

Smooth Rotation of 2-D and 3-D Representations of Terrain: An Investigation Into the Utility of Visual Momentum

Justin G. Hollands, Nada J. Pavlovic, Yukari Enomoto, and Haiying Jiang, Defence Research and Development Canada, Toronto, Canada

Objective: The potential advantage of visual momentum in the form of smooth rotation between two-dimensional (2-D) and three-dimensional (3-D) displays of geographic terrain was examined. **Background:** The relative effectiveness of 2-D and 3-D displays is task dependent, leading to the need for multiple frames of reference as users switch tasks. The use of smooth rotation to provide visual momentum has received little scrutiny in the task-switching context. A cognitive model of the processes involved in switching viewpoints on a set of spatial elements is proposed. **Methods:** In three experiments, participants judged the properties of two points placed on terrain depicted as 2-D or 3-D displays. Participants indicated whether Point A was higher than Point B, or whether Point B could be seen from Point A. Participants performed the two tasks in pairs of trials, switching tasks and displays within the pair. In the continuous transition condition the display dynamically rotated in depth from one display format to the other. In the discrete condition there was an instantaneous viewpoint shift that varied across experiments (Experiment 1: immediate; Experiment 2: delay; Experiment 3: preview). **Results:** Performance after continuous transition was superior to that after discrete transition. **Conclusion:** The visual momentum provided by smooth rotation helped users switch tasks. **Application:** The use of dynamic transition is recommended when observers examine multiple views of terrain over time. The model may serve as a useful heuristic for designers. The results are pertinent to command and control, geological engineering, urban planning, and imagery analysis domains.

INTRODUCTION

The concept of *visual momentum* originated with Hochberg and Brooks (1978), who described film-cutting techniques designed to help an audience maintain spatial understanding of a single scene from different viewpoints. Woods (1984) extended the visual momentum concept to user-computer interaction and defined it in that context as the user's ability to extract and integrate data from multiple, consecutive display windows. A variety of visual momentum techniques have been proposed, including the use of consistent representations, graceful transitions, highlighted anchors, and the continuous display of a world-centered reference map (see Wickens & Hollands, 2000, for a summary).

Some of these approaches to improving visual momentum have been implemented (e.g., Andre,

Wickens, Moorman, & Boschelli, 1991; Aretz, 1991; Bederson et al., 1998; Roth, Chuah, Kerpeljiev, Kolojchick, & Lucas, 1997) and examined empirically (Aretz, 1991; Neale, 1996, 1997; Olmos, Liang, & Wickens, 1997; Olmos, Wickens, & Chudy, 2000; St. John, Smallman, and Cowen, 2002). The use of smooth rotation or animation from one viewpoint to another seems a relatively intuitive method for providing a graceful transition. However, this technique has received scant attention in the literature.

We were particularly interested in the problem of the depiction of geographic terrain. This is an important component of battlespace visualization for command and control (Barnes, 2003) and other work domains, such as geological engineering, urban planning, landscape architecture, and aviation. The range of tasks necessary to perform work in these domains is diverse; sometimes specific

Report Documentation Page				Form Approved OMB No. 0704-0188	
Public reporting burden for the collection of information is estimated to average 1 hour per response, including the time for reviewing instructions, searching existing data sources, gathering and maintaining the data needed, and completing and reviewing the collection of information. Send comments regarding this burden estimate or any other aspect of this collection of information, including suggestions for reducing this burden, to Washington Headquarters Services, Directorate for Information Operations and Reports, 1215 Jefferson Davis Highway, Suite 1204, Arlington VA 22202-4302. Respondents should be aware that notwithstanding any other provision of law, no person shall be subject to a penalty for failing to comply with a collection of information if it does not display a currently valid OMB control number.					
1. REPORT DATE FEB 2008		2. REPORT TYPE		3. DATES COVERED 00-00-2008 to 00-00-2008	
4. TITLE AND SUBTITLE Smooth Rotation of 2-D and 3-D Representations of Terrain: An Investigation Into the Utility of Visual Momentum				5a. CONTRACT NUMBER	
				5b. GRANT NUMBER	
				5c. PROGRAM ELEMENT NUMBER	
6. AUTHOR(S)				5d. PROJECT NUMBER	
				5e. TASK NUMBER	
				5f. WORK UNIT NUMBER	
7. PERFORMING ORGANIZATION NAME(S) AND ADDRESS(ES) Defence R&D Canada Toronto, Human Systems Integration Section, 1133 Sheppard Avenue West, Toronto, Ontario, Canada M3M 3B9,				8. PERFORMING ORGANIZATION REPORT NUMBER	
9. SPONSORING/MONITORING AGENCY NAME(S) AND ADDRESS(ES)				10. SPONSOR/MONITOR'S ACRONYM(S)	
				11. SPONSOR/MONITOR'S REPORT NUMBER(S)	
12. DISTRIBUTION/AVAILABILITY STATEMENT Approved for public release; distribution unlimited					
13. SUPPLEMENTARY NOTES					
14. ABSTRACT					
15. SUBJECT TERMS					
16. SECURITY CLASSIFICATION OF:			17. LIMITATION OF ABSTRACT Same as Report (SAR)	18. NUMBER OF PAGES 15	19a. NAME OF RESPONSIBLE PERSON
a. REPORT unclassified	b. ABSTRACT unclassified	c. THIS PAGE unclassified			

judgments of relative position are necessary, and at other times users need to get a sense of the overall shape of the terrain.

The relative effectiveness of 2-D and 3-D displays of geographic terrain depends on the judgment task (for summaries, see St. John, Cowen, Smallman, & Oonk, 2001, or Wickens & Hollands, 2000). Whereas 2-D renderings are generally useful for judging relative position because the normal viewing angles minimize distortion, the advantage of 3-D views is in shape and layout understanding because they integrate all three dimensions, and allow for features otherwise invisible in a 2-D view to be depicted (St. John et al., 2001; Wickens & Thomas, 2000). This implies that to perform a variety of tasks, the observer will need both 2-D and 3-D views.

In many contexts, civilian and military, an observer needs to switch tasks frequently while viewing geospatial information. Whereas the display can be changed to fit the task at hand, the mental transition from one display to another may still be difficult. Abruptly changing frames of reference (switching a view from 2-D to 3-D or vice versa) may cause disorientation. However, a gradual transition between 2-D and 3-D views incorporating animation of viewpoint during task switching should provide visual momentum and alleviate the mental transition problem.

St. John et al. (2002) evaluated several methods for combining 2-D and 3-D views on terrain, including smooth rotation from one view to the other (which they called "time morph"). They used an antenna placement task, which has both general shape understanding and specific relative position components. In the time morph condition, participants pressed and held the space bar to initiate the rotation from 2-D to 3-D; releasing the space bar produced the opposite rotation. This occurred within a single trial. Although participants using the time morph showed faster solution times and exceeded a time limit less frequently than in other conditions (e.g., overlays and side-by-side arrangements), the differences failed to reach conventional significance levels.

