

**EXPLICIT ATTENTION TO ALLOCENTRIC VISUAL LANDMARKS IMPROVES
MEMORY-GUIDED REACHING**

Lina Musa

A THESIS SUBMITTED TO THE FACULTY OF GRADUATE STUDIES
IN PARTIAL FULFILLMENT OF THE REQUIREMENTS FOR THE
DEGREE OF MASTER OF ARTS

GRADUATE PROGRAM IN PSYCHOLOGY
YORK UNIVERSITY
TORONTO, ONTARIO

July, 2021

© Lina Musa, 2021

ABSTRACT

This thesis explores the influence of explicit (instruction-dependent) allocentric landmarks on reaching performance. I investigated this effect by testing participants' reaching performance in two instruction conditions (egocentric vs. allocentric), but with similar sensory and motor requirements. In both conditions, participants fixated their gaze near the centre of a display aligned with the right shoulder, and an LED target briefly appeared (alongside a visual landmark) in one visual field. After a mask/memory delay period, the landmark re-appeared in the same or opposite visual field. In the allocentric condition, participants were instructed remember the initial location of the target relative to the landmark, and to reach relative to the shifted landmark. In the egocentric condition, participants were instructed to ignore the landmark and point toward the remembered location of the target. The motor execution was equalised by having participants point relative to a landmark shifted to the opposite visual field on 50% of trials in the allocentric task and anti-point on 50% of trials in the egocentric task. When the landmark stayed within the same visual field, the allocentric instruction yielded significantly more accurate and precise pointing than the egocentric instruction, despite identical visual and motor conditions. Likewise, when the landmark shifted to the opposite side, pointing was significantly better following the allocentric instruction (compared to motor-matched anti-reaches). These results show that in the presence of a visual landmark, memory-guided pointing improves when participants are explicitly instructed to point relative to the landmark. This suggests that explicit attention to a visual landmark better recruits allocentric coding mechanisms, that can augment implicit egocentric visuomotor transformations.

DEDICATION

*This thesis is dedicated to the strongest woman I know.
Your fight for your life has been my motivation and my courage.
Your hope in the darkest of times has allowed me to see the hope within myself.
To my loving mother.*

ACKNOWLEDGEMENTS

This year has been a year of reflection on a global scale. It has shown us that everything our life revolved around is temporary, that hope is important because it makes the present moment less difficult to bear, but most importantly that our priority will always be to keep our friends and family safe. When taking time to reflect on the years I have been in academia, my first thought is all the people that made a difference in my life. I am sincerely grateful and immensely fortunate to have met all the right teachers and mentors at the right moment. First and foremost, I would like to thank my supervisor, Dr. J. Douglas Crawford. Thank you for the giving me the opportunity to pursue my research interests. Your work ethic and dedication to your research has been a great inspiration. Thank you for your patience, kindness, support, wisdom and encouragement. Above all, thank you for taking me under your wing despite my personal challenges. This year had some highs and very low lows for me, but your compassion throughout this period made it possible for me to pursue my goals.

Specials thanks to my colleagues, the Crawford lab members, whom I am pleased to now call my friends. Thanks to Gaelle and Sharimini who have shown me that a work-life is balance is possible to achieve if your lab mates can be your concert-going mates, to George, Haider and Amir for sharing great conversations and my love for caffeine, to Bianca for being a career and fitness inspiration, and to Vishal and Veronica for their valuable advice. I could not have asked for a better lab environment to conduct research in.

I would also like to thank Dr. Xiaogang Yan for his technical and theoretical support and Saihong Sun for her programming support. Thank you for facilitating my experiments and sharing your invaluable expertise with me. I have learned a lot from you and tremendously enjoyed our conversations within and outside of the lab setting. Thanks to Dr. Mark Adkins for his guidance with my statistical analysis. In addition, thank you to Dr. Joe DeSouza and Dr. Mazyar Fallah for giving me my first opportunities to work in a research setting and inspiring me to pursue research at a graduate level. I feel privileged to count myself among those who have conducted research at the Centre for Vision Research. I would like to thank the members of my examining committee, Dr. Jennifer Steeves, Dr. Erez Freud and Dr. Denise Henriques.

Last but not least, thank you to all the individuals who participated in my experiments and to the reader of this manuscript for your interest in my research.

TABLE OF CONTENTS

ABSTRACT	ii
DEDICATION	iii
ACKNOWLEDGEMENTS	iv
TABLE OF CONTENTS	v
LIST OF FIGURES AND TABLES	vii
CHAPTER 1: GENERAL INTRODUCTION	1
1.1 Introduction	2
1.2 Egocentric vs Allocentric Coding of Space.....	4
1.3 Egocentric Mechanisms for Reach	7
1.4 Allocentric Representations in Reach.....	9
1.5 Neural Mechanisms of Allocentric and Egocentric Spatial Representations.....	13
1.6 Aims of the Current Study	18
 CHAPTER 2: INSTRUCTION ALTERS THE INFLUENCE OF ALLOCENTRIC LANDMARKS IN A REACH TASK.....	 20
2.1 Abstract	21
2.2 Introduction	22
2.3 Material and Methods.....	25
2.3.1 Participants	25
2.3.2 Apparatus	25
2.3.3 Task and Stimuli	28
2.3.4 Experimental Design	29
2.3.4a Egocentric Instructions Condition	30
2.3.4b Allocentric Instructions Condition	31
2.3.5 Sample Size Analysis	33
2.3.6 Data Analysis	33
2.4 Results	37
2.4.1 Reaching Variance	37
2.4.2 Reaching Accuracy	41

2.4.3 Reaction Time	49
2.5 Discussion	52
2.5.1 Effect of Landmark Instruction on Reaching Precision and Accuracy.....	53
2.5.2 Comparison to Previous Studies Utilizing Allocentric Landmarks.....	57
2.5.3 Possible Physiological Mechanisms	59
2.5.4 Conclusion	62
 CHAPTER 3: GENERAL DISCUSSION	 63
3.1 Summary	64
3.2 Contributions to Spatial Representation Literature	64
3.3 Limitations	66
3.4 Unresolved Questions and Future Directions	67
3.5 Applications of Cuing to Individuals with Neurological Impairment	68
3.6 Conclusion	71
REFERENCES	72
APPENDICES	83
Appendix A: Supplementary Figure 1	83
Appendix B: Supplementary Figure 2	84
Appendix C: Supplementary Figure 3	85
Appendix D: Author Contributions	86

LIST OF FIGURES AND TABLES

Figure 1. <i>Spatial Encoding Strategies</i>	5
Figure 2. <i>Major Projections to the Two Visual Streams of Processing</i>	15
Figure 3. <i>Experimental Set-up</i>	27
Figure 4. <i>Experimental Stimuli and Paradigm</i>	32
Figure 5. <i>Typical Eye and Finger Trajectories</i>	36
Figure 6. <i>Reach Endpoint Ellipses</i>	39
Figure 7. <i>Reach Endpoint Ellipse Areas</i>	40
Figure 8. <i>Regression Plots of Pro-Egocentric Task</i>	43
Figure 9. <i>Regression Plots of Same-Allocentric Task</i>	44
Figure 8. <i>Regression Plots of Anti-Egocentric Task</i>	45
Figure 9. <i>Regression Plots of Opposite-Allocentric Task</i>	46
Figure 12. <i>Regression Plots of Allocentric vs Egocentric Reach Endpoints</i>	47
Figure 13. <i>Constant Reaching Error</i>	48
Figure 14. <i>Reaction times</i>	50
Table 1. <i>Summary of Ex-Gaussian Model</i>	51

CHAPTER 1
GENERAL INTRODUCTION

1.1 Introduction

Real life scenarios perpetually involve interactions with objects in the surrounding environment. Whether it be for performing our daily tasks, such as reaching and grasping a cup of coffee, or the skilled use of an instrument, such as playing a piano, we inexorably rely on brain mechanisms for goal-directed action. This is accomplished through a complicated process, whereby sensory information from multiple modalities, including visual information, must be effectively encoded and then transformed into motor action (Crawford et al., 2004; Shadmehr & Wise, 2005). Research has shown that when pointing or reaching to a visual target, motor behaviour must be generated based on the exact location of the target (Darling et al., 2007; Desmurget et al., 1998). However, the visual environment is not always stable, interruptions to the visual system are a frequent challenge. Targets for goal-directed action can become occluded, their position might change, their position relative to oneself might change due to self motion and they may no longer be visible in a dark environment. Additionally, gaze might become needed for another aspect, or an action might have to be accomplished too quickly to keep track of target position. To illustrate these scenarios, imagine you are in a crowded concert, and you want to retrieve your phone from your roommate who politely offered to hold onto it while you finish drinking your beverage. The location of your phone on their hand can become occluded by someone's backpack in the row in front of you, your roommate's position might have shifted slightly by the pushing crowd, similarly your position could have been shifted by the excited crowd, the performer might have changed the ambiance of the next song to darker environment that makes it difficult for you to see your phone, you might be too interested in the performance so you direct your gaze to the performer while retrieving your phone and finally you might have to retrieve your phone very quickly before the crowd pushes you. To bridge these discontinuities,

target location is preserved through an internal spatial representation that retains information about the unseen target, which is utilized by feedforward movement plans to accomplish goal-directed actions.

Spatial reference frames can be of two main forms: egocentric, based on body-centred coordinates or allocentric, a map-like representation of objects relative to a stable visual landmark in the surrounding environment (Howard & Templeton, 1966; Vogeley & Fink, 2003). For example, if want to describe your location to your roommate in the crowded concert, you can either describe it an egocentric frame of reference “I’m straight ahead, you’re looking right at me” or in allocentric frame of reference “I’m next to the stage”. Allocentric reference frames allow for more flexibility in the unseen target position as compared to egocentric reference frames. In an egocentric reference frame, the momentarily unseen target must remain in the same position relative to the observer for a movement to be reliably executed towards it. Conversely, in an allocentric reference frame, if the target maintains the same relative relationship to a stable landmark and its position changes relative to the observer or the observer moves, the target location will still be known relative to the landmark and can in theory still be reached to accurately and reliability. As an example, when typing on a computer keyboard, if one knows the relative positions of the “1” and “9” keyboard keys on their laptop, they can locate the position of “9” number key and finish typing the word “COVID-19”, even when their device has been shifted to a new location so their roommate can have a look at the latest news with them. Thus, allocentric cues allow for actions to be executed in a satisfactory manner, even when one’s position relative to the target changes. Despite the flexibility in target position afforded by an allocentric encoding strategy, an important aspect of allocentric encoding that is evident in this example is that allocentric information might need to be deliberately encoded. To utilize the “1”

keyboard key to remember the position of the “9”, one might need to choose to do so due.

Although the task utility of an allocentric strategy is imminent, it is conceivable that in such a scenario where the target is embedded in a crowded display and the landmark is neither proximal nor salient, one might need to effortfully remember a target position relative to the pre-selected landmark, as opposed to relying on implicit mechanisms of allocentric encoding. The question then arises, is a deliberate allocentric encoding strategy worth the effort? Would remembering the position of “9” key only be a better strategy of accurately and reliably locating it or would one benefit from to rules or instructions to remember and utilize the relative position of the “1” key to press the “9”?

1.2 Egocentric vs. Allocentric Coding of Space

Projections of the external world to the retina are inherently egocentric and consequently spatial processing is at least initially egocentric. Early visual cortical areas contain a topographic map of the retina that facilitates target encoding relative to the fovea. It has been shown that visual targets can become encoded in an egocentric reference frame, relative to the eyes, head or hand (Crawford et al., 2004; Henriques et al., 1998; Lemay & Stelmach, 2005). Many studies have found that healthy participants can reach and point with reasonable accuracy to remembered targets, based solely on egocentric cues (Blohm & Crawford, 2007; Crawford et al., 2004). Nonetheless, it has been repeatedly established that human have a normal capacity to use visual landmarks to store spatial information (Carrozzo et al., 2002; Krigolson & Heath, 2004; Lemay et al. 2004; Gentilucci et al., 1996).

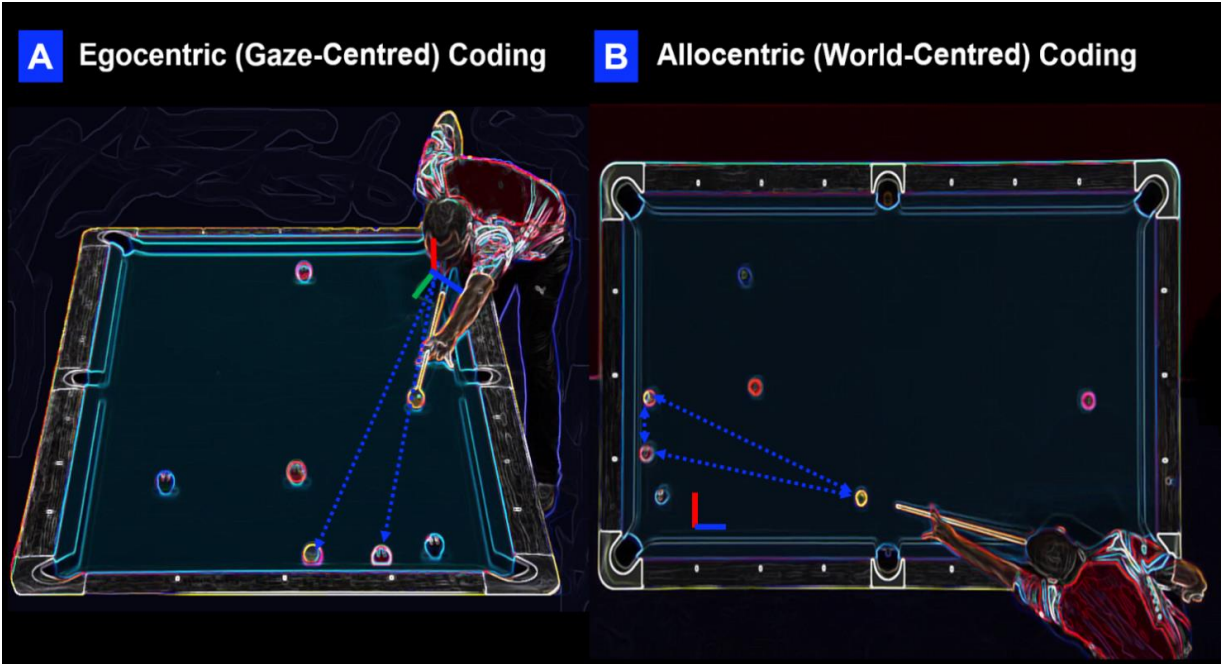


Figure 1. Spatial Encoding Strategies. **(A)** Illustration of gaze-centred egocentric encoding in a game of pool. Blue dotted lines show absolute positions of pool balls from a gaze/eye-centred egocentric perspective, whereas solid lines (x: blue, y: red, z: green) indicate the dimensions of spatial encoding. **(B)** Illustration of allocentric encoding in a game of pool, represented from a bird's eye view. Blue dotted lines show the relative positions of the pool balls in 2D (x: blue, y: red).

A variety of extrinsic factors influence the extent to which allocentric information is used, these include the proximity of the landmark, the stability of the landmark, task requirements, and the length of the memory delay. Various studies have found that the proximity of a visual landmark influences the extent it is utilized to improve reaching performance (Aagten-Murphy & Bays, 2019; Clark et al., 2007; Keller et al., 2005). Clark et al. (2007) investigated the influence of visual backgrounds on reaching performance and found that proximal and medial backgrounds were associated with a decrease in endpoint variability, while distal backgrounds were not. Aagten-Murphy and Bays (2019) similarly found that additional (non-target) stimuli improved precision when they were in close proximity to the target but had not affect on precision when they were too far to be used as landmarks. Other studies have shown that the perceived stability of a landmark can influence the extent to which allocentric information is used. Byrne and Crawford (2010) showed that landmark stability implicitly influences the extent to which allocentric information is utilized in a cue-conflict experiment, where there was a conflict between egocentric and allocentric information about target position. Relatively stable landmarks with low vibration amplitudes, that shifted their position to smaller degree when presented for a second time, had a more profound influence on final reach endpoints. Reach endpoints were biased in the direction of the landmark shift. The importance of task utility of a landmark has been established in a naturalistic cue-conflict experiment that involved 3D objects in virtual reality (Fiehler et al., 2014). It was found that allocentric cues were weighed more (reached endpoints were biased further in the direction of the shifted allocentric cue), based on their task relevancy in addition to their reliability (Fiehler et al., 2014). Studies investigating the influence of timing on allocentric encoding have consistently found that allocentric information tends to dominate over egocentric information after a long memory delay (Chen et al., 2011; Hu &

Goodale, 2000; Westwood et al., 2000,). When there is a longer memory delay between viewing a target and movement onset, participants reach less precisely to the target in an egocentric reference frame but with a close level of precision in an allocentric reference frame (Chen et al., 2011). Intrinsic factors like age have been found to influence how allocentric information is utilized. Lemay et al. (2004) showed that older adults used a different strategy when pointing relative to an allocentric cue, they had longer movement times, although there was no difference in how well allocentric cues are utilized by younger and older adults.

1.3 Egocentric Mechanisms for Reach

Allocentric information must ultimately be transformed into egocentric information to execute a movement, which is controlled by the musculoskeletal system and fixed in skeletal coordinates. To plan a movement, target locations that are initially in allocentric reference frame must be coded in a gaze-centred egocentric frame, that can be converted to a head, body, or mixed frame before precise motor commands are computed (Blohm et al. 2009; Flanders et al., 1992). The mechanism of conversion from an allocentric to an egocentric frame of reference is still not entirely understood. It has been shown that despite the stability of allocentric information, it is converted to egocentric information at the first possible opportunity (Chen et al., 2011; Chen et al., 2018).

Staged transformations are required to convert the internal representation of the target position into a muscular contraction of an effector, this involves transformations between internal egocentric reference frames (Crawford et al., 2004; Pouget & Snyder 2000). The location of the target must first be encoded in gaze-centred coordinates. In gaze-centered coordinates, the spatial location of the target is maintained as the retinal distance between the

current gaze direction (fixation) and the target location. The posterior parietal cortex (PPC) is thought to represent the position of objects relative to their location on the retina (Batista et al., 1999; Crawford et al., 2004). The representation of a visual target in gaze-centred coordinates also depends on the orientation of the eye and head (Medendorp et al., 2005). To orchestrate a movement towards a target, an effector-based reference frame is required to compare the difference between the current and desired position. For eye movements, an eye-centered reference frame is used (Sparks 1989; Klier et al. 2001). Gaze-centred representations are thought to be spatially updated across eye movements (saccades) (Henriques et al., 1998; Medendorp & Crawford, 2002). Spatial updating across saccades plays an important role in the memory of reach targets and movement planning (Batista et al., 1999; Crawford et al., 2004). Some studies have shown that gaze-centered updating persists, even after long delays (Fiehler et al., 2011). However, it has been associated with systematic pointing errors in the direction of gaze shift, where participants overshoot the remembered target location (Henriques et al., 1998).

Whereas, for reaching movements, the position of the arm must be estimated by integrating proprioceptive and visual signals (when the arm/object is seen) from the sensory periphery. Information about the target position must be compared with information about the current hand position. A movement vector is created by computing the estimated arm location (position of the fingertip or joints) from the desired target location (Khan et al. 2007, Vesia et al. 2008). This process is arguably believed to also occur in gaze-centred coordinates, since the hand position signals were found to be represented in gaze-centered coordinates within the PPC, even in the absence of vision (Blohm & Crawford, 2007; Buneo et al., 2002). The extrinsic movement vector must then be converted into a joint-based intrinsic motor command. Information in gaze-centred coordinates is transformed into shoulder-centered coordinates by combining them with current

eye and head positions (Henriques et al., 1998; Snyder, 2000). A shoulder-centered plan is required because the shoulder is the insertion point of the arm and arm movement direction needs to be specified relative to the spatial position and orientation of the shoulder. This is believed to be housed in the premotor cortex since the premotor cortex houses a representation of wrist movement, independent of wrist orientation (Scott, 1997).

1.4 Allocentric Representations in Reach

The utilization of allocentric information for movement has been tested in behavioural experiments where both type of object representations (egocentric and allocentric) are available for the brain to use. These types of experiments are typical of most real situations, where targets are not viewed in isolation and salient objects in the environment can be a source of allocentric information. There are different ways this scenario has tested in behavioural experiments, but I will distinguish them based on whether allocentric encoding was implicit, where landmarks were available during a task, but they did not need to be acknowledged to complete the task, or explicit, where landmarks were available during a task and they had be deliberately used to complete a task. The landmark shift paradigm is an example of an implicit allocentric encoding task. Participants are asked to remember the position of a target, surrounded by landmarks. During a short delay, the target disappears and the landmarks are subtly shifted, unbeknownst to participants. Participants must then reach to the remembered location of the target. If the participants only use an egocentric representation to remember target location, then they should accurately reach to the target position. If they use an allocentric representation, then they should reach inaccurately, and their reach endpoint should be impacted by the landmark shift. Evidence from behavioural research and imaging has shown that egocentric and allocentric information are

weighted based on their reliability based on a Bayesian account (Burgess 2006; Byrne & Crawford, 2010). Reaching biases introduced by shifted, unstable landmarks were quantitatively lower, in agreement with the prediction of a reliability-dependant weighting model that was parameterized on response variability from control experiments (Byrne & Crawford, 2010).

