

Editor: Charles L. Folk

Special Issue
Reward Guides Visual Attention: Selection, Learning
and Motivation
Guest editors
Brian A. Anderson, Leonardo Chelazzi and Jan Theeuwes

Value-driven attentional capture is modulated by spatial context

Brian A. Anderson

To cite this article: Brian A. Anderson (2015) Value-driven attentional capture is modulated by spatial context, Visual Cognition, 23:1-2, 67-81, DOI: [10.1080/13506285.2014.956851](https://doi.org/10.1080/13506285.2014.956851)

To link to this article: <http://dx.doi.org/10.1080/13506285.2014.956851>



Published online: 23 Sep 2014.



Submit your article to this journal [↗](#)



Article views: 348



View related articles [↗](#)



View Crossmark data [↗](#)



Citing articles: 3 View citing articles [↗](#)

Value-driven attentional capture is modulated by spatial context

Brian A. Anderson

Department of Psychological and Brain Sciences, Johns Hopkins University, Baltimore, MD, USA

When stimuli are associated with reward outcome, their visual features acquire high attentional priority such that stimuli possessing those features involuntarily capture attention. Whether a particular feature is predictive of reward, however, will vary with a number of contextual factors. One such factor is spatial location: for example, red berries are likely to be found in low-lying bushes, whereas yellow bananas are likely to be found on treetops. In the present study, I explore whether the attentional priority afforded to reward-associated features is modulated by such location-based contingencies. The results demonstrate that when a stimulus feature is associated with a reward outcome in one spatial location but not another, attentional capture by that feature is selective to when it appears in the rewarded location. This finding provides insight into how reward learning effectively modulates attention in an environment with complex stimulus–reward contingencies, thereby supporting efficient foraging.

Keywords: Selective attention; Spatial attention; Reward learning; Contextual learning.

By selectively attending to certain stimuli and not others, organisms prioritize information in the environment for perceptual processing, determining which stimuli guide decision making and action. Attentional selection has long been characterized as arising from the interplay between goal-directed (e.g., Folk, Remington, & Johnston, 1992) and salience-driven mechanisms (Theeuwes, 1992, 2010; Yantis & Jonides, 1984). In order to promote survival and well-being, however, it is also important that the attention system selects stimuli associated with reward (Anderson, 2013). Recent evidence shows that attentional priority is modulated by the reward associated with visual stimuli (e.g., Della Libera & Chelazzi, 2009; Kiss, Driver, & Eimer, 2009; Krebs, Boehler, &

Please address all correspondence to Brian A. Anderson, Johns Hopkins University, Department of Psychological & Brain Sciences, 3400 N. Charles St., Baltimore, MD 21218-2686, USA. E-mail: bander33@jhu.edu

This research was supported by the NIH [grant numbers F31-DA033754 and R01-DA013165].

Woldorff, 2010; Raymond & O'Brien, 2009), and that the receipt of high reward strongly primes attentional selection (e.g., Della Libera & Chelazzi, 2006; Hickey, Chelazzi, & Theeuwes, 2010a, 2010b). When a stimulus feature is learned to predict a reward outcome, a bias to attend to stimuli possessing that feature develops such that these stimuli will involuntarily capture attention even when physically nonsalient, currently task irrelevant, and no longer associated with reward (e.g., Anderson, Laurent, & Yantis, 2011a, 2011b; Anderson & Yantis, 2012, 2013; Qi, Zeng, Ding, & Li, 2013). This automatic orienting of attention to stimuli previously associated with reward has been referred to as *value-driven attentional capture* (Anderson et al., 2011b).

Whether a stimulus feature is predictive of reward will vary according to contingencies that govern which reward-associated objects tend to be found in which contexts. For example, when foraging for food, red berries are likely to be found close to the ground in bushes, whereas yellow bananas are often found above the ground in treetops. Such location-based contingencies are known to have a strong influence on search strategy that is largely implicit. Searched-for targets are found more efficiently when they appear within a familiar spatial configuration of stimuli, a phenomenon referred to as *contextual cueing* (Chun & Jiang, 1998). Attention is biased towards locations that have been more likely to contain a target in the past, despite a lack of reported awareness of this target–location relationship (Jiang & Swallow, 2013; Jiang, Swallow, Rosenbaum, & Herzog, 2013).

An attentional bias for a particular region of space can also arise as a result of associative reward learning. When selecting a target stimulus in a particular location is associated with a comparatively large reward, targets subsequently appearing in that location are more quickly and accurately reported even when rewards are no longer available (Chelazzi et al., 2014; Sawaki & Raymond, this issue 2014). Although associative reward learning can influence attention to both stimulus features and spatial locations, whether the attention system is sensitive to the confluence of these two sources of visual information in predicting reward (i.e., reward is contingent upon a particular feature appearing in a particular location) is unknown.

