

From Feature to Dimension:

Generalization of Value-Driven Attentional Capture into a Multidimensional World

Jennifer Bu

Princeton University

This thesis was submitted to Princeton University in partial fulfillment of the requirements for the Degree of Bachelor Arts in Psychology.

Adviser: Yael Niv

April 2017

Honor Pledge

I pledge my honor that this paper represents my own work in accordance with University regulations.

/s/ Jennifer Bu

Acknowledgements

To Professor Yael Niv, who has unfailingly offered her insight and mentorship throughout every step of this independent work journey. From the moment I joined her lab, Yael has always encouraged me to put in my best effort and has given me the independence to pursue my interests. I cannot be thankful enough for the support Yael has provided as an adviser, both in guiding my research and in fostering a fun and engaging lab atmosphere. She has given me everything I could ask for in an undergraduate research experience, and has inspired me to continue pursuing research in the future. Her passion for her field is contagious, and her commitment to her students is truly remarkable.

To Angela Radulescu, who has advised my research project through thick and thin, and who has quite literally taught me all the research methods and skills that went into this thesis. Where do I even begin? Before I met her, I had never heard of an ANOVA or IRB, and I had never even seen a line of MATLAB! I will genuinely miss all the meetings we spent mulling over puzzling data or brainstorming new experimental designs, and it is for all these challenges that I continue to fall in love with research over and over again. This thesis would not have been possible without her steadfast support.

To Nick Turk-Browne, who co-advised this thesis and who always had thoughtful commentary to share: we will be continuing this conversation soon! To the rest of the Niv Lab, with whom I have learned so much through casual interactions and from their lively engagement in lab meetings: I thoroughly enjoyed working in the company of you all! To the Princeton Psychology Department, who gave me the opportunity to pursue my passions: thank you for letting me join the department during my junior winter, and for supporting me ever since. It was one of the best decisions I could have made during my time here at Princeton. A special thank

you to Casey Lew-Williams, Lauren Emberson, and Justin Junge, who have each taught me, shared their experiences with me, and inspired me in their own way to pursue a career in this field.

To the Office of the Dean of the College and the Luce Summer Program fund for providing me with the financial support to pursue my research.

To my friends, who have always stood by me and with whom I will always stand by. Nicole, Natsuko, Joyce, and Joy: the walk from PNI to Terrace, which I must have made over 100 times now, was always better because I knew I could see your smiling faces at the end of it! To Jennifer, Angelica, and Alice: food and lab buddies forever. To Annie, Em, and Amy: Scully 240 is love.

Finally, to my sister, parents, and grandparents back home in California. You have been with me since day one, and if I listed everything that you have done for me, it would far exceed the length of this thesis. Words can do no justice.

I am forever grateful for all of these individuals who have brought me to where I am today, and who have continued to show me the way forward.

Jennifer Bu

Abstract

Some material in this abstract was taken from Bu et al. (2017).

Reward can help teach selective attention what is relevant in a complex world. Previous studies have demonstrated that specific features consistently associated with high reward capture attention more than those associated with low reward. How reward influences attentional capture at the level of entire feature dimensions remains unclear. In this thesis, I test the hypothesis that learning to attend to highly rewarding features in one dimension generalizes to attentional capture by unrewarded features within that same dimension. I used two variants of a behavioral paradigm named the Token Task. In this task, participants learned to associate certain features from the dimensions of color and orientation with high reward, and others with low reward. Participants were randomly assigned to either a “color” or an “orientation” group, and all highly rewarded features were drawn from the corresponding dimension. To probe attention, I occasionally interrupted the learning task with a visual search task. Each search array contained both a color singleton and an orientation singleton, one of which was randomly chosen each trial to be the search target. I predicted that a learned dimensional attention bias would facilitate pop out and thus search for singleton targets in the dimension containing the highly rewarded feature. The findings suggest that reward can drive attentional capture at the dimensional level, at least for color. Additionally, the findings provide evidence of separate learning styles during the task that can differentially drive attentional capture. An experimental replication with variation demonstrated that the design of the visual search array is crucial to detecting an effect of dimensional attention.

Table of Contents

Acknowledgements.....	3
Abstract.....	5
Chapter 1 Introduction	8
1.1 Selective Attention in Everyday Life.....	8
1.2 What Drives Selective Attention?.....	8
1.3 Reward and Attention.....	10
1.4 Mechanisms of Value-Driven Attention.....	13
1.5 Generalizability into a Complex World.....	15
1.6 From Feature to Dimension.....	16
Chapter 2 General Methods	19
2.1 The Token Task.....	19
2.1.1 Collect Trials and Reward Structure.....	20
2.1.2 Visual Search Trials as a Probe of Dimensional Attention.....	22
2.1.3 Predict Trials as a Measure of Learning.....	23
2.2 Experimental Overview.....	24
2.2.1 Hypothesis.....	24
2.2.2 The Present Experiments.....	25
Chapter 3 Experiment 1: The Token Task	26
3.1 Methods.....	26
3.1.1 Participants.....	26
3.1.2 Stimuli.....	26
3.1.3 Procedure.....	29
3.1.4 Statistical Analyses.....	31
3.2 Results.....	31
3.2.1 Trial Accuracy.....	31
3.2.2 Learning Performance.....	32
3.2.3 Reward and Attention.....	34
3.2.4 Exploratory Analyses.....	36
3.3 Discussion.....	42

Chapter 4	Experiment 2: A Replication with Variation	47
4.1	Methods.....	47
4.1.1	Participants.....	47
4.1.2	Stimuli.....	48
4.1.3	Procedure and Statistical Analyses.....	50
4.2	Results.....	50
4.2.1	Trial Accuracy.....	50
4.2.2	Learning Performance.....	51
4.2.3	Reward and Attention.....	52
4.2.4	Exploratory Analyses.....	53
4.3	Discussion.....	55
4.3.1	Reconciling Experiments 1 and 2.....	56
4.3.2	Lessons Learned.....	59
Chapter 5	General Discussion	61
5.1	The Effect of Learning on Dimensional Attention.....	61
5.2	Experimental Limitations.....	64
5.3	Future Directions.....	68
References		70
Appendix A: Debriefing Questionnaire		78
Form for Collaboration in Senior Thesis Work		79
Approval Form for Undergraduate Research Involving Animals		81
Approval Form for Undergraduate Research Involving Human Subjects		82

Chapter 1

Introduction

Some material in this chapter was adapted from Bu (2016).

1.1 Selective Attention in Everyday Life

If we were aware of everything that our sensory systems were capable of processing, we would be paralyzed by the sheer overload of information we would receive at any given time. The real world environment is complex, and processing it in a manner most practical for survival presents a significant challenge to the brain. The brain's attentional network, in a process called selective attention, provides a solution by actively determining which stimuli will receive additional cognitive processing. We engage in selective attention almost constantly like when we read in a crowded coffee shop or immediately spot a deer in the woods. The commonality in each of these cases is that attention is prioritizing certain stimuli, like a book or a deer, at the expense of surrounding and less relevant ones. How is the brain able to so quickly decide what is relevant, and direct attention accordingly?

1.2 What Drives Selective Attention?

The conventional framework of selective attention consists of an endogenous “top-down” mechanism in which attention is voluntarily deployed to a goal, and an exogenous “bottom-up” mechanism in which attention is involuntarily captured by the most salient, or noticeable, stimuli (Yantis, 2000). This dichotomy suggests that attentional priority is primarily determined by two factors: the extent to which a stimulus perceptually “pops out” in relation to other stimuli based on physical properties such as color contrast or sudden appearance (Theeuwes, 1992; Yantis & Jonides, 1984), and the extent to which attention is strategically directed in accordance with

one's internal goals (Bricolo et al., 2002). However, factors outside of the conventional dichotomy can drive selective attention as well. Attentional selection also depends on the subjective significance that a stimulus has acquired over time through learning and experience. For example, a gun in a visual scene is easier to search for than a tree, and a burger is easier to search for than a couch (Oca & Black, 2017). A snake in the grass is more likely to receive attentional priority than a chair on the lawn (Flykt, Esteves, Ohman, Flykt, & Esteves, 2011), and a twenty-dollar bill lying on the street is more likely to capture our attention than a crinkled flyer (Anderson, Laurent, & Yantis, 2011). To illustrate the example of the dollar bill on the ground, the light green paper is not particularly salient based on its physical features, and we are not actively searching for dollar bills when we walk on the street. Indeed, factors such as threat, emotion, reward, and pleasure can all spontaneously influence the deployment of selective attention independent of the classical exogenous and endogenous framework.

Zhao, Al-Aidroos, and Turk-Browne (2013) identified an even subtler factor that drives selective attention. They found that statistical learning spontaneously biases attention toward the location, feature, or dimension of a regularity, once again independent of goal relevance or stimulus salience. They demonstrated this through a series of experiments involving the presentation of shape streams that contained regularities in a location, feature, or dimension. Attentional bias was probed through reaction times in a visual search task. Critically, whenever the target of the search matched the location, feature, or dimension with the regularity, visual search times were faster; whenever the three sources of regularity served as distractors, search times were slower. This showed that statistical learning of regularity biases attentional priority.

Clearly, selective attention is a flexibly adaptive and computationally complex process, and it is driven by daily experiences from the world in an effort to constantly learn how to attend

(Dayan, Kakade, & Montague, 2000). In this thesis, I seek to expand Zhao et al.'s experimental design and findings on the effects of statistical learning on dimensional attention to a body of literature suggesting that reward learning also drives attention. Specifically, I explore the mechanisms behind how reward can teach selective attention what is worth attending to in the environment, and I focus on the phenomenon of value-driven attentional capture.

1.3 Reward and Attention

It has been well established that consequences in the form of reward or punishment play a crucial role in shaping decision-making and overt behavior. For example, Thorndike's Law of Effect (1911) states that any behavior followed by reward is likely to be repeated, and any behavior followed by less favorable consequences is likely to be stopped. Reinforcement learning is defined as the process of learning to predict consequences and subsequently optimize favorable interactions with the environment (Sutton & Barto, 1998). Indeed, factoring in outcomes when assessing the relevance of a behavior or stimulus is highly adaptive because it favors the production of rewarding behavior at the expense of other competing behaviors and choices.

Just as reward teaches us to consciously select optimal behaviors, so reward also teaches attention to select which stimuli are most relevant and deserve more priority (Chelazzi et al., 2013). The mechanisms by which reward drives visual selective attention have undergone much investigation in recent years. There is evidence through priming that reward may have an immediate effect on the deployment of visual selective attention. Della Libera and Chelazzi (2006) utilized the negative reward-priming paradigm, which entails the presentation of a prime display immediately followed by a probe display. The prime display consists of a visual task involving a relevant stimulus, and a distracting stimulus that interferes with the visual scene and

must be ignored. Negative priming occurs when subjects take a longer time to complete the task in the following probe display when the new target matches the distractor that was ignored in the previous priming task, indicating a lingering effect of an inhibitory attentional mechanism.

Critically, the authors found that reward value modulates the negative priming paradigm, with the priming effect being robust when the prime task was highly rewarded and absent when lowly rewarded. A similar reward priming phenomenon was also observed for objects, as opposed to specific shape stimuli: highly rewarded visual search for an object category either facilitated or impaired search when the category was present as a target or distractor, respectively, in the subsequent search trial (Hickey et al., 2015). Based on these studies, reward plays an important role in rapidly and involuntarily influencing the deployment of visual selective attention.

The effects of reward on attention are not limited to immediate reward priming tasks. Della Libera and Chelazzi (2009) expanded their earlier findings on priming to a more general phenomenon about how reward can persistently influence attention, demonstrating the same attentional capture paradigm across several days. An initial visual search task required participants to search for targets, some accompanied by high reward and other by low reward, and was carried out over three training sessions on consecutive days. Five days later, in a test phase, participants performed the same task in the absence of reward, yet still showed lingering effects of the target's previous reward contingencies through reaction time: a previously high-reward shape became easier to select for and more difficult to reject, and the opposite was true for previously low-reward targets.

The phenomenon by which previously rewarded stimuli involuntarily capture attention, independent of motivation, has been termed *value-driven attentional capture* (Anderson et al., 2011). In Anderson et al.'s experiment, high reward and low reward were associated with a basic

stimulus feature – red or green – during an initial training phase. In the test phase, highly rewarded stimuli again either facilitated or interfered with visual search efficiency as previously described. Most importantly, the reward-associated stimuli were neither task-relevant nor salient in the test phase trials, suggesting that associated reward value inherently increases the salience and attentional priority of a stimulus. Additional studies on value-driven attentional capture by the same authors have revealed that the effect is modulated by the magnitude of the reward (Anderson et al., 2011) and that the effect may persist for as long as seven to nine months following reward learning without any additional training (Anderson & Yantis, 2013). One limitation to these experiments, which rely on training reward associations all at once and then following up with a block of attention probes at varying points in time, is that it is difficult to parse out the competing influences of memory as opposed to reward learning. The question then remains: to what extent is the modulation of selective attention by reward driven by memory of the reward contingency, or by the very process of reward learning and prediction error? Are these two different types of attentional learning by reward, or do they overlap in many ways? In this thesis, I present surprising findings to argue that the process of reward learning and prediction error, not memory and rule finding, is crucial for driving value-driven attentional capture.

Overall, these studies converge on the conclusion that value-driven attentional capture is a robust, persistent, and involuntary aspect of selective attention. Such a mechanism has great implications for clinical conditions in which attention tends to be uncontrollably and abnormally biased toward certain critical stimuli. In the context of reward, this is most relevant for addiction, as a number of studies have demonstrated that addicted individuals have a strong tendency to attend to addiction-related objects in their environment, possibly because of an exaggerated

association of that object with reward that may increase its salience (Field & Cox, 2008).

However, similar attentional biases have also been found towards pertinent stimuli in phobia (Mogg & Bradley, 2006) and eating disorders (Dobson & Dozois, 2004). Certainly, gaining a clearer understanding of the mechanisms behind which attention is flexibly deployed based on reward can illuminate not only learning and decision-making processes in everyday life, but also abnormal conditions in which similar attentional biases are implicated.

1.4 Mechanisms of Value-Driven Attention

Thus far, the evidence on reward and visual selective attention has come from a variety of approaches and research designs. Some studies have demonstrated an immediate effect of reward through priming in the subsequent trial (e.g. Hickey et al., 2015), others have demonstrated a lingering effect up to days following a long reward-association training session (e.g. Anderson et al., 2011) and still more have probed at the process of attentional learning itself by inserting attention probe trials amongst the reward training trials (e.g. Zhao et al., 2013). How might the underlying mechanisms be different, or are they ultimately the same?

One limitation was that in most previously discussed studies, attentional capture typically transferred from a top-down goal-driven search task in the training phase to a bottom-up “pop out” search task in the test phase. This left the possibility remaining that reward modulation is specific to a particular mechanism of attention. Lee and Shomstein (2014) addressed this when they found that the effects of reward on attention also transfer in the reverse direction from bottom-up to top-down attention, suggesting that reward value prioritizes stimuli in a manner independent of specific attentional mechanisms. One possible explanation is that reward value enhances target salience in an integrated priority map that guides both bottom-up and top-down attention; a priority map is a topographic map of sensory input that combines the representation

of salience and relevance when stimuli compete for attentional selection (Fecteau & Munoz, 2006).

These findings are in line with fMRI evidence suggesting that value influences visual salience through early sensation and perception processing areas of the brain. For example, stimulus-reward associations directly modulate activation levels in the early visual cortices responsible for representing low-level stimulus features, such as V1 (Serences, 2008). Furthermore, rewarded stimulus features sharpen the population response of neurons that specifically encode the feature (Serences & Saproo, 2010); this specialized population encoding of a feature is most well-known with orientation (Hubel & Wiesel, 1962). Indeed, Hickey et al. (2010) proposed the anterior cingulate cortex (ACC) as the connection between reward and visual salience when they found that responses to reward feedback in the ACC could predict visual biases by reward. Anderson et al. (2014) supplemented these findings by proposing that both the caudate nucleus of the basal ganglia and the extrastriate visual cortex provide an attentional priority signal sensitive to reward history and irrelevant of endogenous attention. While the caudate nucleus is known to represent objects based on location and reward history (Yamamoto et al., 2013) and the extrastriate cortex is involved in upper stages of vision, the authors propose connections between the two regions (Seger, 2013) that could contribute to attentional capture by reward.

While fully understanding the neural mechanisms of value-based attentional priority still requires further investigation, an overarching conclusion based on the fMRI studies is that reward enhances the representation of the rewarded feature in relevant visual areas. Areas of the brain responsible for processing reward might modulate these early sensory processing regions.

1.5 Generalizability into a Complex World

Despite recent findings that connect value-driven attentional capture to generalized neural mechanisms of sensation, perception and reward, behavioral evidence on the generalizability of the phenomenon is limited. Value-driven attentional capture is most commonly studied when bound to a particular feature; for example, in both their training and test phases, Anderson et al. (2011) use a red circle or a green circle. Likewise, Della Libera and Chelazzi (2009) tested attention using identical shapes, and Laurent et al. (2014) find that a specific orientation can also capture orientation when the same stimulus was previously associated with reward.

Repeating the same stimulus for both the reward-association training phase and the attention-probing training phase puts vast constraints on the external validity of these previous findings. In the real world, settings are significantly more complex than in the laboratory. The same stimulus is unlikely to always look the same due to environmental factors like lighting and perspective. How, then, can mechanisms of value-driven attentional capture be extended to everyday experiences and perceptions of stimuli such as houses, cars, and faces? An early investigation of this question revealed that stimulus-reward associations of a particular feature generalized to different stimuli that shared the same feature: for example, a red circle and a red letter (Anderson et al., 2012). Additional studies have expanded the capture of attention by reward to object categories rather than specific features, providing evidence through reward priming (Hickey et al., 2015) and MEG (Donohue et al., 2016).

