

Spatial Representation and Reasoning for Human-Robot Collaboration

**William G. Kennedy, Magdalena D. Bugajska, Matthew Marge, William Adams,
Benjamin R. Fransen, Dennis Perzanowski, Alan C. Schultz, and J. Gregory Trafton**

Naval Research Laboratory, 4555 Overlook Ave. SW, Washington, DC 20375
{wkennedy, magda, adams, fransen, dennisp, schultz, trafton}@aic.nrl.navy.mil, mmarge@gmail.com

Abstract

How should a robot represent and reason about spatial information when it needs to collaborate effectively with a human? The form of spatial representation that is useful for robot navigation may not be useful in higher-level reasoning or working with humans as a team member. To explore this question, we have extended previous work on how children and robots learn to play hide and seek to a human-robot team covertly approaching a moving target. We used the cognitive modeling system, ACT-R, with an added spatial module to support the robot's spatial reasoning. The robot interacted with a team member through voice, gestures, and movement during the team's covert approach of a moving target. This paper describes the new robotic system and its integration of metric, symbolic, and cognitive layers of spatial representation and reasoning for its individual and team behavior.

Introduction

Reconnaissance, or RECON, is the essential first step of any military action whether it is setting up a defensive position or planning an attack. Within a U.S. Marine Corps reconnaissance unit, a RECON team, Marines operate in pairs and always within sight of each other to ensure mutual support. The core competencies for this type of mission include spatial reasoning, perspective-taking, and covert communications. In order to provide effective support within a RECON team, future tactical mobile robots must have credible competencies in all of these areas.

How any of these core abilities should be achieved is still subject of a debate in the community. For example, one of the many spatial representations could be used to perform spatial reasoning (Montemerlo, Roy, and Thrun 2003; Schultz, Adams, and Yamauchi 1999). The decision of which reasoning algorithm to use is usually based on its convenience, computational efficiency, and robustness. Trafton, et al. (2005a) has suggested that another aspect to be considered while making this choice is how well the system works together with a person. We were guided by

the representation hypothesis that suggests that a system that uses representations and processes or algorithms similar to a person's will be able to collaborate with a person better than a computational system that does not. Furthermore, such a system will be less likely to exhibit unreasonable behavior, which is a sure benefit in any strategic domain.

Our principal goal in this project is to show how the scientifically-principled integration of computational cognitive models can facilitate human-robot interaction and, specifically, how different spatial representations need to be integrated for coherent human-robot interaction (HRI). Note that this paper reports a systems approach to HRI; psychological studies and usability tests will be performed in the future. In addition, our engineering goal in this project is to create a system that can covertly approach a moving target with a team member in a laboratory scenario inspired by the Marine RECON task.

Laboratory RECON Scenario

In the research presented in this paper, we introduce a reconnaissance task that requires a robot and a human to work together to covertly track and approach a moving target (a human or robot). See Figure 1.



Figure 1. Robot, Target (standing), Team Member (crouching), and Objects in the Laboratory Environment

The target continually moves either to random locations or in a predefined path that is not known to the human-

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 2007		2. REPORT TYPE		3. DATES COVERED 00-00-2007 to 00-00-2007	
4. TITLE AND SUBTITLE Spatial Representation and Reasoning for Human-Robot Collaboration				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) Naval Research Laboratory, Navy Center for Applied Research in Artificial Intelligence (NCARAI), 4555 Overlook Avenue SW, Washington, DC, 20375				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 Twenty-Second Conference on Artificial Intelligence (AAAI-07), 22-26 July 2007, Vancouver, Canada					
14. ABSTRACT How should a robot represent and reason about spatial information when it needs to collaborate effectively with a human? The form of spatial representation that is useful for robot navigation may not be useful in higher-level reasoning or working with humans as a team member. To explore this question, we have extended previous work on how children and robots learn to play hide and seek to a human-robot team covertly approaching a moving target. We used the cognitive modeling system, ACT-R, with an added spatial module to support the robot's spatial reasoning. The robot interacted with a team member through voice, gestures, and movement during the team's covert approach of a moving target. This paper describes the new robotic system and its integration of metric, symbolic, and cognitive layers of spatial representation and reasoning for its individual and team behavior.					
15. SUBJECT TERMS					
16. SECURITY CLASSIFICATION OF:			17. LIMITATION OF ABSTRACT Same as Report (SAR)	18. NUMBER OF PAGES 6	19a. NAME OF RESPONSIBLE PERSON
a. REPORT unclassified	b. ABSTRACT unclassified	c. THIS PAGE unclassified			

