

Shortcutting from self-motion signals: quantifying trajectories and active sensing in an open maze

Reviewed Preprint

v2 • September 20, 2024

Revised by authors

Reviewed Preprint

v1 • May 9, 2024

Jiayun Xu, Mauricio Girardi-Schappo, Jean-Claude Béique, André Longtin, Leonard Maler 

Department of Cellular and Molecular Medicine, University of Ottawa, Ottawa, Ontario, Canada, K1H 8M5 •

Department of Physics, University of Ottawa, Ottawa, Ontario, Canada, K1N 6N5 • Departamento de Física,

Universidade Federal de Santa Catarina, 88040-900, Florianópolis, Santa Catarina, Brazil • Brain and Mind Institute,

University of Ottawa, Ottawa, Ontario, Canada, K1H 8M5 • Center for Neural Dynamics and Artificial Intelligence,

University of Ottawa, Ottawa, Ontario, Canada, K1H 8M5

 https://en.wikipedia.org/wiki/Open_access
 Copyright information

Abstract

Animals navigate by learning the spatial layout of their environment. We investigated spatial learning of mice in an open maze where food was hidden in one of a hundred holes. Mice leaving from a stable entrance learned to efficiently navigate to the food without the need for landmarks. We developed a quantitative framework to reveal how the mice estimate the food location based on analyses of trajectories and active hole checks. After learning, the computed “target estimation vector” (TEV) closely approximated the mice’s route and its hole check distribution. The TEV required learning both the direction and distance of the start to food vector, and our data suggests that different learning dynamics underlie these estimates. We propose that the TEV can be precisely connected to the properties of hippocampal place cells. Finally, we provide the first demonstration that, after learning the location of two food sites, the mice took a shortcut between the sites, demonstrating that they had generated a cognitive map.

eLife assessment

This **fundamental** work provides creative and thoughtful analysis of rodent foraging behavior and its dependence on body reference frame-based vs world reference frame-based cues. It compellingly demonstrates that a robust map, capable of supporting taking novel shortcuts, is learned primarily if not exclusively based on self-motion cues, which has rarely if ever been reported outside of the human literature. The work, which will be of interest to a broad audience of neuroscientists, provides a rich discussion about the role of the hippocampus in supporting the behavior that should guide future neurophysiological investigations.

<https://doi.org/10.7554/eLife.95764.2.sa3>

Introduction

Animals must learn the spatial layout of their environment to successfully forage and then return home. Local (1) or distal (2) landmark(s) are typically important spatial cues. Landmarks can act as simple beacons, as stimuli for response learning or for more flexible “place” learning strategies (2–5). Experimental studies addressing this issue have used one (6) or many (7) distal landmarks that provide allocentric coordinates for learning trajectories from a start site to a hidden location. A second source of spatial information derives from idiothetic cues that provide a reference frame via path integration of self-motion signals (8, 9). Research on the interaction of landmark and self-motion cues has concluded that a “spatial map” can be generated by self-motion cues anchored to stable landmarks (8, 10–12). This spatial map can link start sites to important locations via efficient trajectories. However, it is not clear whether such a map is sufficient for flexible navigation using entirely novel routes.

We designed a circular open maze where navigation trajectories are unconstrained, the mice’s allocentric reference frame is determined by large visual cues and food is hidden in one of 100 holes. We found that mice did not use the landmarks but required only minimal cues, namely a stable start location, self-motion signals and active sensing (hole checks), to learn to find food. Trajectories and hole checking strategies dramatically changed as the mice learned to efficiently navigate to the hidden food. Averaged trajectory directions converged to a “mean displacement direction”, and hole checking became concentrated near the expected food hole. These analyses resulted in a computed target estimation vector (TEV) that closely approximated the most direct route between the start location and food. When, after learning, the mice were transferred to another start site, their TEVs also rotated and then pointed to the “rotationally equivalent location” (REL) instead of the target, therefore demonstrating that they were utilizing self-motion cues for navigation. We propose ways to link the mean displacement and hole checking components of the TEV to the properties of hippocampal place cells and suggest a neural equivalent of the TEV that might be computed by downstream hippocampal targets.

The cognitive map hypothesis proposes that animals can learn a metric map of their environment and use it to flexibly guide their navigation, that is, they can take shortcuts, detours and novel routes when needed (7, 13, 14). It is further proposed that the hippocampus and closely connected cortical areas generate a cognitive map (15, 16). We found that, after training to find food in two distinct sites, mice were able to take an unrehearsed shortcut between the sites on a final probe (*i.e.*, no food) trial. The trajectories and hole check analyses demonstrated that the computed TEVs were accurate and efficient. To the best of our knowledge, this is the first demonstration that a stable start location and self-motion cues are sufficient for a rodent to compute a cognitive map. It is not clear whether current theories based on electrophysiological studies of the hippocampus and associated cortices can account for the cognitive map we have observed (15, 17).

Results

We designed a behavioural framework, the Hidden Food Maze (HFM), to tease apart the roles of idiothetic and allothetic cues for navigation in a foraging task (Fig. 1A). Mice are trained to find a food reward hidden inside one hole (the target) out of 100 in a large, circular open maze (120 cm in diameter). The target hole has a “rotationally equivalent location” (REL) in each of the arena’s quadrants (Fig. 1B; see “REL” in Methods for a description). In the mouse’s perspective, the displacement from the entrance location to the target is the same as from the entrance in any quadrant to their respective REL. The HFM has several unique features such as 90° rotational

symmetry to eliminate geometric cues and control for distal cues, movable home cage entranceways, and experimenter-controlled visual cues (see Methods). More importantly, our framework gives the mice freedom to explore the arena with minimum stress, yielding variable trajectories due to the nonstereotyped environment. These features allow us to control whether mice orient to allothetic or idiothetic cues, enabling us to investigate the respective role of landmarks versus path integration in spatial navigation.

We first analyzed the effects of randomizing the entrance location taken by the mouse on successive trials (**Fig. 1C**), versus maintaining a static entrance throughout training (**Fig. 1D**). We quantified the geometric and kinetic features of the trajectories and the behavior of hole checking throughout the arena. The variability of the trajectories resulted in the development of statistical methods to directly infer the mice estimation of the target location (the “target estimation vector”, or TEV). We showed that it coincides with the actual target vector only for the static entrance protocol. In order to eliminate the effect of visual cues, we performed rotated probe trials after training in static entrance, where mice went either to the target (following landmarks) or to the REL (ignoring landmarks; see **Fig. 1D**).

Finally, we investigated the features of the mice trajectories and active sensing behavior when trained on two targets consecutively under the static entrance protocol (**Fig. 1E**). In this case, the TEV method was employed to demonstrate the learning of both target directions and the use of shortcuts, one manifestation of a cognitive map (14). See Methods and the Supplementary Material for the definition of all measured quantities and details on experimental protocols.

Pretraining trials

Initially, naïve mice were introduced to our maze. In agreement with previous work (18), we observed that mice initially explored the maze by following the wall (thigmotaxis) on round trips of increasing complexity before they explored the maze center. Our pretraining protocol induced the mice to explore the arena center (see Fig. S1A, Methods). Thigmotaxis was still observed after pretraining (see, e.g., **Figs. 2A,B** and **3A,B**), though now the mice spent more time in the central regions of the maze. In the Static trial experiments, there was not a distinct class of thigmotactic trajectories, but rather a continuum of trajectory distances from the wall. Furthermore, the mice were initially equally likely to take trajectories at any angle from the start-to-food line as evidenced by the variation of the TEV deviation from the target vector in trial 1 of **Fig. 4I**.

Mice used self-motion sensory input but did not utilize 2D wall cues for spatial learning

Random Entrance experiments

Mice were trained with their cage being moved randomly between quadrant entrances from trial to trial without directly manipulating the animals. Thus, they never entered from the same start position for consecutive trials, making the wall pictures their only reliable orienting landmarks. These cues were 15×15cm black shapes on a white background (see Methods). The configuration of the 4 cues varied relative to the entrance, but the food was always placed in the same hole and nearest the X-shaped landmark (**Fig. 2**, Methods). The food-restricted mice showed persistent complex search trajectories throughout the learning period (**Fig. 2A, B**). The latency to reach the food target was significantly negatively correlated with trial number (statistics in **Fig. 2D**) but with a non-significant difference between the first and last trials ($N=8$; First trial mean±SD = 151.34 ± 148.91 s; Last trial mean±SD = 80.38 ± 43.86 s; $P=0.2602$). Despite the decrease in latency, mice did not appear to exhibit spatial search strategies (**Fig. 2C**) given that the time spent in

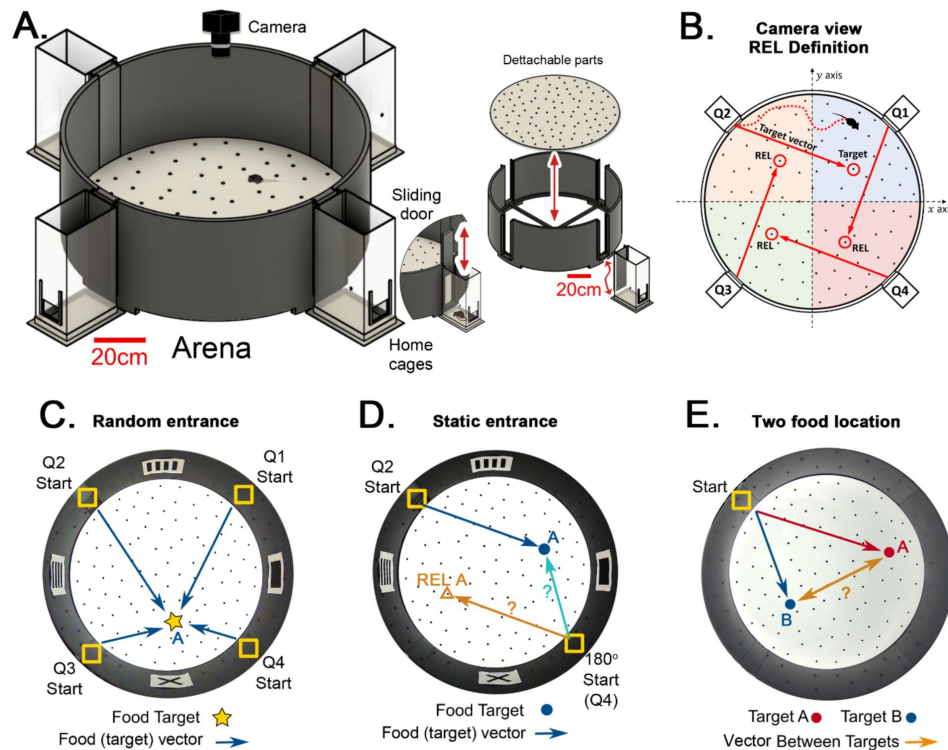


Figure 1.

Hidden Food Maze and experimental setup.

A. The floor of the arena is 120cm in diameter, and the walls are 45cm tall. Note 20 cm scale bar in this panel. The home cage has a 10.5×6.5cm² floor area. The door slides upward (mice enter without handling). The floor is washed and rotated between every trial to avoid predictable scratch marks and odor trails. The subfloor containing food is not illustrated. **B.** Camera view of a mouse searching for hidden food (target, pointed by the target vector). The REL of the target is marked for each entrance (from the mouse's perspective, the displacement from the start in each quadrant to its respective REL is the same for any entrance; e.g., "70cm forward + 30cm to the left"; and it is equal to the displacement from the trained start to target). **C.** "Random entrances" experiment. Mice enter from any of the four entrances randomly over trials to search for food ("A"-labeled star) always in front of the X landmark. Arrows show the four possible displacements. **D.** "Static entrances" experiment. Mice start from the same entrance (labeled "Start") in every trial to search for food in front of the same landmark. Blue/cyan arrows=food vector (start → food); Orange arrow=REL vector (start → REL). After training, the start position in a probe trial can be rotated ("180° Start") to check whether mice follow idiothetic (start → REL; ignoring landmarks) or allothetic (start → A; following landmarks) cues; going via the REL vector is regarded as evidence of path integration. **E.** "Two food location" experiment. Mice start from a static entrance to search for food (red vector to target A). Afterwards, mice are trained to find food in a different location (blue vector to target B). After learning both targets, a probe trial (i.e., a trial without food) is designed to check whether mice can compute shortcuts from B to A (B-A vector, orange arrow).

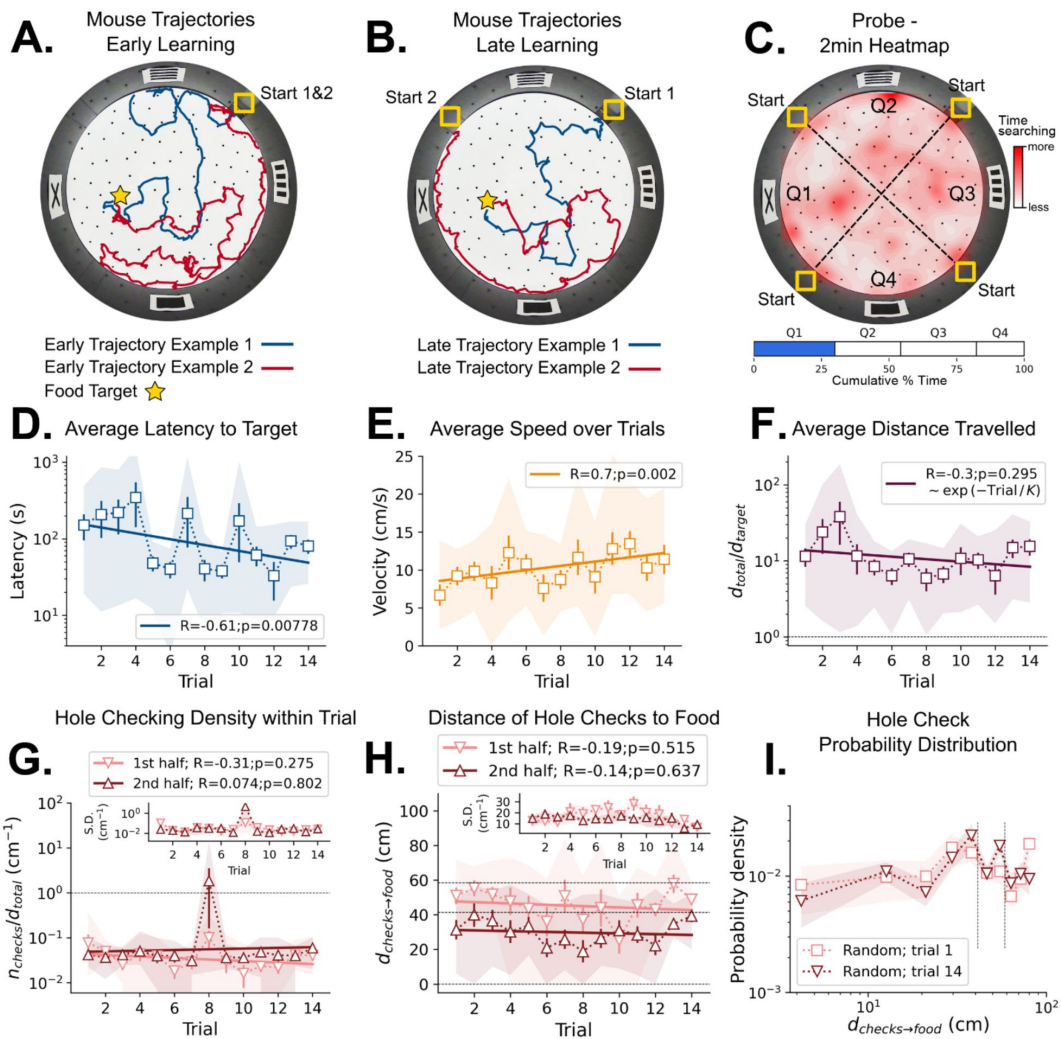


Figure 2.

Mouse spatial learning with random entrances.

A. Two examples of mouse search trajectories during early learning (trial 3) when the entrance changes from trial to trial. They are irregular and vary unpredictably across trials. **A, B.** Star = target location. Yellow Square = entrance site. **B.** Two examples of mouse search trajectories during late learning (trial 14) after starting from different entrances. Trajectories look as irregular as in early trials. **C. (Top)** Heatmap of the first 2min of a probe trial done after trial 18 (red=more time in a given region). **(Bottom)** Mice spent about the same time (25%) in each of the four sectors, regardless of being close to the target (blue) or to its REL (white). **D.** Some significant reduction in latency to reach target is seen across trials ($p = 0.008$; $N = 8$). **D-I:** Error bars = S.E. Shaded area = data range. **E.** Some significant increase in speed is seen ($p = 0.002$; $N = 8$). **F.** Average normalized distance traveled to reach target ($d_{total}/d_{target} = 1$ is optimal; $p = 0.05$, $N = 8$). **G.** Hole-checking density (number of hole checks per distance traveled) in each half of the trajectory. The density remains constant for both halves and across trials, suggesting that mice remained uncertain as to the food location. **G-inset.** The S.D. of the density over the mice sample remains constant for both halves ($N = 8$). **H.** The average distance of the checked holes to the food $d_{(checks \rightarrow food)}$ remains almost constant across trials. Horizontal lines are just guides to the eye. **I.** The probability density of the distance of hole checks to the food $d_{(checks \rightarrow food)}$ for the first and last learning trials (the corresponding averages over trials are in panel H). The density remains unaltered. Vertical dotted lines mark the same distances as the horizontal lines in panel H.

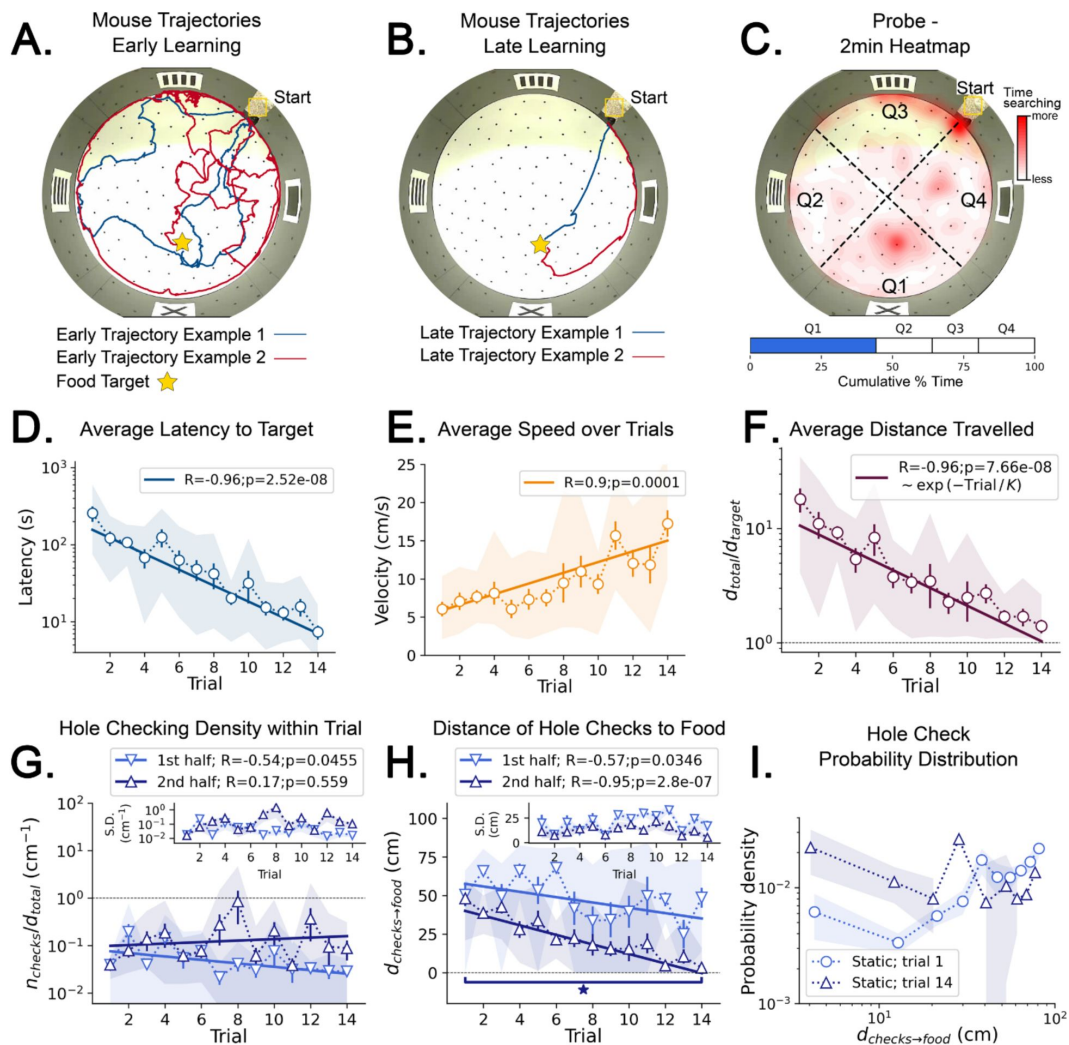


Figure 3.

