

Suppression of a salient distractor protects the processing of target features

William Narhi-Martinez, Blaire Dube, Jiageng Chen, Andrew B. Leber, & Julie D. Golomb

Department of Psychology, The Ohio State University

Author Note

Correspondence concerning this article should be addressed to William Narhi-Martinez, Department of Psychology, The Ohio State University, 1835 Neil Ave, Columbus, OH 43210, United States. Email: narhi-martinez.1@osu.edu

FEATURE PROTECTION

Abstract

We are often bombarded with salient stimuli that capture our attention and distract us from our current goals. Decades of research has shown the robust detrimental impacts of salient distractors on search performance and, of late, in leading to altered feature perception. These feature errors can be quite extreme, and thus, undesirable. In search tasks, salient distractors can be suppressed if they appear more frequently in one location, and this learned spatial suppression can lead to reductions in the cost of distraction as measured by reaction time slowing. Can learned spatial suppression also protect against visual feature errors? To investigate this question, participants were cued to report one of four briefly presented colored squares on a color wheel. On two-thirds of trials, a salient distractor appeared around one of the nontarget squares, appearing more frequently in one location over the course of the experiment. Participants' responses were fit to a model estimating performance parameters and compared across conditions. Our results showed that general performance (guessing and precision) improved when the salient distractor appeared in a likely location relative to elsewhere. Critically, feature swap errors (probability of misreporting the color at the salient distractor's location) were also significantly reduced when the distractor appeared in a likely location, suggesting that learned spatial suppression of a salient distractor helps protect the processing of target features. This study provides evidence that, in addition to helping us avoid salient distractors, suppression likely plays a larger role in helping to prevent distracting information from being encoded.

FEATURE PROTECTION

Introduction

Effective attentional control requires not only accurate guidance towards relevant stimuli and locations, but also the ability to suppress irrelevant, but salient, stimuli. Driving to work requires attending to approaching stoplights while preventing ourselves from being distracted by things like text messages. The goal of external selective attention is to determine what in our environment is worthy of further processing (i.e., working memory encoding). However, an irrelevant stimulus may appear that is so salient that we are unable to ignore it, causing us to select the irrelevant stimulus for attention and hindering our ability to efficiently achieve our current goal. Such incidental shifts in attention are known as attentional capture.

Much of the work on attentional capture has focused on the spatiotemporal impacts of attentional capture using measures of reaction time and forced-choice accuracy measures. These studies have converged on the finding that attentional capture can negatively impact behavior: participants take longer to respond and are less accurate when a distractor is present (Pashler, 1988; Theeuwes, 1994; Yantis & Jonides, 1984; for a review see Luck et al., 2021). More recent work, however, suggests that the consequences of attentional capture are broader than once thought. Chen, Leber, and Golomb (2019) measured the consequences of attentional capture on feature perception and recall using a delayed-estimation task with a continuous response modality. They conducted two experiments in which four colored squares were briefly shown, with the target being simultaneously outlined by a bolded white frame. On two-thirds of trials, a salient distractor appeared in the display: four white dots surrounded one of the colored squares, half of the time around one of the nontargets adjacent to the target square. Participants reported the color of the target square on a subsequent color wheel surrounding a post-cue of the target location. Probabilistic mixture modeling (Bays et al., 2009; Zhang & Luck, 2008) was used to analyze the distribution of each participant's responses, which allowed for measurements revealing the specific types of errors a salient distractor could elicit from this continuous

FEATURE PROTECTION

feature space. Not only were participants more likely to guess and respond with less precision when a salient distractor was present, but significant amounts of swapping (selecting the color in the salient distractor location) and repulsion (near-target responses biased slightly away from the color in the salient distractor location) were also observed. These results showed, in addition to basic performance decrements in the presence of a salient distractor, that attentional capture by a distractor can lead to perceptual feature interference (even though the color appearing in the salient distractor location was no more salient than the other colors present).

In the real world, the consequences of feature interference could be detrimental, especially since there is an ever-increasing number of items being designed to capture our attention (e.g., online advertisements, phone notifications, storefront displays, etc.). How are we able to avoid these distractors to effectively navigate our environments? Prior work has shown suppression plays a key role in meeting this challenge (Gaspelin et al., 2015; Gaspelin & Luck, 2018; Sawaki et al., 2012; Sawaki & Luck, 2010). For example, when a particular location is more likely to contain a distractor, observers can learn this statistical regularity over time. They then begin to suppress that high probability distractor location, mitigating the consequences of attentional capture and improving performance. This general finding has been documented whether the high probability distractor location is fixed (Britton & Anderson, 2020; Huang et al., 2021; Kong et al., 2020; Wang & Theeuwes, 2018b, 2018a) or flexible (i.e., when the high probability distractor location is defined via its position relative to another display item; Leber et al., 2016). While information about statistical regularities can be explicit (i.e., directly cued) or implicit (i.e., learned over time), there is some evidence that implicit learning over time can be more effective in mitigating attentional capture effects (Moher & Egeth, 2012; Noonan et al., 2016; Wang & Theeuwes, 2018a).

A number of studies have explored the mechanisms behind this suppression (Failing et al., 2019; Gaspelin et al., 2015; Geng & Duarte, 2021; Gong & Theeuwes, 2021; Huang et al., 2021; Won et al.,

FEATURE PROTECTION

2022), with debates regarding whether suppression effects can be explained better by distractor inhibition or target enhancement (Failing et al., 2019), or if it operates in a proactive (i.e., preemptively suppressing a distractor before selection) or reactive (e.g., “search and destroy”) manner (Chang et al., 2023; Geng & Duarte, 2021; Huang et al., 2021; Kong et al., 2020). Recently, Gong and Theeuwes (2021) characterized a saliency-specific mechanism, while Won et al. (2022) suggested that attentional suppression serves to prevent salient, task-irrelevant information from entering working memory. However, much of the previous literature on experience-driven suppression has relied solely on simple search tasks with limited (often, two) response options, and these studies have largely focused on how suppression may protect against the prolonged behavioral response times characteristic of attentional capture (but see Won et al., 2022, for an investigation of how suppression affects memory precision for a salient distractor over time). As we now know, however, the consequences of distraction extend beyond disruptions to response time: dynamic distraction *also* causes systematic and measurable perceptual errors (Chen et al., 2019). Presently, it is unknown whether experience-driven suppression *also* protects target representations against distractor-induced perceptual errors.

The present study aims to investigate whether – and to what extent – experience-driven spatial suppression protects the processing of the target features. We employ a continuous color report paradigm – rather than reaction time or accuracy measurements – to examine whether a salient distractor appearing in a learned likely location will result in reduced feature interference compared to a distractor appearing in a less likely location. We predict that spatial suppression of a high probability distractor location will not only result in improved overall performance when the salient distractor appears in the likely location (relative to when it appears elsewhere), but predictable distractors may interfere less with feature perception and working memory processes. This protection from feature interference could be evidenced by a reduction in the feature swapping and/or repulsion errors typically elicited by attention capture (Chen et al., 2019).