In the St. John et al. (2002) study, participants performed one task that had different components. In contrast, we were interested in the question of whether smooth transition aided observers as they switched tasks. To examine this question, we used two tasks developed by St. John et al. (2001). A shape-understanding task required judgment of

whether one ground location was visible from another (the *A-See-B* task), and a relative position task required a judgment of which one of two points was of higher altitude (the *A-High-B* task). St. John et al. (2001) found that the *A-See-B* task was performed better with a 3-D display, whereas the *A-High-B* task was performed better with a 2-D topographic map.

A Cognitive Model

We propose that when an observer is shown, in sequence, a pair of displays that depict the same spatial elements from different viewpoints, some element characteristics (e.g., position, distance, height) are retained from the first display and can be used for a judgment with the second display. Thus, when the first display is shown, a cognitive representation of the spatial elements is formed. This representation should assist in the performance of a spatial judgment using the second display. Such representations are similar in some respects to those proposed in theories of spatial reasoning (e.g., Gattis, 2004; Hörnig, Oberauer, & Weidenfeld, 2005).

If there is a discrete change in viewpoint, the observer must determine correspondences between the spatial elements in the cognitive representation and the elements in the second display, which may require mental rotation (Shepard & Metzler, 1971) or similar transformational processes. These processes will require time and may not always be accurate. In contrast, when there is a continuous rotation of viewpoint on the geospatial scene, the cognitive representation is continually updated based upon the visual representation presented on the display.

Thus, one should see a performance advantage with the continuous rotation between viewpoints relative to a discrete shift. By improved performance, we are predicting shortened response times (RTs) or increased accuracy (without a speed-accuracy trade-off) or both shortened RTs and increased accuracy.

Continuous rotation requires time to depict. One is left, therefore, with the problem of how to construct a discrete equivalent. Three approaches seem reasonable. In the first, the observer is shown the second display immediately. In the second approach, the terrain is removed after the first judgment and the observer waits for the second display. The delay is set equal to the time required for the rotation. The time available to prepare (preparation

time) is thus equivalent. In the third approach, the observer is shown continuously moving terrain during the delay, but there is still a discrete shift to the final position.

Each of these approaches maps relatively easily to a realistic scenario. The “immediate” scenario could occur when the second viewpoint had been previously generated and therefore is available to the operator upon demand. The “delay” scenario could result when an operator must open a file to access the new representation or looks away from the display. The third scenario, “preview,” could occur when an observer examines the old display while waiting on an automated system to present the new display. There are likely other possibilities, but these examples show how the discrete approaches mimic real-world situations.

What predictions does the model make for each approach? First, given the formation of a cognitive representation, one should see improvement in performance for the second judgment regardless of transition type. Second, the model predicts better performance for the continuous transition than the discrete transition in each case.

In the immediate approach, the cognitive processes required to establish correspondences between representations in the discrete condition should require processing time and have some probability of error. With continuous transition, these processes should not be necessary. Moreover, continuous transition provides preparation time not available in the discrete condition. Thus, superior performance is predicted for the continuous condition over the discrete condition post-transition.

With the delay approach, observers should have time to prepare in the discrete condition but must still perform the cognitive processes to establish correspondences between representations. Further, any cognitive representation should be subject to decay during the delay period. One would therefore expect superior performance for the continuous condition over the discrete condition.

Finally, with the preview approach, preparation time is equalized and a visual representation is available in the discrete condition. However, given the sudden shift in viewpoint, there is still a need to perform the cognitive processing to establish correspondences between spatial elements. Thus, one should still expect superior performance in the continuous condition.

It is an important goal for applied research to

show that an advantage for a particular display technique can be demonstrated across a range of realistic scenarios, not simply in one particular case. Thus, we performed three experiments for which the design of the control condition (discrete transition) varied across experiments. In all experiments, participants performed tasks in pairs of trials, switching tasks and displays between trials. In Experiment 1, participants were immediately shown the alternate display format. In Experiment 2, a blank screen was shown for a duration equal to that used for the continuous transition. In Experiment 3 the terrain was shown rotating as in the continuous case, but then “snapped” to the final orientation to allow us to examine the role of preview.

For all experiments, the model predicts that (a) second-trial performance should be superior to first-trial performance and (b) continuous transition should produce better second-trial performance than discrete transition. Further, we predicted that the relative 2-D/3-D advantages observed by St. John et al. (2001) should be replicated in our experiments.

GENERAL METHOD

Participants

All participants had normal or corrected-to-normal vision. They were recruited from Defence Research and Development Canada - Toronto (DRDC Toronto) and the nearby community. Participants were financially compensated for their participation.

Stimuli and Apparatus

There are many possible ways to depict geographic terrain in both 2-D and 3-D formats. Rather than developing new 2-D and 3-D displays, we chose to use the same displays and tasks as St. John et al. (2001) so that our results could be compared with at least one set of relevant studies. The 2-D map is probably the most common method for representing terrain, and maps often use color coding to show terrain elevation. The 3-D view at a 45° angle is a commonly used default for showing terrain models of various types. Color coding was not necessary to depict elevation for the 3-D view because height was portrayed explicitly.

Ten different terrain models were created from digital terrain elevation data (DTED) using Creator/TerrainPro (Multigen-Paradigm, 2001a) modeling tools. Each model represented a 13,351- × 11,288-m region of the state of Wyoming. (An

11th terrain model was created for practice trials.) Four pairs of A and B points were chosen from a central 11,600- × 10,600-m area of each terrain model. Point A was picked randomly from this area, and a Point B was found that was at least 500 m different in altitude and at least 2,000 m distant. For two of the pairs, A was higher than B (*A-High-B-Yes* pairs); for the other two A was lower than B (*A-High-B-No* pairs). For one of the *A-High-B-Yes* pairs, Point B could be seen from Point A; for the other, Point B could not be seen from Point A. The same was true for the *A-High-B-No* pairs. The terrain models and pair locations were identical for both transition conditions.

The Vega visual simulation system (MultiGen-Paradigm, 2001b) was used to render each terrain model as a 3-D display, and an example is shown in Figure 1. The 3-D display depicted the terrain model at a viewing angle of 45° with respect to the ground plane. One virtual light source illuminated the 3-D display. The light source was set with a 90° azimuth angle and a 50° elevation angle. The rays emitted from this source were collimated to optical infinity (i.e., they were parallel). The ambient color of the light source was gray, and the diffuse light was white. Points depicted as farther away were no darker on average than points depicted as closer to the observer; that is, shading was not used to portray greater distance from the observer. Figure 1 shows the effects of this lighting on a typical terrain.

MICRODEM (Microcomputer Digital Elevation Models; Guth, 2001) was used to create a 2-D display of each terrain model with colored contour lines (see Figure 2 for an example). The ordering of colors from lowest to highest elevation was dark blue, light blue, green, yellow, red, and magenta. A color key was available on the display to define the color coding.

Points were indicated using a dot for both 2-D and 3-D displays and were labeled A or B. See Figures 1 and 2 for examples.

The experiments were conducted in a room with dimmed lighting to accentuate visibility and contrast. Stimuli were presented on a 53-cm (21-inch) CRT monitor at 1280 × 1024 resolution. Key-strokes and RT data were collected by a Windows NT graphics workstation. Participants sat at a comfortable viewing distance.

Design and Procedure

Each experiment had a $2 \times 2 \times 2 \times 2$ within-subjects design with transition (continuous vs. discrete), display (2-D vs. 3-D), task (*A-See-B* vs. *A-High-B*), and trial (first vs. second trial in each pair) as independent variables. Dependent measures were RT and accuracy (proportion correct).

