

Bayesian Decision Theory

Perception

G89.2223

Larry Maloney
New York University

Every life is a series of decisions



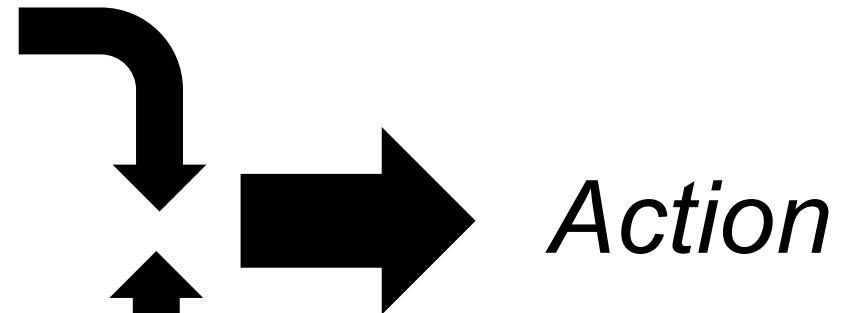
Anatomy of a Decision

Feasibility

Probability, Frequency,
Uncertainty, etc.

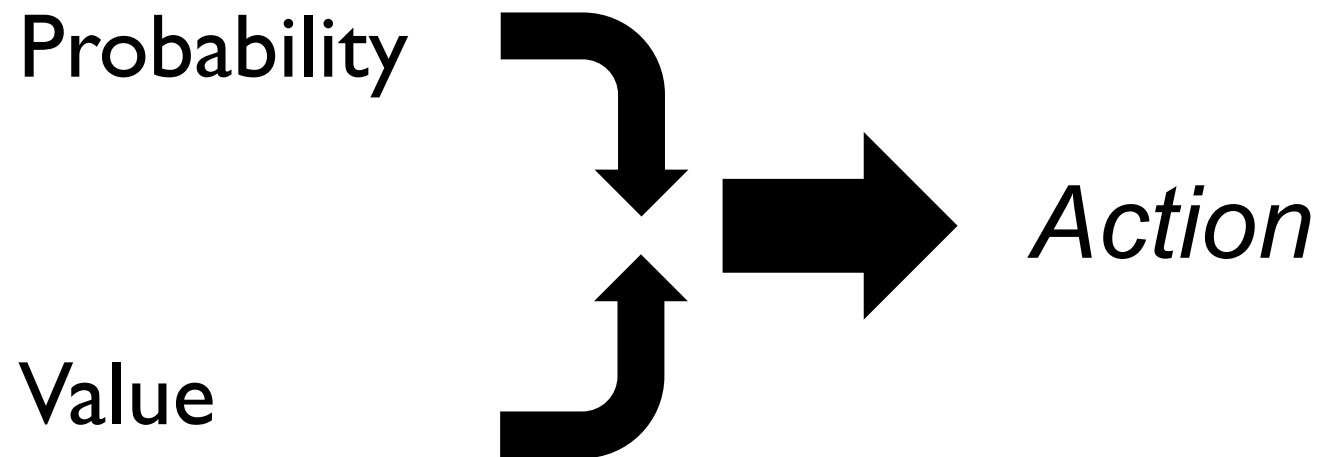
Desireability

Reward, Utility, etc



*Biology
Psychology
Economics
Political Science
Neuroscience*

Stochastic Form