FEATURE PROTECTION

Methods

Sample

Participants were recruited from The Ohio State University campus and received either course credit or payment for their time. To determine our sample size, we pre-registered an optional stopping rule based on a sequential Bayes factor design. This method of determining sample size has been demonstrated as an effective method that protects interpretability of the results and does not induce statistical bias or require penalization for checking (Rouder, 2014; Schönbrodt & Wagenmakers, 2018). Our stopping rule was pre-registered on OSF as follows (condensed here, see <https://osf.io/ys3kc> for full description): According to a sequential analysis of the Bayes factor for the swap effect in Experiment 1 of Chen et al., (2019), ‘strong evidence’ ($BF_{10} > 10$; Lee & Wagenmakers, 2014) was observed by their 20th participant. Therefore, we set our minimum sample size to 20. After that point, we collected data in 8-participant intervals until sufficient evidence for or against our main comparison of interest was reached (or a maximum of $N=60$). Our main comparison of interest concerned whether the swap effect for the salient distractor in the probable location was significantly different compared to the salient distractor in a control location (see **Analyses**). We set our thresholds to $BF_{10} > 6$ as sufficiently in support of the alternative and $BF_{10} < 1/6$ as sufficiently in support of the null model.

Our stopping-rule threshold was reached at 52 participants (29 female, 22 male, 1 non-binary, aged 18-36). Data from 8 additional participants who completed the experiment were excluded prior to analysis for not maintaining fixation on at least 75 trials in each condition of interest.

Setup

Each participant was seated and placed their head against a chin and forehead rest 60cm away from the monitor. The 62cm LCD monitor’s resolution was adjusted to display a 4x3 presentation window (resolution: 1280x960, refresh rate: 200Hz) and was color calibrated with a Minolta CS-100

FEATURE PROTECTION

colorimeter. Stimuli were generated using MATLAB (Mathworks) and the Psychophysics Toolbox (Brainard, 1997; Kleiner et al., 2007; Pelli, 1997) on a Windows computer. Eye position was recorded using an Eyelink 1000 eye-tracker (SR Research).

Procedure

As displayed in *Figure 1*, every trial began with a black fixation cross appearing at the center of a grey background. Once this cross had been fixated (eye position accurately maintained within a 2° radius) for a consecutive 700ms, it changed into a black dot, and placeholders (four thin, white frames) appeared outlining the locations of the upcoming stimuli. Fixation had to be maintained for an additional 300ms. If fixation was broken during this time (>2° deviation), the cross would re-appear, and this loop would continue until fixation was properly maintained for the entire 1000ms. This two-stage fixation period allowed us to maximize the number of usable trials, as we would exclude any trials from analyses in which fixation was broken following this period.

Once consistent fixation was achieved, the fixation dot remained on-screen while the stimulus array was presented for 50ms. The stimulus array was four colored squares (each sized 2° x 2°, centered at an eccentricity of 4°), which appeared in upper left, upper right, lower left, and lower right corner positions. The color of the squares varied on every trial. The color of the upper left square was chosen

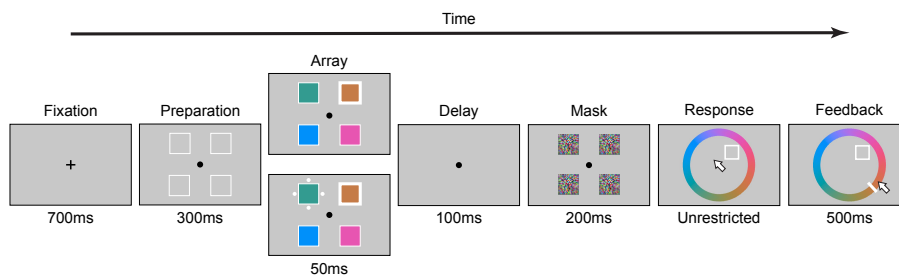


Figure 1. Experimental procedure (not drawn to scale). On every trial, participants were shown four colored squares and instructed to report the color within the target square (outlined by the bold white frame, also post-cued during presentation of the color wheel). On two-thirds of trials, four white dots would appear around one of the nontarget locations (salient distractor). This salient distractor would appear in one particular location on 62.5% of distractor-present trials, with this high probability distractor location being counterbalanced across participants.

FEATURE PROTECTION

randomly from 180 possible color values that were evenly distributed along a color wheel in CIE L*a*b color space ($L = 60$, $a = 22$, $b = -1$, radius = 50). The colors of the squares in the upper right and lower left were then selected to be exactly 90° and -90° away in color space, direction randomly assigned on each trial, from the color in the upper left. The lower right square was always 180° away in color space from the color in the upper left. The target square was indicated by surrounding its location with a bold, thick frame for the duration of the stimulus array. The stimulus array was followed by a blank delay screen for 100ms, followed by four scrambled-color square masks (each mask was a 22×22 grid of randomly generated colors created prior to the start of the experiment, with each mask always appearing in the same location for the duration of the experiment) for 200ms. Afterwards, the response screen appeared, consisting of a color wheel centered on the screen (diameter = 6.5° , width = 1°) displaying all 180 possible color values, along with a white frame post-cue in the target location to remind participants which location they should try and recall the color from. After making their selection by clicking on a color, a white feedback line appeared over the correct color for 500ms before proceeding to the next trial.

On two thirds of trials, a salient distractor (four white dots) would appear around one of the nontarget locations. To create a *high probability distractor location*, one of the four stimulus locations was pre-determined (counterbalanced across participants) to contain the salient distractor on 62.5% of distractor-present trials, with the appearance of the salient distractor evenly split among the other three locations on the remaining (12.5% each) distractor-present trials. Note that target appearance was also unevenly distributed amongst the four locations, with the target appearing in the high probability distractor location less frequently (16.7% of trials) compared to each of the other three locations (27.8% of trials). We opted for this design with differential target regularities to more easily balance other factors, since previous research has shown that spatial suppression is mainly driven by statistical learning of distractor location probabilities, not target activation (Failing et al., 2019), and similar

Deleted: 37.5%

Deleted: of

Deleted: Additionally, the

Formatted: Font color: Text 1

Deleted: . T

Deleted: ed

Deleted: on

Deleted: with

Deleted: containing the target on

Deleted:

Deleted:

Deleted: Despite these

Deleted: for ease of

Deleted: ing

Formatted: Font color: Text 1

Formatted: Font color: Text 1

Deleted: being present alongside our main manipulation of interest regarding the salient distractor probabilities, we believe it more likely that the differences in distractor location will drive any suppression we measure, based on

Deleted: (e.g., Chang et al., 2023; Failing et al., 2019). In particular, (Failing et al., 2019) showed that learned spatial suppression cannot be explained by target probability. In their first experiment, no suppression was observed when the target was less likely to appear at the high-probability location and the distractor appeared equally often at each location. However, in their second experiment, when the distractor was more likely to appear at the high probability location and the target probability was kept equal at each location, spatial suppression was observed. These results suggest ...

Deleted: attributable to

Formatted: Font: 11 pt, Font color: Text 1

Deleted: .

