
Freespace Supports Metacognition for Navigation

Susan L. Epstein

Department of Computer Science

Hunter College and The Graduate Center of The City University of New York

New York, NY 10065

susan.epstein@hunter.cuny.edu

Abstract

In a new environment, people identify, remember, and recognize where they can comfortably travel. This paper argues that a robot navigator too should learn and rely upon a mental model of unobstructed space. Extensive simulation of a controller for an industrial-strength robot demonstrates how metacognition applied to a model of unobstructed space resolves some engineering challenges and provides resilience in the face of others. The robot plans and learns quickly, considers alternative actions, takes novel shortcuts, and interrupts its own plans.

1 Introduction

Despite early successes with a purely reactive autonomous robot navigator [1], the complexities of even simple indoor spaces soon indicated that more was required. AI researchers introduced planning, which requires a model of the robot’s environment. Roboticists responded with SLAM, a suite of control methods that simultaneously generates a map and identifies the robot’s precise location as allocentric coordinates within it (*localization*) [2, 3]. SLAM mapping, however, is slow because it focuses on architectural detail and repeatedly seeks to resolve conflicting sensor reports probabilistically. A traditional navigation controller feeds the resultant metric map into either a graph search or randomized search algorithm to formulate a plan that avoids the detected obstacles [4, 5]. Such planning requires a fine-grained map plus flawless sensors and actuators. In contrast, even ferrets learn their way around a new environment quickly and take novel shortcuts [6], without flawless actuators and sensors.

Our position is that the relevant ground truth for navigation is not obstacles but *freespace*, the spatial affordances that facilitate movement. Early work demonstrated that rodents appear to learn a *cognitive map*, a relational model of their environment [7]. Neuroscientists have since identified structures in the hippocampus and entorhinal cortex that support mapping and localization, record distances between locations, and represent space and time during recall [8–16]. Asked to sketch a large, complex, virtual-reality world through which they had just navigated, human subjects’ drawings were more topological than metric [17, 18]. They focused on connectivity that only roughly approximated distances and angles.

We argue here that a freespace model naturally supports *metacognition*, in particular, *when and how to exploit problem-solving strategies* [19]. Our experiments place an industrial-strength robot in a large, complex indoor space (henceforward, *world*) at an architecturally-intended entry point (e.g., an elevator door). The robot first explores a new world, and is then tasked with an ordered, user-specified set of randomly selected targets to visit, each within a fixed *step limit* of decisions. The worlds are detailed floorplans of museums and academic buildings that are notorious for their ability to perplex human visitors. Their tight corners, interior columns, and narrow doorways also challenge the robot. The robot’s controller, built for ROS [20], learns and exploits a freespace model that generalizes over its percepts retrospectively. With metacognition and algorithms honed for speed, it assesses its

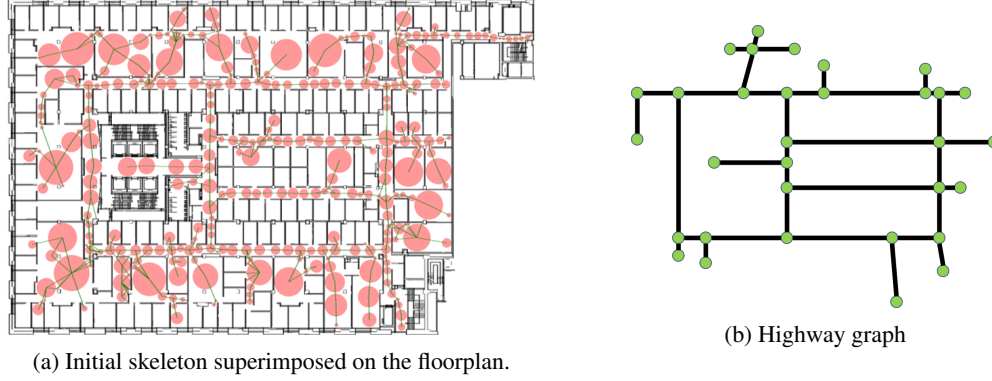


Figure 1: Freespace representations for the fifth floor of a $110 \times 70m$ academic building.

own performance, appropriately defers decisions, intervenes during execution of its own plans, and continues to refine its model of its world as it becomes a reliable navigator.