Statistical Decision Theory



Abraham Wald



John von Neumann



David Blackwell



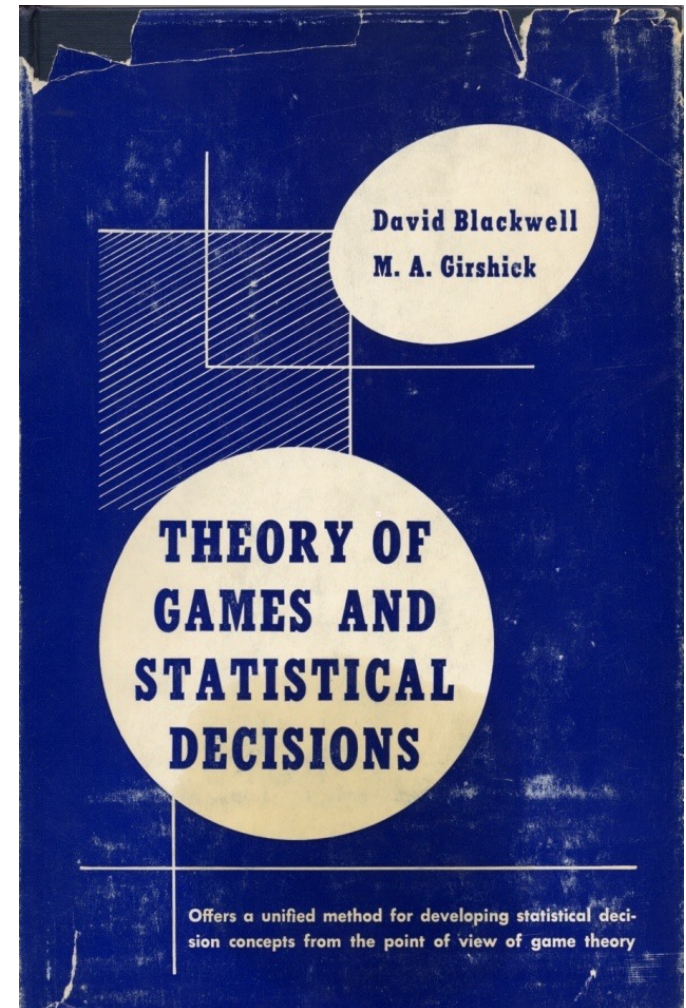
Frank P Ramsey



O Morgenstern



M. A. Girschick



1954

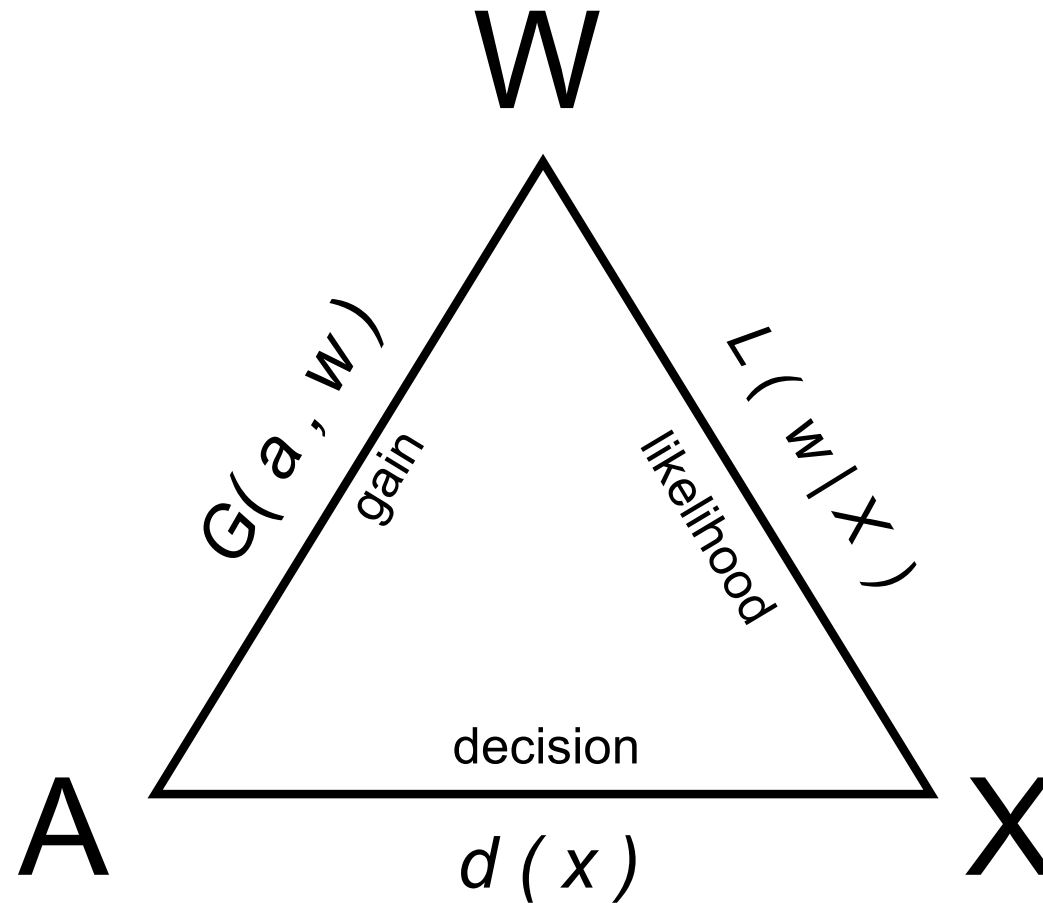
Three Elements of SDT

$W = \{w_1, w_2, \dots, w_m\}$ *possible states of the world*

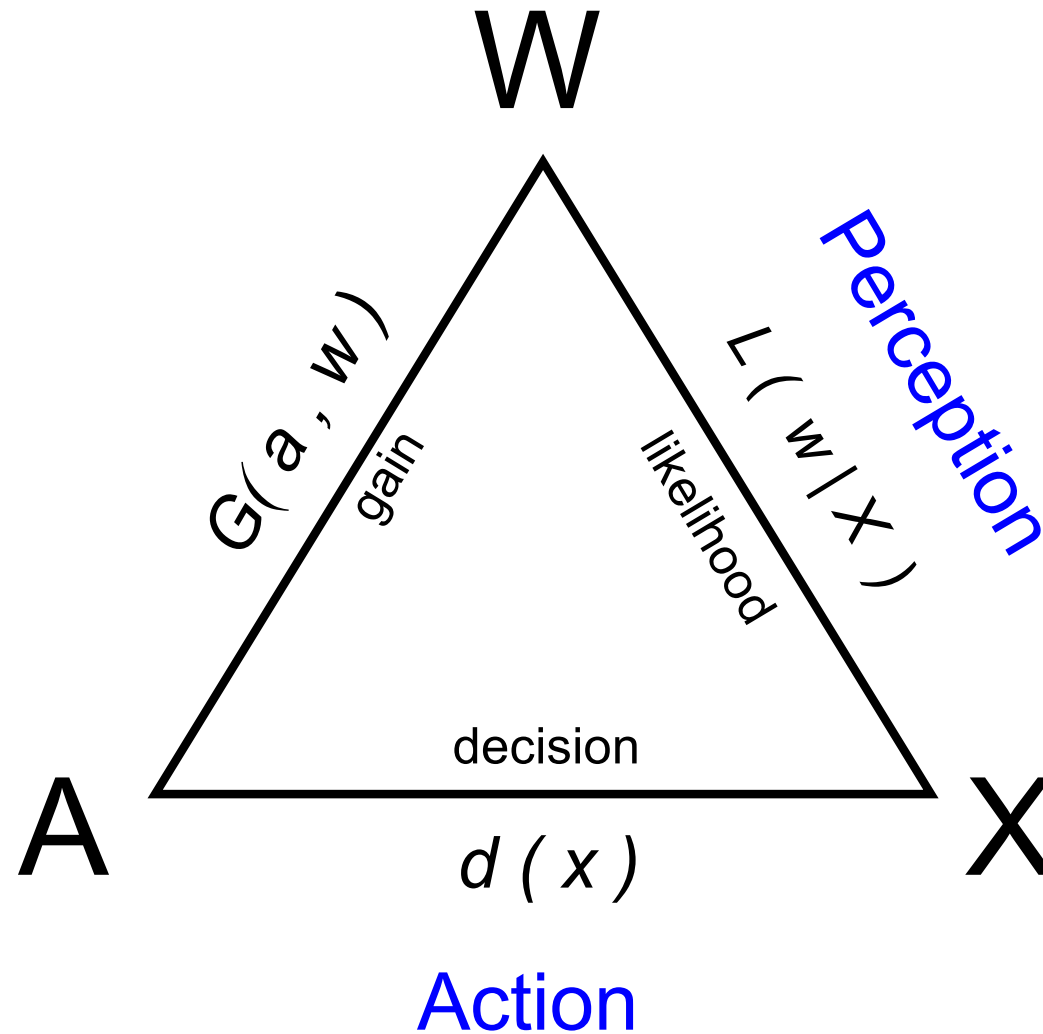
$A = \{a_1, a_2, \dots, a_p\}$ *possible actions*

