, we trained CNNs to perform LIM (i.e. output the direction of the target) by minimizing the standard cross-entropy loss function used in image classification tasks. Other than the training objective, all aspects of the optimization algorithm, CNN architecture, initialization, and training data were the same as those used to train the VAM. We trained one task-optimized model for each VAM, using data from the same participants ($n = 75$ task-optimized models).

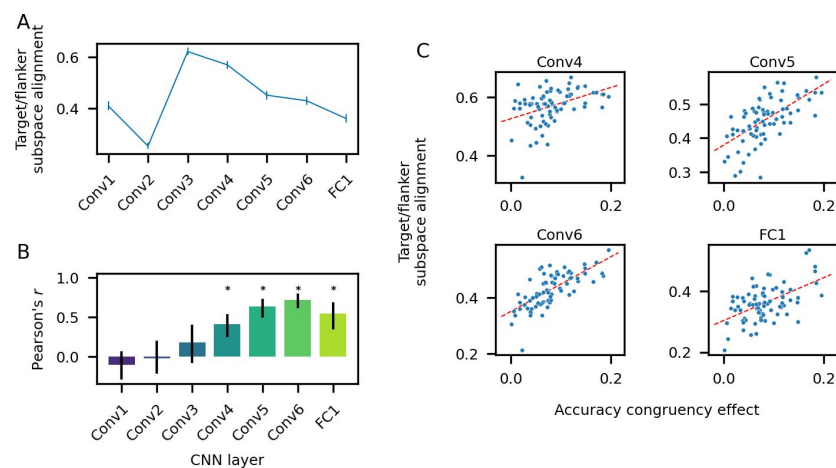


Figure 5

Orthogonality of target/flanker subspaces predicts accuracy congruency effects.

A) Target/flanker subspace alignment averaged across models. **B)** Pearson's correlation coefficient between target/flanker subspace alignment and accuracy congruency effect calculated across models. **C)** Target/flanker subspace alignment vs. accuracy congruency effect for layers Conv4-FC1. Each point corresponds to one model; the red line is the linear best-fit. For all panels, $n = 75$ models. Error bars in panels A-B correspond to bootstrap 95% confidence intervals. Asterisks in panel B indicate a significant Pearson's r (adjusted p -value < 0.05 , permutation test with $n = 1000$ shuffles, Bonferroni correction for 7 comparisons).

We first compared the behavioral outputs of the task-optimized models and the VAMs. We found that the task-optimized models did not exhibit an accuracy congruency effect (**Fig. 6A**). Note we could not examine the RT congruency effect since the task-optimized models do not generate RTs. Thus, simply training the model to perform the task is not sufficient to reproduce a behavioral phenomenon widely-observed in conflict tasks. This challenges a core (but often implicit) assumption of the task-optimized training paradigm, namely that training a model to do a task well will result in model representations that are similar to those employed by humans.

Above, we showed that VAMs with greater orthogonalization of target and flanker information exhibit smaller accuracy congruency effects (**Fig. 5B**), providing evidence that the relative geometry of task-relevant and irrelevant representations is a critical determinant of the degree of flanker interference at the behavioral level. Given that the task-optimized models do not exhibit an accuracy congruency effect, we therefore expect that these models would exhibit a higher degree of orthogonalization of target and flanker information. Consistent with this idea, we found that the task-optimized models had lower target/flanker subspace alignment (i.e. higher orthogonalization) for all network layers beginning with the third convolutional layer (**Fig. 6B**).

A direct consequence of the training paradigms used to train the VAMs is that these models are encouraged to capture dependencies between the stimulus features and behavior (RTs and accuracy), while the task-optimized models are not. As a result, the VAMs may learn richer representations of the stimuli, since a variety of stimulus features—layout, stimulus position, flanker direction—influence behavior (**Fig. 2**). To investigate this possibility, we compared the dimensionality of target representations between the VAMs and task-optimized models using the participation ratio metric discussed above.

We found that the dimensionality of the VAM and task-optimized model representations was nearly identical for the first four convolutional layers (**Fig. 6C**). In contrast, for the final two convolutional layers, the VAMs exhibited a substantially more pronounced expansion of dimensionality than the task-optimized models. In the final fully-connected layer, dimensionality decreased sharply for both types of models, and was somewhat higher for the VAMs.

The increased dimensionality of VAMs' target representations in later network layers is consistent with the view that these models must learn more complex representations of the stimuli in order to successfully capture stimulus-behavior dependencies. It is also noteworthy that the most striking difference between the dimensionality of the VAMs and task-optimized models occurs during the latter part (layers Conv5-Conv6) of the expansion phase of the hunchback-shaped dimensionality profile discussed above and observed in prior work [1], [47]. In these layers, single-unit information for target and flanker direction—the primary task features—decreases, while population-level decoding of these features remains high (**Figs. 3B-C** & **4B-C**). As discussed above, this dissociation implies a transition from a simple representation of target/flanker direction with separate populations of directionally-tuned units to a more complex and distributed code.

To determine whether the task-optimized models exhibited this change in coding properties, we quantified the single-unit information and population-level decoding accuracy for target/flanker direction in these models. Relative to the VAMs, the task-optimized models had substantially higher single-unit information for both target and flanker direction in layers Conv5-Conv6 (**Fig. 6D**). The task-optimized models also showed marginally more or roughly equivalent single-unit information for stimulus position in these layers relative to the VAMs, but had less single-unit information for stimulus layout (**Fig. S8A**).

In contrast, population-level decoding accuracy for target/flanker direction and stimulus position was similar between the task-optimized models and VAMs in layers Conv5-Conv6, though decoding accuracy for stimulus layout was notably lower for the task-optimized models (**Fig. 6E** & **S8B**).

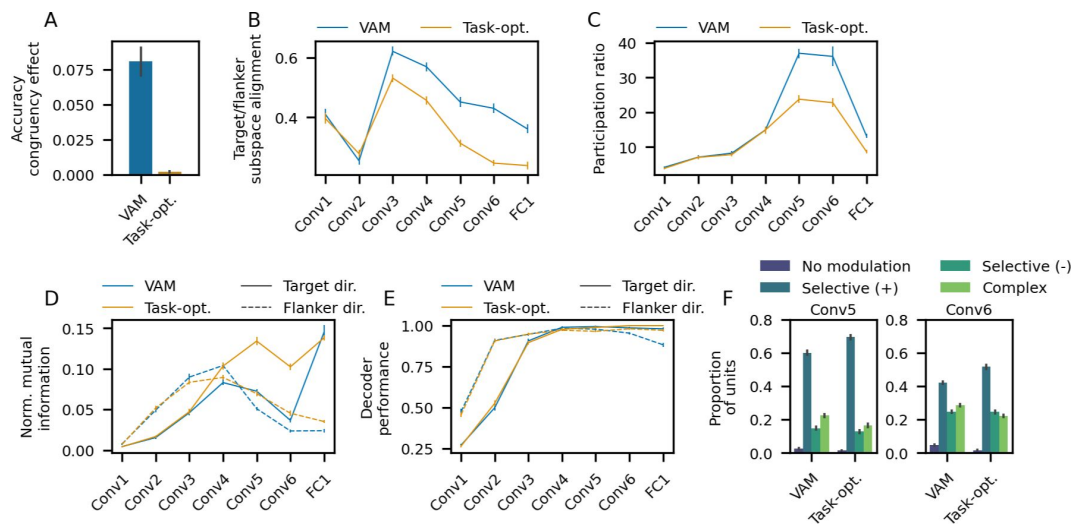


Figure 6

Comparison of VAMs and task-optimized models.

A) Accuracy congruency effect. **B)** Target/flanker subspace alignment. **C)** Dimensionality of target representations, as measured by the participation ratio of the target-centered feature covariance matrix. **D)** Normalized mutual information for target/flanker direction conveyed by single units, averaged across units. Mutual information was normalized by the entropy of the target/flanker direction distribution. **E)** Decoding accuracy of target/flanker direction. **F)** Proportion of units exhibiting selectivity for target direction in layers Conv5-Conv6. All panels show the average across $n = 75$ task-optimized models and $n = 75$ VAMs; error bars correspond to bootstrap 95% confidence intervals. The VAM data shown in panels A-F is the same as that shown in **Figs. 2E**, **5A**, **3D**, **3C**, **4C**, **3B**, **4B**, and **3E**, respectively.

These results suggest that the task-optimized models maintained a simpler code for target/flanker direction in the later convolutional layers relative to the VAMs, primarily relying on separate populations of directionally-tuned units. Consistent with this idea, the task-optimized models had a higher proportion of simple selective (+) directionally-tuned units and a lower proportion of complex units in layers Conv5-6 relative to the VAMs (Fig. 6F).

Discussion

The dominant models of decision-making, while providing a good quantitative description of psychophysical data, do not incorporate biologically plausible models of the perceptual processes that are essential for many behaviors [17]. On the other hand, neural network models of vision, while capturing core properties of the primate visual system, do not account for many results from behavioral experiments [2], [5], [18], [42], [44]. The VAM addresses both of these limitations by integrating neural network models for visual processing with traditional decision-making models in a unified probabilistic framework that can be fitted to visual stimuli and RT data. Leveraging large-scale data from a task with rich visual stimuli, we demonstrate that our framework captures complex dependencies between stimulus features and behavior at the level of individual participants. We also illustrate how congruency effects—a core behavioral phenomenon observed in conflict tasks—can be explained in terms of the visual representations of the model. Finally, we document several key differences between the representations learned by models fitted to human behavioral data (the VAMs) and those learned by models trained only to do the task.

To perform the task and capture the behavioral data, the VAMs learned representations that—like the mammalian cortex—extract task-relevant information from raw visual inputs. Our analyses of these representations revealed several discrete processing phases that are fruitfully discussed in relation to the changes in target representation dimensionality we observed. Across the layers of the VAM's CNN, target dimensionality exhibited a prominent hunchback shape profile, corresponding to an initial protracted phase of dimensionality expansion followed by abrupt dimensionality compression in the final network layers. An analogous expansion and subsequent compression of object representation dimensionality has been documented in CNNs trained to classify images and along the rat visual-cortical hierarchy [1], [47], with some notable differences from our work that we highlight below.

The initial phase of dimensionality expansion can be explained in part by the pruning of low-level stimulus features (e.g. contrast and luminosity) that are correlated across stimuli in the dataset [1]. These correlations are particularly strong in our dataset since we used the same background image for each stimulus, resulting in low-dimensional activity in the initial network layers. We speculate that the initial increase in target representation dimensionality (i.e. layers Conv1-Conv4) results from the removal of these correlations, analogous to a whitening transformation [1].

In the early and middle convolutional layers (Conv1-Conv4), we found that the population-level decoding and single-unit information for both task-relevant (target direction) and distracting information (flanker direction, stimulus layout, and position) increased, occurring in parallel with the subtle expansion of dimensionality we observed in these layers. These trends are broadly consistent with observations that the mammalian visual cortex encodes increasingly abstract stimulus properties (e.g. object identity) in downstream brain regions [24], [28], [66]. The fact that information for distracting stimulus features increased in these layers is especially noteworthy, given that this information—the flanker direction, most prominently—impairs performance on the task [16], [73]. While it is tempting to speculate that the VAMs learned to extract distracter information because they were fitted to behavioral data, where the impairments in task performance manifest, the task-optimized models also exhibited increased distracter

information in these layers. This suggests an alternative explanation in which the model first extracts more granular, superordinate stimulus representations that group together stimulus features with similar statistics. For example, the initial increase in encoding of target and flanker direction may reflect the representation of a unitary “directional signal” that facilitates the eventual separation of target and flanker direction representations that occurs in later layers. This idea is consistent with our observation that target/flanker subspaces are highly aligned in the middle convolutional layers and become progressively more orthogonal in later network layers.