Bucker and Theeuwes (2017) provide strong evidence for the applicability and presence of value-driven attentional capture in real life when they demonstrate that Pavlovian reward learning, as compared to the instrumental association of reward with a task as commonly seen in the literature, underlies the phenomenon of value-driven attentional capture. Given that

Pavlovian learning constantly and automatically occurs in every day life (Bouton, 2007), this also provides evidence for value-driven attentional capture as a mechanism that underlies the role of selective attention and reward in everyday experience.

These findings suggest a versatile role for reward learning in modulating attention, but questions still remain unanswered. For example, how does value-driven attentional capture by a feature account for learning in a multidimensional world?

1.6 From Feature to Dimension

Dimensional attention is one way in which value-driven attentional capture can generalize into a multidimensional world, but this has not been extensively explored. To clarify, *dimensional attention* is a broader attention towards an entire category of features, such as color; *feature-based attention* is attention towards a specific element within a dimension, such as red. The effect of reward on dimensional attentional capture represents promising area of research for a number of reasons.

First, the phenomenon of dimensional attention has been demonstrated multiple times in the literature, but just not in the context of reward. For example, Zhao et al. (2013) , as previously described, found that by imbuing regularity in one dimension – for example, having red always follow blue in a stream of colored shapes – they were able to create an attentional bias towards the dimension with regularities (color overall). Harris, Becker, Remington, and Harris (2015) showed that, even when participants were cued to search for only one feature, any singleton feature within that whole dimension subsequently captured their attention. Muller, Reimann, and Krummenacher (2003) used a similar attentional paradigm to likewise demonstrate that observers cued to a specific feature showed attentional capture by any feature within that dimension during a visual search array.

Furthermore, dimensional attention has crucial implications for understanding the mechanisms behind reward learning and how complex details of the environment are represented in the brain. At any given time, only a few dimensions in a complex environment – such as color, shape, or texture – are relevant for obtaining reward. For example, our selective attention toward the multidimensionality of cars must be versatile: in searching for a taxi, only color is relevant whereas attending to speed and distance is imperative for crossing the road. Reinforcement learning algorithms have notoriously suffered from what Bellman (1957) has referred to as the “Curse of Dimensionality,” becoming less efficient as the dimensionality of stimuli in the environment increases. Directing attentional priority toward some dimensions and not others is a promising solution to this problem. Consistent with this, Niv et al. (2015) demonstrated that reinforcement learning does rely on attentional mechanisms that reduce the dimensionality of a complex environment into only those dimensions that are behaviorally relevant.

Finally, intuitively, we know that dimensional attention and learning do happen implicitly. For example, if we are learning that red is rewarding, we are simultaneously learning that a color is rewarding.

In light of this discussion, this thesis seeks to fill the gap in the literature by bridging dimensional attention with the phenomenon of value-driven attentional capture. Attentional capture by a feature alone is not enough to generalize to a multidimensional world. I ask the questions: can reward learning spontaneously influence dimensional attention? To what extent does learning about a highly rewarding feature lead to attentional capture by not only the feature, but also the whole dimension of that feature? I hypothesized that learning about high reward features in one dimension will bias subjects’ attention towards other low reward features within the same dimension. Addressing these questions will reveal the extent to which the spontaneous

deployment of attention is flexible and can exist at multiple levels of abstraction, from features to dimensions to objects. Moreover, value-driven attentional capture at the dimensional level will provide a broad and lower level mechanism by which the reduction of dimensionality in reinforcement learning can occur.

In order to investigate whether feature-based reward learning biases dimensional attention, I conducted two behavioral experiments on human subjects, utilizing a Pavlovian reward learning paradigm and a visual search array in order to probe for dimensional attention. The methods, findings, and discussions of these two experiments are presented in the following chapters.

Chapter 2

General Methods

Some material in this chapter was adapted from Bu (2016).

In this chapter, I provide an overview of the general methods for the two experiments reported in this thesis, focusing on the experimental design and statistical analyses common to both experiments.

2.1 The Token Task

The Token Task was first developed by Angela Radulescu and based on the experimental design by Zhao et al. (2013). For the experiments in this thesis, I used a set of stimuli, called “tokens,” which varied along the dimensions of color and orientation (fig. 2.1). The task consisted of three types of trials: collect trials, visual search array trials, and predict trials (fig. 2.2). The purpose of the collect trials was to present a visual stream of tokens and their associated reward, imbuing certain visual features with high value and others with low value over time. Array trials were intermittent probes for the effect of reward on selective attention and attentional capture, critical to the hypothesis of the experiment. Predict trials assessed how successfully participants learned the reward contingencies after each block of collect and array trials. Together, these trials enabled me to investigate how reward learning influences the way we direct our attention in the future.

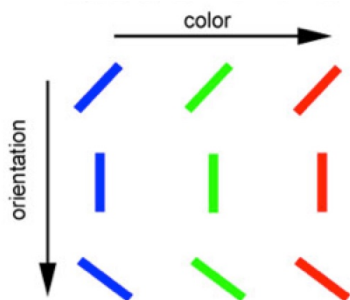


Figure 2.1. An example set of 9 “tokens” varying along two dimensions, with 3 features within each dimension: color (red, green, blue) and orientation (45°, 90°, 135°).

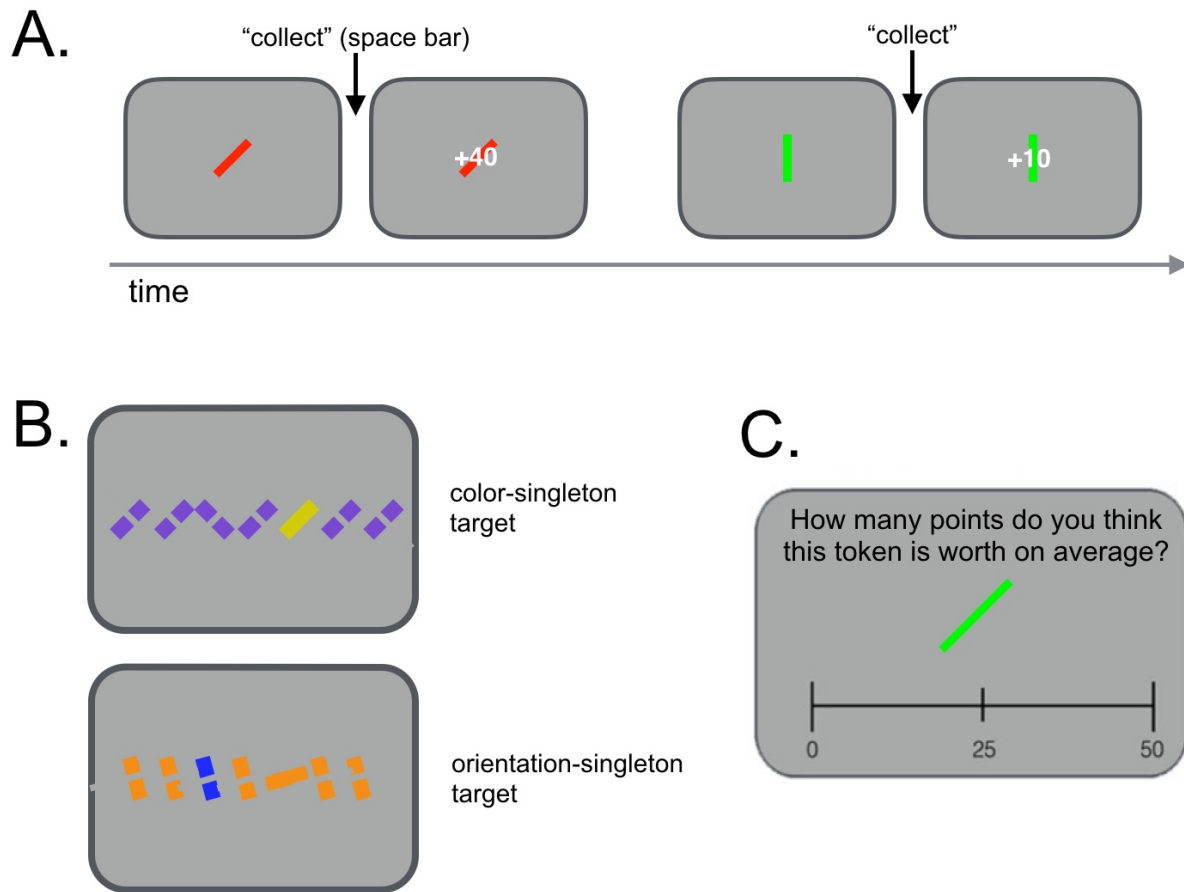


Figure 2.2. Trial types throughout the experiment. **A)** *Collect trials*. Subjects pressed space bar to view the point value of each token. **B)** *Visual search array trials* (adapted from Experiment 1 in this figure) interrupted collect trials to probe attention. Each contained one color singleton and one orientation singleton, and dimensional attention was measured based on which singleton was the target. **C)** *Prediction trials*. Subjects dragged the center token to a value on the bar, and clicked to select their prediction of the token’s point worth.

2.1.1 Collect Trials and Reward Structure

Collect trials (fig. 2.2A) comprised the majority of the experimental task. For each collect trial, subjects viewed a token in the center of the screen and “collected” each token by pressing the space bar on the keyboard. Upon collecting the token, a point value of either “+10” (low reward) or “+40” (high reward) was overlaid in white above the token; subjects were told that the more points they collected from these tokens, the larger their cash bonus would be. For each

block of trials, tokens with one specific feature (e.g., red tokens) provided high reward most of the time, while tokens without this feature provided low reward most of the time. For the rest of this thesis, the feature predicting high reward in every block will be referred to as the *relevant feature*. The relevant feature changed after every block in the task, and participants were informed between each block that the relevant feature would change. Participants were randomly assigned to either a “color” group or an “orientation” group, and all rewarded features were drawn from the corresponding dimension (e.g. for subjects in the color group, the highly rewarded feature in each block would always be a color such as red or blue). This dimension will be referred to as the *relevant dimension*. Subjects were not informed of the relevant dimension or that only one feature at a time predicted high reward, and they had to learn these rules over time.

The design of the collect trial in this task was strictly conducive to Pavlovian reward learning. In addition to what has been previously discussed in the introduction of this thesis, the use of Pavlovian reward learning was a noteworthy design choice because it allowed me to directly test whether simply imbuing certain stimuli with value would lead to an attentional bias toward the dimension, independent of any existing task set or instrumental behavior. It also allowed me to ensure that the reward learning experience of every subject was comparable.

Several aspects of this experimental design were modified from the original Token Task. Visual reward overlaid above the stimulus was chosen due to evidence that suggested that the simultaneous presentation of a stimulus and its associated reward was an important factor in guiding visual representation in the brain (Roelfsema et al., 2010). A point reward system was chosen because it allows for the flexible manipulation of reward, such that high reward tokens (+40) were of a considerably greater magnitude than low reward tokens (+10); attentional capture by rewarded stimuli is modulated by the relative size of the reward (Anderson et al.,

2011). The only possible reward values for each individual token were “+10” and “+40” in order to create a clear distinction between low and high reward. Finally, previous experiments involving the Token Task used an 80%/20% probabilistic feedback system, meaning that tokens with the relevant feature provided high reward with 80% probability and low reward with 20% probability. Tokens without the relevant feature provided high reward with 20% probability and low reward with 80% probability. However, as a result of this, subjects often struggled to learn the relevant rules guiding reward and its structure. For the experiments in this thesis, a 90%/10% probabilistic feedback system was chosen to facilitate learning while simultaneously maintaining a small prediction error to prolong the learning process.

2.1.2 Visual Search Trials as a Probe of Dimensional Attention

Visual search array trials would occasionally interrupt the stream of collect trials (fig. 2.2B). The array always contained two singletons amidst homogenous distractors; one singleton would pop out solely due to its orientation, while the other would pop out solely due to its color. Furthermore, one singleton would always appear on the right side of the screen, and the other would always appear on the left side of the screen. Both would always be equidistant from the center. All tokens in the array except for one of the singletons contained a small gap in the center. The subject’s task was to press either “X” or “M” as quickly and accurately as possible whenever a search array appeared, to report whether the gapless token was on the left side of the screen or on the right side of the screen. “X” was always for “left”, while “M” was always for “right”; this was due to their relative positions on the keyboard, which made the keyboard mappings and the search task quick and intuitive.

As the gapless token was always one of the singletons, the target for these search trials alternated between the orientation singleton and the color singleton (a example of each type of

target is shown in fig. 2.2B). There was an equal amount of color singleton target and orientation singleton target trials throughout the experiment. The timing of array trials was designed such that they were unpredictable but would appear more frequently near the end of each stream of collect trials (around once every 4 trials, as opposed to around once every 5 trials in the beginning), in order to better probe for an attentional bias after subjects learned which features predict high reward.

The use of two competing singletons in a search array to probe for dimensional attention is novel and was motivated by a number of findings in the selective attention and attentional capture literature. On the most intuitive level, a competing singleton design depends on the biased-competition theory of attention. Objects are constantly competing for cortical representation and cognitive processing, and this theory suggests that attention's role is to bias this competition in favor of more potentially relevant objects at the expense of others (Duncan, 1996). When faced with two highly salient singletons in the visual field that differ only in their pop out dimension, an acquired dimensional attention bias could influence which singleton will compete better and capture attention first, manifesting in subject reaction times. This conclusion is further supported by evidence that attentional capture is strongly weighted towards a single object rather than triggered equally by multiple objects in visual scene (Grubert & Eimer, 2017) and that top-down processes can mediate attentional capture by a salient singleton (Kiss, Jolicoeur, Dell, & Eimer, 2009).

2.1.3 Predict Trials as a Measure of Learning

At the end of each block of collect and array trials, participants were presented with prediction trials (fig. 2.2C). All nine tokens seen in the collect trials were presented one by one, and subjects were asked to guess the average point value of each token in the set from out of a

range from 0 to 50. This continuous scale was designed to assess both how well the subjects had learned the relevant rule for the block and how confident subjects were in what they had learned (less confident subjects would be more likely to guess around 20 to 30, as opposed to values at either end of the scale). Because of the 90%/10% reward contingency, the correct average worth was set at 13 for tokens without the relevant feature of that particular block (low value) and 37 for tokens with the relevant feature (high value). To motivate learning, answers were rewarded bonus points up to 100 per prediction based on how close they were to the real worth of the token. Participants were given aggregate feedback out of total points possible at the end.

2.2 Experimental Overview

2.2.1 Hypothesis

In this thesis, I investigate whether reward can influence attentional capture at the level of entire feature dimensions. Specifically, I test the hypothesis that learning to attend to highly rewarding features in a single dimension results in attentional capture for other unrewarded features within that same dimension, as compared to features in other dimensions.

The visual search trials are a critical test of this hypothesis. A learned dimensional attention bias should both facilitate pop out and search for a singleton target in the dimension containing the highly rewarded feature. If learning about rewarding features in one dimension creates an attentional bias towards that dimension, then I would expect subjects in the color group to be faster at finding color singleton targets than at finding orientation singleton targets, and I would expect subjects in the orientation group to find orientation targets faster than color targets. Statistically, I would expect an interaction between target dimension and group in a 2x2 analysis of variance (ANOVA, fig. 2.3).

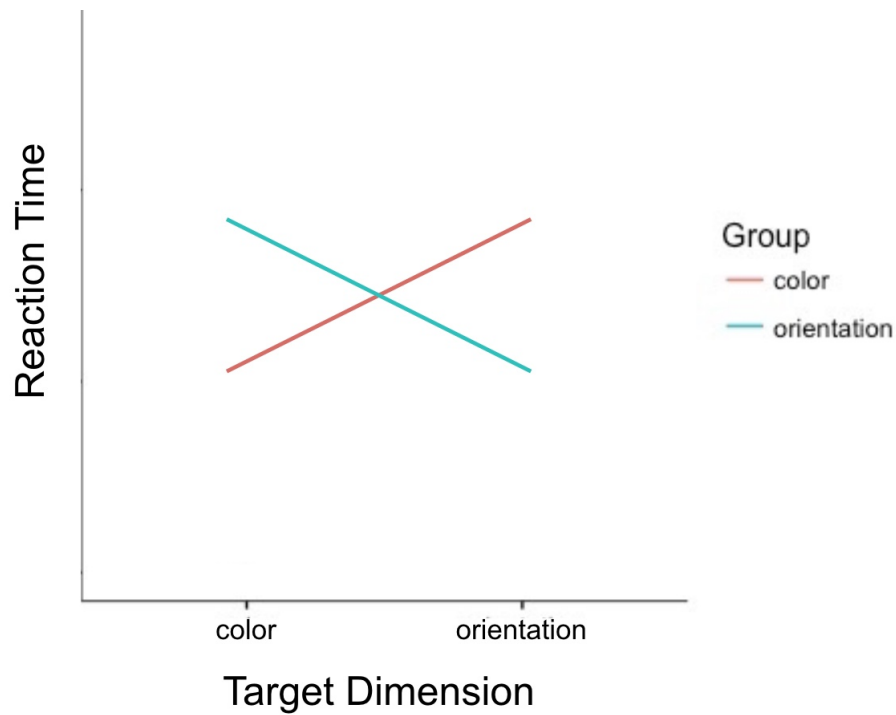


Figure 2.3. Hypothesized results for reaction times of visual search trials, predicting a 2-way interaction between group (rewarded dimension throughout the experiment) and target dimension.

2.2.2 The Present Experiments

In the following two chapters, I present two variants of the Token Task. I begin by describing in full the methods, results, and discussion of the first experimental run (chapter 3). Then, I discuss how the findings from this study were used to inform and modify the Token Task for a replication with variation (chapter 4). A general discussion (chapter 5) synthesizes and reconciles the findings from these two experiments, identifies potential limitations, and proposes pathways for future investigation.