Mouse spatial learning with static entrances.

A. Two examples of mouse search trajectories during early learning (trial 3). They are irregular and variable similarly to those in the random entrance experiments. **A, B.** Star = target location. Yellow Square = entrance site. **B.** Two examples of mouse search trajectories during late learning (trial 14). They go directly towards the food or go along the wall before turning to the food, creating variation across mice and trials. **C. (Top)** Heatmap of the first 2min of a probe trial done after trial 14 (red=more time in a given region). **(Bottom)** Mice spent almost 50% of the time within 15cm radius of the target (blue) compared to the RELs (white). **D.** Latency dramatically decreases ($p < 10^{-7}$; $N=8$). **D-I:** Error bars = S.E. Shaded area = data range. **E.** Speed significantly increases during trials ($p = 0.0001$; $N=8$). **F.** Normalized distance to reach target ($d_{total}/d_{target}=1$ is optimal) becomes almost optimal ($p < 10^{-7}$; $N=8$). **G.** Hole-checking density over distance in each half of the trajectory. It significantly decreases in the first half ($p = 0.05$), and stays constant in the second. **G-inset:** The S.D. of the density is larger in the second half. **H.** The average distance of the checked holes to the food $d_{(checks \rightarrow food)}$ decreases for both halves of the trajectory. After learning, the hole checks happen closer to the food ($d_{(checks \rightarrow food)}$ is almost zero), although there are more checks per distance. **I.** The probability density of the distance of hole checks to the food $d_{(checks \rightarrow food)}$ for the first and last trials (the corresponding averages over trials are in panel h). After learning (trial 14), the density is larger closer to the food, a feature that does not appear in the random entrance experiments.

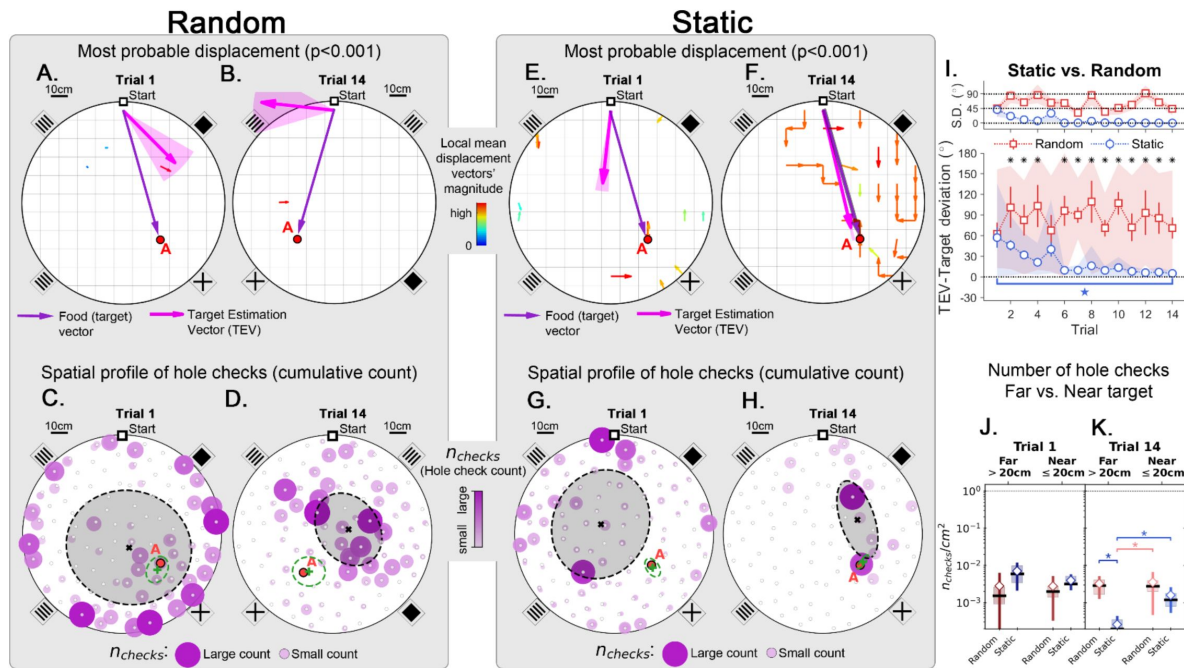


Figure 4.

Trajectory directionality and active sensing for random and static experiments.

Arenas on the top row (mean displacement vector – see color scale between panels b and e) correspond to the ones immediately below them (hole checking spatial distribution); the red “A” label marks the target (food site), which is pointed by the food (target) vector (purple arrow). **Top row (A,B,E,F)**: the color and arrows indicate the most probable route taken (red=more probable; only $p < 0.001$ displacements shown; pink arrow=inferred target position, or TEV; shaded pink sector=S.D. of TEV; see Methods, and Fig. S7). **Bottom row (C,D,G,H)**: spatial distribution of hole-checks; size and color of circles=normalized frequency that a hole was checked (larger pink circles=higher frequency); Black ellipse (x =mean): covariance of spatial distribution. Green ellipse ($+$ =mean): covariance of spatial distribution restricted to ≤ 20 cm of the target. **Random entrance** experiments ($N=8$; panels **A,C**: trial 1; **B,D**: trial 14): regardless of training stage, no significant preferred routes and the TEV does not point to target (**A,B**); hole checks are randomly distributed throughout the arena, and shift from the walls (c) to near the center (d) after learning. **Static entrance** experiments ($N=8$; panels **E,G**: trial 1; **F,H**: trial 14): after learning (f) the TEV and significant displacements go straight to the target (although individual trajectories are variable); and hole checks align along the start-target path (h). **Panel I**: deviation between the TEV (pink arrow) and the target vector (purple arrow) illustrated in top panels. Directionality is quickly learned (static case). **Panels J,K**: hole-check area density corresponding to the spatial profiles in bottom panels. Density after learning is larger near the target (static case), supporting the path integration hypothesis. Asterisks/star: $p < 0.05$ (paired t-test). Note the presence of more significant displacements in late learning for static entrances only, and the associated alignment of the TEV and food vector.

each quadrant during the probe trial is similar (target quadrant=28.78%, other quadrants=24.93%, 28.04%, 18.25%; 1-Way ANOVA: $P=0.070574$, $F\text{-ratio}=2.71621$; see Behavioral Analysis in Methods). This suggests latency decreased for reasons independent of spatial learning.

We found a corresponding significant positive correlation of speed and trial number (**Fig. 2E**) with non-significant differences between first and last trials ($N=8$; First trial mean $\pm$ SD = 6.68 ± 3.79 cm/s; Last trial mean $\pm$ SD = 11.40 ± 5.07 cm/s; $P=0.0704$). Since the distance from start to target (d_{target}) varies over trials due to changing entrance location, we defined the normalized trajectory length (**Fig. 2F**). It is the ratio between total traveled distance (d_{total}) and d_{target} ; $d_{\text{total}}/d_{\text{target}}=1$ is optimal. There is no significant correlation of this quantity and trial number ($P=0.295$). The first and last normalized trajectory lengths were not significantly different either ($N=8$; First trial mean $\pm$ SD = 11.79 ± 8.25 ; Last trial mean $\pm$ SD = 15.51 ± 9.92 ; $P=0.4241$).

Fitting the function $d_{\text{total}} = B \cdot \exp(-\text{Trial}/K)$ reveals the characteristic timescale of learning, K , in trial units (**Fig. 2F**). We obtained $K=26 \pm 24$ giving a coefficient of variation (CV) of 0.92. The mean, $K=26$, is therefore very uncertain and far greater than the actual number of trials. Thus, we hypothesize that the mice did not significantly reduce their distance travelled (**Fig. 2A,B,F**) because they had not learned the food location – the decrease in latency (**Fig. 2D**) was due to its increased running speed and familiarity with non-spatial task parameters.

We next examined hole checking as an active sensing indicator of the mouse's expectation of the food location (see Methods, Fig. S2A,B and the Shortcut Video (Video 1)) for hole-check detection after spatial learning). Even though hole check counts increase with the total trajectory length and latency (Fig. S4B,C), we expected that the mice would increase the number of hole checks near the expected food location, or along their path toward it. Thus, we split the trajectories into two halves of same duration to identify whether hole checks would increase in the second half (being closer to the target). However, the hole checking density (ratio between number of checks and distance travelled) remains constant between the first and second halves of any trial (**Fig. 2G**), suggesting mice did not anticipate where the food is placed. Mice also did not increase their hole sampling rate as they approached the target as shown by the distance of hole checks to the target, $d_{\text{checks} \rightarrow \text{food}}$ (**Fig. 2H,I**), again suggesting no expectation of its location. Mice do not appear to learn a landmark-based allocentric spatial map.

Static Entrance Experiments

We tested whether the mice might acquire an allocentric spatial map if, contrary to the random entrance case, they had a single stable view of the landmarks and could associate one cue configuration with the food location (see Methods). For this, each mouse entered the arena from the same entrance during the training trials (hence “static entrance”). Mice greatly improved their trajectory efficiency after training (**Fig. 3A,B**), with distance and latency to food plateauing around trial 7. During early learning, mice showed random search trajectories that transformed to stereotyped trajectories towards food by late learning (**Fig. 3B**). Typical trajectories were either (i) directly to the food, (ii) initially displaced towards the maze center then returning to a food-oriented trajectory, or (iii) initially along the wall and then turning and heading to the food (**Fig. 3B**).

During the probe trial, mice spent the most time searching in the target quadrant (44.26%; **Fig. 3C**), compared to the time spent in the other quadrants, respectively 19.38% ($P=0.007$), 16.55% ($P=0.0065$), 19.81% ($P=0.007$; p-values obtained for a pairwise t-test and adjusted for multiple comparisons via the Benjamini-Hochberg method). The latency to target strongly decreased over trials (**Fig. 3D**) and the last trial latency was significantly reduced compared to the first ($N=8$; First trial mean $\pm$ SD = 255.39 ± 153.28 s; Last trial mean $\pm$ SD = 7.38 ± 4.03 s; $P=0.004$). The speed dramatically increased over trials (**Fig. 3E**) and the ending speed was significantly greater than the initial speed ($N=8$, First trial mean $\pm$ SD = 6.04 ± 2.48 cm/s; Last trial mean $\pm$ SD = 17.25 ± 4.64

cm/s; $P=0.0002$). The normalized trajectory length strongly negatively correlated with trial number (**Fig. 3F**) and the final normalized length greatly reduced ($N=8$; First trial mean $\pm$ SD = 18.02 ± 11.36 ; Last trial mean $\pm$ SD = 1.40 ± 0.50 ; $P=0.0061$), becoming almost optimal (1).

Now the fitting of the function $d_{total}=B \exp(-\text{Trial}/K)$ yielded $K=5.6\pm0.5$ with a CV = 0.08; the mean is therefore a reliable estimate of total distance travelled. We interpret this to indicate that it takes a minimum number of $K=6$ trials for learning the distance to the target (see also Fig. S4D,E,F,G). Learning is still not complete because it takes 14 trials before the trajectories become near optimal.

Hole checking density varied across trajectories and differed between early and late learning trials. The density significantly decreased with trial number in the first half of trajectories ($R=-0.54$; $P=0.0455$, **Fig. 3G**) but not the second half ($R=0.17$; $P=0.559$, **Fig. 3G**). The hole checks also happened significantly closer to the target in the second half of the trajectory (**Fig. 3H**; First trial mean $\pm$ SD = 48.06 ± 11.50 cm; Last trial mean $\pm$ SD = 3.13 ± 5.70 cm; $P=3\times10^{-6}$). A remarkable effect of learning is that the probability of checking holes increases as the mouse approaches the food (**Fig. 3I**). These results (**Fig. 3H,I**) suggest that the mice may be predicting the food location and checking their prediction as they approach the estimated food location (see also Figs. S4 and S5 for other quantities and scatter plots of trajectory features vs. hole checking). Hole checking near the food site also implies that the mice are not using odorant or visual cues to sense the precise food location, but instead rely on a memory-based estimate of the food-containing hole (see also **Fig. 4G,H**). This motivated us to look at the spatial configuration of these variables to get a detailed picture of the random vs. static condition.

Revealing the mouse estimate of target position from behavior

Although quantitative, the analysis so far only showed a broad comparison between random versus static entrance experimental conditions. We developed a method to map out the directions that the mice stepped towards more frequently from each particular location in the arena, yielding a *displacement map* (see Statistical Analysis of Trajectories in Methods). This map can generate a precise estimate of the mean directionality of the trajectories for each experimental condition (**Fig. 4A,B,E,F**, P -values in Fig. S7A,B,C,D). Additionally, we also made a frequency plot of the checks for each particular hole in the arena, yielding the spatial distribution of hole checks (**Fig. 4C,D,G,H**). These two maps are combined to give the Target Estimation Vector (TEV;). The TEV is interpreted as the mouse's estimate of the target position (pink vectors in **Fig. 4A,B,E,F**).

In random entrance training, no significant directionality was found in the arena (**Fig. 4A,B**), again consistent with a random search. The spatial distribution of hole checks migrated from near the walls towards the center of the arena but remained random and displaced from the target (**Fig. 4C,D**, and also Fig. S8), consistent with the broad results of **Fig. 2**. Conversely, there was a clear and significant flow of trajectories for mice trained with static entrance (**Fig. 4E,F**; $P<0.0001$; see Methods for the definition of the p -values assigned to directions). The hole checks reflected this and became solely concentrated along the route from start to target (**Fig. 4G,H**). In the static entrance condition, the entropy (uncertainty) of the number of hole checks near (<20 cm) the target was lower than that for the far (>20 cm) condition (Fig. S7F,G). In other words, mice were more consistent in the number of their hole checks near the target compared to far from the target, suggesting that they had an internal representation of their proximity to the food site.

We compared the deviation between the TEV and the true target vector (that points from start directly to the food hole; **Fig. 4I**). While the random entrance mice had a persistent deviation between TEV and target of more than 70° , the static entrance mice were able to learn the direction of the target almost perfectly by trial 6 (TEV-target deviation in first trial mean $\pm$ SD = $57.27^\circ \pm 41.61^\circ$; last trial mean $\pm$ SD = $5.16^\circ \pm 0.20^\circ$; $P=0.0166$). A minimum of 6 trials is sufficient for learning both the direction and distance to food (**Fig. 4I**) (**Fig. 3F**) (see Discussion). The kinetics of

learning direction to food are clearly different from learning distance to food since the direction to food remains stable after Trial 6 while the distance to food continues to approach the optimal value.

Under the path integration hypothesis, it is expected that error accumulates as the mouse walks (15 [↗](#)). This motivated us to investigate the frequency of hole checks per area near the target (within 20 cm) vs. far (further from 20 cm away), **Fig. 4J,K** [↗](#). The density of hole checks in both conditions remained constant for random entrance mice, regardless of training duration. However, static entrance mice had significantly higher density of hole checks near the target (**Fig. 4K** [↗](#); P-values: random-far vs. static-far=0.005; static-far vs. random-near=0.03; random-near vs. static-near=0.01). The mean position of hole checks near (≤ 20 cm) the target is interpreted as the mouse estimated target (**Fig. 4C,D,G,H** [↗](#); green + sign=mean position; green ellipses = covariance of spatial hole check distribution restricted to 20cm near the target). This finding together with the displacement and spatial hole check maps (**Figs. 4F** [↗](#) and **4H** [↗](#), respectively) corroborates the heatmap of time spent in the target quadrant (**Fig. 3C** [↗](#)), suggesting a positive correlation between hole checks and time searching (see also Fig. S4C). Thus, we use the distance from this point to the entrance as the magnitude of the TEV. This definition reveals that the TEV is nearly coincident with the direct route between start location and the food-containing hole (pink vectors in **Fig. 4F** [↗](#)), consistent with the near optimal normalized trajectory length previously discussed (**Fig. 3F** [↗](#)).

Ruling out landmarks

These results demonstrated that mice can learn the spatial location of the food target from a static entrance. We next investigated whether the performance gain can be attributed to the cues by manipulating the mouse's entrance. Well-trained static entrance mice were subject to rotation of their home cage by 180°, +90°, or -90° for probe trials in order to start from a different location. We used a different cohort of N=8 mice for each rotation. The mice were not directly handled during the rotation procedure. If the mice were using the distal landmarks, they would still be expected to find the correct food location. If they relied on self-motion cues, they should travel to the REL, since this is the spot where the target would have been if the mouse had been trained from that entrance. In other words, going to the REL means that the trajectory is anchored to the mouse's start point and not to the wall cues. As illustrated in **Fig. 1B** [↗](#), the displacement between each start-REL pair is exactly equal from the mouse's point of view (e.g., going 70 cm forward then 30 cm to the left reaches the respective REL of the target).

Following all rotations, mice searched at the REL instead of the correct location for the food reward regardless of the arena having wall cues (**Fig. 5A,B,C** [↗](#)), meaning that they were not using landmarks to compensate for rotation. In fact, the trajectory kinetics of the 180°-rotated probe for reaching the REL target (mean±SD latency = 4.80 ± 3.06 s; speed = 14.60 ± 5.21 cm/s; normalized distance to REL = 1.96 ± 0.90) is indistinguishable (within one SD or less) from the kinetics of the last trial of static entrance for reaching the true target. The criterion for reaching the REL target was getting within 5cm of its position (the average inter-hole distance is 10cm).

We also show the displacement maps (mean trajectory directionality), hole check spatial profile and TEV for the 180°-rotated probe trial (**Fig. 5C** [↗](#)). The TEV now points to the REL, and the hole checks accumulate along the route to the REL. These maps are also very similar to the static entrance maps in **Fig. 4** [↗](#). Thus, the mice failed to use allothetic cues for navigation in the novel start location test.

Further controls

It is possible that mice, like gerbils (6 [↗](#)), choose one landmark and learn the food location relative to that landmark (e.g., the "X" in **Fig. 3B** [↗](#)); in this case, the mice would only have to associate their initial orientation to that landmark without identifying the image on it. We used two

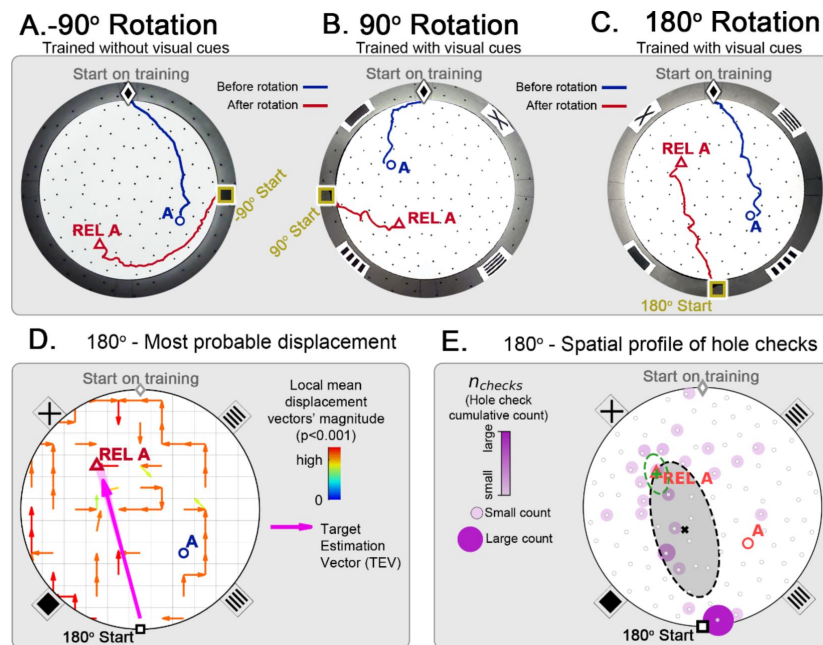


Figure 5.

Changing start position after training in static protocol.

Mice are trained in the static entrance protocol to find food at the target labeled “A” (blue circle), and a probe trial is executed with mice entering from a rotated entrance after 18 trials. Panels **A,B,C** show the comparison between trajectories from the last learning trial (blue) versus the probe (red). The training was performed without landmarks (**A**: N=8, -90° rotation) and with landmarks (**B**: N=8, 90° rotation; **C**: N=8, 180° rotation). In all instances, mice ignored landmarks and went to the REL location (“REL A” label, red triangle), something that is expected under the path integration hypothesis. **Panel D**: trajectory directionality analysis and TEV (pink arrow; shaded sector: S.D.) show that significant paths of all mice ($p < 0.001$; N=8; see Methods) point to the REL-A location in the same way that it pointed to the target without rotated entrance in **Fig. 4F**. **Panel E**: the spatial distribution shows that hole checks accumulate along the start-REL vector, instead of the start-target vector of the case without rotation in **Fig. 4H**. Black ellipse ($\bar{x}$ =mean): covariance of hole check distribution. Green ellipse ($\bar{+}$ =mean): covariance of the data within 20cm of the REL-A location. This suggests that mice follow trajectories anchored to their start location (idiothetic frame of reference).

additional controls to assess this possibility. Mice were trained without any cues on the wall, eliminating the possibility that they provided orienting guides. The learning curve of mice with or without the 2D wall cues were not significantly different (Fig. S6A,B).