In the last convolutional layers (Conv5-Conv6), we observed a large expansion in target representation dimensionality, which was notably less pronounced in the task-optimized models compared to the VAMs. The increase in dimensionality may be partly due to an increase in mixed selectivity for multiple stimulus features at the single-unit level, similar to the single-unit coding properties observed in primate PFC [64]. In general agreement with this idea, we observed a reduction in the proportion of units with selectivity for particular target directions in these layers, and a concomitant increase in units that were modulated by multiple target directions. Relative to the VAMs, the task-optimized models showed a higher proportion of directionally-tuned units in these layers. A notable advantage of high-dimensional neural codes composed of units with mixed selectivity is that many different input-output relationships can be implemented with simple decoders [10], [51], [64]. Thus, the higher dimensionality and more complex target direction coding observed in the VAMs relative to the task-optimized models may reflect the fact that the VAMs are trained to capture a potentially large number of dependencies between stimulus attributes and behavior, while the task-optimized models need only extract the target direction. Given the impressive capacity of primates to learn complex tasks and arbitrary stimulus behavior relationships, and the abundance of high-dimensional representations across the cortex, models trained to capture the richness of stimulus-behavior dependencies—such as the VAM—may result in better models of cortical processing than optimizing models to perform tasks.

The increase in dimensionality in layers Conv5-Conv6 may also result from suppression of correlations induced by covarying high-level stimulus features—in particular, target and flanker direction—a transformation that we conjectured to operate on low-level stimulus features in the earlier network layers. Consistent with this, the largest reduction in target/flanker subspace alignment occurs simultaneously with the largest expansion in target representation dimensionality, between layers Conv4 and Conv5. Given that the VAMs and task-optimized models exhibit a comparable reduction in target/flanker subspace alignment between these layers, the larger dimensionality that we observed in the VAMs may be driven primarily by other factors, such as the increase in units with mixed selectivity discussed above.

One of the key strengths of our modeling framework is that neural representations of stimulus features can be directly related to behavioral outputs. We focused in particular on congruency effects since they are ubiquitously observed in conflict tasks and highlight an inherent limitation of selective attention, namely that humans cannot completely filter out distracting information. Congruency effects are often interpreted in the context of a “dual-process” model in which automatic or impulsive processing arising from distracting stimulus information and more controlled or deliberate processing of task-relevant information compete for control over behavior [62], [78], [83]. Congruency effects result from the incomplete suppression of incorrect response activation in the automatic pathway. Building on this framework, a variety of models have been proposed that successfully capture congruency effects and other behavioral phenomena observed in conflict tasks [9], [78], [82].

Our work differs from these prior modeling efforts in two key ways. First, we did not attempt to “build in” a predetermined implementation of congruency effects. Rather, we fitted a relatively unstructured neural network model to human flanker task data and examined how congruency effects emerged [29]. Second, we explicitly modeled how task-relevant (and task-irrelevant) representations are extracted from raw visual stimuli with a biologically-plausible model of the

mammalian visual system (a CNN). These features of the VAM allowed us to explore a space of possible explanations for congruency effects that are not readily investigated with prior conflict task models. Each CNN layer is composed of many simple neuron-like processing units, enabling a population-level description of task-relevant and irrelevant representations. Prior connectionist models of conflict tasks are also implemented as networks of interacting processing units [9], but are much reduced in scale relative to the VAM, precluding a characterization of stimulus feature representations in terms of their high-dimensional geometry as we pursue here. Another relevant attribute of CNNs is that the successive convolution and pooling operations in each layer extract increasingly abstract, task-related representations from raw visual inputs, a hallmark of the mammalian ventral visual system. This enabled a rich characterization of the intermediate visual representations that sequentially transform raw visual stimuli to behavioral outputs.

Leveraging these features of the VAM framework, we investigated potential explanations for congruency effects in terms of the population-level activity and representation geometry in each network layer. Motivated by the dual-process model of conflict tasks mentioned above, we initially sought evidence for the suppression of task-irrelevant information, as operationalized by population-level decoding accuracy and single-unit information for flanker direction. While we did find evidence of greater suppression in later vs. middle convolutional layers for both of these metrics, variability in the magnitude of this suppression across models was not related to variability in congruency effects.

Instead, we found that the relative geometry of target and flanker representations was a critical determinant of congruency effects. In particular, models with smaller accuracy congruency effects had more orthogonal (or less aligned) representations of targets and flankers in later network layers, proximal to behavioral outputs. This finding is consistent with the idea that orthogonalization of task-relevant/irrelevant representations shields the task-relevant information from distracters, and parallels findings from neural recording studies in both humans and animals that representations for different sources of information can be orthogonalized in order to prevent interference [19], [32], [40], [52], [65], [84].

To elaborate on this idea, in models with more aligned (less orthogonal) target/flanker representations in intermediate layers, the task-relevant (target) subspace is effectively “contaminated” by task-irrelevant (flanker) information. In the absence of any corrective process, the task-relevant subspace propagates a mixture of both correct (target) and incorrect (flanker) direction signals forward through the network. At the final readout layer of the network, assuming that the target subspace is well-aligned with the readout weights [53], flanker signals within the target subspace then contribute to the drift rate for the (incorrect) flanker direction. This increases the probability of an error, such that models with more aligned target/flanker representations exhibit larger accuracy congruency effects. A corollary of this proposition is that the absolute amount of flanker information in the network—equivalently, the degree to which flanker information has been suppressed—is not necessarily predictive of congruency effects. This may account for our observations that the measures of suppression we considered are not correlated with accuracy congruency effects across models, and that representations for flanker direction do not appear to be strongly suppressed even in later network layers, as evidenced by high decoding accuracy for flanker direction in both the VAMs and task-optimized models.

One apparent limitation of the VAM as presented here is that it does not have dynamics, which seem to be required to explain some observations in the flanker task and related conflict tasks. For example, the RTs on incongruent error trials are typically faster than error RTs for congruent trials and RTs for correct trials [83], an effect that we confirm is also present in the flanker task variant studied here. This observation can be explained by a “shrinking attention spotlight” in which response activation from flankers starts high and diminishes over time, resulting in a higher proportion of errors for faster RTs [82]. While the primary VAMs analyzed in this paper did not capture these error patterns, we show that the simple modification of training separate

VAMs on each RT quantile produced error patterns that closely resemble those observed in the participant data. The representations learned by these models could in principle be compared, allowing one to investigate the mechanisms underlying these error phenomena. Future work may aim to incorporate true dynamics into the visual component and decision component of the VAM with recurrent CNNs [22], [49] and the task-DyVA model [29], respectively.

In summary, the VAM is a probabilistic model of psychophysical data that captures how raw sensory inputs are transformed into the abstract representations that guide decisions. Raw (pixel-space) visual stimuli are processed by a biologically-plausible neural network model of vision that outputs the parameters of a traditional decision-making model. Each VAM is fitted to data from a single participant, a feature that allowed us to study how individual differences in behavior emerge from differences in the “brains” of the models. To this end, we found that models with smaller congruency effects had more orthogonal representations for task-relevant and irrelevant information. While we chose to use a CNN to model visual processing, we note that the VAM is not limited to this choice: other sensory encoding models, such as those based on transformer architectures [15], can be readily swapped in to replace the CNN with minimal changes to the underlying VAM implementation. Similarly, the LBA decision-making model we employed is easily replaced with other decision-making models that have a closed-form or easily approximated likelihood, such as the diffusion-decision model and leaky competing accumulator model [43], [48], [60], [80]. In this way, the VAM provides a general probabilistic framework for jointly fitting interpretable decision-making models and expressive neural network architectures of sensory processing with psychophysical data.

Methods

Datasets

We used deidentified Lost in Migration gameplay data from 75 Lumosity users (participants) to train the models included in this paper. The mean (SD) age of this sample at the time of signup was 56.4 (18.6) years; 46.7% identified as female, 48.0% identified as male, and 5.3% did not report their gender. All analysis and modeling was done retrospectively on preexisting Lumosity data that were collected during the normal use of the Lumosity program (at home). All participants consented to the use and disclosure of their deidentified Lumosity data for any purpose as laid out in the Lumosity Privacy Policy (www.lumosity.com/legal/privacy_policy). No statistical methods were used to predetermine sample sizes, though our sample sizes are comparable to or larger than those used in related modeling work [8], [82].

The included participants were selected from a larger pool that met certain inclusion criteria. The selection from this larger pool was done at random, with the following exception: we included all participants under the age of 40 ($n = 19$) to ensure adequate representation of younger participants. The criteria used to define the larger participant pool were as follows: we required that participants had signed up as Lumosity users between June 28, 2015 and June 30, 2020 (inclusive); that they were between the ages of 18 and 89 at the time of signup (inclusive); that their country of origin was Australia, Canada, New Zealand, or the United States; that their preferred language was English; and that they were not employees of Lumos Labs, Inc. Accounts created for research purposes were also excluded. We also required that all of a given participant's Lost in Migration gameplays were done on the web (as opposed to mobile) platform. Finally, we required that participants had at least 200 Lost in Migration gameplays and at least 25,000 trials starting with their 50th gameplay. One additional participant with near chance accuracy (33%) on Lost in Migration was excluded from the larger participant pool.

Trials with very short and very long RTs were excluded and did not count toward the 25,000 trial minimum. Specifically, we excluded trials less than or equal to 250ms and trials classified as outliers from a criterion based on the median absolute deviation from the median (the MAD): trials with an absolute deviation from the median RT of more than 10 times a given participant's MAD were excluded.

The final datasets from each participant used for model training consist of the first 25,000 non-outlier trials starting with their 50th gameplay. Data from the first 49 gameplays were excluded to reduce learning-related variability. Approximately 50% of trials were congruent (vs. incongruent).

Modeling framework

Linear ballistic accumulator (LBA) model

The decision-making component of our modeling framework is the LBA model [8], tailored to Lost in Migration. Evidence for each of the four possible response directions is accumulated linearly and independently at a response-specific drift rate d_k . Commitment to a decision occurs when one of the accumulators reaches a fixed threshold parameter b . The RT on a given trial is the duration of this evidence accumulation process plus a constant non-decision time parameter t_0 that includes both sensory processing and motor execution. Response variability comes from two sources. First, on each trial, the drift rate d_k for each response $k \in \{1, 2, 3, 4\}$ is sampled independently from a Gaussian distribution with mean v_k and a common SD s . We fix s to 1 for all models to ensure that the LBA parameters are identifiable [11], [25]. Second, the initial evidence for each accumulator is sampled independently from a uniform distribution on the interval $[0, A]$, where A is a model parameter.