Value-driven attentional selection is not limited to cases in which the properties of the stimulus and context match what has been rewarded in the past. Rather, the influence of associative reward learning on attention has been shown to be capable of transferring across stimuli and contexts. In the study by Sawaki and Raymond (this issue 2014), the observed location bias was evident even for stimuli appearing at the previously high-reward location that were themselves never rewarded. In another study in which comparatively high reward was associated with a stimulus feature (colour), different objects possessing that colour were preferentially attended in a different experimental task (Anderson, Laurent, & Yantis, 2012). Such generalization of value-based attentional priority can be adaptive, allowing the organism to leverage prior learning in newly

encountered contexts. However, as previously discussed, the reward value of a particular feature may vary reliably across spatial locations. When this is the case, can the value-driven attentional bias for a particular stimulus feature be location dependent? Is the attention system only sensitive to the aggregated value of a stimulus feature, abstracted from where it appears in the visual field, or is value-driven attentional priority for stimulus features modulated by learning about the locations in which a particular feature is predictive of reward?

In the present study, participants experienced a training phase in which targets of a particular colour were only rewarded when they appeared on a particular side of the display. In Experiment 1A, participants searched for a red target that was only followed by reward when presented on either the left or right side of the display. In the test phase, I examined whether value-driven attentional capture by a red stimulus would be specific to when that stimulus appeared in the location in which it was previously rewarded. Experiment 1B tested this same idea, but with two target colours each of which was only rewarded when appearing on a different side of the display (red on right, green on left, or vice versa). In this latter case, neither target colour nor target location was itself predictive of reward, which could only be predicted by the conjunction of target colour and target location. In both experiments, value-driven attentional capture by a previously reward-associated feature was found to be modulated by whether that feature appeared in a location within which it was rewarded during training.

METHODS

Experiment 1A

Participants. Sixteen participants were recruited from the Johns Hopkins University community. All reported normal or corrected-to-normal visual acuity and normal colour vision.

Apparatus. A Mac Mini equipped with Matlab software and Psychophysics Toolbox extensions (Brainard, 1997) was used to present the stimuli on an Asus VE247 monitor. The participants viewed the monitor from a distance of approximately 70 cm in a dimly lit room. Manual responses were entered using a standard keyboard.

Training phase.

Stimuli. Each trial consisted of a fixation display, a search array, and a feedback display (Figure 1A). The fixation display contained a white fixation cross ($0.8^\circ \times 0.8^\circ$ visual angle) presented in the centre of the screen against a black background, and the search array consisted of the fixation cross surrounded by six coloured circles (each $3.1^\circ \times 3.1^\circ$), three on each side of fixation. The middle of the three shapes on each side of the display was

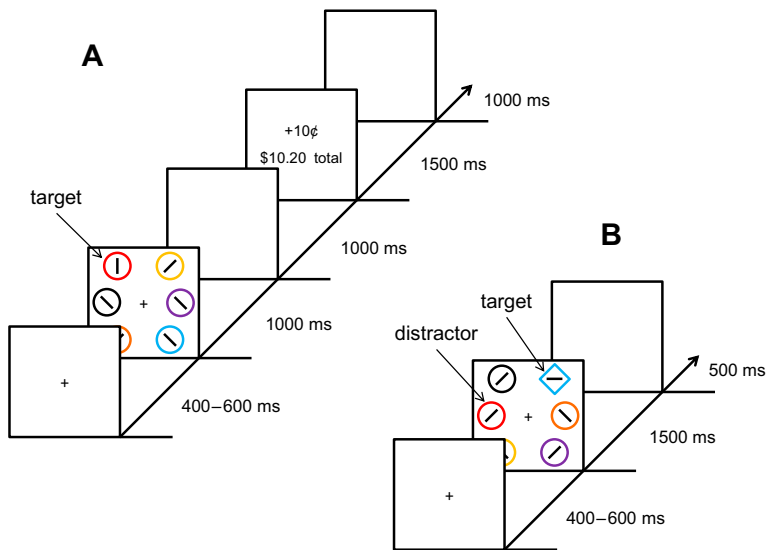


Figure 1. Sequence and time course of trial events. (A) Targets during the training phase were defined by colour, and participants reported the identity of the line segment inside of the target (vertical or horizontal) with a key press. Correct responses were followed by the delivery of monetary reward feedback, which varied based on the combination of target colour and target location. (B) During the test phase, the target was defined as the unique shape, and no reward feedback was provided. On half of the trials, one of the nontarget items—the distractor—was rendered in the colour of a formerly rewarded target. To view this figure in colour, please see the online issue of the Journal.

presented 7.3° centre-to-centre from fixation, and the two outer shapes were presented 5.7° from the vertical meridian, 5.5° above and below the horizontal meridian.