2 Perception and planning with the freespace model

Our robot’s only sensor is a laser rangefinder mounted slightly above floor level. At each decision point, it provides a *view*, a vector of 660 distances to the nearest obstruction within a $25m$ range and across a human-like 220° field of view. Given a view, our robot controller, SemaFORR, learns a *region*, a circle of freespace whose center is the robot’s location and whose radius is the minimum detected distance. Views are saved until the robot completes travel to a target. SemaFORR then retrospectively revises the traveled path into a *trail*, a simplified sequence of decision points in freespace that would have been consecutively visible. Together, the regions and *subtrails* (trail segments that connect consecutively visited regions) form a *skeleton*, as in Figure 1(a). SemaFORR curates regions to keep them disjoint. It also identifies points where the robot crosses a region’s circumference, and generalizes over those close together on the same region to form a *door*, an arc in freespace.

Landmarks are often cited as significant to human navigators [21]. Even when no longer present (e.g., repurposed or demolished), some landmarks are often included by people in directions as if they were still there. Those landmarks are significant because their location connects freespace. SemaFORR learns such a landmark as a *conveyor*, a small square of freespace. In a grid superimposed on the world’s footprint, a conveyor is a cell frequently intersected by trails. (For additional images and associated learning algorithms, see [22].)

Because planning requires a world model, SemaFORR explores to develop a preliminary version before it travels to user-specified targets in a new world. This one-time active learning is modeled on young animals; when set down in a large novel space for the first time, they follow extended walls [23]. SemaFORR’s interpretation of this behavior is to pursue exceptionally long distances reported in its view, ones orthogonal to the robot’s forward orientation. As the robot travels along such a *candidate*, SemaFORR notes and curates additional possibilities for subsequent exploration [24]. This self-directed, time-limited travel produces a *highway graph*, such as Figure 1(b). It represents long, relatively narrow extents in freespace that overlap at *intersections*. Both the highway graph and the skeleton are orders of magnitude smaller than a traditional planning graph based on a sufficiently fine occupancy grid. Thus, graph-search planning is faster.

Neurological evidence suggests that the human brain represents navigational plans in small spatial networks hierarchically [25], and represents continuous navigation experience as a sequence of discrete abstract concepts that may vary individually but maintain their fundamental nature and chronological order [26]. SemaFORR models its world as a hierarchy of continuous freespace elements. A plan is an ordered list drawn from the highest levels of that hierarchy, a sequence of opportunities in freespace. Given a target, SemaFORR initially plans in the highway graph, from the intersection nearest the robot’s current location to the intersection nearest the target. Because intersections are sparse, SemaFORR can embellish a plan with entry and exit details from the evolving,

more detailed skeleton. Once a plan to the target is established, SemaFORR executes it with a limited action set of forward moves and clockwise and counterclockwise rotations. In a sense-decide-act loop, it retrieves a view, selects an action to address the current plan step, and sends motor commands to execute that action.

SemaFORR makes decisions with a **cognitive architecture that integrates multiple reasoning methods** [27]. It begins with a sequence of reactive rules: if the target is directly in front of the robot, move toward it; do not further consider any move that would come too close to an obstacle or reorient the robot in a recent direction; otherwise, execute the next step in the plan. When multiple actions would satisfy the next plan step (e.g., forward moves of different lengths), the architecture consults a set of heuristics that score those actions by their idiosyncratic preferences. Some heuristics are commonsense (e.g., “get closer to the target”) while others exploit the spatial model (e.g., “use a door to enter that next region”). The architecture selects any action with maximum heuristic support. (See [22] for further details.)

3 Metacognition contends with plan failure

When the percepts and locations that a planner anticipates do not occur, its plan can fail. Sensors and actuators are subject to environmental interference (e.g., reflection or surface changes) and to perceptual errors, some of which may become embedded in the model itself. Our experiments deliberately introduce random, realistic sensor and actuator errors. Thus, even if SemaFORR believes it has a correct plan and can execute its next action, that may not be the case.

SemaFORR’s use of freespace supports metacognition that contends with plan failure in four ways: context-sensitive action selection, operationalization with a spatial hierarchy, local exploration, and reactive planning. SemaFORR is resilient to plan failure in part because, rather than traditional discrete waypoints, its plan steps are continuous freespace elements and so can be accomplished in infinitely many ways. For example, when the next plan step is a region, any location in it is acceptable. This metacognition conserves computational resources; it anticipates minor sensor and actuator errors and postpones decisions until execution provides relevant, and possibly more accurate, information.

