

On the Role of Working Memory in Spatial Contextual Cueing

Susan L. Travis, Jason B. Mattingley, and Paul E. Dux
The University of Queensland

The human visual system receives more information than can be consciously processed. To overcome this capacity limit, we employ attentional mechanisms to prioritize task-relevant (target) information over less relevant (distractor) information. Regularities in the environment can facilitate the allocation of attention, as demonstrated by the spatial contextual cueing paradigm. When observers are exposed repeatedly to a scene and invariant distractor information, learning from earlier exposures enhances the search for the target. Here, we investigated whether spatial contextual cueing draws on spatial working memory resources and, if so, at what level of processing working memory load has its effect. Participants performed 2 tasks concurrently: a visual search task, in which the spatial configuration of some search arrays occasionally repeated, and a spatial working memory task. Increases in working memory load significantly impaired contextual learning. These findings indicate that spatial contextual cueing utilizes working memory resources.

Keywords: contextual cueing, working memory, attention, dual tasking

Our visual world is not random. Rather, it is a relatively stable environment in which particular objects and events predict the presence and spatiotemporal dynamics of other objects and events. When we search for misplaced keys, for example, we safely predict that they will be resting somewhere on a surface rather than, say, drifting in the air or dangling from a tree. Based on such predictions, we can prioritize some locations over others. This knowledge promotes more efficient search and attenuates the adverse influence of information processing limitations.

Contextual information, or the environment in which a stimulus appears, enhances search efficiency (e.g., Biederman, Mezzanotte, & Rabinowitz, 1982; Eckstein, Drescher, & Shimozaki, 2006; Hollingworth & Henderson, 1998; Palmer, 1975) and, specifically, the orienting of attention (Hoffmann & Kunde, 1999; Kunde & Hoffmann, 2005). Generally, targets are detected more efficiently when they are placed in a predictable context than when they are placed in an unpredictable context. For example, in natural scenes, observers use their knowledge about the regularities in the environment to facilitate target detection; it is easier to find a mailbox on the side of the street than it is to find the same mailbox in the middle of a field (Biederman et al., 1982).

Chun and Jiang (1998) demonstrated that the context in artificial scenes could also enhance target detection. In their seminal work on context effects in visual search, observers looked for a T-shaped target in arrays of L-shaped distractors. Half the trials contained repeat search arrays; that is, the specific distractor and target layouts remained consistent throughout the experiment, whereas the target's identity changed pseudorandomly. This consistency gave predictive information about the target's location but not the required response. On the other half of the trials, the spatial layouts did not repeat and, thus, those layouts could not be used to predict the target's location. Observers were not alerted to this manipulation, and in a follow-up test, observers performed at chance when asked to classify search arrays as repeated or novel. Despite the observers' apparent inability to recognize repeat arrays, over time, those observers detected targets faster for repeat search arrays than for novel search arrays. This phenomenon is referred to as contextual cueing. The authors concluded that people can learn contextual information to facilitate visual search and that this learning occurs automatically and largely without awareness (but see Smyth & Shanks, 2008, who provide some evidence of limited explicit knowledge in contextual cueing).

The behavioral mechanisms that underlie contextual cueing effects are not entirely clear. Chun and Jiang (1998) contend that as the consistency of repeated arrays does not predict target identity, any learning of the repeated arrays is perceptually based rather than response-based or effector-based learning (for research on perceptual-based learning, see Deroost, Coomans, & Soetens, 2009; Remillard, 2009). Other researchers (e.g., Kunar, Flusberg, & Wolfe, 2008), however, suggest that although search times are faster for repeated arrays than for novel arrays, most of this difference is due to the facilitation of other processing components, such as a lowering of the decision threshold, rather than being a consequence of context facilitating attentional guidance.

Recent evidence suggests that the long-term memory system plays an important role in contextual learning. The hippocampus is well known for its involvement in encoding long-term declarative

This article was published Online First May 28, 2012.

Susan L. Travis, School of Psychology, The University of Queensland, St. Lucia, Queensland, Australia; Jason B. Mattingley, School of Psychology and Queensland Brain Institute, The University of Queensland; Paul E. Dux, School of Psychology, The University of Queensland.

Susan L. Travis is now at the Queensland Brain Institute, The University of Queensland.

This research was supported by an Australian Research Council Discovery Grant and Fellowship (DP0986387) awarded to Paul E. Dux and by an Australian Research Council Laureate Fellowship (FL110100103) awarded to Jason B. Mattingley.

Correspondence concerning this article should be addressed to Paul E. Dux, School of Psychology, The University of Queensland, McElwain Building, St. Lucia, Queensland 4072, Australia. E-mail: paul.e.dux@gmail.com

memories (Squire, 1992), but it has also been implicated in spatial and contextual learning (Cohen & Eichenbaum, 1993; Hirsh, 1974; O'Keefe & Nadel, 1978). Chun and Phelps (1999) had amnesic patients with hippocampal and medial temporal lobe damage perform a spatial contextual cueing task. They found that these patients failed to show a contextual cueing effect. Converging evidence comes from a recent neuroimaging study, which demonstrated that the hippocampus is sensitive to repeated versus novel contexts (Greene, Gross, Elsinger, & Rao, 2007).

The studies described above clearly implicate long-term memory processes in contextual cueing, but human memory is not a unitary construct. Decades of research have demonstrated that different cognitive subsystems serve distinct memory operations (Atkinson & Shiffrin 1968; Eichenbaum & Cohen, 2001; Schacter, Wagner, & Buckner, 2000; Squire, Stark, & Clark, 2004; Tulving, 1985). For example, Atkinson and Shiffrin's (1968) multistore memory model distinguishes between short-term and long-term memory. Baddeley and Hitch (1974) proposed that the short-term retention of information is served by working memory (WM), a capacity-limited but flexible memory system that commonly involves processing task-relevant information for storage into long-term memory. Indeed, learning commonly involves processing task-relevant information in WM before it is stored in long-term memory.

An important question, therefore, concerns the role of spatial WM in spatial contextual cueing. If spatial contextual cueing depends on WM, then exhausting this resource should reduce the contextual cueing effect. Specifically, a difficult spatial WM task that draws heavily on spatial WM capacity should leave limited residual capacity for processing the spatial configurations in the search task. An easy spatial WM task, on the other hand, should leave considerable spare capacity, thus permitting simultaneous processing of repeated spatial configurations. Contrary to this prediction, however, a recent study has suggested that contextual learning in visual search may be unaffected by increases in concurrent spatial WM load.

Vickery, Sussman, and Jiang (2010) performed a series of experiments aimed at examining the potential contribution of various WM tasks on contextual learning in visual search. Each of their experimental paradigms involved an initial training phase and a subsequent testing phase. For example, in one experiment, the training phase required participants to remember the location of two sequentially presented dots while they searched for a target in a repeated search array. After responding to the search task, two dots were presented simultaneously. On matched trials, these dots occupied the same spatial locations as the sequentially presented dots, whereas on mismatched trials, one of the dots occupied a different location. Participants completed this task under three WM load conditions: a high-load condition (in which two dots were presented), a low-load condition (in which one dot was presented), and a no-load condition (in which no dots were presented). After training, participants performed a search-only testing phase in which half the search arrays were repeated from the training session and half were previously unseen, randomly generated arrays. Participants detected targets faster in the repeat search arrays than in the novel search arrays. Importantly, this contextual cueing effect was equivalent in magnitude under all load conditions, suggesting that participants learned the repeated spatial configurations even when WM resources were reduced. On

the basis of these findings, Vickery et al. suggested that spatial WM and contextual cueing draw on distinct processing resources.

Although Vickery et al. (2010) concluded that contextual learning does not depend on WM, three methodological issues in their study warrant further investigation. The first two issues concern their WM manipulation. First, their high-WM condition may not have exhausted WM resources fully. Luck and Vogel (1997; see also Cowan, 2001; Henderson, 1972; Posner, 1969; Sperling, 1960; Vogel, Woodman, & Luck, 2001) suggested that visual WM capacity is approximately four items. In the study of Vickery et al., observers held, at most, only two dots in WM. Thus, despite performing a spatial WM task, some residual capacity might have been available to process contextual information in the search arrays. Second, even if the load manipulation used by Vickery et al. did exhaust WM resources in their participants, it may be that this particular aspect of WM does not affect contextual learning, whereas others do. Specifically, the simultaneous presentation of dots in the retrieval phase of the WM task employed by Vickery et al. might have encouraged participants to encode the two dots as a single "barbell" object comprising two dots connected by a line. This object-based strategy might have freed spatial WM resources, thus allowing participants to encode contextual information from the search task. Similarly, Woodman and Luck (2004) demonstrated that processing spatial information interferes with visual search, whereas simply maintaining visual information in WM does not. Thus, processing spatial information in WM may interfere with contextual learning, as both tasks involve allocating attention and coordinating and maintaining spatial information.