robot team. However, the target's position is always available to the human/robot team. The target has a limited field of view that determines when it can see the members of the human/robot team.

The goal of the human/robot team is to use knowledge of the target's position, the target's field of view, and obstacles in the environment to follow the target and to get as close as possible to the target while remaining as hidden as possible. The covertness part of the goal causes the team members to minimize their visibility to the target. The requirement to approach the target prevents the team from finding a single, covert hiding place and staying there.

This scenario provides challenges in spatial reasoning and modeling of the behavior of the target to predict its behavior rather than having a static, spatial reasoning problem as in earlier research. We will discuss the design of our StealthBot system intended to meet these challenges and the behavior of the StealthBot in a team environment.

StealthBot System Overview

The StealthBot system will be discussed in three layers or tiers similar to those used by other researchers (Bonasso et al 1997; Montemerlo, Roy, and Thrun 2003). The three layers are a hardware layer with sensors and effectors, a spatial support layer, and a cognitive layer as shown in Figure 2. The next section will focus on the non-spatial components. The spatial components will be discussed in detail in subsequent sections.

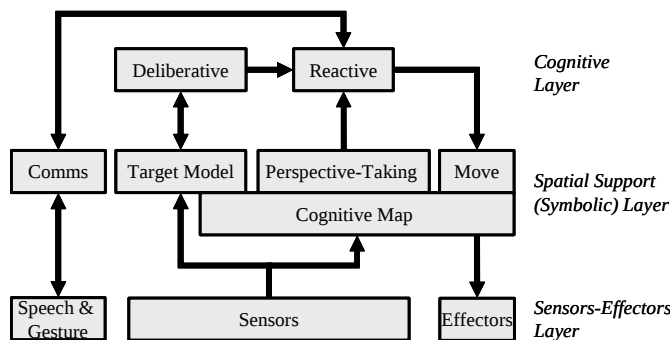


Figure 2. StealthBot Three Layer Architecture

Non-Spatial Components

The non-spatial components are the basic robot hardware, its speech recognition system, and its gesture recognition system.

Robot Hardware

The robot is a commercial iRobot B21r. It is an upright cylinder with a zero-turn-radius drive system and an array of range and tactile sensors. The CMVision package (Bruce, Balch, and Veloso 2000) provides simple color blob detection using a digital camera mounted on the robot.

The color marker was used as the identifier for characteristics of an object: the target was orange, the team member green, and stationary objects blue. The bearing to objects was determined from its location in the camera image, while the range was obtained using a laser rangefinder. In addition, a high-fidelity stereo camera system was added to allow for gesture recognition.

The robot's mobility capabilities, including map building, self-localization, path planning, collision avoidance, and on-line map adaptation in changing environments, were introduced previously as the WAX system (Schultz, Adams, and Yamauchi 1999). Additional details of the robot's basic systems are provided in a previous paper (Trafton et al. 2006).

Speech Recognition

To provide the StealthBot with the capability of handling verbal commands, if needed, ViaVoice™ is used for speech recognition. For this scenario, a very simple list of BNF (Backus-Naur Form) grammar definitions was compiled. With this speech capability enabled, the human team member can order the StealthBot to "Attention", "Stop," "Assemble", (i.e., "Come here"), "As you were" (i.e., "Continue"), and "Report". The first four of these were taken from the U.S. Marine Corps Rifle Squad manual, FMFM 6-5. The final command, to report, was added to allow the StealthBot to share its knowledge with its team member. Since the verbal interaction with the StealthBot is rather simple, no further natural language processing was required for this task, although we do have more advanced capability. Spoken input need not always be supplied and interaction with the human team member may be based solely on gestures.

