

Supplementary Information

We provide supplementary materials accompanying our paper, which include seven additional figures. Figure S1 demonstrates a drop in task-relevance in the scenario where a target is statistically indistinguishable from a distractor. Figure S2 provides additional analysis that supports modeling decisions, and Figure S3 visualizes model covariates. Figure S4 reports the results of using inferences from the adaptive computation instead of the fixed-resource architecture to compute heuristic covariates in the probes experiment. Figures S5 and S6 visualize the results of grid search over the load and importance hyperparameters, respectively.

Additional Analysis of Study 1

Here we report supplemental analyses for the probe detection study.

Using Adaptive Computation Inferences to Compute Heuristic Probe Covariates

Instead of using inferences from the fixed-resource particle filter to compute heuristic model covariates, here we report the explained variance (R^2) when using inferences from the adaptive computation architecture: Target center $R^2 = .61[.50, .70], p < .001$; distance to origin $R^2 = .35[.14, .36], p < .001$; nearest distractor $R^2 = .25[.15, .36], p < .001$. These are largely similar to those reported in the main text.

For completeness, the rest of the model comparisons are included below (Figure S4). In pairwise comparisons, adaptive computation surpassed the amount of variance explained by these accounts (direct bootstrapped hypothesis testing of attention centroid - target center: $\delta R^2 = .10[.02, .17], p = .005$, distance to origin: $\delta R^2 = .34[.24, .44], p < .001$, and nearest distractor: $\delta R^2 = .43[.31, .54], p < .001$). Importantly, the explanatory power of these domain-specific heuristics was subsumed by our model. Adaptive computation explained a significant portion of variance after residualizing each heuristic in turn (target center: $R^2 = .27[.15, .39]$, bootstrapped one-tailed hypothesis testing $p < .001$; distance to origin: $R^2 = .53[.43, .63], p < .001$; nearest distractor: $R^2 = .61[.51, .70], p < .001$). In stark contrast, however, the explanatory power of each heuristic greatly diminished after residualizing shared variance with adaptive computation (target center: $R^2 = .04[.01, .11], p = .023$; distance to origin $R^2 = .01, p = .70$; nearest distractor $R^2 = .03[.01, .09], p = .033$), with the strength of the residualized explained variance by our model surpassing the residualized explained variance by each heuristic (resid. attention centroid - target center: $R^2 = .23[.06, .39], p = .005$; origin: $R^2 = .26[.14, .38], p < .001$; and nearest distractor: $R^2 = .25[.08, .38], p = .002$, pairwise direct bootstrapped hypothesis testing comparisons).

Considering Only Correct Tracking Trials

Here we report the univariate fits of all models and their bootstrapped comparisons when considering probe detection hit rates from responses where subjects perfectly tracked all four targets.

Attention centroid $R^2 = .54[.44, .63]$. Target center $R^2 = .42[.31, .53], p < .001$; distance to origin $R^2 = .24[.13, .34], p < .001$; nearest distractor $R^2 = .15[.07, .25], p < .001$. While these are lower values than those reported in the main text, the overall rank-order remains the same.

In pairwise comparisons, adaptive computation surpassed the amount of variance explained by these accounts (direct bootstrapped hypothesis testing of attention centroid - target center: $\delta R^2 = .12[.04, .19], p = .001$, distance to origin: $\delta R^2 = .31[.21, .40], p < .001$, and nearest distractor: $\delta R^2 = .38[.27, .48], p < .001$). Importantly, the explanatory power of these domain-specific heuristics was subsumed by our model. Adaptive computation explained a significant portion of variance after residualizing each heuristic in turn (target center: $R^2 = .21[.11, .32]$, bootstrapped one-tailed hypothesis testing $p < .001$; distance to origin: $R^2 = .40[.29, .51], p < .001$; nearest distractor: $R^2 = .45[.35, .55], p < .001$). In stark contrast, however, the explanatory power of each heuristic greatly diminished after residualizing shared variance with adaptive computation (target center: $R^2 = .01[.01, .05], p = .292$; distance to origin $R^2 = .01[.00, .04], p = .724$; nearest distractor $R^2 = .01[.00, .03], p = .503$), with the strength of the residualized explained variance by our model surpassing the residualized explained variance by each heuristic (resid. attention centroid - target center: $R^2 = .20[.08, .31], p = .001$; origin: $R^2 = .20[.10, .31], p < .001$; and nearest distractor: $R^2 = .20[.10, .31], p < .001$, pairwise direct bootstrapped hypothesis testing comparisons).