To address these concerns, we had participants perform a demanding spatial WM task that required participants to remember the locations and the temporal sequence of up to four dots. We reasoned that having participants remember *both* the location and the temporal order of the dots would encourage a spatial strategy rather than an object-based strategy.

The third methodological issue concerns the approach used by Vickery et al. (2010) to measure contextual learning. In their study, contextual cueing was assessed exclusively under single-task conditions (after the dual-task training phase). This approach allows an examination of whether contextual cueing draws on WM resources during learning (with the caveats described above) but does not allow a determination of whether increases in spatial WM load affect the *expression* of contextual learning. Jiang and Leung (2005) found that a set of repeated distractors failed to enhance search times when these items did not share a feature with the attentional set, but once they did, these items immediately facilitated reaction time (RT). Thus, although the expression of contextual learning may depend on the attentional set maintained by participants, actually learning the spatial arrays does not.

Nissen and Bullemer (1987) also found a distinction between the expression of learning and learning itself when they investigated whether increases in WM load disrupt motor learning in a serial reaction time (SRT) task. In a typical SRT task, participants respond to visual stimuli presented at one of several spatial locations. Stimuli are presented sequentially, in either a repeated sequence or a random sequence. Despite being unaware that some sequences repeat, participants typically recall repeated sequences faster than random sequences. Interestingly, Nissen and Bullemer found that this effect vanished when participants performed a tone-counting task concurrently. But Frensch, Lin, and Buchner

(1998) demonstrated that although learning was not expressed under dual-task conditions, it was immediately evident when testing switched to the single-task condition. Thus, it appears that WM may be involved in the expression of implicit learning in the SRT task but not in the actual learning of the motor sequence itself (but see Shanks, Rowland, & Ranger, 2005, who provided evidence that sequence learning does not occur under dual-task conditions). Importantly, it is unclear whether this effect also occurs for contextual learning, as contextual cueing does not involve motor learning. As Vickery et al. (2010) tested contextual learning exclusively under single-task conditions, their investigation did not examine the role of WM in the expression of contextual learning. Separating learning from the expression of learning requires testing contextual learning under both single- and dual-task conditions.

On the basis of the considerations outlined above, the present study had two aims: to investigate whether contextual cueing draws on spatial WM resources and, if so, to determine at what level of processing WM load has its effect. In Experiment 1, we tested contextual learning in visual search under dual-task conditions, wherein a spatial WM task required participants to retain two items (low load) or four items in WM. The aim was to investigate whether contextual cueing is attenuated when spatial WM resources are exhausted. In Experiment 2, we examined at what level of processing WM load has its effect by examining contextual learning under single-task conditions similar to those employed in the study of Vickery et al. (2010). In Experiment 3, we tested contextual learning under both single- and dual-task conditions.

Experiment 1

In the first experiment, we addressed two questions regarding the WM manipulation used by Vickery et al. (2010). The first was whether their WM load manipulation was sufficient to moderate contextual cueing. To address this issue, we increased the number of items in the WM task from two to four in the high-load condition and from one to two in the low-load condition. We reasoned that holding four items in spatial WM would provide a more stringent test of whether contextual learning relies on this limited resource (cf. Cowan, 2001; Henderson, 1972; Luck and Vogel, 1997; Posner, 1969; Sperling, 1960; Vogel et al., 2001).

A second question was whether the components of the dual task employed by Vickery et al. (2010) might have encouraged participants to adopt a strategy of perceptually grouping pairs of dots into single objects ("barbells": two dots joined by a line; Feldman, 1997, 1999), thus potentially reducing WM load (Patterson, Bly, Porcelli, & Rypma, 2007). A grouping strategy is likely to reduce WM demands, as grouping dots into a single object should require fewer resources to consolidate than encoding each dot separately. To reduce the likelihood of any such strategy, we presented the probe dots sequentially in different spatial locations during initial encoding and subsequent matching. Probe dots in the matching display were presented either in the same sequence as in the encoding phase (matched trial) or in a different sequence (mismatched trial). We reasoned that participants should find it more difficult to adopt a grouping strategy with such sequential displays than with displays in which the dots are visible concurrently, as used by Vickery et al. Feldman (1997, 1999) found that simultaneously presented dots that form a collinear line tend to cohere into

a single object. Presenting dots sequentially may reduce the likelihood that the dots will be encoded using such strategies.

Method

Participants. Twenty-two students (11 men, 11 women; mean age = 19.68 years, $SD = 2.12$) from The University of Queensland volunteered for course credit. All reported normal or corrected-to-normal vision, and all were naïve to the purpose of the study. Each provided informed written consent.

Stimuli and apparatuses. Participants were tested in a dimly illuminated room. Stimuli were presented on a 24-inch LCD monitor with a resolution of $1,920 \times 1,200$ and a refresh rate of 60 Hz. Stimulus delivery and response recording was controlled using a Dell PC running Cogent software (Cogent 2000 toolbox: Functional Imaging Laboratory, Institute of Cognitive Neuroscience, and Wellcome Department of Imaging Neuroscience) in Matlab under Windows XP. The background color was set to gray (red-green-blue [RGB]: 128, 128, 128) across all conditions. Participants sat unrestricted at a viewing distance of approximately 57 cm.

Search task. The search task stimuli comprised one T-shaped target ($1^\circ \times 1^\circ$) rotated either clockwise or anticlockwise by 90° and 16 L-shaped distractors ($1^\circ \times 1^\circ$) shown upright (0°) or rotated by 90° , 180° , or 270° . The target and distractors were presented in Gill Sans MT Bold font. All search items were positioned randomly at the intersections of an imaginary 6×6 grid ($18^\circ \times 21^\circ$), and all search items appeared in white (RGB: 255, 255, 255). Participants were required to report whether the target was oriented to the "left" or to the "right" as quickly and as accurately as possible. Feedback was not given for the search task.

Working memory task. The WM stimuli consisted of either two or four white dots (RGB: 255, 255, 255; $1.4^\circ \times 1.4^\circ$) that were presented sequentially for 100 ms each, separated by a 400 ms blank interval. Dots could appear at 16 possible locations, and dots were positioned on two imaginary concentric circles, with any dot occupying one of eight possible equidistant locations on either circle. Dot positions on the inner circle were positioned at an eccentricity of 5° from the center of the screen, whereas dots on the outer circle were positioned 10° from the center of the screen. After the dots in the initial encoding display appeared, the search array of Ts and Ls was presented until participants made a key press. Immediately following this response, the dots in the probe display were presented. In half of the trials, the probe dots were presented at the same locations and in the same sequence as the dots in the initial encoding display (matched trials). In the other half of trials, the probe dots were presented in the same locations but in a *different* sequence than in the encoding display (mismatched trials). Participants were asked to judge whether the order of appearance of the probe dots was the same as, or different from, the sequence in the encoding display. There was no time pressure on this task, and feedback was given after participants responded.

Procedure and design. Each participant completed 24 trials of practice before undertaking the testing phase, which consisted of 20 blocks of 24 trials. Before the testing phase, 12 search arrays were randomly generated and were used as the repeat search arrays throughout the testing phase. Each block contained 12 repeat arrays and 12 new, randomly generated arrays (novel arrays). Each repeat-array was presented once in each block, for a total of 20

presentations. Repeat arrays were not presented in the practice block.

The spatial configuration of target and distractors remained consistent for each repeat array, that is, the target location and the distractor locations and orientations were invariant for each presentation. The target's orientation, however, was pseudorandomly selected for each presentation. This allowed the configuration of the repeat search arrays to predict the target's location, but not its orientation.