Mice that were well trained in the lighted maze were given trials in complete darkness. These mice still showed the same learning curve compared to untrained mice or mice following rotation (Fig. S6C,D).

Finally, to completely rule out both landmark and possible olfactory cues, we trained and tested mice in total darkness. Head direction tuning may be impaired in sighted mice navigating in darkness (19), but spatial learning was not affected in our experiments. Mice can also use localized odor sources as landmarks for spatial learning (20) and we therefore attempted to eliminate such cues (see Discussion and Methods). The mice were still able to learn the food location in the absence of visual and olfactory cues (Fig. S6E,F). These experiments demonstrate that the mice can learn to efficiently navigate from a fixed start location to the food location using path integration of self-motion cues.

Furthermore, the location of the targets in all instances of the experiment was carefully chosen to avoid cues from the arena's circular geometry. If the target was placed near the center of the arena, the learning task would be trivial and would involve little spatial inference. Conversely, if the target was placed too close to the wall, thigmotaxis would be sufficient to get near the food. Also, the line going through the center of the arena that is perpendicular to the line connecting start to center has to be avoided (see Fig. 1B for reference), otherwise the system would be symmetric to rotation or could be simplified to a left-or-right choice that would involve little spatial learning. Finally, the target cannot be placed near the start positions. Thus, in the static entrance experiments, the chosen target position is always more than 35cm away from the wall, and more than 60cm away from the start. In the random entrance case, the target is placed more than 30cm away from the wall, and more than 40cm away from the closest starting location.

Mice can use a shortcut to navigate between remembered targets

Our results established that mice learn the spatial location of a food reward using path integration of self-motion cues. Have the mice learned a flexible cognitive map? A test of cognitive mapping would be if mice could take a short cut and navigate from one remembered location to another along a novel, unreinforced path. The “two food location experiment” set out to test this possibility by training mice to travel to food from their home cage to two very differently located sites (chosen according to the rules mentioned in the previous section). The two locations were trained sequentially so that one site was trained first (target A) followed by training on the second site (target B). A probe trial without food after training target A was designed to check if mice had learned it successfully. Only one hole was filled with food in any given training trial. In a second probe trial after training on B, mice were tested to see whether they can travel between the two remembered locations (food is absent; probe B-A trial) despite no prior training nor reinforced experience on the short cut route (Fig. 1E).

Mice successfully learned the location of target A, evidenced by their direct search trajectories and significant decreases in distance traveled (Fig. 6A), with the same performance as in the static entrance case. During the learning of target A, the distance of the mouse to the future location of target B was always larger than that expected by chance, *i.e.*, to a randomly chosen location in the maze (see Methods, Fig. S10A,B). In fact, the trajectories stayed, on average, approximately 40 cm away from the forthcoming target B position. In other words, the mice did not learn, by chance, direct unreinforced routes from “near B to A”. In Fig. 6A, for example, the mouse passed within 16 cm of the future target B site in an early learning trial, but this close approach does not appear to constitute a B->A trajectory. The first probe trial (without food) confirmed that mice had learned the location of target A. After not finding food in A, the mice re-initiated a random search (Fig. 6B).

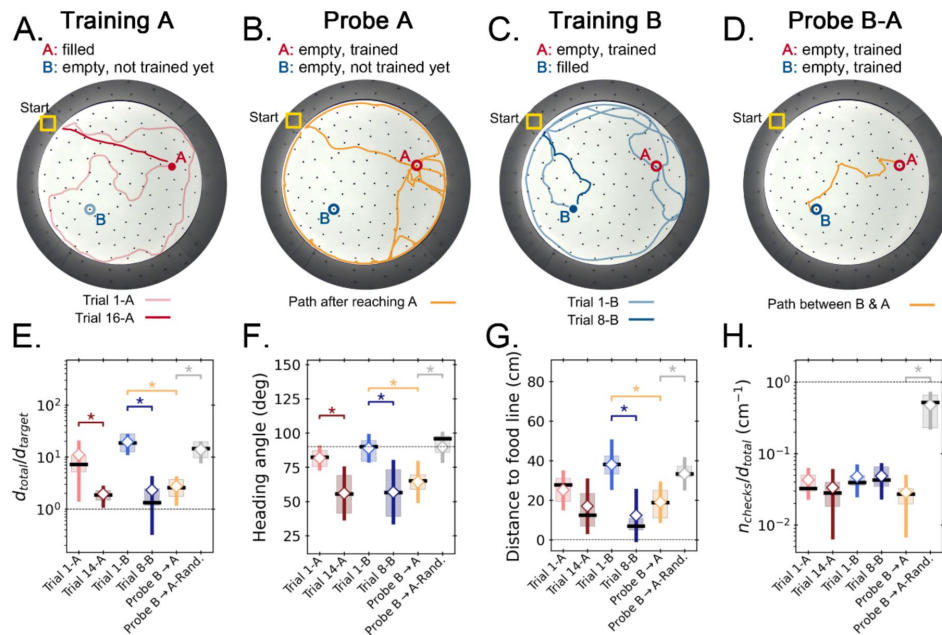


Figure 6.

Two food location experiment.

Panels A-D. Trajectory exemplars of four sequential stages of the experiment (all trials done with static entrance, $N=8$): **(A)** training target A (trials 1-A and 16-A for early and late learning, respectively); **(B)** probe A (no food is found, triggering a random search); **(C)** training target B (keeping A empty; trials 1-B and 8-B for early and late learning); **(D)** probe B-A (where both targets are trained and empty, and the mice take a shortcut from B to A; see Fig. S8 for all exemplars). Filled circles=filled target hole; empty circles=empty target. In the A and B learning stages, the trajectories evolve from random to going straight from start to the respective target. **Panels E-H.** Standard boxplot statistics of learning versus probe (diamonds are averages; asterisks: $p<0.05$ in a paired t-test comparison). Quantities are defined in Fig. S4A. Significant differences between early and late learning were observed for the traveled distance **(E)**, heading angle **(F)**, and distance to the food line **(G)**. Density of hole checks **(H)** remained nearly constant, as expected. In all instances, the values of all quantities in the B-A probe resembled the values of the late learning trials, whereas the randomized B-A probe (gray) had values that resembled early learning, suggesting the B-A behavior is not random. In the Probe B → A trials, the “food line” is the straight line that connects B to A, along which the reference distance d_{target} is measured between A and B. In the other trials, the food line is a straight line from start to the specific target, either A or B, along which d_{target} is measured between start and target.

After the probe trial, mice were trained to learn target B. During this training, 3 mice (Mouse 33, 35, 36) first learned a route to target B but, on subsequent searches, they would sometimes first check the site of the previous target A and then travel to the correct target B along a direct short cut route (see Discussion). Nevertheless, on average the mice kept farther from the previous site A than the expected by chance (Fig. S10A,B). A first visit to B then to A was never observed during training because mice go back home after getting the food in B.

After learning the location of target B (Fig. 6C), food was also omitted from this target for Probe B-A. As illustrated in Fig. 6D and the Shortcut video (Video 1) (all mice shown in Fig.S8B), 5 of 8 mice were observed to go from B to A via varying trajectories (namely, mice numbered 33, 35, 36, 57 and 59). Two (number 58 and 60) of the remaining three mice went first to A and subsequently to B via direct or indirect routes; these mice had previously taken direct routes to B, suggesting that going to A first is not due to a learning deficiency. The remaining one, mouse 34, went from B to the start location and then, to A. This mouse had previously taken the B-A route during training. In all cases, the mice clearly remembered the location of both targets, even though target A had not been presented or rewarded for 4 days.

The geometric and kinetic features of these experiments for early and late learning of each target, and for Probe B-A, are presented in Fig. 6E-H. For defining the trajectory quantities in the Probe B-A trial, the “start” position is taken as target B, and the “target” is the A site (e.g., the normalized trajectory length is the ratio between total traveled and direct B-A distances). The quantities in probe B-A are statistically indistinguishable from late learning of targets (i.e., trials 14-A and 8-B). Conversely, all the Probe B-A values are significantly different both from trial 1-B and from randomized B-A trajectories (see Methods; $P < 0.05$; except for the density of hole checks which is statistically equal for all considered trials, Fig. 6H). Therefore, the behavior of the mice when going from B to A in the probe is compatible with the behavior of an animal that acquired spatial memory about the trajectory, even though this trajectory was never reinforced. This suggests that the mouse computed the shortcuts without prior experience (see Shortcut video (Video 1)) for an example of shortcutting and associated hole checks).

We calculated the displacement and hole check maps and the TEV for the whole training protocol, including the probe B-A (Fig. 7). Again, the data show that, after training for A, the TEV pointed directly to A and hole checks accumulated along its path (Fig. 7A,B). The beginning of training for B (trial 1-B) generated random search patterns, while the TEV and hole checks tended towards A. After learning B, we see that the TEV and hole checks completely shifted to it (Fig. 7C,D,E,F). Remarkably, the probe B-A trajectories revealed a strong directed flow with a TEV pointing from B to A, whilst hole checks visibly accumulated along this route (Fig. 7G,H). The TEV-target deviation remained close to zero in late learning and probe B-A (Fig. 7I), whereas the area density of hole checks is increased near the target compared to far only after learning and in probe B-A (Fig. 7J). These maps suggest the emergence of a cognitive map guiding the mice when taking the B to A unrehearsed shortcut routes.

As a stringent control, we performed a two-food-locations experiment with a 180°-rotated probe B-A trial (Fig. 8; $N=8$). Landmarks were never present (neither for training nor for probe). In the rotated probe B-A, the mice went to the REL of B and then to the REL of A. The displacement map, hole check spatial distribution and TEV now all pointed from REL B to REL A (instead of B to A). This is consistent with both the independent experiments of the rotated probe with static entrance training and of the two-food location. Mice can thus take a novel short cut and have therefore computed a cognitive map based on a fixed starting point and self-motion cues alone.

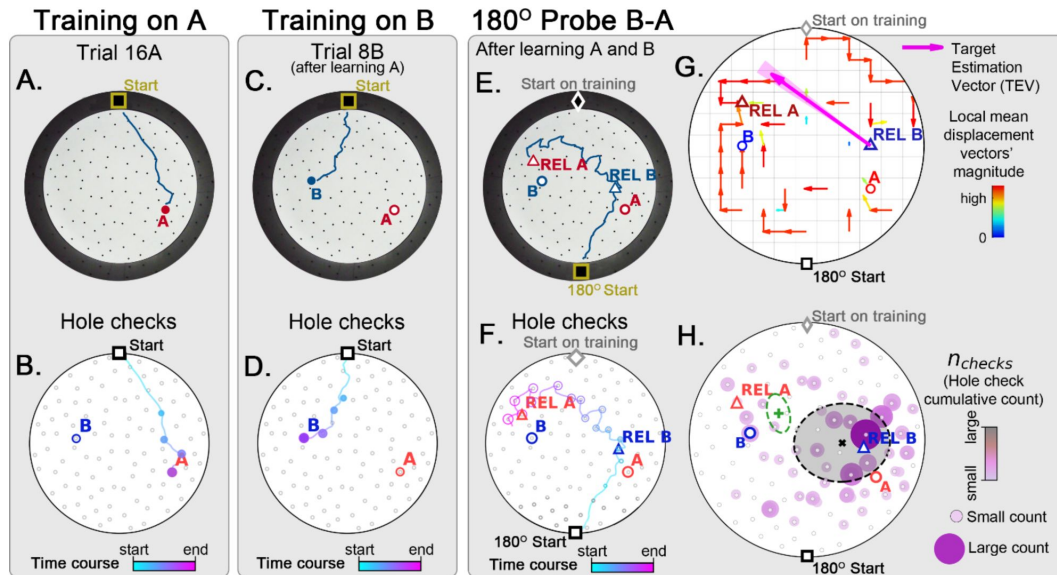


Figure 8.

Two food location training with 180° rotated probe.

Mice (N=8) are trained to find food in the target labeled “A” (red circle, panel **A**), and in target labeled “B” (blue circle, panel **B**) afterwards; then a 180° rotated probe trial (no food; panel **c**) is realized to check whether the mice are able to generalize and take the shortcut from REL-B (blue triangle) to REL-A (red triangle), instead of the A and B targets. No landmarks are present. Panels **A,C,E** show exemplars of trajectories in three stages of the experiment, and **B,D,F** show their time course and hole-check locations marked with circles that increase with elapsed time. **G**. Trajectory directionality analysis and TEV (pink arrow; shaded sector: S.D.) show that significant paths ($p<0.001$; N=8; see Methods) point from REL-B to REL-A in the same way that it pointed from B to A without rotated entrance in **Fig. 7G**. **Panel H**: the spatial density shows that hole checks accumulate along the REL-B to REL-A direction, instead of the B-A direction in the case without rotation in **Fig. 7H**. Black ellipse ($\times$ =mean): covariance of hole check density. Green ellipse ($+$ =mean): covariance of the same data restricted to ≤ 20 cm of the REL-A location. This suggests that mice follow shortcut trajectories anchored to their start location (idiothetic frame of reference).

Discussion

We have shown that mice can learn the location of a hidden food site when their entrance to an open maze remains the same across trials. Trajectories initially appeared random and took on average, over 100 seconds and covered a distance of about 10-fold that of the direct route between start and food sites (**Fig. 3**). After a maximum of 14 trials, mice reached asymptotic performance taking, on average, <10 s to reach the food site and with trajectories reaching a near optimal distance: 1.4 times greater than the direct route (**Fig. 3F**). In the spatial learning trials, mice frequently checked holes for food along their start to food trajectories. Analyses of hole checks demonstrated the distribution of hole check sites was also modified during learning (**Fig. 4G,H**). Hole check density decreased with learning in the first half of trajectories but remained constant during the second half. At higher spatial resolution, hole checks occurred most frequently closer to the food site after learning (**Fig. 4K**). Finally, the number of holes sampled near the food site became more consistent with learning (lower entropy: see Fig. S7G). We conclude that, given a static entrance, the mice can learn an accurate estimate of food location that guides their trajectories and associated hole check distribution.

Similar results have been reported for rats trained to find food in a virtual reality (VR) 2D spatial navigation task (21). During spatial learning, the reward check rate increased as the rat approached the reward edge for both visual and auditory reward location cues. Navigation (trajectories) and reward checking were, however, in register only for distal visual cues. In our experiments, hole checking at a reward site was not dependent on distal visual cues (see REL discussion below) suggesting that a different source (or sources) of spatial sensory input permits registration of trajectories and hole checking during real world (RW) spatial learning.

What cues are required for learning the food location?

A number of exogenous (tactile, localized and global visual, odorant) and self-motion (optic flow, proprioceptive and vestibular) cues might contribute to the spatial learning we observed. For any cue to be effective, it must be stable under the experimental perturbations we impose, and therefore provide unambiguous information about the location of the food reward relative to the mouse's starting point. Below we describe which cues might be relevant for spatial learning.

Allothetic cues

Tactile

There were no obvious scratches on the maze floor, but we cannot exclude fine scratches detectable by the mice. We wash the floor and randomly turn it between trials so that any scratches or odor trails are not consistently correlated with the food location. We also performed a specific control experiment with an exhaust fan to eliminate food scent. We conclude that tactile cues are likely not used for spatial learning under our experimental conditions.

Localized visual landmarks

Rodent vision is important for tasks such as navigating complex environments, finding shelter, prey capture and predator avoidance (see (22) for a review). Previous behavioral and physiological studies suggested that our wall cues could be resolved by the mouse visual system (23–27). It was therefore surprising that the mice were unable to find the food site when its start location was randomly switched between trials. A comparison of static versus random start sites after learning revealed that, unlike the static case, there was no reduction of trajectory length or increase of hole checks near the food site in the random start experiments. We considered whether the mice were learning slowly and simply needed more trials to learn the random start

site task. We therefore checked if there was any improvement in target estimation over 4 successive trials. As shown in Fig. S8, actual hole checks over the 4 trials are no different from randomized hole checks suggesting that the mice are randomly choosing holes to check.

The random start task requires the mouse to learn the relative location of food with respect to a changing view of the landmarks. Rats have been shown to be able to solve this task in the Morris Water Maze (MWM) (7 [↗](#)). Previous studies showed that mice can use distal cues for spatial learning under different experimental conditions, such as in a one-dimensional T-maze, linear track, or when there are multiple highly salient cues (7 [↗](#), 28 [↗](#)–31 [↗](#)). Our random start task would therefore seem to have all the requirements for spatial learning. The simplest explanation is that, unlike some of the cited papers, the cues we used were not sufficiently salient.

An alternate factor might be the lack of reliability of distal spatial cues in predicting the food location. The mice, during pretraining trials, learned to find multiple food locations without landmarks. In the random trials, the continuous change of relative landmark location may lead the mice to not identifying them as “stable landmarks”. This view is supported by behavioral experiments that showed the importance of landmark stability for spatial learning (32 [↗](#)–34 [↗](#)) and that place cells are not controlled by “unreliable landmarks” (35 [↗](#)–38 [↗](#)). Control experiments without landmarks (Fig. S6A,B) or in the dark (Fig. S6C–F) confirmed that the mice did not need landmarks for spatial learning of the food location.

Arena geometry

The mice might use the distance and bearing from its start location to the three other maze entrances. Triangulation might then enable it to estimate the food location based on such global features and continuously update its position by using the relative location of the entrances. Our control experiments (Fig. S6C,D,E,F) showed that such information is not needed for spatial learning but do not rule out that this information is an additional cue available to the mice under more natural foraging conditions.

Odor cues

The mouse’s cage will be saturated by its own odor. Transferring the mouse to a different cage might enable it to recognize the change in its starting point. However, our maze was designed to eliminate this “self-cage” odor cue as each mouse stays in its own home cage during the entire experiment. In the Random trials, the entire cage + mouse is moved to a different entrance on each trial. In the Static trials, the cage + mouse is only moved to a different entrance on the REL trials.

Odor cues might come from odor trails left during a preceding trajectory to food, from the odor of the hidden food or from an odor gradient emerging from the mouse’s home as it leaves to find food. The odor trails were reduced by washing the floor between trials. Any remaining odour was made unreliable as the floor was rotated between trials. A subfloor beneath the maze floor was covered with crumbled fragments of the same food within the capped hole. We assume that the entire maze was saturated with the food odor and that this would mask the odor coming from the accessible food. Mice would search in holes adjacent to the food without successfully locating the food, suggesting no detectable odour at that distance. In addition, the mouse ran to and searched at the food hole on probe trials confirming that it did not require a food odor to find the food hole.

Mice can also use localized odor sources as landmarks for spatial learning (20 [↗](#)), and might therefore use odor gradients emanating from their home cage as a cue. We reduced this potential cue by applying negative pressure via exhaust fans covering the cage top that were turned on 5 minutes before an experiment commenced and its full air volume evacuated 30 times/minute (150 evacuations) to equilibrate it to the room air. After the door to the maze was opened, the fans induced negative pressure to draw air from the maze into the cage and thereby reduce a potential

outward gradient (see Methods). The mice were still able to learn the food location after such reduction of olfactory cues (Fig. S6E,F) suggesting that odor gradients are not required for spatial learning under our experimental conditions.

Idiothetic cues

Optic flow cues

Although the mice can take varying trajectories, the TEV suggests that they continuously update their knowledge of the direction from start to food throughout their trajectories. This further suggests that the neurons coding for head direction are essential for spatial learning in our maze. Heading direction can be derived from optic flow signals (39, 40). In the static entrance experiments, the heading direction is constantly updated throughout the entire run (Fig. S4G), having the mouse veer toward the food only in the last segment of the trajectory (Fig. S4Giii). Hippocampal neurons of rats randomly foraging in a real world or virtual reality environment can develop directional responses based on rich and stable optic flow signals and without requiring vestibular input (41). In our maze, landmarks and home cage openings may drive optic flow signalling irrespective of whether they act as local visual cues. This is consistent with the hypothesis that visual input can guide navigation via a central retinal stream for landmarks and a peripheral retinal stream for optic flow input (42). Our in the dark experiments suggest that other sensory input might also drive the head direction network.

Vestibular and proprioceptive cues

The vestibular system encodes natural motion (43, 44) and contributes to the generation of head direction responses (45, 46). We hypothesize that, in the “dark” experiments, vestibular cues are a major signal supporting spatial learning. In agreement with VR studies (47, 48), we hypothesize that contextual visual cues of the starting point and arena boundary anchor directional information derived from optic flow, vestibular and perhaps proprioceptive signals to fixed environmental features.