Visual accumulator model (VAM): overview

The VAM is a generalization of the LBA model in which the drift rate means $v^{(i)} = \{v_1^{(i)}, v_2^{(i)}, v_3^{(i)}, v_4^{(i)}\}$ on trial i depend on the stimuli $s^{(i)}$ through a CNN. For a dataset with N trials, the task data are $\mathbf{x} = \{\text{response times } \mathbf{t}, \text{choices } \mathbf{c}, \text{stimuli } \mathbf{s}\}$, where $\mathbf{t} = \{t^{(i)}\}_{i=1}^N$, $\mathbf{c} = \{c^{(i)}\}_{i=1}^N$ and $\mathbf{s} = \{s^{(i)}\}_{i=1}^N$. We adopt a Bayesian framework and model the joint density of the task data $\mathbf{x}$ and LBA parameters $\boldsymbol{\theta} = \{b, A, t_0\}$ as:

$$p(\mathbf{x}, \boldsymbol{\theta}) = \prod_{i=1}^N p(t^{(i)}, c^{(i)} | v^{(i)}, b, A, t_0) p(b, A, t_0), \quad (1)$$

$$v^{(i)} = \text{CNN}_{\boldsymbol{\zeta}}(s^{(i)}), \quad (2)$$

where the drift rate means $v_{(i)}$ depend on a CNN with parameters $\boldsymbol{\zeta}$ (described below). The factor on the left of Equation (1) is the likelihood of the LBA model, which has a closed-form solution [8]. The factor on the right is the prior distribution over the LBA parameters, specified as a standard multivariate Gaussian $p(\boldsymbol{\theta}) \sim N(\mathbf{0}, \mathbf{I})$. We express the joint density more compactly as:

$$p(\mathbf{x}, \boldsymbol{\theta}; \boldsymbol{\zeta}) = p(\mathbf{t}, \mathbf{c} | \text{CNN}_{\boldsymbol{\zeta}}(\mathbf{s}), b, A, t_0) p(b, A, t_0). \quad (3)$$

To fit the VAM, i.e. learn the posterior distribution over the LBA parameters $p(\boldsymbol{\theta} | \mathbf{x})$ and simultaneously optimize the CNN parameters $\boldsymbol{\zeta}$, we apply automatic differentiation variational inference (ADVI) [38]. Rather than attempting to sample from the true posterior directly as in Markov chain Monte Carlo (MCMC) methods, variational inference introduces an approximate

posterior density $q(\theta; \phi)$ and minimizes the Kullback-Leibler (KL) divergence from $p(\theta | \mathbf{x})$ to $q(\theta; \phi)$ by optimizing ϕ . Here, we specify the approximate posterior $q(\theta; \phi)$ as a multivariate Gaussian with mean μ and unconstrained covariance matrix Σ (thus $\phi = \{\mu, \Sigma\}$).

We use variational inference since it scales well to large datasets and can handle complicated models (such as the VAM), in contrast to MCMC methods. Leveraging automatic differentiation software (JAX [6]), we automate the calculation of the derivatives of the variational objective (described below) with respect to both the CNN parameters ζ and variational parameters ϕ .

Variable transformations

Note that the LBA parameters $\theta = \{b, A, t_0\}$ are restricted to be nonnegative, while $\mu \in \mathbb{R}^3$ and Σ is restricted to be positive semidefinite. However, to optimize the variational objective (defined below), we require that ϕ and θ have the same support [38]. To achieve this, we transform ϕ and θ to both have support on the real line. Specifically, we reparameterize Σ using the Cholesky decomposition as $\Sigma = \mathbf{L}\mathbf{L}^T$, where $\mathbf{L}$ is a lower-triangular matrix with entries in $\mathbb{R}$ (so that now $\phi = \{\mu, \mathbf{L}\}$). For the LBA parameters, in addition to mapping them to the real line, we must also enforce the constraint that the threshold b is always greater than the parameter A controlling the range of the initial evidence distribution. We enforce this by defining $\tilde{b} = b - A$ and taking the log of $\tilde{b}$, A , and t_0 to map them to the real line [11]. We define the transformed parameters as $\theta^* = \{b^*, A^*, t_0^*\} = \{\log(b - A), \log A, \log t_0\}$ and the transformation that maps θ to θ^* as T .

VAM objective function

To fit the VAM, we maximize the evidence lower bound (ELBO):

$$\mathcal{L}(\phi, \zeta) = \mathbb{E}_{q_\phi(\theta)}[\log p(\mathbf{x}, \theta; \zeta) - \log q(\theta; \phi)].$$

This is the standard ELBO optimized in variational inference, with an additional dependence on the CNN parameters ζ . The ELBO is a lower bound on the marginal likelihood of the data $p(\mathbf{x})$, and maximizing the ELBO is equivalent to minimizing the KL divergence from $p(\theta | \mathbf{x})$ to $q(\theta; \phi)$. Writing the ELBO in terms of the transformed variables θ^* , we have:

$$\mathcal{L}(\phi, \zeta) = \mathbb{E}_{q_\phi(\theta^*)}[\log p(\mathbf{x}, T^{-1}(\theta^*); \zeta) + \log |\det J_{T^{-1}}(\theta^*)| - \log q(\theta^*; \phi)]. \quad (4)$$

The term $\log |\det J_{T^{-1}}(\theta^*)|$ is a Jacobian adjustment for the transformation T^{-1} that maps θ^* to θ , required to ensure that the transformed density integrates to one. For the transformation T^{-1} , with θ^* vectorized as $[b^*, A^*, t_0^*]$, the Jacobian is given by:

$$J_{T^{-1}}(\theta^*) = \begin{bmatrix} b^* & A^* & 0 \\ 0 & A^* & 0 \\ 0 & 0 & t_0^* \end{bmatrix}.$$

Taking the log of the absolute value of the determinant of $J_{T^{-1}}(\theta^*)$ gives the required adjustment:

$$\log |\det J_{T^{-1}}(\theta^*)| = \log |b^* A^* t_0^*|. \quad (5)$$

We will apply stochastic gradient ascent to maximize the ELBO objective function given by Equation (4), using Monte Carlo (MC) estimates of the expectation and AD to calculate gradients with respect to ϕ and ζ . However, the gradients of the ELBO with respect to ϕ cannot be calculated directly by AD, since the expectation in Equation (4) is taken with respect to $q(\theta^*; \phi)$, which depends on ϕ . We work around this by applying the *reparameterization trick* [33], [61], which

expresses θ^* as a differentiable transformation of an auxiliary noise variable $\epsilon \sim p(\epsilon)$, where $p(\epsilon)$ does not depend on ϕ . In particular, we set $p(\epsilon)$ to be a standard multivariate Gaussian, $p(\epsilon) = N(\epsilon; \mathbf{0}, \mathbf{I})$, and define the transformation $\tilde{\theta}^* = \mathbf{L}\epsilon + \mu$, where μ and $\mathbf{L}$ are the mean and covariance parameters of the Gaussian approximating density $q(\theta^*; \phi)$.

Now we estimate the expectation in Equation (4) with MC samples from $p(\epsilon)$. Since the expectation no longer depends on ϕ , we can use AD directly on these MC estimates to obtain unbiased gradients of the ELBO. We estimate the ELBO for each data point using independent MC samples, since this reduces the variance of the gradients relative to using the same set of MC samples for a given batch of data [35]. Thus the ELBO objective for the i th data point is estimated by:

$$\tilde{\mathcal{L}}^{(i)}(\phi, \zeta) = \frac{1}{L} \sum_{l=1}^L \log p(x^{(i)}, T^{-1}(\theta^{*(i,l)}); \zeta) + \log |\det J_{T^{-1}}(\theta^{*(i,l)})| - \log q(\theta^{*(i,l)}; \phi), \quad (6)$$

where $\theta^{*(i,l)} = \mathbf{L}\epsilon^{(i,l)} + \mu$, $\epsilon^{(i,l)} \sim p(\epsilon)$, and $x^{(i)} = \{t^{(i)}, c^{(i)}, s^{(i)}\}$.

VAM inference algorithm and other training details

The VAM training/inference algorithm is summarized in **Algorithm 1**. The size of the training set was 16,250 samples (65% of the total dataset) for each model. An additional validation set of 3,750 samples (15%) was used to monitor model training. The remaining 5,000 samples (20%) were used to evaluate model performance (the holdout set). We set the batch size M to 256 and the number of MC samples used to estimate the ELBO L to 10. We used the Adam optimizer [34] with the following hyperparameters for model training: learning rate = $1e-3$, $\beta_1 = 0.9$, and $\beta_2 = 0.999$. The same random seed was used to initialize the parameters of all models. Models were trained using a single NVIDIA GeForce RTX 3060 Ti GPU. Each model took approximately one hour to train.

For a small fraction of attempted model training runs (10/85 = 11.8%), the model either failed to exceed chance accuracy (3/10 models) or converged to a state in which almost all trials had exclusively negative drift rates (7/10 models). These models were not included in summary analyses.

Convolutional neural network (CNN)

We used the same seven-layer CNN architecture for all models (six convolutional layers followed by one fully-connected hidden layer). The number of channels/units used in the seven layers was as follows: 64, 64, 128, 128, 128, 256, 1024. The output layer (after the fully-connected layer) has four channels, one for each drift rate. Each convolutional layer used a 3×3 pixel kernel (stride 1, same padding). Each convolutional layer was followed by a ReLU nonlinearity, instance normalization [79], and a max-pooling layer (2×2 pixel window, stride 2), in that order. The fully-connected hidden layer was followed by a ReLU nonlinearity and a dropout layer (dropout rate set to 0.5).

To speed up model training, the first two convolutional layers were initialized with the parameters from a larger CNN trained on an image classification task. Specifically, the pretrained model used a 16-layer VGG architecture and was trained on the ImageNet dataset [12], [70]. These first two layers were trainable (i.e. not fixed to their initial values). The weights of the other convolutional layers, fully-connected layer and output layer were initialized randomly using the LeCun normal initializer [36]; the biases were initialized with zeros.

Data: $\mathbf{x} = \{\text{RTs } t^{(i)} \dots t^{(N)}, \text{ choices } c^{(i)} \dots c^{(N)}, \text{ stimuli } s^{(i)} \dots s^{(N)}\}$
 $(\phi, \zeta) \leftarrow$ Initialize CNN parameters ζ and approximate posterior parameters ϕ
while *not converged* **do**
 $\mathbf{t}^M, \mathbf{c}^M, \mathbf{s}^M \sim \mathbf{x}$ (sample random minibatch of size M from full dataset)
 $\mathbf{v}^M \leftarrow \text{CNN}_{\zeta}(\mathbf{s}^M)$ (calculate drift rate means)
 $\boldsymbol{\epsilon}^M \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$ (draw $M \times L$ samples from the standard multivariate Gaussian)
 $\boldsymbol{\theta}^{*M} \leftarrow \mathbf{L}\boldsymbol{\epsilon}^M + \boldsymbol{\mu}$ (reparameterize)
 $\tilde{\mathcal{L}}^M(\phi, \zeta) \leftarrow$ calculate ELBO using equation (6)
 $\mathbf{g}^M \leftarrow \nabla_{\phi, \zeta} \tilde{\mathcal{L}}^M(\phi, \zeta)$ (calculate gradients using AD)
 $(\phi, \zeta) \leftarrow$ update parameters using the Adam optimizer and $\mathbf{g}^M$
end

Algorithm 1

VAM inference algorithm

Image preprocessing and data augmentation

The image stimuli used to train the VAM differed from the original Lost in Migration stimuli in a few ways. Information about the current score and time remaining visible at the top of the game window was removed. We also used the same blue background image for all stimuli (the background in the original Lost in Migration changes over the course of the gameplay). Finally, the stimuli were resized from 640×480 pixels (width×height) to 128×128 pixels.

We used data augmentation techniques commonly used in other image classification training paradigms to improve the generalization ability of the models. Specifically, for each batch of training data, each image was independently and randomly translated by a small amount. The size of the translation (in pixels) was drawn from a uniform distribution on the interval [0,1] (vertical) and [0,2] (horizontal), rounded to the nearest pixel (horizontal and vertical translations were sampled independently). We also applied a variation of a random elastic image deformation [68] as implemented in the Augmax Python package [26]. Specifically, we used the Warp transformation in Augmax and set the strength parameter to 3 and the coarseness parameter to 32. The transformation was applied to each image independently with a probability of 0.75.

Task-optimized models

The task-optimized CNN models were trained by minimizing the cross-entropy loss function, where the correct label was defined as the direction of the target bird. Other than the loss function, all aspects of the training process were the same as those used for the VAM (CNN architecture, optimizer settings, initialization, etc.). For each VAM we trained, we trained one task-optimized model using the same training data ($n = 75$ task-optimized models).