Additional Heuristic Models

Here we consider several heuristic models in addition to those in the main text. We consider a new univariate covariate, the “distractor center”, which is the average distance from the probed target to each target. While this covariate explained a significant amount of variance in human probe detection rates ($R^2 = .36[.24, .45], p < .001$), the attention centroid provided an overall stronger explanation ($\delta R^2 = .34[.24, .45], p < .001$).

Load bin	Model	Model fit		Comparisons to attention centroid	
		R^2	p-value	δR^2	p-value
Low	Attention centroid	.67	< .001	-	-
Low	Origin + Target center	.63	< .001	.00	.460
Low	Nearest distractor + Target center	.66	< .001	.00	.499
Low	Distractor center + Target center	.66	< .001	.02	.306
Moderate	Attention centroid	.74	< .001	-	-
Moderate	Origin + Target center	.64	< .001	.02	.326
Moderate	Nearest distractor + Target center	.70	< .001	.05	.213
Moderate	Distractor center + Target center	.68	< .001	.08	.094
High	Attention centroid	.70	< .001	-	-
High	Origin + Target center	.53	< .001	.13	.027
High	Nearest distractor + Target center	.57	< .001	.08	.045
High	Distractor center + Target center	.62	< .001	.16	.013

Table S1: Comparison of “composite” heuristic models across load bins.

Supplementary Figures

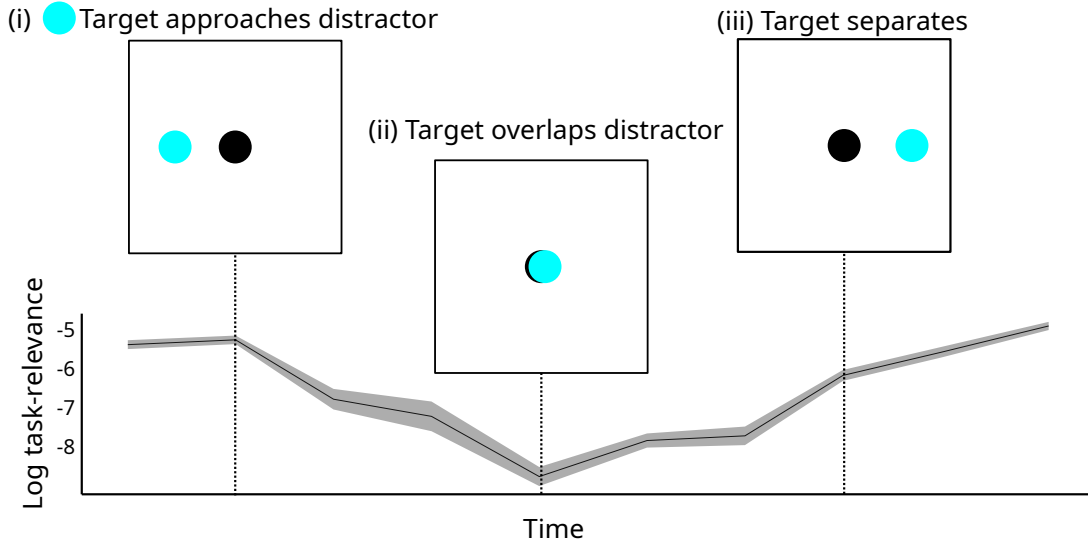


Figure S1: **Demonstrating drop in task-relevance during complete target-distractor overlap.** The adaptive computation architecture was run on this demo scene (for 10 independent runs) with a target moving from left to right, and a stationary distractor at the center. As the target (s_k) completely overlaps the distractor (panel ii, in the center of the scene), adaptive computation calculates a drop in task-relevance, ∇_k . This is primarily due to a drop in $\delta_k \pi$, since the complete overlap causes the target probability to become under determined (in this example, there are no past motion features in the observation to help distinguish between the moving target and stationary distractor). In other words, even applying perceptual computations (C_k) the target probability would not vary from 50%. This pattern in task-relevance would allow adaptive computation to “give up” on this target, and prioritize other targets that would benefit from additional processing. All error bars show standard error of the mean. (Note, in all of the behavioral experiments in the paper, objects did not overlap)