Stable implicit landmarks (landmarks that do not change their location relative to the target) allow for reliable computation of spatial relationships between targets and cues without effort from participants. Stable landmarks have been shown to improve performance, depending on when they are visible during a task. Landmarks could be visible all time during a task, as in the study by Krigolson and Heath (2004), where a continually visible landmark was used while only varying target visibility. Participants saw a target that was always visible, occluded at movement onset or 0-1 sec before the movement initiation cue. In the presence of a visual background, endpoints were more accurate and less variable in all target visibility conditions. Other studies have utilized transient landmarks. In Obhi and Goodale's (2005) study, participants were asked to reach to remembered visual targets presented with landmarks that either were only available during encoding or absent. They found that response variance was lower when landmarks were present, despite only being present during encoding. However, Aagten-Murphy and Bays (2019) found that when landmarks are present only during encoding, they have no beneficial effect. They utilized conditions where the landmark was continuously present during performance, disappeared after stimulus presentation and reappeared at the time of the probe, and conditions where the landmark was only present during encoding. They found that the presence of visual landmark substantially improved the localization of remembered stimuli in the landmark's vicinity only in the first two conditions, evidenced by improved precision of response (Aagten-Murphy & Bays, 2019). Byrne et al. (2010) also observed better performance due

landmarks that were available during stimulus presentation and re-presented before response, specifically on retinal magnification errors (RME). RMEs have been observed when performing memory-guided movements to isolated peripheral targets, where only egocentric information can be used. In those conditions, participants tend to exaggerate target eccentricity (Henriques & Crawford, 2000). Byrne et al. (2010), found that allocentric information can help reduce RMEs when gaze is directed to the right of the target, for right-handed participants.

Information about landmark visibility-based performance bears importance to considerations of when to present a landmark to influence movement during a behavioural task, but conclusions about when allocentric information is needed are varied. Chen et al. (2014) found that when only allocentric information is available to guide reach, it is converted into egocentric information at the first possible chance. This finding suggests that allocentric information calibrates information about target location during encoding and is then discarded. Conversely, Byrne et al. (2010) found that when participants made an intervening saccade, between target offset and reach, the influence of landmarks on RME was based on the final gaze direction. This observation implies that reduction in RME in the presence of visual landmarks did not update across eye movements, the brain appears to wait until the last possible moment to combine egocentric and allocentric information. This outcome, along with other experiments that have shown that landmarks are beneficial when they appear right before movement execution demonstrate that presence of a visual landmark facilitates the movement planning process, rather than simply stabilizing the target location. Velay and Beaubaton (1986) found that allocentric cues are even needed for online motor control mechanisms. Allocentric information in the form of a visual background led to the greatest improvement in performance when continuously

visible throughout the reach movement as opposed to only visible at movement onset or not available at all (Velay & Beaubaton, 1986).

Evidently, in experimental situations comparable to real life, where landmarks are stable and allocentric information can be combined with egocentric information, studies have pervasively shown that landmarks facilitate goal-directed action. This is especially true when the landmark is visible during the movement planning, in addition to encoding. However, it is still not clear if one form of spatial encoding leads to better performance than the other. Such comparisons are only feasible when allocentric encoding is explicit, and allocentric information is no longer congruent with egocentric information, as in the landmark shift paradigm. Only a few studies have tested this scenario. Thaler and Tod (2009) found that participants were better at egocentric encoding than allocentric encoding when participants were asked to point to the tip of line segment using their unseen hand. The stimuli were a solid line vector that was projected from marker indicating the position of their unseen hand, to be used to guide the direction of their hand movement, and a transient line segment, that they used to determine the distance they had to move their hand. There were 4 conditions: an eye-head condition, where the line segment and solid line vector were aligned, a hand-centred egocentric condition, where the line segment was not aligned with the hand vector but still projected from the egocentric midline and two allocentric conditions, where the line segment was either translated or rotated and translated. They found that pointing accuracy and reliability was highest in the eye-head egocentric condition, followed by the hand-centred egocentric condition and then the allocentric conditions.

On the other hand, Lemay et al. (2004) quantitatively compared participants' performance in an allocentric task, where a stimulus array shifted its position on the screen between encoding and recall but maintained the same relation relative to a rectangular frame and an

egocentric task, where there was no rectangular frame or shift in target position. Performance in the allocentric condition was found to be less variable than the egocentric condition. Chen et al. (2011) tested participants in an experiment involving variable delays and qualitatively compared participants' precision in an explicit allocentric task, to an egocentric task and an allocentric to egocentric conversion task. Delays were found to have a less detrimental effect on the allocentric task, suggesting that performance in allocentric task was better than the egocentric task.

Byrne and Crawford (2010) found no quantitative difference in precision between explicit allocentric and egocentric encoding in a reaching task, but when they compared spatial updating in both reference frames, they found that reaching precision in allocentric coordinates was better. Spatial updating in gaze-centered coordinates (fixation shifts) influenced precision in the egocentric control condition, but not the allocentric control condition. Spatial updating in the allocentric coordinates (landmark shifts), did not influence reaching precision in the allocentric task. Hence, despite contradictory evidence, an explicit allocentric seems to prevail over an explicit egocentric one when there is either a time delay or visuospatial transformation.

1.5 Neural mechanisms of Egocentric and Allocentric Spatial Representations

The cortical basis of visual representations is well explained by Goodale and Milner's highly influential 1992 perception-action model, where visual processing is proposed to be divided into a ventral stream for perception and dorsal stream for action. Visual perception is the conscious experience of the environment that can be expressed in subjective reports and tested using visual judgement tasks. The essential role of visual perception is thought to be for the representation of the visual environment, rather than storing information about objects for goal-directed action. The functional distinction of the two streams in this model is best explained by

their main output areas in the brain, rather than their input. Both streams arise from the same input area, the early visual area (V1), but the ventral stream projects to the infero-temporal cortex (IT) through visual areas V2, V3, V4 and the posterior inferior temporal area (TEO), whereas the dorsal stream projects to the posterior parietal cortex (PPC) via visual area V2, the medial temporal area (MT) and the medial superior temporal area (MST) (among other routes). The perception-action model does not imply that perception is only conducted by the ventral stream. An important aspect of this model is that both streams process perceptual information about objects, including their orientation, size and location, but this information is then used to carry out different functions.

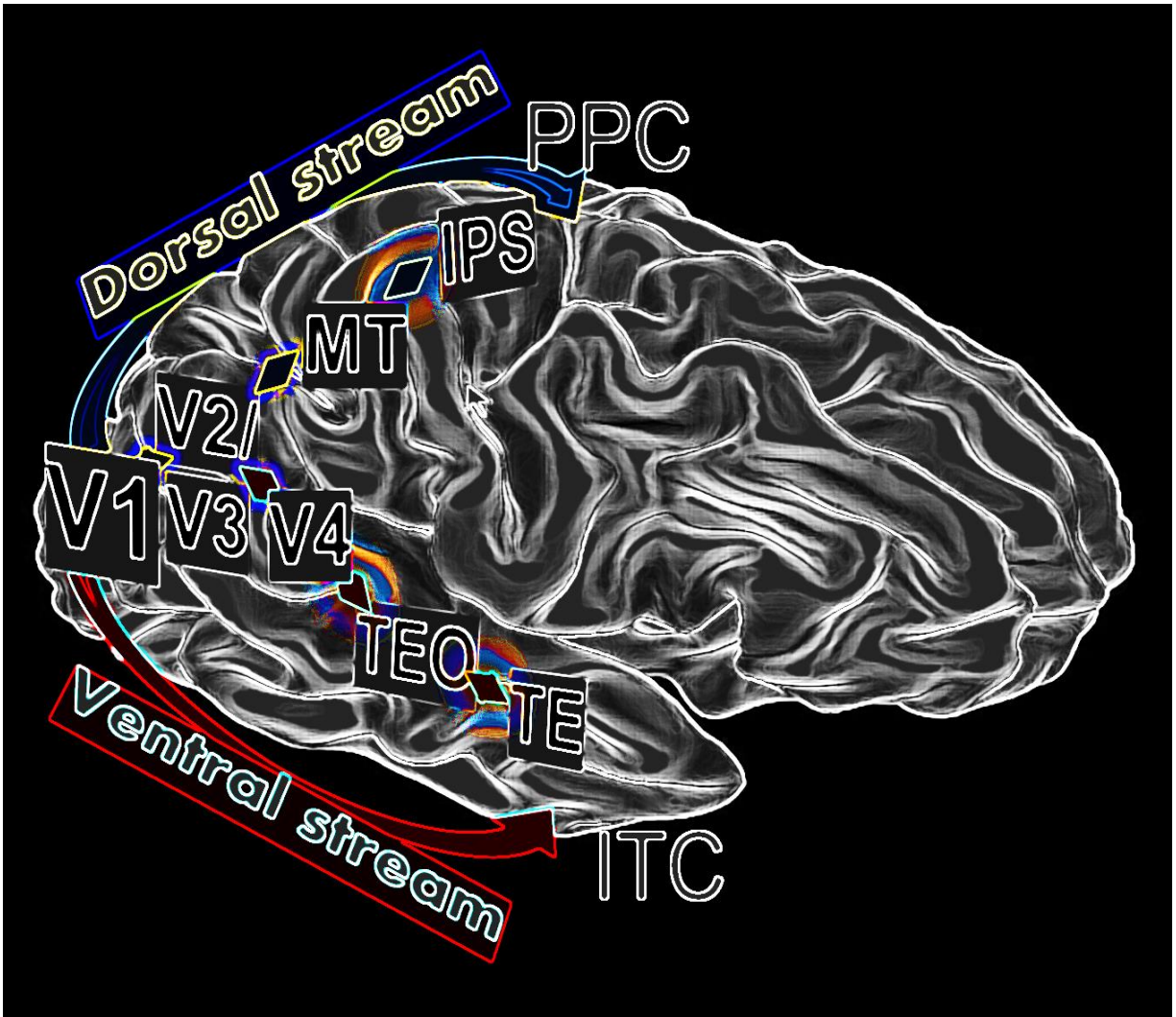


Figure 2. Major projections to the two visual streams of processing. The dorsal stream is shown in blue, and the ventral stream shown in red. V1-V4: visual areas; MT: middle temporal area; IPS: intra-parietal sulcus; PPC: posterior parietal; TEO: posterior inferior temporal area; TE: Inferior temporal area; ITC: Inferotemporal cortex. Diamond shaped connectors in the figure simplify projections between brain areas, which involve more complex interconnections.

Neuropsychological evidence vastly supports the perception-action model, particularly in the form of a classic double dissociation of deficits. Double dissociation is when an impairment in one brain area impairs one function and spares another, while impairment in another brain area demonstrates the opposite pattern of impairment. Double dissociation has been observed in two patient groups with lesions in either the dorsal or ventral stream. Patients with optic ataxia have impaired vision for action and spared visual perception, while patients with visual form agnosia have impaired visual perception accompanied by no visuomotor deficits. Optic ataxia is due to unilateral or bilateral damage to the posterior parietal lobe, leading to contralateral or bilateral reaching errors, respectively. Previous studies have shown that optic ataxia can occur without perceptual damage (Goodale et al., 1994; Jakobson & Goodale, 1991). A study examining an individual with optic ataxia in their ability to grasp objects found that, while the participant with optic ataxia was able to indicate properties of objects, such as size orientation and shape, they were unable to accurately orient their hand, scale their grip or accurately place their hand on grasp points (Jakobson & Goodale, 1991).

Patients with visual form agnosia generally have bilateral damage to the lateral parts of the occipital lobes, whereas the dorsal stream remains intact. D.F. is arguably the most well-known patient with visual form agnosia, who had extensive damage to the lateral occipital complex. She had difficulty identifying familiar objects and determining their orientation or shape but can maneuver objects with the correct orientation and accurately scale her grip aperture when grasping objects. In a seminal paper, Goodale et al. (1991) tested D.F. on a perceptual matching task, where she had to rotate a card to match the orientation of a visible slot. Her performance was barely above chance, while her performance was quite similar to normal controls when she was asked to post the card into the slot. D.F.'s preserved visuomotor control of manual

movements is believed to be due to her preserved dorsal stream, which facilitates her utilization of object qualities such as orientation and size to accurately perform movements. According to Goodale and Milner (1992), these findings suggest that perceptual judgements and motor actions involve functionally and neurologically distinct processing streams.

Schenk (2006) suggested instead, the card posting task requires an egocentric representation of visual and haptic space, whereas the perceptual matching task requires an allocentric representation of space. The dorsal and ventral stream are accordingly conceptualized as egocentric and allocentric processing streams. Spatial information about the absolute location of a target is encoded in egocentric coordinates for online motor command execution throughout the dorsal stream and ventral projections to the inferior temporal cortex represent increasing complex information about object properties, with decreasing sensitivity to spatial location and little retinotopic organization. Rather than being lost, information is increasingly represented as relationships between objects in the environment. The dissociation between egocentric and allocentric spatial representation is rationalized by emerging evidence from neuroimaging and behavioural studies. Various human neuroimaging studies have found that the dorsal parieto-frontal network of regions is associated with memory-guided reaches in an egocentric reference frame. The PPC and the dorsal premotor cortex (PMd) were found to be involved in reach planning in a memory delay period between viewing a target and reaching towards it (Connolly et al. 2003; Chen et al., 2014). Other regions that have been identified are the medial intraparietal sulcus (IPS) and precuneus, located in the superior parietal occipital cortex (Connolly et al. 2003; Beurze et al., 2009). When brain activation was tested in the presence of gaze shift, it was found that reach activity altered by gaze direction, suggesting that these regions are involved in a gaze-centred egocentric coding (DeSouza et al., 2000). Neuroimaging studies have frequently

identified correlates of allocentric coding for spatial judgment cognitive tasks and manual distance judgments tasks in the ventral stream, that are part of a widely distributed network (Fink et al., 2000; Thaler & Goodale 2011). Landmark centered coding has been found to be present in the inferior temporal gyrus (ITG) and the inferior occipital gyrus (IOG) (Chen et al., 2014). Hippocampal structures have been implicated in allocentric spatial encoding, in relation to spatial navigation studies, but there is also some evidence for its implication the coding of visual space (Byrne et al., 2007). Despite these findings, the neural basis of allocentric coding are far more poorly understood than the frameworks for egocentric coding. Due to the difficulty of disentangling interactions between egocentric and allocentric reference frames, it is still not clear if there are separable representations of non-retinotopic stimuli are encoded within distinct neural populations and pathways or whether they are instances of timely manipulation of egocentric information. The topic of whether there are separate allocentric reference frames is still being keenly debated.

1.6 Aims of the Current Study

Landmarks have a stabilizing effect on performance. This has been observed repeatedly; when participants had to wait longer to touch the position of a remembered target or when they had to perform some spatial transformation before precisely executing a movement to a remembered target (Byrne & Crawford, 2010; Chen et al., 2011). Yet, the practical implication of this landmark advantage needs to be investigated further. Does this entail that when landmarks are available when encoding target position, we will immediately see a performance improvement, or do we have to choose to attend to them? We investigated this effect in our study

by utilizing instructions. Whilst landmarks were always available during the presentation of stimuli, participants were given instructions that encouraged them to use them or ignore them.

The first goal of the current study was to determine the influence of instructions on performance when participants are asked to remember the position of the target only or remember the position of the target relative to the landmark. Unlike previous explicit encoding studies, the landmark is also present in the egocentric task and can implicitly influence encoding. The stimuli (target and landmark) that were presented were the same, but the tasks differed in the instructions for target encoding and motor execution. The second goal was to determine the influence of explicit allocentric cues on reaching performance when motor execution must be completed in the same or opposite visual field than stimulus encoding. Rather than requiring a gaze-shift to update the visual field, the landmark was represented in the same/opposite visual field and the egocentric task had to be executed in the same/opposite visual field. We anticipated that instruction based allocentric encoding will be more advantageous and facilitate performance, especially when more difficult spatial transformations are necessary to point to target with accuracy and precision.

CHAPTER 2
INSTRUCTION ALTERS THE INFLUENCE OF ALLOCENTRIC LANDMARKS IN A
REACH TASK

(A manuscript in preparation for submission)

Lina Musa, Xiaogang Yan, and J. Douglas Crawford

2.1 Abstract

The presence of an allocentric landmark can have both explicit (instruction-dependent) and implicit influences on reaching performance. However, it is not known how the instruction itself (to rely either on egocentric versus allocentric cues) influences memory-guided reaching. Here, 13 participants performed a task with two instruction conditions (egocentric vs. allocentric), but with similar sensory and motor conditions. Participants fixated their gaze near the centre of a display aligned with their right shoulder while an LED target briefly appeared (alongside a visual landmark) in one visual field, in the two instruction conditions. The landmark then re-appeared in the same or opposite visual field after a brief mask/memory delay period. When pointing in the allocentric condition, participants had to remember the initial target location relative to the landmark and reach relative to the shifted landmark. Whereas in the egocentric condition, participants had to ignore the landmark and point towards the remembered location of the target. To equalize motor aspects (when the landmark shifted opposite), subjects were instructed to anti-point (point opposite to the remembered target) on 50% of the egocentric trials. In both pointing scenarios, whether reaches were executed in the same or opposite visual field, participants performed better (more accurately and precisely) in the allocentric condition. We also observed a visual field effect, where performance was worse overall in the right visual field but pointing in the allocentric condition showed the highest advantage in performance over the egocentric condition. These results entail that explicit attention to a visual landmark better recruits allocentric coding mechanisms that can augment implicit egocentric visuomotor transformations. These allocentric coding mechanisms also appear to be recruited in a hemi-spherically asymmetric manner.

2.2 Introduction

Humans use spatial reference frames to retain information about a target that is not immediately visible for goal-directed action. An exact memory representation of the target is needed to execute accurate movements towards the target location. Spatial reference frames can be thought of as coordinate systems that map object locations in space (Soechting & Flanders, 1992). The visual system is known to utilize two spatial reference frames, an observer-centered, egocentric reference frame and a world-centered allocentric reference frame, that is anchored to reliable objects in the environment called landmarks (Byrne et al., 2007; Howard & Templeton, 1996; Vogeley & Fink, 2003). The brain typically integrates the two reference frames and has been shown to do so in Bayesian manner (Byrne & Crawford, 2010; Fiehler et al., 2014). However, circumstances do arise where one spatial reference must be preferentially used, especially in non-naturalistic settings. Behavioural tasks as such manipulate the extent to which an alternative reference frame permits a solution, or participants may be instructed to use one reference frame over the other (Byrne & Crawford, 2010; Chen et. al, 2011, Lemay et. al, 2004). An opportunity that arises in these kinds of tasks, is asking which of the spatial reference frames leads to better performance?

Various studies have looked at the influence of spatial reference frames in the context of goal-directed action, such as making a saccade, reaching, or pointing towards a seen or remembered target. Previous findings have shown that targets can be remembered reasonably well in an egocentric reference frame, in the absence of visual landmarks (Henriques et al., 1998; Lemay & Stelmach, 2005; McIntyre et al., 1997; Vindras & Viviani, 1998). In studies that involve the integration of egocentric and allocentric information, reach endpoint errors or

another metric of performance is used to investigate the influence of a landmark that can be either present throughout the task, or at some points within the task. It has been generally found that humans reach and point more accurately in the presence of both egocentric and allocentric cues (Byrne et al., 2010; Krigolson & Heath, 2004; Redon & Hay, 2005). These studies typically manipulate allocentric information implicitly, through the presence of visual background or an allocentric landmark adjacent to a target, which participants do need to deliberately use. These studies are analogous to most normal situations, where allocentric landmarks are stable and agree with egocentric cues, so that their integration leads to the best estimate of target location. When egocentric information and allocentric information are put into conflict, it has been found that they are optimally integrated on the basis of their reliability (Byrne & Crawford, 2010).