Gesture Recognition

To maintain the covert nature of StealthBot operations for this laboratory-based RECON scenario, gesture-based communications (Perzanowski et al. 1998) was integrated because it makes covert communications possible, i.e., without broadcasting sound or electromagnetic signals. A gesture identification module has been incorporated into the StealthBot to identify gestures based on hand motion and position (Fransen et al. 2007). The gesture-based system currently recognizes the same commands as the speech recognition system, "Attention," "Stop," "Assemble," "As you were," and "Report."

Spatial Representation and Reasoning

A spatial reasoning capability is essential to covert operations in urban environments and it is strongly determined by the underlying representation. A great many spatial representations have been suggested for human and artificial navigation, communication, and reasoning. These representations include survey and route representations (Taylor and Tversky 1992), egocentric information (Previc

1998), metric representations, qualitative representations (Forbus 1993), and topographic representations. Our approach has been to use efficient computational representations, such as metric representation within our robotic system, until the point where person-interaction is needed. At which point, these representations are converted into a more abstract and “team member friendly” format. Interestingly, the approach of many roboticists has been to take egocentric information and convert it into an exocentric representation in a series of iterations for external display to a person (for example, Schultz, Adams, and Yamauchi 1999), and occasionally for the robot’s own navigation/reasoning.

Spatial information is generated and used differently in each of three layers of the StealthBot: the basic robot sensors and effectors layer, the spatial support layer, and the cognitive layer implemented by an ACT-R cognitive model (Anderson and Lebiere 1990). An appropriate type of spatial representation is used at each level of the architecture often requiring the translation of information between different representations. We discuss when and why our system integrates different representations below.

Sensors and Effectors Layer: Metric Information

The sensors and effectors generate and make use of egocentric metric information. The metric information includes numerical values for range and bearing. This egocentric metric information is then converted into both an egocentric and exocentric evidence grid. The exocentric evidence grid is considered a long-term map of the world, while the egocentric evidence grid is considered a short-term perception map. Localization occurs by registering the egocentric representation within the exocentric (Hiatt et al. 2004; Schultz and Adams 1998). The metric information in this layer is precise but noisy and the system has been shown to deal with the noise effectively (Schultz and Adams 1998). The metric information is primarily used by the robot for navigation and collision avoidance. This layer also receives motion commands to move the robot to map coordinates and turn the robot to face a specific map location. Object avoidance and getting to a specified coordinate location is handled by this layer.

These representations are not considered cognitively plausible. However, they are a fundamental part of our core robotic system (e.g., they are proven, fast, and efficient at navigation and obstacle avoidance). The metric information is converted into symbolic information and a “cognitive map” in the spatial support layer to facilitate cognitively plausible reasoning in the cognitive layer which in turn facilitates human-robot interaction.

Spatial Support Layer: Symbolic Information

The spatial support layer provides the interface between the robot’s hardware and cognitive layers. Metric information from the sensors is translated into a cognitive map. This layer also analyzes target motion and provides symbolic information modeling the target’s motion to the

cognitive layer. Within this layer, the StealthBot’s visibility by the target is determined based on a clear line-of-sight between the target and the robot. Thus, there are three components associated with passing information from the sensors to the cognitive layer: the cognitive map, the tracking of target motion, and visibility determination. In the opposite direction, from the cognitive layer to the effectors, this layer converts the cognitive information into a metric representation to be used by the robot’s effectors.

The Cognitive Map. The cognitive map is our implementation of the hypothesis that people represent space in a qualitative manner. The cognitive map is created and maintained based on the information from the metric layer. Objects are placed in a 2-D grid based on their metric information. However, the map does not maintain the precise metric location of objects. To support spatial reasoning, the cognitive map is used to provide relationships between objects not easily available in a symbolic representation alone. For example, only knowing that a target is left of a building and a Marine is also left of a building, does not automatically provide information about the relative position of the target to the Marine.