After reading a brief description of the experiment and signing an informed consent form, participants performed a block of practice trials (with the practice terrain model). In the *A-High-B* task, participants indicated whether Point A was higher

A-See-B

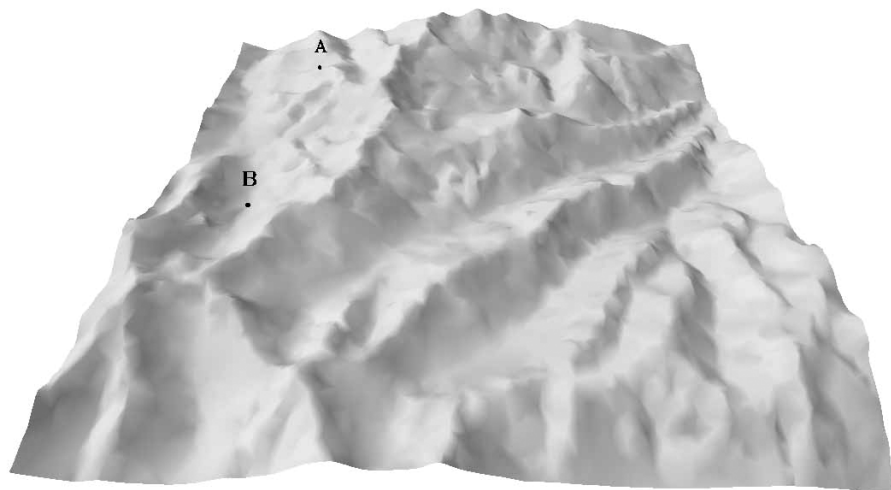


Figure 1. Example of 3-D displays used in Experiments 1 and 2. (A can see B in the figure.)

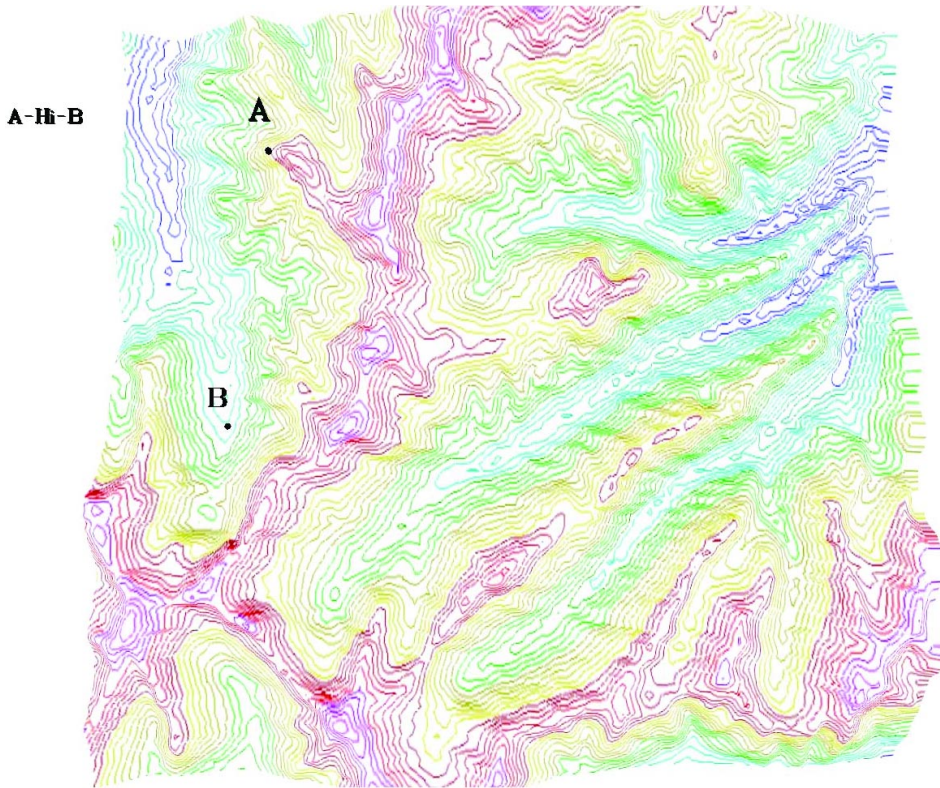


Figure 2. Example of 2-D displays used in all experiments. (A is higher than B in the figure.)

than Point B. In the A-See-B task, participants indicated whether they could see Point B if they were standing at Point A. Feedback (correct or incorrect response, stated verbally by the experimenter) was provided on practice trials but not on experimental trials.

Participants performed the tasks in pairs of trials. The terrain model and A-B points were the same within each trial pair. However, there was a switch in the display type and task across trials in the pair. This led to four possible display-task sequences (2-D-High→3-D-See; 2-D-See→3-D-High; 3-D-High→2-D-See; and 3-D-See→2-D-High).

For each terrain, one of the four A-B pairs was randomly chosen. Four trial pairs were created by using this A-B pair in each of the four display-task sequences. As there were 10 terrains, this produced a block of 40 trial pairs. The process was repeated, randomly choosing from the remaining A-B pairs until four blocks of 40 trial pairs were created. Participants performed the four trial pairs for each terrain consecutively, and the order of the terrains was randomized within blocks. The ordering of

terrain models and trial pairs across blocks was unique for every participant. (There was a slight difference in assignment of A-B pairs to display-task sequences for Experiment 2. This is described in the Experiment 2 Method section.)

This process was followed for continuous and discrete transition conditions. The order of continuous and discrete transition conditions was counterbalanced across participants. Thus there were 160 trial pairs (4 blocks × 40 trial pairs per block) for each transition condition, which meant that each participant completed a total of 320 trial pairs (or 640 trials) during the session.

In the *continuous transition* condition, a smooth rotation of the viewpoint took place. A fade-in/fade-out process occurred prior to the rotation when transitioning from the 2-D display to the 3-D model viewed from above. Fading occurred after the rotation when transitioning from 3-D to 2-D. Shading was added when fading into the 3-D display (and removed when fading out from the 3-D display). The rotation took approximately 3 s. The A-B points were visible during the transition. In the *discrete*

transition condition, the terrain model was shown sequentially, first using the 2-D display and then the 3-D display (or vice versa). Details of the continuous and discrete transitions varied across experiments and are therefore described in the appropriate Method sections.

At the start of each trial, a task prompt (“A-See-B” or “A-Hi-B”) was shown at the same time as the terrain (see Figures 1 and 2). RT was measured from display onset until the participant responded. On the second trial in the continuous condition, the task prompt appeared when the terrain had finished rotating. In this case, RT was measured from the time the task prompt was displayed until the participant responded. The time for any rotation or delay prior to the trial was not included.

Participants initiated each pair of trials by pressing the space bar. The participant’s response on the first trial initiated the transition. For each trial in the pair, the participant responded by pressing a key marked “Y” for *yes* or “N” for *no* (the 1 or 2 key on the numeric keypad), and the participant was asked to respond as quickly and accurately as possible.

Each experimental session took about 90 min to complete. Participants could take breaks between blocks of trials and were debriefed after the session.