The influence of allocentric cues can also be separated to some extent experimentally. Behavioural studies comparing the fidelity of egocentric and allocentric representations in memory-guided reaching have found conflicting outcomes. Using a shifted landmark paradigm, Lemay et al. (2004) report not only less variability in the allocentric task, but also similar to variability in the allocentric-egocentric integration condition. Thaler and Tod (2009) found reaching to be worse relative to a rotated and translated landmark. Explicit allocentric tasks can also involve instructions to encode targets in an allocentric reference frame. In such scenarios, participants are not just explicitly using target-to-landmark distance judgements because they are needed to facilitate accurate reaching, they are intentionally choosing to do so, even when other strategies are available. One such task investigated the stability of spatial reference frames due to time delays and found that egocentric spatial representations decay rapidly, whereas allocentric information can be maintained for several seconds (Chen et al., 2011). Byrne & Crawford (2010) utilized explicit instructions in an egocentric and allocentric control task, to compare egocentric

and allocentric task reliabilities, but found no difference in variability of performance. However, they did find that the reliability of egocentric information degraded with the extent of gaze-shift, whereas allocentric reliability was not influenced by the magnitude of change in landmark position (Byrne & Crawford, 2010).

Previous studies either indirectly compared allocentric and egocentric conditions or utilized an egocentric task where the allocentric cue is absent. In our study, we investigated the influence of an instruction (to use or not use a landmark) in a reach task where the allocentric cue was present in both conditions. By utilizing the same stimulus display for the allocentric and egocentric tasks while varying the instructions, we were able to dissociate the explicit effect of a landmark from its implicit influence during encoding. It has been previously demonstrated that landmark-centered encoding is a less noisy process and more stable over a time delay (Byrne et al., 2010; Chen et al., 2010). Thus, we predict that reliance on an allocentric landmark should improve performance. This should be especially true when egocentric information becomes less reliable, such as when there is a spatial updating component (as in an anti-reach task). We found that 1) reaching was less variable and more accurate in the allocentric instruction tasks than egocentric instruction tasks. However, the influence of a landmark on precision in performance was visual field dependent. 2) the beneficial effect of allocentric encoding is more pronounced when there was spatial updating and participants had to respond in the visual field opposite to encoding.

2.3 Material and Methods

2.3.1 Participants

13 individuals (7 males and 6 females; ages 20 – 33) gave their informed consent to participate in the study. All participants were right-handed with no neuromuscular deficits, based on self report. Participants also reported normal or corrected normal vision and intact color vision. Data from 3 participants was excluded from further analysis because they were not able to maintain adequate fixation for a critical duration, resulting in a total of 10 participants (5 males and 5 females). All participants were naïve to the purpose of the experiment and were given monetary compensation for their time. The experimental procedures were approved by the York Human Participants Review Subcommittee.

2.3.2 Apparatus

The experiment was conducted in complete darkness. Participants sat behind a table, on a chair with an adjustable height (Fig. 1). The height of the chair was adjusted so that their right shoulder bone (acromion) was aligned with the centre of a LED panel. Their head was stabilized on a personalized bite bar made of a dental impression compound (Kerr Corporation, Orange, CA). Participants' right hand was positioned on a button box on the table-top, directly in front of them. The button box was used to control the pace of the experiment and was the designated start position for reaching. A customized ring with a 3x3 array of Infrared-emitting diodes (IREDS) was attached to participants' right index finger and their 3D-position was continuously recorded by two OptoTrack 3020 tracking systems (Northern Digital, Waterloo, Ontario, Canada). Gaze direction was monitored from only the right eye through the EyeLink II infrared eye-tracking system, (SR Research, Mississauga, Ontario, Canada), that was mounted to a bite bar stand.

Visual stimuli were presented in a customized 43 channel, wooden LED panel (307 mm x 161 mm), placed 50 cm away from their eye position. Horizontally arranged 1.15-degree dots of light were created by drilling holes 1 cm in diameter & 1 cm apart and illuminating them by inserting 3 mm LEDs. The LED panel was positioned 152 mm to the right and 91 mm below the centre of the bite bar stand. At that position, the centre of the LED panel closely coincided with the shoulder bone position of most participants. Audio instructions were delivered from a speaker, at the beginning of each trial and when they had to touch the LED panel. Two 40-watt desk lamps, placed on either side of the desk, were turned on every 10 trials (2 min), to eliminate potential dark adaptation.

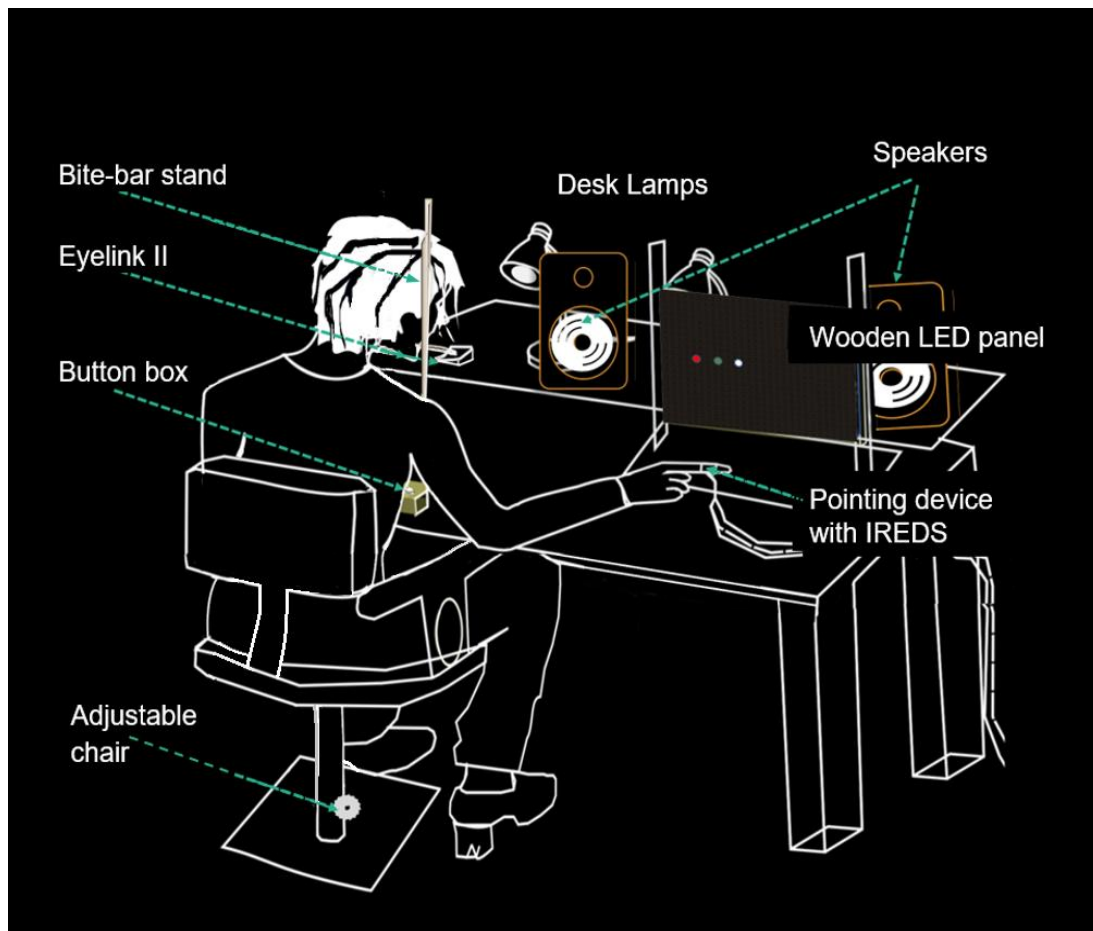


Figure 3. Experimental set-up. From left to right and top to bottom: custom made metallic bite-bar stand with personalized dental impression bite plate (not visible in the figure), used to stabilize participants' head position. Eyelink II cameras (SR Research, Mississauga, Ontario, Canada) were dismantled from the headset and installed on the bite-bar stand. The left camera is occluded in this view of the set-up. Black button box (yellow in the figure), on the desk immediately in front of participants was used to control the pace of the experiment. Participants' height was adjusted relative to the fixed bite-bar height using a metallic screw on a rigid chair with an adjustable height. Two 40W desk lamps illuminated the dark room during breaks and every 3 trials. Audio instructions were played from two desktop speakers. A custom made 307 mm x 161 mm wooden panel, fitted with light emitting diodes (LEDs) was used to display stimuli. The pointing device was customized ring with a 3x3 array of Infrared-emitting diodes (IREDS) that continuously relayed signal to two OptoTrack 3020 tracking systems (Northern Digital, Waterloo, Ontario, Canada), on either side of the room (not shown in the figure).

2.3.3 Task and stimuli

Before the start of the experiment, participants completed a set of calibration procedures, that were also conducted in complete darkness. Eye calibration was done through sequential presentation of a 5-dot display, consisting of a white dot at the centre of the LED panel, along with 4 evenly spaced white dots at the corners of the LED panel (See Fig. 2A). Participants sat head-fixed, left hand rested on their lap and right index finger on the button box, while they followed the 5-dot calibration sequence with their eyes. They then completed two OptoTrak calibration sessions. Finger-tip position was calibrated by having participants point a pre-calibrated cross rigid body with IREDs to the tip of their right index finger. A screen calibration was done by having participants successively point to the 4 corner positions in the LED display. IRED position data in the OptoTrak intrinsic coordinate system was combined offline with known coordinates for the calibration dot positions to create a linear mapping between IRED positions and the screen coordinates. This procedure allowed for the conversion of finger-tip position into screen relevant coordinates.

The experiment involved reaching and pointing using the right index finger in complete darkness. Participants wore a black glove that further visually obscured hand position. The task involved touching the remembered location of a transient visual target that always appeared simultaneously with a landmark, as accurately as possible. Task relevant visual stimuli (the target, landmark and fixation) were all circular, arranged in a horizontal array of 25 dots that were 1cm in diameter and 1cm apart (1.15 degree when viewed from the centre of gaze). All the stimuli displayed by the LED panel are annotated in Fig. 2A. Briefly, the LED panel displayed a green-coloured target, in one of 9 locations, to the right or the left of the centre of the LED panel (18 locations in total), starting at a 4.57 degree distance from the centre. Of the 9 target locations

(at the right and left), the 3 central locations were designed to also be illuminated a red colour and were used to display the landmark. Thus, the red landmark appeared at either 7.96, 9.09 or 11.31 degrees, to the right or left of the centre of the LED panel, there were 6 in total. A white fixation dot appeared in one of 7 locations in the centre, the central fixation also corresponded with the centre of the LED panel and the remaining 3 were to its right and left. The central fixation and 20 white dots of light on either side of the display were illuminated to create a mask (40 in total), they were located halfway in between target locations, flanking both sides of each target, but located above and below it. The same stimuli were used for two experimental conditions: an egocentric and an allocentric instructions condition. The conditions were randomly interleaved, to eliminate practice and fatigue effects. The protocol for each condition will be summarized in-depth, in the next section. The protocol was adapted from Chen et al. (2013), where participants were asked to perform a similar task using a gaze-centred egocentric reference frame in an MRI scanner.

2.3.4 Experimental Design

Participants completed the task in 3 blocks of 72 trials and were allowed adequate time for breaks, during which the two lamps were illuminated: they had a 5 min rest period in between blocks and four, inter-spaced 2 min breaks within blocks. Individual trials lasted 12.05 s in both instruction conditions, but the experiment was self-paced, subsequent trials were initiated through a button press. On average, participants took 21.65 mins to complete an entire block. The order of trials within each block was randomly interleaved. It was also randomized between blocks and every participant. To familiarize participants with the task, they were given 3 practice blocks of 10 trials that were not recorded. The first two blocks were made up of each instruction condition separately and the last block was randomly interleaved.

2.3.4a Egocentric Instructions Condition

In the egocentric instructions condition, participants always heard “Reach to target” from the speakers, at the beginning of the task. These instructions prompted participants to remember the location of the peripheral green target relative to the screen midline that was aligned with their shoulder, while ignoring all else. The white fixation dot then randomly appeared in one of the 6 off-centre positions, it remained in that location for 2s and during the next step. Participants knew to maintain fixation throughout the task since they were instructed to do so prior the beginning of the task and had practice fixating throughout the training blocks. In the following step, the target and landmark simultaneously flashed for 2s, on the same side relative to the centre of the LED panel. Their position was pseudorandom: the 72 trials within each block are exhaustive of all the landmark and target position combinations (and instruction conditions), each trial is random instance of those potential combinations. Subsequently, the fixation dot shifted to the centre of the LED panel, and stayed there for the remainder of the experiment. This step is the same in all trials. By having the fixation dot shift from its original position, the possibility of using the fixation dot as a landmark is removed. The 40-dot mask appeared briefly, for 150 ms, to eliminate the potential of after images and was followed by a brief delay period of 400 ms. The landmark then reappeared for 2 s, always at a different location than it was initially viewed. This was immediately followed by one of two audio instructions. When they heard “Target”, participants were required to reach-to-touch the LED panel at the same exact position of the remembered target but when they heard “Opposite”, participants were required to reach-to-touch the LED panel at the mirror opposite position, relative to the screen midline. Participants were allowed a response period of 4 s, after which they freely initiated the next trial by clicking the button.

2.3.4b Allocentric Instructions Condition

This condition was similar to the egocentric instructions condition, but participants now heard “Reach relative to cue” from the speakers. In this task, participants were required to remember the position of the green target, relative to the red landmark. When the landmark reappeared in a shifted position and they heard “Reach”, participants had to reach-to-touch based on the remembered target to landmark distance and the updated landmark position. Since the landmark always shifted to a different position than it was initially, the final location of the target could not be inferred from its initial position relative to the screen midline (egocentric distance). Thus, in the allocentric condition, participants had to use the landmark position to encode and store the target location.

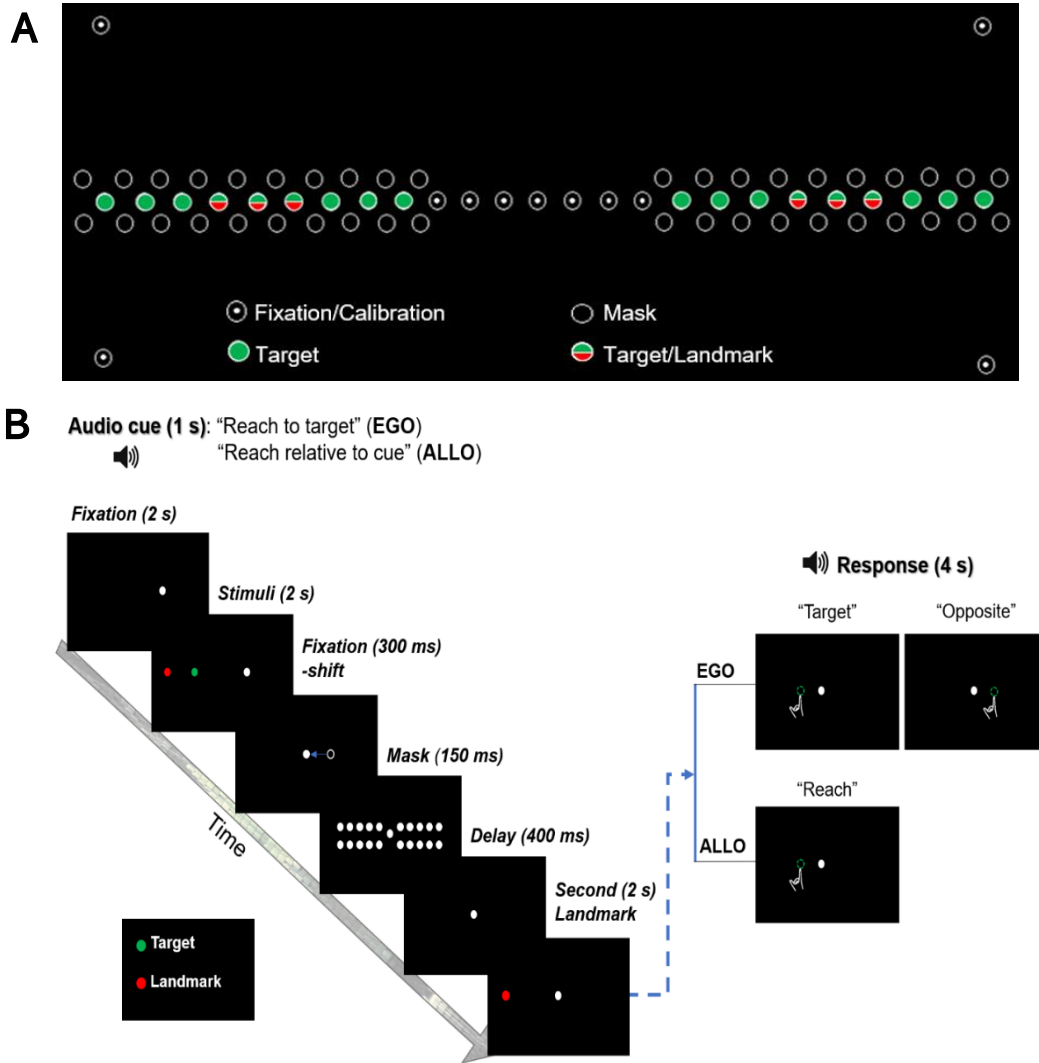


Figure 4. Experimental Stimuli and Paradigm. **(A) Stimuli.** The stimuli were displayed in a horizontal array. The central fixation, off-centre fixation and eye calibration points were displayed by white LEDs. The green target was randomly displayed in one of the 18 (9 L, 9 R) LEDs, 3 circles (4.57°) from screen centre. The first red landmark simultaneously appeared on the same side of the screen as the target, displayed by one of the 3 LEDs in the middle of the 9 target LEDs (half red, half green circles in the figure), while the second red landmark randomly appeared in the remaining landmark positions. **(B) Paradigm.** The order of a typical trial is shown in the figure. Each trial began with an audio instruction, where participants were instructed to remember the spatial location of the target, or the position of the target relative to the landmark. Participants fixated on a white circle, while the green target and red landmark briefly appeared. They then executed an eye movement to the screen centre. After a mask and a short delay, the red landmark reappeared at a different spatial location. This was followed by a response period. In the egocentric task, participants heard "Target", when they were required to point to exact spatial location of the target and "Opposite" when they were required to point to the mirror opposite. In the allocentric task, participants heard "Reach" when they had to reach to the remembered position of the target, but relative to the second landmark.

2.3.5 Sample Size Analysis

Sample size was calculated using the SIMR R package, which calculates the power for generalized linear mixed models from the LME4 package (Green & MacLeod, 2016). The power calculations were based on Monte Carlo simulations. Reaching variance was obtained from a pilot study of 2 individuals and 80 observations for each instruction condition (egocentric and allocentric). A mixed effects model was fitted on pilot data, to obtain an estimated effect size of 0.267 for reaching variance. Using the obtained effect size and the simulation package, the sample size was gradually increased to achieve a power of 80%. The sample size required to achieve this power was 12.

2.3.6 Data Analysis

All data obtained from OptoTrak and EyeLink was analysed offline using a custom software written in MatLab R2019a (Math Works). A program was written to generate a mapping between the Optotrak coordinates and screen-relative coordinates of the fingertip. This was done by utilizing the data exported from the screen calibration session and known screen coordinates of the 4 calibration points on the screen corners. This mapping was then utilized to obtain reach endpoints in screen coordinates for analysis. The movement kinematics were inspected to ascertain that participants followed instructions well. The red line in Figure 3 shows the 2D- reach performance of one participant in the LED screen coordinates, from the participant's point of view. The panel on the left shows the 2D data for the egocentric task in the "Opposite" condition and right shows the 2D data for the allocentric task, when the landmark appeared in the opposite hemifield. Only trials where participants did not make any anticipatory reaches (reaches initiated before the audio instructions) were included for further analysis. Similarly, the data from the 5-dot eye calibration was used to create a mapping of eye position in

screen coordinates. 2D eye movements are shown in green, in Figure 3. Maintaining fixation was important throughout the task but a critical aspect of task performance was for them to have not looked at the landmark (during the allocentric task) when presented for the second time. If that were the case, participants would have converted the spatial location of the target into a gaze-centered egocentric reference frame, earlier than required. A program was written to automatically average eye fixation, from the period of the second landmark to the go signal. Trials were excluded if eye variation was greater than 2 degree in the horizontal direction for more than 20% of time. Only trials where the participants met the criteria for reaching and eye fixation were analysed for reaching performance. Overall, 10 participants who had at least 173 viable trials (80%) were included in the study. Participants were asked to maintain fixation while reaching. However, different oculomotor strategies were found to be used. Some trials showed that participants stabilized their gaze at the fixation while reaching, while some trials showed that participants broke fixation and made a saccade to the remembered target location. Participants used the two different strategies to different extents, but they were not excluded for that reason. Van Pelt & Medendorp (2007) found that reaching error is the same, regardless of whether participants maintained fixation or made a saccade towards the vicinity of the remembered target location.