Repeat search arrays were allocated to one of three load conditions: high load (remember 4 dots), low load (remember 2 dots), or no load (ignore dots). This allocation remained consistent across blocks. High- and low-load trial sequences are shown in Figure 1. Each block was subdivided into three subblocks: high load, low load, and no load. The order in which the subblocks were presented was randomized for each block. At the beginning of each subblock, instructions were presented at the center of the screen. The instructions were "remember 4 dots" for high-load subblocks and "remember 2 dots" in the low-load subblocks. On half of the no-load subblocks, participants were instructed to "ignore 2 dots," and on the other half, they were instructed to "ignore 4 dots."

Participants pressed the spacebar to start each trial. The trial began with a fixation cross for 400 ms, followed by the two or four dots. A visual search array was then presented until participants

responded. They responded by pressing the *L* key for a target oriented to the left and the *R* key for a target oriented to the right. Immediately following the response, the probe dots were presented. On high- and low-load trials, participants judged whether the dots were presented in the same sequence as in the initial encoding display. They responded by pressing the *S* key for match trials and the *D* key for mismatch trials. On no-load trials, participants ignored the dots, and at the end of the trial, participants were instructed to "press any key" to continue.

Explicit search-array recognition test. After completing the testing phase, participants were asked to perform an explicit search-array recognition test. They were informed that half of the search arrays had been repeated throughout the experiment. In a forced-choice format, participants were asked to decide whether a given display was a repeat array or a novel array. Participants completed four blocks of 24 trials (12 repeat and 12 novel search arrays) and responded by pressing the *R* key (repeat) or the *N* key (novel). This was a nonspeeded task, and feedback was not given.

Results and Discussion

Data from four participants were excluded from the analysis, in one case because overall search performance was below 60%, in another two because the participants used incorrect response keys

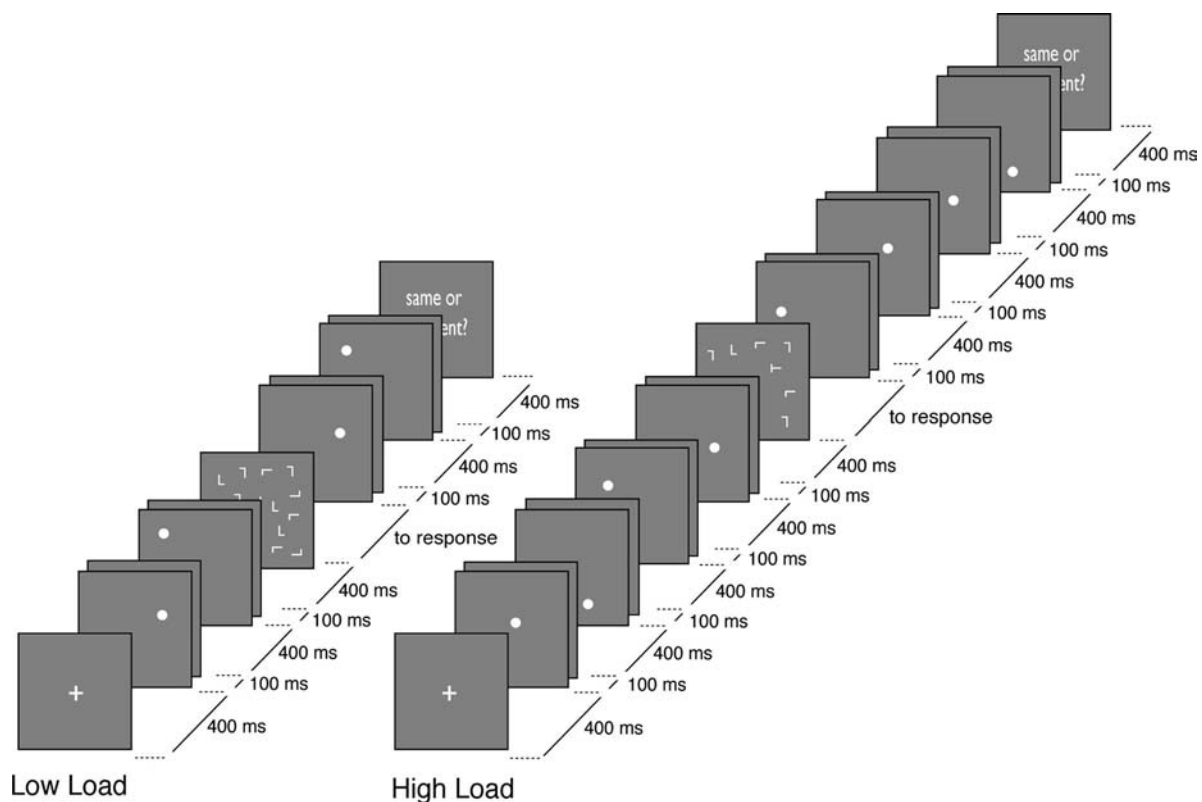


Figure 1. Schematic of the task sequence for low-load and high-load trials. On matched trials, the dots were presented in the same location and sequence before and after the search task (the low-load example here shows a matched trial). On mismatched trials, dots presented after the search task were presented in the same locations but in a different sequence from the dots presented before the search task (the high-load example here shows a mismatched trial). The search task always consisted of 16 distractors (Ls) and 1 target (T; not drawn to scale).

on more than four blocks of trials, and in another due to a technical error. Thus, the data from the remaining 18 participants were analyzed.

Accuracy in the WM task was well above chance (50%) for both the low-load WM task (94.17%), $t(17) = 47.00$, $p < .001$, and the high-load WM task (88.96%), $t(17) = 26.20$, $p < .001$. Importantly, a paired-samples t test revealed that accuracy was significantly higher in the low-load condition than in the high-load condition, $t(17) = 5.71$, $p < .001$, indicating that the high-load WM task was more difficult than the low-load WM task.

Accuracy on the search task was near the ceiling for all load conditions (no load, 98.13%; low load, 98.92%; and high load, 98.72%) and did not differ significantly across the three conditions, $F(1, 17) = 1.99$, $p = .15$. Mean reaction times (RTs) were calculated for each participant in each condition, excluding incorrect responses. To reduce the influence of outliers, a screening protocol was applied: Reaction times greater than 3 standard deviations above the participant's mean for each condition were removed, after which the standard deviation was recalculated. This process was completed a maximum of three times (2% of trials were removed using this outlier criterion).

To increase statistical power, the 20 blocks of the search task were binned into 10 epochs. Mean RT was then subjected to a 2 (context) $\times$ 3 (load) $\times$ 10 (epoch) repeated-measures analysis of variance (ANOVA). Results revealed a significant main effect of context, $F(1, 17) = 6.67$, $p = .019$. The mean RT for repeat search arrays was on average 57 ms faster than the mean RT for novel search arrays, indicating that participants learned the spatial configurations of the repeated arrays. There was also a significant main effect of load, $F(2, 34) = 5.47$, $p = .009$. Follow-up t tests, assessed with a Bonferroni correction (adjusted $p = .0167$), revealed that the RT was faster in the no-load condition than in the low-load condition, $t'(17) = 3.35$, $p = .004$, but did not reach the corrected significance level for comparisons between no-load and high-load conditions, $t'(17) = 2.38$, $p = .029$, or between low-load and high-load conditions, $t'(17) = 0.16$, $p = .878$. A significant main effect of epoch was also revealed, $F(9, 153) = 31.64$, $p < .001$. The mean RT was 308 ms faster in Epoch 10 than in Epoch 1, indicating a reliable practice effect across the task. Importantly, there was a significant interaction between epoch and context, $F(9, 153) = 2.45$, $p = .012$, indicating that the contextual cueing effect increased over time (from 18 ms in Epoch 1 to 111 ms in Epoch 10). In addition, a significant two-way interaction between load

and context was revealed, $F(2, 34) = 3.43$, $p = .044$. This interaction was followed-up with Bonferroni corrected t tests to assess the difference between repeat and novel arrays for each level of the WM load factor (adjusted p value = .0167). In the no-load conditions, mean RT was significantly faster for repeat arrays ($M = 1,056$, $SD = 185$) than for novel arrays ($M = 1,196$, $SD = 210$), $t(17) = 3.50$, $p = .003$ (see Figure 2). By contrast, RTs did not differ between repeat ($M = 1,218$, $SD = 224$) and novel arrays ($M = 1,233$, $SD = 180$) under low-WM load, $t(17) = 0.42$, $p = .678$, or between repeat ($M = 1,212$, $SD = 300$) and novel arrays ($M = 1,229$, $SD = 218$) under high-WM load, $t(17) = 0.43$, $p = .675$. Thus, it appears that participants were unable to learn contextual information in visual search when performing a concurrent WM task.

