

UC Merced

Proceedings of the Annual Meeting of the Cognitive Science Society

Title

'Decoding' the locus of spatial representation from simple localization errors

Permalink

<https://escholarship.org/uc/item/185176jj>

Journal

Proceedings of the Annual Meeting of the Cognitive Science Society, 43(43)

Authors

Yousif, Sami R

Keil, Frank

Publication Date

2021

Copyright Information

This work is made available under the terms of a Creative Commons Attribution License, available at <https://creativecommons.org/licenses/by/4.0/>

Peer reviewed

‘Decoding’ the locus of spatial representation from simple localization errors

Sami R. Yousif (sami.yousif@yale.edu)

Department of Psychology, Yale University, New Haven, CT, 06520 USA

Frank C. Keil (frank.keil@yale.edu)

Department of Psychology, Yale University, New Haven, CT, 06520 USA

Abstract

Representing a location in space requires two things: an anchor point, and a code (or coordinate system) to define other locations relative to that anchor point. Recent work has shed light on the latter, providing evidence that the default ‘format’ of visuospatial representation may be polar coordinates (i.e., angle/distance relations). Yet the former remains a topic of debate. For example, a classic distinction in the realm of spatial navigation research pits representation relative to landmarks against representation relative to boundaries. Here, we exploit the polar format of spatial representations to propose a new method for assessing the locus of spatial representation. Specifically, we show that from simple localization errors we can infer the anchor point from which observers localized a target point. We highlight a few basic demonstrations of this method and discuss possible applications for further research on spatial representation.

Keywords: Space; spatial representation; format; polar; landmark

Introduction

Consider your favorite restaurant. Where is it located? The answer is not obvious: You cannot simply define its location in absolute terms. To answer this question, you must first refer to some other location in space, whether that be a familiar landmark or a nearby street (or perhaps some global coordinates, which are themselves defined relative to other points in space). The simple demonstrations highlight a key aspect of spatial representation: that all locations in space must be represented relative to some other location.

There is debate regarding the nature of spatial representation. What coordinate system, if any, supports spatial representation (Yang & Flombaum, 2018; Yousif & Keil, 2021)? What role do landmarks vs. boundaries play in spatial representation (Bullens et al., 2010; Doeller & Burgess, 2008; Doeller et al., 2008)? Is the ‘cognitive map’ Euclidean or network-like (Kuipers et al., 2003; O’Keefe & Nadel, 1978; Tolman, 1948; Warren et al., 2017; Werner et al., 2000)? Naturally, the answers to some of these questions may be related: if the cognitive map is network-like, it may be easier to imagine it being supported by polar coordinates (i.e., angle distance relations) than Euclidean coordinates. And perhaps a network-like cognitive map is more likely to depend on landmarks than boundaries.

Here, we exploit one aspect of spatial representation to explore another. Recent work analyzed errors made in a simple localization task to infer the coordinate system observers use place an object (Yousif & Keil, 2021). In brief, this method depends on an inference based on correlations of errors across

trials. The method assumes that if space is represented with a given coordinate system, errors in each dimension of that coordinate system should be random, or uncorrelated (if they are in fact represented as independent dimensions). In this way, the presence of error correlations indicates that a given coordinate system is an unlikely candidate for representation. Conversely, the absence of error correlations indicates that a given coordinate system is a good candidate for representation. Across several studies, observers defaulted to a polar coordinate system (although they used other coordinate systems flexibly, depending on the space; Yousif & Keil, 2021; see also Huttenlocher et al., 1991; Yousif et al., 2020).

This method may uncover more than merely *how* observers represent space; it may also reveal the locus of that representation. In prior work (Yousif & Keil, 2021), all analyses are conducted relative to other objects present in the display. But suppose the presence of other objects in the display was unknown. Would it be possible to infer the locations of the landmarks? In theory, the answer is yes: if the absence of correlation is evidence of representation, then one could ‘search’ the space for the point with the lowest error correlations. In other words, it should be possible to do a ‘searchlight’ analysis (akin to searchlight analyses in fMRI; Etzel et al., 2013) to *find* the regions of low error correlation — and from those low correlations infer the locus of representation.