Models trained on separate RT quantiles

For the conditional accuracy function and delta plot results presented in **Fig. S4C-D**, we trained models on data from $n = 10$ randomly selected participants using a modified training paradigm. For each participant, we calculated the quintiles (i.e. the 0.2, 0.4, 0.6, and 0.8 quantiles) of their RT distribution and trained a separate VAM on the trials within each RT quantile bin (five models per participant). As for the VAMs used in the original training paradigm, 65% of the data within each RT quantile bin was used for training, 15% was used to monitor model training, and the remaining 20% was designated as holdout data used to evaluate model performance. We generated RTs and choices from the holdout stimuli for each of the five models from a given participant, then combined these outputs to compare with the participant's data.

For a subset of attempted training runs ($7/57 = 12.3\%$, 7 different participants), the model failed to converge (near chance accuracy for 6/7 models, 69% accuracy for 1/7 models). In these cases, we trained a new model from the same participant and RT bin, but used a different random seed to initialize the model parameters and partition the training set vs. holdout set (all such retrained models converged successfully). We did not observe a clear association between training success rate and RT quantile bin (number of failed models in RT quantile bins 1-5: 0, 2, 1, 2, 2).

Analysis methods: overview

We analyzed all of the VAMs after 12,800 parameter update steps (corresponding to the 200th training epoch), by which point training had converged. We analyzed all of the task-optimized models after 1,600 parameter update steps (the 25th training epoch) since these models converged much more quickly. The VAMs trained on data from a single RT quantile bin were all analyzed after 6,500 parameter update steps (the 500th training epoch, given the smaller training set size) since these models also converged more quickly than those trained with the entire dataset. All analyses were conducted using a holdout set of $n = 5000$ LIM stimuli/RTs/choices (i.e., data that was not used to train the models). To generate RTs and choices from the trained models, the LIM

stimuli from the holdout set were provided as inputs to the CNN, which output mean drift rates for each stimulus. We denote the drift rate mean for the k^{th} accumulator on trial i as $v_k^{(i)}$. We used the LBA model to sample RTs and choices using these mean drift rates, where the drift rate for the k^{th} accumulator on trial i , $d_k^{(i)}$, is sampled from a normal distribution with mean $v_k^{(i)}$ and SD = 1. For a very small fraction of trials (mean $\pm$ s.e.m. percent excluded trials: $0.080 \pm 0.011\%$, $n = 75$ models), all of the sampled drift rates were negative, and thus the RT and choice were undefined. These trials were excluded from all analyses.

All error bars shown in the figures correspond to bootstrap 95% confidence intervals calculated using 1000 bootstrap samples.

Behavior analyses

The RT congruency effect was defined as the mean RT on incongruent trials minus the mean RT on congruent trials, where only correct trials were included in the calculation. The accuracy congruency effect was defined as the accuracy on congruent trials minus the accuracy on incongruent trials.

To calculate the RT delta plots for a given model/participant, we first calculated the RT deciles (0.1, 0.2, ..., 0.9 quantiles) separately for congruent and incongruent trials, using only correct trials. Within each RT decile, we then calculated the RT congruency effect and mean RT (average of congruent and incongruent mean RTs), forming a delta plot for that participant/model. These mean RTs and congruency effects were then averaged across participants.

To calculate the conditional accuracy functions for a given model/participant, we calculated the accuracy and mean RT of trials within each RT quintile (0.2, 0.4, 0.6, and 0.8 quantiles) separately for congruent and incongruent trials. As above, these measures were then averaged across participants.

The models/participants included in the analysis of how RT varied with stimulus layout and horizontal/vertical position were those that exhibited significant modulation of RT by one or more of these stimulus features. Significant modulation was determined by running an ANOVA on the RTs in each stimulus feature bin (e.g. for each layout or each horizontal position bin), using a threshold p -value of 0.05. For horizontal stimulus position, we used bins of width = 50 pixels in the original 640×480 pixel window space, except for the leftmost and rightmost bins which had width = 60 pixels. For vertical stimulus position, we used bins of width = 25 pixels.

The number of participants exhibiting significant modulation of RT was 60 for stimulus layout, 72 for horizontal position, and 69 for vertical position (out of 75 total). To determine whether a given participant exhibited significant modulation of accuracy by a given stimulus feature, we used a chi-squared test for equality of proportions across stimulus feature bins (threshold p -value = 0.05). No participants exhibited significant modulation of accuracy for any of the stimulus features.

LBA parameters

For analyses of the fitted LBA parameters t_0 , b , and A , we used the maximum a posteriori (MAP) estimates of these parameters, corresponding to the mean parameter vector of the learned Gaussian posterior density $q(\theta; \phi)$. For analyses of the mean target and flanker drift rates, we provided the LIM stimuli from the holdout set as inputs to the fitted CNN, which generates the mean drift rates as its outputs. The non-target/non-flanker (other) drift rates were calculated by averaging the drift rates from the two non-target/non-flanker accumulators on incongruent trials or the three non-target accumulators on congruent trials.

CNN features

The model features used in the analyses of CNN representations were derived from the unit activations in each layer elicited by the holdout image set. The activations were processed just after the ReLU nonlinearity. The features in the convolutional layers are defined as the maximum value of each channel (i.e., the maximum was calculated across the spatial dimensions) [27], yielding a $N \times K_l$ feature matrix, where N is the number of images in the holdout set (5000) and K_l is the number of active channels in layer l . For the fully-connected layer, the features were defined as the $N \times K_l$ matrix of activations.

Only incongruent trials were used for all of the analyses involving the CNN features described below, for the following reasons. For the selectivity and tolerance analyses in which we decoded target or flanker direction from the model features in each layer, described in more detail below, note that the target and flanker directions are by definition the same on congruent trials. As such, a classifier trained to decode target direction from congruent trials could achieve perfect accuracy using information in the model features attributable to the flankers, rather than targets. Using only incongruent trials ensured that such cross-contamination could not occur. We used only incongruent trials for all of the other analyses of the model features simply for convenience.

Stimulus feature decoding

To assess how well particular stimulus features could be decoded from the activity in each layer, we trained a linear SVM to classify the values of that stimulus feature using the $N \times K_l$ activation (feature) matrix in each layer. The activation matrix was standardized before training the classifiers, such that each column had zero mean and unit variance. Horizontal and vertical stimulus positions were discretized to enable classification using the same bins as were used in the behavioral analyses.

The classification task was done in a standard “one-vs-rest” setting: for each value of a given stimulus feature, one sub-classifier was trained on a binary classification task with the chosen value as one class and all other values as the other class, yielding one classifier for each value of the stimulus feature (e.g. four for target direction, seven for stimulus layout). To determine the decision of the combined classifier for a given image, we generated predictions from each sub-classifier and assigned the decision to the sub-classifier with the largest (most confident) prediction. We assessed the overall decoding accuracy of each SVM on a separate test image set. Note that both the training set and test set for the SVMs were subsets of the holdout image set (i.e., the images that were not used to train the VAMs). The SVMs were trained using the LinearSVC model in the scikit-learn Python package [54] with the squared hinge loss function and L2 regularization with the penalty parameter C set to 1.0 (the default settings).

Tolerance

To assess the tolerance of model representations to variation in a given stimulus feature (flanker direction, stimulus layout, horizontal position, and vertical position), we trained a linear SVM to classify target direction using stimuli with the chosen stimulus feature fixed to one value (the training context), and assessed the generalization performance of the SVM on stimuli that contained all other values of that stimulus feature (the generalization context). For example, to assess tolerance to stimulus layout, we trained one SVM to classify target direction using stimuli with the vertical line layout, and assessed the generalization performance of that classifier on stimuli with the six other layouts. We trained one such SVM for each of the seven stimulus layouts, and averaged the generalization performance across these seven SVMs to derive the overall generalization performance measure for a given model and network layer. Other details of the classifiers were the same as those used for the decoding analyses described above.

Target/flanker subspace alignment. To calculate the target/flanker subspace alignment metric, we first defined target and flanker subspaces using the SVM classifiers that we trained for the target/flanker decoding analyses. Specifically, for each of the four classifiers, which were trained to classify stimuli as a given target or flanker direction vs. all other target or flanker directions using the model features from a given layer, we extracted the vector orthogonal to the decision hyperplane. In agreement with prior work [4], [40], we refer to these vectors as decoding vectors for a given target/flanker direction. Let $\mathbf{x}_{k,\text{targ}}^T$ and $\mathbf{x}_{k,\text{flnk}}^T$ denote the target decoding row vector and flanker decoding row vector for the k th direction, respectively. We define the matrices formed with these four decoding vectors filling the rows as the target subspace matrix $\mathbf{X}_{\text{targ}}$ and the flanker subspace matrix $\mathbf{X}_{\text{flnk}}$. These matrices have dimensions $4 \times K_l$, where K_l is the number of feature dimensions for layer l . Each matrix therefore spans a subspace of the full K_l -dimensional feature space. Our goal is to determine whether the target and flanker subspaces are orthogonal.

To do so, we employ principal angles between subspaces [31], [88], which generalizes the more intuitive notion of angles between lines or planes to arbitrary dimensions. To calculate the principal angles, we require orthonormal bases for the target and flanker subspaces, which we determine using a reduced singular value decomposition (SVD):

$$\mathbf{X}_{\text{targ}} = \mathbf{U}_{\text{targ}} \mathbf{\Sigma}_{\text{targ}} \mathbf{V}_{\text{targ}}^T, \quad \mathbf{X}_{\text{flnk}} = \mathbf{U}_{\text{flnk}} \mathbf{\Sigma}_{\text{flnk}} \mathbf{V}_{\text{flnk}}^T.$$

The rows of the $4 \times K_l$ matrices $\mathbf{V}_{\text{targ}}^T$ and $\mathbf{V}_{\text{flnk}}^T$ form an orthonormal basis for the target and flanker subspaces, respectively. The cosines of the principal angles are given by the singular values of $\mathbf{V}_{\text{targ}} \mathbf{V}_{\text{flnk}}^T$. The average of these singular values is our subspace alignment metric, which ranges from zero (completely orthogonal subspaces) to one (completely aligned/parallel subspaces).

Participation ratio

We measured the dimensionality of target representations for a given layer with the participation ratio (PR) [21], defined as:

$$\text{PR}_l = \frac{\left(\sum_{i=1}^{K_l} \lambda_i \right)^2}{\sum_{i=1}^{K_l} \lambda_i^2},$$

where K_l is the number of features in layer l and $\lambda_1 \geq \dots \geq \lambda_i \geq \dots \geq \lambda_{K_l}$ are the eigenvalues of the target-centered feature covariance matrix for layer l . The target-centered features were obtained by subtracting the centroid of the feature matrix for each target direction from the corresponding trials in the feature matrix [56].

The participation ratio is a continuous measure of dimensionality ranging from 1 to K_l . The minimum ($\text{PR}_l = 1$) is obtained when all of the feature variance is concentrated in a single dimension, such that $\lambda_i = 0$ for $i \geq 2$. The maximum ($\text{PR}_l = K_l$) is obtained when the variance is evenly spread across the K_l feature dimensions, such that all K_l eigenvalues are equal.

Mutual information

The activity of each unit in response to the holdout image set was discretized into 10 equally-sized bins, yielding a unit-specific activation distribution $p_X(x)$. Stimulus features (horizontal/vertical position, layout, target/flanker direction) were discretized if the values were continuous, or used as is if not, yielding a stimulus feature distribution $p_Y(y)$. Discretization of the continuous

variables was done as described in the decoding methods section. The joint probability mass function of the unit activity and stimulus feature is denoted by $p_{(X,Y)}(x,y)$. For a given unit and stimulus feature, the mutual information is given by:

$$I(X;Y) = \sum_{x,y} p_{(X,Y)}(x,y) \log \frac{p_{(X,Y)}(x,y)}{p_X(x)p_Y(y)}.$$

The mutual information for a given stimulus feature was normalized by the entropy of the feature distribution to facilitate comparisons between the features [47]. The entropy is given by:

$$H(Y) = - \sum_y p_Y(y) \log p_Y(y).$$

Single-unit modulation by target direction

To identify units that were modulated by target direction, we ran a one-way ANOVA on the z-scored activity of each unit, where each group in the ANOVA was determined by the activity of that unit in response to stimuli for a given target direction (the activity of each unit was z-scored across the stimuli). Units with an ANOVA p -value < 0.001 were defined as significantly modulated by target direction. These units were further split into three subtypes based on their degree of selectivity and sign of modulation: selective (+), selective (-), and complex units.