FEATURE PROTECTION

suppression effects have been found for contexts that contain either balanced or unbalanced target probabilities (Chang et al., 2023).

Before beginning the experiment, participants were instructed to always report the color of the target stimulus on that trial (the color appearing in the location of the bold, white frame). No mention

regarding the four white dots (salient distractor) was made. however, we strongly emphasized to participants that their target would be outlined in a bold white frame, the location of which was the only relevant stimulus to attend to in the array. (We also included a post-cue redisplaying the target location during the response stage, so it was unlikely participants would simply mistake the distractor cue for the target cue; Chen et al. (2019) verified with confidence reports that distractor-induced swap errors were made with high confidence even in the presence of the post-cue, suggesting that swap errors were due to disrupted color-location binding and not participants actively encoding the distractor color.)

Participants performed 10 practice trials (7 of which contained a salient distractor; all 10 excluded from analyses) before starting the main experiment. Participants then completed 1080 total trials (five blocks of 216 trials) within 1.5 to 2 hours.

Following the completion of the experiment, a series of exit questions appeared on-screen for participants to answer. The first question (EQ1) asked, "Overall, on what percentage of trials did the four white dots seem to appear?", and participants were asked to respond by selecting one of the number keys 1 through 9, corresponding to 10% to 90% of trials, in 10%-increments. Next, participants were asked (EQ2), "Did one location seem to be indicated more frequently by the four white dots?", and they were told to press the Y-key for "yes" and the N-key for "no". The next question (EQ3) asked, "Regardless of how you answered the previous question, take a guess at the location that the four white dots appeared in the most by pressing the corresponding number key.", and a black frame appeared in each of the previous four stimulus locations numbered 1 through 4. The final question (EQ4) asked, "Did the target location ever appear to change within a trial?", for which participants were again instructed

Deleted: in

Deleted: where

Deleted: are

Deleted: balanced versus unbalanced

Deleted: .

Formatted: Font color: Text 1

Formatted: Font: 11 pt, Font color: Text 1

Formatted: Font: 11 pt, Font color: Text 1

Deleted:

Formatted: Font color: Text 1

Commented [GJD1]: We need to mention the blocks here now so the reference to them later makes sense. Just confirming that subjects actually did this in 5 blocks, right? Not that we arbitrarily divided into 5 for the timecourse analysis?

Also, if this is correct, then that sounds like a single block lasted like 20min. Did they break mid-block as well? If so, please add that info.

Formatted: Font color: Text 1

FEATURE PROTECTION

to press the Y-key for “yes” and the N-key for “no”. This final question was included as an exploratory measure to see if participants may have been confused about which location was initially the target.

Deleted: attention was captured enough that

Deleted: (e.g., if they initially confused the four white dots for the target cue, then the target post-cue would have appeared to be indicating a different location)

Analyses

We restricted our analyses to trials where the distractor and target were in adjacent positions (to accommodate the analysis below), or where the distractor was absent but the target was in or adjacent to the high probability distractor location. We excluded any trials in which fixation was broken ($>2^\circ$ deviation from the fixation dot) during the stimulus array.

For our main analyses, we focused on three types of trials, depicted in *Figure 2A*. In the “High Probability Distractor” condition (185-300 trials per subject, depending on fixation exclusions), the salient distractor appeared in its likely location on that trial. For ease of reference, we refer to this likely distractor location as the High Probability Distractor (HPD) location. In the “Low Probability Distractor” condition (83-120 trials per subject, depending on exclusions), the salient distractor appeared in one of the other locations (and the Target was also not in the HPD location). In the “Distractor Absent: Control” condition (190-270 trials per subject, depending on exclusions), there was no salient distractor and the target appeared in a control (not HPD) location. In a set of secondary analyses, we also analyzed an additional condition, “Distractor Absent: target in HPD” condition (75-90 trials per subject, depending on exclusions), when the salient distractor was absent and the target appeared in the high probability distractor location. Our design also resulted in additional trial types (e.g., distractor present: target in HPD) that we did not analyze due to their infrequency (<75 trials per subject) not allowing for sufficient trials to model.

Formatted: Font color: Text 1

Deleted: a control

Deleted: a control

Formatted: Font color: Text 1

Deleted: (high probability distractor location)

For every trial, the angular distance on the color wheel between the reported color and the target color was calculated as the response error. This error was then aligned such that the target color was centered at 0° and the reported color could be a maximum of $\pm 180^\circ$ away. On distractor-present

FEATURE PROTECTION

conditions of interest, in the actual display, the salient distractor color could have been located either +90° or -90° from the target on the color wheel; for the analyses, we re-aligned the direction of response errors on the -90 trials so that the salient distractor would always be coded as +90° and the control nontarget (the square located diagonal to the distractor location) as -90° in our analyses. This allowed us to label response errors with a positive sign as being ‘towards’ the distractor location’s color and response errors with a negative sign as ‘away’ from the distractor location’s color within the distractor-present conditions of interest. On half of the distractor-absent trials (randomly selected), the sign of the response error was flipped to match the distractor-present trials’ realignment process and eliminate any selection confounds driven by color direction on the color wheel.

For each condition, each participant’s distribution of response errors was then fit with a probabilistic mixture model (*Formula 1* for the distractor-present conditions and *Formula 2* for the distractor-absent conditions) estimating five parameters: γ estimated for the proportion of random guesses (a uniform distribution); β_{sal} estimated the probability of misreporting the nontarget in the salient distractor location on distractor-present trials (or β_{ntA} for one of the control nontargets on distractor-absent trials; i.e., a von Mises distribution centered at +90°); β_{nt} estimated the probability of misreporting the control nontarget on distractor-present trials (or β_{ntB} for the other control nontarget on distractor-absent trials; i.e., a von Mises centered at -90°); and the probability of reporting the target (a von Mises distribution with a flexible mean μ , and concentration κ) was estimated by $1 - \beta_{sal} - \beta_{nt} - \gamma$ for distractor-present conditions and $1 - \beta_{ntA} - \beta_{ntB} - \gamma$ for distractor-absent conditions.

$$p(\theta) = (1 - \beta_{sal} - \beta_{nt} - \gamma)\phi_{\mu,\kappa} + \beta_{sal}\phi_{90^\circ,\kappa} + \beta_{nt}\phi_{-90^\circ,\kappa} + \gamma\left(\frac{1}{2\pi}\right) \quad (Formula 1)$$

$$p(\theta) = (1 - \beta_{ntA} - \beta_{ntB} - \gamma)\phi_{\mu,\kappa} + \beta_{ntA}\phi_{90^\circ,\kappa} + \beta_{ntB}\phi_{-90^\circ,\kappa} + \gamma\left(\frac{1}{2\pi}\right) \quad (Formula 2)$$

Commented [GJD2]: FYI it’s easier to read if we make the whole new term red rather than letter by letter. The highlighting doesn’t have to be a pure track changes, just something to highlight new things for the reviewers