Chapter 3

Experiment 1: The Token Task

The goal of this experiment was to investigate if and how feature-based reward learning in the Token Task drives selective attention and attentional capture toward whole feature dimensions. Though previous variants of the Token Task have been run in the past as described in Bu (2016), this is the first experiment to use a competing singleton search array design to probe for dimensional attention.

3.1 Methods

3.1.1 Participants

Thirty-five members of the Princeton University community (28 undergraduates; 26 females and 9 males; age range 19-30 years; mean age of 21.7 years) participated in the experiment and were paid a flat rate of \$8 for half an hour plus a bonus of \$2 based on performance throughout the experiment. Due to technical failure, three subjects were excluded from analysis. This left a total of 32 participants (24 females and 8 males; age range 19-30 years; mean age of 21.9 years). All subjects reported normal or corrected-to normal visual acuity and color vision. Study materials and procedures received approval from the Princeton University Institutional Review Board (IRB), and subjects provided informed consent.

3.1.2 Stimuli

Visual stimuli were presented with a grey background using MATLAB 2014A and the Psychophysics Toolbox on a Dell XPS monitor. The stimuli used throughout the experiment consisted primarily of two sets of nine tokens each (fig. 3.1).

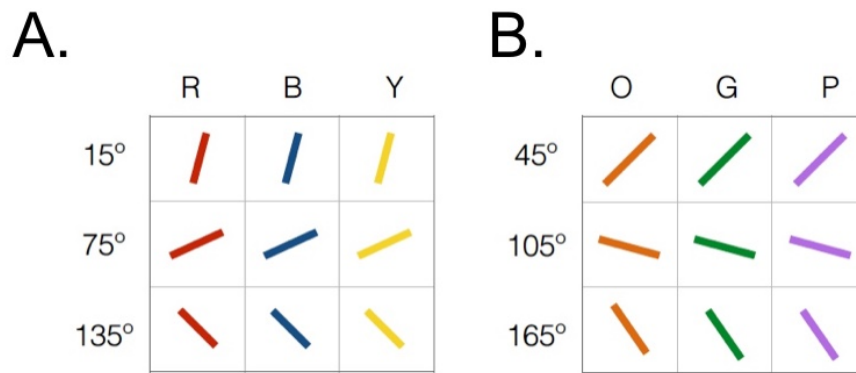


Figure 3.1. Two sets of tokens used in the experiment. Tokens varied along two dimensions, and there were 6 possible features for each dimension. **A)** The 9 tokens used for the collect trials. These tokens and their features would be associated with reward throughout the experiment. **B)** The 9 tokens used as non-singleton distractors in the visual search trials. To avoid potential confounds, these distractor tokens had features completely distinct from those used in the collect trials.

Tokens in the first set were used exclusively for collect trials (fig. 3.1A). These nine tokens and six features, consisting of three from each dimension (color: red, blue, and yellow; orientation: 15°, 75°, and 135°), were imbued with varying levels of reward throughout the experiment. Tokens in the second set (fig. 3.1B) were used only as distractors in the visual search trials and were characterized by six features different from those seen in the collect trials (color: orange, green, and purple; orientation: 45°, 105°, and 165°). This was because in previous runs of the Token Task, I found one factor that could potentially influence the reaction times in visual search was the presence of rewarded features in the distractors. For example, if subjects learn that red is the relevant feature, then a search array consisting of a green target with red distractors could lead to potentially unpredictable effects such as enhanced distractor filtering (Lee & Shomstein, 2014) that could interfere with or confound the effects of a dimensional attention bias. As a result, I wanted to ensure that distractors in the visual search array did not contain any features that could be associated with reward. The specific colors and orientation angles in each set were chosen to create the maximum possible differentiation between stimuli.

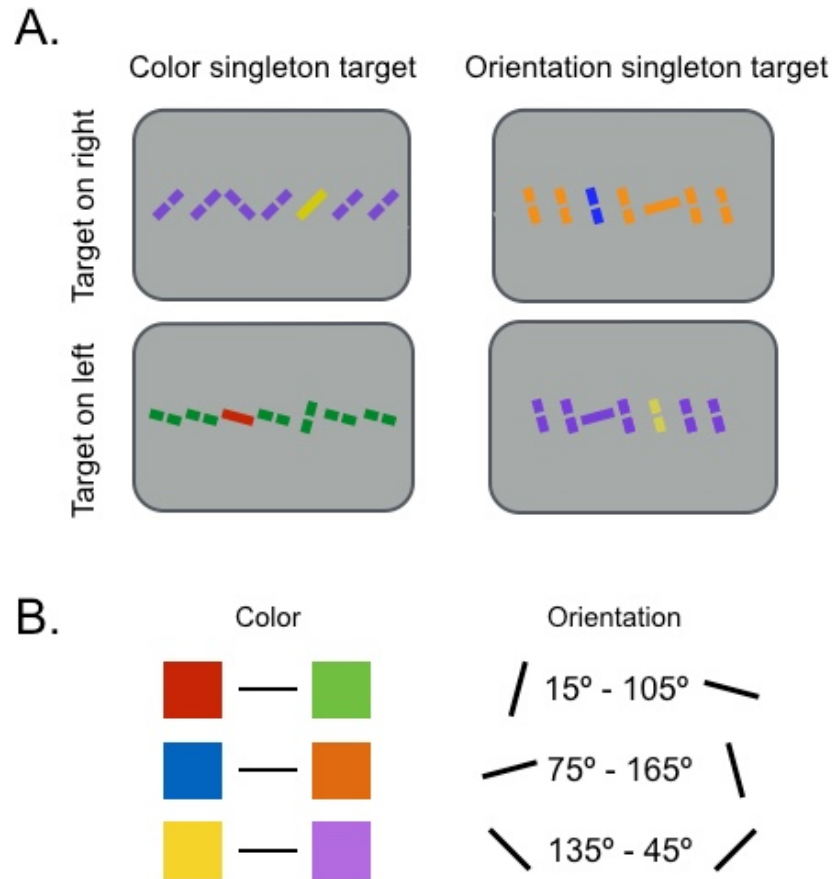


Figure 3.2. *A)* Four examples of the visual search arrays used in this experiment. Note how the singleton features are a hybrid of the collect and distractor tokens in fig. 3.1.

B) The pairings used to determine the features of singletons and distractors in the visual search trials, in order to both normalize and maximize singleton pop outs.

Stimuli in the visual search trials were designed with careful attention to the features of the competing singletons and distractors (fig. 3.2A). Singletons were always a hybrid of features from the collect and distractor trials; the feature in the pop out dimension of the singleton was always a reward-associated one from the collect trials, while the feature in the other dimension was the same as the distractors. In other words, the color singleton would always be red, blue, or yellow, and the orientation singleton would always be 15°, 75°, or 135°. I chose this design because it would create an array in which all features were homogenous except for the orientation feature in the orientation singleton and the color feature in the color singleton. Half of

the time, the singleton in the rewarded dimension (e.g. the orientation singleton for the orientation group) would be in the relevant feature for that particular block of trials; the other half of the time, it would be in the two other low-reward features. This was in order to allow me to later analyze whether attentional capture by the relevant feature differs significantly from attentional capture by non-relevant features in the rewarded dimension.

In the visual search array, every color singleton feature was paired with a distractor color maximally distant from it on the color wheel, and every orientation singleton feature was paired with a distractor orientation 90° away (fig.3.2B, can be observed in 3.2A). This was to preclude, for example, the possibility that color pop out in one array trial would be weaker with a yellow singleton among orange distractors, compared to another trial with a yellow singleton among purple distractors.

Finally, the visual search array in this experiment always consisted of a 7x1 array of tokens. The singletons would always appear at the third and fifth position; the dimensions of the singletons and the location of the target alternated between these two positions.

3.1.3 Procedure

Before beginning the task, participants received on-screen instructions informing them how to complete each trial, including three examples of collect trials and one example of a visual search array trial. To facilitate learning of the relevant rules for reward, participants were explicitly told that the orientation and color of the tokens could help determine their worth: “It will be useful to pay attention to the point outcome of each token, because every now and then, we will test you on how well you can predict the values of certain tokens... The key to predicting successfully is to pay attention to the visual features of each token, as they might indicate how valuable it is.” Furthermore, to alert subjects to the possibility of probabilistic feedback, they

were told that the same token might give different points on different trials, but that on average some were worth more than others. Following the instructions, participants were given an abridged game as a practice round before beginning the experiment.

Altogether, the experiment consisted of 6 blocks or “games” in total, such that each possible feature in the collect trials within the rewarded dimension served as the relevant feature two times. Each game consisted of a series of 108 collect trials interrupted by 24 array trials and followed by 9 predict trials at the end. The inter-trial interval was always 0.5 seconds with a timeout after 2 seconds for collect trials, 1 second with a timeout after 2 seconds for array trials, and 0.5 seconds with a timeout after 5 seconds for predict trials. A schematic of one game is shown in figure 3.3. After the experiment, subjects were debriefed and given a questionnaire (appendix A) assessing their engagement with the task and their knowledge of the relevant rules to reward. The whole procedure took 30 minutes.

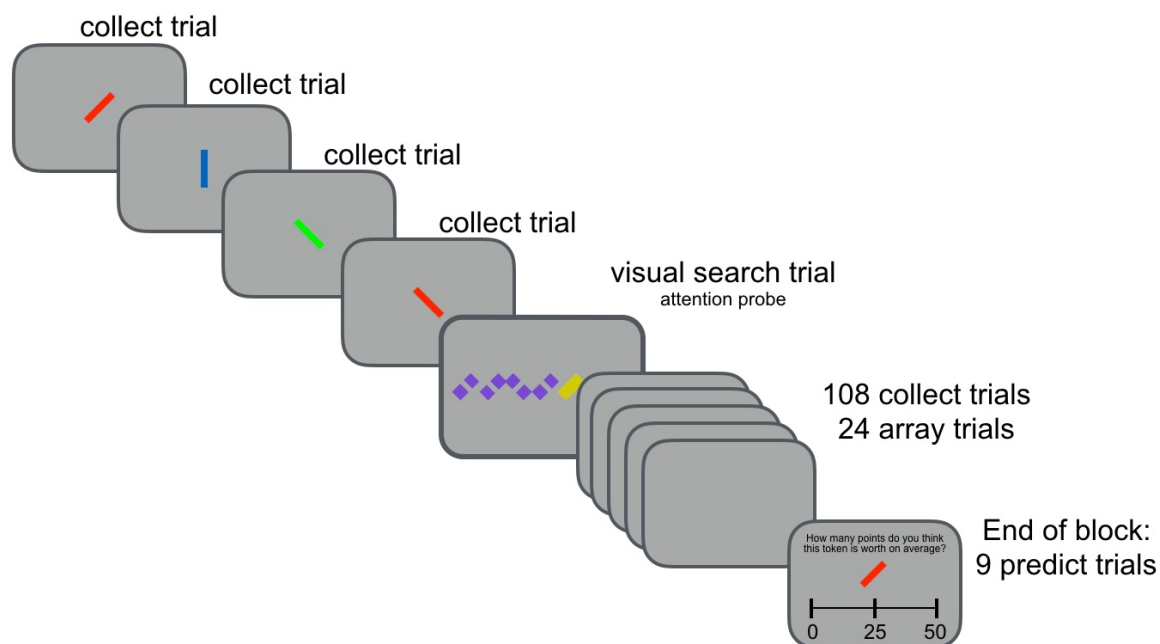


Figure 3.3. Schematic of one block, or “game”, of the experiment out of nine games total.

3.1.4 Statistical Analysis

All data were recorded and processed using MATLAB. Statistical analyses and tests on the data were carried out using R, with $\alpha = 0.05$. When stated, within-subject error bars were calculated using the procedure described in Cousineau (2005).

3.2 Results

3.2.1 Trial Accuracy

Due to the dependence of this experiment on a sensitive measure of reaction time, a preliminary descriptive analysis of subject performance on all trials was conducted to identify and exclude subjects with unusually poor performance (table 3.1). Collect trial accuracy was quantified by the percentage of tokens that a subject collected throughout the experiment; subjects were instructed at the beginning of the experiment to collect all the tokens. Visual search trial accuracy was quantified as the percentage of search targets that subjects correctly identified as being on the left or right side of the screen. Subjects were told to answer as quickly and as accurately as possible. Finally, predict trial accuracy was quantified as the percentage of trials that a subject guessed within six points of the correct value. Participants meeting any of the exclusion criteria, making their performance for that particular type of trial a strong statistical outlier, were excluded. Based on this method, three subjects were excluded, leaving 29 subjects total remaining for further analyses (15 in the color group and 14 in the orientation group).

Trial Type	Accuracy Measure	Subject mean (SD)	Exclusion Criteria
Collect trial	% tokens collected	99.84% (0.55%)	< 98.74%
Visual search trial	% correct response	94.74% (11.93%)	< 70.75%
Predict Trial	% correct answers	84.84% (12.83%)	< 59.18%

Table 3.1: overall subject accuracy on collect, visual search, and predict trials. Exclusion criteria were determined using a statistical definition of outlier: the sample mean plus or minus two times the standard deviation.

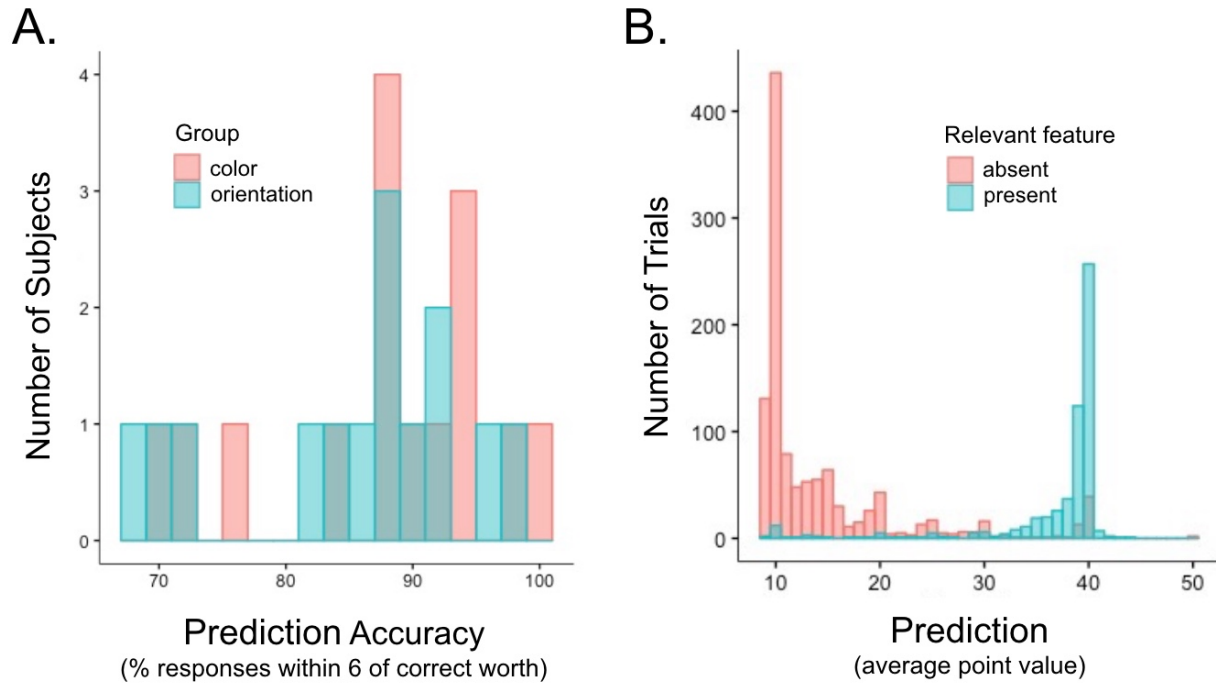


Figure 3.4: Assessment of subject learning performance based on prediction trials. **A)** Histogram of subjects' overall prediction accuracy, quantified as the % of predictions within 6 of the token's correct worth. The correct value of tokens with the relevant feature present was set at 37 points, due to the probabilistic feedback previously discussed. The correct value of tokens without the relevant feature was 13 points. **B)** Overlapping histograms showing the distribution of all subject predictions for tokens with (high reward) and without (low reward) the relevant feature.

3.2.2 Learning Performance

Overall, participants showed robust evidence of learning throughout the task. Each participant's accuracy was measured as a percentage of how many correct predictions they made within a range of six of the correct token worth out of the total number of prediction trials. The distribution of the prediction accuracies for the 29 subjects, split by their assigned group, is shown in figure 3.4A. The median subject prediction accuracy was 88.9% with a range from 68.5 to 100%, meaning that all subjects (except for one outlier, who was excluded as previously discussed) performed above chance. This is complemented by an aggregated distribution of the specific predictions that subjects made for both high reward tokens, in which the relevant feature was present, and for low-reward tokens, in which the relevant feature was absent (fig. 3.4B).

With the exception of a few misses, the vast majority of predictions were able to distinguish between high reward (blue, on the right) and low reward (red, on the left) tokens. The most common predictions were 10 for low reward tokens and 40 for high reward tokens. It is clear that participants were able to distinguish between relevant and non-relevant tokens and had learned the relevant rule for determining reward, even if they did not take into account the probabilistic 90%/10% feedback. Furthermore, subjects learned the relevant rule for determining high reward quickly (fig. 3.5). Even by the end of the first game, most subjects were able to predict with accuracy what the point value of each token was.

