

Spatial Abstraction for Autonomous Robot Navigation

Susan L. Epstein

Department of Computer Science, Hunter College and The Graduate Center of The City University of New York, New York, New York, USA

Anoop Aroor

Department of Computer Science, The Graduate Center of The City University of New York, New York, New York, USA

Matthew Evanusa

Department of Computer Science, Hunter College of The City University of New York, New York, New York, USA

Elizabeth I. Sklar

Department of Computer Science, University of Liverpool, Liverpool, UK

Simon Parsons

Department of Computer Science, University of Liverpool, Liverpool, UK

Corresponding author:

Susan L. Epstein, susan.epstein@hunter.cuny.edu

001-212-772-5213 (telephone)

001-212-772-5219 (fax)

Keywords Spatial model, Spatial abstractions, Qualitative reasoning, Robot navigation

Spatial Abstraction for Autonomous Robot Navigation

Abstract Optimal navigation for a simulated robot relies on a detailed map and explicit path planning. This approach is problematic for real-world robots, whose sensors and actuators are subject to noise and error, and whose environment may be dynamic. This paper reports on autonomous robots that rely on local spatial perception, learning, and commonsense rationales instead. Despite realistic actuator error, learned spatial abstractions form a model that supports effective travel.

Introduction

SemaFORR is a navigation system that learns and reasons about *spatial abstractions*, which remove perceived but irrelevant details from spatial information. The thesis of our work is that spatial abstractions learned from local sensing during travel can support effective, autonomous robot navigation. The principal result reported here is that our approach, without planning or a map, quickly produces efficient travel in a variety of spaces (*worlds*).

In our robots’ two-dimensional worlds, maps are unreliable or unavailable, and landmarks may be absent, obscured, or obliterated. Such spaces include complex office buildings, warehouses, and search and rescue environments. *SemaFORR* makes decisions with a cognitive architecture that relies on reactive and heuristic procedures based on simple rationales and spatial abstractions that describe where the robot has been. As it travels, the robot learns *affordances*, abstractions that facilitate movement and represent the world. These include unobstructed areas, useful transit points, and route segments. Together they form a *model* that represents the world but is not a map.

This paper reports on *SemaFORR*’s performance in worlds with different connectivity. The next two sections describe the abstractions, the robot, and

how *SemaFORR* abstracts its experience and reasons about it. Subsequent sections provide experimental methods, results, related research, and a discussion.

Robots and Learned Abstractions

A robot’s *position* in a world is its *heading* (the direction it faces) and its *location* (real-valued planar coordinates). A robot’s experience is a sequence of *decision points* where it senses, decides, and then acts. “Sense” extracts a partial view of the world through a *wall register* of 10 limited-range sensors. They calculate the distance to the nearest wall, as in Figure 1(a).

SemaFORR’s spatial abstractions are regions, trails, and conveyors. A *region*, such as that in Figure 1(a), is an area without permanent obstructions (e.g., walls). Wherever the robot senses, it learns a region: the circle centered at its location with radius equal to the shortest sensed distance. Whenever the robot crosses a region’s perimeter, that point becomes an *exit* from the region (shown here as a dot). The regions and their exits form a *skeleton* for the space the robot explores, as in Figure 1(b). A *leaf region* has exits only to one other region. (With perfect knowledge a leaf region is a dead-end.) A *trail*, such as that in Figure 1(c), is a revision of a path that reached a target. Travel along any subsequence of a trail in either direction should be reliable. Finally, *conveyors*, such as those in Figure 1(d), are small areas regularly used in successful travel. A conveyor represents a useful, target-independent transit point.