Our cognitive map is used to support such high-level, symbolic reasoning about the space. It facilitates the robot reasoning about the relative locations of the target, team member, itself, and the objects in the environment and then good places to hide in the current and future states of the environment.

In our system, the information passed to the cognitive layer from the cognitive map consists only of the identifier of the object nearest to the target and the spatial relationship of the target to that object, such as “north-of” “box2”, and the analogous information about the object nearest to the StealthBot. The distances and relationships generated from the cognitive map are a symbolic (near, far, etc.) and are based on cognitive map coordinates, not their original metric information. The use of symbolic distances can result in ambiguity as to which object is the closest. The cognitive layer must deal with this ambiguity and does so by having specific rules for these situations.

The cognitive plausibility of cognitive maps has been the subject of some debate (Tolman 1948; Tversky 1993; Previc 1998). The prevalent view seems to be that it takes people time and effort to build a cognitive map; it is not an “automatic” process. However, from a computational perspective, the translation of metric data to a 2-D cognitive map is relatively straightforward. A similar translation into a complex cognitively plausible 3-D egocentric representation is not currently available, either within ACT-R or within the general cognitive science community. Our relatively simple cognitive map has both computational efficiency and cognitive plausibility, although we acknowledge our representation is not optimized for either.

Modeling Target Motion. This intermediate layer also develops symbolic knowledge concerning the movement of the target. The target’s movement is currently modeled as a straight line and its current direction is classified as: none

(i.e., not moving), north, north-east, east, south-east, etc. The duration of the target's movement in one of these directions is also available. When a change is detected, the cognitive layer is given the length of the track that ended with the change and the new track's heading.

The StealthBot models the target's movement on three levels. The target's current course and speed are based on sensor input and are referred to as the target's current tactic. A series of tactics is treated as a strategy and strategies are combined as necessary to accomplish missions. Using this information, the StealthBot can reason about the target's strategy and mission.

We considered using Kalman filters (Kalman 1960) based on their success in tracking movement in robotics environments. However, we decided not to use that technique because Kalman filters do not allow access to internal components of the representation and we need to reason with the target's trajectory and changes to its trajectory for cognitively plausible spatial reasoning.

Visibility Determination. The spatial support layer also affords the spatial aspect of perspective-taking in the form of an evaluation of the visibility of the StealthBot by the target. When requested by the cognitive layer, line-of-site and target field-of-view calculations are made based on the current model of the target's motion and the cognitive map. The result is provided to the cognitive layer. Currently, we assume that the target has perfect vision over a 180-degree field forward in the direction of movement. This is similar to ACT-R's visual module (Anderson and Lebiere 1990).

Interacting with the ACT-R Model. We chose not to directly modify ACT-R 6.0 (<http://act-r.psy.cmu.edu>) to implement this spatial module. Instead, this layer indirectly provides spatial representation and processing in support of higher-level spatial reasoning by the cognitive model. To pass the information to the cognitive model, we inserted chunks directly into the declarative memory of the ACT-R system. This was done when either the target moved enough to be in a different cognitive map cell or there was a change in its direction. The cognitive map itself is not passed to the cognitive layer nor is it directly accessible by that layer. ACT-R productions react to the change and reconsider what action the StealthBot should take.

Cognitive Layer: ACT-R Cognitive Model

We have built a cognitive model of what we believe is plausible for high-level spatial representation and reasoning. The cognitive model is implemented in ACT-R which has a long and successful history of representing and matching human cognition. The ACT-R cognitive model has pre-loaded declarative and procedural knowledge and learns new knowledge from interactions with the environment.

Declarative Knowledge. The declarative knowledge is represented as chunks of information with symbolic attribute slots and values. The information from the spatial support layer is provided to the cognitive layer as declarative memory chunks. The chunk representing a

change in the target's location includes both exocentric and egocentric information from the lower layer. The exocentric information, i.e., externally referenced information, is the object closest to the target and the target's relative bearing from it. The egocentric information is the object closest to the StealthBot and the target's relative bearing from it. Determining these references is an example of the translation of information between different spatial representations necessary for higher-level reasoning.