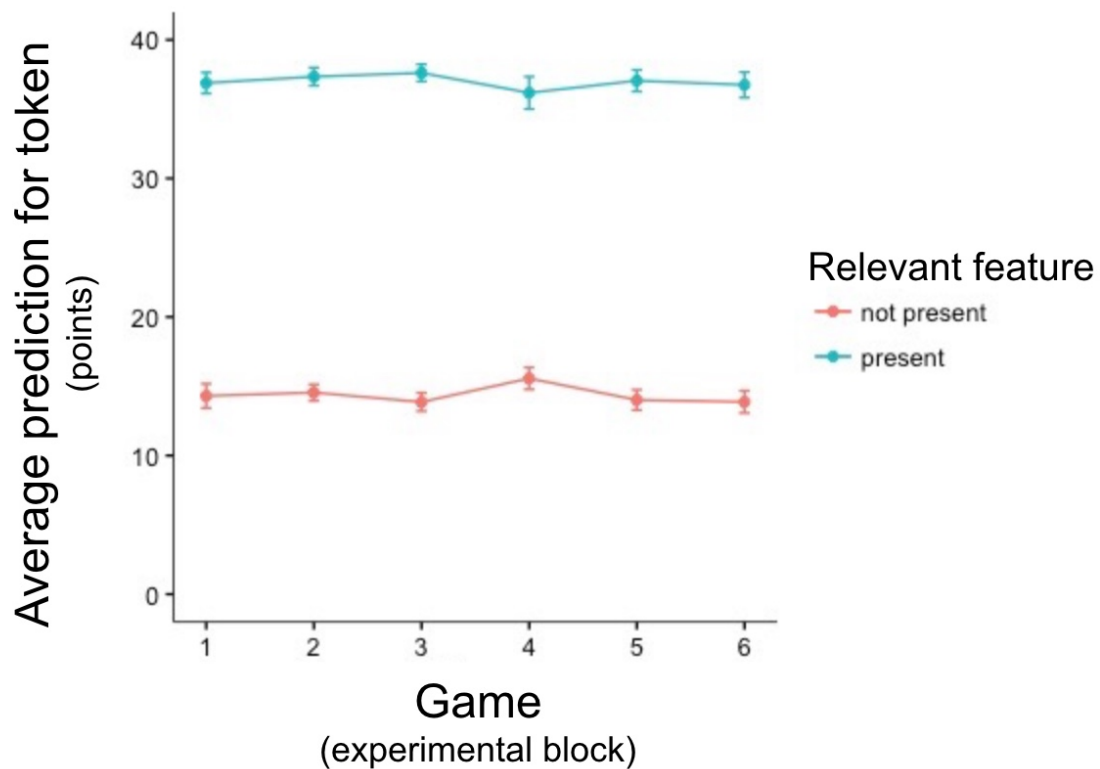


Figure 3.5: A learning curve, quantified by average prediction for tokens with and without the relevant feature in each game, aggregated across subjects. 1 is the first game or block of the experiment, and 6 is the final block of the experiment. The correct prediction for tokens with the relevant feature not present was 13, and 37 for tokens with the relevant feature present.

Finally, these results are consistent with written feedback acquired from subjects in a questionnaire following the study. In these debriefing questionnaires, subjects were prompted with the question, “Did you notice that a particular characteristic was more likely to indicate whether tokens were worth more points? If so, which?” All but three subjects out of twenty-nine total were able to correctly identify the dimension that was more relevant to predicting reward. Moreover, when asked to rate the difficulty of the experiment on a scale of 1 (very easy) to 5 (very difficult), the average response was 2.6 ($SD = 0.91$), leaning on the side of easy.

3.2.3 Reward and Attention

The critical test of my hypothesis depended on subject response times to the visual search trials. If feature-based reward learning generalizes to attentional capture at the dimensional level, then we would expect the group that a subject was assigned to (color or orientation) to differentially affect their response times to color and orientation singleton targets. Specifically, subjects should be faster to respond to singleton targets in the rewarded dimension compared to singleton targets in the non-rewarded dimension. To ensure that reaction times represented a search for the correct singleton target, all visual search trials in which subjects responded incorrectly (e.g., left when the gapless token was on the right side of the screen) were excluded.

Visual search trial reaction times (fig. 3.6) were analyzed with a 2 (group: color, orientation; between subjects) $\times$ 2 (target dimension: color, orientation; within subjects) mixed-effects analysis of variance (ANOVA). Consistent with our hypothesis, there was a trending interaction between target dimension and group, $F(1, 27) = 2.876, p = 0.10, \eta^2 = 0.096$. A post-hoc paired t-test revealed that subjects in the color group were significantly faster to find color singleton targets as opposed to orientation targets, $t(14) = 2.311, p < 0.05, d = 0.22$. There was also a main effect of target dimension, $F(1, 27) = 5.200, p < 0.05, \eta^2 = 0.161$.

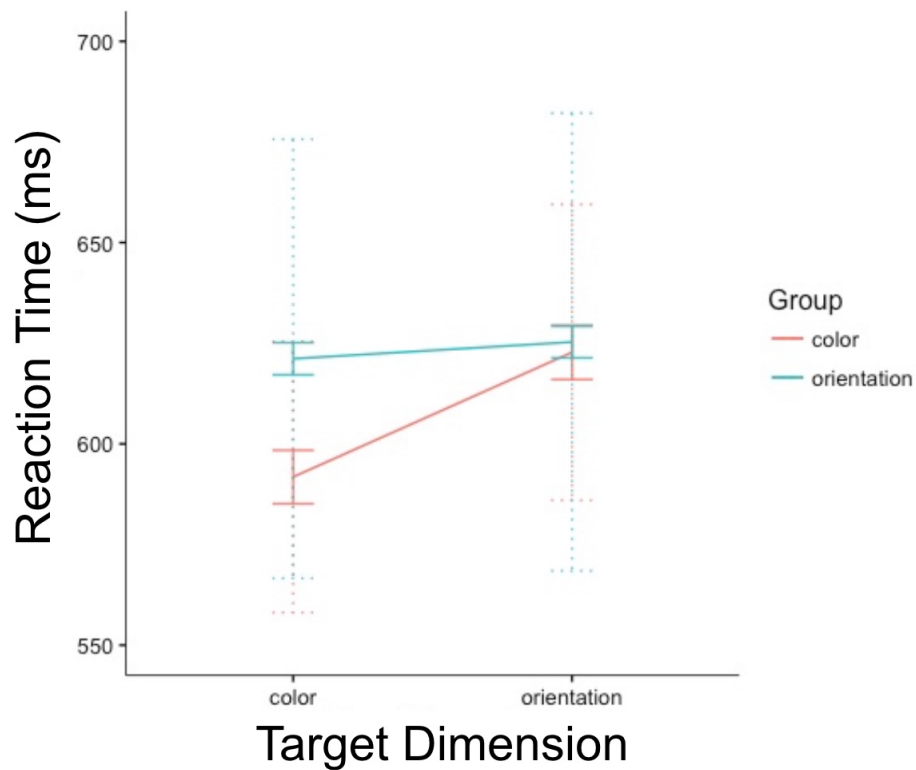


Figure 3.6: An interaction plot showing reaction times to visual search trials, broken down by group and target dimension. Since this was a mixed effect design with both a within subjects variable (target dimension) and a between subjects variable (group), two types of error bars were needed for each point. The solid lines are within-subject standard error (Cousineau, 2005), and the dotted lines are between subject standard error. N=15 color group subjects, N=14 orientation group participants.

The presence of the trending two-way interaction between group and target dimension, particularly as it manifests in the color group compared to the orientation group, provides evidence of an effect of attentional capture at the dimensional level for color, but not for orientation. In the following three sections, I sought to better understand the factors that may have contributed to this effect. I present three additional analyses of the reaction time data based on the time course of the experiment, the presence of the relevant feature, and learning performance as measured by prediction accuracy.

3.2.4 Exploratory Analyses

Analysis of Attentional Effects by Experimental Time Course

I hypothesized that feature-based reward learning leads to attentional capture at the level of entire feature dimensions. However, reward learning does not usually happen immediately; prediction errors will dynamically update learned value over time (Sutton & Barto, 1998). For this experiment, a number of collect and search trials might be expected to elapse before subjects learn the reward structure and rules that predict reward, which then might lead to a learned attentional bias. This delay might manifest in the reaction time data in two ways.

First, subjects are more likely to have successfully learned that features from one dimension consistently predict high reward by the latter half of the experiment. Thus, I predicted for this analysis that the dimensional attention bias effect might be stronger in the second half of visual search trials that appeared in the experiment, compared to the first half of visual search trials. However, a 2 (time course: first half of search trials, second half of search trials; within subjects) $\times$ 2 (group: color, orientation; between subjects) $\times$ 2 (target dimension: color, orientation; within subjects) ANOVA did not reveal a significant three-way interaction, $F(1, 27) = 0.468, p = 0.50, \eta^2 = 0.017$, whereas the two-way interaction between target dimension and group remained robust, $F(1, 27) = 2.993, p < 0.10, \eta^2 = 0.173$. Moreover, there was a main effect of time overall, $F(1, 27) = 6.450, p < 0.05, \eta^2 = 0.173$. Thus, participants got faster at responding to visual search trials in the latter half of the experiment, possibly suggesting a practice effect (Donovan & Radosevich, 1999), but the time course of the experiment was unable to account for differences in attentional capture.

Second, since the highly rewarded feature changed in every game of the experiment, subjects need to take time to newly learn which feature predicts high reward in every game. It is

possible then, that the hypothesized interaction between group and target dimension becomes robust in the second half of visual search trials *within each game*, as compared to the beginning trials of each game. However, a 2 (time course: first half search trials of each game, second half search trials of each game; within subjects) x 2 (group: color, orientation; between subjects) x 2 (target dimension: color, orientation; within subjects) revealed almost exactly the same trend as the first time course analysis. There was no significant three-way interaction $F(1, 27) = 0.220, p = 0.642, \eta^2 = 0.008$, whereas the two-way interaction between target dimension and group remained trending, $F(1, 27) = 2.921, p < 0.10, \eta^2 = 0.10$. Thus, there was no evidence to conclude that time course of learning, and of the experiment overall, significantly impacted the critical interaction observed between group and target dimension.

Analysis of Attentional Effects by Presence of Relevant Feature

Another factor to account for is the presence of the highly rewarded feature in the visual search trials. The pop out features of the singletons in the visual search trials were always one of the features that subjects saw during the collect trials (e.g. blue, yellow, red, 15°, 75°, 135°). As a result, in every game, some singletons in the visual search trials had the highly rewarded feature. Attentional capture by a highly rewarded feature is a phenomenon that has been well established in the literature, and would not provide any evidence of generalization to a dimensional attention bias (Anderson et al., 2011). Thus, in order to demonstrate an attentional capture effect by the whole dimension, this portion of trials alone cannot drive the interaction effect that we observed. The visual search trials containing singletons with low-reward features (e.g., the relevant feature is absent) must also contribute to the effect, or at least not differ significantly from those trials with the relevant feature present. Part of the reason I chose to include visual search trials with the relevant feature was to be able to test for whether this difference exists.

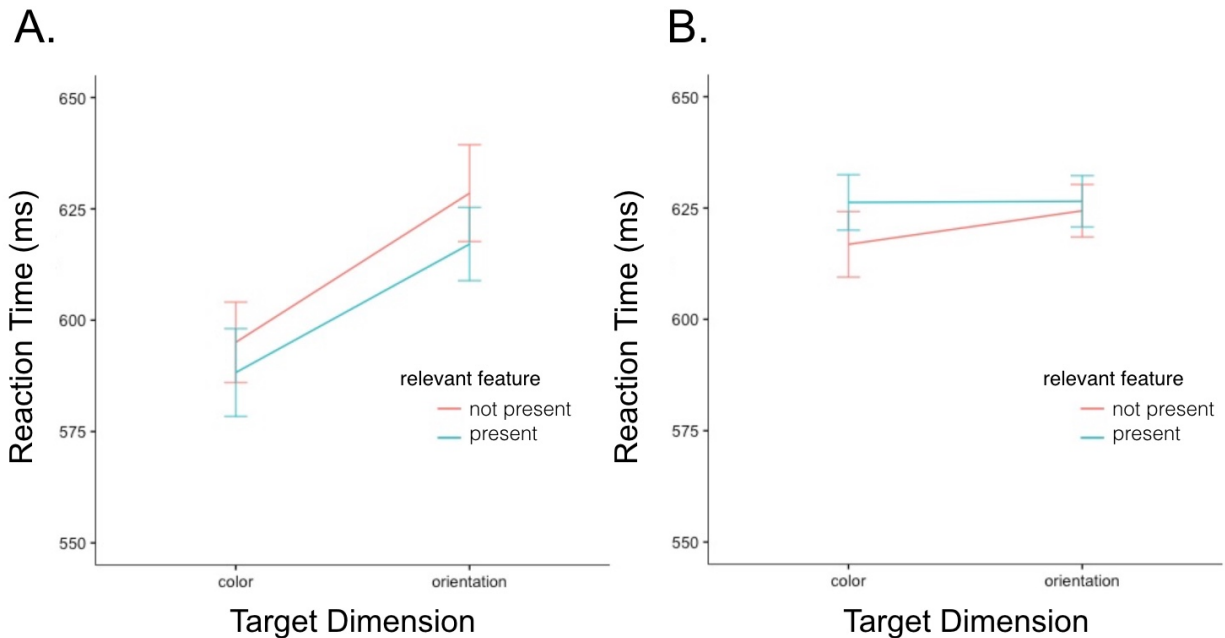


Figure 3.7: Interaction plots showing reaction times to visual search trials, broken down by target dimension and the presence of the relevant (highly rewarded) feature for **A) subjects in the color group** and **B) subjects in the orientation group**. Standard error, corrected again for within-subjects comparison in each graph, was used. N=15 color group subjects (**A**), N=14 orientation group participants (**B**).

The breakdown of the reaction time data by the presence or absence of the highly rewarded feature is shown in figure 3.7. Even with within-subjects error bars, it can be observed that there is no significant difference between visual search trials with and without the highly rewarded feature. In fact, for the orientation group, the difference actually trends in the opposite direction of what we might expect assuming value-driven attentional capture by a feature. Additionally, a 2 (presence of relevant feature: absent, present; within subjects) x 2 (target dimension: color, orientation; within subjects) x 2 (group: color, orientation; between subjects) ANOVA did not reveal a three-way interaction, $F(1, 27) = 0.011$, $p = .918$, $\eta^2 < 0.001$, despite a trending interaction between target dimension and group as previously noted.

Thus, there is no evidence to conclude that the presence of the relevant feature in the visual search trials could account for the attentional bias that we observed, particularly with

regards to the color group and the color target singletons. This is consistent with our hypothesis for an effect of attentional capture at the dimensional level, as opposed to attentional capture at the level of individual features.

Analysis of Attentional Effects by Prediction Accuracy

Finally, I predicted that the better subjects learned the features that predicted high reward throughout the task, the stronger they would show an effect of an attentional bias at the dimensional level. This is an intuitive prediction; a subject who has successfully learned which features predict high reward is more likely to be biased toward those features than a subject who is struggling to learn or may not be aware of what predicts high reward best.

In order to perform a continuous analysis between subjects' learning performance and attentional capture effect by a dimension, the two variables must first be quantified. Learning performance was measured and quantified using subject accuracy on the prediction trials; specifically, prediction accuracy was defined as the percentage of trials in which a subject was able to successfully predict the token's true value within 6 of the correct value (e.g., for low-reward tokens, a prediction from 10-16 points was considered correct; for high-reward tokens, a prediction from 34-40 points). If a subject guessed outside of these ranges, it is likely that they did not learn or were not confident that one specific feature always predicted high reward.

The attentional capture effect was quantified by calculating the average reaction time in each subject for visual search trials with a color singleton target and for visual search trials with an orientation singleton target:

Capture effect

$$= \text{Mean}(RT_{\text{target in unrewarded dimension}}) \\ - \text{Mean}(RT_{\text{target in rewarded dimension}})$$

Based on our dimensional attention hypothesis, reaction times with a target in the unrewarded dimension should be longer than reaction times with a target in the rewarded dimension. Thus, the more positive the capture effect is for subjects in either the color or orientation group, the more attentional bias they showed toward the rewarded dimension. The more negative the capture effect is, the more their attentional bias trended towards the unrewarded dimension, contrary to what we would predict.

Surprisingly, a correlational analysis between learning performance and attentional capture effect showed a significant negative correlation, $r = -0.55$, $p < 0.005$ (fig. 3.8). This negative correlation remained robust even when the potential outlier with a capture effect of around 150 ms was removed, $r = -0.45$, $p = 0.01$ and also when the color and orientation group subjects were correlated separately, $r = -0.77$, $p < 0.001$ for color and $r = -0.54$, $p < 0.05$ for orientation.

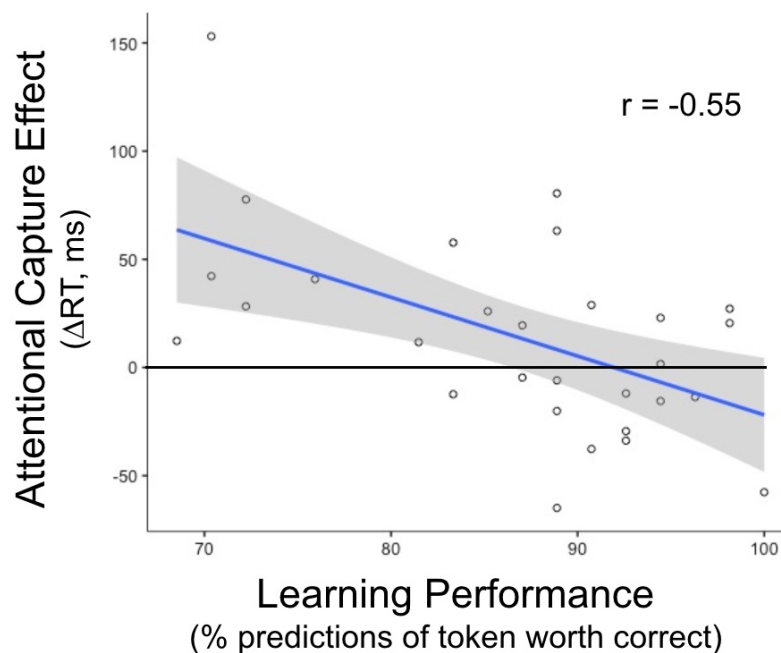


Figure 3.8: A scatterplot summarizing the findings of the correlation analysis between subject learning performance and attentional capture effect. The blue line indicates the best linear fit and the black line indicates 0 (no difference between RT for singleton targets in the rewarded versus unrewarded dimension, below which subjects showed an attentional bias in the direction opposite than what we would expect (e.g. a color group subject shows faster RTs for orientation singletons compared to color singletons)).