The target was a red circle, exactly one of which was presented on each trial. The colour of each nontarget circle was drawn from the set {green, blue, pink, orange, yellow, white} without replacement. A white bar appeared inside each of the six circles; for the target it was oriented either vertically or horizontally, and for each of the nontarget circles it was tilted at 45° to the left or to the right (randomly determined for each nontarget). The feedback display indicated the amount of monetary reward earned on the current trial, as well as the total accumulated reward.

Design. The target appeared in each of the six possible stimulus positions equally often. Correct identification of the oriented bar within the target was followed by a reward of 10¢ when the target appeared on one side of the display (right or left, counterbalanced across participants) and 0¢ feedback when it appeared on the other side.

Procedure. The training phase consisted of 360 trials, which were preceded by 48 practice trials. Each trial began with the presentation of the fixation display for a randomly varying interval of 400, 500, or 600 ms. The search array then appeared and remained on screen until a response was made or 1000 ms had elapsed, after which the trial timed out. The search array was followed by a blank screen for 1000 ms, the reward feedback display for 1500 ms, and a blank 1000 ms intertrial interval (ITI).

Participants made a forced-choice target identification by pressing the “z” and the “m” keys for the vertically and horizontally oriented bars within the targets, respectively. Correct responses were followed by monetary reward feedback in which either 10¢ or 0¢ was added to the participant’s total earnings, depending on the location of the target as outlined earlier. Incorrect responses were followed by feedback in which the word “Incorrect” was presented in place of the monetary increment, and responses that were too slow (i.e., no response before the trial timed out) were followed by a 500 ms 1000 Hz tone and no monetary increment (i.e., just the total earnings were presented in the feedback display).

Test phase.

Stimuli. Each trial consisted of a fixation display, a search array, and a feedback display (Figure 1B). The six shapes now consisted of either a diamond among circles or a circle among diamonds, and the target was defined as the unique shape. On a subset of the trials, one of the nontarget shapes was rendered in the colour of a formerly reward-associated target from the training phase (referred to as the *valuable distractor*); the target shape was never the colour of a target from the training phase. The feedback display only informed participants if their prior response was correct or not.

Design. Target identity, target location, distractor identity, and distractor location were fully crossed and counterbalanced, and trials were presented in a random order. Thus, both the target shape and the distractor shape varied unpredictably from trial to trial. Red (i.e., valuable) distractors were presented on 50% of all trials; the remaining trials contained no red stimulus (distractor absent trials).

Procedure. Participants were instructed to ignore the colour of the shapes and to focus on identifying the oriented bar within the unique shape using the same orientation-to-response mapping. The test phase consisted of 480 trials, which were preceded by 32 practice (distractor absent) trials. The search array was followed immediately by nonreward feedback (the word “Incorrect”) for 1000 ms in the event of an incorrect response (this display was omitted following a correct response) and then by a 500 ms ITI; no monetary rewards were given in the test phase, and the task instructions made no reference to reward. Trials timed out after 1500 ms. As in the training phase, if the trial timed out, the computer

emitted a 500 ms 1000 Hz tone. Upon completion of the experiment, participants were paid the cumulative reward they had earned in the training phase.

Exit question. At the conclusion of the test phase, participants were asked to select which of three statements they believed best described the reward contingencies in the training phase (see [Appendix](#)).

Data analysis. Only correct responses were included in all analyses of RT, and RTs more than three *SDs* above or below the mean of their respective condition for each participant were trimmed. This resulted in a reduction of <1% of all trials.

Experiment 1B

Participants. Twelve new participants were recruited from the Johns Hopkins University community. All reported normal or corrected-to-normal visual acuity and normal colour vision.

Apparatus. The apparatus was identical to that used in Experiment 1A.

Training phase.

Stimuli. Each trial consisted of a fixation display, a search array, and a feedback display as in Experiment 1A ([Figure 1A](#)). Experiment 1B differed in that the target was now defined as the red or green circle, exactly one of which was present in the display on each trial. The colour of each nontarget circle was drawn from the set {cyan, blue, pink, orange, yellow, white} without replacement.

Design. Each colour target appeared in each location equally often. The amount of reward that could be earned on each trial was determined by the conjunction of target colour and target location. Each colour target was rewarded 10¢ for correct identification when it appeared on a particular side of the display, and 0¢ when appearing on the other side of the display. For each participant, one colour target (counterbalanced across participants) was rewarded when appearing on the right side of the display while the other was rewarded when appearing on the left side of the display—therefore, neither colour nor location alone predicted reward, but reward was predicted by the conjunction of target colour and location (e.g., red on the right and green on the left).

Procedure. The procedure was identical to that of Experiment 1A, with the exception that the training phase consisted of 480 trials and correct responses were followed by 10¢ or 0¢ according to the contingencies outlined earlier.

Test phase.

Stimuli. Each trial consisted of a fixation display, a search array, and a feedback display as in Experiment 1A ([Figure 1B](#)). All that differed in

Experiment 1B was that the valuable distractor was now equally often red and green (rather than only red), and cyan was included in the colour set as in the preceding training phase.