Current Study

The present paper explores whether this ‘searchlight’ approach can reveal the locus of spatial representation. In three initial demonstrations, observers localized points in a space comprised of only a single landmark. Different sets of observers were shown different landmarks. We show that we can reliably infer where the landmark was located — based solely on localization errors observers made throughout the task. In one additional experiment, we explore how this method may apply in more complex spatial environments, i.e., ones with more than one landmark. Finally, we discuss how this method may be applied to other aspects of spatial cognition.

Experiment 1: ‘Decoding’ landmarks

Here, four separate groups of observers completed a visual matching task, modeled after that used by Yousif & Keil (2021). In the simplest version of this task, observers saw a grey dot paired with a separate blue dot in one corner of the screen. In the opposite corner of the screen, observers saw another grey dot, but no blue dot. Observers were instructed to place a new blue dot near the other grey dot, such that the relative position of the new

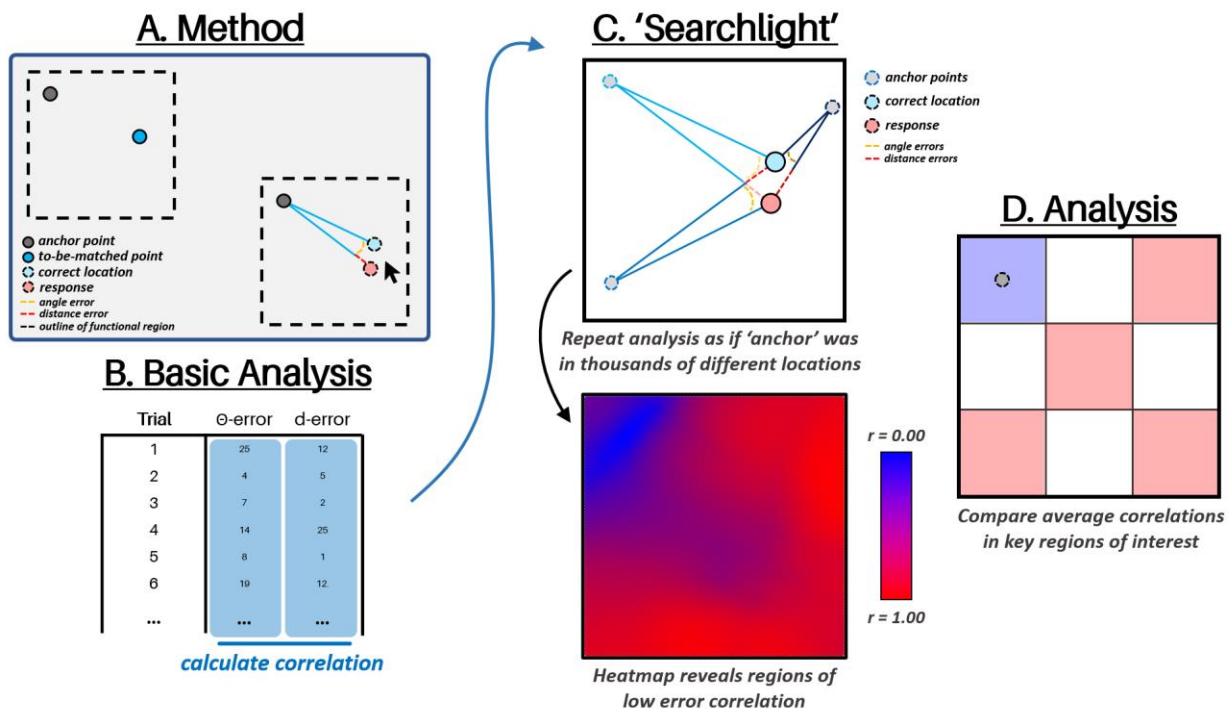


Figure 1. (A) A caricature of the method. In this example, observers must place a point to match the initially present blue point in the top left. The dotted grey lines represent the ‘functional’ region in which points can appear. (B) A representation of the error correlation analysis. Here, we are analyzing correlations between errors in the two dimensions of polar coordinates. In the basic analysis, each subject therefore has a unique correlation value that is used for subsequent analyses. (C) A representation of the ‘searchlight’ analysis. Here, error correlations are calculated relative to thousands of points in the space, as if each was the anchor point relative to which observers made their localization judgments. The resulting heatmap shows the regions of lowest correlation (in blue) indicating the likely presence of a landmark in that location. (D) A depiction of the primary analyses.

blue dot relative to the second grey dot matched the relative position of the already-visible blue dot to its grey dot counterpart. See Figure 1 for an example display. The goal of the present study is to determine whether we can use a ‘searchlight’ error correlation analysis to infer the location of the landmark (grey dot) based solely on an analysis of observers’ localization errors.