EXPERIMENT 1: IMMEDIATE SCENARIO

For Experiment 1 participants were immediately shown the alternate display format on the second trial of a pair in the discrete transition condition.

Method

Participants. We ran 22 participants (12 men and 10 women), aged 18 to 53 years.

Design and procedure. In the continuous condition, the viewpoint was continuously shifted from a position centered directly above the terrain (2-D) until the angle between ground level and the line of sight was 45° (3-D). The viewpoint rotation occurred between the first and second trials. The opposite sequence was used for 3-D to 2-D. The height of the viewpoint above the terrain was constant. In the discrete transition condition, the terrain model was shown sequentially, first using the 2-D display and then the 3-D display (or vice versa). There was no visible delay between the views.

Results

Response time. A mean RT for correct trials was computed for each participant in each condi-

tion. The mean RT data were submitted to a $2 \times 2 \times 2$ within-subjects ANOVA with transition, display, task, and trial serving as independent variables. (Accuracy was generally high, and so there were always sufficient trials to compute a mean RT. Thus, all ANOVAs conducted on RTs were balanced.) In general, we report all effects relating to hypotheses first. Then we report all other significant effects not superseded by higher order interactions.

There was a main effect for trial, $F(1, 21) = 60.63$, $MS_E = 3.328$, $p < .0001$. RTs were shorter on the second trial. There was an interaction between transition and trial, $F(1, 21) = 18.61$, $MS_E = 0.953$, $p < .0005$. Continuous transition produced shorter RTs than did discrete transition for the second trial in a pair (but not the first), as shown in Figure 3. There was also an interaction between display and task, $F(1, 21) = 55.59$, $MS_E = 0.283$, $p < .0001$. RTs were shorter with the 3-D than with the 2-D display, but the difference was greater for the A-See-B task (4.13 s for 2-D vs. 2.96 s for 3-D) than for the A-High-B task (3.28 s for 2-D and 2.97 s for 3-D).

Accuracy. Each trial was scored as correct or incorrect. The proportion of correct trials was computed for each participant in each condition. These data were submitted to a $2 \times 2 \times 2 \times 2$ within-subjects ANOVA. There was a main effect for trial, with participants more accurate on the second trial, $F(1, 21) = 29.37$, $MS_E = 0.0017$, $p < .0001$.

There was a main effect for transition, $F(1, 21) = 6.52$, $MS_E = 0.0025$, $p < .05$. Continuous transition produced greater accuracy than discrete transition, as shown in Figure 3. There was an interaction between transition and display, $F(1, 21) = 5.45$, $MS_E = 0.001$, $p < .05$. The continuous advantage was larger for the 3-D display (mean accuracies were .867 and .845 for continuous and discrete, respectively) than for 2-D (.875 and .869, respectively). There was also an interaction among task, display, and trial order, $F(1, 21) = 6.90$, $MS_E = 0.0016$, $p < .05$. The 3-D display produced greater accuracy for the A-See-B task, but the 2-D display produced greater accuracy for the A-High-B task, especially on the first trial. Mean values are shown in Table 1.

Discussion

As the model predicted, performance was better on the second trial than on the first. Also as predicted, participants were faster and more accurate

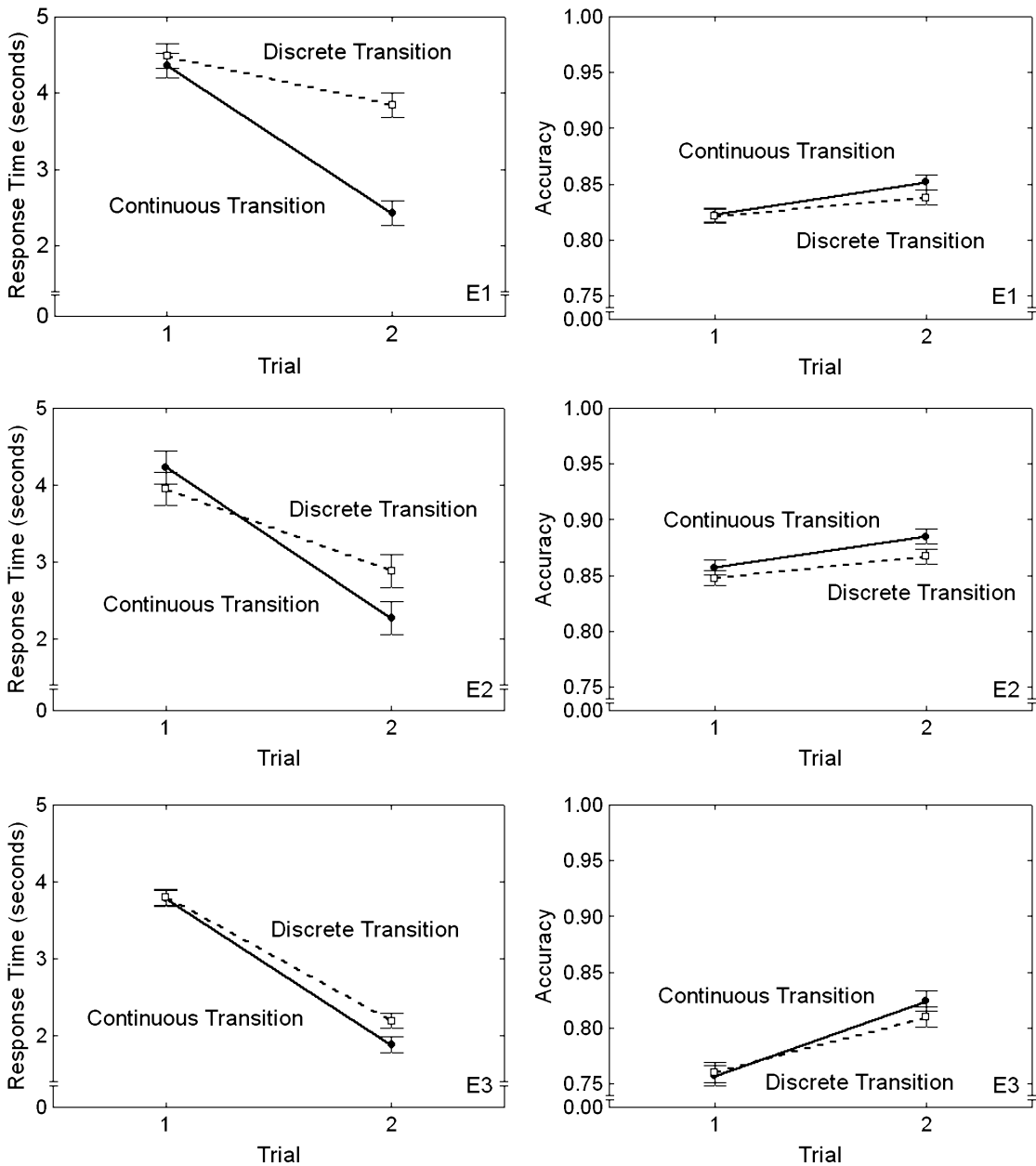


Figure 3. Response time and accuracy as a function of transition and trial for Experiments 1 through 3 (E1, E2, and E3). Error bars indicate 95% confidence intervals based on the within-subjects standard error of the mean in all graphs (Loftus & Masson, 1994).

after the continuous transition than after the discrete transition.