Using a customized MatLab program, movement start was marked at 20% maximum resultant velocity (V_{max}), while movement end was marked at 8% of V_{max} . Reach endpoint was found by averaging 30% of points at 8% V_{max} and at 1.15 times the minimum finger to screen distance (distance in the z direction). The screen relative coordinates of reach endpoints in the horizontal (x) and vertical (y) dimension were used to compute the variable error, a measure of the distance of reach endpoints from the mean final position. R 3.6.3 (R Core Team,

2017) was used to create 95% confidence ellipse of the scatter of reach endpoints. The area of the ellipse was found by first computing eigenvalues (σ_1 , σ_2) of the covariance matrix of reach endpoints. The eigenvalues were then used to derive half the lengths of the of the semi-major (principal axis) and semi-minor axes (orthogonal to the principle axis) of the 95% confidence ellipse and used to compute the ellipse area (formula shown bellow). The ellipse areas were then used to compare the egocentric and allocentric instructions conditions, using paired t-tests.

$$\text{Ellipse area} = \pi * \text{sqrt}(5.991 * \sigma_1) * \text{sqrt}(5.991 * \sigma_2)$$

The horizontal reach endpoints were used to analyse their response accuracy, computed in allocentric and egocentric coordinates. The data was analysed in egocentric coordinates by computing the distance between reach endpoint to screen midpoint and in allocentric coordinates by computing the distance between the reach endpoint to the second landmark. Participants' response in the control session, that involved visually guided reaching (their correct response), was compared to their performance in the task conditions using a monotonic predictive model using the brms package (Bürkner, 2017). Their constant error was also computed by finding the difference in their reach endpoint to the actual target location and overshoot error was found by averaging the magnitude and direction of the reaching error (positive to the right, negative to the left). A mixed effects model was used to analyse the reaching error (lme4 package) and interaction contrasts (emmeans package) were used to analyse the constant reaching error. Their reaction times were fit with an ex-Gaussian model (using the retimes package) and an repeated measure ANOVA was used compare parameters from the fitted model. Post-hoc analyses were Bonferroni corrected.

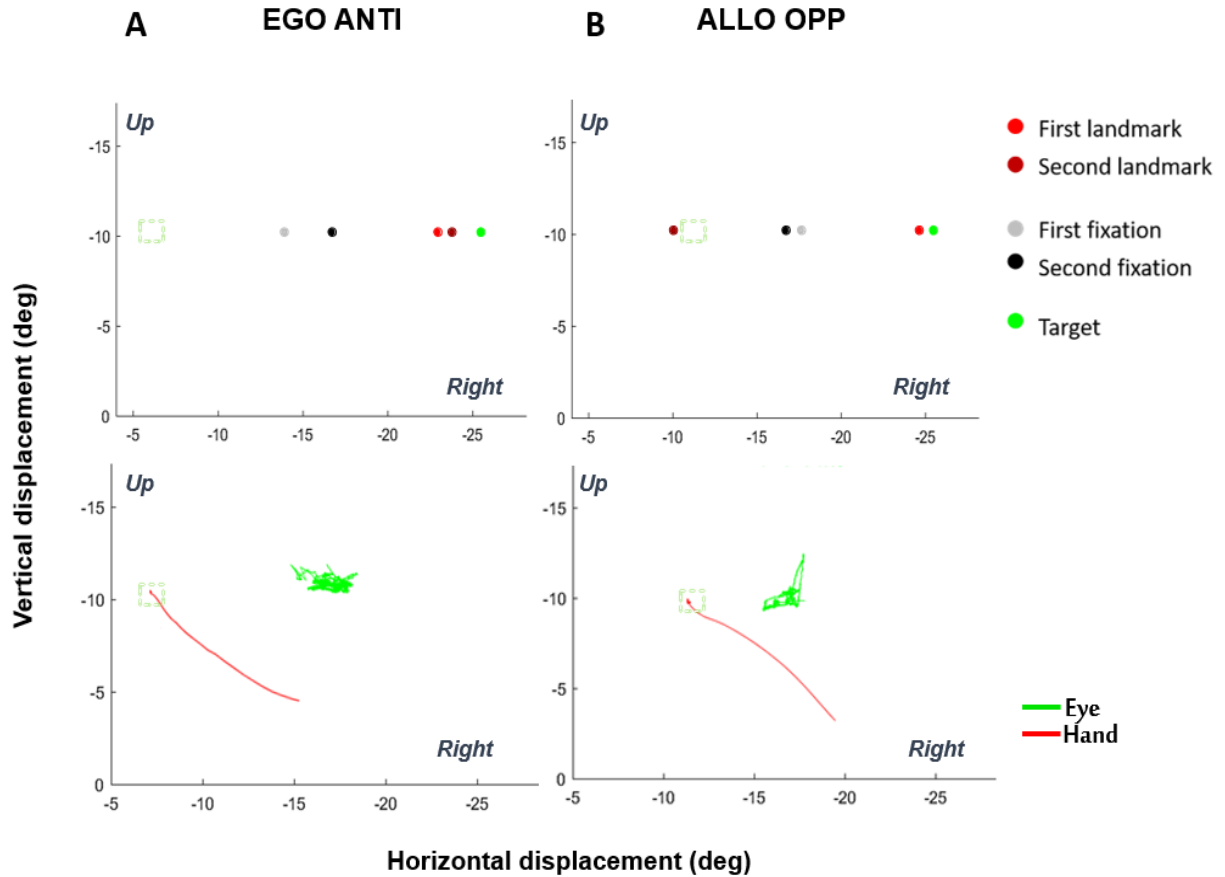


Figure 5. Typical eye and finger trajectories. **(A)** Anti-egocentric trial, where the participant was instructed to touch mirror opposite spatial location of the remembered target position. **(B)** Opposite allocentric condition, where the landmark appeared in the opposite side of the fixation. *Top:* the spatial location of the stimuli. The dashed square outlines the expected target position. *Bottom:* The green line shows the 2D gaze location, while the red line shows the 2D finger trajectory, immediately before touching the screen. The dashed green square shows the expected target location.

2.4 Results

The purpose of my analysis was to determine if there is a difference in the utilization and execution of feedforward plans, using exclusive egocentric or allocentric information for reaching and touching a transient visual stimulus. I was also interested in determining if performance is impacted by spatial updating, whether executing movement to the remembered target in the same/opposite visual field would influence performance in a different way in both spatial information conditions. Finally, I also determined if there was a difference based on visual field (right or left) on performance in each condition. In order to do so, I looked at the difference of precision, accuracy, reaction time and movement time. Before conducting the analysis, I excluded trials using the criteria stated before. This included trials where participants did not maintain adequate fixation and trials where participants executed a movement to the opposite side than instructed or executed a movement before the Go signal. A total of 2016 trials were recorded from the 10 participants, 216/ participant. Participants completed on average 72.90 ± 20.58 egocentric instructions trials, 36.20 ± 9.98 of trials were trials in which the target was viewed and the movement was executed in the same visual field while 36.70 ± 10.60 were the movement was executed in the opposite visual field. Participants also completed 68.40 ± 19.82 allocentric trials, 32.70 ± 7.81 of trials were on the same visual field and 35.70 ± 12.99 of trials were in the opposite visual field.

2.4.1 Reaching variance

2D reach endpoint ellipses were all calculated at each target location for all participants, using their endpoints from individual trials. Figure 6A shows ellipses (unfilled circles) overlaid on top of the reach endpoints (filled circles), plotted for an individual participant. The plots for the egocentric task are shown in the left panel while the plots for the allocentric task in the right

panel. The top/bottom panels in both task conditions show trials were reaches were executed to the same/opposite visual fields. Figure 6B shows the plots mean ellipses of participants in all task conditions. The mean ellipses at each target location and condition were constructed by averaging covariance matrices of the corresponding ellipses of individual participants. The eigenvalues of the mean covariance matrix were used to derive the length of major and minor axes of the mean ellipse. The eigenvector associated with the largest eigenvalue of the mean covariance matrix was derived and used to find the associated major axis direction of the mean ellipse. The areas of the average ellipses for each condition are shown in Figure 7A. Figure 7B shows the ellipse areas averaged over the left and right visual field targets.

The influence of 2 (spatial instructions: Ego/Allo) x 2 (spatial updating: same/opposite) x 2 (visual field: left/right) on precision, quantified by the ellipse areas, was analysed using a generalised linear mixed model. The outcome of the analysis revealed 2 two-way interactions between spatial instructions and visual field, $t_{(219)} = 2.1376$, $p = .0337$ and spatial updating and visual field $t_{(219)} = 2.0121$, $p = .0451$. Post-hoc interaction contrasts revealed a conditional main effect of spatial instructions. The ellipse areas of the allocentric instruction conditions were significantly lower than the ellipse areas of the egocentric instruction conditions, but only in the right visual field, $t_{(219)} = 1.994$, $p = .0360$. There was no significant difference between the instruction conditions in the left visual field. This outcome reveals that allocentric instructions lead to improved precision, but it was dependent on the visual field of response. Post hoc analysis of the second interaction effect revealed a difference of ellipse areas between the spatial updating conditions was significantly higher in the right visual field than in the left visual field, $t_{(219)} = 2.099$, $p = .0307$. Thus, only when participants responded in the right visual field, did we see a significant reduction precision when pointing anti/opposite.

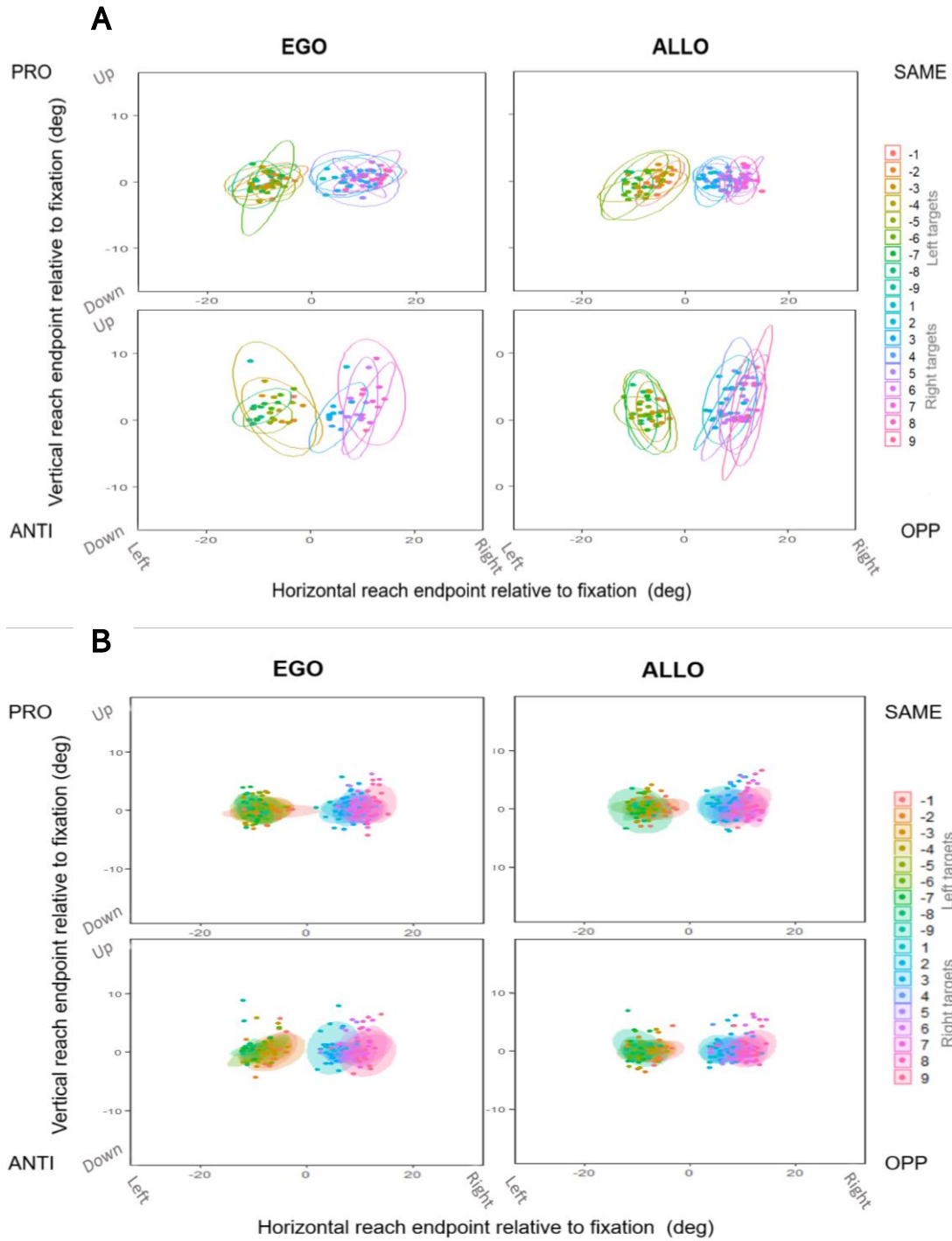


Figure 6. Reach Endpoint Ellipses. **(A)** Single participant 95% confidence ellipses. The small filled circles are reach endpoints relative to the central fixation. **(B)** Mean 95% confidence ellipses of 10 participants. The filled circles show the mean reach endpoints relative to central fixation. Ellipses and reach endpoints are all color coded by initial target spatial location. The left panels show the egocentric conditions, pro (top) and anti (bottom), and the right panels show the allocentric conditions, same (top) and opposite (bottom).

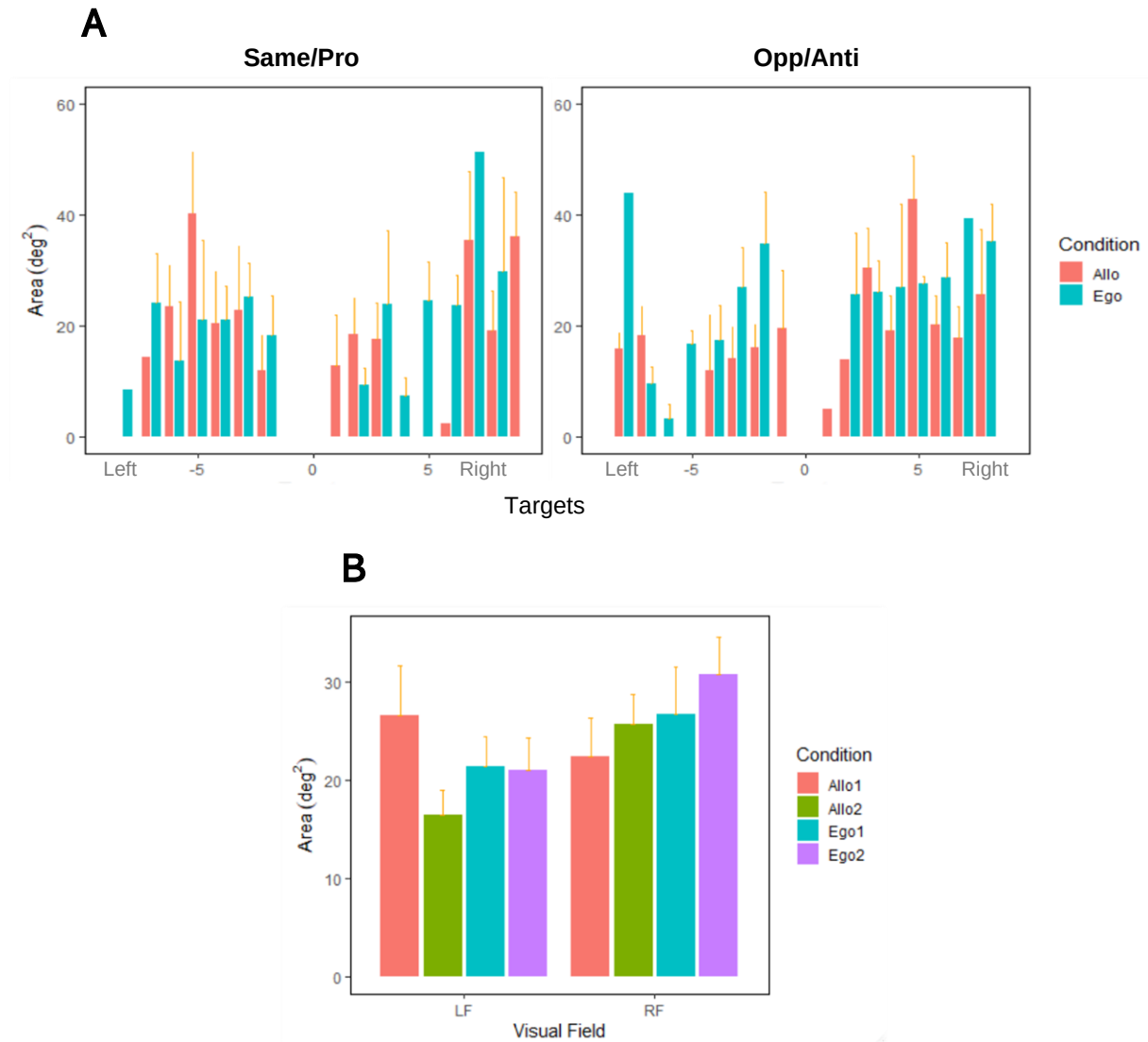


Figure 7. Reach Endpoint Ellipse Areas. **(A)** Bar plots of individual target locations. The green and red bars depict the egocentric and allocentric conditions, respectively. The left panel shows the pro-/same conditions, whereas the right panel shows the anti/opposite conditions. **(B)** Bar plot of ellipse areas averaged over right/left visual field targets. The bottom plot summarises the observed trends in the left and right visual field for each condition. Egocentric and allocentric conditions are denoted by Ego/Allo in the figure legend. The numbers represent spatial updating conditions, 1: pro-/same; 2: anti/opposite. The error bars in the plots show the SE of the mean.

2.4.2 Reaching accuracy

To quantify the accuracy of response, a measure of participant's reach endpoints was first obtained in egocentric coordinates, relative to the screen centre and in allocentric coordinates, relative to the landmark. These two outcomes were analysed separately. In egocentric coordinates, data was analysed through pairwise comparisons of monotonic regression models. The predictors were target locations and the response variables were reach endpoints. Since target locations were bimodally distributed, a model robust to violations of the normality assumption was necessary for analysing the data. The fit of three models on control data were compared: a model with linear predictors (linear regression model), a model with ordinal predictors (monotonic model) and model with nominal predictors. The monotonic model had a significantly better fit, indicated by the lowest Bayesian estimate of log pointwise predictive density. The outcome of this analysis did not reveal any significant effects. For the reach endpoints in allocentric coordinates, the best fitting model was a linear model, thus pairwise comparisons of linear regression models were done. Similarly, analysis in allocentric coordinates revealed no significant results.

The accuracy of response was also quantified by calculating participants' horizontal reaching error. In this analysis, reach endpoint errors were found relative to the absolute target locations, rather than performance at baseline. The constant error was found, utilizing only the magnitude of the reaching error, as well as the overshoot error, which also takes into account the direction of the error (whether participants pointed to the left or right of the targets on average). The difference in reaching error across 2 (spatial instruction conditions) x 2 (spatial updating conditions) and 2 (visual fields) was analysed using a generalized linear mixed model. The analysis of overshoot reaching error revealed a three-way interaction effect, $t_{(2008)} = 1.99$, $p = .$

0457. The interaction effect was analysed by post-hoc interaction contrasts. A conditional effect was found in overshoot error across spatial updating conditions (opp – same and anti – opp), between the allocentric and egocentric condition. In the right visual field, there was a significant difference in the way participants performed in the spatial updating conditions between the allocentric condition and egocentric condition, $t_{(2008)} = 2.30$, $p = 0.015$. Pairwise t-tests revealed that there was no significant difference due to spatial updating (opp-same) in the allo condition but there was a significant difference due to spatial updating in the egocentric condition $t_{(2008)} = 1.89$, $p = 0.045$. Whereas, in the left visual field, there was no significant difference between egocentric and allocentric conditions, during spatial updating. The outcome of this analysis demonstrates that participants tended to overshoot more when they had to anti-point in the right visual field, but there was no increase in overshoot error when they had to point relative to landmark in the right visual field. The presence of landmark mitigates overshoot error, only in the right visual field. The constant pointing error was analysed in a similar manner and the outcome of this analysis revealed a two-way interaction effect of instruction conditions x spatial updating. Post-hoc contrasts revealed a significant difference between the instruction conditions, but only when participants had to point anti/opp, $t_{(2008)} = 2.16$, $p = 0.023$. These results demonstrate that allocentric instructions best improved pointing accuracy when spatial updating was part of the task.