Formatted: Font color: Text 1

Formatted: Font color: Text 1

Formatted: Font color: Text 1

Formatted: Font color: Text 1

Formatted: Font color: Text 1

Formatted: Font color: Text 1

Formatted: Font color: Text 1

Formatted: Font color: Text 1

FEATURE PROTECTION

The model was fit to individual participant data for each condition of interest by applying the Markov chain Monte Carlo method using MemToolbox (Suchow et al., 2013). Kolmogorov–Smirnov tests were then run on all main model fittings to ensure good fits to the raw data (all p values > .1). The best-fitting parameter estimates obtained for each subject and condition were compared in JASP software (Version 0.11.1) and MATLAB (Mathworks) using one- and two-way repeated measures ANOVAs, along with paired- and one-sample two-tailed t-tests. Our main comparisons of note involved (1) comparisons of basic performance indicators, including the parameter estimates for random guessing (γ) and standard deviation ($SD = \sqrt{1/\kappa}$), and (2) comparisons of systematic feature errors, specifically feature swap errors indicated by comparing the probability of nontarget reports of the salient distractor vs control colors (β_{sal} vs β_{nt}) and distortion errors indicated by mean shifts (μ) deviating from 0.

While we have chosen to use a probabilistic mixture model (Bays et al., 2009; Zhang & Luck, 2008) for our main data analyses, we recognize the criticisms of this type of model, particularly in comparison to the target confusability competition (TCC) model (Schurgin et al., 2020; Williams et al., 2022). However, we note that we are not drawing conclusions based on an assumption that the parameters for guess rate and response precision reflect independent theoretical entities. Our main focus is on the swap rate parameter, and it has been shown that in cases where overall memory strength is high (e.g., in the current study the probability of reporting the target is greater than .9, on average), there is general agreement between swapping estimates obtained from a standard mixture model and the TCC-Swap model, according to Williams et al. (Williams et al., 2022).

Formatted: Indent: First line: 0.5"

Results

On average, 7% of trials were discarded due to fixation broken across the 52 participants.

FEATURE PROTECTION

Attentional Capture: Basic Performance Indicators

Figure 2B shows the performance measures for our three main conditions of interest. As pre-registered, we first tested the basic premise that the salient distractors captured attention and impaired performance in the control (Low Probability Distractor) condition. Indeed, we measured a significantly higher probability of random guessing (γ parameter) on “Low Probability Distractor” trials compared to “Distractor Absent: Control” trials, $t(51) = 5.980$, $p < .001$, $d = .829$, $BF_{10} = 6.72 \times 10^4$, as well as significantly worse precision (higher SD parameter), $t(51) = 4.582$, $p < .001$, $d = .635$, $BF_{10} = 6.85 \times 10^2$. These results indicate that the presence of a salient distractor hampered overall performance, mirroring the analogous comparisons reported by Chen et al. (2019).

We next examined whether the learned spatial suppression manipulation was effective, by comparing these basic performance indicators of attention capture on trials where the distractor was in the likely location versus a control location. The guess rate was indeed significantly lower on “High Probability Distractor” trials compared to “Low Probability Distractor” trials, $t(51) = -3.696$, $p < .001$, $d = -.513$, $BF_{10} = 49.28$. SD was also significantly lower on “High Probability Distractor” trials compared to “Low Probability Distractor” trials, $t(51) = -2.589$, $p = .012$, $d = -.359$, $BF_{10} = 3.06$. Together, these results suggest that spatial suppression was occurring when the salient distractor appeared in the likely location, leading to overall improved performance on those trials relative to when the salient distractor appeared in one of the less likely locations.

Formatted: Font color: Text 1

Deleted: a

Deleted: control

Formatted: Font color: Text 1

FEATURE PROTECTION

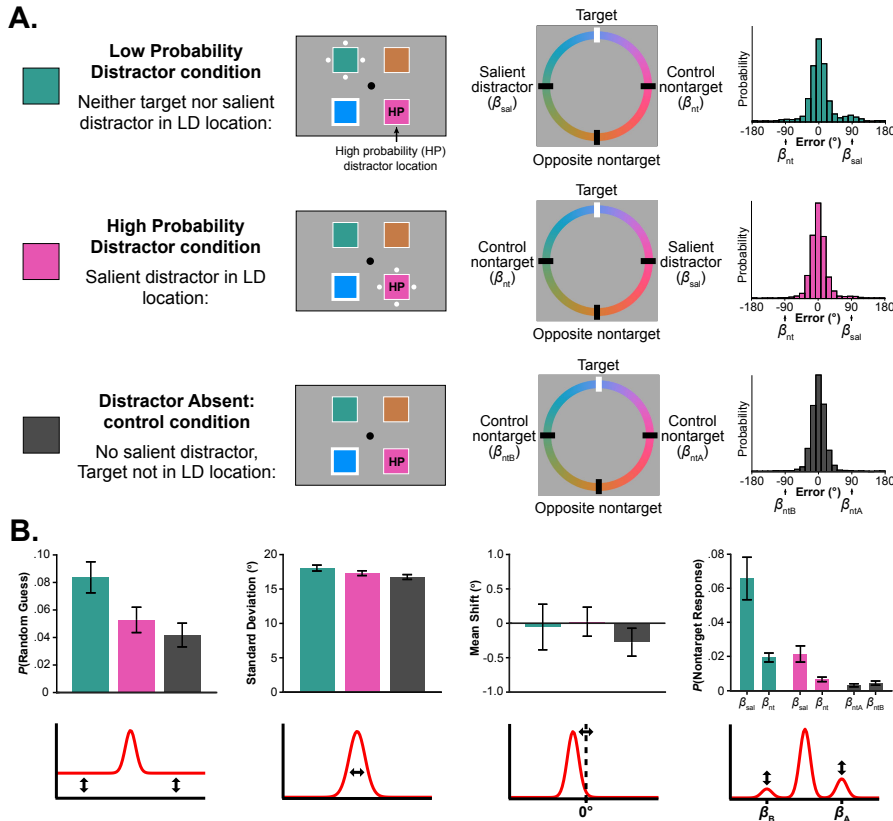


Figure 2. Probabilistic mixture results for main conditions of interest. A.) Schematics of “Low Probability Distractor”, “High Probability Distractor”, and “Distractor Absent: Control” conditions, illustrating example stimulus arrays based on the position of the salient distractor relative to the high probability (HP) distractor location. Nontargets are labeled in physical space (left) and color wheel space (right), for these illustrative examples. Response histograms collapsed across participants are shown for each condition at the right. All histograms are plotted as response errors relative to the correct target color (0° error), aligned with the salient distractor at +90° error when present. B.) Mean maximum likelihood parameter estimates for: probability of random guesses (γ), SD ($\sqrt{1/\kappa}$), mean shift (μ), and probability of nontarget responses (β). Cartoons illustrating each parameter in the model are shown in red below each plot. In the distractor present condition, the nontarget in the distractor location is represented by β_{sal} , while β_{nt} represents the control nontarget; a negative mean shift indicates a biasing of target responses away from the color of the nontarget in the distractor location. Error bars indicate standard error from the mean, $N=52$.