There was no significant interaction between epoch and load, $F(18, 306) = 1.46$, $p = .103$, and no significant three-way interaction between epoch, load, and context, $F(18, 306) = 1.21$, $p = .247$.

Overall accuracy on the explicit search-array recognition test (55%) was numerically close to, but statistically greater than, chance, $t(17) = 3.69$, $p < .002$. However, this difference was not reliable across blocks. Although accuracy rates were above chance in Block 2 (56%), $t(17) = 2.68$, $p = .016$, and Block 3 (58%), $t(17) = 3.20$, $p = .005$, accuracy did not differ from chance in Block 1 (55%), $t(17) = 1.68$, $p = .112$, or in Block 4 (52%), $t(17) = 0.73$, $p = .473$. These results are consistent with Smyth and Shanks (2008), who suggest that extended recognition tests can sometimes expose evidence of limited, explicit contextual knowledge.

These results demonstrate that spatial contextual learning is significantly reduced under concurrent WM load. Thus, the results of Experiment 1 differ from those of Vickery et al. (2010), who found that WM load did *not* affect contextual learning. The discrepancy between our results and those of Vickery et al. might be due to differences in the WM manipulation. In the current experiment, WM load was increased in two ways. First, we increased the number of elements to be remembered. Our high-load condition included four dots and our low-load condition included two dots, whereas in the study of Vickery et al. the WM tasks comprised two dots and one dot, respectively. Thus, the current low-load condition was numerically equivalent to the high-load condition of Vickery et al. However, we found no evidence of contextual learning under low WM load, suggesting that the discrepancy

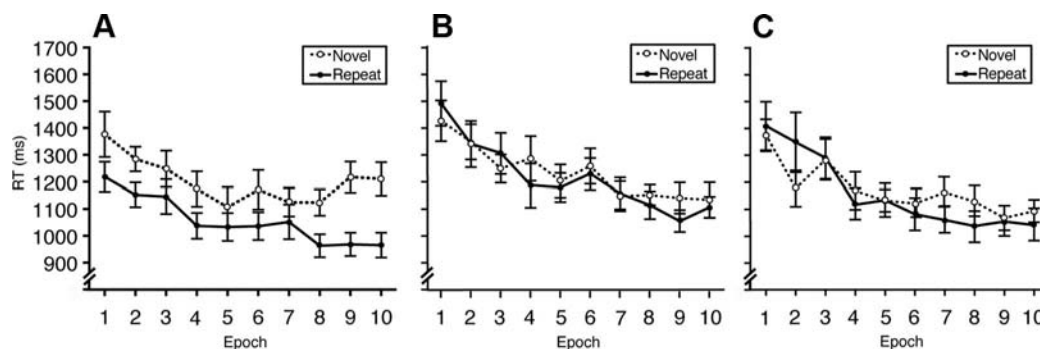


Figure 2. Mean reaction times (RTs) for visual search over epoch in the no-load (A), low-load (B), and high-load (C) conditions in Experiment 1. Error bars represent 1 standard error of the mean.

between our results and those of Vickery et al. are not merely due to differences in the number of elements to be remembered in the WM task.

A second critical difference between the two studies is that Vickery et al. (2010) presented WM items sequentially in the initial encoding phases and concurrently in the recognition phases, whereas we presented dot probes sequentially during *both* initial encoding and subsequent matching, requiring participants to remember both the spatial locations and the temporal order of the dot probes. It is possible that encoding such temporal information in multielement arrays requires access to a limited-capacity resource that is also required for spatial contextual learning.

Despite these subtle methodological differences, there is an alternative explanation for the discrepancy between our findings and those of Vickery et al. (2010): WM load may have interfered with the *expression* of learning across repeat search arrays rather than with contextual learning per se. Thus, WM load may not have influenced contextual cueing in the experiments of Vickery et al. because in their paradigm, the expression of learning never overlapped with the WM manipulation. Past research with SRT tasks has suggested that motor learning is not expressed under dual-task conditions, even though the learning effect can be shown to have taken place when testing is switched to a single-task condition (Frensch et al., 1998). This suggests that WM may play an important role in the expression of learning in SRT tasks, but not in the acquisition of motor sequence knowledge.

In Experiment 2, we tested the hypothesis that although contextual learning can arise under WM load, as suggested by Vickery et al. (2010), the learning effect cannot be expressed while participants are engaged in a concurrent WM task.

Experiment 2

In Experiment 2, we addressed the possibility that participants in Experiment 1 learned contextual information from the visual search array but that this learning was not expressed under dual task conditions. To do this, we separated the initial *training phase*, during which participants performed the search task with or without a WM load task, from a subsequent *testing phase*, during which participants performed the search task in isolation (i.e., without any concurrent WM task). If performing a concurrent WM task affects the expression of learning, but not learning itself, then a significant contextual cueing effect should arise when participants perform the visual search task in isolation (during the testing phase), after having been exposed to the dual task during the training phase.

Method

Participants. Twenty-three students (9 men, 14 women; mean age = 18.26 years, $SD = 1.66$) from The University of Queensland volunteered for course credit. All participants reported normal or corrected-to-normal vision, and participants were naïve to the purpose of the study. Each participant provided written informed consent.

Stimuli and apparatuses. The stimuli and apparatuses were similar to those used in Experiment 1, except that we excluded the low-WM load condition. Since the results of Experiment 1 showed no statistical difference in RTs between repeat search arrays per-

formed under low- and high-WM load, we reasoned that excluding the low-WM load condition would have the benefit of increasing statistical power without reducing the scope of our investigation.

Procedure and design.

Training phase. The training phase in Experiment 2 was similar to the testing phase in Experiment 1, with some important exceptions. First, participants did not complete a practice block before training. Second, 24 search arrays were randomly generated and used as the search arrays throughout the training session. Each search array in the training phase was repeated once in each block, for a total of 20 presentations. Thus, unlike Experiment 1, novel arrays were not presented during the training phase. Half of the search arrays were assigned to the no-WM load condition, and the other half were assigned to the high-WM load condition. In no-load trials, four dots were presented but could be ignored.

Testing phase. The training phase was followed by a testing phase. In the testing phase, participants performed four blocks of 48 search-only trials; crucially, there was no concurrent WM task during this phase. Half of the search arrays were repeat search arrays from the training phase, whereas the other half were previously unseen, randomly generated search arrays.

Explicit search-array recognition test. The explicit search-array recognition test was identical to that used in Experiment 1, except that participants performed two blocks of 48 trials (24 repeat, 24 novel).

Results and Discussion

Data from five participants were excluded from the analysis because their overall accuracy on the WM task (2 participants) or the search task (3 participants) was less than 60%. Data from the remaining 18 participants were included in the group analysis.

In the training phase, accuracy on the WM task was 85%, which was well above chance, $t(17) = 17.78$, $p < .001$. Accuracy on the search task was also high for both the no-load (95.05%), $t(17) = 33.05$, $p < .001$, and the high-load trials (96.16%), $t(17) = 41.65$, $p < .001$. Search accuracy did not differ between these conditions, $t(17) = 1.55$, $p = .139$. In addition, RTs did not differ significantly between high-load ($M = 1,209$, $SD = 273$) and no-load ($M = 1,144$, $SD = 161$) conditions, $t(17) = 1.12$, $p = .280$.

Overall accuracy in the testing phase was 96%, which was significantly above chance (50%), $t(17) = 67.80$, $p < .001$. Filtering outliers with the same procedure as that described in Experiment 1 resulted in the removal of 3% of the trials.