If SemaFORR cannot immediately satisfy a plan step, it *operationalizes* it, that is, it replaces the step with a sequence of freespace structures lower in the model’s hierarchy. Any substitution serves as the current step only until a higher-level freespace element in the plan is sensed. If no action satisfies an intersection, SemaFORR replaces that step with a highway that connects to it. If no action satisfies a highway, by construction there is a sequence of regions equivalent to that highway. SemaFORR substitutes that sequence and, as long as possible, moves the robot from one region there to the next. If at some decision point the robot’s view forbids movement into the next region, SemaFORR replaces it with the subtrail known to connect the current region to the next one. SemaFORR can then move the robot to any point on that subtrail and follow along it to reach some location in the next region. Note that only the lowest level of this hierarchy is discrete freespace; all other levels provide multiple options.

Operationalization is opportunistic metacognition because it makes real-time decisions in context. As implemented in SemaFORR it not only provides more detailed guidance at run time, but also recognizes and takes “novel shortcuts” that disrupt plan execution, an ability often cited as a hallmark of intelligence [6]. A novel shortcut travels along the third side of a triangle when the navigator has only previously traveled along the other two. SemaFORR takes novel shortcuts in region sequences. If, for example, the current plan specifies movement from region R_1 to R_2 and then to R_3 , the robot in R_1 may actually sense a point in R_3 . If so SemaFORR will ignore R_2 in its plan and move the robot directly into R_3 .

Local exploration is SemaFORR’s metacognitive response to knowing what it does not know. In such large worlds, some targets go uncovered by the initial freespace model, particularly early ones. As a result, SemaFORR may only be able to formulate a partial plan, one that attempts to reach the vicinity of the target but stops short of it. After SemaFORR executes such a plan it can rely only on its rules and commonsense heuristics. Local exploration then intervenes. It identifies candidates that come near the target, left over from initial exploration or glimpsed from views during the just-executed plan. This target-driven exploration is somewhat naive, since “near by” might well be on the other side of a long wall. Local exploration halts when it reaches the step limit or the target. SemaFORR

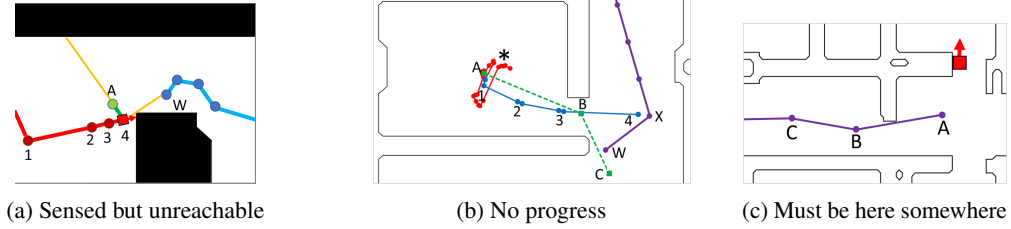


Figure 2: Metacognitive responses to three common situations

then refines all freespace elements (except the highway graph) based on what it has experienced and addresses the next target.

Reactive planning is metacognition that intervenes in the current plan when triggered by a particular situation. SemaFORR has three such planners. Each of them repositions the robot one action at a time to extricate itself from a recognized difficulty, and cedes control once some action would advance a future step in the current plan. The first reactive planner addresses a plan step that the robot can sense but to which SemaFORR forbids movement. In Figure 2(a), while there was freespace in front of the robot, SemaFORR repeatedly moved it toward W along the (red) path $1 \rightarrow 2 \rightarrow 3 \rightarrow 4$. At 4 the robot could still sense W (orange ray) but all forward moves would have come too close to the black obstacle. This reactive planner signaled and, in a sequence of actions, repositioned the robot to A along the ray to its left, from which it could reach W . At A SemaFORR then resumed control, turned right, and moved toward W . The second reactive planner addresses repeated confinement in a small area. When recently sensed freespace fails to increase, this reactive planner signals SemaFORR. Once in control, it rotates 360° to identify new freespace. If it senses none, it extracts the robot from the confined space along a subtrail derived from the experiment’s path history. In Figure 2(b), the robot could not sense W on its current (purple) plan and had been near the asterisk for many (red) decision points. This reactive planner gained control and intended to take the (green) path $A \rightarrow B \rightarrow C$ to escape. The robot easily reached A and moved toward B from 1 to 3, but just after 3 it sensed X , a step later than W on its original plan. SemaFORR resumed control and moved the robot toward X . The third reactive planner addresses the robot’s limited field of view. SemaFORR’s planner only addresses the location of the robot, not its orientation. Thus a plan may specify nearby freespace that the sensor does not detect. In Figure 2(c), the robot’s (purple) plan $A \rightarrow B \rightarrow C$ began close by, at A , but the robot had its back to it. This reactive planner triggered, rotated the robot 90° to its right, sensed A , and ceded control to SemaFORR.