Commented [GJD3]: The figure still says LD location (on the left side descriptions). Also, if we’re using “HPD” for the nae of the 4th condition, we should maybe label the pink square HPD, not HP.

Commented [GJD4]: HP or HPD? We should be consistent

FEATURE PROTECTION

Feature Interference Errors

Our primary question of interest is whether suppression of the high probability distractor location protects targets from feature interference errors. We compared our distractor present conditions (“High Probability Distractor” vs “Low Probability Distractor”) and their rates of nontarget misreports (β_{sal} vs β_{nt}) using a repeated-measures ANOVA. We observed a significant main effect of condition, $F(1, 51) = 31.739$, $p < .001$, $\eta^2 = .106$, condition model $BF_{10} = 7.97 \times 10^2$, such that misreport errors were greater for low than high probability distractors; and a significant main effect of nontarget, $F(1, 51) = 16.153$, $p < .001$, $\eta^2 = .122$, nontarget model $BF_{10} = 3.28 \times 10^3$, such that misreports more frequently reflected the salient distractor than the control nontarget. Importantly, there was a significant interaction, $F(1, 51) = 8.795$, $p = .005$, $\eta^2 = .032$, with strongest evidence for the condition + nontarget + condition $\times$ nontarget model $BF_{10} = 3.29 \times 10^7$ compared to the null model (next strongest evidence: condition + nontarget model $BF_{10} = 8.20 \times 10^6$), indicating that the difference between β_{sal} and β_{nt} was greater in the “Low Probability Distractor” condition compared to the “High Probability Distractor” condition. Follow-up simple effect t-tests found the β_{sal} vs. β_{nt} comparison was significant for both conditions (“Low Probability Distractor”: $t(51) = 3.756$, $p < .001$, $d = .521$, $BF_{10} = 58.42$, “High Probability Distractor”: $t(51) = 3.321$, $p = .002$, $d = .461$, $BF_{10} = 17.86$). Together, these results suggest that distractor-induced swap errors were present in both conditions, but there was indeed a significant difference in the rate of these errors depending on whether the salient distractor appeared in its likely location or elsewhere. In other words, suppression of the likely salient distractor location reduced the likelihood of swap error feature interference.

As a purely exploratory analysis for potential learning effects, we aggregated trials across all participants to increase power so that we could analyze each of the five blocks separately. Swaps to the salient distractor decreased over time, as expected based on distractor habituation literature (e.g., Turatto et al. 2018). However, the rate of decrease was similar for both distractor conditions, and the

Deleted: An

Deleted: block-wise analyses in which

Deleted: collapsed

Formatted: Indent: First line: 0.5"

Deleted: showed that s

Formatted: Font: 11 pt, Font color: Text 1

Deleted: in

Deleted: . However

FEATURE PROTECTION

difference between “High Probability Distractor” and “Low Probability Distractor” swaps was evident even within the first block of trials (216 trials), suggesting these statistical regularities were learned relatively quickly.

Finally, we assessed target distortion errors, by comparing the target distribution’s mean parameter estimate to the true, aligned value of 0. One-sample t-tests did not show a significant mean shift in either the “High Probability Distractor”, $t(51) = .109$, $p = .914$, $d = .015$, $BF_{10} = .15$, nor “Low Probability Distractor”, $t(51) = -.164$, $p = .871$, $d = -.023$, $BF_{10} = .15$, conditions. The latter result was surprising to us, as we had anticipated observing a repulsion effect for the mean of the target distribution when the distractor was present in a low-probability location, based on previous work (Chen et al., 2019). Possible explanations for the lack of a repulsion effect are explored in the **Discussion** section below.

Secondary Analysis: Distractor Absent Conditions Comparison

The primary results above showed that participants spatially suppressed the high probability distractor location, and that this suppression aided in preventing feature interference by reducing swap errors induced by the salient distractor. As a secondary question, we can ask what additional effects this spatial suppression may have on feature perception on trials when the *target* appeared in the high probability distractor location, as prior studies examining (non-learning related) suppression have reported worse memory performance at the location of a suppressed salient distractor (e.g., Gaspelin et al., 2015; Gaspelin & Luck, 2018). To assess this, we compared distractor absent conditions which differed only on whether the target was in the high probability distractor location (“Distractor Absent: target in HPD”) or a control location (“Distractor Absent: Control”; *Figure 3*). A paired-samples t-test revealed no significant difference in guess rate between the “Distractor Absent: target in HPD” condition and the “Distractor Absent: Control” condition, $t(51) = -1.531$, $p = .132$, $d = -.212$, $BF_{10} = .45$. There was,

Deleted: within the first block there was already a visible

Deleted: (

Deleted: control

Deleted:)

FEATURE PROTECTION

however, a significantly higher SD in the “Distractor Absent: target in HPD” condition compared to the “Distractor Absent: Control” condition, $t(51) = 2.399$, $p = .020$, $d = .333$, $BF_{10} = 2.04$, suggesting worse overall response precision when the target appeared in the high probability distractor location (but see Schurgin et al., 2020, for arguments against interpreting guess rate and SD as separate parameters).

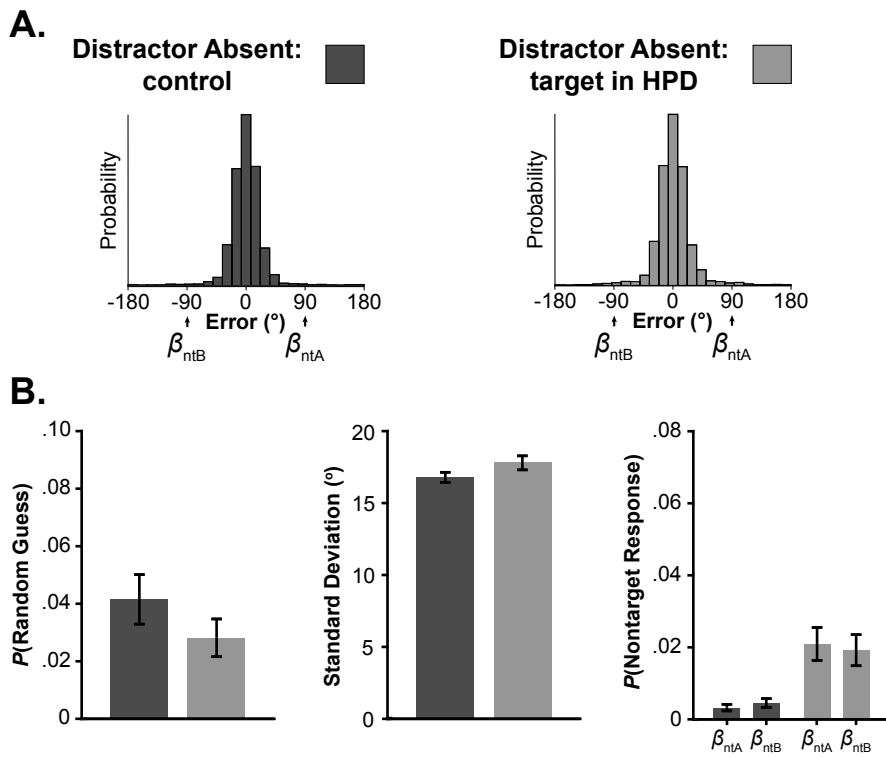


Figure 3. Probabilistic mixture results for distractor absent conditions. A.) “Distractor Absent: Control” (target not in the high probability distractor [HPD] location) and “Distractor Absent: target in HPD” condition response histograms collapsed across participants. B.) Mean maximum likelihood parameter estimates for: probability of random guesses (γ), SD ($\sqrt{1/\kappa}$), and probability of nontarget responses (β). Error bars indicate standard error from the mean, $N=52$.