The mean RT from the testing phase was subjected to a 3 (learning condition: novel vs. no-load vs. high-load) $\times$ 4 (block) repeated-measures ANOVA. Results revealed a significant main effect of learning condition, $F(2, 34) = 8.21$, $p = .001$ (see Figure 3). This main effect was followed-up with Bonferroni corrected t tests for three comparisons (adjusted $p = .0167$). As expected, the mean RT for search arrays learned under no-WM load in the training phase was significantly faster ($M = 1,238$, $SD = 286$) than the mean RT for novel search arrays ($M = 1,352$, $SD = 245$), $t(17) = 3.78$, $p = .001$. Crucially, the mean RT for search arrays learned under high-WM load in the training phase was now also significantly faster ($M = 1,232$, $SD = 240$) than the mean RT for novel search arrays, $t(17) = 3.48$, $p = .003$. There was no significant difference between the RTs for search arrays learned

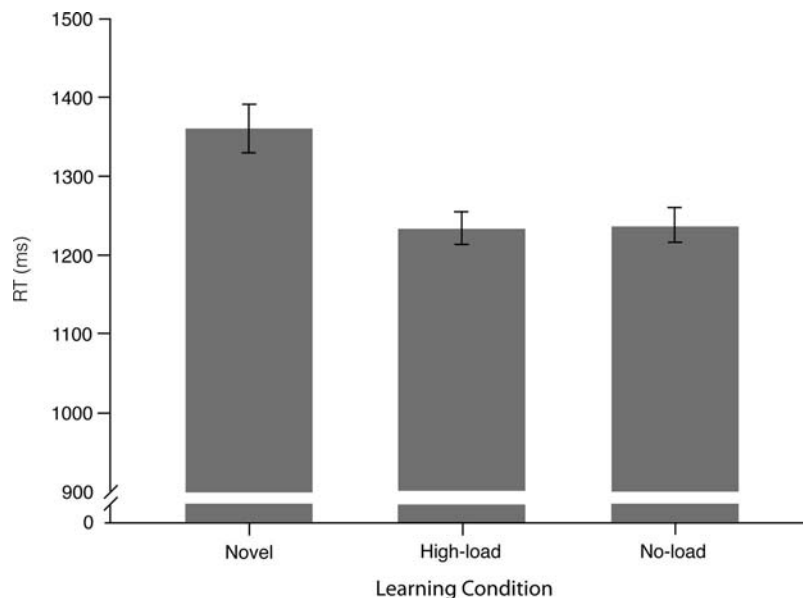


Figure 3. Mean reaction times (RTs) for visual search in the testing phase as a function of learning condition. Error bars represent 1 standard error of the mean.

under high-load conditions and those learned under no-load conditions, $t(17) = 0.18$, $p = .861$.

Results also revealed a main effect of block, $F(3, 51) = 3.72$, $p = .017$. Follow-up t tests, assessed with a Bonferroni correction (adjusted $p = .008$), revealed that the RT in Block 1 was significantly faster than the RT in Block 3, $t'(17) = 3.19$, $p = .005$ (see Table 1 for means and standard deviations). No other comparisons reached significance. If rapid learning occurred during the testing phase, we would expect to see faster RTs in later epochs, compared with earlier epochs. We did not find this pattern of results, suggesting that the learning effect seen in the testing phase is due to learning that occurred during the training phase. The interaction between learning condition and block was not significant, $F(6, 102) = 1.77$, $p = .112$.

Mean accuracy on the explicit search-array recognition test was 53%, which was statistically above chance, $t(17) = 2.98$, $p = .008$. As in Experiment 1, however, this effect was not consistent across blocks. Accuracy was significantly above chance in Block 1 (55%), $t(17) = 2.39$, $p = .029$, but not in Block 2 (53%), $t(18) = 1.72$, $p = .104$, again suggesting an inconsistent level of explicit recognition of the repeat search arrays.

Table 1
Mean Reaction Times (and Standard Deviations) in the Testing Phase of Experiment 2

Learning condition	Block			
	1	2	3	4
Novel	1,260 (277)	1,370 (262)	1,390 (321)	1,386 (280)
No-WM load	1,203 (305)	1,172 (288)	1,317 (336)	1,260 (324)
High-WM load	1,154 (232)	1,293 (340)	1,253 (280)	1,227 (251)

Note. WM = working-memory.

These findings suggest that spatial contextual learning is robust to increases in WM load, as originally proposed by Vickery et al. (2010), but contrary to the results of Experiment 1, in which we found no evidence for contextual learning under or high- or low-WM load. The findings from Experiments 1 and 2 might suggest that contextual learning of visual search arrays can occur under varying levels of WM load, but having participants perform the search task concurrently with a WM load task prevents the expression of this learning. There is, however, an alternative explanation for these findings that warrants further investigation.

If our findings in Experiment 1 and Experiment 2 differ because the expression of contextual learning was only impaired by WM load in the former experiment, we should expect to see a difference in search performance between the no-load (single-task) and load (dual-task) conditions in the two training phases. That is, the RTs should be faster for repeat arrays performed under single-task conditions than those performed under dual-task conditions. We observed this difference in Experiment 1, but not in Experiment 2, suggesting that the training phases were not equated between the two experiments.

A possible explanation for this difference is that in Experiment 1, both novel and repeat arrays were presented during the learning phase, but this was not the case in Experiment 2 (i.e., novel arrays in Experiment 2 were only included in the testing phase). It is conceivable that the task in Experiment 1 was made more difficult by the interleaving of novel search arrays, thereby contributing to the absence of a contextual cueing effect under dual-task conditions in Experiment 1. Experiment 3 was designed to test this possibility.

Experiment 3

Experiment 3 was designed to distinguish between two alternative explanations of the results from Experiment 1 and Experiment

2: the *expression of learning* account and the *novel-array interference* account. The expression of learning account suggests that contextual learning survives increases in WM load, but performing a concurrent WM task impairs the expression of that learned information. The novel-array interference account predicts that interspersing repeat and novel arrays during the initial training phase contributes to an increase in overall task demands, and thus, contextual cueing does not survive increases in WM load. To distinguish between these two explanations, we combined the methods of Experiments 1 and 2. We included novel and repeat search arrays in the initial training phase (as per Experiment 1), during which participants performed the search task under no-, low-, or high-WM load. In addition, we included a subsequent testing phase in which participants performed the search task without any concurrent WM task (as per Experiment 2). We hypothesized that contextual cueing would be reduced in both the training and testing phases, when repeat arrays were learned under a WM-load in which repeat and novel arrays were interleaved.

Method

Participants. Twenty-four participants (12 men, 12 women; mean age = 23.30 years, $SD = 3.56$) from the psychology research participation scheme at The University of Queensland volunteered for the experiment. All participants reported normal or corrected-to-normal vision and were naïve to the purpose of the study. Each participant provided written informed consent, and each was paid \$15 for his or her time.

Procedure and design.

Training phase. The training phase in Experiment 3 was identical to the testing phase in Experiment 1.

Testing phase. The testing phase in Experiment 3 was similar to the testing phase in Experiment 2, except that participants performed four blocks of 24 search-only trials. As in Experiment 2, half of the search arrays in the testing phase were repeat search arrays (from the training phase), whereas the other half were previously unseen, randomly generated search arrays.

Explicit search-array recognition test. The explicit search-array recognition test was identical to that used in Experiment 1.

Results and Discussion

Data from three participants were excluded from the analysis because their accuracy on the high-WM task (2 participants) or the

search task (1 participant) was less than 60%. Data from the remaining 21 participants were included in the group analysis.

Accuracy in the WM task was well above chance (50%) for both the low-load WM task (90.36%), $t(20) = 23.18$, $p < .001$, and the high-load WM task (83.96%), $t(20) = 15.25$, $p < .001$. A paired-samples t test revealed that accuracy was significantly higher in the low-load condition than in the high-load condition, $t(20) = 7.11$, $p < .001$, indicating that the high-load WM task was more difficult than the low-load WM task. Accuracy on the search task was also well above chance (50%) in the no-load (98.10%), $t(20) = 103.33$, $p < .001$, low-load (98.93%), $t(20) = 150.08$, $p < .001$, and high-load conditions (98.57%), $t(20) = 129.06$, $p < .001$. Search accuracy differed between these conditions, $F(2, 40) = 4.45$, $p = .018$. Participants were significantly more accurate on a search performed under low-WM load than one performed under a no-WM load, $t(20) = 2.97$, $p = .008$. Search accuracy performed under a high-WM load was not significantly different from search accuracy performed under a no-WM load, $t(20) = 1.44$, $p = .166$, or low-WM load, $t(20) = 1.64$, $p = .117$. Filtering outliers with the same procedure as that described in Experiment 1 resulted in the removal of 2% of trials.