A fixed start location and self-motion cues are required for spatial learning

We also considered that mice, unlike rats (7), may be unable to learn the random entrance task because it requires associating four different spatial cue configurations with a food location. This would hold when the visual cues were local. In the static entrance task, the mice might be learning a unique configuration of the cues or even the relative location of a single cue and the food location (6). In the REL experiments we therefore first fully trained the mice with a static entrance and, in a probe trial, randomly switched them to another entrance rotated by 90°, 180°, or -90°. Unlike rats in the MWM (7), the mice ran to the REL location (Fig. 5) clearly demonstrating that they assume that their start location has not changed and that the local visual cues did not guide the mice. We hypothesize that the optic flow signals emanating from the distinct landmarks, derived from low resolution peripheral retina (42), are invariant to and thus cannot differentiate between the mouse’s start sites. Additional control experiments made without landmarks or in darkness also showed that visual input is not required for spatial learning in our maze (Fig. S6). We hypothesize that the minimal conditions for the mice learning the heading direction from start to food is based on self-motion cues (optic flow, proprioception and vestibular) and one fixed local cue – the start site.

Combining trajectory direction and hole check locations yields a Target Estimation Vector

Here we follow a review by Knierim and Hamilton (12) that hypothesized independent mechanisms for extraction of target direction versus target distance information. Our data strongly supports their hypothesis. Target direction is nearly perfectly estimated at trial 6 (Fig. 4I and Results). The deviation of the TEV from the start to food vector is rapidly reduced to its minimal value (5.16°) and with minimal variability (SD=0.20°). Learning the distance from start to food is also evident at trial 6 but only reaches an asymptotic near optimal value at trial 14 (Fig. 3F). The learning dynamics are therefore very different for target direction versus target distance. As noted below, the food direction is likely estimated from the activity of head direction cells. The neural mechanisms by which distance from start to food is estimated are not known (but see (49)).

Averaging across trajectories gave a mean displacement direction, an estimate of the average heading direction as the mouse ran from start to food. The heading direction must be continuously updated as the mice runs towards the food (which is suggested in Fig. S4G), given that the mean displacement direction remains straight despite the variation across individual trajectories. Heading direction might be extracted from optic flow and/or vestibular system and be encoded by head direction cells. However, the distance from home to food is not encoded by head direction signals.

The mice, after learning with static entrances, made the greatest number of hole checks in the vicinity of the food-containing hole. We therefore used the mean of the “near food” hole check distribution (see Fig. 4) to give a magnitude to the displacement direction, thus generating the TEV. With learning, the TEV rapidly converged to closely align with the direct start-to-food vector (Fig. 4E,F,I). In REL trials, the calculated TEV pointed to the REL despite the lack of food at the hole check sites (Fig. 5). A close analysis of individual trajectories revealed that, while some closely aligned to the TEV and went directly to food, others deviated from the TEV/direct route and then returned to it (Fig. 4F and S7D). The standard assumption is that deviations from a direct route are due to path integration errors (50). It is not obvious why errors only occur on some routes. A second possibility is that the deviations are intended and meant to prevent route predictability and therefore predation (51). It is not clear why the mouse does hole checks if it is only reducing route predictability. A plausible hypothesis is that the mice deliberately deviate from the TEV in order to continue exploring for food-containing holes, *en route* to exploit the food reward. We hypothesize, following Knierim and Hamilton (12), that a path integration mechanism operates continuously to return the mouse to its current TEV estimate no matter what the reason for the deviations from the TEV.

We hypothesize that path integration over trajectories is used to estimate the distance from start to food. The stimuli used for integration might include proprioception or acceleration (vestibular) signals as neither depends on visual input. Our conclusion is in accord with a literature survey that concluded that the distance of a target from a start location was based on path integration and separate from the coding of target heading direction (12). Our “in the dark” experiments reveal the minimal stimuli required for spatial learning – an anchoring starting point and directional information based on vestibular and perhaps proprioceptive signals. This view is in accord with recent studies using VR (47, 48). Under more naturalistic conditions, animals have many additional cues available that can be used for flexible control of navigation under time or predation pressure (52).

Implications for theories of hippocampal representations of spatial maps

The TEV is learned using self-motion cues alone and we hypothesize that it guides locomotion to the food hole via a path integration mechanism. This conclusion is in accord with an extensive literature that emphasizes the importance of self-motion cues and path integration for spatial learning (8, 53). The conclusion is also not surprising given studies on weakly electric fish that demonstrate that, with a fixed initial location, active sensing and self-motion cues are sufficient for learning the location of a food site in the dark (49, 54), and that accumulation of path integration error degrades performance as a function of trajectory length (55, 56). Given the presumed accumulation of error by the mammalian path integration mechanism (50), it is generally assumed that proximal and/or distal exogenous cues must calibrate the putative path integrator in order to determine not only target direction but also target distance from a start site (8, 12, 53, 57). Path integration gain of hippocampal neurons is a plastic variable that can be altered by conflicts between self-motion cues and cues and feedback from landmarks (57). The independence of the TEV from landmark cues (REL experiments) again demonstrates that, in absence of such “conflicts”, self-motion provides consistent cues for path integration and spatial map formation.

The use of hole checking to compute a mouse’s estimate of the food location allows us to refine these conclusions. The TEV provides not only an estimate of the food location, but also of the distance from start to food site; this is especially clear in the REL experiments where, given the lack of food, a mouse’s search is clearly centered at the expected food location (Fig. 5). Sophisticated analyses have been used to link the spatial coding neurons of entorhinal, subicular and hippocampal neurons to behavioral studies on the interaction of exogenous and self-motion signals (15, 58). Here, we provide strong behavioral evidence to support the role of hippocampal place cells in encoding the trajectories and food site locations observed in our study. Two studies suggest that, in the rat, CA1 place fields will remain stable in the absence of visual cues. Many place cells responses observed in the presence of visual cues will remain after these cues are removed; the authors conclude that self-motion cues are sufficient to maintain normal place fields (11). Experiments with blind rats have shown vision is not necessary for the development of normal firing of hippocampal place cells (59). Together, these studies suggest that mice CA1 cells will exhibit place fields in our open maze. Analyses of spatial learning in VR versus RW (rats) suggested that distal visual cues, and self-motion (i.e., proprioceptive, vestibular) cues may be required to activate place cells representing allocentric space. In the absence of distal visual cues, CA1 cells preferentially encoded distance traveled during learned trajectories toward a food goal (60, 61).

Rats can learn to navigate, in VR, from different start locations to a hidden goal using very salient distal visual cues (62). Unlike RW spatial learning, CA1 pyramidal cells then only weakly encoded allocentric spatial information (place fields) but instead, primarily encoded trajectory distance and head direction – thus, these cells may comprise a possible neural implementation of the behavioral TEV. These experiments demonstrate great flexibility in the hippocampal encoding of trajectories and location both across pyramidal cells and, for individual cells, across the learned trajectory. An important question is how CA1 pyramidal cells will discharge as a mouse runs along the stereotyped trajectories learned with only self-motion cues as in our experiments. A stringent prediction of the discussion above is that mouse CA1 cells activated along trajectories towards the hidden food should be activated at equivalent locations in REL experiments, and in response to the same multiplexed cues: trajectory distance, head direction and allocentric location.

A subset of hippocampal (CA1) neurons in the bat were reported to encode the direction and/or distance of a hidden goal (63). The vectorial representation of goals by such cells could be the substrate of the start-to-food location trajectories we have observed. Recently described CA1

convergence sink (ConSink) place cells (64) may also have the properties needed to account for the learned start-to-food trajectories we observed. ConSink cells are directional place cells that can encode local direction towards a goal. ConSink cells will, with training, shift their direction tuning to a new goal. The ConSink population vector average, like the TEV, then points directly from start to goal. In both cited studies, landmarks were present, and it is unknown whether such cells will be found in the absence of visual cues.

ConSink-like cells were reported in the CA1 of a mouse foraging in an open field, but their direction tuning was not pointed towards a singular goal (65). We hypothesize that in the pretraining foraging phase of our spatial learning task, ConSink cells will also have random direction tuning. Upon fixed start location training, we hypothesize that the ConSink direction tuning will become aligned with the mouse's trajectories and their population average will closely approximate the computed behavioral TEV. The TEV and the ConSink cell population average are statistical measures derived from behavioral and electrophysiological data respectively. An important question is whether an explicit TEV is computed and defines the spatial map guiding the mouse's food-finding trajectories. An equally plausible hypothesis is that the "spatial map" remains a distributed computation in the CA1 targeted neural networks. Experimental tests of these alternatives address an essential question: how is spatial information represented in neural networks?

Numerous studies have reported that goal sites are overrepresented by CA1 place cells (4). The requirements for such overrepresentation are that there are stereotyped trajectories directed towards an invisible memory-based goal associated with reward (4). The stereotyped trajectories and the hidden memory-based goal of our study imply that these requirements are met. Interestingly, the "goal-related place cells" are activated before the goal is reached (4), just as hole checks mostly occur as the mice approach the food containing hole from any direction (Fig. 4F,H). We hypothesize that, after learning, CA1 place cells will overrepresent the maze region containing the goal location, and that their place fields therefore overlap the hole check sites surrounding the hidden food hole.

Behavioral time scale plasticity (BTSP) has been proposed to be the cellular mechanism that generates new CA1 place fields at important locations, including those associated with reward (66–68). BTSP operates up to a ~3 s time frame. If BTSP is operating during spatial learning in our maze, there will be excessive place fields within the <3 s search time before it finds the food hole. BTSP is bidirectional and can result in CA1 place fields translocating with experience (68) and the location of CA1 cell place fields may therefore evolve during static entrance training. We note that the cited experiments were done with virtual movement constrained to 1D and in the presence of landmarks. It remains to be shown whether similar results obtain in our unconstrained 2D maze and with only self-motion cues available.

The putative emergent CA1 place fields might be randomly distributed but might also be connected with the "special" hole check locations. We plotted the evolution of <3 s from target hole checks in the static and random entrance experiments (Fig. S9). With spatial learning (static entrance), temporally "close to target" hole checks increase relative to temporally distant hole checks and converge to the target site. Active sensing is critical for electric fish spatial learning (49), and is also known to potentiate or induce CA1 cell place fields (69). We hypothesize that the persistent <3s hole checks near the food site that increase during training will, via the BTSP mechanism, drive the emergence and translocation of CA1 cell place fields so that they accumulate centered on checked holes near the rewarded food site (4). This leads to the prediction that a place cell discharging during a hole check near the food site should also discharge after the mouse start site has been rotated and it checks an empty REL hole.

Shortcutting – Evidence for a cognitive map derived from self-motion signals

The O'Keefe and Nadel text (16) connected hippocampal place cells to the abstract concept of a cognitive map developed by Tolman (13), and this linkage has been generally supported with few dissenting views (70–73). The criteria for a cognitive map are of animals taking unrehearsed shortcuts (14), detours (13) or novel routes (7). In this literature animals typically have both landmark and self-motion cues available for spatial learning. To the best of our knowledge, unrehearsed shortcut behavior using only self-motion cues and a fixed start location has only been shown in humans (53, 74). Our result on shortcutting after spatial learning based entirely on a fixed start location and self-motion cues (Fig. 7) is therefore the first behavioral demonstration of a rodent cognitive map learned without exogenous cues and using the strict Tolman definition. The TEV for the shortcut trajectories well approximates the direct route between the Site B and Site A food locations (Fig. 7G,H) demonstrating the accuracy of the putative cognitive map food location estimate.

The shortcut trajectory from Site B to Site A (Fig. 7G,H) might be formally computed in two ways. For the 3 mice that had taken the Site A to Site B route during training, the following vector arithmetic is required for the final probe trial when it went from Site B to Site A:

$$\overrightarrow{TEV}_{A \rightarrow B} = \overrightarrow{TEV}_{Start \rightarrow B} - \overrightarrow{TEV}_{Start \rightarrow A},$$

where $\overrightarrow{TEV}_{B \rightarrow A} = -\overrightarrow{TEV}_{A \rightarrow B}$. For the 4 mice that had never taken the Site A to Site B route during training, the following vector arithmetic is required for the final Site B to Site A shortcut:

$$\overrightarrow{TEV}_{B \rightarrow A} = \overrightarrow{TEV}_{Start \rightarrow A} - \overrightarrow{TEV}_{Start \rightarrow B}.$$

The TEVs for Start → A and Start → B appear to be still remembered by the mice in the probe trial, since the accumulation of hole checks at both sites is still evident (Fig. 7H). The hypothesized ConSink place cells directed to Targets A and B and the accumulation of cells with place fields at both sites will therefore, as described above, still encode the trajectories to each location. To our knowledge, neurons that might compute the required vector arithmetic have not been identified in any part of the rodent brain.

We hypothesize that the TEVs estimated by neural networks downstream of goal vector cells, CA1 ConSink cells and goal location place field cells will be used to compute the shortcut trajectories. Discovering the networks that do these putative computation(s) may provide insight into the neural bases of the spatial cognitive map.

Author Contributions

All experiments were carried out in the L.M. laboratory at the University of Ottawa. J.X. ran all experiments, analysed learning and contributed to experiment design, writing the manuscript and preparing figures. M.G-S. developed analysis for trajectories and hole checks, and carried out the analyses plus associated figures and contributed to writing the manuscript. J-C. B. contributed to writing the manuscript. A.L. contributed to developing trajectory analysis methods and writing the manuscript. L.M. developed the conceptual framework for the experiments and contributed to and supervised all aspects of the experiments and data analyses and wrote the manuscript.

Acknowledgements

We would like to thank Dr. Érik Harvey-Girard for technical support, and William Moldenhauer de Jesus for joining the short cut experiment video with the hole check data. This work was supported by the Canadian Institutes for Health Research Grant # 153143 to AL and LM and J-C B, a Brockhouse award (493076–2017) to AL and LM, an NSERC award (RGPIN/06204-2014) to AL, an NSERC award to LM (RGPIN/2017-147489-2017) and a grant from the Krembil Foundation to AL, LM and J-C B.

Data availability

All codes and data are available in <https://github.com/neuro-physics/mouse-cogmap> 

Declaration of interests

The authors declare no competing interests.

Material and Methods

Animals

All animals were housed in the University of Ottawa Animal Care and Veterinary Services (ACVS) facility. C57Bl/6 wild-type male and female mice were ordered from Charles River, arriving at 8-9 weeks old. Mice were individually housed in 12h light/12h dark cycles (lights on at 11:00PM EST, lights off at 11:00AM EST). Animals had food and water available *ad libitum*. The temperature of the room was kept at 22.5°C and the humidity was 40%. Mice were habituated in ACVS facilities for 1 week and began testing when they were 10-11 weeks old. Testing of each cohort took place over the course of approximately 2 weeks, 1 hour after the lights turned off. Subjects weighed 22-27 grams at the start of behavioural training. Both male and female mice were used; the same results were obtained in both sexes and were therefore pooled. All animal procedures were conducted with the approval of the University of Ottawa's Animal Care Committee and in accordance with guidelines set out by the Canadian Council of Animal Care.

Apparatus Design & Setup

The Hidden Food Maze (HFM) is a framework that trains mice to search for a food reward hidden in an open, circular arena (**Fig. 1** ). The protocol for the task was inspired by an open maze protocol used to study electric fish spatial learning ([54](#) ) and by the Cheese Board spatial task in which rats are placed inside an arena with many holes in the floor, one of which contains a food reward ([75](#) ). Unlike the Kesner et al task, the HFM does not require the animal to be handled, and the pattern of holes is arbitrary rather than grid-like. This task design was also chosen to mirror the setup of the Morris Water Maze (MWM) ([7](#) ) , which can test allothetic or idiothetic navigation. Like the MWM and the electric fish arena, mice are searching for a food reward location in a circular environment after starting from one of four starting home locations. External landmarks can, if desired, be placed on the maze walls or local cues placed in the maze. In contrast to the MWM, this spatial learning task is a dry maze that is food motivated, which is

more naturalistic and less stressful than motivation by the aversion to swimming and fear of drowning. Notably, the task has specific design features to control for against unwanted cues, such as odours, visual cues, and handling.

The maze has a removable floor that is washed & rotated between trials to eliminate odor trails which mice from previous trials might leave behind. The circular floor is 120 cm in diameter and has 100 holes (1.2 cm diameter) randomly dispersed throughout the surface, with 25 holes in each quadrant. The distance between each hole is, on average, 10 cm. The pattern of the holes is rotationally symmetrical, so the pattern looks the same regardless of whether the floor is rotated 90°, 180°, or 270° with respect to the mouse's entrance. This ensures that the mouse will have the same initial view of the holes regardless of which entrance it starts from and how the floor is rotated. Each hole is encircled by a 1 cm plastic rib, sticking downwards, which can be capped at the bottom to hold food that remains invisible from the surface. Thus, the mice cannot discern the contents of the hole just by looking across the floor from their home, but instead need to approach the hole and look inside. The surface of the floor is sanded to be matte to avoid generating reflections from the lights which might interfere with the cameras or distract the mice.

The maze floor is circular and uniform from the inside so as not to provide any directional cues. The floor of the arena is encircled by tall black walls. The walls are made of solid black PVC plastic which forms a cylinder around the maze and is open at the top. The walls are 1 cm thick and 45 cm tall, which is tall enough so the mice cannot jump out. The walls are symmetrical and designed to eliminate geometric cues that would give away directional information from asymmetries in the environment shape. Mice can use odors as landmarks for spatial learning (20) and we wanted to eliminate local food odour from a filled hole as a landmark. The arena is resting on a subfloor that contains crumbled food; food odor will diffuse through the open holes and saturate the maze thus masking the odor from a food-filled hole.

The home cages attach to the main arena by being slotted into each entrance. The home cages are 27 cm x 16.5 cm x 45 cm and are open at the top. They contain a food hopper and a water bottle feeder. The dimensions of the home cage are based on commercial mouse cages and comply with the Canadian Council on Animal Care mouse housing standards.

When moving mice to a different starting quadrant, the home cages are designed to slide out from the maze and into a new entrance so the mice do not need to be handled and will therefore not be stressed. The home cages include their own doors, separate from the doors that provide entry into the maze, so the home cage can be freely moved and the mice remain securely confined. The detachable home cages allow the mice to be moved to different starting locations, allowing us to control against navigational strategies that rely solely on response learning or path integration from a fixed starting point.

Mice can perform the entirety of the trial without experimenter handling. The maze doors are designed to slide upwards and provide access to the main arena. When the trial is finished, the experimenter can slide the doors back into place and the mouse is confined to its home cage once again. To not disturb the mice, the doors are left open while they navigate.

Four LED flashlights were aimed at a white ceiling in order to create dim, diffuse lighting throughout the maze. The illuminance is measured to be 50 lux at the surface of the arena. Dim light was used to allow the mice to properly see all visual cues as well as the maze details. This procedure is assumed to not perturb the nocturnal cycle of mice (76, 77). Black curtains surrounded the maze to prevent the interference of non-controlled visual cues.

We used additional experiments to control for possible visual or odorant cues from the open home door. We did some experiments in total darkness using infrared LEDs with emission spectra detected by our camera. In order to exclude any potential olfactory cues emanating from the open

cage door we did experiments where odor absorbent kitty litter was placed on the cage floor and two connected exhaust fans (AC Infinity MULTIFAN S5, Quiet Dual 80mm USB Fan) placed above the home cage. At their lowest (quiet) setting, the fans drew air from the home cage and blew it to the outside; replacement air from the outside came in via small openings at the bottom of the home cage. The home cage air volume was evacuated 30 times/minute, effectively equilibrating the cage and open maze air before the doors were slid upwards. Fans were turned on for 5 minutes (150 evacuations) before the trial started and for the duration of the trial. After the door was opened, the fans induced negative pressure to eliminate diffusion of odors from the home cage to the maze. The two fans were placed over each home cage and turned on to eliminate a potential directional noise cue.

Behavioural Training

Pre-training

Five days before training, mice are transferred from standard animal facility cages to the experimental home cages for habituation. On the first day, mice are allowed to habituate to their new home cages. On the second day, the door dividing the home cage and the arena is removed and each mouse is given free access to explore the empty arena for 10 minutes. There are no extra-maze landmarks present during pretraining. Animals are food restricted, with 10% of their body weight in food given back each day which they could consume *ad libitum*. On day 3, randomly chosen 50% of the holes in the arena are filled with food treats (a piece of Cheerio), and each mouse is given 10 minutes to forage for food. This is repeated on day 4, where 25% of the arena's holes are filled with food treats. On day 5, only four holes placed at the maze center contained treats. Mice pass the pre-training stage when they successfully found treats at all locations within 20 minutes. Mice mostly confined their search to circular trajectories near the wall of the maze for Days 3 and 4 but, after training with food near the maze center (Day 5), they checked holes throughout the maze (Fig. S1A). Training in the spatial learning task then commenced.

Visual Cues + Randomized Entrances Training

The mice are trained to locate a food reward that has a fixed relationship with four visual cues on the walls. The protocol mimics classic Morris Water Maze setups to test allocentric landmark-based learning. One of the holes in the area is capped from the bottom with a food-containing insert that is not visible from the surface. Extra-maze landmarks (described in the Visual Cues section) are placed on the walls of the maze to serve as location cues. At the start of the trial, the door is slid upwards so the mice can enter the maze, and the mice are given a maximum of 20 minutes to find the target hole. The home door is kept open to permit the mice free access to return to their home cage during the trial; preliminary experiments indicated that closing the door perturbed the mice. The door is re-inserted after the mouse has found the food reward and returned home. If the mouse has not returned home by itself after 1 minute of finishing the food reward, it is gently guided back by the experimenter. There is an inter-trial interval of approximately 20 minutes.