FEATURE PROTECTION

Next, we examined the probability of misreporting a nontarget color between these two distractor absent conditions. We compared our distractor absent conditions (“Distractor Absent: target in HPD” vs “Distractor Absent: Control”) and their rates of nontarget misreports (β_{ntA} and β_{ntB}) using a repeated-measures ANOVA. We observed a significant main effect of condition, $F(1, 51) = 14.499$, $p < .001$, $\eta^2 = .134$, such that misreport errors were greater in the “Distractor Absent: target in HPD” condition, and the strongest evidence for the condition model $BF_{10} = 6.98 \times 10^3$ compared to the null model (next strongest evidence: condition + nontarget model $BF_{10} = 1.08 \times 10^3$). The main effect of nontarget was not significant, $F(1, 51) = .052$, $p = .820$, $\eta^2 = 1.87 \times 10^{-4}$, nontarget model $BF_{10} = .15$; nor was there a significant interaction, $F(1, 51) = .281$, $p = .598$, $\eta^2 = .001$, condition + nontarget + condition $\times$ nontarget model $BF_{10} = 2.39 \times 10^2$. Significant increases in overall nontarget reports but not random guesses in the “Distractor Absent: target in HPD” condition (*Figure 3*) suggests that suppression of the high probability distractor location may have led participants to allocate relatively more of their attention to the control locations, so they were more likely to select a control nontarget color instead of the target when the target appeared in the high probability distractor location compared to when it appeared elsewhere.

Exit Questions

On average, participants reported that the four white dots appeared on 45% of trials (EQ1), significantly less than the true frequency of 67%, $t(51) = -6.731$, $p < .001$, $d = -.933$, $BF_{10} = 8.82 \times 10^5$. Accordingly, most participants (37/52) responded “no” when asked if they noticed the salient distractor being biased to one location (EQ2), and only 10 of 52 participants correctly identified their high probability distractor location (EQ3), which did not significantly differ from chance (25%) according to a binomial test ($p = .423$). Finally, 30/52 participants answered “yes” when asked if the target probe ever changed location between the memory array and response screen (EQ4). Importantly, none of these exit

Formatted: Font color: Text 1

Formatted: Font color: Text 1

Deleted: most participants (

Deleted:)

FEATURE PROTECTION

question results showed significant interactions with our key findings (all p-values > .22), suggesting that the resulting effects of suppression we measured did not depend on conscious awareness of the distractor or its manipulated regularities.

Deleted: our

Discussion

Salient distractors can not only slow our reaction times (Folk et al., 1992; Pashler, 1988; Theeuwes, 1994; Yantis & Jonides, 1984), but they also interfere with our perception of a nearby target's features (Chen et al., 2019). We investigated how distractor suppression could potentially attenuate the detrimental effects of attentional capture on feature processing. We found that experience-based suppression of a high probability distractor location led to decreased interference from a salient distractor appearing in that location. This was evidenced by lower guessing rates, SD, and swap rates when a salient distractor appeared in a high probability distractor location, compared to when it appeared in a less high probability distractor location. These results expand upon the extent of experience-driven suppression's role in cognitive processes beyond benefits to reaction time and simple accuracy measures (Britton & Anderson, 2020; Huang et al., 2021; Kong et al., 2020; Leber et al., 2016; Wang & Theeuwes, 2018b, 2018a).↓

Here, we showed that the suppression of a salient distractor aided in protecting target feature processing. In addition to overall performance improving when the distractor was suppressed, there was a significantly lower likelihood of mistakenly reporting the color (swap errors) in the suppressed location. This suggests that the reduction in capture by the salient distractor reduced the ability for co-located features to inadvertently enter working memory and potentially interfere with the target representation. The idea that the features of suppressed distractors are less likely to be processed and enter working memory is supported by the findings of Won et al. (2022). Won et al. (2022) studied a different type of distractor suppression: rather than learned spatial suppression of a predictable

Deleted: While we have chosen to use a probabilistic mixture model While we have chosen to use a probabilistic mixture model (Bays et al., 2009; Zhang & Luck, 2008) for our main data analyses, we recognize the criticisms of this type of model, particularly in comparison to the target confusability competition (TCC) model for our main analyses, proponents of the target confusability competition (TCC) model have criticized its use (Schurgin et al., 2020; Williams et al., 2022). However, we note that we are not drawing conclusions based on an assumption that the parameters for guess rate and response precision reflect independent theoretical entities. Our main focus is on the swap rate parameter, and it has been shown that in cases where overall memory strength is high (e.g., in the current study the probability of reporting the target is greater than .9, on average), there is general agreement between swapping estimates obtained from a standard mixture model and the TCC-Swap model, according to Williams et al. Notably, however, we would anticipate general agreement of our results with the TCC model in this case based on the following: 1) we are not attempting to separately interpret the estimates of random guessing and response precision, and we focus on swap rates for the critical effects, and 2) because overall memory strength is strong (the probabilities of guessing and misreport rates each average less than .1 across conditions), we would expect comparable results between our mixture model swapping estimates and the TCC-Swap model (Williams et al., 2022).

FEATURE PROTECTION

distractor location, their suppression was based on distractor frequency (the idea that an initial unexpected distractor captures attention more than a repeated distractor in a context where distractors are frequently present). Moreover, whereas our study focused on distractor-induced interference for target feature reports, Won et al. (2022) probed memory for the salient distractor feature itself. Their design combined a typical singleton search task with a one-shot memory probe for the color of the salient distractor after a varying number of trials that differed across participants. Their results showed that, in addition to search times for the target decreasing over time, feature report performance for the color of the salient distractor also decreased over time. In other words, those participants who saw more salient distractors - before being asked to report what one looked like - had worse feature memory compared to participants who saw fewer salient distractors before the memory probe (Won et al., 2022). These findings suggest that the features of a suppressed salient distractor are less likely to be processed, at least in the case of spatially-generic distractor suppression via repeated exposure context. The results of the current study suggest that experience-driven suppression of a spatial location expected to contain a distractor also results in reduced processing of that distractor's features. Moreover, we provide novel evidence for an additional consequence of this effect: that learned spatial suppression also benefits encoding and recall of the target features.