Design. Target identity, target location, distractor identity, and distractor location were fully crossed and counterbalanced, and trials were presented in a random order. Half of the trials contained a valuable distractor (red or green nontarget), and half did not (distractor absent trials). Red and green distractors were presented equally often on distractor present trials (i.e., each colour on 25% of all total trials), with each colour distractor appearing equally often in each of the six possible stimulus positions.

Procedure. The procedure was identical to that of Experiment 1A.

Exit question. At the conclusion of the test phase, participants were asked to select which of six statements they believed best described the reward contingencies in the training phase (see [Appendix](#)). Due to experimenter error, one of the participants was not administered the exit question.

Data analysis. Only correct responses were included in all analyses of RT, and RTs more than three *SDs* above or below the mean of their respective condition for each participant were trimmed. This resulted in a reduction of <2% of all trials.

RESULTS

Experiment 1A

Training phase. Participants were not significantly faster, $t(15) = 1.19$, $p = .255$, or more accurate, $t(15) = 0.17$, $p = .867$, to report the target when it appeared on the rewarded compared to the unrewarded side of the display (see [Table 1](#)). This is consistent with previous findings and suggests that in simple search tasks such as the one used here, top-down goals favour targets regardless of reward value (e.g., Anderson et al., 2011a, 2012; Anderson, Faulkner, Rilee, Yantis, & Marvel, 2013). As the training task emphasized accuracy in order to

TABLE 1
Mean response time and accuracy by target location in the training phase,
separately for each experiment

	<i>Experiment 1A</i>		<i>Experiment 1B</i>	
	<i>Unrewarded</i>	<i>Rewarded</i>	<i>Unrewarded</i>	<i>Rewarded</i>
Response time (ms)	545	538	580	583
Accuracy	96.0%	96.2%	95.1%	94.4%

obtain reward, participants may also have responded conservatively, making RT a potentially insensitive measure to detect value-based effects. Most importantly, however, the training phase provided participants with the opportunity to experience the experimental reward contingencies, and the effect of this experience on involuntary attentional selection was examined in the test phase.

Test phase. A repeated measures analysis of variance (ANOVA) with distractor condition (absent, unrewarded location, rewarded location) as a factor revealed a marginally significant main effect, $F(2, 30) = 2.72, p = .082, \eta_p^2 = .153$ (see Figure 2). Planned orthogonal comparisons revealed that RT was significantly slower when the distractor was presented in a location in which it was previously rewarded compared to the other two conditions (averaged together), $t(15) = 2.43, p = .028, d = 0.61$, which did not significantly differ, $t(15) = 0.34, p = .736$. Thus, the red distractor captured attention when presented in a location in which it was previously rewarded, but not when it appeared in a location in which it was never rewarded. There was no main effect of distractor condition evident in accuracy, $F(2, 30) = 1.02, p = .375$ (91.7%, 91.1%, and 92.2% across the absent, unrewarded, and rewarded distractor conditions, respectively).

Collapsing across distractor condition, participants were not significantly faster to report the shape target in the test phase when it appeared on the side of the display in which the red target was rewarded during training, mean difference = 7 ms, $t(15) = 0.96, p = .352$. This suggests that a purely spatial bias, independent of feature information, was weak to nonexistent. Instead, the combination of feature and location had an especially strong effect on attentional selection, above and beyond either alone.

Experiment 1B

Training phase. Participants were not significantly faster, $t(11) = -0.98, p = .347$, or more accurate, $t(15) = -1.00, p = .337$, to report the target when it

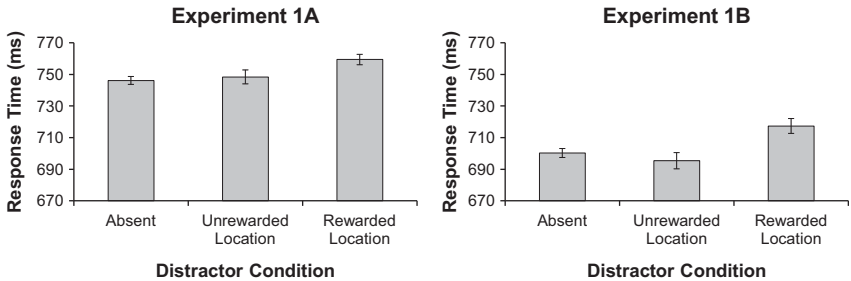


Figure 2. Mean response time by distractor condition in the test phase, separately for each experiment. Error bars reflect the within-subjects SEM.

appeared on the side of the display in which its colour was rewarded, mirroring the results from Experiment 1A (see Table 1).