Participants Twenty observers were recruited via Amazon Mechanical Turk completed each task in exchange for monetary compensation (80 observers total).

Design Observers completed 48 trials of the visual matching task, in which dots appeared in random locations within a 400 pixel by 400 pixel region. In one experiment, the visible anchor point (grey dot) was presented in the bottom left of the functional region (as shown in Figure 1); in another, it was presented in the center; in another, it was presented in the top right; and in a final experiment, two separate anchor points were presented in opposite corners.

Procedure On each trial, observers simply had to place the missing shapes to match the relative location of the corresponding shape, by moving and then clicking with the mouse (see Figure 1a). The missing shape appeared upon mouse-click, at which point observers could click additional times or drag and drop the dot to change its location. Once observers were

satisfied with the missing object’s location, they pressed a key to submit their response. If a response was recorded, then the display was replaced with a blank screen for 0.5s, after which the next trial began. If no response was recorded within 7s, then the next trial automatically began (after .5s), and that trial was randomly shuffled back into the trial sequence. Observers were repeatedly instructed to ensure that they responded within the time limit.

Each observer completed 48 trials. The location of each target object on each trial was randomized within and across observers. Observers completed two representative practice trials (the data from which were not recorded) before beginning the task.

Analysis The key analysis is a correlation between errors in various spatial dimensions. Unlike Yousif & Keil (2021), only correlations among polar errors are analyzed. Further, these analyses are conducted relative to the entire visible space, rather than a single location (i.e., the same analysis will be conducted iteratively at 6561 points throughout the space).

First, the original polar coordinates (i.e., the true polar coordinates of the initially visible blue dot relative to its counterpart grey dot) are calculated relative to a fixed point in the 400 x 400 space. Then, the polar coordinates of the new point (i.e., the point placed by the participant) are calculated relative to the same fixed point. The error in each dimension is calculated as

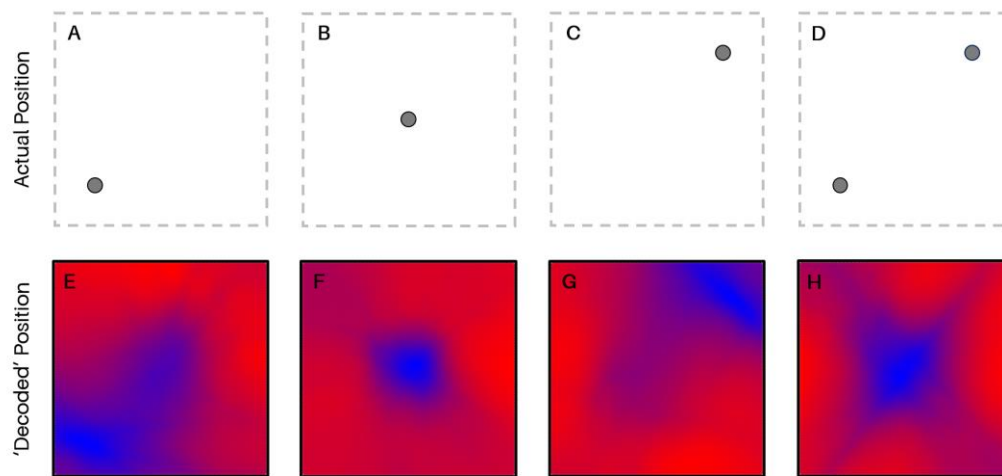


Figure 2. A depiction of the location of the anchor points in each experiment for Experiment 1 (A, B, C, and D). The dotted grey lines correspond to the ‘functional’ region of space in which points could appear; these lines were not visible to observers. The resulting heatmaps from the ‘searchlight’ error correlation in each experiment are also displayed (D, E, F, and G). Regions of lower correlation appear in blue; regions of higher correlation appear in red.

the absolute difference between the original value and the new value in each dimension. Therefore, for each trial, there is a measure of error in the angle dimension and in the distance dimension. Then, at the subject level, the correlation between the errors in these two dimensions is calculated (across the 48 trials). This is done iteratively for each observer. Then, an average correlation value is calculated by averaging the correlation values across observers.