Accuracy was also higher in the continuous transition condition on Trial 1 (before the transition). Participants took longer on Trial 1 in the continuous condition than in the discrete condition. The strategic advantage is unclear, but perhaps the abrupt sequence of immediate display switch-

es in the discrete condition led to a faster pace subjectively, producing generally faster responses at the expense of accuracy.

The accuracy results showed the expected pattern with respect to the relative effectiveness of 2-D and 3-D displays across tasks: Accuracy was greater with the 2-D than the 3-D display for the A-High-B task, but the reverse was true for the

TABLE 1: Experiment 1 – Mean Accuracy as a Function of Task, Display, and Trial

Task	2-D		3-D	
	Trial 1	Trial 2	Trial 1	Trial 2
A-High-B	.966	.978	.885	.948
A-See-B	.769	.775	.790	.802

Note. The within-subjects standard error of the mean was .006.

A-See-B task (greater accuracy with 3-D than 2-D). RTs were also smaller with the 3-D displays for the A-See-B task, consistent with predictions. However, an RT advantage was also found for the 3-D displays in the A-High-B task.

EXPERIMENT 2: DELAY SCENARIO

In Experiment 2, a blank screen was shown for 3 s between the two trials in the discrete transition condition. This duration was approximately equivalent to the animated rotation in the continuous transition condition. With the delay in the discrete condition, observers have time to prepare for the second display but still must perform the cognitive processes to establish correspondences between spatial elements once the second display is shown. One would therefore expect superior performance for the continuous condition post transition.

Participants. We ran 42 participants (22 men and 20 women), aged 19 to 49 years. No participant served in Experiment 1.

Procedure. For Experiment 2, there were eight pairs of A-B points for each terrain model. Two of these pairs were used for each possible response sequence (YY, YN, NY, NN). One pair was used for the 2-D–3-D sequence, and the other was used for the 3-D–2-D sequence, resulting in 16 unique combinations of trial pairs for each terrain model. As there were 10 models, this produced 160 trial pairs.

Results

Response time. RTs were averaged and analyzed as in Experiment 1. There was a main effect for trial, $F(1, 41) = 96.86, MS_E = 2.88, p < .0001$. RTs were shorter on the second trial. There was an interaction between transition and trial, $F(1, 41) = 65.17, MS_E = 1.08, p < .0001$. As Figure 3 shows, continuous transition produced shorter RTs than did discrete transition for the second trial in a pair (but not the first).

There was also an interaction between transition and task, $F(1, 41) = 4.34, MS_E = 0.79, p < .05$. Continuous transition shortened RTs more for the A-See-B task (3.58 s for continuous vs. 4.50 s for discrete) than for the A-High-B task (3.20 s for continuous vs. 3.83 s for discrete). Finally, there was an interaction among display, task, and trial, $F(1, 41) = 7.31, MS_E = 0.67, p < .01$. RTs were shorter for 3-D displays regardless of task, but this difference was greater for the A-See-B task on the first trial of a pair. Mean values are shown in Table 2.

Accuracy. Accuracy scores were computed and analyzed as in Experiment 1. There was a main effect for trial, with participants more accurate on the second trial, $F(1, 41) = 32.46, MS_E = 0.0028, p < .0001$. An interaction between transition and trial just failed to reach conventional significance levels, $F(1, 41) = 3.77, MS_E = 0.0017, p = .059$. As shown in Figure 3, the advantage for continuous

TABLE 2: Experiment 2 – Mean Response Time (in Seconds) as a Function of Task, Display, and Trial

Task	2-D		3-D	
	Trial 1	Trial 2	Trial 1	Trial 2
A-High-B	4.423	3.387	3.727	2.530
A-See-B	5.543	3.820	3.994	2.795

Note. The within-subjects standard error of the mean was 0.089.

transition appears larger for the second trial in a pair. There was an interaction between transition and task, $F(1, 41) = 7.02$, $MS_E = 0.0037$, $p < .05$. Continuous transition increased accuracy for the A-High-B task (.94 for continuous vs. .92 for discrete) but not for the A-See-B task (.73 for continuous vs. .74 for discrete). There was also an interaction among display, task and trial, $F(1, 41) = 8.59$, $MS_E = 0.0018$, $p < .01$. Accuracy was generally greater for 2-D displays, but especially for the A-High-B task on the first trial. Mean values are shown in Table 3.

Discussion

As predicted, second-trial performance was superior to first-trial performance in Experiment 2. Also as predicted, posttransition performance was superior with continuous transition. Participants were faster and there was a trend toward greater accuracy in the continuous condition post-transition. Given the RT advantage without any trade-off in accuracy, it appears that the smooth rotation provided improved visual momentum between consecutive displays. This implies that smooth rotation does more than simply provide preparation time for the subsequent judgment.

Continuous transition produced an RT advantage for both tasks, although it was greater for the A-See-B task, and there was no trade-off with accuracy. That is, accuracy was not lower with continuous transition: Indeed, it was higher for the A-High-B task. Generally the continuous transition was advantageous regardless of task.

The predicted effects of task dependency were only partially supported in Experiment 2. RTs were shorter for 3-D displays in the A-See-B task, and accuracy was greater for 2-D displays in the A-High-B task, consistent with predictions. Further, these effects were larger on the first trial of a pair. One might expect to see a larger effect of display type when observers have not yet seen the same terrain and A-B points in the other display

format. However, RTs were also shorter for 3-D displays in the A-High-B task, and accuracy was greater for 2-D displays in the A-See-B task, contrary to predictions. In combination these results suggest a speed-accuracy trade-off, with 3-D displays producing less accurate but faster processing than 2-D displays, regardless of task. We will further consider the task dependency results for all experiments in the General Discussion.

EXPERIMENT 3: PREVIEW SCENARIO

In the preview scenario, the observer is shown continuously moving terrain during the delay. Thus, preparation time is equalized and a visual representation is available in the discrete condition. However, there is a sudden shift in viewpoint after the preview and just before the second task begins. According to the model, there is a need to perform cognitive processing to establish correspondences in the discrete condition, given the sudden change in viewpoint. Thus, one would still expect superior performance in the continuous condition, in which such processing should not be necessary.

In Experiments 1 and 2, the rotation from 2-D to 3-D (and vice versa) occurred only in depth. In Experiment 3, the terrain was also rotated in the azimuth so that the viewpoint for the 3-D display was aligned with an imaginary line connecting Points A and B. This was done to make the 3-D display more immersive or egocentric (Wickens & Hollands, 2000) and to provide a method for equalizing the rotation time in the continuous and discrete conditions.

Method

In general, Experiment 3 followed the general method, but key exceptions will be noted.

Participants. We ran 24 participants (12 men and 12 women) with normal or corrected-to-normal vision, recruited from DRDC Toronto and the nearby

TABLE 3: Experiment 2 – Mean Accuracy as a Function of Task, Display, and Trial

Task	2-D		3-D	
	Trial 1	Trial 2	Trial 1	Trial 2
A-High-B	.968	.970	.870	.919
A-See-B	.751	.768	.698	.722

Note. The within-subjects standard error of the mean was .005.

community. Participants were financially compensated for their participation. None served in Experiment 1 or 2.