We first selected units that had significantly higher or lower activation for one direction relative to the other three directions, as assessed by Tukey's HSD test ($p < 0.05$). Within this population, some units had both significantly higher activation for one direction relative to the other three *and* significantly lower activation for one direction relative to the other three. Units that did vs. did not have this property were handled separately. In the former (simpler) case, the units that only had significantly higher activation for one direction relative to the other three were defined as selective (+) units; the units that only had significantly lower activation for one direction relative to the other three were defined as selective (-) units.

In the latter (more complicated) case, we compared the magnitude of the activation for the positive and negative modulation directions with a rank-sum test. Units for which the magnitude of activity for the positive modulation direction was significantly greater ($p < 0.05$) than the magnitude of activity for the negative modulation direction were defined as selective (+) units. Analogous criteria were used to define selective (-) units (with the signs reversed). Units within this pool with rank-sum p -value > 0.05 were defined as complex units. Finally, units that were significantly modulated by target direction but that did not meet any of the criteria described above were also defined as complex units.

Data and code availability

All of the code (<https://github.com/pauljaffe/vam>) and data (<https://doi.org/10.5281/zenodo.10775514>) used to train the VAMs and reproduce our results are publicly-available without restrictions.

Acknowledgements

We thank M. Steyvers, M. Robinson, D. Yamins and members of the NeuroAILab, the Lumos Labs research team, and members of the Poldrack Lab for helpful conversations and comments on this work. We also thank A. Kaluszka for providing graphics files of the Lost and Migration stimuli. The authors received no specific funding for this work.

Author contributions

P.I.J. designed research, performed research, and wrote the paper. P.I.J. and G.X.S.R. analyzed data. P.I.J., G.X.S.R., R.J.S., P.G.B., and R.A.P. edited the paper. R.J.S. and R.A.P. provided resources. R.J.S., P.G.B., and R.A.P. supervised research and provided input at all stages.

Competing interests

R.J.S. is employed by Lumos Labs, Inc. and owns stock in the company. P.I.J., G.X.S.R., P.G.B., and R.A.P. have no competing interests.

Supporting Information

References

- [1] Ansuini A., Laio A., Macke J. H., Zoccolan D. (2019) **Intrinsic dimension of data representations in deep neural networks** *Adv. Neural Inf. Process. Syst* **32**
- [2] Baker N., Lu H., Erlikhman G., Kellman P. J. (2018) **Deep convolutional networks do not classify based on global object shape** *PLoS Comput. Biol* **14**
- [3] Ben-David B. M., Eidels A., Donkin C. (2014) **Effects of Aging and Distractors on Detection of Redundant Visual Targets and Capacity: Do Older Adults Integrate Visual Targets Differently than Younger Adults?** *PLoS One* **9**
- [4] Bernardi S., Benna M. K., Rigotti M., Munuera J., Fusi S., Salzman C. D. (2020) **The geometry of abstraction in the hippocampus and prefrontal cortex** *Cell* **183**:954–967
- [5] Bowers J. S., Malhotra G., Dujmovic M., et al. (2022) **Deep Problems with Neural Network Models of Human Vision** *Behav. Brain Sci* :1–74
- [6] Bradbury J., Frostig R., Hawkins P., et al. (2018) **JAX: Composable transformations of Python+NumPy programs**
- [7] Brincat S. L., Siegel M., Nicolai C. v., Miller E. K. (2018) **Gradual progression from sensory to task-related processing in cerebral cortex** *Proc. Natl. Acad. Sci. U. S. A* **115**

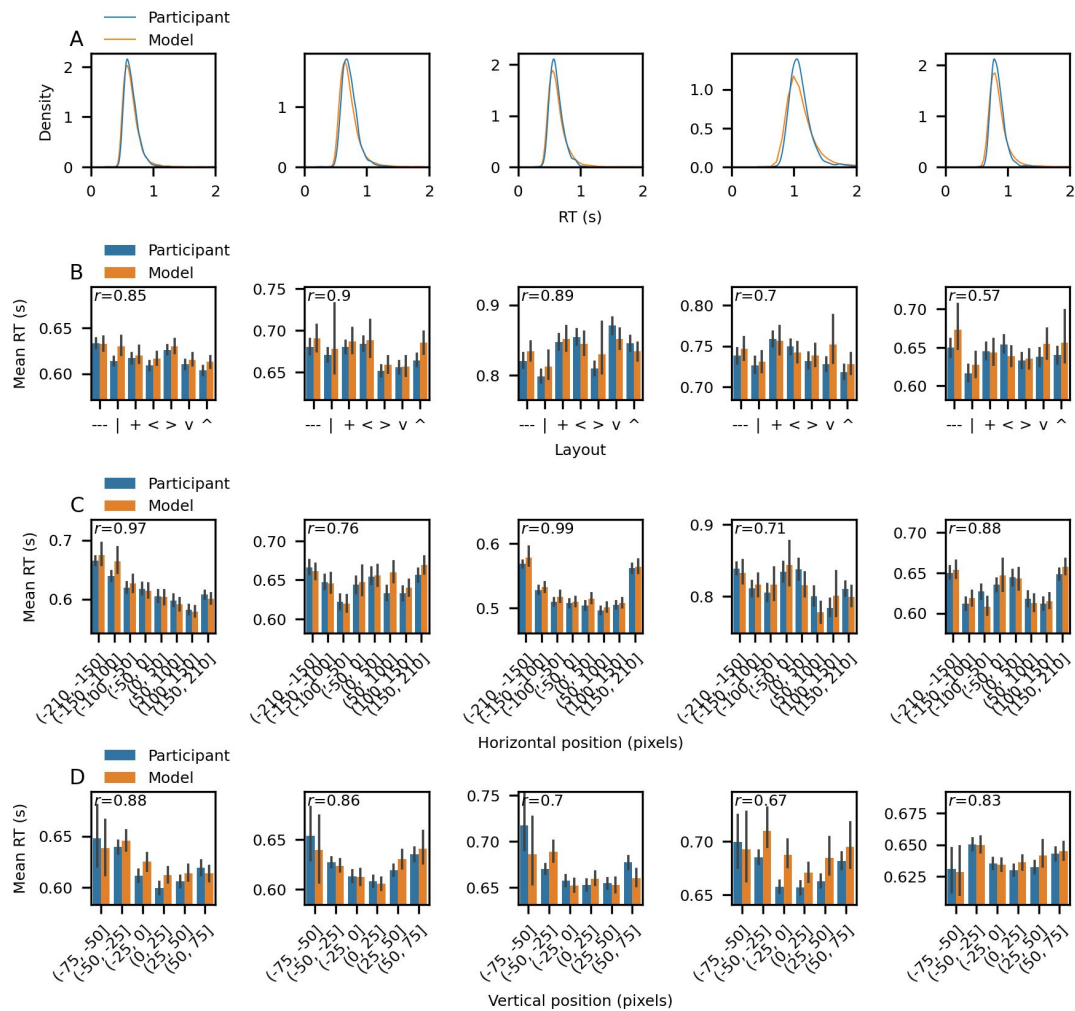


Fig. S1

Example model/participant RT distributions and dependence of RTs on stimulus features.

A) Example model/participant RT distributions (all trials). **B)** Examples of model/participant mean RT vs. stimulus layout. **C)** Examples of model/participant mean RT vs. horizontal stimulus position (negative values: left of center). **D)** Examples of model/participant mean RT vs. vertical stimulus position (negative values: above center). For all panels, error bars correspond to bootstrap 95% confidence intervals.

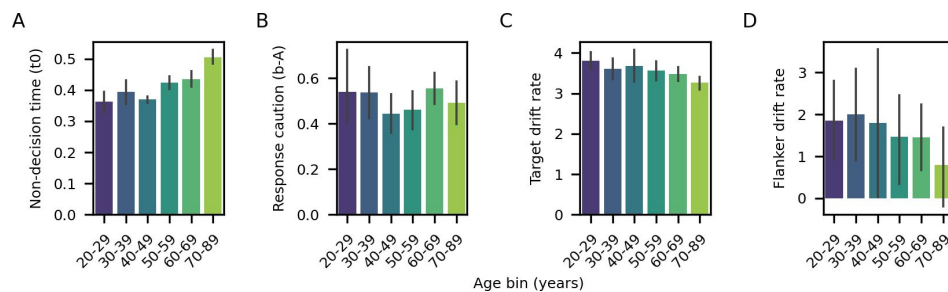


Fig. S2

Age dependence of LBA parameters.

For all panels, we tested age-dependence with a oneway ANOVA and report Bonferroni-adjusted p -values, corrected for 4 comparisons ($n = 75$ models). We also report adjusted p -values from a post-hoc comparison of the 20–29 vs. 70–89 age groups conducted with Tukey's HSD. Error bars correspond to bootstrap 95% confidence intervals. **A)** Non-decision time parameter t_0 ($F(5, 69) = 13.3, p < 1e-7$). Tukey's HSD for 20–29 vs. 70–89 age groups: $p < 1e-8$. **B)** Response caution ($b - A$; $F(5, 69) = 0.49, p = 1.0$). **C)** Mean target drift rate ($F(5, 69) = 3.4, p = 0.026$). Tukey's HSD for 20–29 vs. 70–89 age groups: $p = 0.002$. **D)** Mean flanker drift rate ($F(5, 69) = 0.72, p = 1.0$).

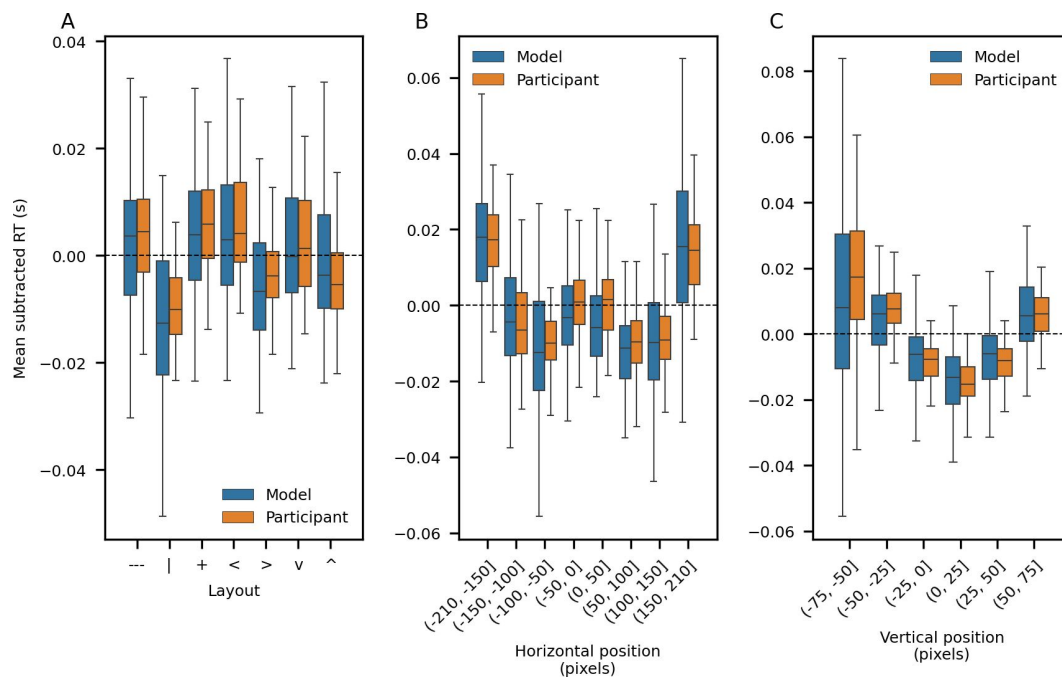


Fig. S3

Dependence of RTs on stimulus layout and position.