Test phase. A repeated measures ANOVA with distractor condition (absent, unrewarded location, rewarded location) as a factor revealed a significant main effect, $F(2, 22) = 4.73$, $p = .020$, $\eta_p^2 = .301$ (see Figure 2). As in the prior experiment, planned orthogonal comparisons revealed that RT was significantly slower when the distractor was presented in a location in which it was previously rewarded compared to the other two conditions (averaged together), $t(11) = 2.78$, $p = .018$, $d = 0.80$, which did not significantly differ, $t(11) = -0.71$, $p = .495$. There was no main effect of distractor condition evident in accuracy, $F < 1$ (92.8%, 93.5%, and 92.6% across the absent, unrewarded, and rewarded distractor conditions, respectively). Thus, even with more complex contingencies in which only the combination of a particular colour in a particular location predicts reward, value-driven attentional capture is selective for when this combination matches what has been rewarded in the past.

Combined analysis

Collapsing across experiment, the location in which a distractor feature had been rewarded had a robust influence on RT in the test phase, $F(2, 54) = 7.02$, $p = .002$, $\eta_p^2 = .208$. RT was slower when a distractor was presented in a location in which it was previously rewarded compared to when the very same stimulus was presented in a location in which it was never rewarded, $t(27) = 2.70$, $p = .012$, $d = 0.51$; while the former captured attention when compared to distractor absent trials, mean difference = 15 ms, $t(27) = 4.49$, $p < .001$, $d = 0.85$, the latter did not, mean difference = -1 ms, $t(27) = -0.18$, $p = .861$.

Exit question. In Experiment 1A, 11 of the 16 participants indicated that the rewards were random, three indicated the correct contingency, and two indicated the incorrect contingency. In Experiment 1B, seven out of 11 participants indicated that the rewards were random, two indicated the correct contingency, and two indicated an incorrect contingency. Across both experiments, the number of participants indicating the correct contingency was less than what would be expected by random guessing.

DISCUSSION

The present study demonstrates that when a stimulus feature (in this case, colour) is associated with a reward outcome in one spatial location but not another, value-driven attentional capture by a stimulus possessing that feature is modulated by the location within which it appears. Specifically, when the combination of feature and location match what has been rewarded in the past,

value-driven attentional capture by that feature is observed. In contrast, when that same feature appears in a location within which it has gone unrewarded, it does not produce evidence of attentional capture.

In Experiment 1A, a single target feature was selectively rewarded on one side of the display during training. In the test phase of this experiment, stimuli possessing this feature only captured attention when appearing in the previously rewarded location. Such selectivity could be explained by either a bias to attend to a particular feature appearing in a particular spatial position, or two separate biases, one for the reward-associated feature and one for the reward-associated location, working in tandem to guide selection. However, in Experiment 1B, each of the two target-defining features and each of the two sides of the display was alone unpredictable of reward, which could only be predicted from the confluence of a particular feature in a particular location. Thus, the selectivity of value-driven attentional capture in the test phase of this experiment can only be explained by a bias that is more narrowly tuned to specific combinations of feature and location information.

Interestingly, in Experiment 1A, the observed value-based attentional bias was found to be specific to the previous target-defining feature. Although rewards were only delivered for stimuli appearing on one particular side of the display, a more general bias to attend to that region of space was not found to significantly benefit the processing of a shape-defined target. On the surface, this conflicts with previous studies reporting attentional biases for stimuli appearing in previously reward-predictive locations (Chelazzi et al., 2014; Sawaki & Raymond, this issue 2014). There are differences in the experimental design used in the present study that likely contributed to this difference. First, the target feature during training was consistent across trials in Experiment 1A, making the bound representation of colour and location equally as predictive of reward as location alone. Second, the target during the test phase was defined by its relative salience (shape singleton), encouraging a broad distribution of attention across the entire stimulus array. The fact that the previously rewarded target feature captured attention in the test phase of Experiment 1A, but only when appearing in a particular location, demonstrates that feature-based attentional biases arising from reward history can be modulated by spatial context, a conclusion corroborated by Experiment 1B.

The findings from the present study provide the first evidence that value-driven attentional priorities can be sensitive to contextual information. Rather than associate a colour with reward without regard to spatial context, which would have produced equivalent attentional capture across all spatial locations, the attentional priority for colour as a function of reward history was contingent upon where that feature appears in visual space. This contrasts with other studies demonstrating the ability of value-based attentional priorities to generalize across stimuli, locations, and tasks (Anderson et al., 2012; Sawaki & Raymond, this issue 2014). A critical difference between the present study and these previous

studies is that in these previous studies reward was entirely predicted by either feature or location alone. Thus, it appears to be the case that the attention system defaults to context-general representations of stimulus value when contextual information is itself nonpredictive of reward, but is capable of incorporating contextual information when such information predicts whether a feature will be rewarded or not. In this sense, organisms are poised to exploit previous reward learning in newly encountered contexts, but can appropriately limit the influence of that learning based on context when doing so is supported by the reward structure, thereby avoiding overgeneralization of learning.