Mean RTs in the training phase were subjected to a 3 (load) $\times$ 2 (context) $\times$ 10 (epoch) repeated-measures ANOVA (see Figure 4). Results revealed a significant main effect of context, $F(1, 20) = 29.73$, $p < .001$. The mean RT for repeat search arrays was on average 136 ms faster than the mean RT for novel search arrays, indicating that on average, participants learned the spatial configurations in the repeat search arrays. There was also a significant main effect of epoch, $F(9, 180) = 26.06$, $p < .001$. The mean RT was 416 ms faster in Epoch 10 than in Epoch 1, indicating a reliable practice effect across the task. The main effect of load did not reach significance, $F(2, 40) = 2.15$, $p = .130$, nor did the interaction between context and epoch, $F(9, 180) = 1.78$, $p = .075$. All other interaction terms failed to reach significance (F s < 1.22 and p s $> .249$).

Overall accuracy in the testing phase was 98.29%, which was significantly above chance (50%), $t(20) = 116.00$, $p < .001$. Data from the testing phase was filtered for outliers using the same procedure as that described in Experiment 1, resulting in the removal of 2% of trials. The mean RT from the testing phase was subjected to a 4 (learning condition) $\times$ 4 (block) repeated-measures ANOVA. Results revealed a significant effect of learning condition, $F(3, 60) = 3.73$, $p = .016$ (see Figure 5). This main

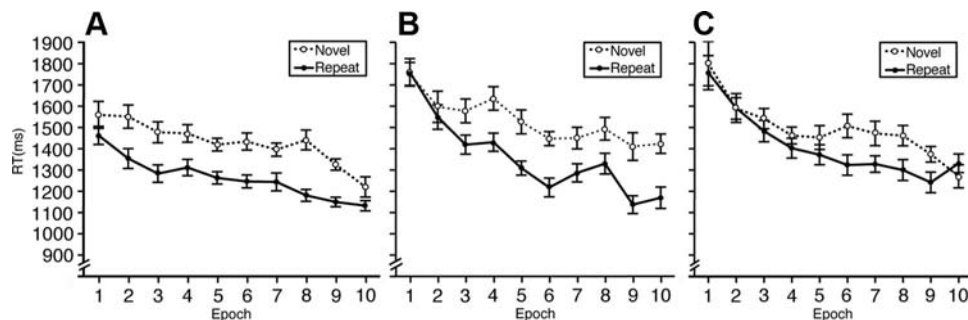


Figure 4. Mean reaction times (RT) for visual search over epoch in the no-load (A), low-load (B), and high-load (C) conditions in Experiment 3. Error bars represent 1 standard error of the mean.

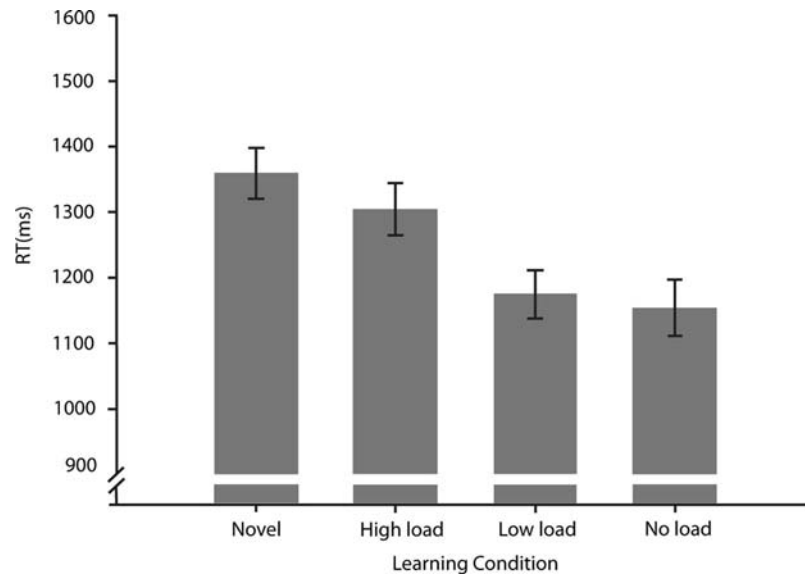


Figure 5. Mean reaction times (RTs) for visual search in the testing phase as a function of learning condition. Error bars represent 1 standard error of the mean.

effect was followed-up with Bonferroni corrected t tests for three comparisons (adjusted $p = .0167$). As expected, mean RTs for search arrays learned under a no-WM load ($M = 1,152.86$, $SD = 253.89$) were significantly faster than mean RTs for novel arrays ($M = 1,357.75$, $SD = 366.31$), $t'(20) = 2.91$, $p = .009$. Mean RTs for search arrays learned under a low-WM load ($M = 1,173.72$, $SD = 286.39$) were also significantly faster than mean RTs for novel search arrays, $t'(20) = 3.24$, $p = .004$. Interestingly, there was no significant difference between mean RTs for search arrays learned under a high-WM load ($M = 1,303.06$, $SD = 354.29$) and novel search arrays, $t'(20) = 0.88$, $p = .388$, indicating that participants were unable to learn contextual information when learning occurred under high-WM load.

The results from the training phase may appear somewhat inconsistent with those from the testing phase. In the training phase, we did not find a significant interaction between load and context, suggesting that increases in WM load had no effect on the contextual cueing effect. In contrast, in the testing phase, we found that performance on the search task differed across learning conditions. To reconcile these seemingly inconsistent findings we conducted three follow-up tests (with Bonferroni corrections) on the data from the training phase. Results revealed that the mean RT in the no-WM load condition was significantly faster for repeat arrays ($M = 1,265$, $SD = 321$) than for novel arrays ($M = 1,423$, $SD = 262$), $t'(20) = 3.35$, $p = .003$. Likewise, in the low-WM load condition mean RT was significantly faster for repeat arrays ($M = 1,326$, $SD = 351$) than for novel arrays ($M = 1,500$, $SD = 342$), $t'(20) = 4.09$, $p = .001$. Critically, however, mean RT in the high-WM load condition did not differ significantly between repeat ($M = 1,409$, $SD = 405$) and novel arrays ($M = 1,487$, $SD = 333$), $t'(20) = 1.45$, $p = .164$. Thus, the results of the training and testing phases are indeed comparable; both results indicate that increases in WM load can impair learning processes involved in contextual cueing.

To ensure that the findings from the testing phase were not due to learning during testing, we conducted further t tests using a Bonferroni correction (adjusted $p = .008$), to investigate performance across blocks in the testing phase. Results revealed that the RT in Block 1 was significantly faster than the RT in Block 4, $t'(20) = 3.25$, $p = .004$, with all other comparisons failing to reach significance (see Table 2 for means and standard deviations). Thus, the learning effect observed in the testing phase was due to learning that occurred during the training phase and could not be attributed to rapid learning during the testing phase.

Overall accuracy on the explicit search-array recognition test was 52%, which did not differ from chance (50%), $t(20) = 1.02$, $p = .321$. This finding was consistent across blocks. In Block 1 accuracy was 50%, $t(20) = 0.18$, $p = .863$; in Block 2, accuracy was 49%, $t(20) = 0.32$, $p = .750$; in Block 3, accuracy was 53%, $t(20) = 1.37$, $p = .186$; and in Block 4, accuracy was 55%, $t(20) = 1.72$, $p = .100$. Again, these results suggest that much of the learning effect is implicit.

Experiment 3 was designed to distinguish between the *expression of learning* and the *novel-array interference* accounts of the

Table 2
Mean Reaction Times (and Standard Deviations) in the Testing Phase of Experiment 3

Learning condition	Block			
	1	2	3	4
Novel	1,267 (387)	1,236 (389)	1,460 (550)	1,356 (435)
No-WM load	1,074 (233)	1,192 (380)	1,140 (275)	1,198 (414)
Low-WM load	1,075 (334)	1,097 (359)	1,237 (357)	1,294 (552)
High-WM load	1,279 (400)	1,288 (378)	1,218 (304)	1,425 (519)

Note. WM = working-memory.

apparently discrepant results of Experiments 1 and 2. In Experiment 3, contextual learning was influenced by WM load in both the dual-task training phase and the single-task testing phases. These findings do not support the expression of learning account and instead suggest that the novel-array interference hypothesis provides a better explanation for the different findings of the first two experiments.

It should be noted that although Experiment 1 and the training phase in Experiment 3 were methodologically identical, the outcomes for the two experiments differed. In Experiment 1, we found no evidence for contextual cueing under low-WM load, but we found a contextual cueing effect under low-WM load in Experiment 3. Participants in Experiment 1 were 1st year undergraduate students, whereas participants in Experiment 3 were from a psychology research participation scheme. It may be the case that individuals from the latter group had more experience participating in experiments than did 1st year undergraduate students. Nevertheless, the fact that we found no differences in performances between repeat and novel arrays under high-load conditions in Experiment 3 is clear evidence that contextual cueing and WM draw on common resources.