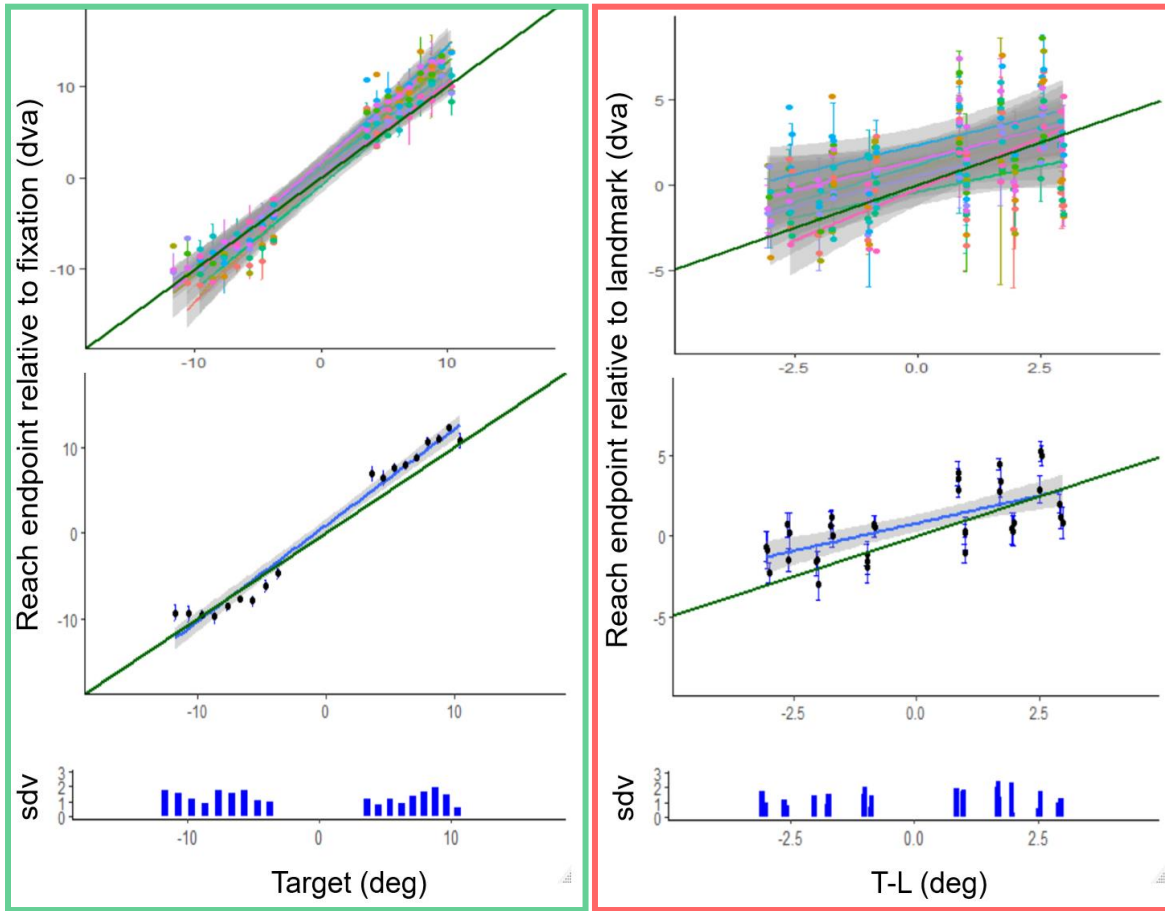


Figure 8. Regression Plots of Pro-Egocentric task. Green panel: egocentric coordinates. Red panel: allocentric coordinates. Top-green: Participants' reach endpoint relative to fixation, averaged at each target location and plotted against target-to-fixation distances. Top-red: Participants' reach endpoints relative to first landmark position averaged at each target-to-landmark distances and plotted against target-to-landmark distances. Error bars show standard deviations of reach endpoints and the shaded grey outline shows the standard error of the estimate. Bottom-green: Mean reach endpoints relative to fixation of 10 participants, averaged at each target location and plotted against target positions relative to fixation. Bottom-red: Mean reach endpoints relative to first landmark positions of 10 participants, averaged at each target-to-landmark distance and plotted against target-to-landmark distance. Error bars show SE of the mean and the shaded grey line shows the standard error of the estimate. Bar chart shows the average standard deviation (horizontal variable error). Green line in figures is the line of unity. Target-to-landmark distances are the computed distances of all permutations of target to landmark positions, to the right and left of the central fixation.

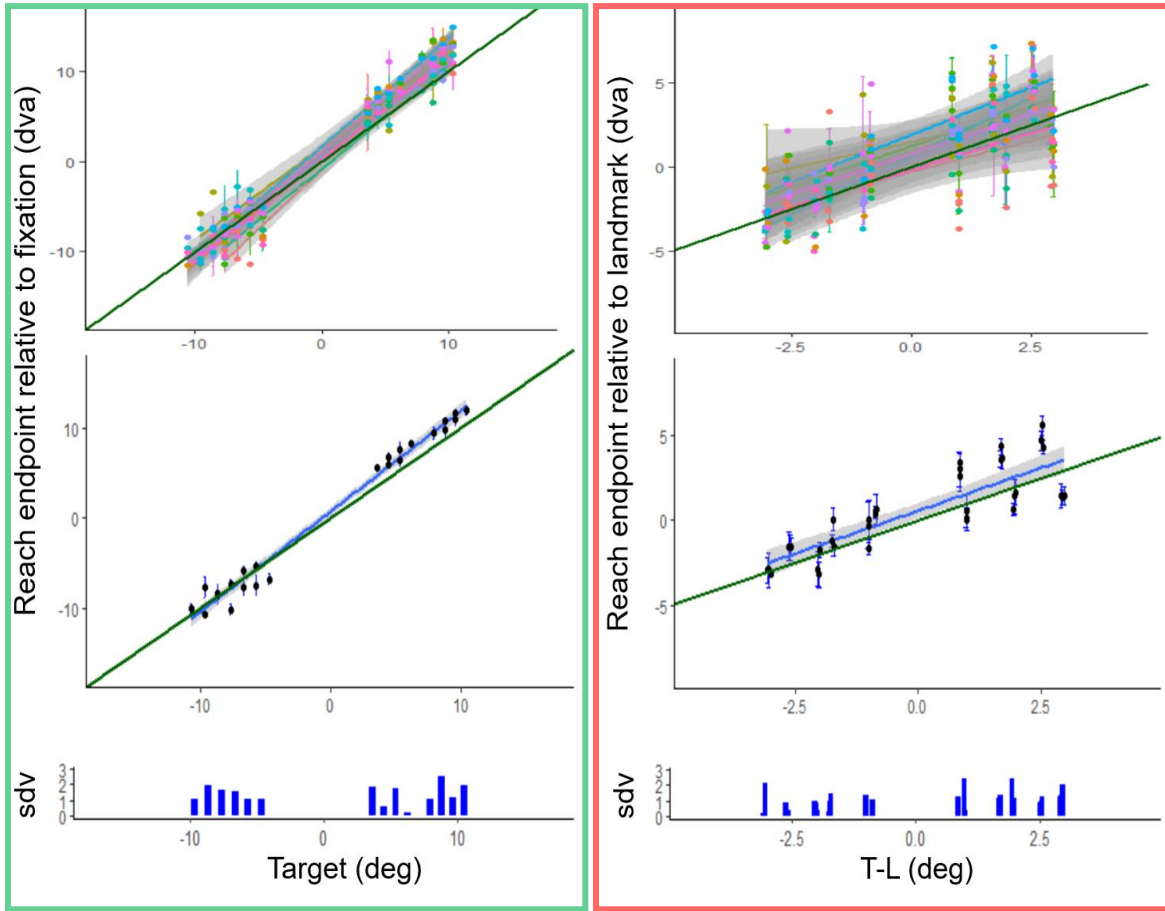


Figure 9. Regression Plots of Same-Allocentric Task. Green panel: egocentric coordinates. Red panel: allocentric coordinates. Top-green: Participants' reach endpoint relative to fixation, averaged at each target location and plotted against inferred target positions relative to fixation. Top-red: Participants' reach endpoints relative to second landmark position averaged at each target-to-landmark distance and plotted against target-to-landmark distances. Error bars show standard deviations of reach endpoints and the shaded grey outline shows the standard error of the estimate. Bottom-green: Mean reach endpoints relative to fixation plotted against inferred target to fixation distances. Bottom-red: Mean reach endpoints relative to second landmark positions plotted against target to landmark distances. Error bars show SE of the mean and the shaded grey line shows the standard error of the estimate. Bar chart shows the average standard deviation (horizontal reaching variance). Green line in figures is the line of unity. Targets to the right of the fixation or landmark are positive, whereas targets to left are negative. Inferred target positions are the target positions computed at the shifted landmark location.

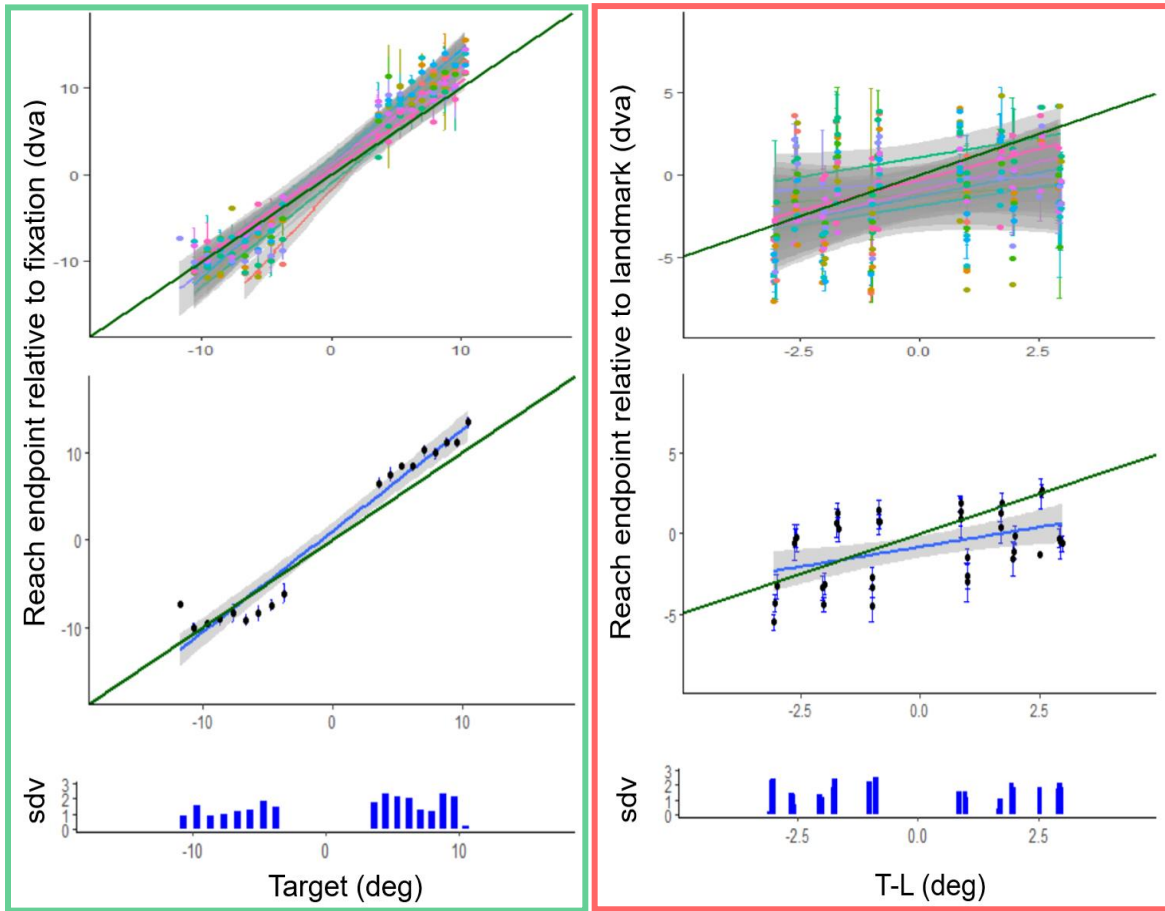


Figure 10. Regression Plots of Anti-Egocentric task. Green panel: egocentric coordinates. Red panel: allocentric coordinates. Top-green: Participants' reach endpoint relative to fixation, averaged at each target location and plotted against inferred target-to-fixation distances. Top-red: Participants' reach endpoints relative to inferred landmark position averaged at each target-to-landmark distance and plotted against target-to-landmark distances. Error bars show standard deviations of reach endpoints and the shaded grey outline shows the standard error of the estimate. Bottom-green: Mean reach endpoints relative to fixation, averaged at each target location and plotted against inferred target positions relative to fixation. Bottom-red: Mean reach endpoints relative to inferred landmark positions, averaged at each target-to-landmark distance and plotted against target-to-landmark distance. Error bars show SE of the mean and the shaded grey line shows the standard error of the estimate. Bar chart shows the average standard deviation (horizontal variable error). Green line in figures is the line of unity. Inferred target and landmark positions are at the mirror opposite location.

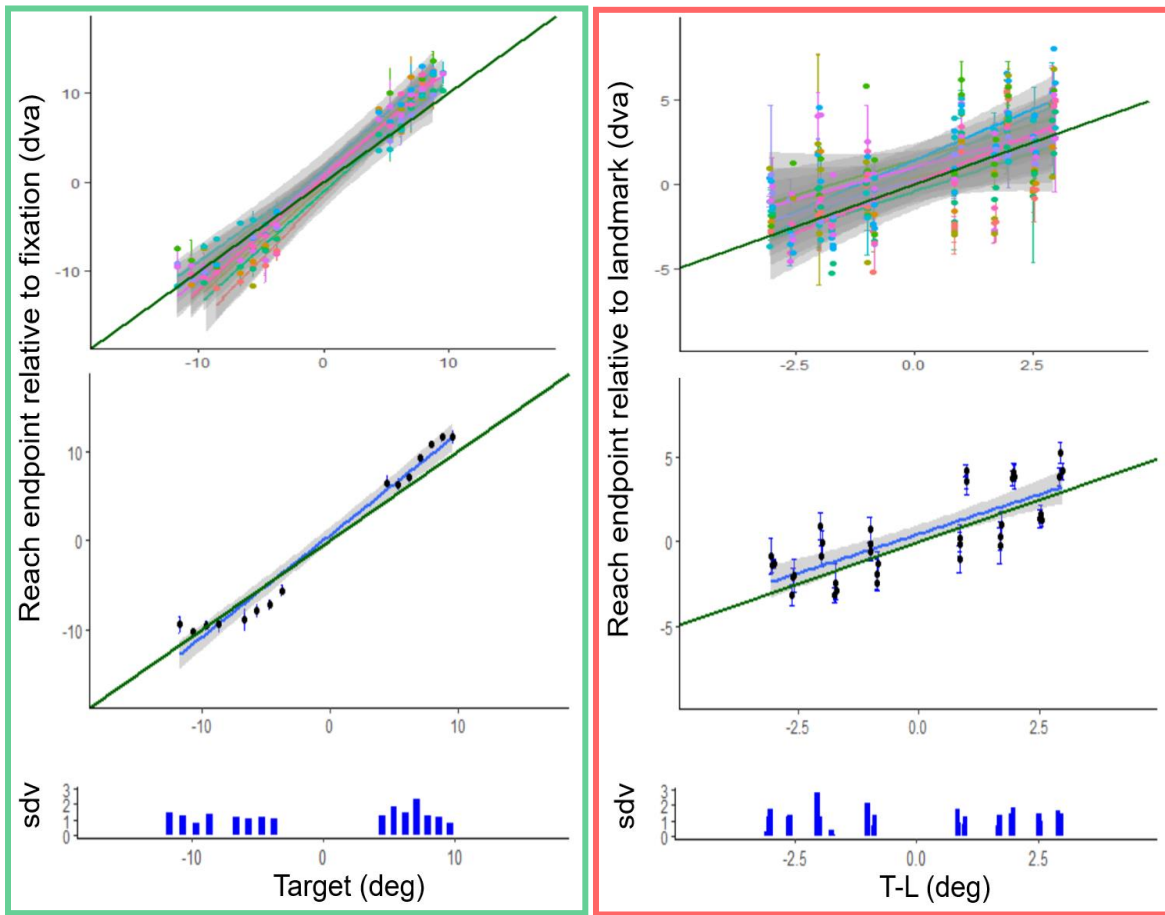
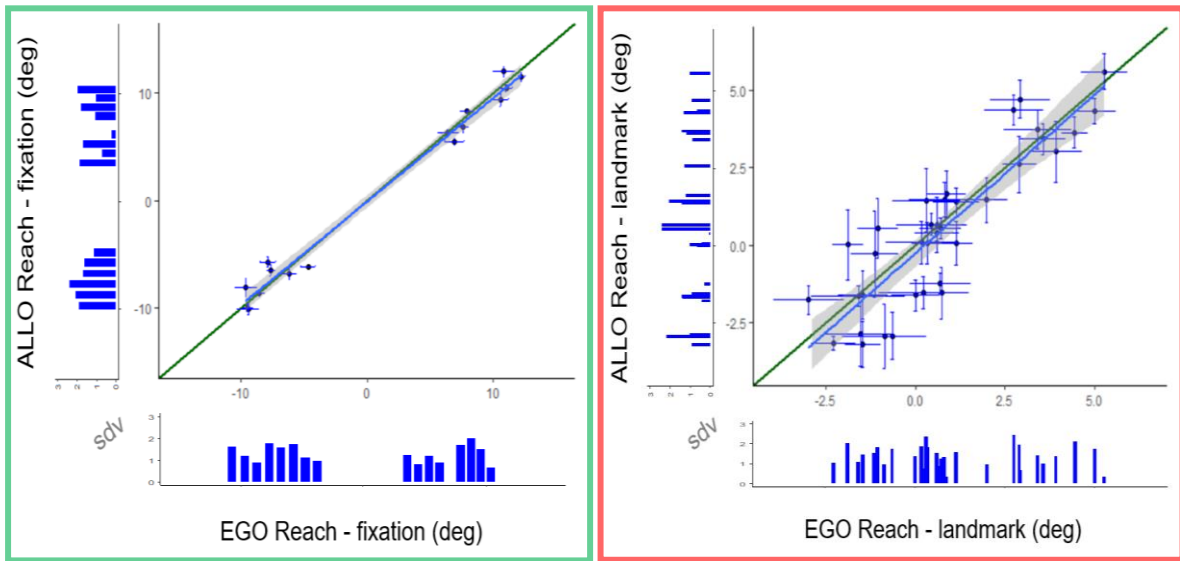


Figure 11. Regression Plots of Opposite-Allocentric task. Green panel: egocentric coordinates. Red panel: allocentric coordinates. Top-green: Participants' reach endpoint relative to fixation, averaged at each target location and plotted against inferred target positions relative to fixation. Top-red: Participants' reach endpoints relative to second landmark position averaged at each target-to-landmark distance and plotted against target-to-landmark distances. Error bars show standard deviations of reach endpoints and the shaded grey outline shows the standard error of the estimate. Bottom-green: Mean reach endpoints relative to fixation plotted against inferred target to fixation distances. Bottom-red: Mean reach endpoints relative to second landmark positions plotted against target-to-landmark distances. Error bars show SE of the mean and the shaded grey line shows the standard error of the estimate. Bar chart shows the average standard deviation (horizontal reaching variance). Green line in figures is the line of unity. Targets to the right of the fixation or landmark are positive, whereas targets to left are negative.