All these abstractions are approximations of a robot’s experience in its world. If the robot never enters a particular area, it will not be included in its model. Regions are disjoint, but can grow or shrink as the robot senses in new locations. A trail is not a

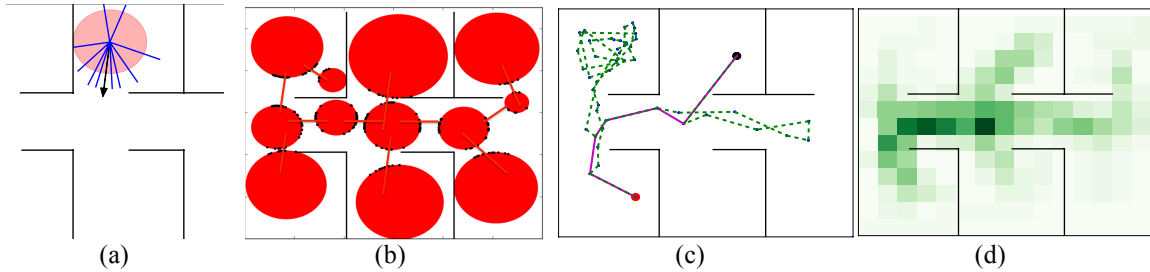


Fig. 1 Spatial abstractions: (a) a wall register with its inferred region (b) a learned, region-based skeleton that treats leaf regions as dead-ends (c) a (dashed) path to the target at top center, with its (solid) inferred trail (d) conveyors, where darker cells are visited more often

perfect path; it reimagines which locations further along its route the robot might have perceived, and moved to, sooner. The algorithm that learns conveyors superimposes on the world a grid with cells 1.5 times the size of the robot’s footprint, and increments a cell’s score each time a trail passes through it.

The real-world robot of interest here is a Surveyor SRV-1 Blackfin, a small, inexpensive robot platform. The Blackfin is subject to actuator error, that is, it may imperfectly execute any chosen action. The work reported here is in simulation, and models the actuator error of Blackfins observed in our laboratory, where larger actions incur larger errors.

Traditionally, robots navigate a shortest path that they plan in a perfect map. Increasingly sophisticated mapping and planning techniques have been developed to contend with such real-world issues as moving obstacles, noisy actuators, and dynamic worlds. Nonetheless, when its plan fails, the robot must repair it or replan.

Reasoning in SemaFORR

FORR (FOR the Right Reasons) is a cognitive architecture for the development of expertise (Epstein 1994). An agent developed within *FORR* inherits its decision mechanism, but specializes it for a particular application area. SemaFORR is a *FORR*-based system that makes robots’ navigation decisions.

FORR’s *decision cycle* repeatedly chooses one action at a time. Input to SemaFORR’s decision cycle is the robot’s position and wall register, learned spatial abstractions, its possible actions and target, and a history of its travel to that target. SemaFORR’s possible actions are 5 forward linear *moves* of various lengths, 10 clockwise or counter-clockwise *turns* of various degrees, and a pause (do nothing). SemaFORR alternately chooses a move or pause on one decision cycle and then a turn or pause on the next. (Intermediate pauses support longer consecutive moves or turns.)

FORR’s “right reasons” are called *Advisors*. Each is a salient preference mechanism for action selection, implemented as a time-limited procedure. During a decision cycle, an Advisor produces *comments* on possible actions. A comment’s numerical *strength* indicates the Advisor’s degree of support or opposition to the action.

SemaFORR’s *tier-1 Advisors* are presequenced, quick, correct, commonsense reasons for action selection. If the target is perceived, they select an action that drives the robot toward it; otherwise they eliminate actions that would cause a collision (with a wall or another robot) or oscillation in place. All remaining possible actions are forwarded to tier 3.

(Tier 2 is the focus of current work on reactive planning.)

SemaFORR’s *tier-3 Advisors* are heuristic and comment together. They are sensible but imperfect reasons for action selection. There are Advisors for “take any large action,” “go where there is room to move,” “go to unfamiliar locations,” “turn to avoid a nearby obstacle,” and “get closer to the target.” When multiple robots navigate simultaneously, there are also Advisors for “go where there is less traffic.”