Initial experiments used two trials per day but this was increased to speed up learning; our analyses and graphs are truncated at 14 trials to permit averaging over all the mice. In most experiments, three trials a day were given, over six days; at the end of each trial the mouse's home is moved to a different entrance. The home location was changed for each trial; the floor was washed and rotated by 90°, 180°, or 270° between trials, independent on whether the home entrance was the same (static trials) or rotated (random trials). The location of the entrance is randomized with the following constraints: no two trials in a row have the same entrance, and all entrances are selected the same number of times so there is no bias. On the 7th day, a probe trial is given where the mouse is allowed to search for food in the absence of any food. Learning continues over 1 more day and, on the 9th day, a reversal trial is given where the visual landmarks

are rotated by 180°. The location of the home cage is randomized before every trial. In this task, the mouse would have to learn the invariant relation between the food hole and up to four visual cues in order to locate the food.

Visual Cues

Four black symbols on white backgrounds: a square, a cross, vertical bars, and horizontal bars. The cues are taped 5 cm above the floor. The square was 15 by 15 cm. The cross consisted of two 2.5 cm wide and 14 cm long bars. The four vertical and four horizontal bars were 2.5 cm wide and 14 cm long. Previous behavioral studies have shown that mice can discriminate the visual stimuli we use (23, 25, 26). Recording of neurons in mouse V1 have additionally shown that orientation selective cells in mice visual cortex can discriminate our visual stimuli (24).

Visual Cues + Static Entrance Training

This protocol allows mice to navigate by potentially using visual cues in cooperation with path integration or by path integration alone. The setup is the same as visual cues training, except mice enter from the same entrance each trial instead of a randomized entrance. Mice are considered well trained once their latency learning curve has plateaued for 3 successive trials. The mice's latencies had plateaued by 14 trials but training continued till 18 trials.

After 18 trials, reversal trials were done with no food present and the mouse's home cage rotated 180°. The mouse was allowed to search for food for 10 minutes. If the mice were able to learn the invariant landmark/food spatial relationship, we would expect them to search at the learned food site. If they had used path integration of self-motion cues to navigate from the fixed entrance to the food, we would expect them to search at the rotationally equivalent location (REL).

Rotationally Equivalent Location (REL)

The four quadrants of the arena's floor are identical. This means that the position of the holes in the second quadrant (Q2, Fig. 1B; also Fig. S4A) is equal to the position of the holes in the first quadrant (Q1) rotated by 90° counter clockwise (CCW) about the center of the arena. Q3 is Q1 rotated by 180° (CCW) and Q4 is Q1 rotated by -90° (i.e., 90° CW). In other words, every hole in Q1 has an REL in each quadrant achieved by the respective rotation. The REL target works in the same way: it is the position (relative to the entrance used in the trial) where the food would have been had the mouse been trained from the entrance it used in the trial. To illustrate this, consider Fig. 1D: we put the food (target) in a particular hole in Q1 for a mouse that is trained to enter from Q2 (such as the example in Fig. 1B) in the static entrance protocol. The vector that points from the mouse's start point (in Q2) to the target in Q1 (i.e., the target vector) is "anchored" to the start position at Q2. For example, in the mouse's perspective, the target vector is equivalent to going forward for 70cm, and then turn left and follow for another 30cm. If in a probe trial, we now let the mouse enter from Q4 instead, there are two options: either the mouse uses the landmarks and searches for the food in Q1 (where the food actually used to be), or it follows the target vector (70cm forward + 30cm left) and goes to the REL location of the target in Q3. Following the target vector is a sign of path integration using self-motion signals (idiothetic cues).

Control Experiments

Rotations following Static Entrances – with and without visual cues

We performed additional tests for the use of visual landmarks in the Visual Cues + Static Entrance protocol. Well-trained mice had their home cages rotated to another quadrant and tested on how well they found the food from a new starting location. Mice were rotated 180°, 90°, and then -90° with respect to their original location (Fig. 5A,B,C). Mice were given 6 consecutive trials at each new location. One control group of mice had the visual cues on the arena wall removed during

rotation training. If mice were able to use visual cues, we would expect them to improve their search efficiency when visual cues were available. Lack of improvement or similar performance compared to the control group without visual cues suggests that the mice do not rely on the type of visual cues we provided for spatial learning and navigation.

Navigating and Training without visual cues and in Darkness

To control for the effects of all visual information, one cohort (4 mice) was trained and tested without any cues (Fig. S6A,B). A second cohort (4 mice) was trained in the light according to the Visual Cues + Static Entrance protocol, then the lights are turned off and the mice are given a trial in darkness (Fig. S6C,D). The lights were opened in between trials so the mice did not acclimatize to the darkness. Mice were tracked using infrared LEDs with emission spectra detected by our camera. The following conditions were tested: darkness during a regular trial with food present in the arena and during a probe trial where there is no food present in the arena; both conditions gave the same results.

We additionally both trained and tested a cohort (4 mice) in darkness and with control of potential odors. A recent study has shown that placing sighted mice in darkness impairs entorhinal cortex head direction cell tuning (19), but here the mice successfully learned to find food and with the same time course of learning (Fig. S6E,F).

Two Food Location Training

This protocol aims to test flexibility of spatial learning using only path integration (N=8 mice). No visual cues are placed on the arena walls. Mice entered the arena from the same entrance each trial. Food was placed in a target well in “Target A” and mice were trained to find this location for 18 trials, 3 trials a day. A probe trial (“Probe A”) was done after trial 18. For trials 19-26, the food was moved to a different location, “Target B”, and mice are trained to find the new location. We were careful to choose Target B so that Target A and B were not symmetric with respect to the mouse’s entrance. A second probe trial (“Probe B-A”) was done after trial 26. The purpose of the “Probe B-A” trial is to check whether mice take a shortcut between B (latest learned target position) and A (first learned target position). For some cohorts, the training continued until trial 34 and a third probe was given.

Two Food Location Training with rotated probe

The mice (N=8) were trained exactly as in the “Two food location” experiment described above. However, the mice start location was rotated by 180° prior to the second probe trial. This protocol joins the “Static entrances with rotated probe” protocol with the “Two food location protocol”, and is designed to provide further support for our path integration hypothesis. Each of the trained targets have their REL counterparts (180° rotated around the center of the arena). Now, the purpose of the rotated “Probe B-A” trial is to check which of the two options will happen: either mice take shortcuts between B (latest learned target position) and A (first learned target position); or they take shortcuts between REL B and REL A, supporting path integration via self-motion cues.

Behavioural Analysis

Path tracking was done with Ethovision XT15 (Noldus) based on contrast detection. Mice were tracked according to 3 body points at 30 frames per second on a 1080P USB camera with OV2710 CMOS. Videos were first recorded on AMCap webcam recording software and imported into Ethovision in order to preserve high quality videos. Trajectory plotting was done in Python.

Latency to target was calculated from when the nose point of the mouse enters the arena until the nose of the mouse enters the target hole. Trajectory analyses are based on the nose point sequence of the mouse. Speed and distance were calculated based on trajectory coordinates exported from

Ethovision.

Search bias during probe trials is calculated by totalling the time a mouse spent within a 30×30cm square centered around the target vs. the RELs in the other 3 quadrants during a 2-minute trial (Fig. 2C, Fig. 3C).

Statistical Tests

Significance testing was conducted in a manner that was appropriate for each dataset. Paired comparisons utilized the t-test for parametric data and the Wilcoxon signed-rank test for non-parametric data. The significance cut-off was set at 0.05. Statistics were calculated either using R (URL <http://www.R-project.org/>) or scipy 1.8.0 in python 3.8.2. The statistical significance testing of the trajectory directionality analysis was developed from first principles, since it involves angular variables (the direction of each vector). It is presented in the “Analysis of Trajectories” section ahead.

Figures showing a quantity evolution over trials [Fig. 2D-I; Fig. 3D-I; Fig. 4I; Fig. S4; Fig. S7E; Fig. S10A,B] have symbols as averages, error bars as standard error, and the shading around the curve corresponds to the full data range (lower and upper shading limits correspond to minimum and maximum sampled data points, respectively).

Boxplot figures [Fig. 4J,K; Fig. 6E-H; Fig. 7I,J; Fig. S7L] have the diamond symbol as the average, the thick black line as the median, the box covering the interquartile range (IQR), and the whiskers extend from the box limits to ± 1.5 IQR in both directions. No outliers were detected in any of these plots.

Randomized trials

For the Two Food Location experiment, we calculated a randomized version of the Probe B-A trial for comparison. It consisted of extracting ten random pieces of the trajectory. Each piece was defined by randomly selecting a pair of points in the trial trajectory that are separated by the B to A distance. Then, the particular quantity of interest was averaged over each piece of the trajectory, and these averages were then averaged to obtain a single value for the “Probe B-A Rand.” condition in Figs. 6 and 7. See Fig. S8 for a definition of the quantities.

Statistical Analysis of Trajectories

Experiment Alignment

The procedure described in this section is applied to calculate trajectory directionality (and hence the target estimation vector, TEV) and the hole checking spatial distributions. Each mouse in the Random Entrance experiment started from a different quadrant of the arena, for each consecutive trial, and the target was fixed relative to the arena (global reference frame). However, in order to increase sampling, we need to coherently align the mice entrances, generating the mouse’s perspective reference frame (see Fig. S1B-F). Each experiment batch was performed with four mice, such that at any given trial, a given mouse entered from quadrant 1 [Fig. S1C], another entered from quadrant 2 [Fig. S1D], the next entered from quadrant 3 [Fig. S1E], and the last entered the arena from quadrant 4 [Fig. S1F]. Most of the experiments, then, consisted of two different batches of four mice that had to be conveniently aligned to increase sampling. In this figure, we aligned all the entrances to the Start point at the top of the arena (trajectories were rotated accordingly), emphasizing the four possible positions of the target from the mouse’s perspective.

We had to find a way to align all the targets in a given trial to one of the four possible positions, such that the experiment stays coherent over trials (i.e., the target randomly switches in between these four positions from trial to trial). For example, in Fig. S1B-F, we show trial 14. The target position for each starting quadrant is labeled with a red letter A and a red circle, whereas the target for the previous trial is labeled with a green circle and a green letter A subscripted with “trial 13”. These two positions (trial and previous trial) must always be different to keep the random characteristic of the experiment.

Notice that the target positions in each of the panels D, E and F in Fig. S1 are simply rotated relative to the target positions in panel C. The target configuration in panel D is 90° clockwise rotated relative to the target configuration in panel C. The configuration in panel E is 180° clockwise rotated, and panel F, 270° clockwise rotated; both angles are relative to the configuration in panel C. Thus, in order to sample all the mice together, we simply rotate the trajectories from the mice that started in quadrant 2 [panel D] counter-clockwise by 90°; the trajectories from the mice that started in quadrant 3 [panel E] are rotated by 180° counter-clockwise; and the trajectories starting from quadrant 4 [panel F] are rotated by 270° counter-clockwise. This reduces all the experiments to the first quadrant, allowing us to sample all the trajectories of all the mice together in each individual trial.

Active sensing and hole-check detection

We developed an algorithm to detect the mouse's behavior of sniffing holes to detect food as a measure of active sensing. The hole checking events are used to infer the mouse's memory and uncertainty about its environment. The procedure described here is applied independently to each trial. We compute the spatial distribution $\mathcal{A}(X_i, Y_i)$ for the count of hole checks for each hole i positioned in (X_i, Y_i) in the arena, accumulated across all mice in a given experiment (Fig. 4C,D,G,H; Fig. 5C; Fig. 7B,E,F,H; Fig. 8C; and Fig. S1B-F; usually $N=8$ mice for each experiment), and normalized by the total count. The frequency of checks in each hole is coded both in the color and size of the filled circles: larger and darker shaded circles correspond to larger number of checks in that particular hole (lighter shaded pink small circles are a small number of checks). The black ellipsis marks the covariance of the $\mathcal{A}(X_i, Y_i)$ spatial distribution (“x” in the figures marking the mean position of hole checks). It is constructed from the eigenvalues and eigenvectors of the covariance matrix C of the hole check coordinates (X_i, Y_i) weighted by $\mathcal{A}(X_i, Y_i)$: the eigenvectors give the directions of the ellipsis semi-axes, whereas the eigenvalues give the width of each semi-axes. The center of the ellipsis (mean of the ellipsis' foci) is aligned with the mean of the spatial distribution.

Under the path integration hypothesis, it is expected that error accumulates as the mouse walks (15). This is consistent with the increase in number of hole checks per unit area near the target (Fig. 4J,K; “near”=less than 20cm away from the target). We consider that the change in number of hole checks measures the variability of the mouse's estimate of the target position.

We therefore calculate the mean and covariance of the distribution of hole check events restricted to within 20cm of the target (again, the green “x”=mean). The distance from the start to this restricted mean is used to scale the TEV vector (see Target Estimation Vector).

The arena design forces the mouse to put its nose very close to the hole to be able to see the food or detect any food odour. With that in mind, we defined two methods for detecting a hole check (see Fig. S2A,B, and the Shortcut video (Video 1)). The two methods are complementary, and the second can find hole-check events missed by the first method: we apply the first method, then perform a visual check on the data to see if there were potential hole-checks that were missed. Then, we apply the second method to capture any remaining events.

The first method is the minimum velocity criterion and provides a high threshold for hole check identification. Four conditions must be satisfied simultaneously to detect an event: (i) the nose of the mouse must be within 3 cm of an arena hole; (ii) the velocity has to be less than 20% its maximum value; (iii) the velocity must be at a minimum; and (iv) the velocity has to have dropped by at least 5 cm/s to reach that minimum.

The second method is more inclusive and defines a hole-checking event by a simple slowing down event, provided that: the nose is within 3 cm of an arena hole, and the slowing down is enough for the velocity to drop past the 20% threshold of its maximum.

The detected events by the application of these two methods in sequence are marked for all trajectory samples of the Probe B-A in Fig. S8B as examples. A particular case is shown in more details in Fig. S2, and in the Shortcut video (Video 1) with the recording of a mouse's performance. It is worth mentioning, there is a clear hole-check missed by our algorithm in the 25 to 27 second range. This is because the velocity of the mouse was already below the 20% maximum before and after the hole check, hence the two sets of criteria defined above were not met for this particular event and it was counted by the manual check. Otherwise, we clearly see that all the other events are captured by sequential use of our two methods.

Trajectory directionality (displacement map)

This analysis is a way to visualize mouse trajectories and directionality across all the mice for each trial. It is related to a velocity map of the arena, except that here, the arrows point in the direction of most probable movement instead of the direction of the velocity. Fig. S2B shows a scheme for how we compute this map, and we detail it below. We call this quantity as the “displacement map”, since it gives the displacement of the mouse for each position in the arena. In this section, we explain how to calculate the displacement map for a single sample of N mice. The procedure here is applied to each individual trial independently (either during learning, or for a probe trial). This map is used for inferring the learned directionality of the target (relative to entrance). The error, average and significance of this quantity is estimated from first principles by a jackknife procedure explained in the next section. In the figures we show the average displacement map vectors that were found to be significant. These maps are then used to compute the target estimation vector (TEV; detailed in the last section).

We start by overlaying a lattice of size L on top of the arena recording. This means that there are L boxes on the x (horizontal) direction, and L boxes on the y (vertical) direction, making a total of L^2 boxes in the lattice. Each box defines a lattice site with coordinates (x, y) , such that x and y are integers between 1 and L . Each mouse corresponds to an independent observation of a trajectory, so we overlay all mice trajectories for each trial separately in the calculations below. Next, we map the coordinates of the trajectories into the lattice coordinates in order to obtain a temporal sequence of visited lattice sites in each trial. We are interested in counting the number of times that each subsequent pair of adjacent lattice sites appears in this sequence, regardless of what mouse it came from. This will be used to define the displacement map, detailed in what follows.

The displacement map $\vec{M}(x, y)$ is a vector field defined on the lattice that assigns to each site the preferential direction of movement. This direction is estimated from the trajectories data. Mathematically, it can be written as $\vec{M}(x, y) = [M_h(x, y), M_v(x, y)]$, where M_h is the horizontal component and M_v is the vertical component. The horizontal component expresses the trend to move to the left or right (along the lattice x -axis), and the vertical component in the perpendicular direction, i.e., “up” or “down” in the top view of the lattice of Fig. S2D (along the lattice y -axis).

Each component is a function of the lattice coordinates, (x, y) . The vector $\vec{M}$ is defined as the spatial gradient of the probabilities to move out of (x, y) towards one of its adjacent sites. Thus, its components are given by

$$\begin{aligned} M_h(x, y) &= P_{\rightarrow}(x, y) - P_{\leftarrow}(x, y), \\ M_v(x, y) &= P_{\uparrow}(x, y) - P_{\downarrow}(x, y), \end{aligned} \quad \text{Eq. (1)}$$

where $P_{\rightarrow}(x, y)$ is the probability of stepping right, *i.e.*, from (x, y) to $(x + 1, y)$; $P_{\leftarrow}(x, y)$ is the probability of stepping left [from (x, y) to $(x - 1, y)$]; $P_{\uparrow}(x, y)$ is the probability of stepping up [from (x, y) to $(x, y + 1)$]; and $P_{\downarrow}(x, y)$ is the probability of stepping down [from (x, y) to $(x, y - 1)$]. The probabilities of stepping out of (x, y) must be normalized, therefore $P_{\rightarrow}(x, y) + P_{\leftarrow}(x, y) + P_{\uparrow}(x, y) + P_{\downarrow}(x, y) = 1$. If there is no trajectory going through a particular site, the probability of going from that site into each of the directions is equal to the “null” probability, $P_1 = 1/4$. Time sequences where the mouse does not move are ignored, since we are only interested in the movement between adjacent sites.

The stepping-out probabilities are computed from the mouse trajectory in the following way. This procedure is applied to all mice overlayed together in each individual trial. First, a step is defined as moving from a box at (x, y) to one of its four adjacent boxes, say the one on the right $(x + 1, y)$. Now, if we wanted to calculate the probability of going right from a position (x, y) , we need to go through the sequence of visited lattice sites looking for (x, y) in an instant followed by $(x + 1, y)$ in the immediately next instant. We count the number of times $S_{\rightarrow}(x, y)$ that this pair appears. The count of the transitions of (x, y) to $(x + 1, y)$, in this example, is made regardless of when these transitions happened during the trajectory. The only condition is that the position (x, y) must be immediately followed by $(x + 1, y)$. The same counting is made for the transitions from (x, y) to $(x - 1, y)$ and stored in $S_{\leftarrow}(x, y)$, from (x, y) to $(x, y + 1)$ in $S_{\uparrow}(x, y)$, and from (x, y) to $(x, y - 1)$ in $S_{\downarrow}(x, y)$.

With these sums of steps performed, we can calculate the probability of stepping right, left, up or down. First, note that initially the chance of going to any direction is P_1 , assuming we know nothing about the trajectories. Now, given that we observed $S_{\rightarrow}(x, y)$ steps going to the right, we need to update the chance of going to the right using the union (*i.e.* sum) of the observed chance $S_{\rightarrow}(x, y)/n_s$ with P_1 , where n_s is the total number of steps counted toward any direction in all lattice sites in a given trial:

$$W_{\rightarrow}(x, y) = \frac{S_{\rightarrow}(x, y)}{n_s} + P_0 - \frac{S_{\rightarrow}(x, y)}{n_s} P_0, \quad \text{Eq. (2)}$$

where $W_{\rightarrow}(x, y)$ is the direction weight of the action “step to the right from (x, y) ”. Alternatively, this can be written as a weighted sum of memory-driven stepping with probability 1 and random stepping with probability P_1 :

$$W_{\rightarrow}(x, y) = \frac{S_{\rightarrow}(x, y)}{n_s} + P_0 \left(1 - \frac{S_{\rightarrow}(x, y)}{n_s} \right).$$

We employ **Eq. (2)** because the mouse could, in principle, have chosen any of the other three directions. This ensures that every time the memory is increased (adding $S_{\rightarrow}/n_s$ to the weight), the random contribution for that step decreases; this justifies the last term in **Eq. (2)**, where the intersection between the memory term and the “null” probability, $S_{\rightarrow}P_1/n_s$, is subtracted from the aforementioned union. One can easily see that the formula guarantees that the null probability is automatically recovered, *i.e.* $W_{\rightarrow}(x, y) = P_3$, when there are zero observed steps in a given direction at position (x, y) , and that $W_{\rightarrow}(x, y) = 1$ if all steps are to the right at position (x, y) .