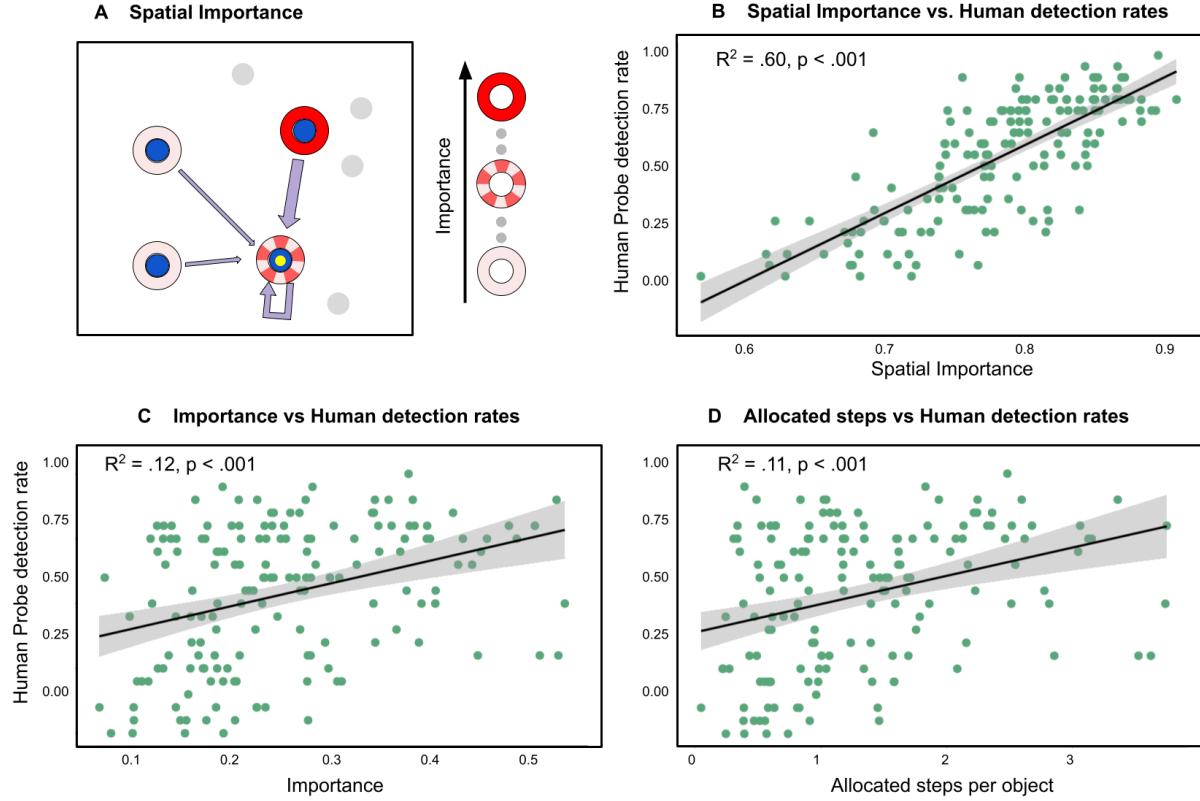


Figure S2: (**A**) An illustration of the multifocal adaptive computation covariate, “spatial importance”, where probe detection rate is proportional to the aggregate importance of nearby objects, with each object’s importance contribution weighted by their distance to the probe. (**B**) Spatial importance explains nearly as much variance as the unimodal attention centroid ($R^2 = .6[.52, .68]$, $N = 160$, where $[l, r]$ correspond to 95% bootstrapped confidence intervals), thus obviating the need for us to commit to a single locus or multifocal spatial deployment pattern. (**C**, **D**) Importance and cycles, which do not incorporate the spatial embedding of the targets, can explain a significant, but rather small amount of variance in probe detection rates ($R^2 = .13[.05, .23]$ and $R^2 = .12[.04, .22]$). Future work should extend adaptive computation with systematic, empirically-informed hypotheses about the mapping between object-level attention and retinotopic deployment of attention.