General Discussion

The current study had two aims: To investigate whether contextual cueing is attenuated when spatial WM resources are exhausted and, if so, at what level of processing WM load has its effect. To address the first aim, we employed a paradigm similar to that of Vickery et al. (2010), in which participants undertook a spatial contextual cueing paradigm while performing a WM task. We addressed concerns about the adequacy of the WM load manipulation in the study of Vickery et al. by increasing the number of items in the WM task and by presenting probe dots sequentially during initial encoding and subsequent matching. To address the second aim, we separated learning from the expression of learning by testing context effects in visual search under dual-task (Experiment 1 and Experiment 3) and single-task conditions (Experiment 2 and Experiment 3).

In Experiment 1, novel and repeat search arrays were interleaved in the training phase, and contextual learning was tested under dual-task conditions. In these trials, we found no evidence of contextual learning for arrays searched under load. In contrast, in Experiment 2, training consisted entirely of repeat search arrays, and contextual learning was tested under single-task conditions. Here, we found that contextual cueing was not influenced by WM load.

The methods used in Experiment 1 and Experiment 2 differed in two ways. First, contextual cueing was probed under dual-task conditions in Experiment 1 and under single-task conditions in Experiment 2. Second, novel and repeat arrays were interleaved in Experiment 1, whereas only repeat arrays were presented in the training phase of Experiment 2. We found that the dual-task condition interfered with contextual cueing in Experiment 1, but not in Experiment 2. Due to the methodological issues described above, however, our results did not allow us to differentiate between accounts based on dual-task interference with the *expression* of contextual learning and an increase in task difficulty brought about by the interleaving of novel and repeat arrays.

Experiment 3 was designed to distinguish between these possibilities. The task involved an initial training phase, in which novel and repeat arrays were interleaved throughout the course of the experiment. During this training phase, participants performed the search task with a concurrent WM task. The training phase was followed by a testing phase, in which participants performed the search task without any concurrent WM task. Thus, we were able to test whether contextual learning was disrupted under both single- and dual-task conditions. We found that contextual learning was reduced under high-WM load in both the training and the testing phases. This finding casts doubt on the expression of learning account. Instead, the results of Experiment 3 indicate that the inclusion of repeat and novel arrays during training increased task difficulty and, when combined with a difficult WM task, reduced the contextual cueing effect.

Past research suggests that WM plays a key role in the active ignoring of task-irrelevant distractor information (Conway, Cowan, & Bunting, 2001; Conway & Engle, 1994; Conway, Tuholski, Shisler, & Engle, 1999; Hasher & Zacks, 1988; Kane, Bleckley, Conway, & Engle, 2001; Rosen & Engle, 1997). Recently, Vogel, McCollough, and Machizawa (2005) used event-related potentials (ERPs) to show that participants with high spatial-WM capacity are better able to filter out irrelevant objects from WM than are participants with low-WM capacity. Participants were required to remember the orientations of one set of colored bars (e.g., red), while ignoring another set of colored bars (e.g., blue). Individuals with high-WM capacity showed similar neural responses when remembering two relevant bars, whether these bars were presented with distractor bars or without distractor bars, suggesting that high-WM individuals were able to filter out the irrelevant bars. By contrast, low-WM capacity individuals showed similar neural responses when remembering four relevant bars and when remembering two relevant bars presented with two distracting bars, suggesting that low-WM individuals attended to both the relevant and the irrelevant bars.

Findings from the above-mentioned studies suggest that WM capacity is important when there is an explicit goal to attend to one set of stimuli and ignore another set. Our findings extend this literature in two ways. First, our findings demonstrate that WM is important even in the absence of explicit instructions or goals that require an individual to focus resources on one stimulus type and ignore another. That is, even when participants were not given explicit instructions to remember the repeat arrays and ignore the novel ones, increases in WM load impaired learning processes involved in contextual cueing. Second, our results extend previous findings, which suggest that WM capacity is important in filtering out meaningless distractor information, by suggesting that WM is also involved in extracting meaningful information from distractors.

It is important to note that WM is multifaceted. We have shown that an additional spatial-WM task can impair contextual learning. WM, however, consists of independent subsystems not only for visuospatial processing but also for verbal and episodic information (Baddeley & Hitch, 1974; Baddeley, 2000). Thus, it remains to be seen whether these latter operations also tap the same resources as those involved in contextual cueing. In addition, investigators have further hypothesized that WM includes the processes of active maintenance, the retrieval of information (Unsworth & Engle, 2007a, 2007b) and, as outlined above, the sup-

pression of distractors (Friedman & Miyake, 2004; Hasher, Lustig, & Zacks, 2007). The present data suggest that spatial-WM load influences at least the encoding, and perhaps even the retrieval, of context information, but further work is needed to understand how distinct stages of WM processing map onto the operations involved in contextual cueing. Indeed, repeated contexts have been shown to enhance not only the guidance of visual attention but also the threshold for response-selection decisions regarding the target (Kunar et al., 2008). It may be that WM resources are involved in one or both of these processes, suggesting that this is fertile ground for future research.

Conclusion

This study employed a spatial WM task with a visual search task to examine whether spatial WM load modulates spatial contextual cueing. We found that increases in WM load significantly impaired contextual learning. These findings indicate that spatial contextual cueing and WM rely on common underlying mechanisms.