SemaFORR also has tier-3 Advisors that comment based on learned spatial abstractions. Some Advisors prefer high-scoring conveyors, particularly those further away. Other Advisors select likely subsequences of a trail, whose decision points serve as attractors, not as a plan. Still other Advisors exploit the skeleton; they support entrance into a region that contains the target, and exit from or avoidance of leaf regions that do not.

Voting sums the comment strengths of all tier-3 Advisors for each possible action, and selects the action with maximum support. Ties are broken at random. See (Epstein et al. 2015) for further details on the reasoning mechanism.

Method

We test navigation in three worlds with different connectivity and a single robot. The task is to *visit* (come within ϵ of) the locations of 5 sets of 40 randomly generated, randomly ordered targets in each world. There is a 250-decision limit per target. Performance metrics include distance travelled, percentage of targets reached, and total time to sense, decide, move, and learn. We average performance over 5 executions of each set of 40 targets, 1000 targets per world in all. Differences cited here are statistically significant at $p < 0.05$.

We compare SemaFORR to *SemaFORR-A**, an ideal navigator for this task. *SemaFORR-A** has an accurate map of the world, discretized by a grid. It plans a shortest path within that map as a sequence of *waypoints*, selects only the smallest moves and turns (to reduce actuator error), and replans when waypoints become inaccessible. To determine how spatial abstractions facilitate navigation, we also tested several *restricted versions* of SemaFORR with reduced sets of tier-3 Advisors. Each restricted version included all commonsense Advisors plus Advisors for only one spatial abstraction: regions, trails, or conveyors.

Results

As Table 1 shows, SemaFORR quickly produces efficient travel without planning or a map. SemaFORR rarely failed to reach a target, and in

Table 1: Performance means and standard deviations in three worlds. SemaFORR-A* plans from a map; SemaFORR is reactive and learns spatial abstractions instead. SemaFORR-T, the version restricted to trails, uses commonsense, exploration, and one spatial abstraction

	SemaFORR-A*		SemaFORR		SemaFORR-T	
<i>World A</i>	μ	σ	μ	σ	μ	σ
Time	1035.89	8.66	1221.19	150.42	1163.75	30.71
Distance	400.06	13.30	854.49	35.73	812.95	30.39
Success rate	100.00	0.00	99.50	2.23	99.50	2.23
<i>World B</i>	μ	σ	μ	σ	μ	σ
Time	884.58	6.02	835.31	92.54	867.23	25.19
Distance	335.14	10.40	554.93	24.93	612.59	26.97
Success rate	100.00	0.00	99.70	1.73	99.80	1.41
<i>World C</i>	μ	σ	μ	σ	μ	σ
Time	1119.93	10.70	1273.69	124.38	1277.95	27.51
Distance	437.87	12.77	798.18	27.18	775.72	27.35
Success rate	100.00	0.00	99.80	1.41	99.60	2.00

world B it is just as fast as SemaFORR-A*. (The apparent speedup in world B is not statistically significant, as are any other differences unaddressed in this section.) In worlds A and C, SemaFORR is only 18% and 14% slower than SemaFORR-A*, respectively. Both systems devote most of their time to movement: 81% for SemaFORR-A* and 82% for SemaFORR. Learning required less than 1% of SemaFORR’s time in every world.

Differences in variance between SemaFORR-A* and SemaFORR are attributable in part to actuator error. Recall that SemaFORR-A* takes only the smallest possible actions, and is therefore subject to less egregious errors. This presumably accounts for much of the difference in their distances as well.

The learned spatial models clearly capture most rooms and hallways in all three worlds. Figure 2 shows models learned from one execution for one setting in each world. Despite actuator error and randomized targets, models learned for any given world are quite similar across settings.

In all three worlds, the restricted version with trails (*SemaFORR-T*) outperformed the other restricted versions. It was faster and traveled less far than the restricted version with regions in world C. SemaFORR-T also always outperformed the restricted version with conveyors, except on time in world C.