This analysis (as explained in the previous paragraph) is conducted relative to 6561 different anchor points throughout the functional space. (There are anchor points every 5 pixels, starting at the point [-200, -200] and reaching to the points [200, 200], functionally resulting in 81 rows and 81 columns of anchor points.) Therefore, a unique average correlation value is calculated for each of those 6561 anchor points. These values will then be summarized via heatmaps. To quantify impressions gleaned from the heatmaps, we will subdivide the space into nine equivalent regions (see Figure 1D) and compare *average* correlations across regions. For example, if the anchor point was present in the top left corner, average correlations for the top left section should be *lower* than the other corner sections. The key statistical analysis will be t-tests between regions, not between individual points.

Note that we analyzed only error correlations in polar dimensions, not Cartesian dimensions (as in Yousif & Keil, 2021). There are two primary reasons for this. First, prior work has clearly demonstrated the use of polar coordinates in simple environments like these (Yousif & Keil, 2021). Despite this, if observers were *not* relying on polar coordinates, then the heatmaps should reveal no systematicity whatsoever. This means that, in practice, the studies here serve as a replication of that prior work; the method here only succeeds if observers are using polar coordinates.

Second, it is impossible to create heatmaps for error-correlations in Cartesian dimensions in the same way. Regardless

of where we assume the Cartesian anchor point is, error-correlations relative to *any* point in space will be the same. Whether Cartesian errors are highly correlated or not, the heatmaps would necessarily be uniform. (The equivalent analysis using Cartesian coordinates would involve rotation, not translation.)

Results and Discussion

Heatmaps for the four experiments are presented in Figure 2. As is evident from the figure, the decoding analysis was capable of identifying the location of the anchor point.

For Experiment 1a, in which the anchor point was in the bottom-left corner, we compare the average correlation for the bottom-left section to the other three corner sections as well as the center section. Each region contained 27x27 separate points relative to which errors were calculated, resulting in 729 total points. T-tests were conducted comparing the 729 points of the target region to the 729 points of the non-target regions. There were indeed significantly lower correlations in the target corner (bottom-left) compared to the other three locations, $t_{s>50}$, $p_{s<.001}$, $d_{s>1.80}$.

For Experiment 1b, in which the anchor point was in the center section, we compare the average correlation for the center section to the four corner sections. There were significantly lower correlations in the relevant section, $t_{s>54}$, $p_{s<.001}$, $d_{s>2.00}$.

For Experiment 1c, in which the anchor point was in the top right section, we compare the average correlation for the top right section with the three other corner sections as well as the center section. Once again, there were significantly lower correlations in the relevant section, $t_{s>78}$, $p_{s<.001}$, $d_{s>2.90}$.

For Experiment 1d, in which there were two separate anchor points in opposite corners of the screen, we tested one key comparison: the diagonal on which the anchor points were situated vs. the opposite diagonal. In other words, we average the correlations for the top-right and bottom-left corners and

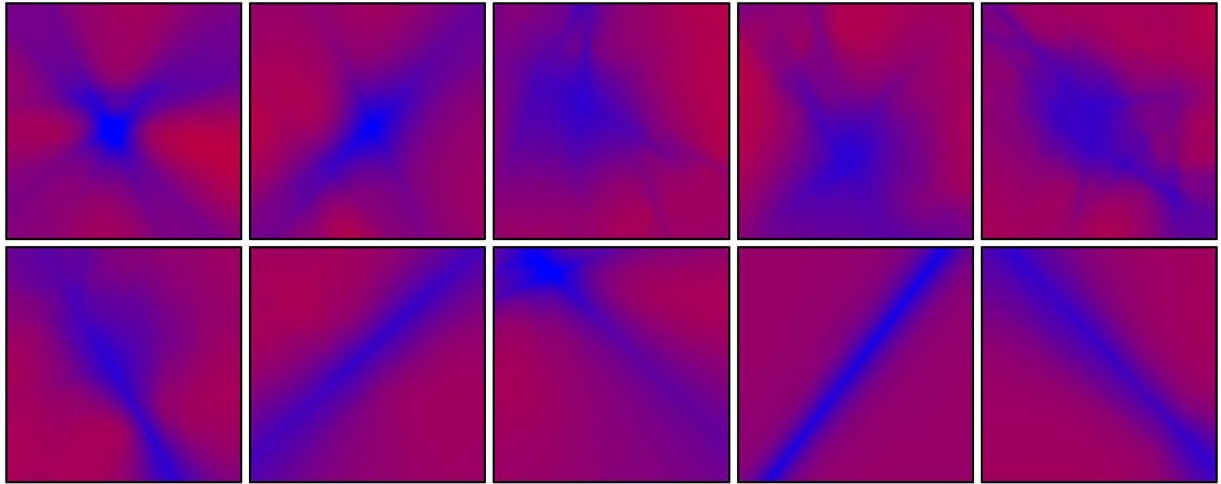


Figure 3. Representative heatmaps from individual observers in Experiment 1d. Regions of lower correlation appear in blue; regions of higher correlation appear in red.