A



B

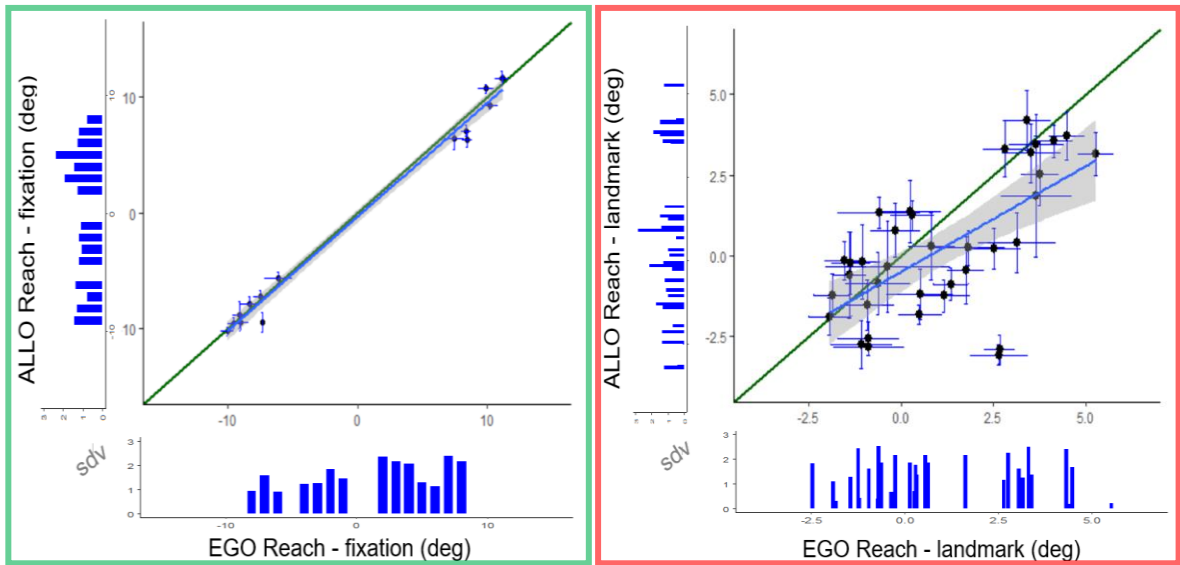


Figure 12 Allocentric vs Egocentric Reach Endpoints. Green panel: egocentric coordinates. Red panel: allocentric coordinates. **(A)** *In the same and pro conditions.* **(B)** *In the opposite and anti conditions.* Reach endpoints to the right of the fixation or landmark are positive, whereas reach endpoints to left are negative. Green line in figures is the line of unity. Error bars show SE of the mean and the shaded grey line shows the standard error of the estimate. Bar chart shows the average standard deviation (horizontal reaching variance)

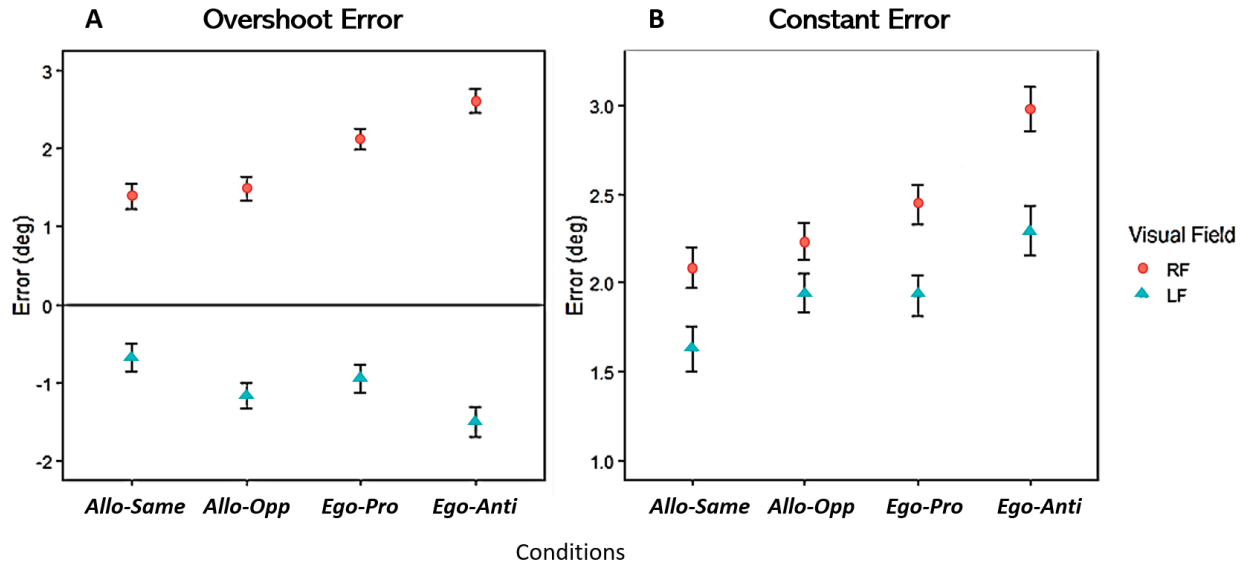


Figure 13. A. *Overshoot Error*. Mean of relative endpoint deviations: positive, to the right and negative, to the left of the expected target locations. B. *Constant error*. Mean of absolute endpoint deviations. The horizontal distance between reach endpoint and expected target location was averaged separately for left and right visual field targets. The right visual field are shown as red circles and left visual field targets are shown as blue triangles. The error bars in the plots show the SE of the mean. Egocentric and allocentric conditions are denoted by Ego/Allo in the figure legend and the spatial updating condition are indicated after the dashed line: pro/same and anti/opposite.

2.4.3 Reaction Time

Reaction time analysed by fitting an ex-Gaussian model on individual participants' data. Values for the mean, standard deviation and the tail of the distribution (tau) were obtained for each participant and averaged across all participants. The average values are reported in Table 1. 3 two way-repeated measures ANOVAs were used to look at the influence of the spatial instructions condition (egocentric/allocentric) and the spatial updating condition (same/opp) on the mean, standard deviation and tau. A strong trend was found for the main effect of the influence of the spatial instructions on the mean of the reaction time distribution, $F_{(3,39)} = 3.899$, $p = 0.057$. Post-hoc comparisons revealed that the mean reaction time in the opposite allocentric condition (value) was significantly lower than the mean reaction time in the anti egocentric condition (value), $t_{(9)} = 3.5117$, $p = 0.0023$. The same analysis was performed for movement times. However, no significant difference was observed for movement times.

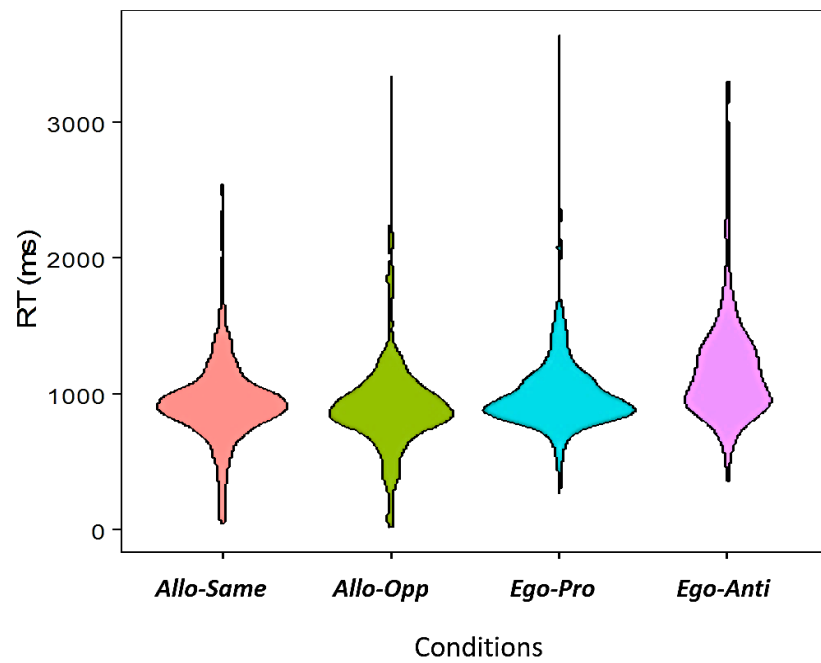


Figure 14. Violin plots of reaction times. Egocentric and allocentric conditions are denoted by Ego/Allo in the figure legend. Egocentric and allocentric conditions are denoted by Ego/Allo in the figure legend and the spatial updating condition are indicated after the dashed line: pro/same and anti/opposite.

Table 1. Summary of Ex-Gaussian Model: Average values of mean, sd and tau of reaction time

CONDITION	REACTION TIME		
	MEAN	SD	TAU
<u>EGOCENTRIC</u>			
SAME	808.20	75.81	204.50
OPPOSITE	902.25	132.18	273.90
<u>ALLOCENTRIC</u>			
SAME	738.36	138.80	216.10
OPPOSITE	743.61	140.22	208.19

2.5 Discussion

Naturally, humans use their gaze to execute goal-directed actions, but it has been previously shown that they produce accurate movements, even without visual feedback (Henriques et al., 1998; Medendorp & Crawford, 2002; Vercher et al., 1994). In such situations, humans can still generate a complex model of feedforward transformations (Blohm & Crawford, 2007). The visual environment is rife with visual landmarks, that stabilize feedforward plans by providing cues for targets to be remembered and reached. Here, we used verbal instructions and along with color to guide participants to use allocentric cues. Yet, the defining properties of a visual landmark do not need to be explicit. In the spatial navigation literature, Caduff and Timpf (2008) detail a trilateral relationship (that has also been described by others) that involves the observer, the object to be used as a landmark and the environment. Their model defines 3 concepts of salience that include bottom-up perception, top-down cognition and the environment. In our task, we've optimized all three aspects landmark saliency: the landmark is red (bottom-up), the remaining stimuli are not red (environment) and participants are instructed to utilize the red dot as a landmark (top-down). We chose red as the landmark color because research on color psychology suggests that longer-wavelength colors, like red and orange are more salient than shorter-wavelength colors like green or blue (Nakshian, 1964). Arguably, a different combination of short-wavelength and long-wavelength colors could have been used for the target and landmark, respectively. However, due to constraints in the available LED colors during experimental design we chose a red-green combination, that is also well discriminated. Trichromatic mammals like humans have an optimized subsystem for comparing the output of middle and long-wavelength sensitive cones, so they can discriminate well between red and green (Nathans, 1999).

In the present study, we showed that deliberately remembering and subsequently reaching to the position of a target relative to a visual landmark is advantageous for reaching performance. Accuracy and precision were generally better when intentionally reaching relative to an allocentric landmark, but this improvement was more apparent when participants had to reach to an updated visual field and in the right visual field, where accuracy and precision were worse overall. This effect could not be ruled out as cued recall, due to the reappearance of landmark in the response visual field right before reaching. In the egocentric task, accuracy and precision were the same, regardless of whether the second landmark was in the same or opposite visual field of response. Improved performance in the allocentric also did not seem to be biased by the utilization of an implicit egocentric strategy. Quite the opposite, if an implicit egocentric strategy were to be utilized in the allocentric task, we would expect the same outcome as a cue-combination experiment. Even when the landmark reappeared in the same visual, it was never in the exact location. Weighing of the egocentric target position would have reduced accuracy in the allocentric task and we would have expected this effect to be exacerbated when the amplitude of landmark shift was lower. Yet, we do not find a difference in accuracy associated with the amplitude of landmark shift.

2.5.1 Effect of Landmark Instruction on Reaching Precision and Accuracy

We found an increase in reaching precision when participants were instructed to reach relative to a visual landmark, that appeared during encoding and just before recall. This effect was dependant on the visual field of motor execution. Motor execution in the right visual field was more precise when target was encoded relative to a visual landmark. In the left visual field, an inconsistent pattern of precision was observed, that was not statistically significant. While participants were the most precise in the Allo-Opp task, they were the least precise in the Allo-

Same task, where had to respond on the same visual field as encoding. A follow-up analysis showed that this observation could be explained by reaction time-precision trade-off. In allocentric conditions, participants reacted faster than the egocentric conditions. Whereas reaction time in the egocentric conditions was positively associated with precision, in the allocentric conditions it was negatively associated with precision (Figure 3C, Appendix C). Although the influence of reaction time on performance in an allocentric task has not been observed in previous studies, Lemay et al. (2004) found that movement time was important in differentiating how older adults and younger adults reach relative to an allocentric stimulus. However, this effect was specific to their study design. Reaching was performed while the allocentric cues were visible. For older adults, a higher movement time was indicative of participants using static visual information to guide their movement online while younger participants relied more on the planning process and took less time. In the present study, movement time was similar in the four conditions.

In the regression analysis, we compared participants' response relative to their noisy baseline performance. While there does appear to be a difference in the Ego-Opp condition, our analysis did not capture this difference. The reason for this could be because the study was underpowered. Although we recruited 13 participants, we only analysed data from 10 participants. The number of participants analysed was marginally lower than the number of participants that were needed based on our power computations, which is 12. Additionally, unlike the regression analysis, the mixed effects analysis was more powerful, since all of the data rather than participants average performance was used in the analysis (while accounting for participants' individual differences using a random intercept model). Our power analysis was

modelled using this more robust analysis, which can also account for the reason the regression analysis was underpowered.

When analysing horizontal reaching errors, the accuracy of response differed in the allocentric and egocentric conditions. Participants tended to overshoot the target position in the egocentric tasks, to a significantly greater extent than the allocentric task. Exaggeration of horizontal retinal eccentricity while fixating and pointing to peripheral targets in complete darkness has been observed in previous studies (Henriques & Crawford, 2000; Henriques et al., 1998), however it was in situations where reaching was based solely on egocentric cues during encoding. In the present study, reaches in the egocentric instructions task were not directed to isolated visual targets, but due to the top-down instructions encoding was directed towards egocentric information only. Although fixation behaviour was not monitored during response, and participants sometimes made fixation errors while pointing, RME have been shown to exist in situations where participants reach while moving their eyes to the remembered location of the target (Van Pelt & Medendorp, 2007). For the egocentric task, exaggeration of target eccentricity was significantly higher in the visual field updating condition, than the no updating condition. Gaze-centered, egocentric remapping is a noisy process (Byrne & Crawford, 2010; Prime et al., 2006), but this has been shown in the context of making a saccade. Saccade-specific remapping of visual information involves utilizing internal signals of motion (efference copies) to update the visual field (Crawford et al., 2004; Sommer & Wurtz, 2008), whereas the current task involves remapping of the target location first and then utilizing the target remapping to direct feedforward movement plans. The two processes are not analogous, but we still anticipated remapping error in this context.

We found that pointing in an allocentric reference frame was especially useful when participants had to point to the opposite visual field. Given that the target maintained the same position relative to the landmark, we hypothesized that pointing in the same or opposite visual field should be comparable. By contrast, the egocentric condition differed based on visual field. In the same visual field condition participants, did not require an additional computation to point correctly but anti-pointing required participants to update their direction of movement. Although an additional computation was needed in the anti-pointing task, it is not clear if this renders it more difficult than the allocentric task overall. Pointing in an allocentric spatial reference frame also requires an additional computation, but to convert the allocentric spatial location of the target into an egocentric command for action. We observed improved performance in the allocentric task overall, despite the need for an additional computation. Thus, it is not clear which egocentric task (pro or anti) compares more closely to the allocentric pointing task. However, it is evident that pointing in the pro or anti conditions in the egocentric task should differ more than pointing in a same or opposite visual field in the allocentric task.

Once again, we observed hemispheric asymmetry in accuracy between the two tasks. The presence of a landmark was more important in mitigating overshoot errors due to spatial updating in the right visual field than the left visual field. The visual field of response was pertinent in this comparison because it has been shown that the effect of landmark on RME depends on the gaze at reach, rather than encoding (Byrne et al., 2010). To that extent, it appears that the beneficial effect of a landmark on RME is maintained, even after visual field updating in the right visual field. Finding an advantage for landmarks in the right visual field could imply that landmarks correct for overshoot error more in the right visual field after spatial updating or it could be that right-to-left visual field updating is associated with more RME and landmarks

correct for overshoot error to the same extent during updating in both visual fields. Yet, we did not find any significant difference between the two egocentric tasks (Ego2 – Ego1) in the right and left visual field, to conclude that right-to-left visual field updating is associated with more error. In contrast to our outcome, Byrne et al. (2010) found that RME was reduced in the presence of visual landmarks, only when gaze is directed to the right of the target. However, our task differed in many ways. Notably, our egocentric task did not have a single peripheral target and the target can appear on either side of the landmark.

2.5.2 Comparison to Previous Studies Utilizing Allocentric Landmarks

Byrne & Crawford (2010) also tested participants in an explicit allocentric task with instructions, but landmark shift was always in the same visual field. Variability of reaching was similar in the egocentric and allocentric tasks. They also observed two unique trends in response variability in the allocentric and egocentric tasks. Shifting the fixation was associated with increased variability in the egocentric task but not the allocentric task, while shifting the landmark did not influence either task. While a fixation shift was part of our task, we do not find any association between amplitude of gaze shift and variability in the egocentric or allocentric tasks. It could be that the gaze shift in their task required more complex computations and was thus associated with more error in spatial updating. In their task, there were two fixation shifts that were 2-dimensional, while our task involved only one, 1D horizontal fixation shift. We similarly do not find an association between the amplitude of the landmark shift (from its mirror opposite position in the visual field updating conditions) and variability in the allocentric tasks.

Our outcome was dissimilar to Thaler and Tod (2009), where reaching precision was found to be lowest in the allocentric task. Yet, they utilized an egocentric task that does permit an implicit allocentric strategy. Participants saw a line segment (landmark) oriented relative to their

right hand. This task condition was most analogous to allocentric-egocentric integration condition. On the other hand, some studies did find improved performance in an allocentric task as opposed to an egocentric task. Lemay et al. (2004) found that reaches based solely on allocentric information were less variable than reaches based solely on egocentric information. In fact, reaching performance in the allocentric task was similar to the allocentric-egocentric integration condition. Chen et al. (2010) found that while reaching precision significantly declined over short intervals in an egocentric reaching task but not an allocentric task.

Reaching performance was inversely related to landmark proximity (Figure 1 A & B, Appendix A). Previous studies have shown that precision of reaching is highest in close vicinity to landmark and lowest further away. In our experiment, we utilized small distances (0.86° - 2.99°) between stimuli. This decrease in performance was not associated with visual crowding during response; the second landmark disappeared prior to response, so participants did not have to reach in close proximity to it. This observed difference may be associated with decreased acuity in peripheral visual vision. Retinal receptive field sizes increase from 1° at the fovea to 6.5° at a 70° degree eccentricity (Ransom-Hogg & Spillmann, 1980). This suggests that allocentric encoding can be affected by limits on visual perception but rather than utilizing wholistic object processing or rule-based (non-spatial rules) small distances can be encoded with significant accuracy and precision.

Although the stimuli for the egocentric and allocentric task were both peripheral, the topographic distance of the target to the fovea ranged from 4.57 - 15.07 degrees to the left and right. In contrast, the topographic target to landmark distance ranged 1.15 - 3.44 degrees. Given this difference in proximity, it is conceivable that remembering landmark-relative spatial distances is an easier task with the caveat being the difficulty of having to attend to distances

peripherally rather than at the centre of fixation. Having to control for this task aspect would require using comparable distances for the egocentric task, while also controlling for the peripheral distance of the target/landmark display. Some considerations are necessary when controlling for topographic distance. Increasing the target to landmark distance to a comparable target-to-fovea distance has limited potential. For example, Aagten-Murphy and Bays (2019) found that memoranda more than 4.7 degrees from the landmark lead to a negligible benefit from allocentric information. The most reasonable control then would be using less peripheral targets and target-to-landmark displays to see if a landmark advantage still exists when these differences are removed.

2.5.3 Possible Physiological Mechanisms

The mechanisms underlying egocentric and allocentric spatial representations remain elusive, but they have been linked to Goodale and Milner's (1992) influential action-perception model. In this model, the so-called dorsal stream that involves posterior parietal regions is associated with visually guided action, whereas temporal regions in the ventral stream are associated with visual perception. The ventral stream is also thought to be associated with delayed action (Goodale et al., 2004). It is believed that egocentric and allocentric representations are discernable in a similar manner, where egocentric visuospatial information is processed in the dorsal stream and allocentric visuospatial information is processed in the ventral stream (Carey et al., 2006; Schenk, 2006). The location of a visual target in an egocentric frame of reference is transiently available for action. The dorsal stream (action system) maintains an online visual representation of a target in an egocentric frame of reference for goal-directed action, that degrades rapidly. Contrarily, allocentric information is more durable over a delay

period. The ventral stream processes visual information in allocentric frame of reference and can retain the target location for a longer duration. This difference in temporal durability has been repeatedly shown in previous studies. In the presence of a visual illusion, delays in memory-grasping have been shown to increase sensitivity to perceptual illusions (Gentilucci et al., 1996; Hu and Goodale, 2000; Westwood et al., 2001). This implies that over longer time scales allocentric information derived from the ventral system (the perceptual system), present in the pictorial illusion has a more profound influence on memory-guided action. Previous studies have in fact shown that egocentric information degrades after a memory delay of a few seconds (Goodale and Milner, 1992; Hu et al., 1999, McIntyre et al., 1997), whereas allocentric information can be retained over longer memory intervals (Hay and Redon, 2006; Krigolson and Heath, 2004; Milner and Goodale, 2006). Chen et al. (2011) directly compared the influence of a delay in a reaching task that was either involved delayed reaching in an allocentric or egocentric frame of reference. They found that while performance significantly degraded over longer delays in the egocentric task, there was lower rate of decay in the allocentric task, once again showing that allocentric information persists longer.

Since it is well known that the memory interval influences the fidelity of memoranda in an allocentric or egocentric frame of reference to a different extent, we do expect that the length of the delay influenced performance in the task. In comparison to Chen et al. (2011) our delay period (2s) was little under the medium delay (3s) that they used to test participants' reaching performance. At a medium delay period, they already saw a significant degradation in precision in the egocentric task compared to a short delay (0s). However, they did not report any significant difference in horizontal pointing accuracy. Although we do anticipate that temporal delay confers an advantage in the allocentric task in our study, we found the most prominent difference

in performance to be in horizontal accuracy. Accuracy was lower in the egocentric task regardless of visual field, although there was a larger difference in the right visual field.