4 Discussion

Metacognition supports effective navigation. SemaFORR’s *accuracy* is gauged by the fraction of targets reached within the step limit. In Figure 1’s world, for example, after 30 minutes of initial exploration, SemaFORR achieved 91.2% *accuracy* within 750 decision steps. On average and including planning time, the robot traveled 83.5m in 94.2 seconds to reach each target, despite a 1m/sec speed limit. In extensive ablation experiments, planning and exploration proved necessary, but it was metacognition that enabled SemaFORR to surpass 70% accuracy [28].

For cognitive plausibility, easy problems should be easy, and harder ones should take a little longer. SemaFORR’s accuracy reflects this. On M5, the fifth floor of New York’s Museum of Modern Art, we gave SemaFORR a 500-step limit and 30 minutes to explore 1585m² of freespace. On 5 target sets in 10 trials, SemaFORR averaged 99.8% accuracy. All but one of our other challenging spaces, including Figure 1, were about 3000m² of freespace, but with substantively different footprints and layouts. Accuracy in them was about 90%. Scaling up, we challenged SemaFORR with the first floor of New York’s Metropolitan Museum of Art, which has 29,707m² of freespace. The robot’s action set and speed remained unchanged; only its initial exploration time and step limit were increased. SemaFORR’s accuracy remained near 90%.

SemaFORR distinguishes between what it experiences and what it has not experienced but believes to be true, such as the difference between freespace it has traversed along its paths and the freespace it surmises from finitely many sensor readings along the way. It speculates about what it perceives based on knowledge about sensor failure, as it does in Figure 2(c). It is sensitive both to its own

inability to progress in Figure 2(b) and to discrepancies between its perception and its ability to move in Figure 2(a). When the planner produces only a partial plan to reach a target, SemaFORR executes it, but recognizes that it is ignorant about the area around its target and begins local exploration.

Given robots’ endemic sensor and actuator errors, precision and complete knowledge of a space may be less important than mapping proponents would have us believe. SemaFORR’s only failure to reach a target in M5 arose from one sensor error late in travel. Moreover, even when initial exploration erred (e.g., the gap in the hallway between the elevators in Figure 1), subsequent target-directed travel soon corrected the model. Although SemaFORR’s final models cover only about half the freespace in these worlds, the robot still reaches about 90% of its targets within the step limit. How the freespace of a world is connected is what matters to a navigator, as long as it can reflect and act on what it knows, what it does not know, what it might discover, and what is reasonable to believe.

Acknowledgements

This work was supported in part by The National Science Foundation under CNS-1625843 and CF-1231216. The author thanks Raj Korpan for his substantial contributions to the development of SemaFORR and Anoop Aroor for the development of its simulation environment.

References

- [1] R. A. Brooks. Intelligence without representation. *Artificial Intelligence*, 47:139–159, 1991.
- [2] John J. Leonard and Hugh F. Durrant-Whyte. Mobile robot localization by tracking geometric beacons. *IEEE Transactions on Robotics and Automation*, 7(3):376–382, 1991.
- [3] Hugh Durrant-Whyte and Tim Bailey. Simultaneous localization and mapping: part i. *IEEE Robotics & Automation Magazine*, 13(2):99–110, 2006.
- [4] Steven M. Lavalle. *Planning Algorithms*. Cambridge University Press, 2006.
- [5] Sertac Karaman and Emilio Frazzoli. Sampling-based algorithms for optimal motion planning. *The International Journal of Robotics Research*, 30(7):846–894, 2011.
- [6] Noam Miller. Taking shortcuts in the study of cognitive maps. *Learning & Behavior*, 49(3):261–262, 2021.
- [7] Edward C. Tolman. Cognitive maps in rats and men. *Psychological Review*, 55(4):189–208, 1948.
- [8] Joshua Jacobs, Christoph T Weidemann, Jonathan F Miller, Alec Solway, John F Burke, Xue-Xin Wei, Nanthia Suthana, Michael R Sperling, Ashwini D Sharan, Itzhak Fried, and Michael J. Kahana. Direct recordings of grid-like neuronal activity in human spatial navigation. *Nature Neuroscience*, 16(9):1188–1190, 2013.
- [9] Liora Las and Nachum Ulanovsky. Hippocampal neurophysiology across species. In *Space, Time and Memory in the Hippocampal Formation*, pages 431–461. Springer, 2014.
- [10] J. O’Keefe and J. Dostrovsky. The hippocampus as a spatial map: Preliminary evidence from unit activity in the freely-moving rat. *Brain Research*, 34:171–175, 1971.
- [11] Torkel Hafting, Marianne Fyhn, Sturla Molden, May-Britt Moser, and Edvard I Moser. Microstructure of a spatial map in the entorhinal cortex. *Nature*, 436(7052):801–806, 2005.
- [12] Trygve Solstad, Charlotte N. Boccara, Emilio Kropff, May-Britt Moser, and Edvard I. Moser. Representation of Geometric Borders in the Entorhinal Cortex. *Science*, 322(5909):1865–1868, 2008.
- [13] Francesca Sargolini, Marianne Fyhn, Torkel Hafting, Bruce L McNaughton, Menno P Witter, May-Britt Moser, and Edvard I Moser. Conjunctive representation of position, direction, and velocity in entorhinal cortex. *Science*, 312(5774):758–762, 2006.
- [14] Lindsay K Morgan, Sean P MacEvoy, Geoffrey K Aguirre, and Russell A Epstein. Distances between real-world locations are represented in the human hippocampus. *Journal of Neuroscience*, 31(4):1238–1245, 2011.
- [15] Dylan M Nielson, Troy A Smith, Vishnu Sreekumar, Simon Dennis, and Per B Sederberg. Human hippocampus represents space and time during retrieval of real-world memories. *Proceedings of the National Academy of Sciences*, 112(35):11078–11083, 2015.