Finally, we use these weights to calculate the probabilities, first defining the total weight of stepping out of (x, y) ,

$$W_{out}(x, y) = W_{\rightarrow}(x, y) + W_{\leftarrow}(x, y) + W_{\uparrow}(x, y) + W_{\downarrow}(x, y), \quad \text{Eq. (3)}$$

and then calculating the probabilities by

$$P_{\rightarrow}(x, y) = \frac{W_{\rightarrow}(x, y)}{W_{out}(x, y)}. \quad \text{Eq. (4)}$$

The same is made for the other directions (left, up and down). **Equation (4)** ensures that the probability of stepping out of (x, y) is always normalized. This is then applied to **Eq. (1)** to obtain the displacement map.

Let us go through the simple example shown in Fig. S2C-F. In this case, the accumulated trajectories in the center box located at $(x, y) = (3, 3)$ are such that there are two passes going up and one pass going down. Thus, each direction has the following weight [given by **Eq. (2)**]: $W_{\uparrow}(3, 3) = \frac{2}{3} + P_0 - \frac{2}{3}P_0 = \frac{3}{4}$, $W_{\downarrow}(3, 3) = \frac{1}{3} + P_0 - \frac{1}{3}P_0 = \frac{1}{2}$, and $W_{\leftarrow}(3, 3) = W_{\rightarrow}(3, 3) = P_0 = \frac{1}{4}$. Applying **Eq. (4)**, we obtain the probabilities $P_{\uparrow}(3, 3) \approx 0.43$, $P_{\downarrow}(3, 3) \approx 0.29$, and $P_{\leftarrow}(3, 3) = P_{\rightarrow}(3, 3) \approx 0.14$. This results in a step map vector $\vec{M}(3, 3) = [0, 0.14]$. This vector represents the preferred direction of movement out of site $(3, 3)$, meaning the mice are likely to pass in the vertical direction, going from bottom to top (hence a positive vertical quantity); and are not picky regarding left or right (hence a null horizontal quantity).

The trajectories of the mice obey box-scaling independently of trial or experimental condition. Box-scaling is a standard method to measure the dimensionality of curves, and is described in Falconer (78). In other words, the number of boxes needed to cover any trajectory scales linearly with the lattice dimension L ; see Fig. S2F. This means that the choice of L is somewhat arbitrary, and we chose $L = 11$ to make each box the size of a few centimeters as expected for the size of a place field from a DG cell (79).

Trajectory directionality statistical significance

The method above gives a single displacement map $\vec{M}(x, y)$ and enables the calculation of the TEV $\vec{D}$ (see below) for each trial of an experiment containing N mice. We have to generate multiple samples for the same experiment using different mice to estimate the significance of the mean displacement direction calculation. We employ a jackknife procedure to achieve this goal. It consists of the “leave one out” rule: this means that a sample of N mice give N unique jackknife samples, each of which containing $N - 1$ unique mice. Then, we get N estimates $\vec{M}_k(x, y)$, with k from 1 to N , by applying the previously described procedure to each jackknife sample of $N - 1$ mice (instead of the whole sample). Finally, we calculate the average, $\vec{M}(x, y) = \frac{1}{N} \sum_{k=1}^N \vec{M}_k(x, y)$ for the displacement map. This is the vector map that we show as small arrows colored from blue to red (**Fig. 4A,B,E,F**; **Fig. 5C**; **Fig. 7A,C,D,G**; **Fig. 8C**; and Fig. S7; we only show statistically significant averages).

The statistical significance of the average $\vec{M}(x, y)$ map in each lattice position (x, y) is built from first principles as follows (see Fig. S3). The main feature of each of the vectors in this map is its direction (*i.e.*, the angle it makes with the positive x-axis). Strong directionality means that the vectors $\vec{M}_k(x, y)$ all pointed roughly in the same direction (*i.e.*, roughly same angles), whereas weak directionality means uniformly distributed angles around the circle (Fig. S3A,B). The stronger the directionality of the sample $\vec{M}_k(x, y)$ the more significant the average $\vec{M}(x, y)$. A simple measure of the spread of the angles (*i.e.*, the directionality strength) is the standard

deviation (S.D.) of the sample angles that $\vec{M}_k(x, y)$ make with the positive x-axis: stronger directionality implies a smaller S.D. (Fig. S3C). Thus, the significance (p) of directionality is how likely it is for a sample of N uniformly distributed angles to display a given S.D. value collectively.

The S.D. of N angles (with zero average) is $\sigma = \sqrt{\frac{1}{N-1} \sum_{k=1}^N \theta_k^2}$, where the θ_k are uniformly distributed from -180° to 180° . If we could determine the probability density function $\varrho(\sigma)$, then the significance would be given by $p(S.D.) = \int_0^{S.D.} \varrho(\sigma) d\sigma$, which is the probability of observing a given S.D. from the sample. In other words, $p(S.D.)$ is the probability that an observed S.D. came from a set of N uniformly distributed angles with zero mean. Thus, given an experimental observation S.D. from the jackknife sample, we will know what is the probability p that the angles from the sample were uniformly distributed in the circle (i.e., have weak directionality). A small p value, therefore, indicates strong directionality (since having small S.D. is very unlikely for uniformly distributed angles – Fig. S3D).

We need to estimate $\varrho(\sigma)$ to be able to calculate p for any S.D. Although it has an integral representation similar to the Chi distribution (Fig. S3E), it is easier to make a numerical estimation: we fix, for instance, $N = 8$ (since we have 8 jackknife samples) and generate 10,000 independent σ values applying the S.D. formula above to N independent uniform angles θ_k from -180° to 180° . The normalized histogram of all the observed σ is a good estimate of $\varrho(\sigma)$ (Fig. S3D). The p value is obtained by numerically integrating $\varrho(\sigma)$ from 0 to the observed S.D. value (Fig. S3E).

Now, it suffices to calculate the S.D. of the angles of the vectors $\vec{M}_k(x, y)$ in the jackknife sample (unbiased to make the average vector $\vec{M}(x, y)$ have $\theta = 0^\circ$). An observed S.D. = 30° corresponds to $p = 10^{-4}$ (i.e., the sample has very strong directionality). We only show in the figures the displacement vectors that have $S.D. \leq 5^\circ$, yielding vanishing $p < 10^{-7}$ (all vectors regardless of significance are shown in Fig. S7 for comparison for random and static entrance experiments). Comparing Fig. S7 to Fig. 4, we notice that, as expected, only vectors from trials that are supposed to have random directionality vanish (i.e., trials for random entrance experiments, and trial 1 for static entrance). The vectors in strongly directed flow trials (such as trial 14 in static entrance experiments) remain. The same happens for the two-food location experiment.

Target Estimation Vector (TEV)

The target estimation vector (TEV) is an estimate of the position of the target based only on the observed trajectories and hole checks of the mice. We write the average TEV as $\vec{D} = \frac{1}{N} \sum_{k=1}^N \vec{D}_k$, and $\vec{D}_k = D \hat{u}_k$ is the TEV for a particular jackknife sample k . The magnitude D is inferred from the hole checks and the direction $\hat{u}_k$ is calculated from the displacement map $\vec{M}_k(x, y)$ as follows. The total sum of the displacement map $\vec{M}_k(x, y)$ over all arena sites (x, y) gives the estimate of the learned direction $\hat{u}_k = \vec{u}_k / |\vec{u}_k|$, with $\vec{u}_k = \sum_{(x,y)} \vec{M}_k(x, y)$. The magnitude D is the distance from start to the average of the hole check distribution restricted to 20cm around the target (i.e., D is the distance from the start to the “x” in the green ellipsis in Fig. 4G,H; Fig. 5C; Fig. 7B,F,H; Fig. 8C; and Fig. S1B-F). In the case of the Probe B-A trial, we use D as the distance from B (instead of “start”) to the restricted average of the hole check distribution around target A. In trials where no preferred direction is detected (such as in the random entrances experiment, or in the first trial after switching from target A to B), the magnitude $|\vec{u}_k|$ is already roughly the distance from start to the arena’s center (expressing the randomness in the displacement map), and hence we take the TEV to be just $\vec{D}_k = \vec{u}_k$ instead.

The error in the magnitude D of the TEV is obtained from the covariance matrix C of the spatial distribution of hole checks restricted to less than 20cm of the target. It is given by $\sigma_D = \sqrt{\lambda_1 + \lambda_2}$, where $\lambda_{1,2}$ are the eigenvalues of C . In the plots [Figs. 4, 5, 7, 8], we simply represent the covariance matrix by the green ellipses in the same way we do for the uncensored distribution

(see Active sensing and hole-check detection). The error in the direction of $\vec{D}$ is simply the S.D. of the angles of $\vec{D}_k$ with respect to the positive x-axis (calculated by shifting $\vec{D}$ to 0°, and making the angles of $\vec{D}_k$ between -180° and 180°, similarly to what is done for the displacement map). This gives the pink shaded circular sector accompanying the TEV in the figures of the displacement maps and the error bars in the TEV-target deviation plot (**Figs. 4I** [↗](#) and **7I** [↗](#)).

Data Availability

All codes and data are available in <https://github.com/neuro-physics/mouse-cogmap> [↗](#)

References

1. Rosenberg M., Zhang T., Perona P., Meister M. (2021) **Mice in a labyrinth show rapid learning, sudden insight, and efficient exploration** *eLife* **10**
2. Chan E., Baumann O., Bellgrove M. A., Mattingley J. B. (2012) **From objects to landmarks: the function of visual location information in spatial navigation** *Frontiers in psychology* **3**
3. Goodman J. (2020) **Place vs. Response Learning: History, Controversy, and Neurobiology** *Front Behav Neurosci* **14**
4. Nyberg N., Duvelle E., Barry C., Spiers H. J. (2022) **Spatial goal coding in the hippocampal formation** *Neuron* **110**:394–422
5. Tolman E. C., Ritchie B. F., Kalish D. (1946) **Studies in spatial learning; place learning versus response learning** *J Exp Psychol* **36**:221–229
6. Collett T. S., Cartwright B. A., Smith B. A. (1986) **Landmark learning and visuo-spatial memories in gerbils** *J Comp Physiol A* **158**:835–851
7. Morris R. G. M. (1981) **Spatial Localization Does Not Require the Presence of Local Cues** *LEARNING AND MOTIVATION* **12**:239–260
8. McNaughton B. L. *et al.* (1996) **Deciphering the hippocampal polyglot: the hippocampus as a path integration system** *J Exp Biol* **199**:173–185
9. Mittelstaedt M. L., Mittelstaedt H. (1980) **Homing by path integration in a mammal** *Naturwissenschaften* **67**:566–567
10. Burgess N. (2006) **Spatial memory: how egocentric and allocentric combine** *Trends Cogn Sci* **10**:551–557
11. Chen G., King J. A., Burgess N., O’Keefe J. (2013) **How vision and movement combine in the hippocampal place code** *Proc Natl Acad Sci U S A* **110**:378–383
12. Knierim J. J., Hamilton D. A. (2011) **Framing spatial cognition: neural representations of proximal and distal frames of reference and their roles in navigation** *Physiol Rev* **91**:1245–1279
13. Tolman E. C. (1948) **Cognitive maps in rats and men** *Psychological review* **55**:189–208
14. Tolman E. C., Ritchie B. F., Kalish D. (1946) **Studies in spatial learning: Orientation and the shortcut** *J Exp Psychol* **36**:13–24
15. McNaughton B. L., Battaglia F. P., Jensen O., Moser E. I., Moser M. B. (2006) **Path integration and the neural basis of the ‘cognitive map’** *Nat Rev Neurosci* **7**:663–678
16. O’Keefe J., Nadel L. (1978) **The Hippocampus as a Cognitive Map**

17. Banino A. *et al.* (2018) **Vector-based navigation using grid-like representations in artificial agents** *Nature* **557**:429–433
18. Fonio E., Benjamini Y., Golani I. (2009) **Freedom of movement and the stability of its unfolding in free exploration of mice** *Proc Natl Acad Sci U S A*
19. Asumbisa K., Peyrache A., Trenholm S. (2022) **Flexible cue anchoring strategies enable stable head direction coding in both sighted and blind animals** *Nature communications* **13**
20. Fischler-Ruiz W., Clark D. G., Joshi N. R., Devi-Chou V., Kitch L., Schnitzer M., Abbott L. F., Axel R. (2021) **Olfactory landmarks and path integration converge to form a cognitive spatial map** *Neuron* **109**:4036–4049
21. Cushman J. D., Aharoni D. B., Willers B., Ravassard P., Kees A., Vuong C., Popeney B., Arisaka K., Mehta M. R. (2013) **Multisensory control of multimodal behavior: do the legs know what the tongue is doing?** *PLoS One* **8**
22. Saleem A. B., Busse L. (2023) **Interactions between rodent visual and spatial systems during navigation** *Nat Rev Neurosci* **24**:487–501
23. Jacobs A. L., Fridman G., Douglas R. M., Alam N. M., Latham P. E., Prusky G. T., Nirenberg S. (2009) **Ruling out and ruling in neural codes** *Proc Natl Acad Sci U S A* **106**:5936–5941
24. Long M., Jiang W., Liu D., Yao H. (2015) **Contrast-dependent orientation discrimination in the mouse** *Sci Rep* **5**
25. Prusky G. T., Douglas R. M. (2004) **Characterization of mouse cortical spatial vision** *Vision Res* **44**:3411–3418
26. Prusky G. T., West P. W., Douglas R. M. (2000) **Behavioral assessment of visual acuity in mice and rats** *Vision Res* **40**:2201–2209
27. Saleem A. B., Diamanti E. M., Fournier J., Harris K. D., Carandini M. (2018) **Coherent encoding of subjective spatial position in visual cortex and hippocampus** *Nature*
28. Chapillon P., Rouillet P. (1996) **Use of proximal and distal cues in place navigation by mice changes during ontogeny** *Dev Psychobiol* **29**:529–545
29. Hebert M., Bulla J., Vivien D., Agin V. (2017) **Are Distal and Proximal Visual Cues Equally Important during Spatial Learning in Mice? A Pilot Study of Overshadowing in the Spatial Domain** *Front Behav Neurosci* **11**
30. Rogers J., Churilov L., Hannan A. J., Renoir T. (2017) **Search strategy selection in the Morris water maze indicates allocentric map formation during learning that underpins spatial memory formation** *Neurobiol Learn Mem* **139**:37–49
31. Youngstrom I. A., Strowbridge B. W. (2012) **Visual landmarks facilitate rodent spatial navigation in virtual reality environments** *Learn Mem* **19**:84–90
32. Biegler R., Morris R. (1996) **Landmark stability: studies exploring whether the perceived stability of the environment influences spatial representation** *J Exp Biol* **199**:187–193
33. Biegler R., Morris R. G. (1993) **Landmark stability is a prerequisite for spatial but not discrimination learning** *Nature* **361**:631–633

34. Biegler R., Morris R. G. (1996) **Landmark stability: further studies pointing to a role in spatial learning** *Q J Exp Psychol B* **49**:307–345
35. Jeffery K. J. (1998) **Learning of landmark stability and instability by hippocampal place cells** *Neuropharmacology* **37**:677–687
36. Knierim J. J., Kudrimoti H. S., McNaughton B. L. (1995) **Place cells, head direction cells, and the learning of landmark stability** *J Neurosci* **15**:1648–1659
37. Save E., Nerad L., Poucet B. (2000) **Contribution of multiple sensory information to place field stability in hippocampal place cells** *Hippocampus* **10**:64–76
38. Zhang S., Schonfeld F., Wiskott L., Manahan-Vaughan D. (2014) **Spatial representations of place cells in darkness are supported by path integration and border information** *Front Behav Neurosci* **8**
39. Burlingham C. S., Heeger D. J. (2020) **Heading perception depends on time-varying evolution of optic flow** *Proc Natl Acad Sci U S A* **117**:33161–33169
40. Horrocks E. A. B., Mareschal I., Saleem A. B. (2023) **Walking humans and running mice: perception and neural encoding of optic flow during self-motion** *Philos Trans R Soc Lond B Biol Sci* **378**
41. Acharya L., Aghajani Z. M., Vuong C., Moore J. J., Mehta M. R. (2016) **Causal Influence of Visual Cues on Hippocampal Directional Selectivity** *Cell* **164**:197–207
42. Saleem A. B. (2020) **Two stream hypothesis of visual processing for navigation in mouse** *Current Opinion in Neurobiology* **64**:70–78
43. Cullen K. E. (2019) **Vestibular processing during natural self-motion: implications for perception and action** *Nat Rev Neurosci* **20**:346–363
44. Mohammadi M., Carriot J., Mackrous I., Cullen K. E., Chacron M. J. (2024) **Neural populations within macaque early vestibular pathways are adapted to encode natural self-motion** *PLoS Biol* **22**
45. Cullen K. E., Taube J. S. (2017) **Our sense of direction: progress, controversies and challenges** *Nat Neurosci* **20**:1465–1473
46. Taube J. S. (2007) **The head direction signal: origins and sensory-motor integration** *Annu Rev Neurosci* **30**:181–207
47. Mao D., Molina L. A., Bonin V., McNaughton B. L. (2020) **Vision and Locomotion Combine to Drive Path Integration Sequences in Mouse Retrosplenial Cortex** *Curr Biol* **30**:1680–1688
48. Yang X., Cacucci F., Burgess N., Wills T. J., Chen G. (2024) **Visual boundary cues suffice to anchor place and grid cells in virtual reality** *Curr Biol*
49. Engelmann J., Wallach A., Maler L. (2021) **Linking active sensing and spatial learning in weakly electric fish** *Curr Opin Neurobiol* **71**:1–10
50. Etienne A. S., Maurer R., Seguinot V. (1996) **Path integration in mammals and its interaction with visual landmarks** *J Exp Biol* **199**:201–209

51. Jun J. J., Longtin A., Maler L. (2014) **Enhanced sensory sampling precedes self-initiated locomotion in an electric fish** *J Exp Biol* **217**:3615–3628
52. Lai A. T., Espinosa G., Wink G. E., Angeloni C. F., Dombeck D. A., MacIver M. A. (2024) **A robotrodent interaction arena with adjustable spatial complexity for ethologically relevant behavioral studies** *Cell Rep* **43**
53. Etienne A. S., Jeffery K. J. (2004) **Path integration in mammals** *Hippocampus* **14**:180–192
54. Jun J. J., Longtin A., Maler L. (2016) **Active sensing associated with spatial learning reveals memory-based attention in an electric fish** *Journal of Neurophysiology* **115**:2577–2592
55. Mirmiran C., Fraser M., Maler L. (2022) **Finding food in the dark: how trajectories of a gymnotiform fish change with spatial learning** *J Exp Biol* **225**
56. Wallach A., Harvey-Girard E., Jun J. J., Longtin A., Maler L. (2018) **A time-stamp mechanism may provide temporal information necessary for egocentric to allocentric spatial transformations** *eLife* **7**
57. Jayakumar R. P., Madhav M. S., Savelli F., Blair H. T., Cowan N. J., Knierim J. J. (2019) **Recalibration of path integration in hippocampal place cells** *Nature* **566**:533–537
58. Savelli F., Knierim J. J. (2019) **Origin and role of path integration in the cognitive representations of the hippocampus: computational insights into open questions** *J Exp Biol* **222**
59. Save E., Cressant A., Thinus-Blanc C., Poucet B. (1998) **Spatial firing of hippocampal place cells in blind rats** *J Neurosci* **18**:1818–1826
60. Aghajan Z. M., Acharya L., Moore J. J., Cushman J. D., Vuong C., Mehta M. R. (2015) **Impaired spatial selectivity and intact phase precession in two-dimensional virtual reality** *Nat Neurosci* **18**:121–128
61. Ravassard P., Kees A., Willers B., Ho D., Aharoni D., Cushman J., Aghajan Z. M., Mehta M. R. (2013) **Multisensory control of hippocampal spatiotemporal selectivity** *Science* **340**:1342–1346
62. Moore J. J., Cushman J. D., Acharya L., Popeney B., Mehta M. R. (2021) **Linking hippocampal multiplexed tuning, Hebbian plasticity and navigation** *Nature* **599**:442–448
63. Sarel A., Finkelstein A., Las L., Ulanovsky N. (2017) **Vectorial representation of spatial goals in the hippocampus of bats** *Science* **355**:176–180
64. Ormond J., O’Keefe J. (2022) **Hippocampal place cells have goal-oriented vector fields during navigation** *Nature* **607**:741–746
65. Jercog P. E., Ahmadian Y., Woodruff C., Deb-Sen R., Abbott L. F., Kandel E. R. (2019) **Heading direction with respect to a reference point modulates place-cell activity** *Nature communications* **10**
66. Bittner K. C., Grienberger C., Vaidya S. P., Milstein A. D., Macklin J. J., Suh J., Tonegawa S., Magee J. C. (2015) **Conjunctive input processing drives feature selectivity in hippocampal CA1 neurons** *Nat Neurosci* **18**:1133–1142

67. Bittner K. C., Milstein A. D., Grienberger C., Romani S., Magee J. C. (2017) **Behavioral time scale synaptic plasticity underlies CA1 place fields** *Science* **357**:1033–1036
68. Milstein A. D., Li Y., Bittner K. C., Grienberger C., Soltesz I., Magee J. C., Romani S. (2021) **Bidirectional synaptic plasticity rapidly modifies hippocampal representations** *eLife* **10**
69. Monaco J. D., Rao G., Roth E. D., Knierim J. J. (2014) **Attentive scanning behavior drives one-trial potentiation of hippocampal place fields** *Nat Neurosci* **17**:725–731
70. Bennett A. T. (1996) **Do animals have cognitive maps?** *J Exp Biol* **199**:219–224
71. Grieves R., Dudchenko P. (2012) **Cognitive maps and spatial inference in animals: Rats fail to take a novel shortcut, but can take a previously experienced one** *Learning and Motivation* **84**:81–92
72. Benhamou S. (1996) **No evidence for cognitive mapping in rats** *Animal Behaviour* **52**:201–212
73. Shamash P., Olesen S. F., Iordanidou P., Campagner D., Banerjee N., Branco T. (2021) **Mice learn multi-step routes by memorizing subgoal locations** *Nat Neurosci* **24**:1270–1279
74. Landau B., Spelke E., Gleitman H. (1984) **Spatial knowledge in a young blind child** *Cognition* **16**:225–260
75. Kesner R. P., Farnsworth G., Kametani H. (1991) **Role of parietal cortex and hippocampus in representing spatial information** *Cereb Cortex* **1**:367–373
76. Kronfeld-Schor N., Dominoni D., de la Iglesia H., Levy O., Herzog E. D., Dayan T., HelfrichForster C. (2013) **Chronobiology by moonlight** *Proc Biol Sci* **280**
77. N. Upham S., Hafner J. C. (2013) **Do nocturnal rodents in the great basin desert avoid moonlight?** *Journal of Mammology* **94**:59–72
78. Falconer K. (2004) **Fractal Geometry: Mathematical Foundations and Application**
79. GoodSmith D., Chen X., Wang C., Kim S. H., Song H., Burgalossi A., Christian K. M., Knierim J. J. (2017) **Spatial Representations of Granule Cells and Mossy Cells of the Dentate Gyrus** *Neuron* **93**:677–690
80. Hair J. F., Black W. C., Babin B. J., Anderson R. E. (2013) **Multivariate Data Analysis**
81. Tomé T., Oliveira M. J. (2015) **Stochastic Dynamics and Irreversibility**

Editors

Reviewing Editor

Noah Cowan

Johns Hopkins University, Baltimore, United States of America

Senior Editor

Kate Wassum

University of California, Los Angeles, Los Angeles, United States of America

Reviewer #1 (Public review):

Assessment:

This important work advances our understanding of navigation and path integration in mammals by using a clever behavioral paradigm. The paper provides compelling evidence that mice are able to create and use a cognitive map to find "short cuts" in an environment, using only the location of rewards relative to the point of entry to the environment and path integration, and need not rely on visual landmarks.