compared that to the average correlations for the top-left and bottom-right corners. (Note that the middle section is excluded from this analysis, as it would be redundant across the comparisons.) Surprisingly, the diagonal along which the anchor points sat had higher error correlations than the opposite diagonal, $t(1457)=28.35$, $p<.001$, $d=.74$. However, note that this effect, although significant, is significantly weaker than those in the previous experiment; note also that this pattern is not obvious from the heatmap itself. To understand this pattern, we created separate heatmaps for each observer (following the same procedure as before, but without averaging across observers). A representative set of these heatmaps are displayed in Figure 3.

As seen in Figure 3, despite the ambiguity in the average heatmaps, the individual observer heatmaps are systematic — although sometimes in opposing directions. Some observers had the lowest correlations on the diagonal with the anchor points, whereas other observers had the lowest correlations on the opposite diagonal. We refrain from overinterpreting these impressions, although we think these results warrant additional investigation given the systematic bimodality of responses.

These experiments serve as a proof of concept: when observers are localizing points relative to a stable location in space, we can ‘decode’ the location of that anchor point using nothing but error correlations. Although results from Experiment 1d were mixed, it is clear that, on an individual level, meaningful systematicity is revealed from this method.

General Discussion

Here, we have proposed a novel ‘searchlight’ method for determining the locus of spatial representation. In three experiments, we showed how this approach can be used to identify landmarks in an environment. In a fourth experiment, we tested this method in a more complex spatial environment — one with multiple landmarks. The results of this latter experiment were mixed, although there were still signs of systematic, interpretable patterns of data. Below, we discuss some key strengths and limitations of this approach and highlight some ongoing research programs that could use a similar approach.

Strengths of the searchlight approach

First, the task employed here is exceedingly simple. The key measure here depends only on localization errors and no other explicit task. Unlike other spatial tasks that rely on relatively explicit measures (e.g., pointing behavior), this method could be framed as either a visual memory or a visual perception task; observers need not view this as a spatial task *per se*.

Second, this method can be straightforwardly adapted to work in both 2D and 3D spatial environments. Although the environments we test here are simple 2D environments on a computer screen, this method could be readily applied to a localization task in both virtual and real-world 3D environments.

Third, the method yields robust, easily appreciable results (at least in simple environments). From the heatmaps alone, it is obvious where the landmarks are located.

Fourth, because this approach is maximally data-driven, it has the potential to reveal counterintuitive, unexpected pattern of results. (Of course, it is hard to predict what such results may be — otherwise we wouldn’t call them ‘counterintuitive’!)

Fifth, this method has the potential to reveal observer-specific variations that may go unnoticed in other spatial tasks (which primarily rely on average behavior across many observers). Although the heatmaps in Experiment 1d produced no conclusive pattern, it is still clear from the heatmaps of individual observers that there are meaningful, systematic patterns for several observers.

Limitations of the searchlight approach

First, this method depends on the use of polar coordinates; the ‘searchlight’ is specifically searching for regions of low polar error correlations. Of course, observers may in some cases rely on other forms of spatial representations (Yousif & Keil, 2021). In such cases, this exact analysis would yield no meaningful results.

Second, this method depends on a correlation analysis that is sensitive. Although results are clear when we average across many trials and many observers, virtually any noise in the data

(i.e., a subset of trials in which observers responded randomly) may yield uninterpretable results. Certain exclusion criteria may eliminate this problem, but for now we have not yet established a clear ‘pre-processing’ plan for determining how data should be filtered.

Third, we view the current version of this method as a prototype. This method could be improved in ways that would allow for its application to numerous other contexts. For example, if better criteria are established for filtering noisy data, then cleaner heatmaps may be obtainable, especially at the level of single observers.