The analysis seems to indicate that the less accurate subjects were in predicting the average point worth of tokens, the stronger their attentional capture effect was by the rewarded dimension. To seek further evidence of this effect, subjects were split into three groups depending on their accuracy: a low accuracy group (prediction accuracy from 68-84%; $N=9$; 4 color and 5 orientation group subjects), a middle accuracy group (84-91%; $N=10$; 5 color and 5 orientation group subjects), and a high accuracy group (91-100%; $N=10$; 6 color and 4 orientation group subjects). A 3 (accuracy: low, mid, high; between subjects) $\times$ 2 (group: color, orientation; between subjects) $\times$ 2 (target dimension: color, orientation) ANOVA revealed a significant three-way interaction, $F(2, 23) = 7.506, p < 0.005, \eta^2 = 0.395$. The interaction between target dimension and group became significant, $F(1, 27) = 6.078, p < 0.05, \eta^2 = 0.21$. There was also a trending interaction between target dimension and accuracy, $F(2, 23) = 2.566, p < 0.10, \eta^2 = 0.18$, as well as a main effect of target dimension, $F(1, 27) = 8.567, p < 0.005, \eta^2 = 0.271$. Overall, the findings of this ANOVA suggest that the variability in subject reaction times can be much better accounted for once subject accuracy is introduced as a factor. The findings from this three-way analysis can be visualized in figure 3.9.

When the low, middle, and high accuracy subject groups were analyzed separately, the attentional capture effect was most robust for the low accuracy subjects, revealing a significant two-way interaction between target dimension and group just as hypothesized, $F(1, 7) = 16.803, p < 0.005, \eta^2 = 0.706$. By contrast, there was no evidence of a two-way interaction in either the middle accuracy subjects, $F(1, 8) = 0.606, p = 0.459, \eta^2 = 0.070$, or the high-accuracy subjects, $F(1, 8) = 0.937, p = 0.361, \eta^2 = 0.105$. This suggests that the low-accuracy subjects were largely responsible for driving the trending effect previously reported in the critical test of the hypothesis (fig. 3.6)

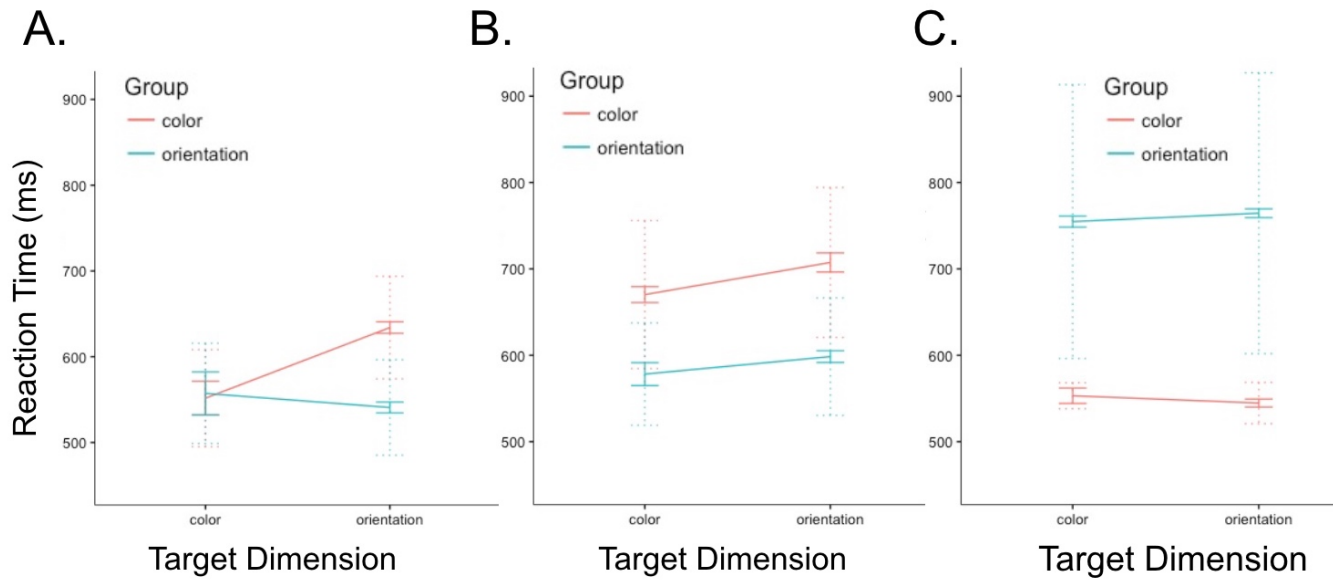


Figure 3.9: Interaction plots showing reaction times to visual search trials, broken down by target dimension and group for *A) subjects in the low accuracy group; B) subjects in the middle accuracy group; and C) subjects in the high accuracy group*. Again, since this was a mixed effect design with both a within subjects variable (target dimension) and a between subjects variable (group), two types of error bars were needed for each point. The solid lines are within-subject standard error (Cousineau, 2005), and the dotted lines are between subject standard error. N=9 low accuracy subjects (*A*), N=10 middle accuracy subjects (*B*), N=10 high accuracy subjects (*C*).

3.3 Discussion

In this experiment, participants viewed a series of “tokens” that varied in the two dimensions of color and orientation and learned the rules that predict which stimuli are worth the highest reward. Specifically, subjects had to figure out that one feature in every experimental block was associated with high reward, and more generally that the highly rewarded features in every block were all drawn from the same dimension (color features if they were assigned to the color group, orientation features if assigned to the orientation group). During this process, visual search trials sporadically appeared, utilizing competing singletons to probe for a dimensional attention bias towards either color or orientation. Analyses of subject predictions of token value

at the end of each experimental block indicated that subjects were largely successful at learning which tokens were worth the highest reward and which tokens were not.

This experiment sought to investigate the effects of reward learning on how our attention is captured and directed in the future. When specific color features were consistently associated with higher reward, subjects were faster to find color singletons even in a low-reward color in a visual search array, compared to orientation singletons. This modulation of attention towards the color dimension by reward was absent when only orientation features were consistently associated with high reward. Here, subjects showed no difference in response times for color compared to orientation singleton targets. In summary, reward seems to be able to drive attentional capture at the dimensional level, at least for color. The present results are largely consistent with my hypothesis, but two main caveats must be addressed.

The first caveat is that there was no observable attentional bias towards orientation in the group where orientation features, not colors, were consistently rewarded. The similarity in the response time for color and orientation singleton targets would suggest that subjects in this group received neither the attentional facilitation that the color group experienced towards color singletons, nor the attentional facilitation that should have been conferred by their own rewarded dimension. The lack of a corresponding attentional bias towards the orientation dimension can at least in part be responsible for why the ANOVA interaction between group and target dimension was trending but not significant. One likely explanation is that the color singletons were inherently more salient than orientation singletons in the visual search trials (for an intuitive feel of this, refer back to the sample arrays in figure 3.2A). This was verbally reported by a number of pilot subjects, and was also evidenced by a main effect of singleton dimension on reaction time, suggesting that subjects were overall faster to find color singletons than orientation

singletons. Indeed, there is evidence in the literature that orientation singletons are more difficult, or require more effort, to identify than color singletons. For example, Holguín et al. (2009) designed a task using EEG in which visual search targets alternated in blocks between color singleton targets, and orientation singleton targets. When orientation singletons acted as targets, they elicited slower reaction times than the color singletons and larger N2p amplitude, an ERP component related to visual selective attention. As a result, it is difficult in the context of this study to determine whether attentional capture does not generalize to the orientation dimension compared to color, or on other hand whether the bottom-up salience of color is simply overriding the more subtle effects of reward towards dimensional attention in the orientation group.

The second caveat to these results was the surprising finding in the exploratory analysis that learning performance, as measured by subject prediction accuracy, was negatively correlated with an attentional capture effect by the rewarded dimension. At first, this finding seems directly in conflict with the present hypothesis: if feature-based reward learning leads to attentional capture by the dimension, then how could it be that the less well subjects learned the rules determining reward, the stronger they demonstrated the hypothesized attentional capture effect? In fact, the negative correlation between prediction accuracy and the attentional capture effect was so robust that “low accuracy” subjects, ranging from 68-84%, were almost entirely responsible for driving the modulation we observed of color singleton target search by reward.

To explain this finding in the framework of my hypothesis, I propose that the negative correlation between subject learning performance and the attentional capture effect provides evidence that different subjects engaged in different types of learning throughout the task. Subjects with higher prediction accuracy, from 85-100%, clearly learned the reward structure of the game very quickly. They may have learned that in every game, tokens with one feature were

more likely to be worth close to 40 points while tokens without the feature were worth close to 10 points. Once they became confident in this rule, continuing to perform well throughout the task no longer required reward learning. One who is aware of the reward structure of the game could determine the highly rewarded feature within just a short stream of collect trials and predict the worth of tokens within close range of the correct value at the end of each game. For these subjects, the task may have become too easy, and factors crucial to reward learning such as prediction error (Rouhani, Norman, & Niv, 2017; Sunny, Manjaly, & Kumar, 2015) may have gone away. This would be consistent with the time course analyses performed, in which there was no evidence to conclude that attentional effects at the beginning of the experiment or each experimental block differed significantly from later trials where subjects had presumably learned the reward associations better.

By contrast, subjects with lower prediction accuracy, from 68-84%, indisputably still learned the general reward structure, as they were able to distinguish between high and low reward tokens the majority of the time. What their lower prediction accuracy might signify is simply that they were less confident or possibly even unaware of the rule that only one feature every game predicted high reward; as a result, their reward learning throughout the task would have been a more continuous process, involving more prediction errors and updating of reward value. The implications of this possibility for our understanding of value-driven attentional capture and for the present task will be further explored in Chapter 5, General Discussion.

Overall, given these caveats, two clear paths for a re-design of the Token Task become apparent. The first is to attempt to create the reward modulation effect of dimensional attention towards orientation, perhaps by increasing the salience of orientation singletons relative to color. This would also have the benefit of exploring whether the effects we observed in this experiment

were robust or generalizable to other types of attention probes. The second possibility is to redesign the reward structure and learning experience of the task to ensure standard reward learning processes, possibly by making the relevant rules predicting reward more difficult to figure out. In the following chapter, I describe an experiment in which I pursued the former option.

Chapter 4

Experiment 2: A Replication with Variation

The goal of this experiment was to perform a replication with variation on experiment 1, and to specifically investigate whether the effects of attentional capture previously tested could generalize to a different type of visual search array. In the previous experiment, we observed an attentional capture effect in which subjects in the color group were able to find color singletons significantly faster than orientation singletons in a visual search array. Subjects in the orientation group did not show any differences in reaction time between color and orientation singleton targets, suggesting that rewarding the orientation dimension did not facilitate search for pop outs in this dimension. I theorized that this was due to the color singletons being inherently and perceptually more salient than orientation. Thus, a secondary goal of this experiment was to redesign the visual search array in order to decrease the salience of the color singleton while increasing the salience of the orientation singleton. My hypothesis for this experiment remains the same: subjects who learn that color features are rewarding should demonstrate a dimensional attention bias towards color singleton targets, while subjects who learn that orientation features are rewarding should demonstrate the same bias towards orientation singleton targets.

4.1 Methods

4.1.1 Participants

Thirty-one new members of the Princeton University community (all undergraduates; 18 females; age range 18-22 years; mean age of 19.8 years) participated in the experiment for course credit plus a \$3 bonus for performance. One subject misunderstood the instructions to the experiment and was excluded, leaving 30 participants (18 females; age range 18-22 years; mean

age of 19.8 years). All subjects reported normal or corrected-to normal visual acuity and color. Study materials and procedures received approval from the Princeton University IRB, and subjects provided informed consent.

4.1.2 Stimuli

The design and presentation of stimuli, or “tokens”, were the same as in experiment 1, with some modifications. First, the features of the nine tokens used exclusively for collect trials (fig. 3.1A) were changed. Colors (green, purple, red) were slightly muted such that they popped out less against the grey background in order to better equate the bottom-up saliency of the color and orientation singletons in the visual search trials. Additionally, one potential factor distinguishing the way these features could be learned and represented in the brain in Experiment 1 was that the naming of the features. For example, the color features were easily named (e.g. red) while the orientation features were not (e.g. 15°). The orientation features were changed to 45°, 90°, and 135° (e.g. horizontal, left diagonal, right diagonal) for this new experiment to reduce this potential difference.

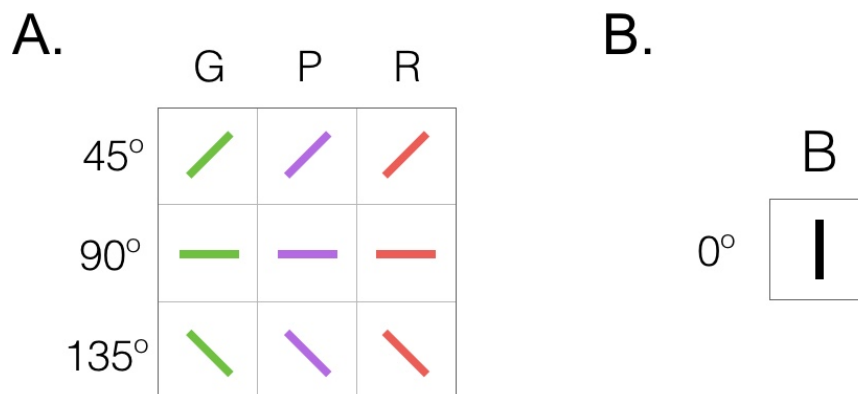


Figure 4.1. Tokens used throughout the experiment. **A)** The 9 tokens used for the collect trials. These tokens and their features would be associated with reward throughout the experiment. **B)** All non-singleton distractors in the visual trials were made the same to control for potential differences in luminance and salience.

I also introduced an additional control in this experiment: all non-singleton distractors would always be the same: black, vertical tokens (fig. 3.1B). This ensured both that distractors did not contain features that might be associated with reward and that some distractors were not more salient or more distracting than others.

As with the previous experiment, a competing singleton design was used for the visual search trial, with one color singleton and one orientation singleton. Again, these singletons only differed from distractors in their pop out feature: color singletons were either green, purple, or red, as in the collect trials, but were otherwise vertical to match the distractors. In comparison, orientation singletons were black to match the distractors and popped out solely due to being oriented 45°, 90°, or 135°.

A 7x7 array was used for the visual search trials in this experiment (fig. 4.2).

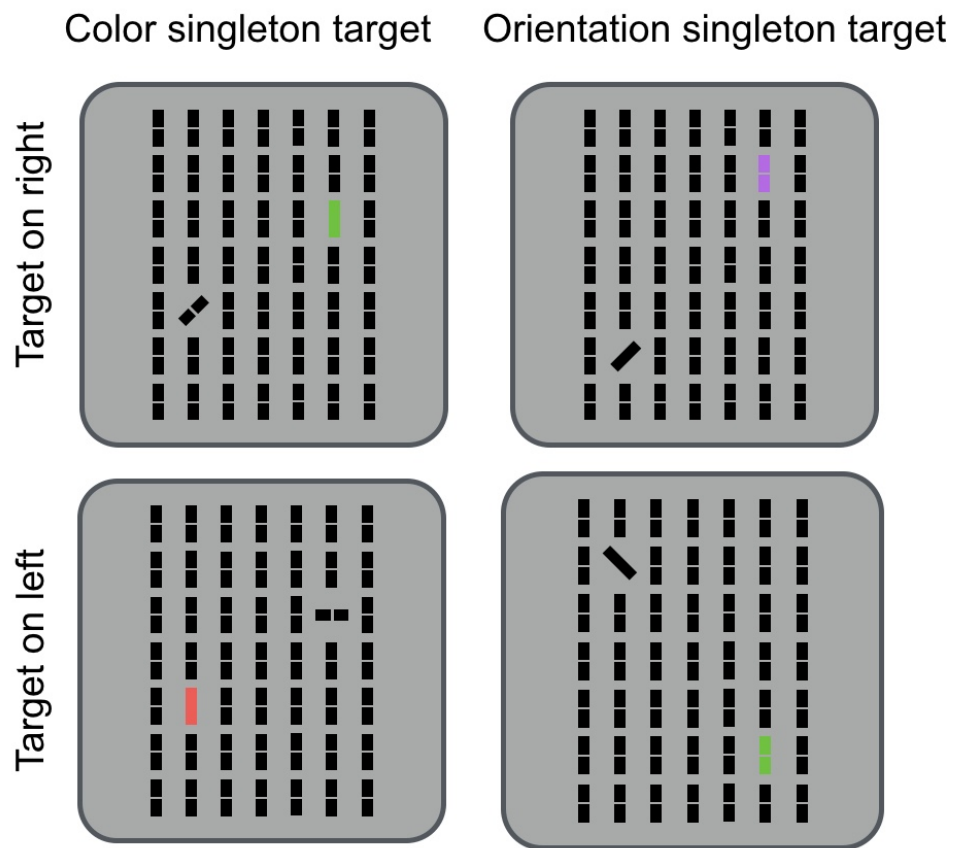


Figure 4.2. Four examples of the visual search arrays used in this experiment.