Stimuli and apparatus. The viewpoint direction (azimuth) for the 3-D display was defined by the vector connecting Points A and B on the terrain. The 3-D display was centered with respect to this vector. An example is shown in Figure 4.

Design and procedure. In the continuous transition condition, the following sequence was used in transitioning from 3-D to 2-D (see Figure 5). First, the 3-D terrain was depicted so that the line of sight was aligned with an imaginary vector connecting Points A and B. Then the terrain was shifted left or right (horizontal translation) so that the terrain was centered in the display. Then the terrain was rotated in the azimuth until the side of the terrain nearest the observer corresponded to the bottom of the 2-D map (azimuth rotation). The terrain was then rotated upwards until the viewpoint was centered directly above the terrain, producing a “God’s eye view” (depth rotation). The height of the viewpoint above the terrain center was constant. The opposite sequence was used to transition from 2-D to 3-D. A-B points were visible during the transition.

In the discrete transition condition, the same sequence of transformations was used, with one

exception: Azimuth rotation was in the direction opposite that which occurred in the continuous case. For example, if the azimuth rotation from the A-B vector was 120° counterclockwise in the continuous condition, then it was 120° clockwise in the discrete condition. This sequence is depicted in Figure 5. Upon reaching this position, the display orientation would immediately switch to the azimuth position aligned with the bottom of the 2-D map. Then the horizontal translation occurred, followed by rotation in depth to produce the God’s eye view. The opposite sequence was used to transition from 2-D to 3-D. The transition took approximately 3.2 s in both the continuous and the discrete conditions.

Results

Response time. RTs were averaged and analyzed as in previous experiments. There was a main effect for trial, $F(1, 23) = 141.93, MS_E = 2.08, p < .0001$. RTs were shorter on the second trial. There was an interaction between transition and trial, $F(1, 23) = 9.11, MS_E = 0.23, p < .01$. As shown in Figure 3, continuous transition produced shorter RTs than did discrete transition for the second trial in a pair (but not the first). There was an interaction among task, display, and trial, $F(1, 23) = 6.68, MS_E = 1.81, p < .05$. RTs were shorter for the

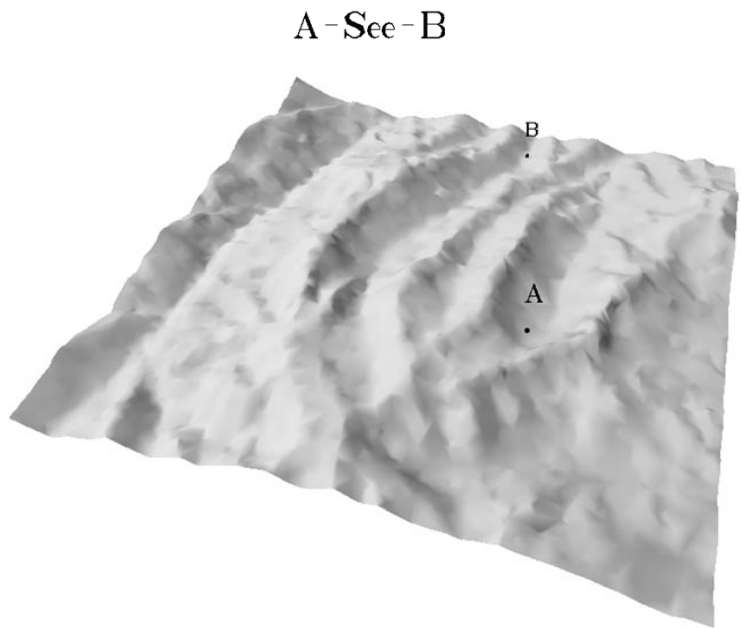


Figure 4. Example of 3-D displays used in Experiment 3. (A cannot see B in this figure.)

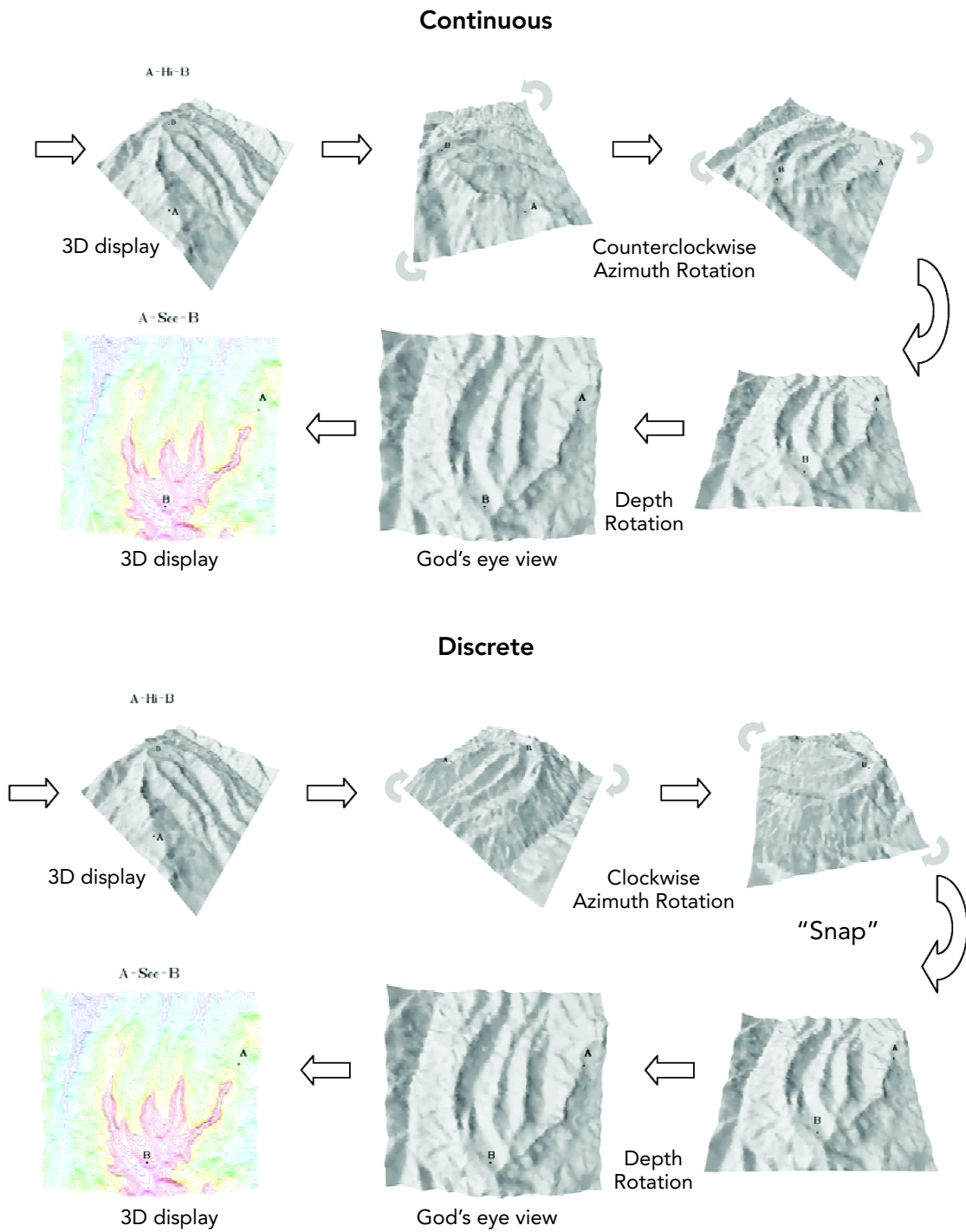


Figure 5. Depiction of smooth transitions between the 2-D and 3-D displays as used in continuous and discrete experimental conditions of Experiment 3. Horizontal translations are not depicted. (A cannot see B, and A is not higher than B in this figure.)