Possible applications

This method could be used address numerous open questions in the fields of spatial cognition and navigation. It can be easily applied both to 2D and 3D spatial tasks, meaning that the range of applications is wide. For instance, perhaps this method could be used to address the use of landmarks vs. boundaries more directly (e.g., Doeller & Burgess, 2008; Doeller et al., 2008). It could also be used evaluate other aspects of spatial representation, e.g., whether spatial representation is supported by Cartesian vs. polar coordinates, or whether spatial representation operates in a metric format or a network-like format (see, e.g., Yousif & Keil, 2021).

This method could also be applied to research questions beyond the domain of spatial cognition. The error correlation analysis presented here could, in theory, be applied to any dimensions that might be represented integrally vs. separably in the mind (Garner & Felfoldy, 1970; for review see Algom & Fitousi, 2016). That said, this approach is flexible enough that we imagine there may be many yet-unforeseen applications.

Conclusion

We have presented a novel method for assessing representational content, which we have used to demonstrate that we can ‘decode’ the locus of spatial representation from mere localization errors. We have also provided some preliminary evidence that this method may be useful in addressing other questions of spatial representation, and we have speculated about how this method could be applied to other research questions (including non-spatial domains). Future work can build on this approach in several ways.

References

Algom, D., & Fitousi, D. (2016). Half a century of research on Garner interference and the separability–integrality distinction. *Psychological Bulletin*, 142, 1352–1383.

Bullens, J., Nardini, M., Doeller, C. F., Braddick, O., Postma, A., & Burgess, N. (2010). The role of landmarks and boundaries in the development of spatial memory. *Developmental Science*, 13, 170–180.

Doeller, C. F., King, J. A., & Burgess, N. (2008). Parallel striatal and hippocampal systems for landmarks and boundaries in spatial memory. *Proceedings of the National Academy of Sciences*, 105, 5915–5920.

Doeller, C. F., & Burgess, N. (2008). Distinct error-correcting and incidental learning of location relative to landmarks and boundaries. *Proceedings of the National Academy of Sciences*, 105, 5909–5914.

Etzel, J. A., Zacks, J. M., & Braver, T. S. (2013). Searchlight analysis: promise, pitfalls, and potential. *Neuroimage*, 78, 261–269.

Gallistel, C. R. (1990). *The organization of learning*. Cambridge, MA: MIT press.

Garner, W. R., & Felfoldy, G. L. (1970). Integrality of stimulus dimensions in various types of information processing. *Cognitive Psychology*, 1, 225–241.

Huttenlocher, J., Hedges, L. V., & Duncan, S. (1991). Categories and particulars: Prototype effects in estimating spatial location. *Psychological Review*, 98, 352–376.

Kuipers, B., Tecuci, D. G., & Stankiewicz, B. J. (2003). The skeleton in the cognitive map: A computational and empirical exploration. *Environment and Behavior*, 35, 81–106.

McNamara, T. P. (1986). Mental representations of spatial relations. *Cognitive Psychology*, 18, 87–121.

Müller, M., & Wehner, R. (1988). Path integration in desert ants, *Cataglyphis fortis*. *Proceedings of the National Academy of Sciences*, 85, 5287–5290.

O’Keefe, J. and Nadel, L. (1978). *The hippocampus as a cognitive map*. Oxford, England: Clarendon Press.

Tolman, E. C. (1948). Cognitive maps in rats and men. *Psychological Review*, 55, 189.

Warren, W. H., Rothman, D. B., Schnapp, B. H., & Ericson, J. D. (2017). Wormholes in virtual space: From cognitive maps to cognitive graphs. *Cognition*, 166, 152–163.

Werner, S., Krieg-Brückner, B., & Herrmann, T. (2000). Modelling navigational knowledge by route graphs. In *Spatial Cognition II* (pp. 295–316). Springer, Berlin, Heidelberg.

Yang, F., & Flombaum, J. (2018). Polar coordinates as the format of spatial representation in visual perception [Abstract]. *Journal of Vision*, 18, 21.

Yousif, S. R., Chen, Y. C., & Scholl, B. J. (2020). Systematic angular biases in the representation of visual space. *Attention, Perception, & Psychophysics*, 82, 3124–3143.

Yousif, S. R., & Keil, F. C. (2021). The Shape of Space: Evidence for Spontaneous but Flexible Use of Polar Coordinates in Visuospatial Representations. *Psychological Science*, 32, 573–586.