Research has shown that egocentric and allocentric information is maintained in a hemispherically symmetric manner (Bremmer, 2000; Medendorp et al., 2003; Merriam et al., 2003). For targets maintained in an egocentric reference frame, making an eye movement that switches a target position from one hemisphere to the other also switches the representation one cerebral hemisphere to another (Medendorp et al., 2003; Merriam et al., 2003). Yet, as discussed earlier, the performance outcomes of due to egocentric and allocentric instructions in our task showed visual field asymmetry. In the previous literature, some studies have demonstrated right hemisphere lateralization of allocentric encoding, which would suggest left visual field dominance. Object-based allocentric information, that is initially processed in the ventral visual stream must enter the dorsal stream parietofrontal loop to influence motor action. This process is believed to be right lateralized. Neuroimaging studies have shown elevated levels of activity in the right posterior parietal cortex in mostly perceptual tasks (allocentric judgement tasks), (Galati et al., 2000; Zaehle et al. 2007). Galati et al. (2000) asked participants to judge the location of a vertical bar either relative to their own midline in the egocentric task or relative to a horizontal bar in an allocentric task. They found that both tasks activated a fronto-parietal circuitry but the allocentric task was even more lateralized to the right hemisphere. Faillenot et al. (1999) have shown a similar right parietal lateralization in an implicit allocentric pointing (and shape-matching) task. When participants were asked to point to the center of a complex object, the right parietal areas were preferentially activated. Our finding of right visual field dominance is in line with the finding of left hemisphere lateralization in the study by Chen et al. (2011), who used a similar task paradigm. Left hemisphere lateralization could be due to interactions between hand

lateralization and hemifield lateralization (Perenin & Vighetto, 1988; Rossetti et al., 2003). It could be that participants attend to their imagined hand while performing the task.

2.5.5 Conclusion

In this study, we were interested in investigating the influence of an explicit, instruction-based allocentric landmark on reaching performance. We found that when instructed to use a landmark, that was always present during encoding in an allocentric vs egocentric instructions task, participants performed better during reaching. Future studies will be needed to determine if there is an advantage of allocentric instructions when the landmark can be implicitly used during both encoding and just before motor execution. In the study that has been described, we are able to show that instructions/ explicit attention can augment implicit egocentric reaching mechanisms. This finding might bear significance to the rehabilitation outcomes of individuals that might have loss of ability to utilize implicit egocentric visuomotor transformations due to brain damage.

Chapter 3

GENERAL DISCUSSION

3.1 Summary

In this thesis, I investigated the influence of allocentric instructions on performance in a delayed pointing task. The study described in the previous chapters shows that explicit attention to an allocentric landmark improves both accuracy and precision when pointing to a remembered target. This effect was stronger when participants had to point opposite to the visual field of encoding, and when they had to point in the right visual field, whether there was spatial updating component or not. These outcomes imply that attending to an allocentric representation compliments complex implicit egocentric visuomotor transformations.

In this chapter, I will discuss the implications of the study and the contribution the study makes to the current understanding of egocentric and allocentric spatial representations in the literature. I will also comment on the limitations of the current study and suggestions for future directions.

3.2 Contributions to the Spatial Representation Literature

Previous studies investigating behavioural aspects of landmark-guided reaching have consistently found reaching and pointing to be most accurate in the presence of both egocentric and allocentric cues (Byrne et al., 2010; Hay & Redon, 2006; Redon & Hay, 2005; Krigolson & Heath, 2004). Landmarks can implicitly influence reaching, such as during memory-guided eye/hand movements in the presence of a visual background (Chakrabarty et al. 2016; Uchimura & Kitazawa, 2013; Mohrmannlendl & Fleischer, 1991) or they can explicitly influence eye/hand movements, in scenarios where the location of a target must be deliberately remembered relative to a stable landmark (Chen et al., 2011; Hay & Redon, 2006; Obhi & Goodale, 2005).

Still, the individual contributions of egocentric and allocentric representations have been of interest, despite being difficult to separate movements in an allocentric frame of reference experimentally. In the absence of allocentric cues, using solely an egocentric frame of reference, targets can be remembered reasonably well (Henriques et al., 1998; McIntyre et al., 1997; Pouget et al., 2002). Cue-conflict experiments dissociate the target position in allocentric frame of reference from an egocentric frame of reference by implicitly shifting the landmark position. The outcomes of such studies suggest that allocentric spatial cues appear to be utilized based on their reliability and subjective judgements of stability, but also based on their task utility (Byrne & Crawford, 2010; Fiehler et al., 2014). Comparison between the fidelity of egocentric spatial reference frames to allocentric reference frames is feasible when only one frame of reference can be used to remember a target of an eye/hand movement. In order to do so, psychophysical studies as such create an “allocentric only” condition, where the target position, is either rotated and/or shifted from the body midline, and can only be accurately be remembered relative to an allocentric landmark. Allocentric representations appear be better maintained over long delays but findings surrounding whether they lead to more accurate and/or precise movements than egocentric cues have been inconsistent (Chen et al., 2011; Lemay & Stelmach, 2005). Additionally, the question still remains of whether the explicit influence of allocentric landmarks differs from their implicit influence.

In our study, we attempted to draw comparisons between egocentric and allocentric spatial reference frames by the use of instructions. Instructions made it possible to disentangle the explicit influence of allocentric landmarks during encoding from their implicit influence. Participants were instructed to either use an allocentric landmark that was always present during encoding or ignore the landmark and only remember the target position in an egocentric

reference frame. Despite attending to only the target position during the egocentric task, the landmark could have had an implicit influence on target encoding. Thus, when instructed to utilize an allocentric encoding strategy (as opposed to an egocentric one), we anticipated a difference in performance based on whether explicit encoding of an allocentric spatial representation is different than implicit encoding. The outcome of our study does reveal such a difference: task performance outcomes show significantly better reaching in an allocentric encoding strategy, that was more pronounced when the task involved a spatial transformation.

3.3 Limitations

Various measures were taken in order to reduce confounds and improve the internal validity of the study. For instance, we optimized the design of our experiment to eliminate motor asymmetry when responding in either the right or left visual field. In order to do so, we aligned the centre of the LED screen midline with the shoulder bone (acromion), rather than aligning it head/centre of gaze. However, in doing so we moved the LED display further than initially intended and the size and distance between stimuli was consequently smaller. Although the distance between adjacent dots was still large enough for participants to easily distinguish between them, supplementary analysis of participants' performance revealed that their performance was slightly better when the target and landmark were further apart. Despite previous findings that allocentric spatial encoding is best in close proximity to a landmark (Aagten-Murphy & Bays, 2019; Krigolson & Heath, 2004), we found that optimizing the minimal distance between stimuli is imperative for better performance.

Another limitation in our study was the timing of stimulus presentation. In the allocentric task condition, the landmark appeared 2s before participants heard audio instructions to reach. This duration was imperative for performance since participants response in the allocentric task

must have been preceded by a visuo-motor transformations to convert the target location from allocentric coordinates to egocentric coordinates of the effector. Yet, participants' responses still show a faster reaction time in the allocentric task. While a faster reaction time was associated with higher accuracy and precision of response in the egocentric task, it was associated with worse performance in the allocentric task. This suggests that adjusting this timing of response might be an important consideration when designing future experiments.

Factors affecting the external validity included handedness and red-green colour blindness. While it was important that participants fit these criteria, it limits generalizations that could be made about visual field asymmetry. Though our findings do not demonstrate a handedness effect (performance was worse in the right visual field, which projects to the dominant left hemisphere), it is unclear if left-handed individuals would show a similar visual field effect. Additionally, we did not exclude participants for these criteria in the most ideal way. Despite excluding participants for these criteria based on a pre-study questions, we did not do any clinical tests of handedness or color blindness. For color-blindness, participants were indirectly tested by having them meet a minimum of 80% accuracy in a pre-study mock experiment. To achieve such accuracy, they must have been able to distinguish the target and landmark to point to the correct side relative to the final landmark position.

3.4 Unresolved Questions and Future Directions

This study made it possible to observe the difference between an explicit instruction-based allocentric spatial encoding, and implicit encoding in the egocentric instructions task. In our task, egocentric and allocentric information were always in conflict since the landmark always shifted to a new position. There is still an opportunity to investigate the influence of explicit instructions in a task where the allocentric cue can implicitly facilitate a correct response. If the landmark

does not shift from its original location, an allocentric spatial encoding strategy would be imperative during the memory delay and response phase, which would place more incentive on implicitly encoding the landmark position. This task would allow for a better understanding of the effect instructions and whether explicitly attending to an allocentric landmark can lead to improved performance. Had there been an experimental condition where egocentric and allocentric information are congruent, I anticipate that they would have been optimally integrated (Byrne et al., 2010) and that performance would have been better than allocentric only or egocentric only condition.

Additionally, our results demonstrate that updating across visual fields in an egocentric reference frame is a different mechanism than updating in an allocentric reference frame. The idea of whether there is an allocentric reference that is independent of an egocentric reference frame has been met some skepticism in the literature (Filimon, 2015). While we did not show that allocentric reference frames are completely viewpoint independent, cartographic map-like, we have shown that visual field updating in an allocentric reference frame is much less flawed than in an egocentric reference frame and accuracy was mostly field independent. Whereas, in an egocentric reference frame, accuracy was considerably impacted by visual field of response, in allocentric reference frame participants had a comparable accuracy in the left and right visual fields.

3.4 Applications of Cuing to Individuals with Neurological Impairments

Previous literature has shown that allocentric cuing can be helpful for individuals with optic ataxia. Optic ataxia is a disorder that results from damage to the posterior parietal cortex (PPC). In the absence of any other sensory cues, patients with optic ataxia have difficulty completing visual-guided reaching movements, while having intact stereoscopic vision, visual

fields, oculomotor control, proprioception, motor abilities and cerebellar function (Pernin & Vighetto, 1988). The deficits they exhibit lead to overshooting or undershooting the target location. Granek et al. (2013) found that a bilateral optic ataxia patient performed much better in a decoupled visuomotor control task (moving a cursor on a screen) when they had to make movements in the horizontal dimension, where the screen boundaries can be utilized as an organized allocentric cue, as compared to diagonal movements, in which the screen boundaries act as a less categorized allocentric cue. In comparison, patients with Balint's syndrome can have optic ataxia accompanied with a visual field deficit. This is often the case because of the close proximity of the occipital cortex and subcortical pathways that relay information from the retina to the occipital lobe, to the parietal cortex. Khan et al. (2013) noted that while an allocentric cue in the form of lighting was useful to a participant with unilateral optic ataxia, due to their symptom of hemianopia oculomotor exploration can limit the influence of allocentric cues in general. In this scenario, allocentric cues are no longer stable and can disappear into the blind field. For those individuals, cuing would be more useful in a comparable controlled situation that involves fixating or only single saccade from the visible field (seen) to the blind field (remembered).

Cuing has also been shown to be beneficial for individuals with neurodegenerative disorders. Motor function is a complicated process whereby the coordination of multiple sensorimotor and motor systems is needed to plan and execute a movement plan. Motor deficits are a known impairment of patients with late Alzheimer's dementia and even more recently patients with early Alzheimer's or mild cognitive impairment (Ghilardi et al. (1999). Ghilardi et al.'s (1999) findings showed that participants with dementia require continuous visual guidance to accurately perform reaching movements, which implies that they had difficulty with

feedforward movement plans. Increasing the complexity of the motor task, such as requiring participants to perform a visuomotor transformation (anti pointing), further exacerbates this deficit and increases the cognitive and sensory processing load (Tippett & Sergio, 2006). Here, we saw that pointing to an allocentric landmark in the opposite visual field was much more accurate, even for healthy participants than anti pointing. It is thus conceivable that cuing can be used to help guided complicated movements. On the other hand, individuals with Parkinson's experience deficits in different characteristics of movement performance that can benefit from cuing. Majsak et al. (1998) found that participants with Parkinsons experienced bradykinesia, not due to decreased forced production or a compensatory strategy for increasing movement accuracy. They were able to reach for a moving ball with increased speed than a stationary ball, which showed that they had difficulty internally driving their motor output. Thus, cuing can serve different purposes for improving the visuomotor behaviour of individuals with neurodegenerative disorders.

Previous literature has suggested that developmental time courses show an early predominance of egocentric encoding and followed by an understanding of allocentric spatial representations (Bullens et al., 2010; Nardini et al., 2006). Wang and Spelke (2002) have speculated that human location coding and navigation primary rely on egocentric coding, and allocentric spatial information only becomes important when individuals cannot maintain a persistent spatial relation with within the surrounding environment and must reorient themselves. Contrary to this, Pasqualotto et al. (2013) found that in comparison to the performance of congenitally blind individuals in a spatial memory task, sighted (blind-folded) and late blind participants preferred to use an allocentric reference frame. Thus, this suggests that visual experience is necessary to develop a preference for an object-based allocentric representation.

Given the marked improvement in memory-guided reaching, it's conceivable that early perceptual training could be beneficial to individuals with a gradual loss of sight or developmental disorders.

3.6 Conclusion

The study I presented in this thesis has shown that landmarks have a profound influence on our behaviour when we are instructed to utilize them. Implicitly, the extent to which landmarks are utilized to influence goal-directed actions has been previously shown to depend on heuristic judgments of reliability, like its perceived stability. This entails, that in some circumstances allocentric information judged as unreliable, will not be utilized despite being accurate. For instance, a vibrating string on a guitar that is perceived to be unreliable is still an accurate allocentric cue to adjacent cues, because it maintains the same average position. Due to the marked improvement in accuracy when utilizing a visual landmark that we have shown in this study, it may be advantageous to deliberately utilize an allocentric strategy.

REFERENCES

- Aagten-Murphy, D., & Bays, P. M. (2019). Independent working memory resources for egocentric and allocentric spatial information. *PLoS Computational Biology*, 1–20.
- Batista, A. P., Buneo, C. A., Snyder, L. H., & Andersen, R. A. (1999). Reach plans in eye-centered coordinates. *Science*, 285(5425), 257-260.
- Beurze, S. M., De Lange, F. P., Toni, I., & Medendorp, W. P. (2009). Spatial and effector processing in the human parietofrontal network for reaches and saccades. *Journal of Neurophysiology*, 101(6), 3053-3062.
- Blohm, G., & Crawford, J. D. (2007). Computations for geometrically accurate visually guided reaching in 3-D space. *Journal of Vision*, 7(5), 4-4.
- Blohm, G., Keith, G. P., & Crawford, J. D. (2009). Decoding the cortical transformations for visually guided reaching in 3D space. *Cerebral Cortex*, 19(6), 1372-1393.
- Bremmer, F. (2000). Eye position effects in macaque area V4. *Neuroreport*, 11(6), 1277-1283.
- Buneo, C. A., Jarvis, M. R., Batista, A. P., & Andersen, R. A. (2002). Direct visuomotor transformations for reaching. *Nature*, 416(6881), 632-636.
- Burgess, N. (2006). Spatial memory: how egocentric and allocentric combine. *Trends in Cognitive Sciences*, 10(12), 551–557.
- Byrne, P. A., Becker, S., & Burgess, N. (2007). Remembering the past and imagining the future: A neural model of spatial memory and imagery. *Psychological Review*, 114(2), 340–375.
- Byrne, P. A., Cappadocia, D. C., & Crawford, J. D. (2010). Interactions between gaze-centered and allocentric representations of reach target location in the presence of spatial updating. *Vision Research*, 50(24), 2661–2670.

- Byrne, P. A., & Crawford, J. D. (2010). Cue reliability and a landmark stability heuristic determine relative weighting between egocentric and allocentric visual information in memory-guided reach. *Journal of Neurophysiology*, 103(6), 3054–3069.
- Caduff, D., & Timpf, S. (2008). On the assessment of landmark salience for human navigation. *Cognitive processing*, 9(4), 249-267.
- Carey, D. P., Dijkerman, H. C., Murphy, K. J., Goodale, M. A., & Milner, A. D. (2006). Pointing to places and spaces in a patient with visual form agnosia. *Neuropsychologia*, 44(9), 1584-1594.
- Carrozzo, M., Stratta, F., McIntyre, J., & Lacquaniti, F. (2002). Cognitive allocentric representations of visual space shape pointing errors. *Experimental Brain Research*, 147(4), 426-436.
- Chakrabarty, M., Nakano, T., & Kitazawa, S. (2017). Short-latency allocentric control of saccadic eye movements. *Journal of neurophysiology*, 117(1), 376-387.
- Chen, Y., Byrne, P., & Crawford, J. D. (2011). Time course of allocentric decay, egocentric decay, and allocentric-to-egocentric conversion in memory-guided reach. *Neuropsychologia*, 49(1), 49-60.
- Chen, Y., Monaco, S., Byrne, P. A., Yan, X., Henriques, D. Y. P., & Crawford, J. D. (2014). Allocentric versus Egocentric Representation of Remembered Reach Targets in Human Cortex. *Journal of Neuroscience*, 34(37), 12515–12526.
- Chen, Y., Monaco, S., & Crawford, J. D. (2018). Neural substrates for allocentric-to-egocentric conversion of remembered reach targets in humans. *European Journal of Neuroscience*, 47(8), 901-917.

- Clark, N., Binsted, G., Heath, M., & Krigolson, O. (2007). The proximity of visual landmarks impacts reaching performance. *Spatial vision*, 20(4), 317-336.
- Connolly, J. D., Andersen, R. A., & Goodale, M. A. (2003). FMRI evidence for a 'parietal reach region' in the human brain. *Experimental brain research*, 153(2), 140-145.
- Crawford, J. D., Medendorp, W. P., & Marotta, J. J. (2004). Spatial Transformations for Eye-Hand Coordination. *Journal of Neurophysiology*.
- Darling, W. G., Seitz, R. J., Peltier, S., Tellmann, L., & Butler, A. J. (2007). Visual cortex activation in kinesthetic guidance of reaching. *Experimental brain research*, 179(4), 607-619.
- Desmurget, M., Pélisson, D., Rossetti, Y., & Prablanc, C. (1998). From eye to hand: planning goal-directed movements. *Neuroscience & Biobehavioral Reviews*, 22(6), 761-788.
- DeSouza, J. F., Dukelow, S. P., Gati, J. S., Menon, R. S., Andersen, R. A., & Vilis, T. (2000). Eye position signal modulates a human parietal pointing region during memory-guided movements. *Journal of Neuroscience*, 20(15), 5835-5840.
- Faillenot, I., Decety, J., & Jeannerod, M. (1999). Human brain activity related to the perception of spatial features of objects. *Neuroimage*, 10(2), 114-124.
- Fiehler, K., Schütz, I., & Henriques, D. Y. (2011). Gaze-centered spatial updating of reach targets across different memory delays. *Vision Research*, 51(8), 890-897.
- Fiehler, K., Wolf, C., Klinghammer, M., & Blohm, G. (2014). Integration of egocentric and allocentric information during memory-guided reaching to images of a natural environment. *Frontiers in Human Neuroscience*, 8, 636.

- Filimon, F. (2015). Are all spatial reference frames egocentric? Reinterpreting evidence for allocentric, object-centered, or world-centered reference frames. *Frontiers in human neuroscience*, 9, 648.
- Fink, G. R., Marshall, J. C., Shah, N. J., Weiss, P. H., Halligan, P. W., Grosse-Ruyken, M., ... & Freund, H. J. (2000). Line bisection judgments implicate right parietal cortex and cerebellum as assessed by fMRI. *Neurology*, 54(6), 1324-1331.
- Flanders, M., Tillery, S. I. H., & Soechting, J. F. (1992). Early stages in a sensorimotor transformation. *Behavioral and Brain Sciences*, 15(2), 309-320.
- Galati, G., Lobel, E., Vallar, G., Berthoz, A., Pizzamiglio, L., & Le Bihan, D. (2000). The neural basis of egocentric and allocentric coding of space in humans: a functional magnetic resonance study. *Experimental brain research*, 133(2), 156-164.
- Gentilucci, M., Chieffi, S., Daprati, E., Saetti, M. C., & Toni, I. (1996). Visual illusion and action. *Neuropsychologia*, 34(5), 369-376.
- Ghilardi, M. F., Alberoni, M., Marelli, S., Rossi, M., Franceschi, M., Ghez, C., & Fazio, F. (1999). Impaired movement control in Alzheimer's disease. *Neuroscience letters*, 260(1), 45-48.
- Goodale, M. A., Jakobson, L. S., & Keillor, J. M. (1994). Differences in the visual control of pantomimed and natural grasping movements. *Neuropsychologia*, 32(10), 1159-1178.
- Goodale, M. A., & Milner, A. D. (1992). Separate visual pathways for perception and action. *Trends in neurosciences*, 15(1), 20-25.
- Goodale, M. A., Milner, A. D., Jakobson, L. S., & Carey, D. P. (1991). A neurological dissociation between perceiving objects and grasping them. *Nature*, 349(6305), 154-156.