A secondary question which the present study may provide some insight into is the ongoing debate over whether learned suppression effects are proactive (Chang et al., 2023; Geng & Duarte, 2021; Huang et al., 2021; Kong et al., 2020). In addition to testing feature reporting depending on the location of the salient distractor, we included a secondary analysis to examine if we would observe any differences between distractor-absent conditions depending on where the target was located. This comparison was conducted to see whether we would find evidence for worse performance when the target appeared in the high probability distractor location, which might have suggested that location was being proactively suppressed. Alternatively, if our main results were driven by reactive suppression,

FEATURE PROTECTION

we would predict performance would not be worse when the target appeared in the high probability distractor location, as on these distractor-absent trials there was no distractor to trigger suppression. While we did observe higher SD and nontarget reporting when the target appeared in the location where the distractor was expected, there was no significant difference in guess rate and, in fact, the average guess rate was numerically lower when the target appeared in the high probability distractor location compared to when it appeared elsewhere. Although the present results do not provide definitive evidence that reduced feature interference is driven by either a proactive or reactive mechanism of suppression, these distractor absent results may support predictions made by the Priority Accumulation Framework (PAF), which proposes that priorities can be assigned and updated for locations over time before an attentional shift is triggered to the highest priority location (Darnell & Lamy, 2022; Lamy et al., 2018). According to the PAF, the learned suppression we measured here can be attributed to deprioritization of the high probability distractor location. This would suggest the other three locations would have relatively higher priority, thus leading to a higher likelihood of attentional selection and explaining the greater tendency to misreport nontargets when the target appeared in the high probability distractor location.

Finally, an unanticipated finding was the absence of any repulsion effects. Chen et al. (2019) observed that attentional capture caused both significant swapping to the salient distractor and a target response distribution shifted significantly away from the salient distractor's color. Therefore, we expected to observe a repulsion effect, at the very least, when our salient distractor appeared in a less likely (control) location. We observed no such effect, however. One potential explanation could be a difference in experimental design: in the current study, to maximize power for our critical distractor probability manipulation, there were no "valid" trial types where the salient cue surrounded the target location, which had occurred on one third of the trials in Chen et al. (2019). Here, the salient cue was always a distractor and never overlapped with the target, making it a more pure "distracting cue" that

FEATURE PROTECTION

participants should always avoid. This design decision regarding the distracting cue could actually represent another learned regularity that resulted in additional distractor suppression. Previous research has shown that distractor regularities beyond its spatial location can be learned about and affect performance, such as features (Stilwell et al., 2019), frequency (Geyer et al., 2008), and prior value (Anderson et al., 2011). Thus, it is possible that learned distractor suppression could be contributing to a protection from feature repulsion errors as well, but future work would be needed to investigate this more definitively.

To conclude, we measured feature encoding and recall under salient distraction, which could be spatially suppressed on a majority of trials by learning where the salient distractor was likely to appear. We found that, in addition to helping us avoid other known distractor costs, suppression plays a larger role in helping to prevent features associated with the distractor from interfering with the processing of our target item.

Acknowledgements

The authors gratefully acknowledge the assistance of India Carter, Jayanth Donthireddy, Haley McIntyre, John McNally, and Veronica Olaker in the recruitment of participants and data collection. This research was supported in part by NSF grant BCS-1848939 (JG & AL) and by an NSERC PDF (BD).

Declarations

- **Funding:** This study was funded in part by NSF grant BCS-1848939 (JG & AL) and an NSERC PDF to BD.
- **Conflicts of interest/Competing interests:** The authors have no relevant financial or non-financial interests to disclose.

FEATURE PROTECTION

- **Ethics approval:** Approval for this study was granted by The Ohio State University Behavioral and Social Sciences Institutional Review Board.
- **Consent to participate:** All participants read and signed a consent form prior to participation.
- **Consent for publication:** Informed consent was obtained from all participants.
- **Availability of data and materials:** Experiment plans are available on OSF and data will be made available there post-publication.
- **Code availability:** Experimental code is available upon request.

Open Practices Statement

This experiment was pre-registered on the Open Science Framework (OSF; <https://osf.io/ys3kc/>) prior to starting data collection. Our original theoretical motivation, sample size stopping rule, exclusion criteria, methods, and analyses can be found there. Any analyses included here that were not listed in the pre-registration are declared as exploratory. Data will also be made available post-publication on OSF (<https://osf.io/gcs4e/>).

FEATURE PROTECTION

References

- Anderson, B. A., Laurent, P. A., & Yantis, S. (2011). Value-driven attentional capture. *Proceedings of the National Academy of Sciences*, 108(25), 10367–10371.
<https://doi.org/10.1073/pnas.1104047108>
- Bays, P. M., Catalao, R. F. G., & Husain, M. (2009). The precision of visual working memory is set by allocation of a shared resource. *Journal of Vision*, 9(10), 7. <https://doi.org/10.1167/9.10.7>
- Brainard, D. H. (1997). The Psychophysics Toolbox. *Spatial Vision*, 10(4), 433–436.
<https://doi.org/10.1163/156856897X00357>
- Britton, M. K., & Anderson, B. A. (2020). Specificity and persistence of statistical learning in distractor suppression. *Journal of Experimental Psychology: Human Perception and Performance*, 46(3), 324. <https://doi.org/10.1037/xhp0000718>
- Chang, S., Dube, B., Golomb, J. D., & Leber, A. B. (2023). Learned spatial suppression is not always proactive. *Journal of Experimental Psychology: Human Perception and Performance*.
<https://doi.org/10.1037/xhp0001133>
- Chen, J., Leber, A. B., & Golomb, J. D. (2019). Attentional capture alters feature perception. *Journal of Experimental Psychology: Human Perception and Performance*, 45(11), 1443–1454.
<https://doi.org/10.1037/xhp0000681>
- Darnell, M., & Lamy, D. (2022). Spatial cueing effects do not always index attentional capture: Evidence for a priority accumulation framework. *Psychological Research*, 86(5), 1547–1564.
<https://doi.org/10.1007/s00426-021-01597-0>

FEATURE PROTECTION

- Failing, M., Wang, B., & Theeuwes, J. (2019). Spatial suppression due to statistical regularities is driven by distractor suppression not by target activation. *Attention, Perception, & Psychophysics*, 81(5), 1405–1414. <https://doi.org/10.3758/s13414-019-01704-9>
- Folk, C. L., Remington, R. W., & Johnston, J. C. (1992). Involuntary covert orienting is contingent on attentional control settings. *Journal of Experimental Psychology: Human Perception and Performance*, 18(4), 1030–1044. <https://doi.org/10.1037/0096-1523.18.4.1030>
- Gaspelin, N., Leonard, C. J., & Luck, S. J. (2015). Direct Evidence for Active Suppression of Salient-but-Irrelevant Sensory Inputs. *Psychological Science*, 26(11), 1740–1750. <https://doi.org/10.1177/0956797615597913>
- Gaspelin, N., & Luck, S. J. (2018). Combined Electrophysiological and Behavioral Evidence for the Suppression of Salient Distractors. *Journal of Cognitive Neuroscience*, 30(9), 1265–1280. https://doi.org/10.1162/jocn_a_01279
- Geng, J. J., & Duarte, S. E. (2021). Unresolved issues in distractor suppression: Proactive and reactive mechanisms, implicit learning, and naturalistic distraction. *Visual Cognition*, 29(9), 608–613. <https://doi.org/10.1080/13506285.2021.1928806>
- Geyer, T., Müller, H. J., & Krummenacher, J. (2008). Expectancies modulate attentional capture by salient color singletons. *Vision Research*, 48(11), 1315–1326. <https://doi.org/10.1016/j.visres.2008.02.006>
- Gong, D., & Theeuwes, J. (2021). A saliency-specific and dimension-independent mechanism of distractor suppression. *Attention, Perception, & Psychophysics*, 83(1), 292–307. <https://doi.org/10.3758/s13414-020-02142-8>