Procedural Knowledge. The procedural knowledge in form of production rules encodes process knowledge on how to:

- handle the environmental information including messages from the StealthBot's team member,
- predict where the target will be in the near future,
- make deductions about where the StealthBot should hide next,
- respond to team member communications, and
- decide whether it is safe to move.

Although the robot continuously monitors the environment for navigational purposes, collecting information for high-level reasoning is a deliberate act initiated by the cognitive level. The cognitive level deliberately looks for new information about the location of the target and communications from the team member.

Based on available and inferred information, the StealthBot decides on the next good hiding place. The StealthBot cognitive model has procedural knowledge encoding the knowledge that if the target is on the north, east, south, or west side of an object, it should hide on the opposite side of the object. We expect that these productions could be learned through experience in this situation, but we encoded them directly.

The StealthBot also decides which object to hide behind, the object closest to the target or the one closest to the StealthBot. StealthBot's choice of a hiding place is based on its information about the target's location, the target's predicted movement, and its own visibility. When the StealthBot decides on a good place to hide, it checks whether it is safe to move there deliberately based on its current visibility. If it is safe to move deliberately and its team member has not directed it to hold its position, the StealthBot starts to move to its desired location. During its movement, it repeatedly checks for changes in the environment and re-evaluates where to move to. If there is a change in the environment, such as the target moves such that another hiding place is preferred, the StealthBot changes where it is going to hide.

Whenever information about the target's tactics or strategy is available, the robot anticipates the target's behavior by simulating within the cognitive map the future state of the world and making a decision to hide next to the object that is closest to the target in the future. The prediction of the target's behavior is based on prior observations of the target and domain knowledge about

patrolling, transiting, and holding strategies and associated tactics.

However, if the StealthBot recognizes, during its movement to its next hiding place, that it is or has become visible, the StealthBot takes immediate (reactive) action to hide using the nearest object. This integrates the modeling of urgent decisions that do not fully consider all information available in the environment similar to research on performing two tasks (Taatgen 2005). These decisions are more direct in deciding where to hide than those discussed earlier. Under this condition, the target's movements are not considered, only the target's position and only how to get out of sight as fast as possible. These productions override the previous thoughtful decision-making about the next hiding place and substitute the appropriate hiding place next to the nearest object. Overall, the deliberate and reactive productions lead to what we consider to be reasonable hiding behavior.

StealthBot Behavior

We have run our StealthBot in a number of scenarios to demonstrate competence in the core competency areas of spatial reasoning, perspective-taking, and covert communications. Figure 3 shows a diagram of one run, which demonstrated the robot's ability to anticipate the target's movement and to dynamically revise its model.

The tracks of the target and the StealthBot are indicated by a sequence of letters: a, b, c, etc. In this run, the target moved right to left above both pillars and the StealthBot started in the southeast. After the initial sensor sweep, the StealthBot located the target north of the "Pillar1". So, it immediately moved to hide south of "Pillar1". Hidden at point b, the StealthBot determined that the target appeared to be moving southwest (toward the lower left) by simulating the target's motion for several steps into the imagined future resulting in the target being south of "Pillar2". The StealthBot decided that north of "Pillar2" would be a good hiding place based on its mission to covertly approach the target. So, the StealthBot began moving to north of "Pillar2". As it began to move from behind "Pillar1", at step c, it inferred based on the target's location and the sensor model that it was visible. It overrode its previous plan and immediately hid south of "Pillar1", step d.

Once there, based on the updated spatial information, the StealthBot was able to revise its belief about the target's behavior. The StealthBot realized that the target was moving more west than southwest and revised its prediction to put the target west of "Pillar2". In anticipation of this future state, the StealthBot began to move towards the east of "Pillar2." It could safely (invisibly) follow the target and hide east of "Pillar1" at step e due to the target's limited field of view. The StealthBot remained at that location for the rest of the run.

Without the ability to anticipate the target's future location, the robot would not have reactively considered this new hiding place until step h. As the target traversed

the field, a purely reactive system would have considered hiding west, south, and only then east of "Pillar2", or worse, only danced around "Pillar1" in the first place.