$X = \{x_1, x_2, \dots, x_n\}$ *possible sensory events*

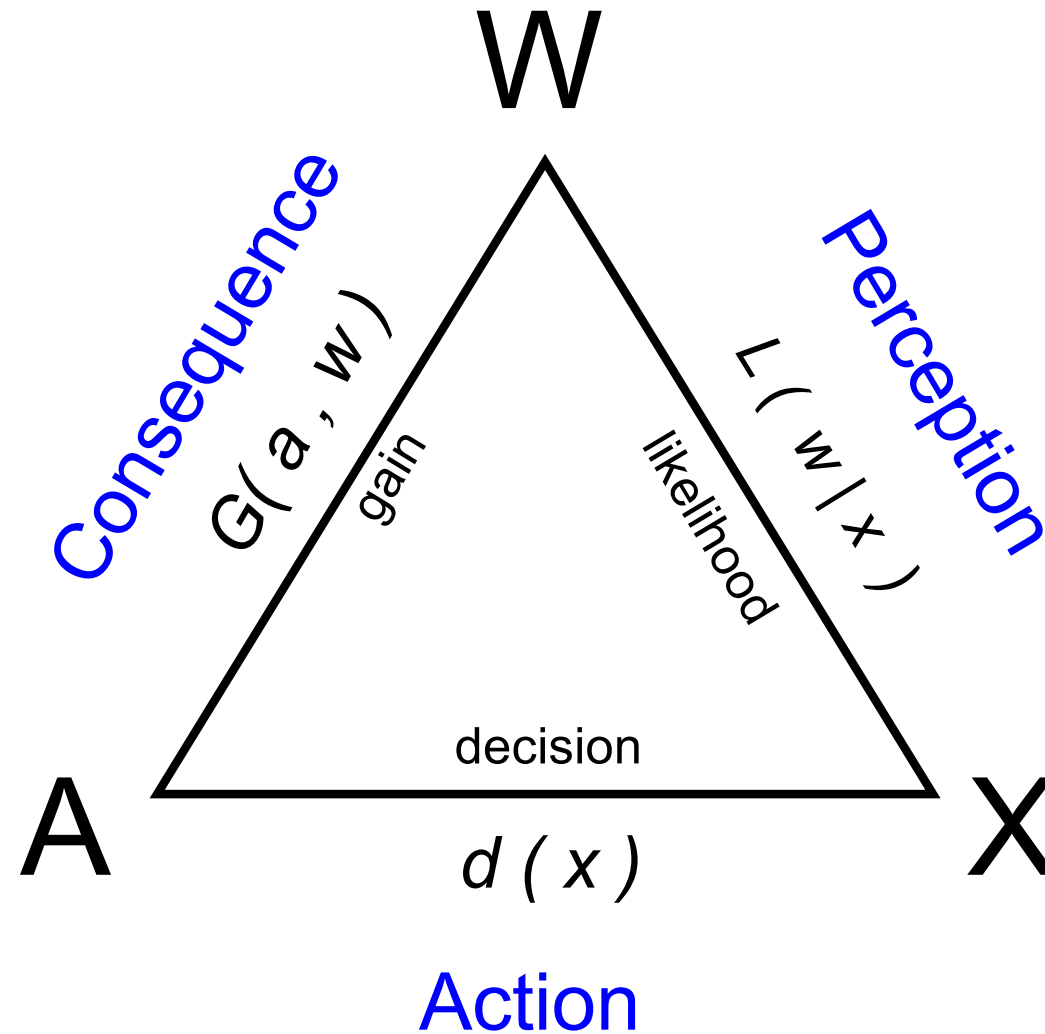
Three Functions of SDT



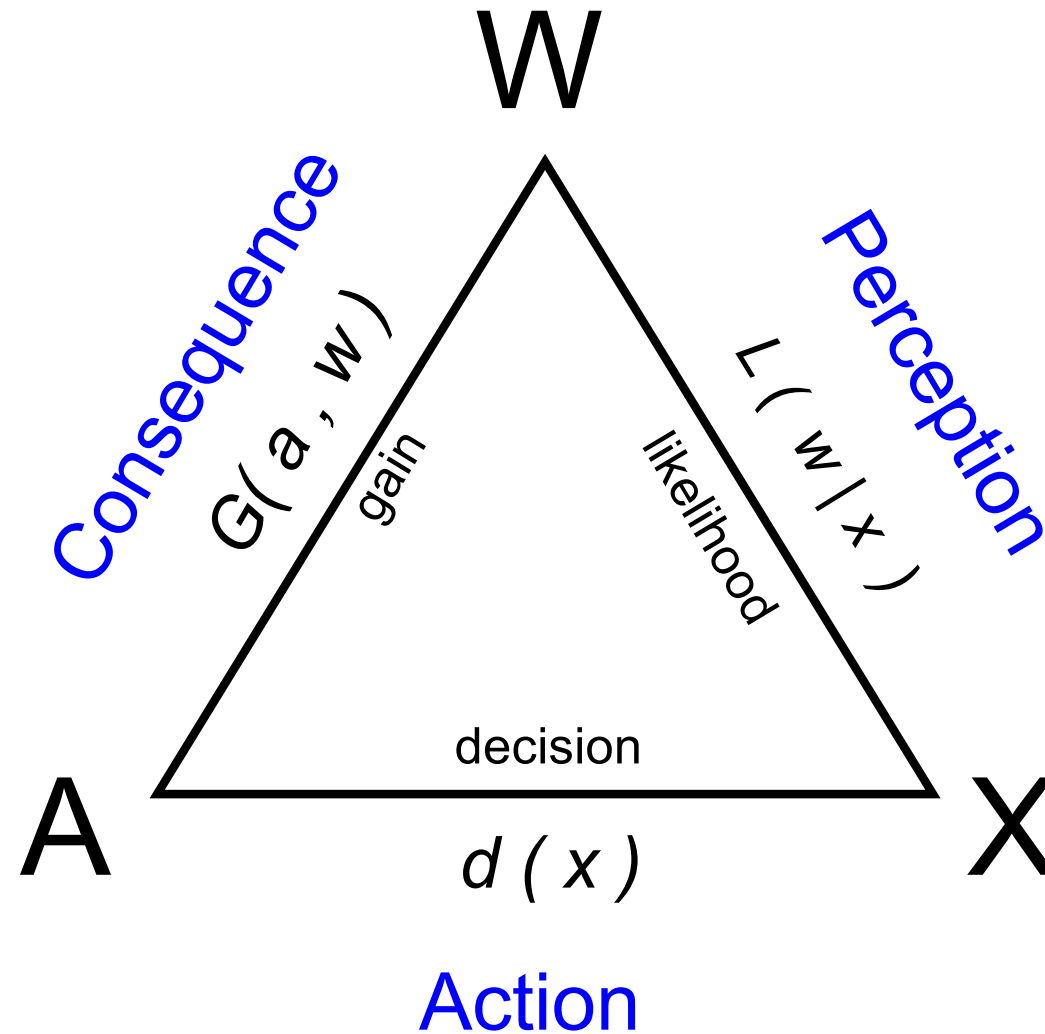
Three Functions of SDT



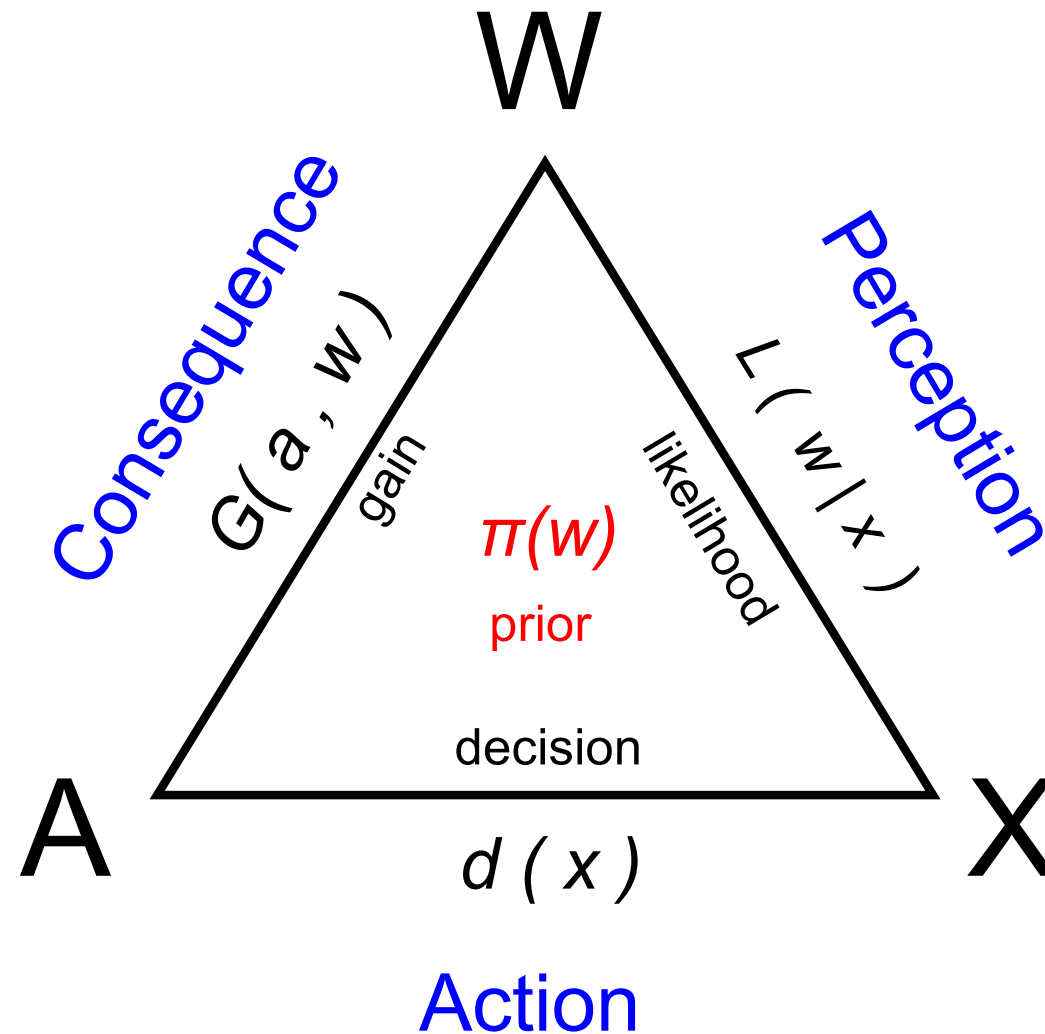
Statistical Decision Theory



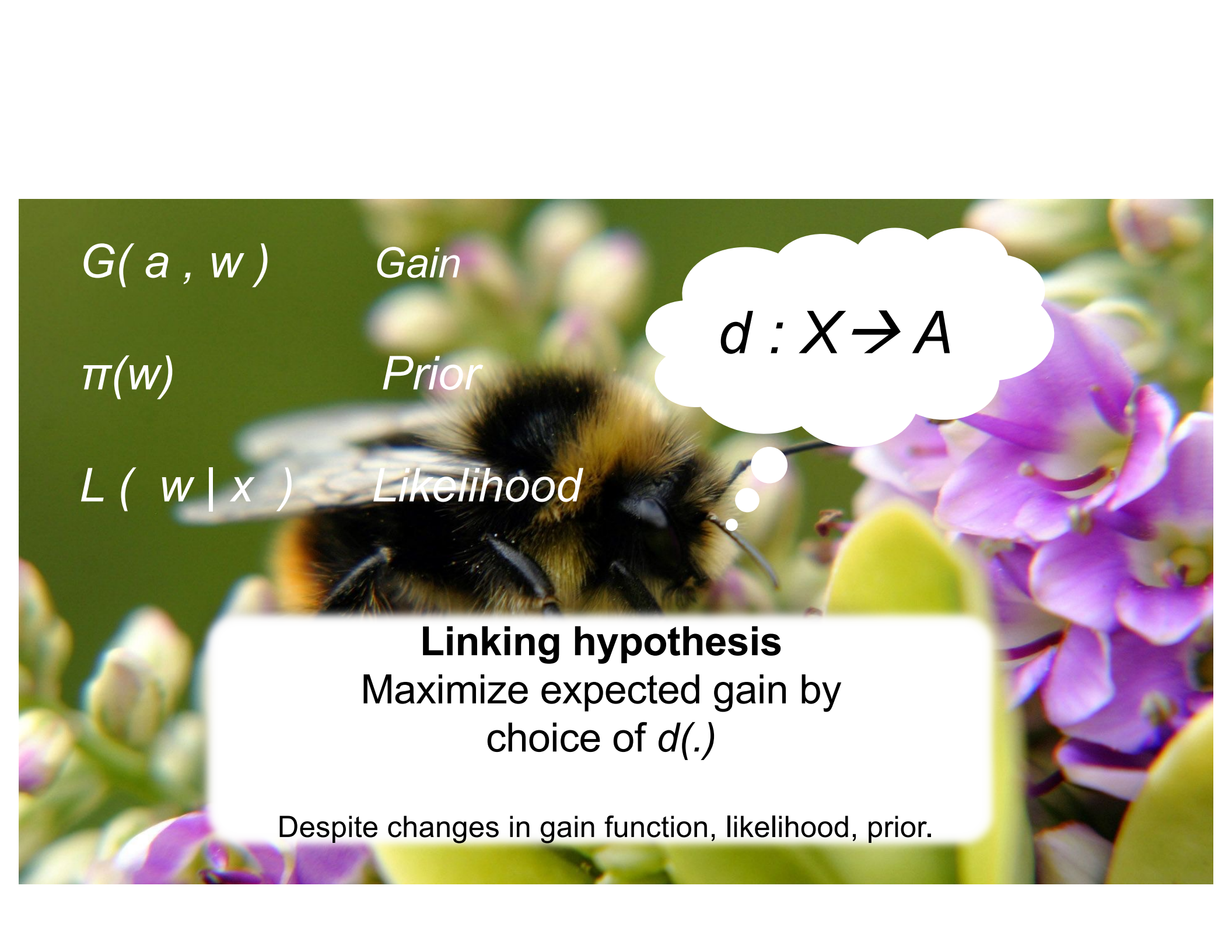
Statistical Decision Theory



Bayesian Decision Theory



Computation

A bumblebee is shown in profile, facing right, on a purple flower. The background is a soft-focus green. Overlaid on the image are mathematical symbols and text. A thought bubble above the bee contains the expression $d : X \rightarrow A$. To the left of the bee, three mathematical expressions are listed vertically: $G(a, w)$, $\pi(w)$, and $L(w | x)$. To the right of each expression is a word: Gain, Prior, and Likelihood. At the bottom, a white box contains the text 'Linking hypothesis' followed by 'Maximize expected gain by choice of $d(.)$ '. Below this box, the text 'Despite changes in gain function, likelihood, prior.' is written.

$G(a, w)$

Gain

$d : X \rightarrow A$

$\pi(w)$

Prior

$L(w | x)$

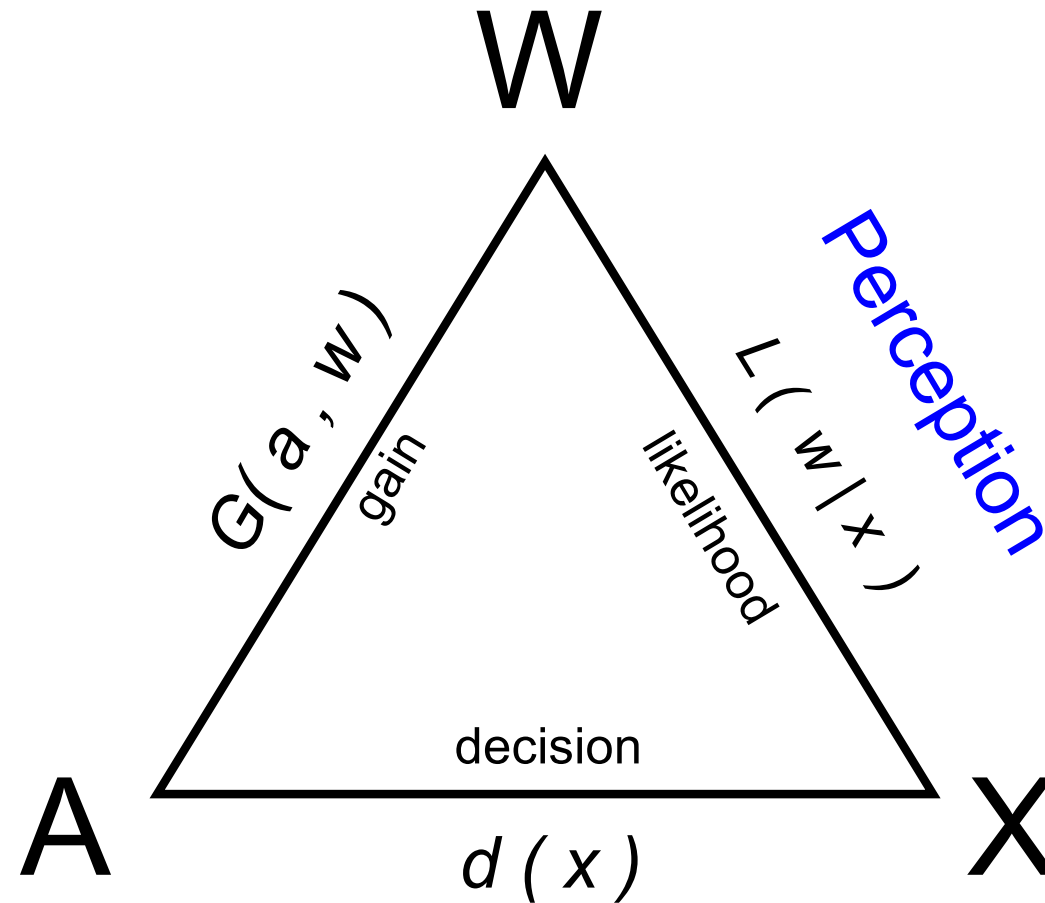
Likelihood

Linking hypothesis

Maximize expected gain by
choice of $d(.)$

Despite changes in gain function, likelihood, prior.

Three Functions of SDT



A Bayesian Game

:

Speeded Reaching



A Bayesian Problem

Trommershäuser *et al.*

Vol. 20, No. 7/July 2003/J. Opt. Soc. Am. A 1419

Statistical decision theory and the selection of rapid, goal-directed movements

Julia Trommershäuser, Laurence T. Maloney, and Michael S. Landy

Department of Psychology and Center for Neural Science, New York University, New York, New York 10003

Received October 1, 2002; revised manuscript received January 30, 2003; accepted February 3, 2003

We present two experiments that test the range of applicability of a movement planning model (MEGaMove) based on statistical decision theory. Subjects attempted to earn money by rapidly touching a green target region on a computer screen while avoiding nearby red penalty regions. In two experiments we varied the magnitudes of penalties, the degree of overlap of target and penalty regions, and the number of penalty regions. Overall, subjects acted so as to maximize gain in a wide variety of stimulus configurations, in good agreement with predictions of the model. © 2003 Optical Society of America

OCIS codes: 330.4060, 330.7310.