FEATURE PROTECTION

- Huang, C., Vilotijević, A., Theeuwes, J., & Donk, M. (2021). Proactive distractor suppression elicited by statistical regularities in visual search. *Psychonomic Bulletin & Review*, 28(3), 918–927. <https://doi.org/10.3758/s13423-021-01891-3>
- Kleiner, M., Brainard, D. H., & Pelli, D. (2007). *What's new in Psychtoolbox-3?* 89.
- Kong, S., Li, X., Wang, B., & Theeuwes, J. (2020). Proactively location-based suppression elicited by statistical learning. *PLOS ONE*, 15(6), e0233544. <https://doi.org/10.1371/journal.pone.0233544>
- Lamy, D., Darnell, M., Levi, A., & Bublil, C. (2018). Testing the Attentional Dwelling Hypothesis of Attentional Capture. *Journal of Cognition*, 1(1), 43. <https://doi.org/10.5334/joc.48>
- Leber, A. B., Gwinn, R. E., Hong, Y., & O'Toole, R. J. (2016). Implicitly learned suppression of irrelevant spatial locations. *Psychonomic Bulletin & Review*, 23(6), 1873–1881. <https://doi.org/10.3758/s13423-016-1065-y>
- Lee, M. D., & Wagenmakers, E.-J. (2014). *Bayesian Cognitive Modeling: A Practical Course*. Cambridge University Press.
- Luck, S. J., Gaspelin, N., Folk, C. L., Remington, R. W., & Theeuwes, J. (2021). Progress toward resolving the attentional capture debate. *Visual Cognition*, 29(1), 1–21. <https://doi.org/10.1080/13506285.2020.1848949>
- Moher, J., & Egeth, H. E. (2012). The ignoring paradox: Cueing distractor features leads first to selection, then to inhibition of to-be-ignored items. *Attention, Perception, & Psychophysics*, 74(8), 1590–1605. <https://doi.org/10.3758/s13414-012-0358-0>

FEATURE PROTECTION

- Noonan, M. P., Adamian, N., Pike, A., Printzlau, F., Crittenden, B. M., & Stokes, M. G. (2016). Distinct Mechanisms for Distractor Suppression and Target Facilitation. *Journal of Neuroscience*, 36(6), 1797–1807. <https://doi.org/10.1523/JNEUROSCI.2133-15.2016>
- Pashler, H. (1988). Cross-dimensional interaction and texture segregation. *Perception & Psychophysics*, 43(4), 307–318. <https://doi.org/10.3758/BF03208800>
- Pelli, D. G. (1997). The VideoToolbox software for visual psychophysics: Transforming numbers into movies. *Spatial Vision*, 10(4), 437–442.
- Rouder, J. N. (2014). Optional stopping: No problem for Bayesians. *Psychonomic Bulletin & Review*, 21(2), 301–308. <https://doi.org/10.3758/s13423-014-0595-4>
- Sawaki, R., Geng, J. J., & Luck, S. J. (2012). A Common Neural Mechanism for Preventing and Terminating the Allocation of Attention. *Journal of Neuroscience*, 32(31), 10725–10736. <https://doi.org/10.1523/JNEUROSCI.1864-12.2012>
- Sawaki, R., & Luck, S. J. (2010). Capture versus suppression of attention by salient singletons: Electrophysiological evidence for an automatic attend-to-me signal. *Attention, Perception, & Psychophysics*, 72(6), 1455–1470. <https://doi.org/10.3758/APP.72.6.1455>
- Schönbrodt, F. D., & Wagenmakers, E.-J. (2018). Bayes factor design analysis: Planning for compelling evidence. *Psychonomic Bulletin & Review*, 25(1), 128–142. <https://doi.org/10.3758/s13423-017-1230-y>

FEATURE PROTECTION

- Schurigin, M. W., Wixted, J. T., & Brady, T. F. (2020). Psychophysical scaling reveals a unified theory of visual memory strength. *Nature Human Behaviour*, 4(11), Article 11.
<https://doi.org/10.1038/s41562-020-00938-0>
- Stilwell, B. T., Bahle, B., & Vecera, S. P. (2019). Feature-based statistical regularities of distractors modulate attentional capture. *Journal of Experimental Psychology: Human Perception and Performance*, 45(3), 419–433. <https://doi.org/10.1037/xhp0000613>
- Suchow, J. W., Brady, T. F., Fournie, D., & Alvarez, G. A. (2013). Modeling visual working memory with the MemToolbox. *Journal of Vision*, 13(10), 9. <https://doi.org/10.1167/13.10.9>
- Theeuwes, J. (1994). Stimulus-driven capture and attentional set: Selective search for color and visual abrupt onsets. *Journal of Experimental Psychology: Human Perception and Performance*, 20(4), 799–806. <https://doi.org/10.1037/0096-1523.20.4.799>
- Wang, B., & Theeuwes, J. (2018a). How to inhibit a distractor location? Statistical learning versus active, top-down suppression. *Attention, Perception, & Psychophysics*, 80(4), 860–870.
<https://doi.org/10.3758/s13414-018-1493-z>
- Wang, B., & Theeuwes, J. (2018b). Statistical regularities modulate attentional capture. *Journal of Experimental Psychology: Human Perception and Performance*, 44(1), 13–17.
<https://doi.org/10.1037/xhp0000472>
- Williams, J. R., Robinson, M. M., & Brady, T. F. (2022). There Is no Theory-Free Measure of “Swaps” in Visual Working Memory Experiments. *Computational Brain & Behavior*.
<https://doi.org/10.1007/s42113-022-00150-5>

FEATURE PROTECTION

- Won, B.-Y., Venkatesh, A., Witkowski, P. P., Banh, T., & Geng, J. J. (2022). Memory precision for salient distractors decreases with learned suppression. *Psychonomic Bulletin & Review*, 29(1), 169–181. <https://doi.org/10.3758/s13423-021-01968-z>
- Yantis, S., & Jonides, J. (1984). Abrupt visual onsets and selective attention: Evidence from visual search. *Journal of Experimental Psychology: Human Perception and Performance*, 10(5), 601–621. <https://doi.org/10.1037/0096-1523.10.5.601>
- Zhang, W., & Luck, S. J. (2008). Discrete fixed-resolution representations in visual working memory. *Nature*, 453(7192), Article 7192. <https://doi.org/10.1038/nature06860>