- [16] L. T Hunt, N. D Daw, P Kaanders, M. A MacIver, U Mugan, E Procyk, A. D Redish, E Russo, J Scholl, K Stachenfeld, C. R. E Wilson, and N Kolling. Formalizing planning and information search in naturalistic decision-making. *Nature Neuroscience*, 24(8):1051–1064, 2021.
- [17] Elizabeth R Chrastil and William H Warren. Active and passive spatial learning in human navigation: Acquisition of survey knowledge. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 39(5):1520, 2013.
- [18] Elizabeth R Chrastil and William H Warren. From cognitive maps to cognitive graphs. *PloS one*, 9(11): e112544, 2014.
- [19] J. Metcalfe and A.P. Shimamura. *Metacognition: Knowing about Knowing*. MIT Press, 1996.
- [20] Stanford Artificial Intelligence Laboratory et al. Robotic operating system. URL <https://www.ros.org>.
- [21] Reginald G Golledge. Human wayfinding and cognitive maps. In *Wayfinding Behavior: Cognitive Mapping and Other Spatial Processes*, pages 5–45. The Johns Hopkins University Press, 1999.
- [22] Susan L. Epstein and Raj Korpan. Planning and explanations with a learned spatial model. In Sabine Timpf, Christoph Schlieder, Markus Kattenbeck, Bernd Ludwig, and Kathleen Stewart, editors, *COSIT*, volume 142 of *LIPICs*, pages 22:1–22:20. Schloss Dagstuhl - Leibniz-Zentrum für Informatik, 2019.
- [23] Elizabeth S Spelke and Sang Ah Lee. Core systems of geometry in animal minds. *Philosophical Transactions of the Royal Society B: Biological Sciences*, 367(1603):2784–2793, 2012.
- [24] Raj Korpan and Susan L. Epstein. Hierarchical freespace planning for navigation in unfamiliar worlds. In Susanne Biundo, Minh Do, Robert Goldman, Michael Katz, Qiang Yang, and Hankz Hankui Zhuo, editors, *Proceedings of the Thirty-First International Conference on Automated Planning and Scheduling, ICAPS*, pages 663–672. AAAI Press, 2021.
- [25] Jan Balaguer, Hugo Spiers, Demis Hassabis, and Christopher Summerfield. Neural mechanisms of hierarchical planning in a virtual subway network. *Neuron*, 90(4):893–903, 2016.
- [26] Chen Sun, Wannan Yang, Jared Martin, and Susumu Tonegawa. Hippocampal neurons represent events as transferable units of experience. *Nature Neuroscience*, 23(5):651–663, 2020.
- [27] Susan L Epstein. For the right reasons: The FORR architecture for learning in a skill domain. *Cognitive Science*, 18(3):479–511, 1994.
- [28] Raj Korpan. *Metareasoning, Opportunistic Exploration, and Explanations for Autonomous Indoor Navigation*. PhD thesis, The Graduate Center, City University of New York, CUNY Academic Works, 6 2021.