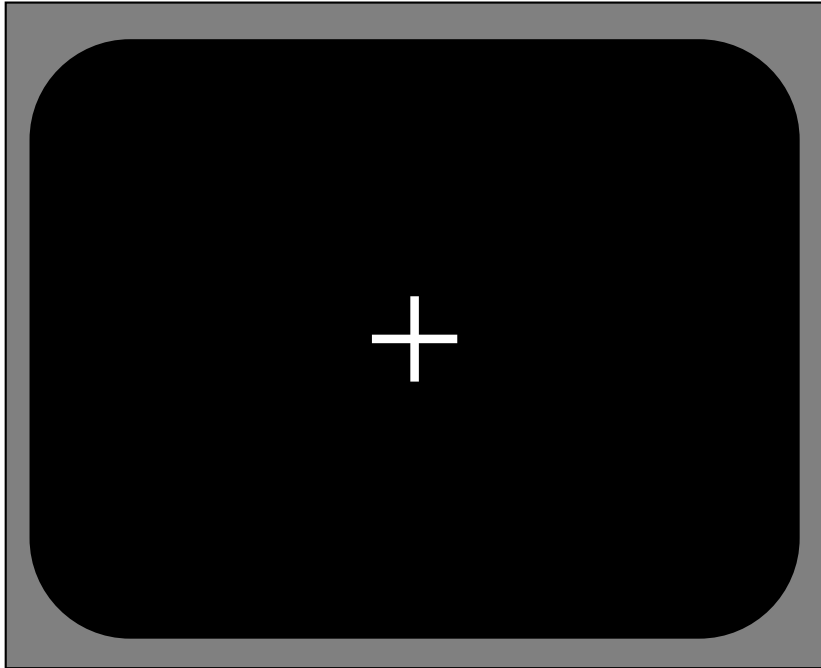
1. INTRODUCTION

Motor responses have consequences. In the first semifinal of the 2002 World Cup, Germany met host South Korea, and the match was decided in the 75th minute by the only goal of the match. Oliver Neuville passed the ball to Germany's play maker Michael Ballack, who scored in his

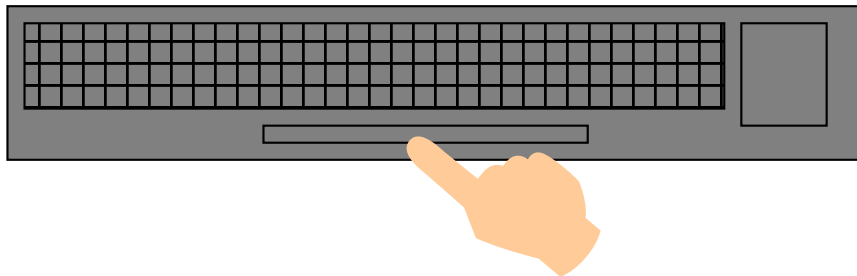
small target region, the subject cannot be certain that a movement aimed at the center of the reward region will not, instead, end up outside the reward region, possibly in the penalty region. To summarize, the subject, in each trial, must select among possible actions and must do so very rapidly. There are explicit monetary penalties associated with the outcome of the action selected, but the



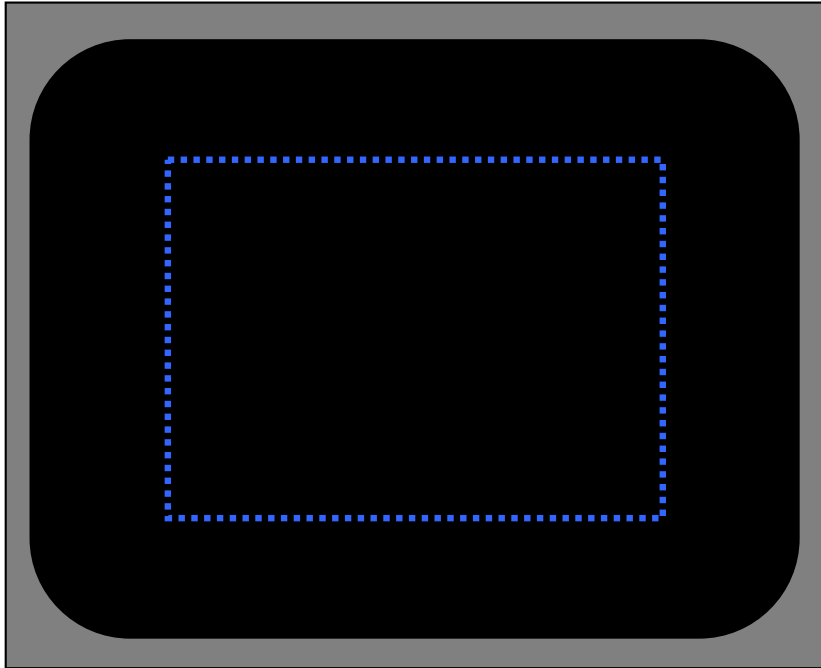
Experimental Task



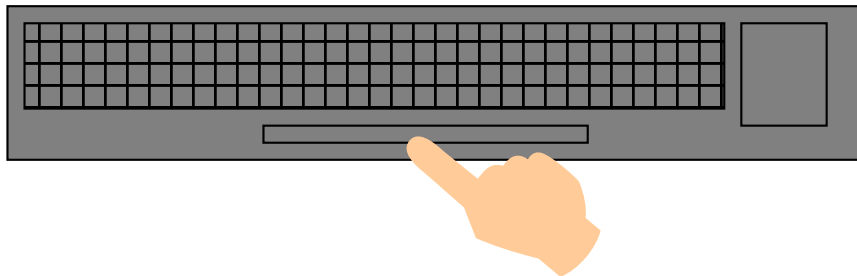
Start of trial:
display of fixation
cross (1.5 s)