For each participant/model, we calculated the mean RT in each stimulus feature bin, then subtracted the average of these mean RTs from each bin. The panels show the average of these centered RTs across all participants with significant modulation of RT for that particular stimulus feature. For all panels, we conducted a one-way ANOVA for both models/participants and report Bonferroni-adjusted p -values, corrected for 3 comparisons ($n = 75$ models). We also report results from post-hoc comparisons between select feature bins conducted with Tukey's HSD. Error bars correspond to bootstrap 95% confidence intervals. **A)** RT vs. stimulus layout (models: $F(6, 53) = 7.43, p < 1e-6$, RTs for the vertical line layout were significantly faster (Tukey's HSD adjusted p -value < 0.05) than RTs from all other layouts except '>'; participants: $F(6, 53) = 23.1, p < 1e-22$; RTs for the vertical line layout were significantly faster than RTs from all other layouts). **B)** RT vs. horizontal stimulus position (negative values: left of center; models: $F(7, 64) = 16.8, p < 1e-18$, RTs for the leftmost and rightmost position bins were significantly slower than RTs from all intermediate position bins; participants: $F(7, 64) = 72.6, p < 1e-73$; RTs for the leftmost and rightmost position bins were significantly slower than RTs from all intermediate position bins). **C)** RT vs. vertical stimulus position (negative values: above center; models: $F(5, 66) = 17.2, p < 1e-14$, RTs for the topmost and bottommost position bins were significantly slower than RTs from the two centermost position bins; participants: $F(5, 66) = 113.3, p < 1e-74$; RTs for the topmost and bottommost position bins were significantly slower than RTs from the two centermost position bins).

Fig. S4

RT delta plots and conditional accuracy functions.

A) RT delta plots for participants and models trained with the standard training paradigm (all RT data; $n = 75$ models/participants). **B)** Conditional accuracy functions for participants and models trained with the standard training paradigm. **C)** RT delta plots for participants and models trained separately on data in each RT quantile, then combined ($n=10$ combined models/participants). **D)** Conditional accuracy functions for participants and models trained separately on data in each RT quantile, then combined. For all panels, error bars correspond to bootstrap 95% confidence intervals.

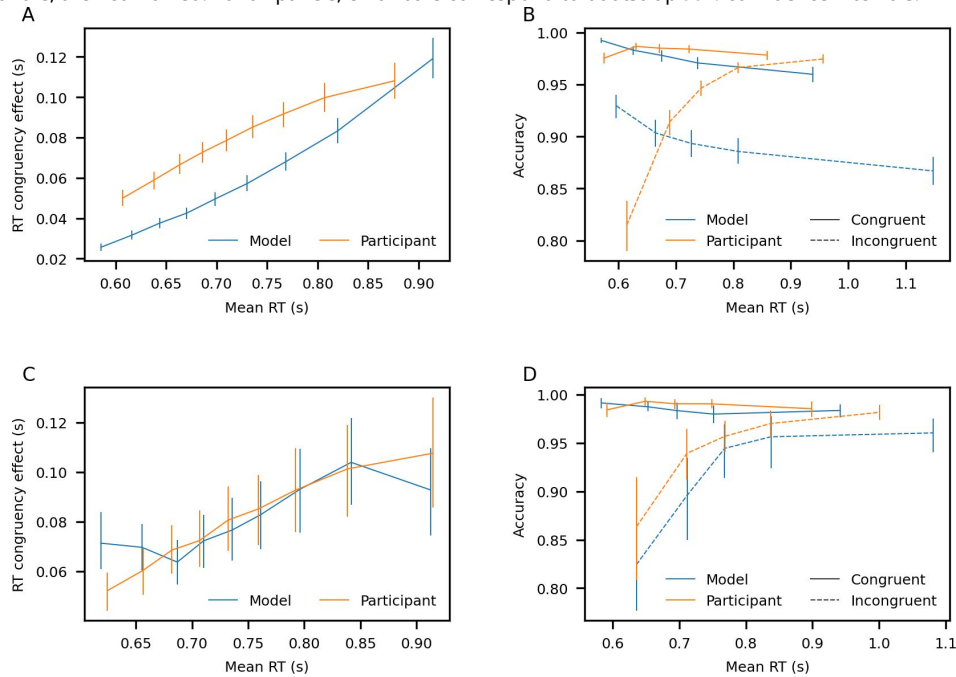


Fig. S5

Activity of all selective (+) units for one example model.

Each row shows the activity of one unit for 100 randomly selected stimuli, sorted by target direction. The activity of each unit was centered and normalized by the activity of the stimulus with the largest magnitude activation. The small number of selective (+) units in layer Conv1 are not shown.

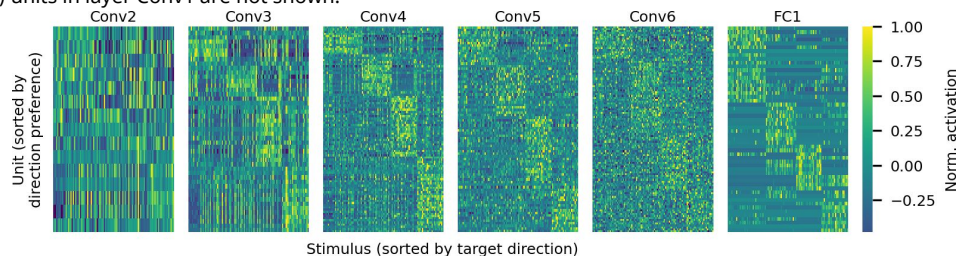


Fig. S6

Absence of correlation between flanker suppression metrics and congruency effects.

All panels show the Pearson's correlation coefficient between the specified suppression and behavior metrics, calculated across models. **A)** Flanker direction decoding accuracy vs. accuracy congruency effect. **B)** Mutual information for flanker direction conveyed by single units vs. accuracy congruency effect. **C)** Flanker direction decoding accuracy vs. RT congruency effect. **D)** Mutual information for flanker direction conveyed by single units vs. RT congruency effect. For all panels, $n = 75$ models, error bars correspond to bootstrap 95% confidence intervals. Asterisks indicate a significant Pearson's r (adjusted p -value < 0.05 , permutation test with $n = 1000$ shuffles, Bonferroni correction for 7 comparisons).

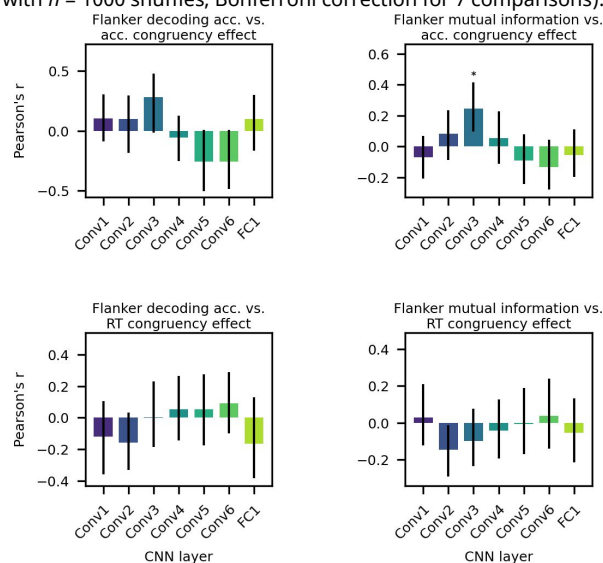
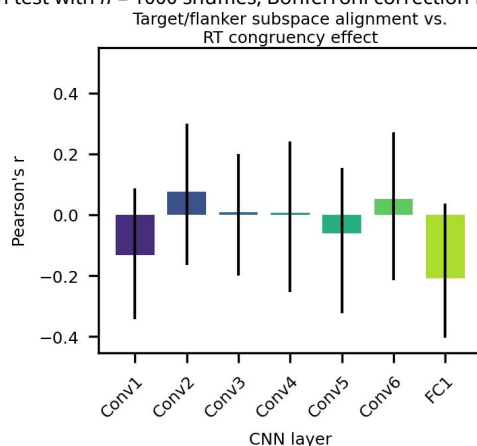


Fig. S7

Absence of correlation between target/flanker subspace alignment and RT congruency effect.

Pearson's correlation coefficient between target/flanker subspace alignment and RT congruency effect across models ($n = 75$ models, error bars correspond to bootstrap 95% confidence intervals). The correlation was not significant for any layer (adjusted p -value > 0.05 , permutation test with $n = 1000$ shuffles, Bonferroni correction for 7 comparisons).



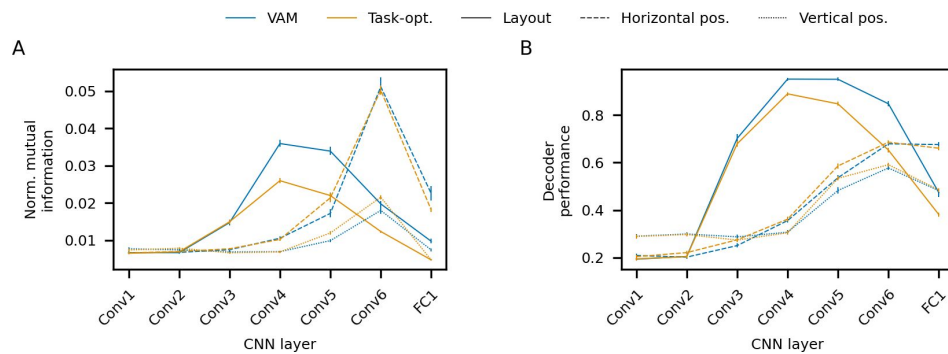


Fig. S8

Additional analysis of VAMs and task-optimized models.

A) Normalized mutual information for stimulus layout and horizontal/vertical stimulus position conveyed by single units, averaged across units. Mutual information was normalized by the entropy of the corresponding stimulus feature distribution.

B) Decoding accuracy of stimulus layout and horizontal/vertical stimulus position. All panels show the average across $n = 75$ task-optimized models and $n = 75$ VAMs; error bars correspond to bootstrap 95% confidence intervals. The VAM data shown in panels A and B is the same as that shown in **Figs. 4C and 4B**, respectively.