Summary:

The authors have designed a novel experimental apparatus called the 'Hidden Food Maze (HFM)' and a beautiful suite of behavioral experiments using this apparatus to investigate the interplay between allothetic and idiothetic cues in navigation. The results presented provide a clear demonstration of the central claim of the paper, namely that mice only need a fixed start location and path integration to develop a cognitive map. The experiments and analyses conducted to test the main claim of the paper -- that the animals have formed a cognitive map -- are conclusive. While I think the results are quite interesting and sound, one issue that needs to be addressed is the framing how landmarks are used (or not), as discussed below, although I believe this will be a straight forward issue for the authors to address.

Strengths:

The 90 degree rotationally symmetric design and use of 4 distal landmarks and 4 quadrants with their corresponding rotationally equivalent locations (REL) lends itself to teasing apart the influence of path integration and landmark-based navigation in a clever way. The authors use a really complete set of experiments and associated controls to show that mice can use a start location and path integration to develop a cognitive map and generate shortcut routes to new locations.

Weaknesses:

There were no major weaknesses identified that were not addressed during revisions.

<https://doi.org/10.7554/eLife.95764.2.sa2>

Reviewer #3 (Public review):

Summary:

How is it that animals find learned food locations in their daily life? Do they use landmarks to home in on these learned locations or do they learn a path based on self-motion (turn left, take ten steps forward, turn right, etc.). This study carefully examines this question in a well designed behavioral apparatus. A key finding is that to support the observed behavior in the hidden food arena, mice appear to not use the distal cues that are present in the environment for performing this task. Removal of such cues did not change the learning rate, for example. In a clever analysis of whether the resulting cognitive map based on self-motion cues could allow a mouse to take a shortcut, it was found that indeed they are. The work nicely shows the evolution of the rodent's learning of the task, and the role of active sensing in the targeted reduction of uncertainty of food location proximal to its expected location.

Strengths:

A convincing demonstration that mice can synthesize a cognitive map for the finding of a static reward using body frame-based cues. Showing that uncertainty of final target location is resolved by an active sensing process of probing holes proximal to the expected location.

Showing that changing the position of entry into the arena rotates the anticipated location of the reward in a manner consistent with failure to use distal cues.

Weaknesses:

The task is low stakes, and thus the failure to use distal cues at most costs the animal a delay in finding the food; this delay is likely unimportant to the animal, and the pre-training procedure is likely to make it clear to the animal's that distal cues are unreliable even if desirable to use. Thus, it is unclear whether this result would generalize to a situation where the animal may be under some time pressure, urgency due to food (or water) restriction, or due to predatory threat, or situations where distal cues are reliable. In such cases, the use of distal cues to make locating the reward robust to changing start locations may be more likely to be observed.

<https://doi.org/10.7554/eLife.95764.2.sa1>

Author response:

The following is the authors' response to the original reviews.

We would like to thank the reviewers and editors for their careful assessment and review of our article. The many detailed comments, questions and suggestions were very helpful in improving our analyses and presentation of data. In particular, our Discussion benefited enormously from the comments.

Below we respond in detail to every point raised.

We especially note that Reviewer #3's small query on "trial where learning is defined to have occurred, we were not given the quantitative criterion operationalizing "learning" - please provide" led to deeper analyses and insights and a lengthy response.

This analysis prompted the addition of a sentence (red) to the Abstract.

"Animals navigate by learning the spatial layout of their environment. We investigated spatial learning of mice in an open maze where food was hidden in one of a hundred holes. Mice leaving from a stable entrance learned to efficiently navigate to the food without the need for landmarks. We developed a quantitative framework to reveal how the mice estimate the food location based on analyses of trajectories and active hole checks. After learning, the computed "target estimation vector" (TEV) closely approximated the mice's route and its hole check distribution. The TEV required learning both the direction and distance of the start to food vector, and our data suggests that different learning dynamics underlie these estimates. We propose that the TEV can be precisely connected to the properties of hippocampal place cells. Finally, we provide the first demonstration that, after learning the location of two food sites, the mice took a shortcut between the sites, demonstrating that they had generated a cognitive map."

Note: we added, at the end of the manuscript, the legends for the Shortcut video (Video 1) and the main text figure legends; these are with a larger font and so easier to read.

Reviewer #1 (Public Review):

Assessment:

This important work advances our understanding of navigation and path integration in mammals by using a clever behavioral paradigm. The paper provides compelling evidence that mice are able to create and use a cognitive map to find "short cuts" in an

environment, using only the location of rewards relative to the point of entry to the environment and path integration, and need not rely on visual landmarks.

Thank you.

Summary:

The authors have designed a novel experimental apparatus called the 'Hidden Food Maze (HFM)' and a beautiful suite of behavioral experiments using this apparatus to investigate the interplay between allothetic and idiothetic cues in navigation. The results presented provide a clear demonstration of the central claim of the paper, namely that mice only need a fixed start location and path integration to develop a cognitive map. The experiments and analyses conducted to test the main claim of the paper -- that the animals have formed a cognitive map -- are conclusive. While I think the results are quite interesting and sound, one issue that needs to be addressed is the framing of how landmarks are used (or not), as discussed below, although I believe this will be a straightforward issue for the authors to address.

We have now added detailed discussion on this important point. See below.

Strengths:

The 90-degree rotationally symmetric design and use of 4 distal landmarks and 4 quadrants with their corresponding rotationally equivalent locations (REL) lends itself to teasing apart the influence of path integration and landmark-based navigation in a clever way. The authors use a really complete set of experiments and associated controls to show that mice can use a start location and path integration to develop a cognitive map and generate shortcut routes to new locations.

Weaknesses:

I have two comments. The second comment is perhaps major and would require rephrasing multiple sentences/paragraphs throughout the paper.

(1) The data clearly indicate that in the hidden food maze (HFM) task mice did not use external visual "cue cards" to navigate, as this is clearly shown in the errors mice make when they start trials from a different start location when trained in the static entrance condition. The absence of visual landmark-guided behavior is indeed surprising, given the previous literature showing the use of distal landmarks to navigate and neural correlates of visual landmarks in hippocampal formation. While the authors briefly mention that the mice might not be using distal landmarks because of their pretraining procedure - I think it is worth highlighting this point (about the importance of landmark stability and citing relevant papers) and elaborating on it in greater detail. It is very likely that mice do not use the distal visual landmarks in this task because the pretraining of animals leads to them not identifying them as stable landmarks. For example, if they thought that each time they were introduced to the arena, it was "through the same door", then the landmarks would appear to be in arbitrary locations compared to the last time. In the same way, we as humans wouldn't use clouds or the location of people or other animate objects as trusted navigational beacons. In addition, the animals are introduced to the environment without any extra-maze landmarks that could help them resolve this ambiguity. Previous work (and what we see in our dome experiments) has shown that in environments with 'unreliable' landmarks, place cells are not controlled by landmarks - <https://www.sciencedirect.com/science/article/pii/S0028390898000537>, <https://pubmed.ncbi.nlm.nih.gov/7891125/>. This makes it likely that the absence of these distal visual landmarks when the animal first entered the maze ensured that the animal does not 'trust' these visual features as landmarks.

Thank you. We have added many references and discussion exactly on this point including both direct behavioral experiments as well as discussion on the effects of landmark (in)stability of place cell encoding of “place”. See Page 18 third paragraph.

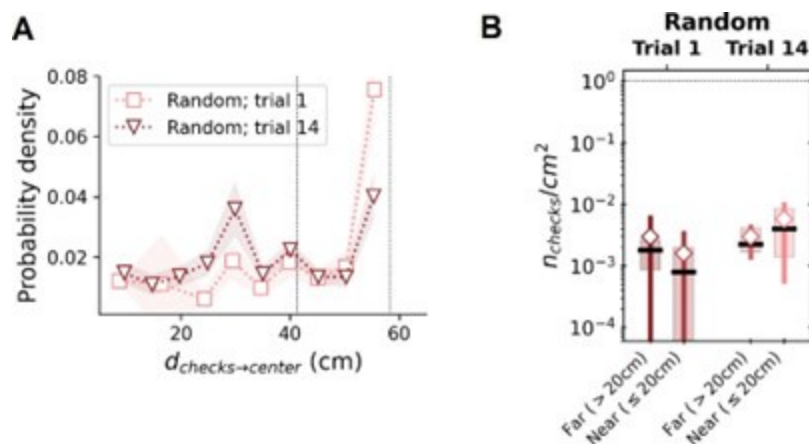
“An alternate factor might be the lack of reliability of distal spatial cues in predicting the food location. The mice, during pretraining trials, learned to find multiple food locations without landmarks. In the random trials, the continuous change of relative landmark location may lead the mice to not identifying them as “stable landmarks”. This view is supported by behavioral experiments that showed the importance of landmark stability for spatial learning (32-34) and that place cells are not controlled by “unreliable landmarks” (35-38). Control experiments without landmarks (Fig. S6A,B) or in the dark (Fig. S6C-F) confirmed that the mice did not need landmarks for spatial learning of the food location.”

(2) I don't agree with the statement that 'Exogenous cues are not required for learning the food location'. There are many cues that the animal is likely using to help reduce errors in path integration. For example, the start location of the rat could act as a landmark/exogenous cue in the sense of partially correcting path integration errors. The maze has four identical entrances (90-degree rotationally symmetric). Despite this, it is entirely plausible that the animal can correct path integration errors by identifying the correct start entrance for a given trial, and indeed the distance/bearing to the others would also help triangulate one's location. Further, the overall arena geometry could help reduce PI error. For example, with a food source learned to be "near the middle" of the arena, the animal would surely not estimate the position to be near the far wall (and an interesting follow-on experiment would be to have two different-sized, but otherwise nearly identical arenas). As the rat travels away from the start location, small path integration errors are bound to accumulate, these errors could be at least partially corrected based on entrance and distal wall locations. If this process of periodically checking the location of the entrance to correct path integration errors is done every few seconds, path integration would be aided 'exogenously' to build a cognitive map. While the original claim of the paper still stands, i.e. mice can learn the location of a hidden food size when their starting point in the environment remains constant across trials. I would advise rewording portions of the paper, including the discussion throughout the paper that states claims such as "Exogenous cues are not required for learning the food location" to account for the possibility that the start and the overall arena geometry could be used as helpful exogenous cues to correct for path integration errors.

We agree with the referee that our claim was ill-phrased. Surely the behavior of the mouse must be constrained by the arena size to some extent. To minimize potential geometric cues from the arena, we carefully analyzed many preliminary experiments (each with a unique batch of 4 mice) having the target positioned at different locations. We added a paragraph to the section “Further controls” where we explain our choice for the target position. Page 12 last paragraph; Page 13 “Arena geometry” paragraph.

Also, following the suggestion from the reviewer, we probed whether the hole checks accumulated near the center of the arena for the random entrance mice, as a potential sign that some spatial learning is going on. In fact, neither the density of hole checks, nor the distance of the hole checks to the center of the arena change with learning: panel A below shows the probability density of finding a hole check at a given distance from the center of the arena; both trial 1 and trial 14 have very similar profiles. Panel B shows the density of hole checks near (<20cm) and far (>20cm) from the arena's center.

Author response image 1.



It also doesn't show any significant differences between trials 1 and 14.

So even though there's some trend (in panel A, the peak goes from 60cm to a double peak, one at 30cm away from the center, and the other still at 60cm), the distance from the center is still way too large compared to the mouse's body size and to the average inter-hole distance (<10cm). These panels are now in the Supplementary Figure S8B.

Finally, we enhanced the wording in our claim. We now have a new section entitled: "What cues are required for learning the food location?". There, we systematically cover all possible cues and how they might be affected by their stability under the perturbation of maze floor rotation.

Reviewer #2 (Public Review):

Summary:

This manuscript reports interesting findings about the navigational behavior of mice. The authors have dissected this behavior in various components using a sophisticated behavioral maze and statistical analysis of the data.

Strengths:

The results are solid and they support the main conclusions, which will be of considerable value to many scientists.

Thank you.

Weaknesses:

Figure 1: In some trials the mice seem to be doing thigmotaxis, walking along the perimeter of the maze. This is perhaps due to the fear of the open arena. But, these paths along the perimeter would significantly influence all metrics of navigation, e.g. the distance or time to reward.

Perhaps analysis can be done that treats such behavior separately and the factors it out from the paths that are away from the perimeter.

In Page 4, we added a small section entitled: "Pretraining trials". Our reference was suggested by Reviewer #3 (noted as "Golani" with first author "Fonio"). Our preliminary experiments

used naïve mice and they typically took greater than 2 days before they ventured into the arena center and found the single filled hole. This added unacceptable delays and the Pretraining trials greatly diminished the extensive thigmotaxis (not quantified). The “near the walls” trajectories did continue in the first learning trial (Fig. 2A, 3A) but then diminished in subsequent trials. We found no evidence that thigmotaxis (trajectories adjacent to the wall) were a separate category of trajectory.

Figure 1c: the color axis seems unusual. Red colors indicate less frequently visited regions (less than 25%) and white corresponds to more frequently visited places (>25%)? Why use such a binary measure instead of a graded map as commonly done?

Thank you; you are completely correct. We have completely changed the color coding.

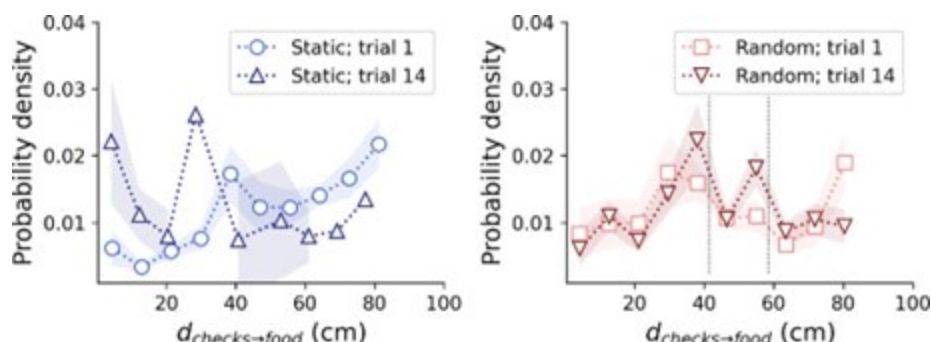
Some figures use linear scale and others use logarithmic scale. Is there a scientific justification? For example, average latency is on a log scale and average speed is on a linear scale, but both quantify the same behavior. The y-axis in panel 1-I is much wider than the data. Is there a reason for this? Or can the authors zoom into the y-axis so that the reader can discern any pattern?

We use logarithmic scale with the purpose of displaying variables that have a wide range of variation (mainly, distance, latency, and number of hole checks, since it linearly and positively correlates with both distance and latency – see new Fig. S4B,C). For example, Latency goes from hundreds of seconds (trial 1) to just a few seconds (trial 14). Similarly, the total distance goes from hundreds of centimeters (trial 1, sometimes more than 1000cm, see answer about the 10-fold variation of distance below) to just the start-target distance (which is ~100cm). These variables vary over a few orders of magnitude. We display speed in a linear axis because it does not increase for more than one order of magnitude.

Moreover, fitting the wide-ranged data (distance, latency, nchecks) yields smaller error in logscale [i.e., fitting $\log(y)$ vs. trial, instead of y vs. trial]. In these cases, the log-scale also helps visualizing how well the data was fitted by the curve. Thus, presenting wide-ranged data in linear scale could be misleading regarding goodness of fit.

We now zoomed into the Y axis scale in Panels I of Fig. 2 and Fig. 3. We kept it in log-scale, but linear Y scale produces Author response image 2 for Figs. 3I and 2I, respectively.

Author response image 2.



Thus, we believe that the loglog-scale in these panels won't compromise the interpretation of the phenomenon. In fact, the loglog of the static case suggests that the probability of hole checking distance increases according to a power law as the mouse approaches the target (however, we did not check this thoroughly, so we did not include this point in the discussion). Power law behavior is observed in other animals (e.g. ants: DOI:

10.1371/journal.pone.0009621) and is sometimes associated with a stochastic process with memory.

1F shows no significant reduction in distance to reward. Does that mean there is no improvement with experience and all the improvement in the latency is due to increasing running speed with experience?

Correct and in the section “Random Entrance experiments” under “Results” (Page 5) we explicitly note this point.

“We hypothesize that the mice did not significantly reduce their distance travelled (Fig. 2A,B,F) because they had not learned the food location - the decrease in latency (Fig. 2D) was due to its increased running speed and familiarity with non-spatial task parameters.”

Figure 3: The distance traveled was reduced by nearly 10-fold and speed increased by by about 3fold. So, the time to reach the reward should decrease by only 3 fold ($t=d/v$) but that too reduced by 10fold. How does one reconcile the 3fold difference between the expected and observed values?

The traveled distance is obtained by linearly interpolating the sampled trajectory points. In other words, the software samples a discrete set of positions, $\vec{r}_t = (x_t, y_t)$ for each recorded instant t . The total distance is

$$d_{total} = \sum^{Latency} |\vec{r}_t - \vec{r}_{t-1}|$$

where $\Delta d_t = |\vec{r}_t - \vec{r}_{t-1}| = \sqrt{(x_t - x_{t-1})^2 + (y_t - y_{t-1})^2}$ is the Euclidean distance between two

consecutively sampled points. However, the same result (within a fraction of cm error) can be obtained by integrating the sampled speed over time v ! using the Simpson method

$$d_{total} = \int_0^{Latency} v_t dt$$

Since Latency varies by 10-fold, it is just expected that, given $d = vt$, the total distance will also

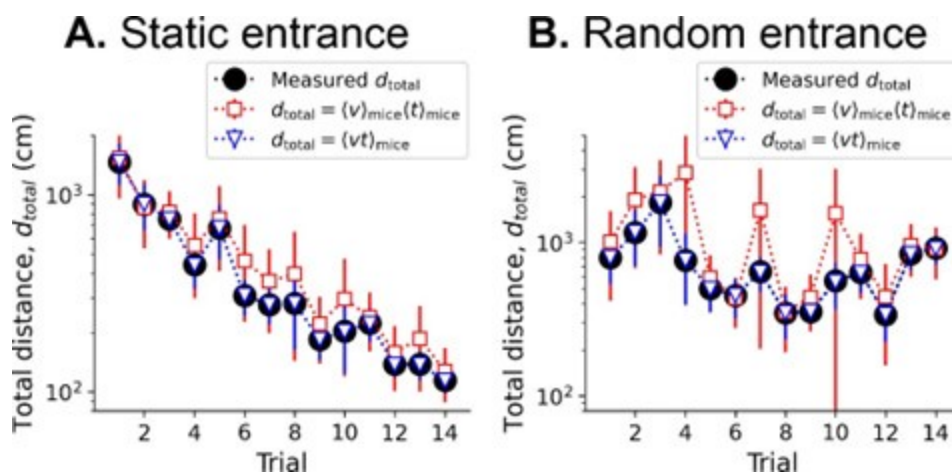
vary by 10-fold (since $v_t = \frac{\Delta d_t}{\Delta t} = \frac{|\vec{r}_t - \vec{r}_{t-1}|}{\Delta t}$ is constant in each time interval Δt ; replacing v ! in

the integral yields the discrete sum above).