The dramatic increase in the size of the visual search array was designed to enhance the pop out of the orientation singleton, as discussed earlier. Orientation singletons tend to be more salient when completely surrounded by a contrasting background (Bogler, Bode, & Haynes, 2013) as opposed to being surrounded only on two sides as in Experiment 1. The competing singletons could only appear in the second and sixth columns of the array and in positions that were equidistant from the center of the screen. The left/right discrimination task for the visual search trials remained the same.

4.1.3 Procedure and Statistical Analyses

The primary focus during this experimental re-design was the visual search array trials. The experimental procedure was the same as in Experiment 1. Additionally, as in experiment 1, all data were recorded and processed using MATLAB. Statistical analyses and tests on the data were carried out using R, with $\alpha = 0.05$, and when stated, within-subject error bars were calculated using the procedure described in Cousineau (2005).

4.2 Results

4.2.1 Trial Accuracy

As previously described, subject performance on collect, visual search, and predict trials was analyzed to determine outliers with unusually poor performance for exclusion (table 4.1). Based on these criteria, four subjects were excluded, leaving 26 subjects total remaining for further analyses (13 in the color group and 13 in the orientation group).

Trial Type	Accuracy Measure	Subject mean (SD)	Exclusion Criteria
Collect trial	% tokens collected	99.82% (0.57%)	< 98.69%
Visual search trial	% correct response	91.98% (13.71%)	< 64.55%
Predict Trial	% correct answers	81.54% (13.79%)	< 53.96%

Table 4.1: overall subject accuracy on collect, visual search, and predict trials. Exclusion criteria for the three types of trials were determined as previously described in experiment 1, based on varying definitions of a statistical outlier.

4.2.2 Learning Performance

Participants again showed rapid and robust evidence of reward learning (fig. 4.3A-C).

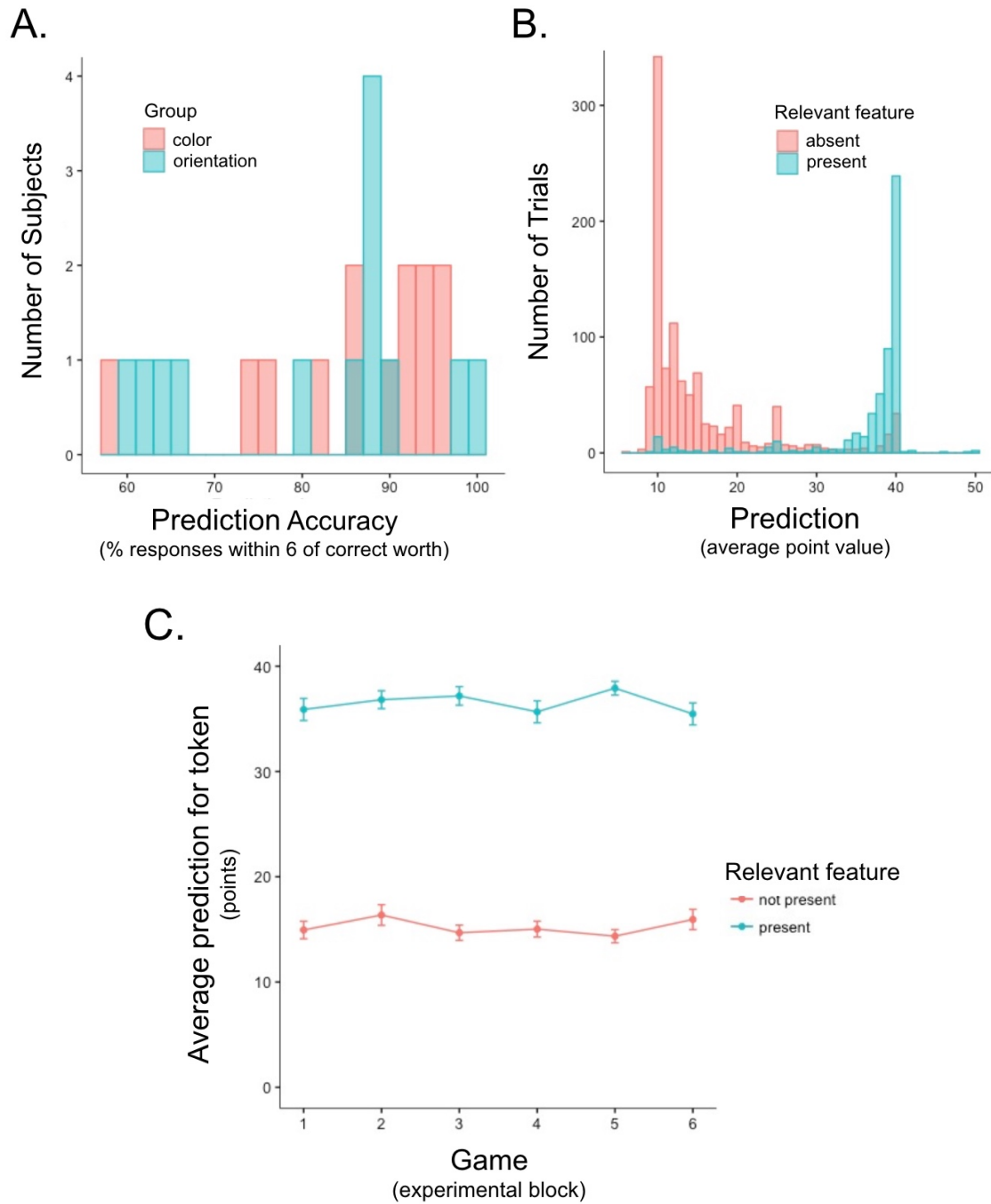


Figure 4.3: Assessment of subject learning performance based on prediction trials. *A)* Histogram of subjects' overall prediction accuracy. *B)* Histograms showing distribution of all subject predictions for tokens with (high reward) and without (low reward) the relevant feature. *C)* An aggregated learning curve over time, quantified by average prediction for tokens with and without the relevant feature in each game.

The subject prediction accuracy (fig. 3.4A) was 87.0% with a range from 57.4 to 100%, meaning that most subjects were able to consistently predict which tokens were highly rewarded and which ones were not. The aggregated distribution of the specific predictions that subjects made for both tokens with and without the relevant feature (fig. 3.4B) also demonstrates that the vast majority of predictions successfully distinguished between high reward (blue, on the right) and low reward (red, on the left) stimuli. Consistent with the previous experiment, the most common predictions were 10 for low reward tokens and 40 for high reward tokens, suggesting that many participants learned the reward structure without taking into account the probabilistic 90%/10% feedback. Finally, the learning curve (fig. 3.4C) demonstrates that most subjects learned the relevant rule for determining high reward quickly, before the end of the first game. In written feedback acquired from a questionnaire following the study, all but two subjects were able to correctly identify the dimension that was more relevant to predicting reward. On average, on a scale from 1 (very easy) to 5 (very difficult), subjects rated the difficulty of the experiment 3.1 ($SD = 1.04$), suggesting that task the task felt neither too difficult nor too easy.

4.2.3 Reward and Attention

Visual search trials in which the subject responded incorrectly were excluded from the following analysis. To test my hypothesis, visual search trial reaction times (fig. 3.6) were analyzed with a 2 (group: color, orientation; between subjects) x 2 (target dimension: color, orientation; within subjects) mixed-effects ANOVA. In direct contrast to the analysis from Experiment 1, there was no evidence of an interaction between target dimension and group, $F(1, 24) = 0.034, p = 0.855, \eta^2 = 0.001$. There was a main effect of target dimension, $F(1, 24) = 13.567, p < .005, \eta^2 = 0.36$, suggesting that subjects were faster overall to find color singletons.

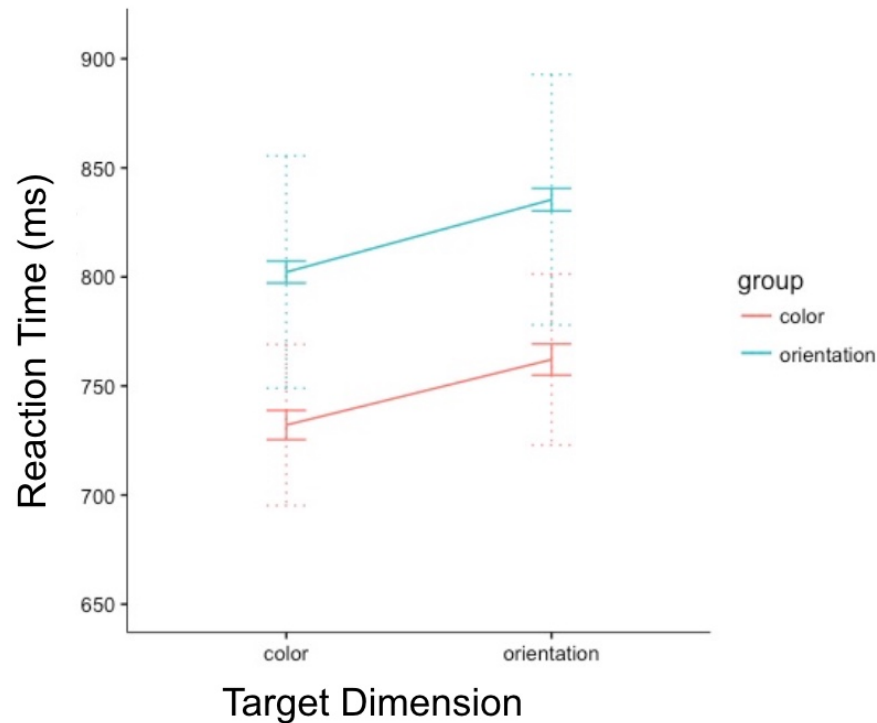


Figure 4.4: Interaction plot showing reaction times to visual search trials, broken down by group and target dimension. Two types of error bars were plotted for each point: the solid lines are within-subject standard error for comparison within group, by target dimension (Cousineau, 2005), and the dotted lines are between subject standard error. N=13 color group subjects, N=13 orientation group participants.

It is clear from the statistics as well as figure 4.4 that the group subjects were assigned to, hence the rewarded dimension throughout the experiment, did not appear to modulate reaction times for color singleton targets as opposed to orientation singleton targets at all. By changing the design of the visual search array, I was unable to replicate the findings from the first experiment. Possible explanations for this will be explored in the following discussion sections.

4.2.4 Exploratory Analyses

Given the lack of an interaction between group and target dimension in these data as well as the robust main effect of target dimension (color singletons capturing attention much more strongly regardless of reward), the analyses previously performed in experiment 1 on the time

course of the experiment, and the correlation between prediction accuracy and attentional effects, did not produce meaningful results.

Analysis by Presence of Relevant Feature

When reaction time data was again broken down by the presence or absence of the highly rewarded feature in every game, a 2 (presence of relevant feature: absent, present; within subjects) x 2 (target dimension: color, orientation; within subjects) x 2 (group: color, orientation; between subjects) ANOVA revealed a significant two-way interaction between the presence of the relevant feature and the group, $F(1, 24) = 15.08, p < 0.001, \eta^2 = 0.385$ in addition to a main effect of target dimension as expected, $F(1, 24) = 10.26, p < 0.005, \eta^2 = 0.300$. Post-hoc paired t-tests revealed that subjects in the color group were slower to respond to visual search trials overall when the relevant feature was present, $t(25) = 3.142, p < 0.005, d = 0.19$. Subjects in the orientation group were faster to respond to visual search trials overall when the relevant feature was present, $t(25) = 3.142, p < 0.05, d = 0.12$. These effects can be visualized in figure 4.5.

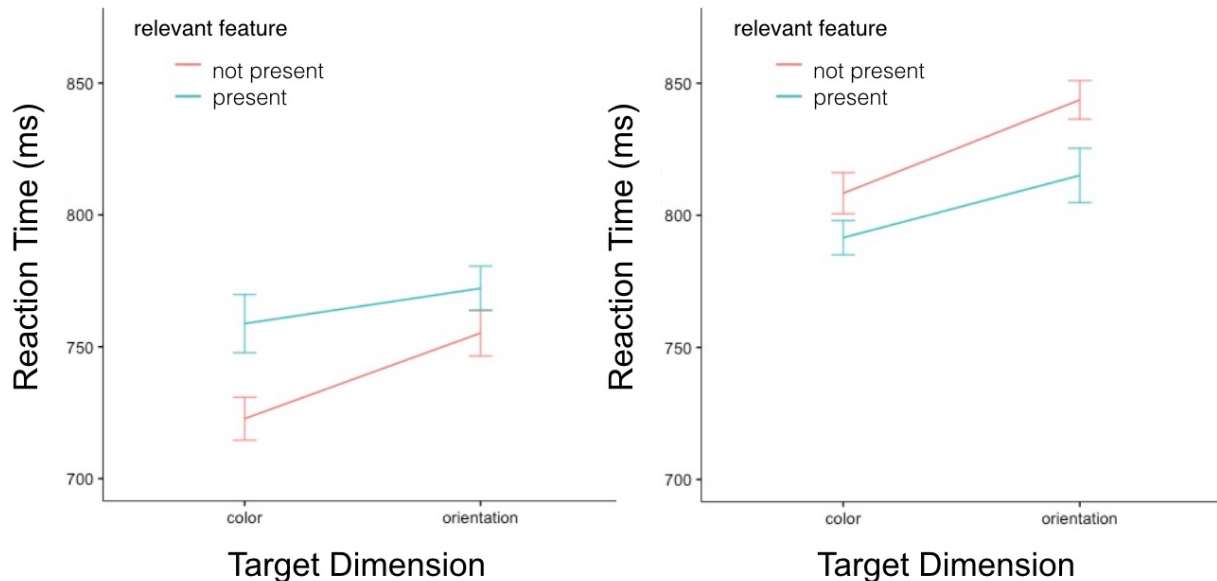


Figure 4.5: Interaction plots showing reaction times to visual search trials, broken down by target dimension and the presence of the relevant (highly rewarded) feature for **A) subjects in the color group** and **B) subjects in the orientation group**. Within subjects standard error was used. N=13 color group subjects (**A**), N=13 orientation group participants (**B**).

4.3 Discussion

The findings of this experiment are puzzling, especially in light of what I found in experiment 1. The learning experience, reward structure, and assessments of learning performance throughout both experiments were kept constant. The only changes made in this experiment were to the design of the visual search trial, in an attempt to create better controls while increasing the salience of the orientation singleton relative to the color singleton. Specifically, the set size of the visual search array was increased from 7x1 to 7x7, colors were muted, orientation angles were adjusted, and all distractor features were made to be black, vertical tokens.

As a result of making one or more of these changes, the reaction time data now failed to produce evidence of an interaction between group (the rewarded dimension) and the target dimension. Rewarding color features or orientation features consistently had no impact on how subjects' attention was deployed toward competing color and orientation singletons in a visual search array. Subjects of both groups were faster to find color singletons than orientation singletons, suggesting that the measures taken to increase orientation pop out were ineffective.

If the subject's group did have an impact on selective attention in this experiment, it was in unpredictable ways. The orientation group was slower overall to respond to the visual search trials, regardless of target dimension. Though this difference was not significant, this effect has been observed before; in the "high accuracy" subjects of Experiment 1 (fig. 3.9C), subjects in the orientation groups were also slower to respond compared to subjects in the color group. Why might rewarding the orientation dimension lead to slower responses than rewarding the color dimension? One potential explanation for this is that rewarding orientation features increases its bottom-up saliency (Brian A. Anderson et al., 2011), but not enough to override the color

singleton's greater physical saliency as previously discussed. Factors such as reward and physical salience are likely to both factor into an integrated attentional salience map (Koene & Zhaoping, 2007), but for the orientation group, both orientation and color receive elevated priority, increasing competition overall for both dimensions. Indeed, this could manifest in increased reaction time: Mounts and Tomaselli (2005) demonstrated using a visual search task that responses were slowest for a target stimulus when it was accompanied by competing high-salience stimuli, compared to when it was accompanied by low-salience stimuli. The color group will not have this same level of competition because increasing the saliency of the color dimension through reward will be congruent with the inherent physical saliency of the color singletons compared to the orientation singletons.

The significant interaction between the presence of the relevant feature and group was also puzzling, and it is both difficult and counterintuitive to explain. Since it is a two-way interaction that aggregates across target dimension, it indicates that the presence of the relevant feature, which is in only one dimension for each subject, would equally affect reaction time when it serves as both a target and a distractor in search trials. This seems surprising given the value-driven attentional capture literature, and calls for replication and future investigation.

4.3.1 Reconciling Experiments 1 and 2

How can the discrepancy in the findings between Experiments 1 and 2 be reconciled? One possibility that must be acknowledged is that the findings from experiment 1 were simply not replicable, and that there was not actually a modulation of dimensional attention towards color by reward. However, this conclusion cannot be made unless a direct replication of the original design is conducted. Another possibility is that learning about orientation features does not generalize to the dimension in the same way that learning about color features does. If color

and orientation are represented differently in the brain, then attempting to increase orientation pop out at the expense of color could greatly reduce ability to observe the hypothesized effect for color. However, assuming that reward can indeed influence attention at the dimensional level and that this experiment simply failed to detect it, the following are a number of proposed factors that explain why the findings of this experiment differed from the findings of experiment 1.

Different subject pools

Experiment 1 used a paid subject pool; participants were volunteers from the Princeton University community who expected to be compensated monetarily for their time. This made the administration of a monetary bonus based on point rewards a much more feasible and natural procedure. By contrast, Experiment 2 used a course subject pool, in which undergraduates at Princeton University participated in order to fulfill a course requirement. Several subjects revealed through verbal report at the end of the experiment that they had not been expecting to receive a bonus for the points they accumulated during the experiment due to these circumstances, even if this was mentioned in the instructions of the experiment. If subjects in this experiment did not correlate performance and the points they accumulated with reward, then it is difficult to say whether they were actually engaging in reward learning processes, for example involving dopamine (Berridge & Robinson, 1998), throughout the experiment. Subsequently, reward in the experiment may have had a weaker impact on attention.