This demonstrates reasonable individual behavior under these conditions. This demonstration suggests that the robot could be a competent team member, exhibiting good, though imperfect, behavior when working alone.

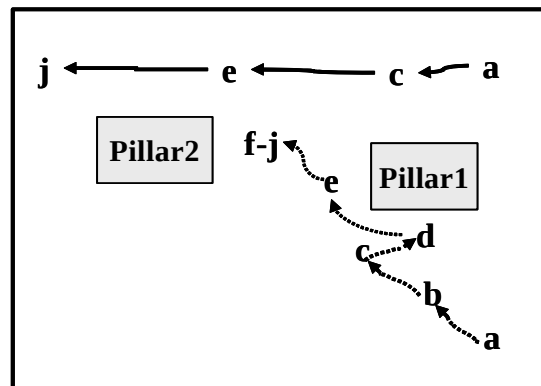


Figure 3. Diagram of Target (solid line) and StealthBot (dotted line) behavior during the scenario. Both began at respective (a) locations. the StealthBot immediately moved to (b) to hide from Target behind "Pillar 1" and then continued towards "Pillar 2", the anticipated ultimate hiding place. It became visible at (c), so it quickly retreated behind "Pillar 1" to (d). Next, the StealthBot left the refuge due to covertness afforded by Target's limited field of view (e) and reached the desired hiding place (f-j).

The primary difference between individual behavior and team behavior is that the robot is able to covertly communicate with its human team member to take advantage of the human's capabilities. For example, a responsible team human team member hiding behind "Pillar 2" and able to see the target's motion could have commanded the robot to stop at step "b" to avoid its becoming visible, and then to resume a few steps later. Furthermore, part of a RECON mission is to gather and report information. To support this, the robot is able to report its understanding of the situation in terms useful to the human. Figure 4 is an example of such a report for the run discussed above.

```
AT TIME e, TARGET IS "north-of"
"pillar2" HEADING W WITH SENSOR OFF.
TARGET APPEARS TO HAVE STRATEGY "NE-NW-
cycle" AND MISSION PATROL. CURRENT PLANS
ARE TO HIDE "east-of" "pillar2".
REPORTING FROM "pillar2".
```

Figure 4. Report from the StealthBot

Discussion

Our StealthBot is able to work independently or with a team member to covertly approach and follow another robot or person in our laboratory RECON scenario under our laboratory conditions, therefore meeting our

engineering goal. Our robotic/computational goal was also met: we integrated a computational cognitive architecture (ACT-R) as the basis for cognitively plausible (at least in parts), spatial reasoning and as the basis for interacting with the other team member using the layers of spatial representation and reasoning discussed in this report.

One of our primary successes in this project was to have a scientifically principled method of integrating multiple spatial representations each useful at their own level of reasoning. Specifically, we used a metric representation primarily for navigation and collision avoidance; this type of representation is a fundamental part of most robotic systems. We used a cognitive map representation as the “glue” between our metric and cognitive layers. Finally, our cognitive representation was based on ACT-R, which allows us to claim the robot reasoned in the cognitive level similarly to how people do. This design allowed us to build a reasonably competent robotic team-member. The human interaction techniques, specifically an improved natural gesture system, allowed a human to interact with the robot as a human would in a covert fashion.

Acknowledgements

Our work benefited from discussions with Dan Bothell, Cynthia Breazeal, Scott Douglass, Farilee Mintz, Linda Sibert, and Scott Thomas. This work was performed while the first author held a National Research Council Research Associateship Award and was partially supported by the Office of Naval Research under job order numbers 55-8551-06, 55-9019-06, and 55-9017-06. The views and conclusions contained in this document should not be interpreted as necessarily representing official policies, either expressed or implied, of the U.S. Navy.