The mechanisms by which spatial context modulate value-driven attentional biases for stimulus features are unclear. One possibility is that the combination of feature identity and spatial position are necessary to generate a bias signal that guides selection. Another possibility is that a reward-associated feature always generates a bias signal regardless of where it is presented, but this bias signal is suppressed when the context of that feature suggests that expected value should be low. Assessment of the processing of nontargets as a function of spatial context, potentially using neuroimaging methods, might provide insight into this issue by allowing for direct measurement of suppression. A related question concerns the locus of the modulation of attentional priority. The observed contextual modulation is consistent with a top-down influence on value-based attentional priority resulting from feedback from higher-level visual representations, but a biasing of stimulus-driven visual input remains an equally tenable explanation. The brain's representation of elementary visual features such as colour is retinotopically organized (e.g., Johnson, Hawken, & Shapley, 2008), and even higher-level representations of complex visual objects are sensitive to the position of these objects in space (e.g., Kravitz, Kriegeskorte, & Baker, 2010; Kravitz, Vinson, & Baker, 2008). To the degree that the observed findings reflect changes in the tuning of stimulus-driven visual processing, value-driven attentional capture should reflect an egocentric, or person-centred, orientation, which is true of attentional biases for high-probability target locations (Jiang & Swallow, 2013). Alternatively, to the degree that the observed modulation of value-driven attention reflects top-down feedback signals, it should be evident for other, more complex forms of contextual information that are not bound to feature representations.

Value-driven attentional capture has been shown to reflect both covert and overt orienting (Anderson et al., 2011a, 2011b; Anderson & Yantis, 2012; Buckner, Belopolsky, & Theeuwes, this issue 2014; Theeuwes & Belopolsky, 2012; Tran, Pearson, Donkin, Most, & Le Pelley, this issue 2014). Indeed, covert and overt attention are interrelated, with covert attention guiding eye movements (e.g., Deubel & Schneider, 1996; Hoffman & Subramaniam, 1995; Thompson & Bichot, 2005). The slowing of RT observed in the present study might reflect contribution from either or both of these selection mechanisms. However, it is important to note that these findings cannot be explained by anticipatory eye

movements and instead reflect reactive mechanisms of control driven by the relationship between stimulus properties and prior learning. No significant bias was observed for targets appearing in the previously reward-associated location during the test phase of Experiment 1A, and both locations were overall equally predictive of reward in Experiment 1B, precluding the selection of a particular location prior to the onset of the stimulus array as an explanation for the spatially specific capture observed in the present study.

Interestingly, participants were largely unable to correctly report which of several reward contingencies were in place during the training phase when provided with a forced-choice question, despite the fact that the actual contingency was 100% predictive of reward. This is consistent with the reward learning that automatically guides attention being implicit in nature, relying on the co-occurrence of visual information and reward feedback rather than the establishment of strategic priorities that persist due to reinforcement, as has been suggested previously and elsewhere (e.g., Anderson et al., 2013; Anderson & Yantis, 2013; Buckner et al., this issue 2014; Della Libera, Perlato, & Chelazzi, 2011; Sali, Anderson, & Yantis, 2014; Tran et al., this issue 2014). However, it should be noted that the evidence provided by the forced-choice question is only suggestive of implicit learning and cannot rule out awareness of the reward contingencies that either extinguished over the course of the test phase or was not sufficiently strong that participants were willing to endorse the correct contingency over random contingencies.

The present study provides insights into how the attention system supports efficient foraging. Rather than relying exclusively on goals and strategies in order to inform when and where to search for particular objects, my findings show that reward learning can automatically guide attention in a way that is sensitive to complex, situationally dependent reward contingencies. By tuning attentional priorities in accordance with the co-occurrence of visual events and reward outcomes, organisms can locate valuable stimuli in the future with minimal effort. Such automatic guidance is surprisingly efficient, taking into account multiple sources of information in biasing selection. This efficiency might help explain why value-driven attention is not easily overridden by goal-directed attentional control mechanisms: The more efficient value-driven attention is, the less of a need there will be for the organism to have to override value-based selection.