We also (nearly) evenly partitioned 40 targets among 3 robots and had them navigate simultaneously. In preliminary testing, their individual learned models are remarkably similar. Figure 3 shows three such models in world A.

Related Work

Human navigators without a metric map do not construct a mental one (Tversky 1993; Zetsche et al. 2009). SemaFORR’s Advisors exploit research on how people perceive, envision, describe, and navigate through space (e.g., (Golledge 1999)). FORR’s use of multiple Advisors mirrors people’s reliance on multiple wayfinding strategies to select

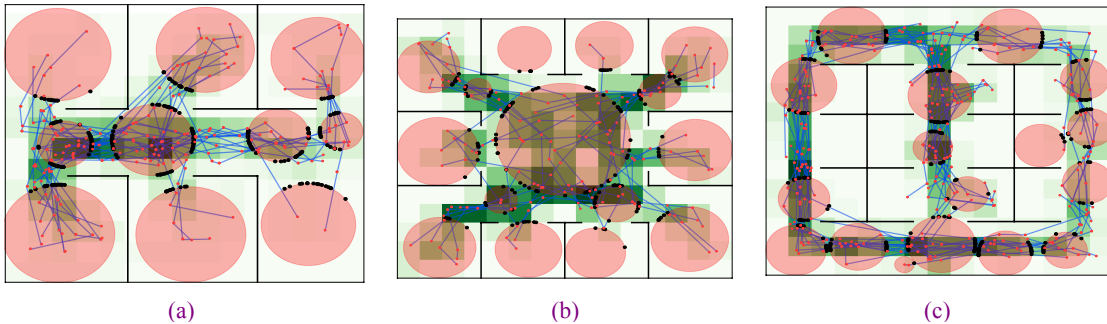


Fig. 2 After navigation to 40 targets, one robot’s learned spatial models, superimposed on their respective maps. Regions are shown as circles with dots along their perimeters for exits. Conveyors are grid squares; darker ones were used more often in successful travel. Trails are lines with dots at decision points. (a) World A’s office space, (b) world B’s rotunda, and (c) world C’s warehouse. Note the learned diagonal conveyor “hallways” in world B and the emphasized perimeter in world C.

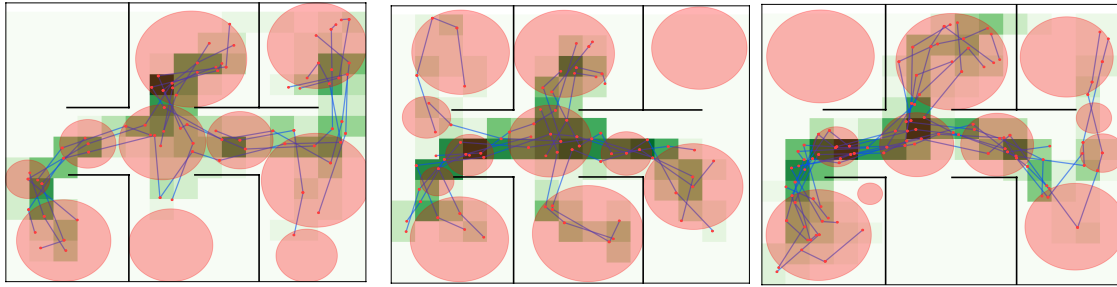


Fig. 3 Spatial models for world A, learned simultaneously by three robots, each with 13 or 14 targets and its own copy of SemaFORR

routes (Takemiya and Ishikawa 2013; Tenbrink et al. 2011).

SemaFORR's wall register is similar to human reference frames (Battles and Fu 2014; Meilinger 2008). Its use of regions (Hölscher et al. 2011; Reineking et al. 2008), conveyors (Meilinger 2008), and trails (Hamburger et al. 2013) is based on human behavior, as is its penchant for exploration (Speekenbrink and Konstantinidis 2014).