The correctness of our kinetic measurements can be simply verified by multiplying the data from the Latency panel with the data from the Velocity panel. If this results in the Distance plot, then there is no discrepancy.

In Author response image 3, we show the actual measured distance, d_{total} , for both conditions (random and static entrance), calculated with the discrete sum above (black filled circles).

Author response image 3.



We compare this with two quantities: (a) average speed multiplied by average latency (red squares); and (b) average of the product of speed by latency (blue inverted triangles). The averages are taken over mice. Notice that if the multiplication is taken before the average (as it should be done), then the product $\langle vt \rangle$ is indistinguishable from the total distance obtained by linear interpolation. Even taking the averages prior to the multiplication (which is physically incorrect, since speed and latency are properties of each individual mouse), yields almost exactly the same result (well within 1 standard deviation).

The only thing to keep in mind here is that the Distance panel in the paper presents the normalized distance according to the target distance to the starting point. This is necessary because in the random entrance experiments, each mouse can go to 1 of 4 possible targets (each of which has a different distance to the starting point).

Figure 4: The reader is confused about the use of a binary color scheme here for the checking behavior: gray for a large amount of checking, and pink for small. But, there is a large ellipse that is gray and there are smaller circles that are also gray, but these two gray areas mean very different things as far as the reader can tell. Is that so? Why not show the entire graded colormap of checking probability instead of such a seemingly arbitrary binary depiction?

Thank you. Our coloring scheme was indeed poorly thought out and we have changed it. Hopefully the reviewer now finds it easier to interpret. The frequency of hole checks is now encoded into only filled circles of varying sizes and shades of pink. Small empty circles represent the arena holes (empty because they have no food); The large transparent gray ellipse is the variance of the unrestricted spatial distribution of hole checks.

Figure 4C: What would explain the large amount of checking behavior at the perimeter? Does that occur predominantly during thigmotaxis?

Yes. As mentioned above, thigmotaxis still occurs in the first trial of training. The point to note is that the hole checking shown in Fig. 4C is over all the mice so that, per mice, it does not appear so overwhelming.

Was there a correlation between the amount of time spent by the animals in a part of the maze and the amount of reward checking? Previous studies have shown that the two

behaviors are often positively correlated, e.g. reference 20 in the manuscript. How does this fit with the path integration hypothesis?

We thank the reviewer for pointing this out. Indeed, the time spent searching & the hole checking behavior are correlated. We added a new panel C to Fig. S4 showing a raw correlation plot between Latency and number of checks.

Also, in the last paragraph of the “Revealing the mouse estimate of target position from behavior” section under “Results”), we now added a sentence relating the findings in Fig. 4H and 4K (spatial distribution of hole checks, and density of checks near the target, respectively) to note that these findings are in agreement with Fig 3C (time spent searching in each quadrant).

“The mean position of hole checks near (20cm) the target is interpreted as the mouse estimated target (Fig. 4C,D,G,H; green + sign=mean position; green ellipses = covariance of spatial hole check distribution restricted to 20cm near the target). This finding together with the displacement and spatial hole check maps (Figs. 4F and 4H, respectively) corroborates the heatmap of time spent in the target quadrant (Fig. 3C), suggesting a positive correlation between hole checks and time searching (see also Fig. S4C).”

"Scratches and odor trails were eliminated by washing and rotating the maze floor between trials." Can one eliminate scratches by just washing the maze floor? Rotation of the maze floor between trials can make these cues unreliable or variable but will not eliminate them. Ditto for odor cues.

The upper arena floor is rotated between trials so that any scratches will not be stable cues. We clarified this in the Discussion about potential cues.

See “What cues are required for learning the food location?”

*"Possible odor gradient cues were eliminated by experiments where such gradients were prevented with vacuum fans (Fig. S6E)" What tests were done to ensure that these were *eliminated* versus just diminished?*

"Probe trials of fully trained mice resulted in trajectories and initial hole checking identical to that of regular trials thereby demonstrating that local odor cues are not essential for spatial learning." As far as the reader can tell, probe trials only eliminated the food odor cues but did not eliminate all other odors. If so, this conclusion can be modified accordingly.

We were most worried about odor cues guiding the mice and as now described at great length, we tried to mitigate this problem in many ways. As the reviewer notes, it is not possible to have absolute certainty that there are no odor cues remaining. The most difficult odor to eliminate was the potential odor gradient emanating from the mouse's home cage. However, the 2 vacuum fans per cage were very powerful in first evacuating the cage air (150x in 5 minutes) and then drawing air from the arena, through the cage and out its top for the duration of each trial. We believe that we did at least vastly reduce any odor cues and perhaps completely eliminated them.

The interpretation of direction selectivity is a bit tricky. At different places in this manuscript, this is interpreted as a path integration signal that encodes goal location, including the Consync cells. However, studies show that (e.g. Acharya et al. 2016) direction selectivity in virtual reality is comparable to that during natural mazes, despite large differences in vestibular cues and spatial selectivity. How would one reconcile these observations with path integration interpretation?

Thank you. We had not been serious enough in considering the VR studies and their implications for optic flow as a cue for spatial learning. We now have a section (*Optic flow cues*) in the Discussion that acknowledges the potential role of such cues in spatial learning in our maze.

However, spatial learning in our maze can also occur in the dark. The next small section (*Vestibular and proprioceptive cues*) addresses this point. We cannot be certain about the precise cues used by the mouse to effectively learn to locate food in our maze, but it will take further behavioral and electrophysiological studies to go deeper into these questions.

An extended discussion is found in the sections entitled “What cues are required for learning the food location” and “A fixed start location and self-motion cues are required for spatial learning”. We may have missed some references or ideas regarding VR maze learning with optic flow signals – the Acharya et al reference was an excellent starting point, and we would be grateful for additional pointers that would improve our discussion of this point.

The manuscript would be improved if the speculations about place cells, grid cells, BTSP, etc. were pared down. I could easily imagine the outcome of these speculations to go the other way and some claims are not supported by data. "We note that the cited experiments were done with virtual movement constrained to 1D and in the presence of landmarks. It remains to be shown whether similar results are obtained in our unconstrained 2D maze and with only self-motion cues available." There are many studies that have measured the evolution of place cells in non- virtual mazes, look up papers from the 1990s. Reference 43 reports such results in a 2D virtual maze.

We understand the reviewer’s concerns with the length of the manuscript. However, both the first and third reviewer did find this extensive section useful. We did not add the many papers on the evolution of place fields in real world mazes simply to prevent even greater expansion of the discussion, but relied on the very thorough review of Knierim and Hamilton instead.

Reviewer #3 (Public Review):

Summary:

How is it that animals find learned food locations in their daily life? Do they use landmarks to home in on these learned locations or do they learn a path based on self-motion (turn left, take ten steps forward, turn right, etc.). This study carefully examines this question in a well-designed behavioral apparatus. A key finding is that to support the observed behavior in the hidden food arena, mice appear to not use the distal cues that are present in the environment for performing this task. Removal of such cues did not change the learning rate, for example. In a clever analysis of whether the resulting cognitive map based on self-motion cues could allow a mouse to take a shortcut, it was found that indeed they are. The work nicely shows the evolution of the rodent's learning of the task, and the role of active sensing in the targeted reduction of uncertainty of food location proximal to its expected location.

Strengths:

A convincing demonstration that mice can synthesize a cognitive map for the finding of a static reward using body frame-based cues. This shows that the uncertainty of the final target location is resolved by an active sensing process of probing holes proximal to the expected location. Showing that changing the position of entry into the arena rotates the anticipated location of the reward in a manner consistent with failure to use distal cues.

Thank you.

Weaknesses:

The task is low stakes, and thus the failure to use distal cues at most costs the animal a delay in finding the food; this delay is likely unimportant to the animal. Thus, it is unclear whether this result would generalize to a situation where the animal may be under some time pressure, urgency due to food (or water) restriction, or due to predatory threat. In such cases, the use of distal cues to make locating the reward robust to changing start locations may be more likely to be observed.

We have added “Combining trajectory direction and hole check locations yields a Target Estimation Vector” a section summarizing our main hypotheses and this section includes noting exactly this point + including the reference to the excellent MacIver paper on “robot aggression”.

The main point here follows the Knierim and Hamilton review and assumes that learning “heading direction” and “distance from start to food” require different cues and extraction mechanisms. “Here we follow a review by Knierim and Hamilton (12) suggesting independent mechanisms for extraction of target direction versus target distance information. Averaging across trajectories gave a mean displacement direction, an estimate of the average heading direction as the mouse ran from start to food. The heading direction must be continuously updated as the mice runs towards the food, given that the mean displacement direction remains straight despite the variation across individual trajectories. Heading direction might be extracted from optic flow and/or vestibular system and be encoded by head direction cells. However, the distance from home to food is not encoded by head direction signals.”

And

“We hypothesize that path integration over trajectories is used to estimate the distance from start to food. The stimuli used for integration might include proprioception or acceleration (vestibular) signals as neither depends on visual input. Our conclusion is in accord with a literature survey that concluded that the distance of a target from a start location was based on path integration and separate from the coding of target heading direction (12). Our “in the dark” experiments reveal the minimal stimuli required for spatial learning – an anchoring starting point and directional information based on vestibular and perhaps proprioceptive signals. This view is in accord with recent studies using VR (47, 48). Under more naturalistic conditions, animals have many additional cues available that can be used for flexible control of navigation under time or predation pressure (51).”.

Furthermore, we added panel G do Fig S4, where we show the evolution of the heading angle along the trajectory, plotted as a function of the trials. We see that the mouse only steer towards the target in the last segment of the trajectory, consistent with having the head direction being continuously updated along the path to the food.

Recommendations for the authors:

Reviewing Editor (Recommendations For The Authors):

All three reviewers agreed during the consultation that the context in which distal cues are described in the manuscript would benefit significantly from refinement. The distal cues may be made completely useless from an ethological perspective e.g. if they are seen as "moving" relative to the entrance point (i.e. if the animal were to think it were entering the same location), then the cues would appear as unstable in the random entrance. As such, they may be so unlike natural experiences as to be potentially

confusing to the animal. Moreover, as reported in some of the reviews, the animals may be using the entrances and boundaries as cues to help refine path integration. The results are still very interesting, but more refinement in the text on the interpretation of cues would greatly improve the manuscript. Thus, we recommend that you revise your manuscript to address the reviews.

Thank you. We agree with this recommendation of the reviewers have greatly expanded our discussion on cue stability as already indicated above.

Should you choose to revise your manuscript, please ensure the manuscript include full statistical reporting including exact p-values wherever possible alongside the summary statistics (test statistic and df) and 95% confidence intervals. These should be reported for all key questions and not only when the p-value is less than 0.05.

Done

Lastly, I want to personally apologize for the long delay in editing this manuscript. All three reviews were unfortunately quite delayed, including my own review. I want to thank you for submitting your work to eLife and hope that we can be more efficient in editing your work in the future.

It was a long review process, but we also appreciate that our article was dense and difficult to read. We tried to be comprehensive in our controls and analyses and we appreciate the considerable effort it must have taken to carefully review our paper.

Reviewer #3 (Recommendations For The Authors):

I quite enjoyed this paper and have some suggestions for further improvement.

First, while I appreciate that the format of the journal has Methods at the end, there are some key details that need to be moved forward in the study for proper appreciation of the results. These include:

(1) Location and size of distal cues.

Done

(2) Use of floor washing between mice.

Done

(3) Use of food across the subfloor to provide some masking of the location of the food reward.

Done

(4) A scale bar on one of the early figures showing the apparatus would be beneficial.

Done for Figure 1 where we also provide arena diameter and area.

(5) Motivational state of the mouse with respect to the food reward (in this case, not food restricted, correct?).

Done

Although we are told the trial where learning is defined to have occurred, we were not given the quantitative criterion operationalizing "learning" - please provide (unless I missed it!).

Thank you. This question turned out to be of importance and led to more detailed analyses and related Discussion. We therefore answer in depth.

We now realize that learning the distance to food versus learning the direction to food must be analyzed separately.

On Page 5 second paragraph we provide a definition of “learning distance to food”.

“Fitting the function $dtotal = B \cdot \exp(-\text{Trial}/K)$ reveals the characteristic timescale of learning, K , in trial units (Fig. 2F). We obtained $K = 26 \pm 24$ giving a coefficient of variation (CV) of 0.92. The mean, $K = 26$, is therefore very uncertain and far greater than the actual number of trials. Thus, we hypothesize that the mice did not significantly reduce their distance travelled (Fig. 2A,B,F) because they had not learned the food location – the decrease in latency (Fig. 2D) was due to its increased running speed and familiarity with non-spatial task parameters.”

On Page 7 second paragraph the same analysis gives:

“Now the fitting of the function $dtotal = B \exp(-\text{Trial}/K)$ yielded $K = 5.6 \pm 0.5$ with a CV = 0.08; the mean is therefore a reliable estimate of total distance travelled. We interpret this to indicate that it takes a minimum number of $K = 6$ trials for learning the distance to the target (see also Fig. S4D,E,F,G).

Learning is still not complete because it takes 14 trials before the trajectories become near optimal.”

Learning of distance to food is evident by Trial 6 but is not complete.

On Page 9 third paragraph we give a very precise answer to time taken to learn the direction from start to food. This was already very clear from Fig. 4I but we had missed the significance of this result.

“We compared the deviation between the TEV and the true target vector (that points from start directly to the food hole; Fig. 4I). While the random entrance mice had a persistent deviation between TEV and target of more than 70°, the static entrance mice were able to learn the direction of the target almost perfectly by trial 6 (TEV-target deviation in first trial $\text{mean} \pm \text{SD} = 57.27^\circ \pm 41.61^\circ$; last trial $\text{mean} \pm \text{SD} = 5.16^\circ \pm 0.20^\circ$; $P = 0.0166$). A minimum of 6 trials is sufficient for learning both the direction and distance to food (Fig. 4I) (Fig. 3F) (see Discussion). The kinetics of learning direction to food are clearly different from learning distance to food since the direction to food remains stable after Trial 6 while the distance to food continues to approach the optimal value.”

Learning the direction from start to food is completely learned by Trial 6.

These analyses led to an addition to the Discussion on Page 20 (following the Heading).

“Here we follow a review by Knierim and Hamilton (12) that hypothesized independent mechanisms for extraction of target direction versus target distance information. Our data strongly supports their hypothesis. Target direction is nearly perfectly estimated at trial 6 (Fig. 4I and Results). The deviation of the TEV from the start to food vector is rapidly reduced to its minimal value (5.16°) and with minimal variability ($\text{SD} = 0.20^\circ$). Learning the distance from start to food is also evident at trial 6 but only reaches an asymptotic near optimal value at trial 14 (Fig. 3F). The learning dynamics are therefore very different for target direction versus target distance. As noted below, the food direction is likely estimated from the activity

of head direction cells. The neural mechanisms by which distance from start to food is estimated are not known (but see (49)).”

We believe that this small addition summarizes the complicated answer to the reviewer’s question and is helpful in better connecting the Knierim and Hamilton paper to our data. However, if the reviewers and editors feel that we have gone too far or that this discussion is not clear, we can remove or alter the extra sentences as per any comments.

Reference #49 is to a review paper on spatial learning in weakly electric fish in the dark (<https://doi.org/10.1016/j.conb.2021.07.002>). The review summarizes data on a neural “time stamp” mechanism for estimating distance from start to food. In this review article, we explicitly hypothesized that rodents might utilize such a time stamp mechanism for finding food. We did not include this in the discussion because it was too distracting and would likely confuse readers but put in the reference in case some readers did want to access the “time stamp” hypothesis for spatial learning in the dark.

Second, the discussion was thoughtful and rich. I particularly enjoyed the segment describing the likely computations of the hippocampus. There are a few thoughts I have for the authors to think about that might be useful to potentially add to the discussion:

“The remaining one, mouse 34, went from B to the start location and then, to A.”

This out-and-back pattern has been seen in the literature, such as multiple papers by Golani (here’s one: <https://www.pnas.org/doi/full/10.1073/pnas.0812513106>). Would the authors speculate, given their suggested algorithm, what the significance of out and back may be? Is there something about the cell’s encoding of direction and distance that requires a return to the start location, and would this be different if representation is based on self-motion versus based on distal cues in an allocentric representation?

We do discuss this for pretraining trials but have no idea what this mouse is doing in this case.

In a low-stakes task environment, for an animal that has a low acuity visual system, where the penalty for not using distal cues is at most some additional (likely enriching in itself to these mice who live a fairly unenriched life in small cages) search/learning/exploration time, perhaps it is not so surprising that body-frame cues are used. Considering the ethology of the animal, if it had multiple exits of an underground burrow, it might need to use distal cues to avoid confusion. The scenario you provide to the animal is essentially a deceptive one where it has no way of telling it is coming out to the arena from a different burrow hole, modulo some small landmarks on an otherwise uniform cylinder of space. This might be asking too much of an animal where the space it would enter normally would not be a uniform cylinder.

What happens with a higher-stakes case? This is clearly a different study, but you may find some recent work with a mobile predatory robot of interest (<https://www.sciencedirect.com/science/article/pii/S2211124723016820>). Visual cues are crucial in the avoidance of threats in this case. Re-routing, as shown by multiple videos of that study, is after a brief pause, and seemingly takes into account the likely future position of the threat.

Done. A fascinating paper that illustrates the unexpected “high level” behavior a rodent is capable of when placed in more naturalistic situations. I think our “two food location” experiments are along the same direction – unexpected rich behavior when the mouse are challenged.

Connected to the low-stakes vs high-stakes point, it might be nice for the paper to discuss situations in which cognitive-map-based spatial problem solutions make sense versus not.

Here is an example of such a discussion, around page 496:

https://www.dropbox.com/scl/fi/ayoo5w4jgnkblgfu7mpad/MacI09a_situated_cog.pdf?rlkey=2qhh89ii7jbkavt6ivevarvdk&dl=0.

Right a very relevant discussion by MacIver. However, when I tried to write it in it took nearly half a page of dense writing to connect to the themes of our article. I figured that the already long discussion will try the patience of most readers and so decided to not include this extra discussion.

Minor points/ queries

Why the increase in sample density at about the 1/4 radius of arena distance? Static, trial 14, Figure 3I, shown also maybe Figure 4 H.

We were also puzzled when this occurred but have no explanation. And there are, in our figures, many other examples of the mice hole checking near their exit site. See next answer.

Why was the hole proximal to start so often probed in 7B?

We were also puzzled when this occurred but have no explanation.

Check Video 1 to exactly see this behavior. The mouse exits its home and immediately checks a nearby hole. It proceeds to Site B (empty) and then Site A (empty) with many hole checks along the way. After leaving Site A, the mouse proceeds to the wall located far from an entrance and does another hole check. The near the wall holes that are checked are in no way remarkable: a) they have never contained food; b) they are rotated between trials, and we wash the floor carefully, so they do not “smell” any particular hole; c) the food on the lower level floor is in no way “clumped” under that hole, etc.

We have discussed this phenomenon quite a lot and LM was able to come up with only one hypothesis for this behavior. In analogy to the electric fish work (responses of diencephalic neurons to “leaving or encountering a landmark”), the “near the entrance” hole check might be an active sensing probe to “time stamp” the exit from home while finding food would “time stamp” the end of a successful trajectory. Path integration between time stamps would then provide the estimate for time/distance from start to food – exactly our hypothesis for weakly electric fish spatial learning in the dark. This hypothesis is exceedingly speculative and so we do not want to include it.

Normally I would cite a line number. Since I do not see line numbers, I will leave it to you to do a search:

"A than the expected by chance" -> "than expected"

Done. I apologize for the lack of line numbers. I have, so far, been unable to get Word to confine line numbers to selected text and not run over onto the Figure Legends. I have put in page numbers and hope this helps.

RW, VR, MWM, etc - please expand the acronym on first use.

Done

It might be interesting to see differences in demand/reliance on active sensing in the individuals who learn the task less well than the animals who learn the task well. If the point is to expunge uncertainty, then does the need for such expunging increase with the poverty of internal representation resolution / fewer decimal places on the internal TEV calculation?

We do have variation in the mice learning time but the numbers are not sufficient for this interesting extension. This is just one of many follow up studies we hope to carry out.

<https://doi.org/10.7554/eLife.95764.2.sa0>