3-D displays in the A-See-B task and on the second trial for the A-High-B task (see Table 4).

Accuracy. The proportion of correct trials was computed and analyzed as in previous experiments. There was a main effect for trial, with participants being more accurate on the second trial, $F(1, 23) = 102.09$, $MS_E = 0.0733$, $p < .0001$. An interaction between transition and trial just failed to reach conventional significance levels, $F(1, 23) = 3.92$, $MS_E = 0.0019$, $p < .06$. As Figure 3 shows, continuous transition appeared to produce higher accuracy than discrete transition for the second trial in a pair. There was an interaction among task, display, and trial, $F(1, 23) = 60.94$, $MS_E = 0.0037$, $p < .0001$. For the A-High-B task, the 2-D display showed an advantage for the first trial only, and for the A-See-B task, display type had no effect (see Table 5).

Discussion

As predicted, posttransition performance was better in the continuous condition. There was no evidence for a speed-accuracy trade-off for this effect – accuracy was not reduced after continuous transition.

RTs were shorter for the 3-D displays in the A-See-B task, consistent with predictions. There was no difference in accuracy between 2-D and 3-D for A-See-B, so there was no speed-accuracy trade-off. For A-High-B, the accuracy results showed a 2-D advantage, consistent with predictions. However, this effect was limited to the first trial. Further, although there was no trade-off on the first trial (no RT difference on the first trial for A-High-B), 3-D produced shorter RTs on the second trial, indicating a trade-off on the second trial (3-D producing less accurate but faster performance than 2-D on the A-High-B task).

GENERAL DISCUSSION

We proposed a model in which a cognitive rep-

resentation of spatial elements is updated with shifts in viewpoint. According to the model, with a discrete change in viewpoint, the observer must determine correspondences between geospatial elements in the display and the cognitive representation. This is not necessary with a continuous change between viewpoints because the cognitive representation is continuously updated. We conducted three experiments in which there was smooth rotation of viewpoint between 2-D and 3-D views of geographic terrain. A different discrete transition condition was used in each experiment.

In the immediate scenario (Experiment 1), the second display was shown immediately after the first. In the delay scenario (Experiment 2), a blank screen was shown for a duration equal to that used for the continuous transition. Finally, in the preview scenario examined in Experiment 3, the terrain was shown for the same amount of time in the discrete condition, but the viewpoint “snapped” to the final orientation. As the model predicted, performance improved more for the continuous transition condition than for its discrete counterpart in each case.

This is not to suggest that having preview or preparation time will not aid performance. In many real-world contexts, it is likely that such factors will co-occur with smooth rotation, given that smooth rotation takes time to portray, and it seems likely that the user will take this time to prepare for subsequent task demands. Indeed, Figure 3 suggests that the difference between continuous and discrete conditions was larger when preview or preparation time was not available in the discrete condition (Experiments 1 and 2) than when it was (Experiment 3). Nonetheless, it appears that visual momentum through smooth rotation improves performance beyond what preview and preparation time can offer.

We predicted that the relative effectiveness of

TABLE 4: Experiment 3 – Mean Response Time (in Seconds) as a Function of Task, Display, and Trial

Task	2-D		3-D	
	Trial 1	Trial 2	Trial 1	Trial 2
A-High-B	3.539	2.055	3.588	1.653
A-See-B	4.640	2.794	3.400	1.650

Note. The within-subjects standard error of the mean was 0.075.

TABLE 5: Experiment 3 – Mean Accuracy as a Function of Task, Display, and Trial

Task	2-D		3-D	
	Trial 1	Trial 2	Trial 1	Trial 2
A-High-B	.973	.971	.747	.946
A-See-B	.669	.684	.646	.667

Note. The within-subjects standard error of the mean was .009.

2-D and 3-D displays in a task-switching context should be task dependent. This prediction was based on a large number of studies (see St. John et al., 2001; Wickens & Hollands, 2000) comparing 2-D and 3-D displays showing a task-dependent relation and, more specifically, on results obtained by St. John et al. (2001), who used the same tasks and similar displays. We summarize the results of our experiments with respect to our task dependency prediction in Table 6.

For the A-See-B task, the results were generally in accord with predictions (3-D beats 2-D). RTs were shorter for 3-D, and accuracy was no worse (except in Experiment 2). Indeed, St. John et al. (2001) obtained a 3-D advantage for RT only, with no difference in accuracy for the A-See-B task. For our A-High-B task, the accuracy results were in accord with predictions (2-D beats 3-D). However, the RT results for A-High-B showed the opposite pattern (3-D beats 2-D). This indicates a speed-accuracy trade-off for A-High-B: Participants were faster but less accurate with the 3-D display than with the 2-D display. St. John et al. (2001) found greater accuracy for 2-D than 3-D for the A-High-B task and no difference in RT. However, closer examination reveals that mean RT was greater for 2-D than 3-D displays in the A-High-B task in their experiment, although it was not a statistically significant difference.

Some previous studies have also shown a speed-accuracy trade-off with respect to 2-D and 3-D displays. For example, Wickens, Liang, Prevett, and Olmos (1996) had participants perform a simulated aircraft landing task and asked them to report on the position of nearby objects (buildings). When participants were asked if an object was above or below the aircraft, Wickens et al. (1996) found that their 3-D display produced less accurate judgments but required less time than their 2-D display condition, which contained separate elevation and God’s eye views.

They attributed the greater RT for the 2-D display to the requirement to examine both views. Although participants in our experiments saw only one 2-D display, they needed to derive altitude information from a color-coded scale for the A-High-B task; this may similarly have been time consuming relative to a line-of-sight estimate. For both experiments, the accuracy of judgments with 3-D displays may have been impaired because of a line-of-sight ambiguity (Wickens et al., 1996).

Conclusions

The use of dynamic transition is recommended when observers switch between two displays that depict the same spatial elements from different viewpoints. We examined a range of discrete scenarios (immediate, delay, and preview) to explore

TABLE 6: Summary of Task-Dependent Predictions and Results

	Dependent Variable	Experiment		
		1	2	3
A-See-B (prediction: 3-D beats 2-D)	Response time	3-D	3-D ^a	3-D
	Accuracy	3-D	2-D	Tie
A-High-B (prediction: 2-D beats 3-D)	Response time	3-D	3-D	Tie/3-D ^b
	Accuracy	2-D ^c	2-D ^c	2-D/Tie

^aEffect interacts with trial (3-D advantage greater on first trial). ^bEffect interacts with trial (3-D advantage on second trial only). ^cEffect interacts with trial (2-D advantage greater on first trial).

the consistency of the advantage across realistic situations and found an advantage for continuous transition in each case.