Difficulty of the task

On average, subjects took around 100 ms longer to respond to visual search trials in this experiment (mean = 785 ms, SD = 240 ms) compared to subjects in experiment 1 (mean = 616 ms, SD = 245 ms). It is likely that increasing the size of the visual search array made searching for targets slightly more difficult. However, we observed in experiment 1 that the potential effect

of reward on dimensional attention is very small, averaging around 30 ms. One possibility is that even a slight increase in the difficulty of the visual search trial could reduce the sensitivity of the attention probe and mask a more subtle effect on reaction times by reward.

Location-based attentional selection

A noteworthy difference between this experiment and experiment 1 was that in experiment 1, the locations of the two singletons were always fixed at 2 positions in the visual search array, directly flanking the center token. By contrast, in experiment 2, the locations of the singletons could be in a number of locations in the 7x7 array, as long as they were equidistant from the center. There is robust evidence in the literature that attentional selection is heavily dependent on location (Hopf, Boelmans, Schoenfeld, Heinze, & Luck, 2002); in the case of experiment 1, it is likely that subjects learned that the target singleton would always appear in one of two locations, and would orient their attention accordingly.

The present hypothesis of the visual search design rests on the assumption that in the face of competing singletons, our selective attention will experience competition and will exhibit a bias based on a subtle effect of attentional priority driven by reward. Experiment 1 possibly creates a cleaner test of this by eliminating varying locations as a factor and relying on local competition between a target and its distractor (Reddy & Vanrullen, 2007). By contrast, the additional effort necessary to make longer-distance eye movements and to search in a much larger array could have overridden an effect of reward-biased competition in Experiment 2.

Larger set size, increased competition, and crowding effects

The use of a larger set size presents other potential problems as well. By surrounding the competing singletons in the visual searches with 47 other black, vertical tokens, the present experiment forced these singletons to compete with a much larger global onset of the whole

display. This is potentially overwhelming and swamps the salience of both singletons (Turatto & Galfano, 2000). As the hypothesis of this thesis relies only on competition of the two singletons with each other, increasing the global competition is largely unnecessary, if not problematic.

Additionally, the use of a much larger and denser search array reduces visibility of the target and subsequently hinders search performance. This is known as the crowding effect, or lateral masking in psychology (Strasburger, Harvey, & Rentschler, 1991). Notably, the effect has been shown to be stronger for the peripheral visual field rather than foveal vision (Bouma, 1970; Flom, Weymouth, & Kahneman, 1963), meaning that it would have a more negative impact on a large search array such as the one used in the present experiment, compared to the smaller line array used in Experiment 1.

Dependence of value-driven attentional capture on context

Finally, Anderson (2015) found that value-driven attentional capture is context specific. He showed that the same stimulus feature does or not capture attention depending on whether it was associated with reward within the context it appears. In the context of the present experiment, the sea of black tokens may have created a drastically different visual scene that could have led the visual system to believe that the reward-associated dimension was appearing in a new context. This differed sharply from Experiment 1 where the line array was much less overwhelming.

4.3.2 Lessons Learned

This experiment demonstrated how carefully a visual search trial must be designed in order to probe for evidence of reward-driven attentional bias towards a dimension. Psychophysics is a crucial consideration when investigating visual selective attention. The experiment was designed to be a controlled improvement upon Experiment 1, but despite only

making changes to the visual search array, the trends in the data were drastically different from the previous experiment. There was almost no interaction at all between reward and the dimension of the singleton target. This could but does not necessarily suggest that the findings reported in Chapter 3 of this thesis were not meaningful or robust. Rather, value-driven attentional capture by the dimension, if it exists, is a subtle effect that could be affected or masked by factors such as learning, stimulus crowding, location, context and competition. Future designs investigating this research question or related ones should pilot and optimize various types of visual search arrays that take these factors into consideration.

Chapter 5

General Discussion

Some material in this chapter was adapted from Bu (2016).

5.1 The Effect of Learning on Dimensional Attention

I began this thesis by asking whether feature-based reward learning generalizes to value-driven attentional capture at the level of entire feature dimensions. Dimensional attention has been demonstrated in behavioral paradigms and can be driven by statistical learning (Harris et al., 2015; Muller et al., 2003; Zhao et al., 2013). Reward is also known to teach selective attention, and dimensional attention plays a significant role in decision-making reinforcement learning paradigms (Chelazzi et al., 2013; Niv et al., 2015). In spite of these findings, the current research on value-driven attentional capture has largely ignored the possibility that dimensional attention could also be at play, and I sought to fill this gap.

In my experiments, subjects had to learn that only features from one dimension, either color or orientation, predicted high reward. I then used a competing singleton visual search array as an attention probe in order to test whether attentional priority can be biased towards faster reaction times in one singleton target dimension compared to the other. I ran two variants of the Token Task behavioral paradigm in order to test my hypothesis, using findings from my first experiment to inform the second.

Feature-based reward learning biases dimensional attention, at least for color

In the first experiment, I found evidence of dimensional attentional capture when colors were consistently highly rewarded, but not when orientations were consistently highly rewarded. Moreover, I found that attentional capture by the high-reward feature in every game did not

differ significantly from attentional capture by the low-reward feature in every game, allowing the data to further support the dimensional attention hypothesis. Whether there is truly a discrepancy between how reward acts on attention towards color and orientation, or if my attention probe simply failed to detect the effect of reward on orientation pop out search, remains an open area for future investigation.

The present experimental design also leaves one question unclear. Participants generally learned two separate rules over the course of the task. The first was that one feature per game would predict high reward. The second was that all highly rewarded features across games were coming from the same dimension. Which learning process drove the effect that I observed? It could be that learning about one feature generalizes to the whole dimension, or that learning about a rewarded dimension leads to attentional capture by that dimension. Future studies could try to disentangle these two questions, possibly by not allowing the relevant feature to switch every game.

The type of learning could influence value-driven attentional capture

Surprisingly, I found in Experiment 1 that learning performance was negatively correlated with an attentional bias towards the rewarded dimension. This potentially introduced a new variable into my experimental question: does the type of learning determine whether value-driven attentional capture by the dimension will be observed? Specifically, does value-driven attentional capture by a dimension require the presence of a continuous reinforcement learning process, using prediction errors to imbue stimuli with varying levels of value? I propose two separate types of learning that may have occurred during the experiment that could explain why low accuracy subjects drove the effect that I observed.

The first type of learning is reward learning as it is classically described (Sutton & Barto, 1998), and is more likely to have been present in the low accuracy subjects. Subjects begin each game not knowing which features are more highly rewarded than others. As they view more tokens with their associated reward and experience prediction errors, they gradually update the values of these tokens over time. With this type of learning, subjects may not have noticed that one feature in every game consistently predicted high reward, or were not confident in this rule, and this would thus manifest in lower prediction accuracy.

The second type of learning is not reward learning but “rule learning.” Subjects with higher accuracy predictions, who are likely to have used this form of learning, may have quickly realized that one relevant feature in every game predicted reward. As a result, they could confirm within just a small number of collect trials what the highly rewarded feature was. There would be no prediction error for the rest of the game, and the subjects would be able to consistently guess close to the correct value of tokens.

If this distinction were true and only subjects with the first type of learning showed evidence of value-driven attentional capture, this could provide valuable insight into how reward drives selective attention. Prediction error has already been shown to be able to account for some associative and perceptual learning (den Ouden, Friston, Daw, McIntosh, & Stephan, 2009; Law & Gold, 2009); it may be that the process of learning about high reward, rather than the acquired value itself, creates a more robust effect on modulating visual selective attention. This conclusion would also be consistent with the findings reviewed in Chapter 1.4, “Mechanisms of Value-Driven Attention,” which directly implicated reward structures of the brain in attentional capture.

Visual search arrays are a sensitive probe for dimensional attention

Experiment 2 showed that by just changing the search array, it is possible to get dramatically different results. I previously proposed a number of potential reasons for why I obtained different results. In short, the distractors, the visual scene, the location of the targets, and the eye movements necessary to perform a search are all worth considering and optimizing.

5.2 Experimental Limitations

In light of all we have learned, several limitations must be acknowledged when considering the findings and conclusions from the present experiments.

Statistics

A number of conclusions in this investigation relied on ANOVA. While this is generally a powerful test for investigating the contribution of categorical variables to a dependent variable, a number of continuous variables were considered as factors in thesis as well. For example, I examined how prediction accuracy and the time course of the experiment may have impacted subject reaction time. The prediction accuracy analysis was particularly critical and relied on arbitrarily splitting the subjects into three separate bins for the ANOVA. While a correlation analysis was performed, it only involved two variables. A reasonable future step with the present data would be to build a linear regression model accounting for all potential variables such as prediction accuracy, time, group, target dimension, and others, to better account for continuous variables and to create a more unified predictor of subject reaction times.

This investigation also relied heavily on a large number of predictions and statistical tests, ranging from three-way and two-way interactions to main effects and t-tests. As a result, I uncovered a lot of surprising findings, including the two-way interaction in the breakdown of reaction time by relevant feature for Experiment 2, and the relationship between prediction

accuracy and attentional capture effect. However, as the alpha threshold would tell us, if I ran twenty comparisons and all of them were significant, one of them is likely to be spurious even if it falls beneath the significance threshold. Corrections for multiple comparisons are one way to address this potential limitation.

Finally, I aimed for around thirty subjects after filtering for performance in each experiment. However, this number was arbitrarily chosen based on the subject number used in a previous and relevant study. Given how small the effect of reward on dimensional attention might be, I may not have had enough power for my statistical analyses, particularly when I proceeded to further break trials down by accuracy, game, and other factors. Power analyses could inform what a target subject number should be in the future.

Color versus Orientation

My hypothesis depended on a critical two-way interaction between group and target dimension. This prediction implicitly assumed that the effect of reward on dimensional attention would be similar for both color and orientation. However, evidence is accumulating that this is actually not the case. Shin and Ma (2017) showed that subjects are able to retrieve irrelevant color features from visual working memory at a much higher rate of precision than irrelevant orientation features. This suggests that color is overall a more memorable irrelevant feature, and is directly in line with my finding in Experiment 1 that the orientation group showed no difference in reaction time between orientation and color singletons. It is likely that in spite of orientation being rewarding, color was still salient and memorable to the orientation group. Moreover, to what extent color and orientation are represented asymmetrically in the visual cortices is still an active area of debate (Girelli & Luck, 1988; Tanigawa, Lu, & Roe, 2010); for example, Brouwer and Heeger (2013) argue that color is clustered categorically in the brain

whereas orientation is not. Given the landscape of this literature, the use of color and orientation as comparable dimensions in the experiment may lead to different and unpredictable effects that limit the data, a potential example being the analysis by presence of relevant feature in Experiment 2.

The use of color as a pop out dimension in this task may be particularly problematic. There is strong evidence that luminance is a specific aspect of color that receives attentional priority (Proulx & Egeth, 2008). Color and luminance have even been shown to have dissociable effects on attentional selection (Morrone, Denti, Spinelli, Moruzzi, & Ardeatina, 2002). In the experiments reported in this thesis, I did not control for the relative luminance of the colors that I used. This could explain why some colors may have popped out more than others, and why color overall may have been more salient than orientation.

Reward System

The reward system in the present experiments was not ideal. The Token Task rewarded subjects using a point system. This is somewhat far removed from actual reward, and is likely to be characterized as a tertiary reinforcer: it predicts money which subsequently predicts more direct rewards (Seward, 1950). Some studies in the attentional capture literature use a cash reward system (e.g., 1 cent for low reward and 5 cents for high reward), but this type of design was not practical for the present experiments due to the large number of trials in one run.

Moreover, though the points were converted to a cash bonus at the end, subjects usually ended up with as many as 14,000 points due to the large number of trials. Thus, the points may have been viewed as a less meaningful reward for subjects. Moreover, the dichotomy between 10 points for low reward and 40 points for high reward might confound the effects of reward learning with the effect of simply seeing a rare occurrence (+40). These limitations all depend on

how token value is learned and represented in each subjects' brain. Potential solutions would be to make the point distribution continuous (for example, have high reward be normally distributed around 40) or to redesign the experiment to convert to a cash reward system.

Reward History and Feature-Based Value-Driven Attentional Capture

Reward history and feature-based value-driven attentional capture are two factors that could potentially confound my conclusion about attentional capture at the dimensional level. For the pop-out singleton features in the visual search trials, I chose to use only features that were associated with varying rewards during the collect trials. This had two advantages: it allowed me to make a direct comparison of attentional capture by high-value features compared to low-value features, and it provided a test of my hypothesis while increasing my chances of observing the hypothesized effect.

The disadvantages of this choice must also be acknowledged. Feature-based value-driven attentional capture, to summarize, is when a feature associated with high reward (e.g., red) receives more attentional priority later. This phenomenon is well established and provides no evidence of generalization to dimensional attention. Given my design, if feature-based value-driven attentional capture occurred for the tokens with the relevant feature present in the visual search array, then they could drive and confound the effect I observed of dimensional attention bias for the color group. Even though a comparison of reaction times for search targets with and without the relevant feature revealed no significant differences, it is possible that feature-based value-driven attentional capture was still present.

Moreover, value-driven attentional capture is an effect that has been shown to persist long after the high-reward stimulus stops being rewarded (Failing & Theeuwes, 2014; Theeuwes & Belopolsky, 2012). It can be observed days or even months after experimental learning first took

place (Anderson & Yantis, 2013). In the present study, each of the three features in the relevant dimension of the collect trial tokens take turns becoming the relevant feature twice altogether per experiment. These are the same features that are used in the visual search trials. It is possible that once a feature becomes the relevant feature even once, its capture effect on attention will persist even when it becomes a low reward feature in the following game. This persistence of value-driven attentional capture could potentially explain the interaction we observed in Experiment 1 between group and target dimension, thus confounding our explanation that attentional capture by the dimension led to the interaction.

5.3 Future Directions

Replicating some of the findings in this thesis would be illuminating in itself, but in order to move forward, some of the aforementioned limitations must be addressed. I propose the following three emphases.

Manipulation of learning

One of the most robust findings from experiment 1 was that prediction accuracy was negatively correlated with the attentional capture effect, and that low accuracy subjects were largely responsible for the potential effect of dimensional attentional capture that I observed. This is intriguing and deserving of further exploration. A future study could possibly reduce the probabilistic reward feedback in the experiment to, for example, 75%/25% to make the learning task more difficult and to increase prediction error for all subjects. If we can again confirm the finding that more prediction error in learning is responsible for driving dimensional attention, this could have vast implications for the mechanisms by which perceptual learning and modulation of visual selective attention occur.

Optimization of the visual search array

A main takeaway from Experiment 2 is that the design of the visual search array is sensitive and crucial for probing dimensional attention. Factors such as size, crowding, and location of singleton targets are all worth considering and optimizing. If color is used as a dimension, then luminance must be taken into account. Li et al. (2014) found that color is organized along the framework of the HSL color space in V4; an ideal control then would be to choose colors for a new experiment that are isoluminant as determined by the HSL color space.

A stronger test of the hypothesis

Reward history and feature-based value-driven attentional capture were identified as potential confounds for the present study. A stronger test of the hypothesis that eliminates these confounds would be to use singletons in the visual search array with new features that subjects have never seen before. If these new features still capture attention without ever being associated with reward, it would provide robust evidence in support of attentional capture by the dimension.

Despite the findings reported in this paper, perhaps the greatest lesson to be learned is that there is still so much more to be done to understand how value-driven attentional capture generalizes into a multidimensional world.

References

- Anderson, B. A., Laurent, P. A., & Yantis, S. (2011). Learned value magnifies salience-based attentional capture. *PLoS ONE*, 6(11), e27926. <http://doi.org/10.1371/journal.pone.0027926>
- Anderson, B. A., Laurent, P. A., & Yantis, S. (2011). Value-driven attentional capture. *Proceedings of the National Academy of Sciences*, 108(25), 10367–10371. <http://doi.org/10.1073/pnas.1104047108>
- Anderson, B. A., Laurent, P. A., & Yantis, S. (2012). Generalization of value-based attentional priority. *Visual Cognition*, 20(6), 1–13. <http://doi.org/10.1080/13506285.2012.679711>.
- Anderson, B. A., Laurent, P. A., & Yantis, S. (2014). Value-driven attentional priority signals in human basal ganglia and visual cortex. *Brain Research*, 1587, 88–96. <http://doi.org/10.1016/j.brainres.2014.08.062>
- Anderson, B. A., & Yantis, S. (2013). Persistence of value-driven attentional capture. *J Exp Psychol Hum Percept Perform*, 39(1), 6–9. <http://doi.org/10.1037/a0030860>
- Berridge, K. C., & Robinson, T. E. (1998). What is the role of dopamine in reward: hedonic impact, reward learning, or incentive salience? *Brain Research Review*, 28, 309–369.
- Bogler, C., Bode, S., & Haynes, J. (2013). Orientation pop-out processing in human visual cortex. *NeuroImage*, 81, 73–80. <http://doi.org/10.1016/j.neuroimage.2013.05.040>
- Bouma, H. (1970). Interaction effects in parafoveal letter recognition. *Nature*, 226(5241), 177–178.
- Bouton, M. E. (2007). *Learning and behavior: A contemporary synthesis*. Sunderland, MA, US: Sinauer Associates.
- Bricolo, E., Giancesini, T., Fanini, A., Bundesen, C., & Chelazzi, L. (2002). Serial attention mechanisms in visual search: A direct behavioral demonstration. *Journal of Cognitive*

- Neuroscience*, 14(7), 980–93. <http://doi.org/10.1162/089892902320474454>
- Brouwer, G. J., & Heeger, D. J. (2013). Categorical clustering of the neural representation of color, 33(39), 15454–15465. <http://doi.org/10.1523/JNEUROSCI.2472-13.2013>
- Bu, J. (2016). *Generalization of value-driven attentional capture into a multidimensional world*. Unpublished manuscript, Department of Psychology, Princeton University, Princeton, NJ.
- Bu, J., Radulescu, A., Turk-Browne, N.B. & Niv, Y. *Feature-based reward learning biases dimensional attention*. Poster to be presented at the Vision Sciences Society Annual Meeting in St. Pete's Beach, FL May 2017.
- Bucker, B., & Theeuwes, J. (2017). Pavlovian reward learning underlies value driven attentional capture. *Attention, Perception & Psychophysics*, 79, 415–428. <http://doi.org/10.3758/s13414-016-1241-1>
- Chelazzi, L., Perlato, A., Santandrea, E., & Della Libera, C. (2013). Rewards teach visual selective attention. *Vision Research*, 85, 58–62. <http://doi.org/10.1016/j.visres.2012.12.005>
- Cousineau, D. (2005). Confidence intervals in within-subject designs: A simpler solution to Loftus and Masson's method. *Tutorials in Quantitative Methods for Psychology*, 1(1), 42–45.
- Dayan, P., Kakade, S., & Montague, P. R. (2000). Learning and selective attention. *Nature Neuroscience*, 3, 1218–1223.
- Della Libera, C., & Chelazzi, L. (2006). Visual selective attention and the effects of monetary rewards. *Psychol Sci*, 17(3), 222–227. <http://doi.org/10.1111/j.1467-9280.2006.01689.x>
- den Ouden, H. E. M., Friston, K. J., Daw, N. D., McIntosh, A. R., & Stephan, K. E. (2009). A dual role for prediction error in associative learning. *Cerebral Cortex*, 19(5), 1175–1185.