References

- Anderson, J.R. and Lebiere, C. 1998. *The Atomic Components of Thought*. Mahwah, NJ, Erlbaum.
- Bonasso, R.P.; Firby, R.J.; Gat, E.; Kortenkamp, D.; Miller, D.; and Slack, M. 1997. Experiences with an architecture for intelligent, reactive agents. *Journal of Experimental and Theoretical Artificial Intelligence* 9(2/3): 237-256.
- Bruce, J.; Balch, T.; and Veloso, M. 2000. Fast and inexpensive color image segmentation for interactive robots. In *Proceedings of the 2000 IEEE/RSJ International Conference on Intelligent Robots and Systems*. Takamatsu, Japan: 2061-2066.
- Forbus, K. 1993. Qualitative process theory: Twelve years after. *Artificial Intelligence* 59: 115-123.
- Fransen, B.; Morariu, V.; Martinson, E.; Blisard, S.; Marge, M.; Thomas, S.; Schultz, A.; and Perzanowski, D. Using vision, acoustics, and natural language for disambiguation. Forthcoming.
- Hiatt, L.M.; Trafton, J.G.; Harrison, A.M.; and Schultz, A.C. 2004. A cognitive model for spatial perspective taking. In *Proceedings of the Sixth International Conference on Cognitive Modeling*. Pittsburgh, PA: Carnegie Mellon University/University of Pittsburgh. 354-355.
- Montemerlo, M.; Roy, N.; and Thrun, S. 2003. Perspectives on Standardization in Mobile Robot Programming: the Carnegie Mellon Navigation (CARMEN) Toolkit. In *Proceedings of the IEEE/RJS International Conference on Intelligent Robots and Systems (IROS 2003)*, Las Vegas: 2436-2441.
- Perzanowski, D.; Schultz, A.C.; and Adams, W. 1998. Integrating Natural Language and Gesture in a Robotics Domain. In *Proceedings of the IEEE International Symposium on Intelligent Control: ISIC/CIRA/ISAS Joint Conference*. Gaithersburg, MD, National Institute of Standards and Technology: 247-252.
- Previc, F. H. 1998. The neuropsychology of 3-D space. *Psychological Bulletin* 124(2): 123-164.
- Schultz, A.C. and Adams, W. 1998. Continuous localization using evidence grids. In *Proceedings of the 1998 IEEE International Conference on Robotics and Automation*, 2833-2939. Leven, Belgium: IEEE Press. .
- Schultz, A. C.; Adams, W.; and Yamauchi, B. 1999. Integrating exploration, localization, navigation, and planning with a common representation. *Autonomous Robots* 6: 293-308.
- Taylor, H.A. and Tversky, B. 1992. Spatial mental models derived from survey and route descriptions. *Journal of Memory and Language* 31: 261-292.
- Tolman, E.C. 1948. Cognitive Maps in Mice and Men. *The Psychological Review* 55(4): 189-208.
- Taatgen, N. 2005. Modeling Parallelization and Flexibility Improvements in Skill Acquisition: From Dual Tasks to Complex Dynamics, *Cognitive Science* 29: 421-455.
- Trafton, J.G., Schultz, A.C.; Bugajska, M.D.; and Mintz, F.E. 2005a. Perspective-taking with robots: experiments and models. In *Proceedings of the International Workshop on Robot and Human Interactions*, 580-584. IEEE Press.
- Trafton, J.G.; Cassimatis, N.L.; Bugajska, M.D.; Brock, D.P.; Mintz, F.E.; and Schultz, A.C. 2005b. Enabling effective human-robot interaction using perspective-taking in robots. *IEEE Transactions on Systems, Man, and Cybernetics* 35(4): 460-470.
- Trafton, J.G., Schultz, A.C.; Perzanowski, D.; Bugajska, M.D.; Adams, W.; Cassimatis, N.L.; and Brock, D.P. 2006. Children and robots learning to play hide and seek. *2006 ACM Conference on Human-Robot Interactions*, Salt lake City, UT, ACM Press: New York. 242-249.
- Tversky, B. 1993. Cognitive maps, cognitive collages, and spatial mental models. In A.U. Frank & I. Campari (Eds.), *Spatial information theory: A theoretical basis for GIS*, 14-24. Berlin: Springer-Verlag.