- Goodale, M. A., Westwood, D. A., & Milner, A. D. (2004). Two distinct modes of control for object-directed action. *Progress in brain research*, 144, 131-144.
- Granek, J. A., Pisella, L., Stemmer, J., Vigliani, A., Rossetti, Y., & Sergio, L. E. (2013). Decoupled visually-guided reaching in optic ataxia: differences in motor control between canonical and non-canonical orientations in space. *PLoS One*, 8(12), e86138.
- Green, P., & MacLeod, C. J. (2016). SIMR: an R package for power analysis of generalized linear mixed models by simulation. *Methods in Ecology and Evolution*, 7(4), 493-498.
- Hay, L., & Redon, C. (2006). Response delay and spatial representation in pointing movements. *Neuroscience letters*, 408(3), 194-198.
- Henriques, D. Y. P., & Crawford, J. D. (2000). Direction-dependent distortions of retinocentric space in the visuomotor transformation for pointing. *Experimental brain research*, 132(2), 179-194.
- Henriques, D. Y., Klier, E. M., Smith, M. A., Lowy, D., & Crawford, J. D. (1998). Gaze-centered remapping of remembered visual space in an open-loop pointing task. *Journal of Neuroscience*, 18(4), 1583-1594.
- Howard, I. P., & Templeton, W. B. (1966). Human spatial orientation.
- Hu, Y., Eagleson, R., & Goodale, M. A. (1999). The effects of delay on the kinematics of grasping. *Experimental Brain Research*, 126(1), 109-116.
- Hu, Y., & Goodale, M. A. (2000). Grasping after a delay shifts size-scaling from absolute to relative metrics. *Journal of Cognitive Neuroscience*, 12(5), 856-868.
- Jakobson, L. S., & Goodale, M. A. (1991). Factors affecting higher-order movement planning: a kinematic analysis of human prehension. *Experimental brain research*, 86(1), 199-208.

- Keller, I., Schindler, I., Kerkhoff, G., Von Rosen, F., & Golz, D. (2005). Visuospatial neglect in near and far space: dissociation between line bisection and letter cancellation. *Neuropsychologia*, 43(5), 724-731.
- Khan, A. Z., Crawford, J. D., Blohm, G., Urquizar, C., Rossetti, Y., & Pisella, L. (2007). Influence of initial hand and target position on reach errors in optic ataxic and normal subjects. *Journal of vision*, 7(5), 8-8.
- Khan, A., Pisella, L., Delporte, L., Rode, G., & Rossetti, Y. (2013). Testing for optic ataxia in a blind field. *Frontiers in human neuroscience*, 7, 399.
- Klier, E. M., Wang, H., & Crawford, J. D. (2001). The superior colliculus encodes gaze commands in retinal coordinates. *Nature neuroscience*, 4(6), 627-632.
- Krigolson, O., & Heath, M. (2004). Background visual cues and memory-guided reaching. *Human movement science*, 23(6), 861-877.
- Majsak, M. J., Kaminski, T., Gentile, A. M., & Flanagan, J. R. (1998). The reaching movements of patients with Parkinson's disease under self-determined maximal speed and visually cued conditions. *Brain: a journal of neurology*, 121(4), 755-766.
- McIntyre, J., Stratta, F., & Lacquaniti, F. (1998). Short-term memory for reaching to visual targets: psychophysical evidence for body-centered reference frames. *Journal of Neuroscience*, 18(20), 8423-8435.
- Medendorp, W. P., Goltz, H. C., Vilis, T., & Crawford, J. D. (2003). Gaze-centered updating of visual space in human parietal cortex. *Journal of Neuroscience*, 23(15), 6209–6214.
- Merriam, E. P., Genovese, C. R., & Colby, C. L. (2003). Spatial updating in human parietal cortex. *Neuron*, 39(2), 361-373.

- Lemay, M., Bertram, C. P., & Stelmach, G. E. (2004). Pointing to an allocentric and egocentric remembered target. *Motor control*, 8(1), 16-32.
- Lemay, M., Gagnon, S., & Proteau, L. (2004). Manual pointing to remembered targets... but also in a remembered visual context. *Acta psychologica*, 117(2), 139-153.
- Lemay, M., & Stelmach, G. E. (2005). Multiple frames of reference for pointing to a remembered target. *Experimental Brain Research*, 164(3), 301-310.
- McIntyre, J., Stratta, F., & Lacquaniti, F. (1997). Viewer-centered frame of reference for pointing to memorized targets in three-dimensional space. *Journal of neurophysiology*, 78(3), 1601-1618.
- Medendorp, P. W., & Crawford, D. J. (2002). Visuospatial updating of reaching targets in near and far space. *Neuroreport*, 13(5), 633-636.
- Medendorp, W. P., Goltz, H. C., & Vilis, T. (2005). Remapping the remembered target location for anti-saccades in human posterior parietal cortex. *Journal of Neurophysiology*, 94(1), 734-740.
- Milner, D., & Goodale, M. (2006). *The visual brain in action* (Vol. 27). OUP Oxford.
- Nakshian, J. S. (1964). The effects of red and green surroundings on behavior. *The Journal of General Psychology*, 70(1), 143-161.
- Nardini, M., Burgess, N., Breckenridge, K., & Atkinson, J. (2006). Differential developmental trajectories for egocentric, environmental and intrinsic frames of reference in spatial memory. *Cognition*, 101(1), 153-172.
- Nathans, J. (1999). The evolution and physiology of human color vision: insights from molecular genetic studies of visual pigments. *Neuron*, 24(2), 299-312.

- Mohrmann-Lendla, H., & Fleischer, A. G. (1991). The effect of a moving background on aimed hand movements. *Ergonomics*, 34(3), 353-364.
- Obhi, S. S., & Goodale, M. A. (2005). The effects of landmarks on the performance of delayed and real-time pointing movements. *Experimental Brain Research*, 167(3), 335–344.
- Pasqualotto, A., Finucane, C. M., & Newell, F. N. (2013). Ambient visual information confers a context-specific, long-term benefit on memory for haptic scenes. *Cognition*, 128(3), 363-379.
- Perenin, M. T., & Vighetto, A. (1988). Optic ataxia: A specific disruption in visuomotor mechanisms: I. Different aspects of the deficit in reaching for objects. *Brain*, 111(3), 643-674.
- Pouget, A., & Snyder, L. H. (2000). Computational approaches to sensorimotor transformations. *Nature neuroscience*, 3(11), 1192-1198.
- Prime, S. L., Niemeier, M., & Crawford, J. D. (2006). Transsaccadic integration of visual features in a line intersection task. *Experimental Brain Research*, 169(4), 532–548.
- Redon, C., & Hay, L. (2005). Role of visual context and oculomotor conditions in pointing accuracy. *Neuroreport*, 16(18), 2065-2067.
- Rossetti, Y., Pisella, L., & Vighetto, A. (2003). Optic ataxia revisited. *Experimental Brain Research*, 153(2), 171-179.
- Schenk, T. (2006). An allocentric rather than perceptual deficit in patient DF. *Nature neuroscience*, 9(11), 1369-1370.
- Scott, S. H., Gribble, P. L., Graham, K. M., & Cabel, D. W. (2001). Dissociation between hand motion and population vectors from neural activity in motor cortex. *Nature*, 413(6852), 161-165.

- Scott, S. H., & Kalaska, J. F. (1997). Reaching movements with similar hand paths but different arm orientations. I. Activity of individual cells in motor cortex. *Journal of neurophysiology*, 77(2), 826-852.
- Shadmehr, R., Wise, S. P., & Wise, S. P. (2005). *The computational neurobiology of reaching and pointing: a foundation for motor learning*. MIT press.
- Sommer, M. A., & Wurtz, R. H. (2008). Brain circuits for the internal monitoring of movements. *Annu. Rev. Neurosci.*, 31, 317-338.
- Soechting, J. F., & Flanders, M. (1992). Moving in three-dimensional space: frames of reference, vectors, and coordinate systems. *Annual review of neuroscience*, 15(1), 167-191.
- Sparks, D. L. (1989). The neural encoding of the location of targets for saccadic eye movements. *Journal of Experimental Biology*, 146(1), 195-207.
- Toni, I., Gentilucci, M., Jeannerod, M., & Decety, J. (1996). Differential influence of the visual framework on end point accuracy and trajectory specification of arm movements. *Experimental Brain Research*, 111(3), 447-454.
- Thaler, L., & Goodale, M. A. (2011). The role of online visual feedback for the control of target-directed and allocentric hand movements. *Journal of neurophysiology*, 105(2), 846-859.
- Thaler, L., & Todd, J. T. (2009). The use of head/eye-centered, hand-centered and allocentric representations for visually guided hand movements and perceptual judgments. *Neuropsychologia*, 47(5), 1227-1244.
- Uchimura, M., & Kitazawa, S. (2013). Cancelling prism adaptation by a shift of background: a novel utility of allocentric coordinates for extracting motor errors. *Journal of Neuroscience*, 33(17), 7595-7602.

- Van Pelt, S., & Medendorp, W. P. (2007). Gaze-centered updating of remembered visual space during active whole-body translations. *Journal of neurophysiology*, 97(2), 1209-1220.
- Velay, J. L., & Beaubaton, D. (1986). Influence of visual context on pointing movement accuracy. *Cahiers de Psychologie Cognitive/Current Psychology of Cognition*.
- Vercher, J. L., Magenes, G., Prablanc, C., & Gauthier, G. M. (1994). Eye-head-hand coordination in pointing at visual targets: spatial and temporal analysis. *Experimental brain research*, 99(3), 507-523.
- Vesia, M., Yan, X., Henriques, D. Y., Sergio, L. E., & Crawford, J. D. (2008). Transcranial magnetic stimulation over human dorsal-lateral posterior parietal cortex disrupts integration of hand position signals into the reach plan. *Journal of Neurophysiology*, 100(4), 2005-2014.
- Vindras, P., & Viviani, P. (1998). Frames of reference and control parameters in visuomanual pointing. *Journal of Experimental Psychology: Human Perception and Performance*, 24(2), 569.
- Vogeley, K., & Fink, G. R. (2003). Neural correlates of the first-person-perspective. *Trends in cognitive sciences*, 7(1), 38-42.
- Wang, R. F., & Spelke, E. S. (2002). Human spatial representation: Insights from animals. *Trends in cognitive sciences*, 6(9), 376-382.
- Westwood, D. A., Heath, M., & Roy, E. A. (2000). The effect of a pictorial illusion on closed-loop and open-loop prehension. *Experimental Brain Research*, 134(4), 456-463.
- Westwood, D. A., McEachern, T., & Roy, E. A. (2001). Delayed grasping of a Müller-Lyer figure. *Experimental brain research*, 141(2), 166-173.

Zaehle, T., Jordan, K., Wüstenberg, T., Baudewig, J., Dechent, P., & Mast, F. W. (2007). The neural basis of the egocentric and allocentric spatial frame of reference. *Brain research, 1137*, 92-103.

APPENDICES

Appendix A: Supplementary Figure 1

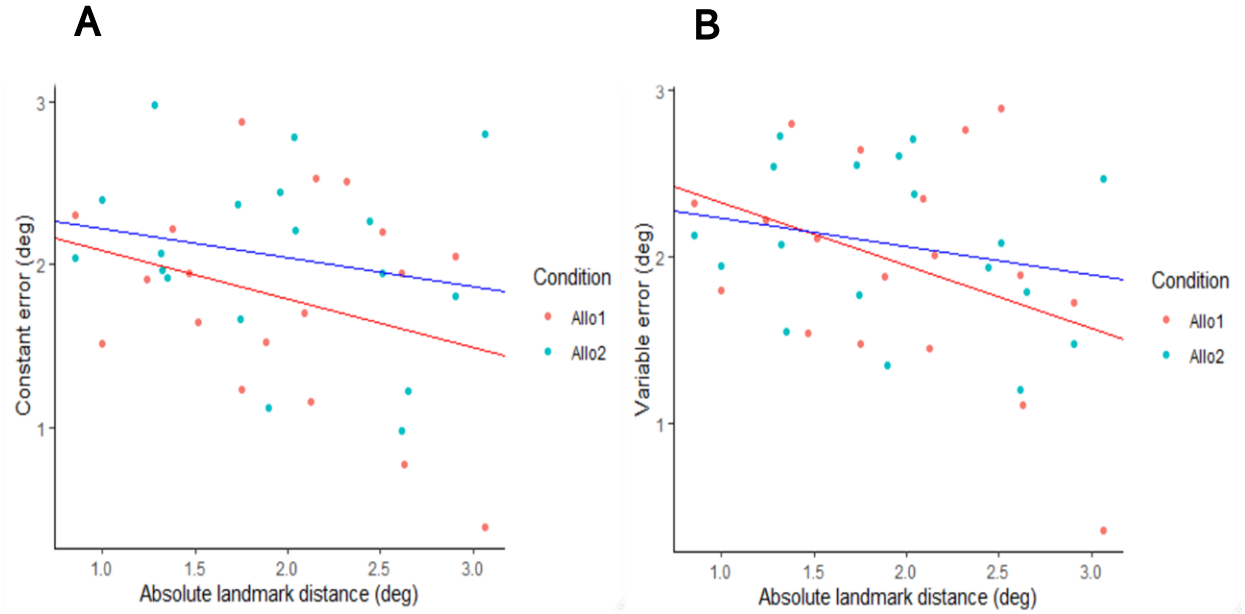


Figure 1. Regression plots of **(A)** *Constant error (accuracy) plotted as function of target to landmark distance.* **(B)** *Variable error in the horizontal direction (average standard deviation) plotted as function of target to landmark distance.* The allocentric condition with no spatial updating is shown in red and the condition with spatial updating is shown in blue.

Appendix B: Supplementary Figure 2

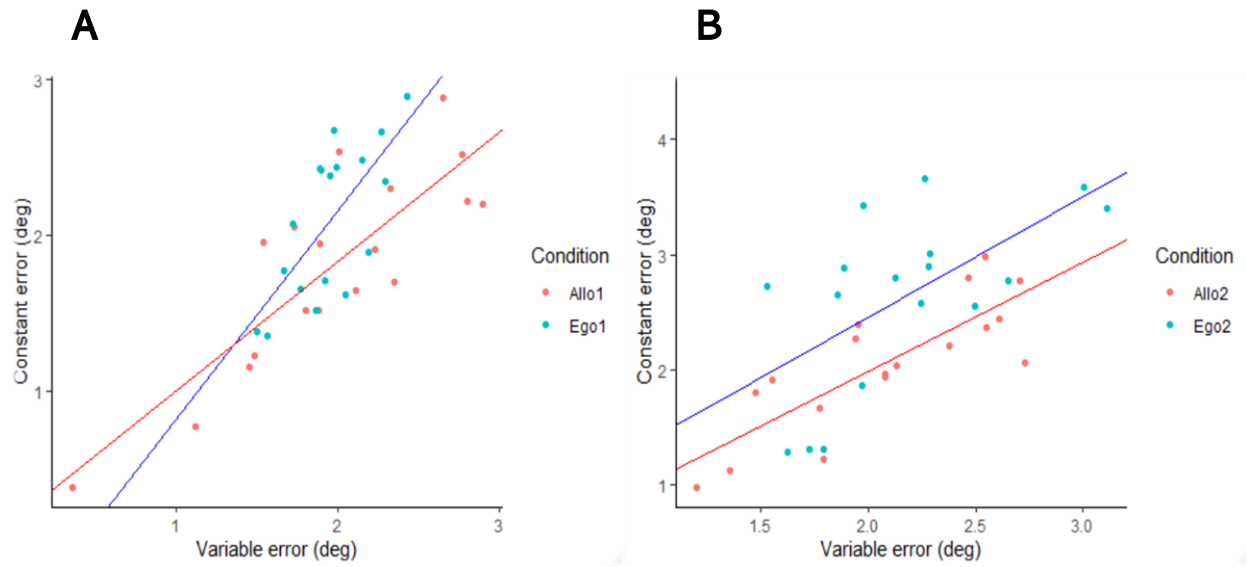


Figure 2. Regression plots of the constant error plotted as function of the variable error **(A)** For the allocentric and egocentric condition with no spatial updating. **(B)** For the allocentric and egocentric condition with spatial updating. Regression lines of the allocentric conditions are shown in red, while the regression plots of the egocentric conditions are shown in blue.

Appendix C: Supplementary Figure 3

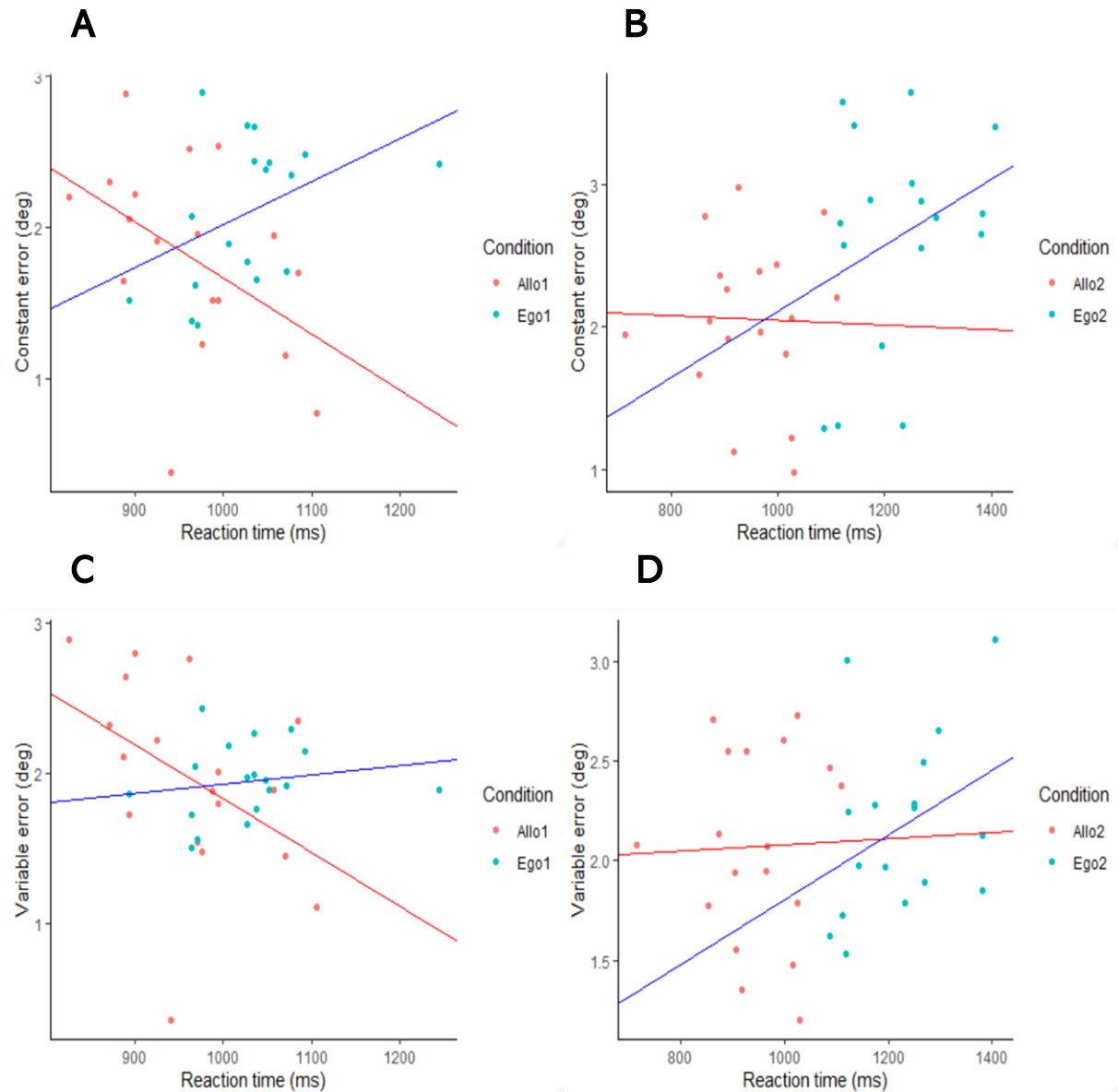


Figure 3. Regression plots of **(A,B)** Constant error (accuracy) plotted as function of reaction time. **(C,D)** Variable error in the horizontal direction (average standard deviation) plotted as function of reaction time. Allocentric conditions are shown in red, and egocentric conditions in blue. Numbers in the figure legends represent spatial updating conditions (1) and no spatial updating (2).

Appendix D: Author Contributions

Lina Musa contributed to experimental design, laboratory preparation, data collection and analysis and wrote the paper. Dr. Yan contributed to laboratory preparation and provided technical support. Dr. Crawford contributed to experimental design, provided advice and feedback on experiments and analyses, editorial comments, and funding support.