Retrieved from <http://dx.doi.org/10.1093/cercor/bhn161>

Dobson, K. S., & Dozois, D. J. A. (2004). Attentional biases in eating disorders: A meta-analytic review of Stroop performance. *Clinical Psychology Review*, 23(8), 1001–1022.

<http://doi.org/10.1016/j.cpr.2003.09.004>

Donohue, S., Hopf, J.-M., Bartsch, M., Schoenfeld, M., Heinze, H.-J., & Woldorff, M. (2016).

The rapid capture of attention by rewarded objects. *Journal of Cognitive Neuroscience*, 28(4), 529–541. <http://doi.org/10.1162/jocn>

Donovan, J. J., & Radosevich, D. J. (1999). A meta-analytic review of the distribution of practice effect: Now you see it, now you don't. *Journal of Applied Psychology*, 84(5), 795–805.

Duncan, J. (1996). Cooperating brain systems in selective perception and action. In T. Inui & J. McClelland (Eds.), *Attention and performance XVI. Information integration in perception and communication* (pp. 549-578). Cambridge, MA: MIT Press.

Failing, M. F., & Theeuwes, J. (2014). Exogenous visual orienting by reward. *Journal of Vision*, 14(5), 6. <http://doi.org/10.1167/14.5.6>

Fecteau, J. H., & Munoz, D. P. (2006). Salience, relevance, and firing: A priority map for target selection. *Trends in Cognitive Sciences*, 10(8), 382–390.

<http://doi.org/10.1016/j.tics.2006.06.011>

Field, M., & Cox, W. M. (2008). Attentional bias in addictive behaviors: A review of its development, causes, and consequences. *Drug and Alcohol Dependence*, 97(1-2), 1–20.

<http://doi.org/10.1016/j.drugalcdep.2008.03.030>

Flom, M. C., Weymouth, F. W., & Kahneman, D. (1963). Visual resolution and contour interaction. *Journal of the Optical Society of America*, 53, 1026–1032.

Flykt, A., Esteves, F., Ohman, A., Flykt, A., & Esteves, F. (2011). Emotion drives attention:

- Detecting the snake in the grass. *Journal of Experimental Psychology*, 130(3), 466–478.
<http://doi.org/10.1037//0096-3445.130.3.466>
- Girelli, M., & Luck, S. J. (1988). Are the same attentional mechanisms used to detect visual search targets defined by color, orientation, and motion? *Journal of Cognitive Neuroscience*, 9(2), 238–253.
- Grubert, A., & Eimer, M. (2017). A unitary focus of spatial attention during attentional capture : Evidence from event-related brain potentials, 13(2013), 1–11.
<http://doi.org/10.1167/13.3.9.doi>
- Harris, A. M., Becker, S. I., Remington, R. W., & Harris, A. M. (2015). Capture by colour: Evidence for dimension-specific singleton capture, 2305–2321.
<http://doi.org/10.3758/s13414-015-0927-0>
- Hickey, C., Chelazzi, L., & Theeuwes, J. (2010). Reward changes salience in human vision via the anterior cingulate. *J Neurosci*, 30(33), 11096–11103.
<http://doi.org/10.1523/JNEUROSCI.1026-10.2010>
- Hickey, C., Kaiser, D., & Peelen, M. V. (2015). Reward guides attention to object categories in real-world scenes. *Journal of Experimental Psychology: General*, 144(2), 264–273.
<http://doi.org/10.1037/a0038627>
- Holguín, S. R., Doallo, S., Vizoso, C., & Cadaveira, F. (2009). N2pc and attentional capture by colour and orientation-singletons in pure and mixed visual search tasks. *International Journal of Psychophysiology*, 73(3), 279–286. <http://doi.org/10.1016/j.ijpsycho.2009.04.006>
- Hopf, J., Boelmans, K., Schoenfeld, A. M., Heinze, H., & Luck, S. J. (2002). How does attention attenuate target-distractor interference in vision? Evidence from magnetoencephalographic recordings. *Cognitive Brain Research*, 15, 17–29.

- Hubel, D. H., & Wiesel, T. N. (1962). Receptive fields, binocular interaction and functional architecture in the cat's visual cortex. *The Journal of Physiology*, 160(1), 106–154.2. <http://doi.org/10.1523/JNEUROSCI.1991-09.2009>
- Kiss, M., Jolicœur, P., Dell, R., & Eimer, M. (2009). Attentional capture by visual singletons is mediated by top-down task set: New evidence from the N2pc component, 45(6), 1013–1024. <http://doi.org/10.1111/j.1469-8986.2008.00700.x>.Attentional
- Koene, A. R., & Zhaoping, L. (2007). Feature-specific interactions in salience from combined feature contrasts: Evidence for a bottom–up saliency map in V1. *Journal of Vision*, 7(6), 1–14. <http://doi.org/10.1167/7.7.6.Introduction>
- Laurent, P. A., Hall, M. G., Anderson, B. A., & Yantis, S. (2014). Valuable orientations capture attention. *Visual Cognition*, 6285(January), 1–14. <http://doi.org/10.1080/13506285.2014.965242>
- Law, C., & Gold, J. I. (2009). Reinforcement learning can account for associative and perceptual learning on a visual-decision task, 12(5), 655–663. <http://doi.org/10.1038/nn.2304>
- Lee, J., & Shomstein, S. (2014). Reward-based transfer from bottom-up to top-down search tasks. *Psychological Science*, 25(2), 466–475. <http://doi.org/10.1177/0956797613509284>
- Li, M., Liu, F., Juusola, M., & Tang, S. (2014). Perceptual color map in macaque visual area V4. *The Journal of Neuroscience*, 34(1), 202–217. <http://doi.org/10.1523/JNEUROSCI.4549-12.2014>
- Libera, C. Della, & Chelazzi, L. (2009). Learning to attend and to ignore is a matter of gains and losses. *Psychological Science*, 20(6), 778–784. <http://doi.org/10.1111/j.1467-9280.2009.02360.x>
- Mogg, K., & Bradley, B. P. (2006). Time course of attentional bias for fear-relevant pictures in

- spider-fearful individuals. *Behaviour Research and Therapy*, 44(9), 1241–1250.
<http://doi.org/10.1016/j.brat.2006.05.003>
- Morrone, M. C., Denti, V., Spinelli, D., Moruzzi, V. G., & Ardeatina, V. (2002). Color and luminance contrasts attract independent attention. *Current Biology*, 12(2), 1134–1137.
- Mounts, J. R. W., & Tomaselli, R. G. (2005). Competition for representation is mediated by relative attentional salience. *Acta Psychologica*, 118, 261–275.
<http://doi.org/10.1016/j.actpsy.2004.09.001>
- Muller, H. J., Reimann, B., & Krummenacher, J. (2003). Visual search for singleton feature targets across dimensions: Stimulus- and expectancy-driven effects in dimensional weighting. *Journal of Experimental Psychology*, 29(5), 1021–1035.
<http://doi.org/10.1037/0096-1523.29.5.1021>
- Niv, Y., Daniel, R., Geana, A., Gershman, S. J., Leong, Y. C., Radulescu, A., & Wilson, R. C. (2015). Reinforcement learning in multidimensional environments relies on attention mechanisms. *Journal of Neuroscience*, 35(21), 8145–8157.
<http://doi.org/10.1523/JNEUROSCI.2978-14.2015>
- Oca, B. M. De, & Black, A. A. (2017). Bullets versus burgers: Is it threat or relevance that captures attention ?, 126(3), 287–300.
- Proulx, M. J., & Egeth, H. E. (2008). Biased competition and visual search: the role of luminance and size contrast. *Psychological Research*, 72, 106–113. <http://doi.org/10.1007/s00426-006-0077-z>
- Reddy, L., & Vanrullen, R. (2007). Spacing affects some but not all visual searches: Implications for theories of attention and crowding. *Journal of Vision*, 7(2007), 1–17.
<http://doi.org/10.1167/7.2.3.Introduction>

- Roelfsema, P. R., van Ooyen, A., & Watanabe, T. (2010). Perceptual learning rules based on reinforcers and attention. *Trends in Cognitive Sciences*, 14(2), 64–71.
<http://doi.org/10.1016/j.tics.2009.11.005>
- Rouhani, N., Norman, K. A., & Niv, Y. (2017). Dissociable effects of surprising rewards on learning and memory. *bioRxiv*. Retrieved from
<http://biorxiv.org/content/early/2017/02/23/111070.abstract>
- Seeger, C. A. (2013). The visual corticostriatal loop through the tail of the caudate: Circuitry and function. *Frontiers in Systems Neuroscience*, 7(December), 104.
<http://doi.org/10.3389/fnsys.2013.00104>
- Serences, J. T. (2008). Value-based modulations in human visual cortex. *Neuron*, 60(6), 1169–1181. <http://doi.org/10.1016/j.neuron.2008.10.051>. Value-based
- Serences, J. T., & Saproo, S. (2010). Population response profiles in early visual cortex are biased in favor of more valuable stimuli. *J Neurophysiol*, 104(104), 76–87.
<http://doi.org/10.1152/jn.01090.2009>
- Seward, J. P. (1950). Secondary reinforcement as tertiary motivation: A revision of Hull's revision. *Psychological Review*. US: American Psychological Association.
<http://doi.org/10.1037/h0058496>
- Shin, H., & Ma, W. J. (2017). Crowdsourced single-trial probes of visual working memory for irrelevant features. *Journal of Vision*, 16(5), 1–8. <http://doi.org/10.1167/16.5.10>.doi
- Strasburger, H., Harvey, L. O. J., & Rentschler, I. (1991). Contrast thresholds for identification of numeric characters in direct and eccentric view. *Perception & Psychophysics*, 49(6), 495–508.
- Sunny, M. M., Manjaly, J. A., & Kumar, N. (2015). Attention capture as a function of prediction

error, (August).

Sutton, R. S., & Barto, A. G. (1998). *Reinforcement Learning : An Introduction*. Cambridge: MIT Press.

Tanigawa, H., Lu, H. D., & Roe, A. W. (2010). Functional organization for color and orientation in macaque V4. *Nature Neuroscience*, 13(12), 1542–1548. <http://doi.org/10.1038/nn.2676>

Theeuwes, J. (1992). Perceptual selectivity for color and form. *Perception & Psychophysics*, 51(6), 599–606. <http://doi.org/10.3758/BF03211656>

Theeuwes, J., & Belopolsky, A. V. (2012). Reward grabs the eye: Oculomotor capture by rewarding stimuli. *Vision Research*, 74, 80–85. <http://doi.org/10.1016/j.visres.2012.07.024>

Turatto, M., & Galfano, G. (2000). Color, form and luminance capture attention in visual search. *Vision Research*, 40, 1639–1643.

Yamamoto, S., Kim, H. F., & Hikosaka, O. (2013). Reward value-contingent changes of visual responses in the primate caudate tial associated with a visuomotor skill. *Journal of Neuroscience*, 33(27), 11227–11238. <http://doi.org/10.1523/JNEUROSCI.0318-13.2013>

Yantis, S. (2000). Goal-directed and stimulus-driven determinants of attentional control. *Control of Cognitive Processes: Attention and Performance Xviii*, 73–103. <http://doi.org/10.2337/db11-0571>

Yantis, S., & Jonides, J. (1984). Abrupt visual onsets and selective attention: Evidence from visual search. *Journal of Experimental Psychology*, 10(5), 601–621.

Zhao, J., Al-Aidroos, N., & Turk-Browne, N. B. (2013). Attention is spontaneously biased toward regularities. *Psychological Science*, 24(5), 667–677.

<http://doi.org/10.1177/0956797612460407>

Appendix A

Post-Study Questionnaire, administered to all subjects following experiments 1 and 2.

Debriefing questionnaire

Thank you for participating in the experiment. We would be grateful if you could answer the following questions.

What did you think this experiment was about?

Did you notice any patterns?

Did you notice that a particular characteristic was more likely to indicate whether tokens were worth more points? If so, which?

Rate (from 1 to 5) how interesting you found the experiment:

Very boring	1	2	3	4	5	Very interesting
-------------	---	---	---	---	---	------------------

Rate (from 1 to 5) how difficult you found the experiment:

Very easy	1	2	3	4	5	Very difficult
-----------	---	---	---	---	---	----------------

How much were you motivated to make correct predictions?

Not at all	1	2	3	4	5	Very much
------------	---	---	---	---	---	-----------

What did you think of the instructions to the experiment?

Not clear	1	2	3	4	5	Very clear
-----------	---	---	---	---	---	------------

Thank you!

Form for Collaboration in Senior Thesis Work

Please use this form to indicate the relationship between previous work and your senior thesis and to indicate whether your thesis involved collaboration with others.

Indicate below whether there is any overlap between your senior thesis and earlier work that you did for junior reports, junior papers, or papers for various courses.

Overlap X

No Overlap

If you checked the box indicating that there is overlap between your senior thesis and previous work, please describe the overlap on a **separate page**, and include it within the thesis after this form.

Readers of your thesis may, if they choose, ask to see earlier papers that you indicate have some overlap with your senior thesis.

Indicate below whether all or part of your thesis resulted from work done collaboratively with one or more other people.

Collaboration X

No Collaboration

If you checked the box indicating that your thesis work was done entirely, or in part, in collaboration with other people, describe the nature of the collaboration and what resulted from it on a separate page, and include it within the thesis after this form.

Statement of Collaboration

This thesis was based off of a project proposal written in the spring of my junior year, and a preliminary literature review that I did in the fall of my junior year. As a result, there is some overlap of material in the introduction and methods with these previous works. The abstract also overlaps with an abstract I wrote and submitted for this project to a national conference.

I completed this thesis in collaboration with Angela Radulescu, my graduate student mentor. We have been brainstorming ideas, analyses, and interpretations of our data together since I joined the lab in my junior fall. Angela has always given me feedback on my written work, presentations, and ideas.

Finally, this thesis would not have been possible without Profs. Yael Niv and Nick Turk-Browne, with whom I often met together with Angela to discuss data interpretation and ways to move forward.

Approval Form for Undergraduate Research Involving Experimental Animals

All research involving experimental animals at Princeton University must receive the prior approval from the Institutional Animal Care and Use Committee (IACUC). The IACUC bases decision about approval on the NRC Guide for the Care and Use of Laboratory Animals. All students conducting research involving animals as part of their junior independent work or senior thesis must receive approval from the IACUC prior to beginning their research. Students should consult first with their advisers about whether the procedures they intend to use are already covered by previously approved submission to the IACUC. The IACUC meets only once a month and it is common for new submissions to require revision before receiving approval so students are strongly encouraged to attend to IACUC issues early in their planning.

Did your Senior Thesis research involve the use of experimental animals?

Yes _____ No X

If you answered “Yes” to the above, you *must* also include a statement at the beginning of your methods section that verifies the work you have done with animals was approved by the Princeton University IACUC.

Lastly, please include this form at the back of your thesis (even if you answered “No”). If you answered “Yes,” please record the IACUC protocol number and date, below.

IACUC # _____ Approval Date _____

Approval Form for Undergraduate Research Involving Human Subjects

The Institutional Review Board for Human Subjects (IRB) is charged by the University Research Board with the task of protecting the interests and rights of human subjects involved in Princeton research. The IRB's responsibility includes the oversight of research conducted by undergraduate as part of their junior independent work and senior thesis work as well as that conducted in fulfillment of course requirements. All students conducting research involving human subjects as part of their junior independent work or senior thesis must receive approval from the IRB prior to beginning their research. Obviously, the sooner students submit their requests to IRB the sooner they will receive this approval. Students should be encouraged to submit their materials to the IRB as soon as possible in the semester. The IRB meets only once a month and it is common for student submissions to require revisions, primarily because of the incompleteness of the original submission, before receiving approval.

Did your Senior Thesis involve research with human subjects?

Yes X No

If your Senior Thesis DID involve research with human subjects, please indicate your IRB Case Number below.

Case # 4452 Approval Date 2/2017