REFERENCES

- Anderson, B. A. (2013). A value-driven mechanism of attentional selection. *Journal of Vision*, 13(3), 1–16. doi:10.1167/13.3.7
- Anderson, B. A., Faulkner, M. L., Rilee, J. J., Yantis, S., & Marvel, C. L. (2013). Attentional bias for non-drug reward is magnified in addiction. *Experimental and Clinical Psychopharmacology*, 21, 499–506. doi:10.1037/a0034575

- Anderson, B. A., Laurent, P. A., & Yantis, S. (2011a). Learned value magnifies salience-based attentional capture. *PLoS ONE*, 6(11), e27926. doi:[10.1371/journal.pone.0027926](https://doi.org/10.1371/journal.pone.0027926)
- Anderson, B. A., Laurent, P. A., & Yantis, S. (2011b). Value-driven attentional capture. *Proceedings of the National Academy of Sciences, USA*, 108, 10367–10371. doi:[10.1073/pnas.1104047108](https://doi.org/10.1073/pnas.1104047108)
- Anderson, B. A., Laurent, P. A., & Yantis, S. (2012). Generalization of value-based attentional priority. *Visual Cognition*, 20, 647–658. doi:[10.1080/13506285.2012.679711](https://doi.org/10.1080/13506285.2012.679711)
- Anderson, B. A., & Yantis, S. (2012). Value-driven attentional and oculomotor capture during goal-directed, unconstrained viewing. *Attention, Perception, and Psychophysics*, 74, 1644–1653. doi:[10.3758/s13414-012-0348-2](https://doi.org/10.3758/s13414-012-0348-2)
- Anderson, B. A., & Yantis, S. (2013). Persistence of value-driven attentional capture. *Journal of Experimental Psychology: Human Perception and Performance*, 39, 6–9. doi:[10.1037/a0030860](https://doi.org/10.1037/a0030860)
- Brainard, D. H. (1997). The psychophysics toolbox. *Spatial Vision*, 10, 433–436. doi:[10.1163/156856897X00357](https://doi.org/10.1163/156856897X00357)
- Buckner, B., Belopolsky, A. V., & Theeuwes, J. (2014). Distractors that signal reward capture the eyes: Automatic reward capture without previous goal-directed selection. *Visual Cognition*.
- Chelazzi, L., Estocinova, J., Calletti, R., Lo Gerfo, E., Sani, I., Della Libera, C., & Santandrea, E. (2014). Altering spatial priority maps via reward-based learning. *Journal of Neuroscience*, 34, 8594–8604. doi:[10.1523/JNEUROSCI.0277-14.2014](https://doi.org/10.1523/JNEUROSCI.0277-14.2014)
- Chun, M. M., & Jiang, Y. (1998). Contextual cueing: Implicit learning and memory of visual context guides spatial attention. *Cognitive Psychology*, 36, 28–71. doi:[10.1006/cogp.1998.0681](https://doi.org/10.1006/cogp.1998.0681)
- Della Libera, C., & Chelazzi, L. (2006). Visual selective attention and the effects of monetary reward. *Psychological Science*, 17, 222–227. doi:[10.1111/j.1467-9280.2006.01689.x](https://doi.org/10.1111/j.1467-9280.2006.01689.x)
- Della Libera, C., & Chelazzi, L. (2009). Learning to attend and to ignore is a matter of gains and losses. *Psychological Science*, 20, 778–784. doi:[10.1111/j.1467-9280.2009.02360.x](https://doi.org/10.1111/j.1467-9280.2009.02360.x)
- Della Libera, C., Perlato, A., & Chelazzi, L. (2011). Dissociable effects of reward on attentional learning: From passive associations to active monitoring. *PLoS ONE*, 6(4), e19460. doi:[10.1371/journal.pone.0019460](https://doi.org/10.1371/journal.pone.0019460)
- Deubel, H., & Schneider, W. X. (1996). Saccade target selection and object recognition: Evidence for a common attentional mechanism. *Vision Research*, 36, 1827–1837. doi:[10.1016/0042-6989\(95\)00294-4](https://doi.org/10.1016/0042-6989(95)00294-4)
- 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, 1030–1044.
- Hickey, C., Chelazzi, L., & Theeuwes, J. (2010a). Reward changes salience in human vision via the anterior cingulate. *Journal of Neuroscience*, 30, 11096–11103. doi:[10.1523/JNEUROSCI.1026-10.2010](https://doi.org/10.1523/JNEUROSCI.1026-10.2010)
- Hickey, C., Chelazzi, L., & Theeuwes, J. (2010b). Reward guides vision when it's your thing: Trait reward-seeking in reward-mediated visual priming. *PLoS ONE*, 5(11), e14087. doi:[10.1371/journal.pone.0014087](https://doi.org/10.1371/journal.pone.0014087)
- Hoffman, J. E., & Subramaniam, B. (1995). The role of visual attention in saccadic eye movements. *Perception and Psychophysics*, 57, 787–795. doi:[10.3758/BF03206794](https://doi.org/10.3758/BF03206794)
- Jiang, Y. V., & Swallow, K. M. (2013). Spatial reference frame of incidentally learned attention. *Cognition*, 126, 378–390. doi:[10.1016/j.cognition.2012.10.011](https://doi.org/10.1016/j.cognition.2012.10.011)
- Jiang, Y. V., Swallow, K. M., Rosenbaum, G. M., & Herzog, C. (2013). Rapid acquisition but slow extinction of an attentional bias in space. *Journal of Experimental Psychology: Human Perception and Performance*, 39, 87–99. doi:[10.1037/a0027611](https://doi.org/10.1037/a0027611)
- Johnson, E. N., Hawken, M. J., & Shapley, R. (2008). The orientation selectivity of color-responsive neurons in macaque V1. *Journal of Neuroscience*, 28, 8096–8106. doi:[10.1523/JNEUROSCI.1404-08.2008](https://doi.org/10.1523/JNEUROSCI.1404-08.2008)