- [8] Brown S. D., Heathcote A. (2008) **The simplest complete model of choice response time: Linear ballistic accumulation** *Cogn. Psychol* **57**:153–178
- [9] Cohen J. D., Servan-Schreiber D., McClelland J. L. (1992) **A Parallel Distributed Processing Approach to Automaticity** *Am. J. Psychol* **105**
- [10] Cover T. M. (1965) **Geometrical and Statistical Properties of Systems of Linear Inequalities with Applications in Pattern Recognition** *IEEE Trans. Electron. Comput* :326–334
- [11] Dao V. H., Gunawan D., Tran M.-N., Kohn R., Hawkins G. E., Brown S. D. (2022) **Efficient Selection Between Hierarchical Cognitive Models: Cross-Validation With Variational Bayes** *Psychol. Methods*
- [12] Deng J., Dong W., Socher R., Li L.-J., Li K., Fei-Fei L. (2009) **ImageNet: A large-scale hierarchical image database** *Proc. IEEE Conf. Comput. Vis. Pattern Recognit* :248–255
- [13] Dezfouli A., Griffiths K., Ramos F., Dayan P., Balleine B. W. (2019) **Models that learn how humans learn: The case of decision-making and its disorders** *PLoS Comput. Biol* **15**
- [14] DiCarlo J. J., Zoccolan D., Rust N. C. (2012) **How Does the Brain Solve Visual Object Recognition?** *Neuron* **73**:415–434
- [15] Dosovitskiy A., Beyer L., Kolesnikov A., et al. (2021) **An image is worth 16x16 words: Transformers for image recognition at scale** *International Conference on Learning Representations*
- [16] Eriksen B. A., Eriksen C. W. (1974) **Effects of noise letters upon the identification of a target letter in a nonsearch task** *Percept. Psychophys* **16**:143–149
- [17] Evans N. J., Wagenmakers E.-J. (2020) **Evidence accumulation models: Current limitations and future directions** *Quant. Meth. Psychol* **16**:73–90
- [18] Fel T., Rodriguez Rodriguez I. F., Linsley D., Serre T. (2022) **Harmonizing the object recognition strategies of deep neural networks with humans** *Adv. Neural Inf. Process. Syst* **35**:9432–9446
- [19] Flesch T., Juechems K., Dumbalska T., Saxe A., Summerfield C. (2022) **Orthogonal representations for robust context-dependent task performance in brains and neural networks** *Neuron* **110**:1258–1270
- [20] Forstmann B. U., Tittgemeyer M., Wagenmakers E.-J., Derrfuss J., Imperati D., Brown S. (2011) **The Speed-Accuracy Tradeoff in the Elderly Brain: A Structural Model-Based Approach** *J. Neurosci* **31**:17–17
- [21] Gao P., Trautmann E., Yu B., et al. (2017) **A theory of multineuronal dimensionality, dynamics and measurement** <https://doi.org/10.1101/214262>
- [22] Goetschalckx L., Govindarajan L. N., Ashok A. Karkada, Ahuja A., Sheinberg D., Serre T. (2023) **Computing a human-like reaction time metric from stable recurrent vision models** *Adv. Neural Inf. Process. Syst* **36**:14–14
- [23] Gottsdanker R. (1982) **Age and Simple Reaction Time** *J. Gerontol* **37**:342–348

- [24] Güçlü U., van Gerven M. A. J. (2015) **Deep neural networks reveal a gradient in the complexity of neural representations across the ventral stream** *J. Neurosci* **35**:10005–10014
- [25] Gunawan D., Hawkins G., Tran M.-N., Kohn R., Brown S. (2020) **New estimation approaches for the hierarchical linear ballistic accumulator model** *J. Math. Psychol* **96**
- [26] Heidler K. (2022) **Augmax**
- [27] Hohman F., Park H., Robinson C., Polo Chau D. H. (2020) **Summit: Scaling deep learning interpretability by visualizing activation and attribution summarizations** *IEEE Trans. Vis. Comput. Graph* **26**:1096–1106
- [28] Hung C. P., Kreiman G., Poggio T., DiCarlo J. J. (2005) **Fast Readout of Object Identity from Macaque Inferior Temporal Cortex** *Science* **310**:863–866
- [29] Jaffe P. I., Poldrack R. A., Schafer R. J., Bissett P. G. (2023) **Modelling human behaviour in cognitive tasks with latent dynamical systems** *Nat. Hum. Behav* **7**:986–1000
- [30] Jong R. D., Liang C.-C., Lauber E. (1994) **Conditional and Unconditional Automaticity: A Dual-Process Model of Effects of Spatial Stimulus-Response Correspondence** *J. Exp. Psychol. Hum. Percept. Perform* **20**:731–750
- [31] Jordan C. (1875) **Essai sur la géométrie à n dimensions** *Bulletin de la Société Mathématique de France* **3**:103–174
- [32] Kaufman M. T., Churchland M. M., Ryu S. I., Shenoy K. V. (2014) **Cortical activity in the null space: permitting preparation without movement** *Nat. Neurosci* **17**:440–448
- [33] Kingma D. P., Welling M. (2013) **Auto-Encoding Variational Bayes** *arXiv*
- [34] Kingma D. P., Ba J. (2017) **Adam: A method for stochastic optimization** *arXiv*
- [35] Kingma D. P., Salimans T., Welling M. (2015) **Variational dropout and the local reparameterization trick** *arXiv*
- [36] Klambauer G., Unterthiner T., Mayr A., Hochreiter S. (2017) **Self-normalizing neural networks** *arXiv*
- [37] Kriegeskorte N. (2015) **Deep neural networks: A new framework for modeling biological vision and brain information processing** *Annu. Rev. Vis. Sci* **1**:417–446
- [38] Kucukelbir A., Tran D., Ranganath R., Gelman A., Blei D. M. (2017) **Automatic differentiation variational inference** *J. Mach. Learn. Res* **18**:1–45
- [39] Kumbhar O., Sizikova E., Majaj N., Pelli D. G. (2020) **Anytime Prediction as a Model of Human Reaction Time** *arXiv*
- [40] Libby A., Buschman T. J. (2021) **Rotational dynamics reduce interference between sensory and memory representations** *Nat. Neurosci* **24**:715–726
- [41] Lindsay G. W. (2021) **Convolutional Neural Networks as a Model of the Visual System: Past, Present, and Future** *J. Cogn. Neurosci* **33**:2017–2031

- [42] Linsley D., Eberhardt S., Sharma T., Gupta P., Serre T. (2017) **What are the Visual Features Underlying Human Versus Machine Vision?** *2017 IEEE International Conference on Computer Vision Workshops (ICCVW)* :2706–2714
- [43] Lo C.-F., Ip H.-Y. (2021) **Modified leaky competing accumulator model of decision making with multiple alternatives: the Lie-algebraic approach** *Sci. Rep* **11**
- [44] Malhotra G., Dujmovic M., Bowers J. S. (2022) **Feature blindness: A challenge for understanding and modelling visual object recognition** *PLoS Comput. Biol* **18**
- [45] Mante V., Sussillo D., Shenoy K. V., Newsome W. T. (2013) **Context-dependent computation by recurrent dynamics in prefrontal cortex** *Nature* **503**:78–84
- [46] Meister M. L. R., Hennig J. A., Huk A. C. (2013) **Signal Multiplexing and Single-Neuron Computations in Lateral Intraparietal Area During Decision-Making** *J. Neurosci* **33**:2254–2267
- [47] Muratore P., Tafazoli S., Piasini E., Laio A., Zoccolan D. (2022) **Prune and distill: Similar reformatting of image information along rat visual cortex and deep neural networks** *Adv. Neural Inf. Process. Syst*
- [48] Navarro D. J., Fuss I. G. (2009) **Fast and accurate calculations for first-passage times in wiener diffusion models** *J. Math. Psychol* **53**:222–230
- [49] Nayeibi A., Bear D., Kubilius J., et al. (2018) **Task-Driven Convolutional Recurrent Models of the Visual System** *arXiv*
- [50] Nettelbeck T., Rabbitt P. M. (1992) **Aging, cognitive performance, and mental speed** *Intelligence* **16**:189–205
- [51] Pagan M., Urban L. S., Wohl M. P., Rust N. C. (2013) **Signals in inferotemporal and perirhinal cortex suggest an untangling of visual target information** *Nat. Neurosci* **16**:1132–1139
- [52] Panichello M. F., Buschman T. J. (2021) **Shared mechanisms underlie the control of working memory and attention** *Nature* **592**:601–605
- [53] Papayan V., Han X. Y., Donoho D. L. (2020) **Prevalence of neural collapse during the terminal phase of deep learning training** *Proc. Natl. Acad. Sci. U. S. A* **117**:24–24
- [54] Pedregosa F., Varoquaux G., Gramfort A., et al. (2011) **Scikit-learn: Machine learning in Python** *J. Mach. Learn. Res* **12**:2825–2830
- [55] Rafiei F., Rahnev D. (2022) **RTNet: A neural network that exhibits the signatures of human perceptual decision making** <https://doi.org/10.1101/2022.08.23.505015>
- [56] Rangamani A., Lindegaard M., Galanti T., Poggio T. A. (2023) **Feature learning in deep classifiers through Intermediate Neural Collapse** *Proc. Mach. Learn. Res* **202**:28–28
- [57] Ratcliff R., Thapar A., McKoon G. (2001) **The effects of aging on reaction time in a signal detection task** *Psychol. Aging* **16**
- [58] Ratcliff R. (1978) **A theory of memory retrieval** *Psychol. Rev* **85**:59–108

- [59] Ratcliff R., McKoon G. (2008) **The Diffusion Decision Model: Theory and Data for Two-Choice Decision Tasks** *Neural Comput* **20**:873–922
- [60] Ratcliff R., Rouder J. N. (1998) **Modeling response times for two-choice decisions** *Psychol. Sci* **9**:347–356
- [61] Rezende D. J., Mohamed S., Wierstra D. (2014) **Stochastic Backpropagation and Approximate Inference in Deep Generative Models** *arXiv*
- [62] Ridderinkhof K. R. (2002) **Activation and suppression in conflict tasks: Empirical clarification through distributional analyses** *Common Mechanisms in Perception and Action: Attention and Performance XIX*
- [63] Ridderinkhof R. K. (2002) **Micro- and macro-adjustments of task set: Activation and suppression in conflict tasks** *Psychol. Res* **66**:312–323
- [64] Rigotti M., Barak O., Warden M. R., et al. (2013) **The importance of mixed selectivity in complex cognitive tasks** *Nature* **497**:585–590
- [65] Ritz H., Shenhav A. (2024) **Orthogonal neural encoding of targets and distractors supports multivariate cognitive control** *Nat. Hum. Behav* :1–17
- [66] Rust N. C., DiCarlo J. J. (2010) **Selectivity and Tolerance (“Invariance”) Both Increase as Visual Information Propagates from Cortical Area V4 to IT** *J. Neurosci* **30**:12–12
- [67] Servant M., Evans N. J. (2020) **A Diffusion Model Analysis of the Effects of Aging in the Flanker Task** *Psychol. Aging* **35**:831–849
- [68] Simard P., Steinkraus D., Platt J. (2003) **Best practices for convolutional neural networks applied to visual document analysis** *Seventh International Conference on Document Analysis and Recognition, 2003. Proceedings* :958–963
- [69] Simon J. (1982) **Effect of an auditory stimulus on the processing of a visual stimulus under single- and dual-tasks conditions** *Acta Psychol* **51**:61–73
- [70] Simonyan K., Zisserman A. (2015) **Very deep convolutional networks for large-scale image recognition** *arXiv*
- [71] Spoerer C. J., Kietzmann T. C., Mehrer J., Charest I., Kriegeskorte N. (2020) **Recurrent neural networks can explain flexible trading of speed and accuracy in biological vision** *PLoS Comput. Biol* **16**
- [72] Steyvers M., Hawkins G. E., Karayanidis F., Brown S. D. (2019) **A large-scale analysis of task switching practice effects across the lifespan** *Proc. Natl. Acad. Sci. U. S. A* **116**:17–17
- [73] Stoffels E. J., Molen M. W. v. d. (1988) **Effects of visual and auditory noise on visual choice reaction time in a continuous-flow paradigm** *Percept. Psychophys* **44**:7–14
- [74] Stroop J. R. (1935) **Studies of interference in serial verbal reactions** *J. Exp. Psychol* **18**:643–662
- [75] Sussillo D., Churchland M. M., Kaufman M. T., Shenoy K. V. (2015) **A neural network that finds a naturalistic solution for the production of muscle activity** *Nat. Neurosci* **18**:1025–1033