Conclusion

Because SemaFORR's learning is both heuristic and dependent on experience, its models may be incomplete and overlook or overly emphasize parts of a world. One subject of current work is a shared model, produced and accessed by a team of robots. This is appropriate for the robots of *HRTeam* (Human-Robot Team) which SemaFORR is ultimately intended to operate (Sklar et al. 2011). The Blackfins in our laboratory now each make decisions with their own copy of SemaFORR-A*.

SemaFORR will eventually explain its decisions to the human member of *HRTeam*. Advisors that support a decision provide readily understandable reasons (e.g., "let's go this way because the other is a dead-end and this should bring us closer to the target"). In summary, robots can learn to navigate well from local perception. With SemaFORR, a robot quickly becomes proficient, and produces a world model that provides important information, both for people and for other robots.

Acknowledgments

This work was supported in part by the National Science Foundation under #IIS-1117000. We thank A. Tuna Ozgelen for the world maps.

References

Battles A, Fu W-T (2014) Navigating Indoor with Maps: Representations and Processes. Paper presented at the 36th Annual Meeting of the Cognitive Science Society, Quebec City.
Epstein SL (1994) For the Right Reasons: The

FORR Architecture for Learning in a Skill Domain. *Cognitive Science* 18:479-511.

Epstein SL, Aroor A, Evanusa M, Sklar E, Parsons S (2015) Navigation with Learned Spatial Affordances. Paper presented at the 37th Annual Conference of the Cognitive Science Society, Pasadena.

Golledge R, G. (1999) Human Wayfinding and Cognitive Maps. In: Golledge R, G. (ed) *Wayfinding Behavior*. Hopkins University Press, pp 5-45.

Hamburger K, Dienelt LE, Strickrodt M, Röser F (2013) Spatial cognition: the return path. Paper presented at the 35th Annual Conference of the Cognitive Science Society, Berlin.

Hölscher C, Tenbrink T, Wiener JM (2011) Would you follow your own route description? Cognitive strategies in urban route planning. *Cognition* 121:228-247.

Meilinger T (2008) The Network of Reference Frames Theory: A Synthesis of Graphs and Cognitive Maps. Paper presented at Spatial Cognition VI, Freiburg, Germany.

Reineking T, Kohlhagen C, Zetsche C (2008) Efficient Wayfinding in Hierarchically Regionalized Spatial Environments. Paper presented at Spatial Cognition VI, Freiburg, Germany.

Sklar EI, Ozgelen AT, Munoz JP, Gonzalea J, Manashirov M, Epstein SL, Parsons S (2011) Designing the *HRTeam* Framework: Lessons Learned from a Rough-and-Ready Human/Multi-Robot Team. Paper presented at the Autonomous Robots and Multirobot Systems workshop, Taipei, Taiwan.

Speekenbrink M, Konstantinidis E Uncertainty and exploration in a restless bandit task. Paper presented at the 36th Annual Conference of the Cognitive Science Society, Austin. pp 1491 - 1496.

Takemiya M, Ishikawa T (2013) Strategy-Based Dynamic Real-Time Route Prediction. In: Tenbrink T, Stell J, Galton A, Wood Z (eds)

- COSIT 2013. LNCS 8116. Springer, pp 149-168.
- Tenbrink T, Bergmann E, Konieczny L (2011) Wayfinding and description strategies in an unfamiliar complex building. Paper presented at the 33rd Annual Conference of the Cognitive Science Society, Boston.
- Tversky B (1993) Cognitive Maps, Cognitive Collages, and Spatial Mental Models. In: Frank AU, Campari I (eds) Spatial information theory: A theoretical basis for GIS, Lecture Notes in Computer Science 716. Springer-Verlag, Berlin, pp 14-24.
- Zetzsche C, Galbraith C, Wolter J, Schill K (2009) Representation of Space: Image-like or Sensorimotor? *Spatial Vision* 22:409-424.