- Kiss, M., Driver, J., & Eimer, M. (2009). Reward priority of visual target singletons modulates event-related potential signatures of attentional selection. *Psychological Science*, 20, 245–251. doi:[10.1111/j.1467-9280.2009.02281.x](https://doi.org/10.1111/j.1467-9280.2009.02281.x)
- Kravitz, D. J., Kriegeskorte, N., & Baker, C. I. (2010). High-level visual object representations are constrained by position. *Cerebral Cortex*, 20, 2916–2925. doi:[10.1093/cercor/bhq042](https://doi.org/10.1093/cercor/bhq042)
- Kravitz, D. J., Vinson, L. D., N., & Baker, C. I. (2008). How position dependent is visual object recognition. *Trends in Cognitive Sciences*, 12, 114–122. doi:[10.1016/j.tics.2007.12.006](https://doi.org/10.1016/j.tics.2007.12.006)
- Krebs, R. M., Boehler, C. N., & Woldorff, M. G. (2010). The influence of reward associations on conflict processing in the Stroop task. *Cognition*, 117, 341–347. doi:[10.1016/j.cognition.2010.08.018](https://doi.org/10.1016/j.cognition.2010.08.018)
- Qi, S., Zeng, Q., Ding, C., & Li, H. (2013). Neural correlates of reward-driven attentional capture in visual search. *Brain Research*, 1532, 32–43. doi:[10.1016/j.brainres.2013.07.044](https://doi.org/10.1016/j.brainres.2013.07.044)
- Raymond, J. E., & O'Brien, J. L. (2009). Selective visual attention and motivation: The consequences of value learning in an attentional blink task. *Psychological Science*, 20, 981–988. doi:[10.1111/j.1467-9280.2009.02391.x](https://doi.org/10.1111/j.1467-9280.2009.02391.x)
- Sali, A. W., Anderson, B. A., & Yantis, S. (2014). The role of reward prediction in the control of attention. *Journal of Experimental Psychology: Human Perception and Performance*, 40(4), 1654–1664.
- Sawaki, R., & Raymond, J. E. (2014). Irrelevant spatial value learning modulates visual search. *Visual Cognition*.
- Theeuwes, J. (1992). Perceptual selectivity for color and form. *Perception and Psychophysics*, 51, 599–606. doi:[10.3758/BF03211656](https://doi.org/10.3758/BF03211656)
- Theeuwes, J. (2010). Top-down and bottom-up control of visual selection. *Acta Psychologica*, 135, 77–99. doi:[10.1016/j.actpsy.2010.02.006](https://doi.org/10.1016/j.actpsy.2010.02.006)
- Theeuwes, J., & Belopolsky, A. V. (2012). Reward grabs the eye: Oculomotor capture by rewarding stimuli. *Vision Research*, 74, 80–85.
- Thompson, K. G., & Bichot, N. P. (2005). A visual salience map in the primate frontal eye field. *Progress in Brain Research*, 147, 251–262.
- Tran, S., Pearson, D., Donkin, C., Most, S. B., & Le Pelley, M. (2014). Cognitive control and counterproductive oculomotor capture by reward-related stimuli. *Visual Cognition*.
- Yantis, S., & Jonides, J. (1984). Abrupt visual onsets and selective attention: Evidence from visual search. *Journal of Experimental Psychology: Human Perception and Performance*, 10, 350–374. doi:[10.1037/0096-1523.10.5.601](https://doi.org/10.1037/0096-1523.10.5.601)

APPENDIX: Questions used to assess awareness of the stimulus–reward contingencies

Which option do you believe best describes the part of the experiment in which you were earning money (please choose only one):

Experiment 1A

- (1) The red circle was generally worth more when it appeared on the right side of the screen
- (2) The red circle was generally worth more when it appeared on the left side of the screen
- (3) How much money I received was random and unrelated to where the red circle appeared

Experiment 1B

- (1) The red circle was generally worth more than the green circle regardless of which side of the screen it appeared on
- (2) The green circle was generally worth more than the red circle regardless of which side of the screen it appeared on

- (3) The two circles were worth the same overall, but one colour was worth more when it appeared on the left side of the screen and the other was worth more when it appeared on the right side of the screen
- (4) Both colour circles were generally worth more when presented on the left side of the screen
- (5) Both colour circles were generally worth more when presented on the right side of the screen
- (6) How much money I received was random and unrelated to both colour and location