References

- Atkinson, R. C., & Shiffrin, R. M. (1968). Human memory: A proposed system and its control processes. In K. W. Spence (Ed.), *The psychology of learning and motivation: Advances in research and theory* (Vol. 2, pp. 89–195). New York, NY: Academic Press. doi:10.1016/S0079-7421(08)60422-3
- Baddeley, A. (2000). The episodic buffer: A new component for working memory? *Trends in Cognitive Sciences*, 4, 417–423. doi:10.1016/S1364-6613(00)01538-2
- Baddeley, A. D., & Hitch, G. (1974). Working memory. In G. H. Bower (Ed.), *The psychology of learning and motivation: Advances in research and theory* (Vol. 8, pp. 47–89). New York, NY: Academic Press. doi:10.1016/S0079-7421(08)60452-1
- Biederman, I., Mezzanotte, R. J., & Rabinowitz, J. C. (1982). Scene perception: Detecting and judging objects undergoing relational violations. *Cognitive Psychology*, 14, 143–177. doi:10.1016/0010-0285(82)90007-X
- Chun, M. M., & Jiang, Y. (1998). Contextual cueing: Implicit learning and memory of visual context guides spatial attention. *Cognitive Psychology*, 36, 28–71. doi:10.1006/cogp.1998.0681
- Chun, M. M., & Phelps, E. A. (1999). Memory deficits for implicit contextual information in amnesic subjects with hippocampal damage. *Nature Neuroscience*, 2, 844–847. doi:10.1038/12222
- Cohen, N. J., & Eichenbaum, H. (1993). *Memory, amnesia, and the hippocampal system*. Cambridge, MA: MIT Press.
- Conway, A. R. A., Cowan, N., & Bunting, M. F. (2001). The cocktail party phenomenon revisited: The importance of working memory capacity. *Psychonomic Bulletin & Review*, 8, 331–335. doi:10.3758/BF03196169
- Conway, A. R. A., & Engle, R. W. (1994). Working memory and retrieval: A resource-dependent inhibition model. *Journal of Experimental Psychology: General*, 123, 354–373. doi:10.1037/0096-3445.123.4.354
- Conway, A. R. A., Tuholski, S. W., Shisler, R. J., & Engle, R. W. (1999). The effect of memory load on negative priming: An individual differences investigation. *Memory & Cognition*, 27, 1042–1050. doi:10.3758/BF03201233
- Cowan, N. (2001). The magical number 4 in short-term memory: A reconsideration of mental storage capacity. *Behavioral and Brain Sciences*, 24, 87–145. doi:10.1017/S0140525X01003922
- Deroost, N., Coomans, D., & Soetens, E. (2009). Perceptual load improves the expression but not learning of relevant sequence information. *Experimental Psychology*, 56, 84–91. doi:10.1027/1618-3169.56.2.84
- Eckstein, M. P., Drescher, B. A., & Shimozaaki, S. S. (2006). Attentional cues in real scenes, saccadic targeting, and Bayesian priors. *Psychological Science*, 17, 973–980. doi:10.1111/j.1467-9280.2006.01815.x
- Eichenbaum, H., & Cohen, N. J. (2001). *From conditioning to conscious recollection: Memory systems of the brain*. New York, NY: Oxford University Press. doi:10.1093/acprof:oso/9780195178043.001.0001
- Feldman, J. (1997). Curvilinearity, covariance, and regularity in perceptual groups. *Vision Research*, 37, 2835–2848. doi:10.1016/S0042-6989(97)00096-5
- Feldman, J. (1999). The role of objects in perceptual grouping. *Acta Psychologica*, 102, 137–163. doi:10.1016/S0001-6918(98)00054-7
- Frensch, P. A., Lin, J., & Buchner, A. (1998). Learning versus behavioral expression of the learned: The effects of a secondary tone-counting task on implicit learning in the serial reaction task. *Psychological Research*, 61, 83–98. doi:10.1007/s004260050015
- Friedman, N. P., & Miyake, A. (2004). The relations among inhibition and interference control function: A latent-variable analysis. *Journal of Experimental Psychology: General*, 133, 101–135. doi:10.1037/0096-3445.133.1.101
- Greene, A. J., Gross, W. L., Elsinger, C. L., & Rao, S. M. (2007). Hippocampal differentiation without recognition: An fMRI analysis of the contextual cueing task. *Learning & Memory*, 14, 548–553. doi:10.1101/lm.609807
- Hasher, L., Lustig, C., & Zacks, R. T. (2007). Inhibitory mechanisms and the control of attention. In A. R. A. Conway, C. Jarrold, M. J. Kane, A. Miyake, & J. N. Towse (Eds.), *Variation in working memory* (pp. 227–249). New York, NY: Oxford University Press. doi:10.1093/acprof:oso/9780195168648.003.0009
- Hasher, L., & Zacks, R. T. (1988). Working memory, comprehension, and aging: A review and a new view. In G. H. Bower (Ed.), *The psychology of learning and motivation: Advances in research and theory* (Vol. 22, pp. 193–225). San Diego, CA: Academic Press. doi:10.1016/S0079-7421(08)60041-9
- Henderson, L. (1972). Spatial and verbal codes and the capacity of STM. *Quarterly Journal of Experimental Psychology*, 24, 485–495. doi:10.1080/14640747208400308
- Hirsh, R. (1974). The hippocampus and contextual retrieval of information from memory: A theory. *Behavioral Biology*, 12, 421–444. doi:10.1016/S0091-6773(74)92231-7
- Hoffmann, J., & Kunde, W. (1999). Location-specific target expectancies in visual search. *Journal of Experimental Psychology: Human Perception and Performance*, 25, 1127–1141. doi:10.1037/0096-1523.25.4.1127
- Hollingworth, A., & Henderson, J. (1998). Does consistent scene context facilitate object perception? *Journal of Experimental Psychology: General*, 127, 398–415. doi:10.1037/0096-3445.127.4.398
- Jiang, Y., & Leung, A. W. (2005). Implicit learning of ignored visual context. *Psychonomic Bulletin & Review*, 12, 100–106. doi:10.3758/BF03196353
- Kane, M. J., Bleckley, M. K., Conway, A. R. A., & Engle, R. W. (2001). A controlled-attention view of working memory capacity. *Journal of Experimental Psychology: General*, 130, 169–183. doi:10.1037/0096-3445.130.2.169
- Kunar, M. A., Flusberg, S. J., & Wolfe, J. M. (2008). Time to guide: Evidence for delayed attentional guidance in contextual cueing. *Visual Cognition*, 16, 804–825. doi:10.1080/13506280701751224
- Kunde, W., & Hoffmann, J. (2005). Selecting spatial frames of reference for visual target localization. *Experimental Psychology*, 52, 201–212. doi:10.1027/1618-3169.52.3.201
- Luck, S. J., & Vogel, E. K. (1997). The capacity of visual working memory for features and conjunctions. *Nature*, 390, 279–281. doi:10.1038/36846
- Nissen, M. J., & Bullemer, P. (1987). Attentional requirements of learning—Evidence from performance measures. *Cognitive Psychology*, 19, 1–32. doi:10.1016/0010-0285(87)90002-8

- O'Keefe, J., & Nadel, L. (1978). *The hippocampus as a cognitive map*. Oxford, England: Clarendon Press.
- Palmer, S. E. (1975). The effects of contextual scenes on the identification of objects. *Memory & Cognition*, 3, 519–526. doi:10.3758/BF03197524
- Patterson, M. D., Bly, B. M., Porcelli, A. J., & Rypma, B. (2007). Visual working memory for global, object, and part-based information. *Memory & Cognition*, 35, 738–751. doi:10.3758/BF03193311
- Posner, M. I. (1969). Abstraction and the process of recognition. In G. H. Bower & J. T. Spence (Eds.), *Psychology of learning and motivation* (Vol. 3, pp. 43–100). New York, NY: Academic Press. doi:10.1016/S0079-7421(08)60397-7
- Remillard, G. (2009). Pure perceptual-based sequence learning: A role of visuospatial attention. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 35, 528–541. doi:10.1037/a0014646
- Rosen, V. M., & Engle, R. W. (1997). The role of working memory capacity in retrieval. *Journal of Experimental Psychology: General*, 126, 211–227. doi:10.1037/0096-3445.126.3.211
- Schacter, D. L., Wagner, A. D., & Buckner, R. L. (2000). Memory systems of 1999. In E. Tulving & F. I. M. Craik (Eds.), *Oxford handbook of memory* (pp. 627–643). New York, NY: Oxford University Press.
- Shanks, D., Rowland, L., & Ranger, M. (2005). Attentional load and implicit sequence learning. *Psychological Research*, 69, 369–382. doi:10.1007/s00426-004-0211-8
- Smyth, A. C., & Shanks, D. R. (2008). Awareness in contextual cuing with extended and concurrent explicit tests. *Memory & Cognition*, 36, 403–415. doi:10.3758/MC.36.2.403
- Sperling, G. (1960). The information available in brief visual presentations. *Psychological Monographs: General and Applied*, 74, 1–30. doi:10.1037/h0093759
- Squire, L. R. (1992). Memory and the hippocampus: A synthesis from findings with rats, monkeys, and humans. *Psychological Review*, 99, 195–231. doi:10.1037/0033-295X.99.2.195
- Squire, L. R., Stark, C. E. L., & Clark, R. E. (2004). The medial temporal lobe. *Annual Review of Neuroscience*, 27, 279–306. doi:10.1146/annurev.neuro.27.070203.144130
- Tulving, E. (1985). How many memory systems are there? *American Psychologist*, 40, 385–398. doi:10.1037/0003-066X.40.4.385
- Unsworth, N., & Engle, R. W. (2007a). The nature of individual differences in working memory capacity: Active maintenance in primary memory and controlled search from secondary memory. *Psychological Review*, 114, 104–132. doi:10.1037/0033-295X.114.1.104
- Unsworth, N., & Engle, R. W. (2007b). On the division of short-term and working memory: An examination of simple and complex spans and their relation to their higher order abilities. *Psychological Bulletin*, 133, 1038–1066. doi:10.1037/0033-2909.133.6.1038
- Vickery, T. J., Sussman, R. S., & Jiang, Y. V. (2010). Spatial context learning survives interference from working memory load. *Journal of Experimental Psychology: Human Perception and Performance*, 36, 1358–1371. doi:10.1037/a0020558
- Vogel, E. K., McCollough, A. W., & Machizawa, M. G. (2005). Neural measures reveal individual differences in controlling access to working memory. *Nature*, 438, 500–503. doi:10.1038/nature04171
- Vogel, E. K., Woodman, G. F., & Luck, S. J. (2001). Storage of features, conjunctions and objects in visual working memory. *Journal of Experimental Psychology: Human Perception and Performance*, 27, 92–114. doi:10.1037/0096-1523.27.1.92
- Woodman, G. F., & Luck, S. J. (2004). Visual search is slowed when visuospatial working memory is occupied. *Psychonomic Bulletin & Review*, 11, 269–274. doi:10.3758/BF03196569

Received September 28, 2011

Revision received March 30, 2012

Accepted April 13, 2012 ■