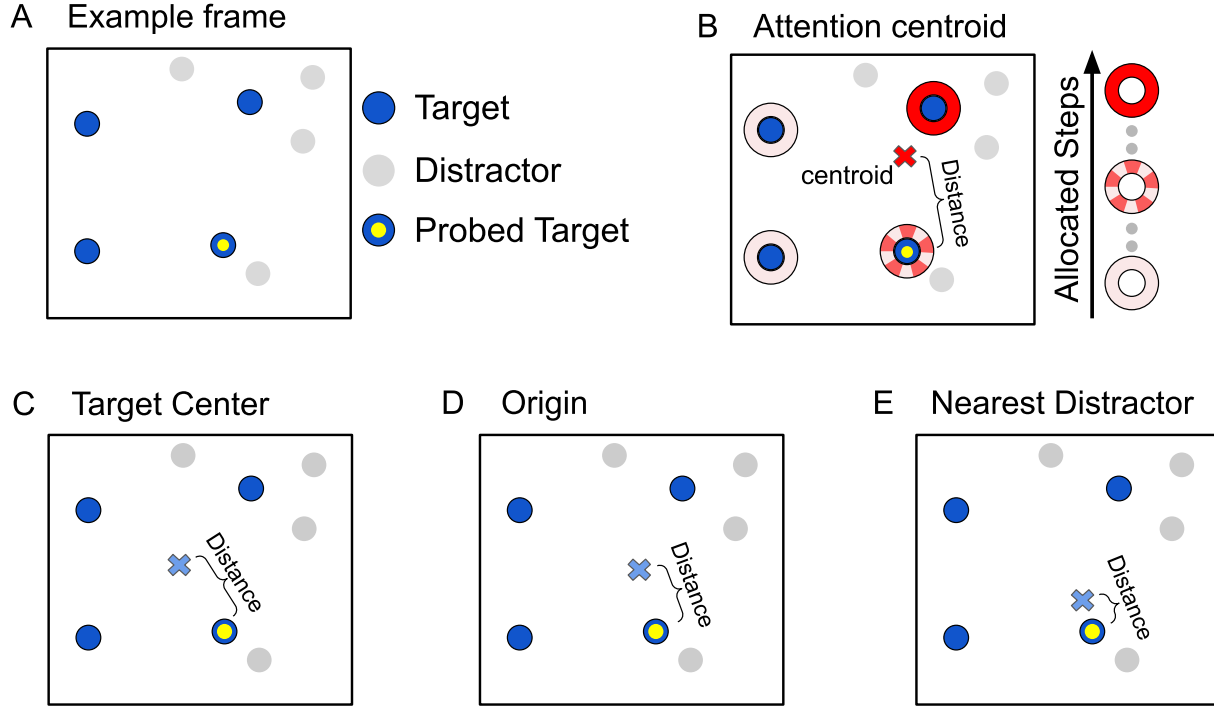


Figure S3: Probe covariates including heuristic models and adaptive computation. (A) An example frame from Exp. 1 (note: the targets and distractors are shown in different colors in this illustration, but for both the behavioral experiment and model simulations, they are visually identical during motion). (B) Illustration of the importance-weighted attention centroid. The target in the top right has the highest sensitivity score due to several nearby distractors, pulling the attention centroid (the red “X”) farther from the probed target (thus predicting a lower probe detection rate) compared to alternatives. (C) The “target center” covariate uses the distances to the unweighted spatial center of targets (the blue “X”). (D) The “origin” covariate uses distances to the constant center of the scene. (E) The “nearest distractor” covariate uses the distances to the weighted spatial center of targets where the weights are based on the distance of each target to its nearest distractor. In this example frame, similar to the target center and origin heuristics, the nearest distractor also overestimates the probe detection rate since the probed target is closer to its nearest distractor compared to the target in the top right. Critically, our general formalization of task-relevance ($\delta\pi \cdot \delta_k S$) automatically takes into account the number and relevance of nearby distractors given the planning objective (target designation).