- [76] Tafazoli S., Safaai H., De Franceschi G., et al. (2017) **Emergence of transformation-tolerant representations of visual objects in rat lateral extrastriate cortex** *eLife* 6
- [77] Taylor J. E. T., Shekhar S., Taylor G. W. (2021) **Neural response time analysis: Explainable artificial intelligence using only a stopwatch** *Appl. AI Lett* 2
- [78] Ulrich R., Schröter H., Leuthold H., Birngruber T. (2015) **Automatic and controlled stimulus processing in conflict tasks: Superimposed diffusion processes and delta functions** *Cogn. Psychol* 78:148–174
- [79] Ulyanov D., Vedaldi A., Lempitsky V. (2017) **Instance normalization: The missing ingredient for fast stylization** *arXiv*
- [80] Usher M., McClelland J. L. (2001) **The Time Course of Perceptual Choice: The Leaky, Competing Accumulator Model** *Psychol. Rev* 108:550–592
- [81] Wang J., Narain D., Hosseini E. A., Jazayeri M. (2018) **Flexible timing by temporal scaling of cortical responses** *Nat. Neurosci* 21:102–110
- [82] White C. N., Ratcliff R., Starns J. J. (2011) **Diffusion models of the flanker task: Discrete versus gradual attentional selection** *Cogn. Psychol* 63:210–238
- [83] Wildenberg W. P. v. d., Wylie S. A., Forstmann B. U., Burle B., Hasbroucq T., Ridderinkhof K. R. (2010) **To Head or to Heed? Beyond the Surface of Selective Action Inhibition: A Review** *Front. Hum. Neurosci* 4
- [84] Xie Y., Hu P., Li J., et al. (2022) **Geometry of sequence working memory in macaque prefrontal cortex** *Science* 375:632–639
- [85] Yamins D. L. K., DiCarlo J. J. (2016) **Using goal-driven deep learning models to understand sensory cortex** *Nat. Neurosci* 19:356–365
- [86] 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** *Proc. Natl. Acad. Sci. U. S. A* 111:8619–8624
- [87] Yang G. R., Joglekar M. R., Song H. F., Newsome W. T., Wang X.-J. (2019) **Task representations in neural networks trained to perform many cognitive tasks** *Nat. Neurosci* 22:297–306
- [88] Zhu P., Knyazev A. (2013) **Angles between subspaces and their tangents** *J. Numer. Math* 21

Editors

Reviewing Editor

Marius Peelen

Radboud University Nijmegen, Nijmegen, Netherlands

Senior Editor

Michael Frank

Brown University, Providence, United States of America

Reviewer #1 (Public Review):

Summary:

This paper introduces a new approach to modeling human behavioral responses using image-computable models. They create a model (VAM) that is a combination of a standard CNN coupled with a standard evidence accumulation model (EAM). The combined model is then trained directly on image-level data using human behavioral responses. This approach is original and can have wide applicability. However, many of the specific findings reported are less compelling.

Strengths:

- (1) The manuscript presents an original approach to fitting an image-computable model to human behavioral data. This type of approach is sorely needed in the field.
- (2) The analyses are very technically sophisticated.
- (3) The behavioral data are large both in terms of sample size ($N=75$) and in terms of trials per subject.

Weaknesses:

Major

- (1) The manuscript appears to suggest that it is the first to combine CNNs with evidence accumulation models (EAMs). However, this was done in a 2022 preprint (<https://www.biorxiv.org/content/10.1101/2022.08.23.505015v1>) that introduced a network called RTNet. This preprint is cited here, but never really discussed. Further, the two unique features of the current approach discussed in lines 55-60 are both present to some extent in RTNet. Given the strong conceptual similarity in approach, it seems that a detailed discussion of similarities and differences (of which there are many) should feature in the Introduction.
- (2) In the approach here, a given stimulus is always processed in the same way through the core CNN to produce activations v_k . These v_k 's are then corrupted by Gaussian noise to produce drift rates d_k , which can differ from trial to trial even for the same stimulus. In other words, the assumption built into VAM appears to be that the drift rate variability stems entirely from post-sensory (decisional) noise. In contrast, the typical interpretation of EAMs is that the variability in drift rates is sensory. This is also the assumption built into RTNet where the core CNN produces noisy evidence. Can the authors comment on the plausibility of VAM's assumption that the noise is post-sensory?
- (3) Figure 2 plots how well VAM explains different behavioral features. It would be very useful if the authors could also fit simple EAMs to the data to clarify which of these features are explainable by EAMs only and which are not.
- (4) VAM is tested in two different ways behaviorally. First, it is tested to what extent it captures individual differences (Figure 2B-E). Second, it is tested to what extent it captures average subject data (Figure 2F-J). It wasn't clear to me why for some metrics only individual differences are examined and for other metrics only average human data is examined. I think that it will be much more informative if separate figures examine average human data and individual difference data. I think that it's especially important to clarify whether VAM can capture individual differences for the quantities plotted in Figures 2F-J.
- (5) The authors look inside VAM and perform many exploratory analyses. I found many of these difficult to follow since there was little guidance about why each analysis was conducted. This also made it difficult to assess the likelihood that any given result is robust and replicable. More importantly, it was unclear which results are hypothesized to depend on the VAM architecture and training, and which results would be expected in performance-optimized CNNs. The authors train and examine performance-optimized CNNs later, but it would be useful to compare those results to the VAM results immediately when each VAM result is first introduced.

(6) The authors don't examine how the task-optimized models would produce RTs. They say in lines 371-2 that they "could not examine the RT congruency effect since the task-optimized models do not generate RTs." CNNs alone don't generate RTs, but RTs can easily be generated from them using the same EAM add-on that is part of VAM. Given that the CNNs are already trained, I can't see a reason why the authors can't train EAMs on top of the already trained CNNs and generate RTs, so these can provide a better comparison to VAM.

(7) The Discussion felt very long and mostly a summary of the Results. I also couldn't shake the feeling that it had many just-so stories related to the variety of findings reported. I think that the section should be condensed and the authors should be clearer about which explanations are speculations and which are air-tight arguments based on the data.

(8) In one of the control analyses, the authors train different VAMs on each RT quantile. I don't understand how it can be claimed that this approach can serve as a model of an individual's sensory processing. Which of the 5 sets of weights (5 VAMs) captures a given subject's visual processing? Are the authors saying that the visual system of a given subject changes based on the expected RT for a stimulus? I feel like I'm missing something about how the authors think about these results.

<https://doi.org/10.7554/eLife.98351.1.sa2>

Reviewer #2 (Public Review):

In an image-computable model of speeded decision-making, the authors introduce and fit a combined CCN-EAM (a 'VAM') to flanker-task-like data. They show that the VAM can fit mean RTs and accuracies as well as the congruency effect that is present in the data, and subsequently analyze the VAM in terms of where in the network congruency effects arise.

Overall, combining DNNs and EAMs appears to be a promising avenue to seriously model the visual system in decision-making tasks compared to the current practice in EAMs. Some variants have been proposed or used before (e.g., doi.org/10.1016/j.neuroimage.2017.12.078, doi.org/10.1007/s42113-019-00042-1), but always in the context of using task-trained models, rather than models trained on behavioral data. However, I was surprised to read that the authors developed their model in the context of a conflict task, rather than a simpler perceptual decision-making task. Conflict effects in human behavior are particularly complex, and thereby, the authors set a high goal for themselves in terms of the to-be-explained human behavior. Unfortunately, the proposed VAM does not appear to provide a great account of conflict effects that are considered fundamental features of human behavior, like the shape of response time distributions, and specifically, delta plots (doi.org/10.1037/0096-1523.20.4.731). The authors argue that it is beyond the scope of the presented paper to analyze delta plots, but as these are central to studies of human conflict behavior, models that aim to explain conflict behavior will need to be able to fit and explain delta plots.

Theories on conflict often suggest that negative/positive-trending delta plots arise through the relative timing of response activation related to relevant and irrelevant information. Accumulation for relevant and irrelevant information would, as a result, either start at different points in time or the rates vary over time. The current VAM, as a feedforward neural network model, does not appear to be able to capture such effects, and perhaps fundamentally not so: accumulation for each choice option is forced to start at the same time, and rates are a static output of the CNN.

The proposed solution of fitting five separate VAMs (one for each of five RT quantiles) is not satisfactory: it does not explain how delta plots result from the model, for the same reason

that fitting five evidence accumulation models (one per RT quantile) does not explain how response time distributions arise. If, for example, one would want to make a prediction about someone's response time and choice based on a given stimulus, one would first have to decide which of the five VAMs to use, which is circular. But more importantly, this way of fitting multiple models does not explain the latent mechanism that underlies the shape of the delta plots.

As such, the extensive analyses on the VAM layers and the resulting conclusions that conflict effects arise due to changing representations across layers (e.g., "the selection of task-relevant information occurs through the orthogonalization of relevant and irrelevant representations") - while inspiring, they remain hard to weigh, as they are contingent on the assumption that the VAM can capture human behavior in the conflict task, which it struggles with. That said, the promise of combining CNNs and EAMs is clearly there. A way forward could be to either adjust the proposed model so that it can explain delta plots, which would potentially require temporal dynamics and time-varying evidence accumulation rates, or perhaps to start simpler and combine CNNs-EAMs that are able to fit more standard perceptual decision-making tasks without conflict effects.

<https://doi.org/10.7554/eLife.98351.1.sa1>

Reviewer #3 (Public Review):

Summary:

In this article, the authors combine a well-established choice-response time (RT) model (the Linear Ballistic Accumulator) with a CNN model of visual processing to model image-based decisions (referred to as the Visual Accumulator Model - VAM). While this is not the first effort to combine these modeling frameworks, it uses this combination of approaches uniquely. Specifically, the authors attempt to better understand the structure of human information representations by fitting this model to behavioral (choice-RT) data from a classic flanker task. This objective is made possible by using a very large (by psychological modeling standards) industry data set to jointly fit both components of this VAM model to individual-level data. Using this approach, they illustrate (among other results) (1) how the interaction between target and flanker representations influence the presence and strength of congruency effects, (2) how the structure of representations changes (distributed versus more localized) with depth in the CNN model component, and (3) how different model training paradigms change the nature of information representations. This work contributes to the ML literature by demonstrating the value of training models with richer behavioral data. It also contributes to cognitive science by demonstrating how ML approaches can be integrated into cognitive modeling. Finally, it contributes to the literature on conflict modeling by illustrating how information representations may lead to some of the classic effects observed in this area of research.

Strengths:

(1) The data set used for this analysis is unique and is made publicly available as part of this article. Specifically, they have access to data for 75 participants with >25,000 trials per participant. This scale of data/individual is unusual and is the foundation on which this research rests.

(2) This is the first time, to my knowledge, that a model combining a CNN with a choice-RT model has been jointly fit to choice-RT data at the level of individual people. This type of model combination has been used before but in a more restricted context. This joint fitting, and in particular, learning a CNN through the choice-RT modeling framework, allows the

authors to probe the structure of human information representations learned directly from behavioral data.

(3) The analysis approaches used in this article are state-of-the-art. The training of these models is straightforward given the data available. The interesting part of this article (opinion of course) is the way in which they probe what CNN has learned once trained. I find their analysis of how distractor and target information interfere with each other particularly compelling as well as their demonstration that training on behavioral data changes the structure of information representations when compared to training models on standard task-optimized data.

Weaknesses:

(1) Just as the data in this article is a major strength, it is also a weakness. This type of modeling would be difficult, if not impossible to do with standard laboratory data. I don't know what the data floor would be, but collecting tens of thousands of decisions for a single person is impractical in most contexts. Thus this type of work may live in the realm of industry. I do want to re-iterate that the data for this study was made publicly available though!

1. While this article uses choice-RT data it doesn't fully leverage the richness of the RT data itself. As the authors point out, this modeling framework, the LBA component in particular, does not account for some of the more nuanced but well-established RT effects in this data. This is not a big concern given the already nice contributions of this article and it leads to an opportunity for ongoing investigation.

<https://doi.org/10.7554/eLife.98351.1.sa0>