How long a transition time is necessary to see the advantage of continuous rotation? The issue of rotation speed appears an open question. Increasing the rotation speed until it fails to help may indicate how quickly the cognitive representation can be updated. Also relatively unexplored is whether allowing the operator to toggle the rotation timing or control the rotation speed or viewpoint would affect the continuous advantage. These would appear to be fruitful avenues for future research.

Beyond the specific experimental findings, there are more general implications of our approach. We introduced a cognitive model for visual momentum that predicts superior performance when task-relevant spatial elements are shown continuously during a viewpoint switch. Although the model clearly requires further validation, it could be used by designers as a heuristic for improving the cognitive transition between multiple displays while a user switches tasks.

For example, an engineer might want to switch from a 3-D skeleton framework to a 2-D elevation of an engine design to determine the exact distance between two bolt locations. The model would predict that the user might have difficulty mapping the two locations with a discrete shift in viewpoint, but a smooth viewpoint rotation should assist. Moreover, the model suggests that other transition methods (e.g., showing a skeleton framework instead of all visual detail, or displaying only task-relevant elements) should also be useful because the cognitive representation is being updated during the transition. In this sense, then, the model specifies why visual momentum (in the form of gradual transition between viewpoints) should be effective in terms of cognitive processing.

In summary, the results and model should be useful for the design of future command and control systems. The results should also have implications for many other domains, such as geographic information systems, virtual environments, and computer-assisted design.

ACKNOWLEDGMENTS

This research was supported by the DRDC Technology Investment Fund. We thank Jerzy Jarmasz and Jocelyn Keillor for comments and Tyler Hause and Matthew Lamb for assistance.

We also thank Harvey Smallman and Mark St. John for helpful comments and for providing details of the stimuli used in their experiments. Some experimental results were presented at a 2002 NATO Research and Technology Organisation workshop in Halden, Norway, and at annual meetings of the Human Factors and Ergonomics Society (2003 and 2004).

REFERENCES

- Andre, A. D., Wickens, C. D., Moorman, L., & Boschelli, M. M. (1991). Display formatting techniques for improving situation awareness in the aircraft cockpit. *International Journal of Aviation Psychology*, 1, 205–218.
- Aretz, A. J. (1991). The design of electronic map displays. *Human Factors*, 33, 85–101.
- Bederson, B. B., Hollan, J. D., Stewart, J., Rogers, D., Vick, D., Ring, L., et al. (1998). A zooming Web browser. In C. Forsythe, E. Grose, & J. Ratner (Eds.), *Human factors and Web development* (pp. 255–266). Mahwah, NJ: Erlbaum.
- Barnes, M. J. (2003). *The human dimension of battlespace visualization: Research and design issues* (ARL Tech. Rep. ARL-TR-2885). Aberdeen Proving Ground, MD: Army Research Laboratory.
- Gattis, M. (2004). Mapping relational structure in spatial reasoning. *Cognitive Science*, 28, 589–610.
- Guth, P. (2001). MICRODEM (Version 5.1) [Computer software]. Annapolis, MD: U.S. Naval Academy, Oceanography Department.
- Hochberg, J., & Brooks, V. (1978). Film cutting and visual momentum. In J. W. Senders, D. F. Fisher, & R. A. Monty (Eds.), *Eye movements and the higher psychological functions* (pp. 293–313). Hillsdale, NJ: Erlbaum.
- Hörmig, R., Oberauer, K., & Weidenfeld, A. (2005). Two principles of premise integration in spatial reasoning. *Memory and Cognition*, 33, 131–139.
- Loftus, G. R., & Masson, M. E. J. (1994). Using confidence intervals in within-subject designs. *Psychonomic Bulletin and Review*, 1, 476–490.
- MultiGen-Paradigm (2001a). Creator (Version 2.5) [Computer software]. San Jose, CA: Author.
- MultiGen-Paradigm (2001b). Vega (Version 2.5) [Computer software]. San Jose, CA: Author.
- Neale, D. C. (1996). Spatial perception in desktop virtual environments. In *Proceedings of the Human Factors and Ergonomics Society 40th Annual Meeting* (pp. 1117–1121). Santa Monica, CA: Human Factors and Ergonomics Society.
- Neale, D. C. (1997). Factors influencing spatial awareness and orientation in desktop virtual environments. In *Proceedings of the Human Factors and Ergonomics Society 41st Annual Meeting* (pp. 1278–1282). Santa Monica, CA: Human Factors and Ergonomics Society.
- Olmos, O., Liang, C. C., & Wickens, C. D. (1997). Electronic map evaluation in simulated visual meteorological conditions. *International Journal of Aviation Psychology*, 7, 37–66.
- Olmos, O., Wickens, C. D., & Chudy, A. (2000). Tactical displays for combat awareness: An examination of dimensionality and frame of reference concepts and the application of cognitive engineering. *International Journal of Aviation Psychology* 10, 247–271.
- Roth, S. F., Chuah, M. C., Kerpedjiev, S., Kolojechick, J., & Lucas, P. (1997). Toward an information visualization workspace: Combining multiple means of expression. *Human-Computer Interaction*, 12, 131–186.
- Shepard, R. N., & Metzler, J. (1971). Mental rotation of three-dimensional objects. *Science*, 171, 701–703.
- St. John, M., Cowen, M. B., Smallman, H. S., & Onk, H. M. (2001). The use of 2D and 3D displays for shape-understanding versus relative-position tasks. *Human Factors*, 43, 79–98.
- St. John, M., Smallman, H. S., & Cowen, M. B. (2002). *Evaluation of schemes for combining 2-D and 3-D views of terrain for a tactical routing task* (Tech. Rep.). San Diego, CA: Pacific Science & Engineering Group.

- Wickens, C. D., & Hollands, J. G. (2000). *Engineering psychology and human performance* (3rd ed.). Upper Saddle River, NJ: Prentice-Hall.
- Wickens, C. D., Liang, C. C., Prevett, T., & Olmos, O. (1996). Electronic maps for terminal area navigation: Effects of frame of reference and dimensionality. *International Journal of Aviation Psychology*, 6, 241–271.
- Wickens, C. D., & Thomas, L. C. (2000). Frames of reference for the display of battlefield information: Judgement-display dependencies. *Human Factors*, 42, 660–675.
- Woods, D. D. (1984). Visual momentum: A concept to improve the cognitive coupling of person and computer. *International Journal of Man-Machine Studies*, 21, 229–244.

Justin G. Hollands is a defence scientist at Defence Research and Development Canada – Toronto. He received his Ph.D. in psychology from the University of Toronto in 1993.

Nada J. Pavlovic is a research technologist at Defence Research and Development Canada – Toronto. She received her M.A.Sc. in mechanical and industrial engineering (human factors) from the University of Toronto in 2006.

Yukari Enomoto received her M.A.Sc. in systems design engineering from the University of Waterloo, Waterloo, Ontario, in 2006. She resides in Toronto.

Haiying Jiang is a senior consultant at the Global Delivery China Center, Hewlett-Packard Company, Shanghai, China. He received his M.S. in motor control and learning from the University of Illinois at Chicago in 1997.

Date received: December 22, 2004

Date accepted: September 7, 2007