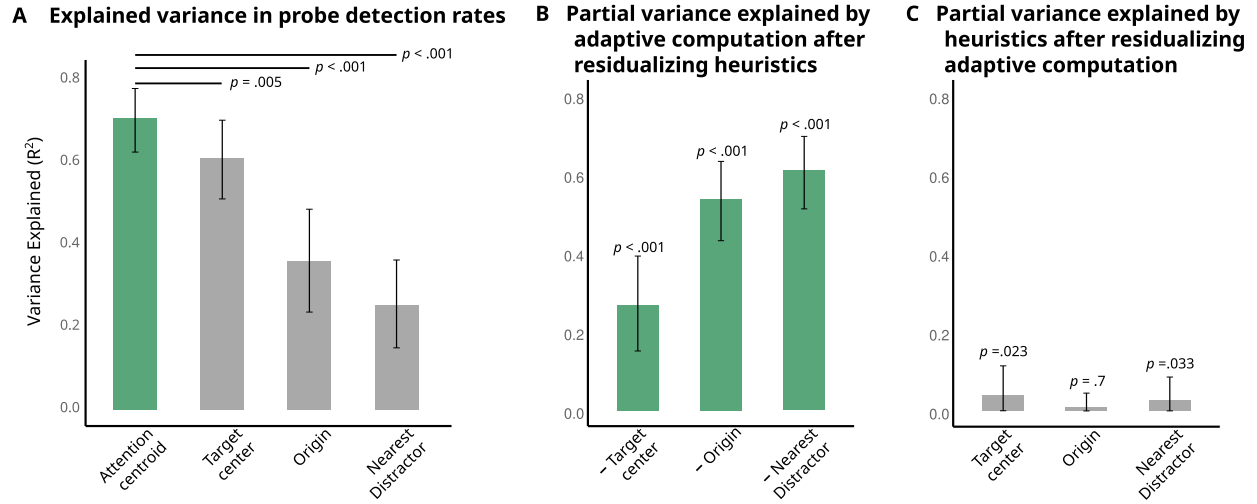


Figure S4: **Calculating probe covariates using adaptive computation architecture** (A) Compared to several MOT-specific heuristic models (gray), the attention centroid (green) explains detection rates to a greater degree, reaching the noise ceiling of the dataset (error bars are 95% confidence intervals). (B) The attention centroid subsumes the predictive power of these heuristics, significantly explaining detection rates after regressing away variance shared with each heuristic (significance calculated using one-tailed bootstrapped hypothesis testing). (C) Critically, the predictive power of each heuristic greatly diminishes after residualizing shared variance with the attention centroid, establishing adaptive computation as a new mechanism of human attention that effectively captures the dynamics of object-level selection present within trials.

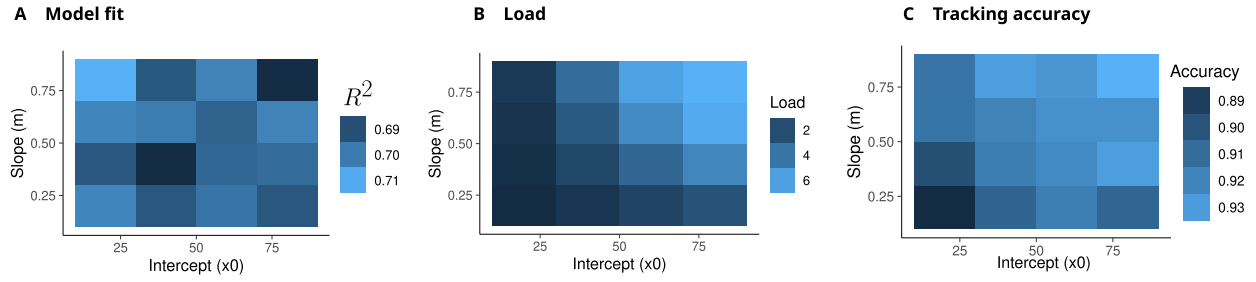


Figure S5: Grid search over load curve parameters (intercept x_0 , and slope m). (**A**) Distribution of explained variance (R^2) of the attention centroid model over human probe detections, ranging from $R^2 = [.68, .71]$. (**B**) Distribution of average load (rejuvenation moves per object, particle, time point). This serves as a sanity check, as both parameters should increase load. (**C**) Distribution of average tracking accuracy, which is proportionally impacted by load.

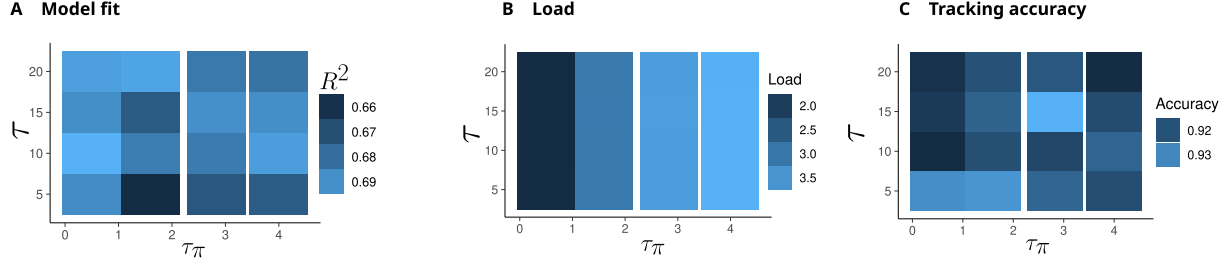


Figure S6: Grid search over importance parameters (softmax temperature τ , and expected reward temperature $\tau\pi$). (**A**) Distribution of explained variance (R^2) of the attention centroid model over human probe detections, ranging from $R^2 = [.66, .70]$. (**B**) Distribution of average load (rejuvenation moves per object, particle, time point). This serves as a sanity check, as load should not vary as a function of τ (y-axis). However, load should vary as a function of $\tau\pi$ as increasing that value should inflate $\delta_k\pi$ and thus ∇ . (**C**) Distribution of average tracking accuracy.