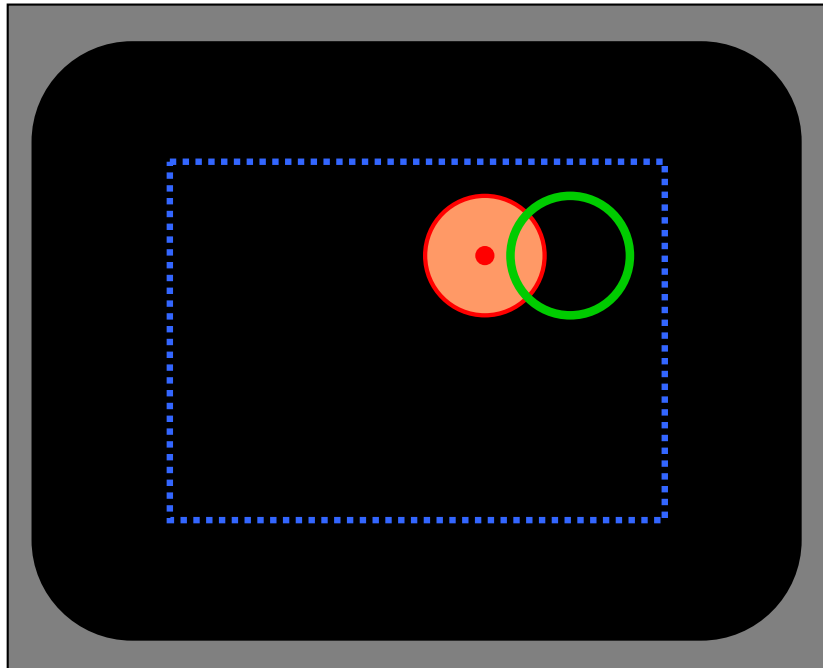
Experimental Task



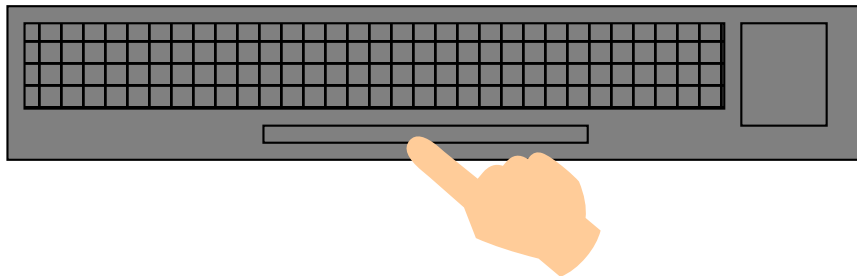
Display of response area,
500 ms before
target onset
(114.2 mm x 80.6 mm)



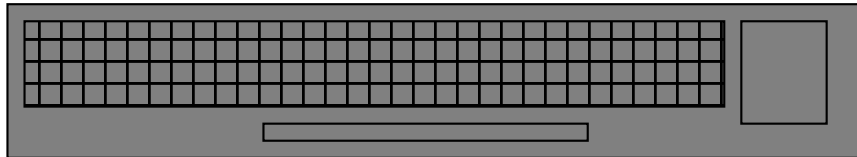
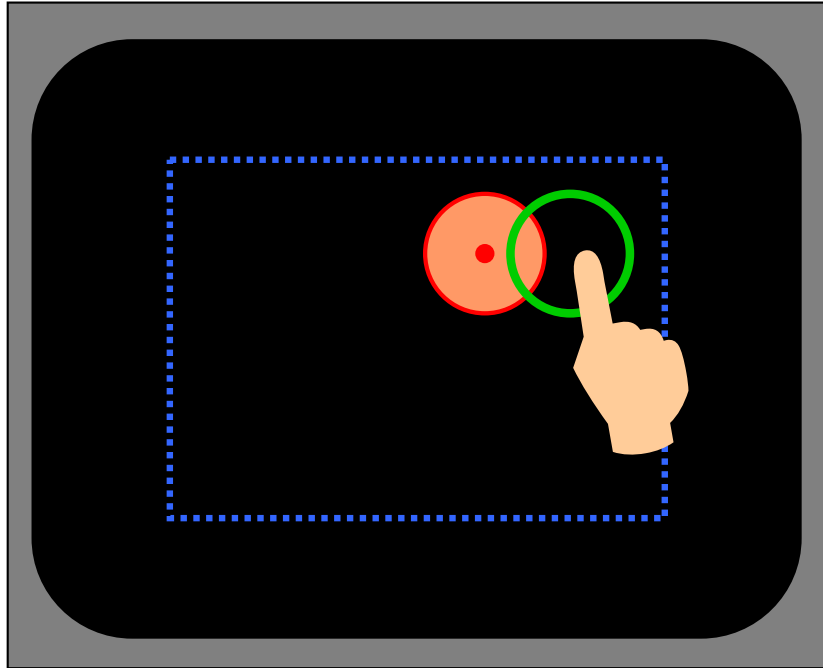
Experimental Task



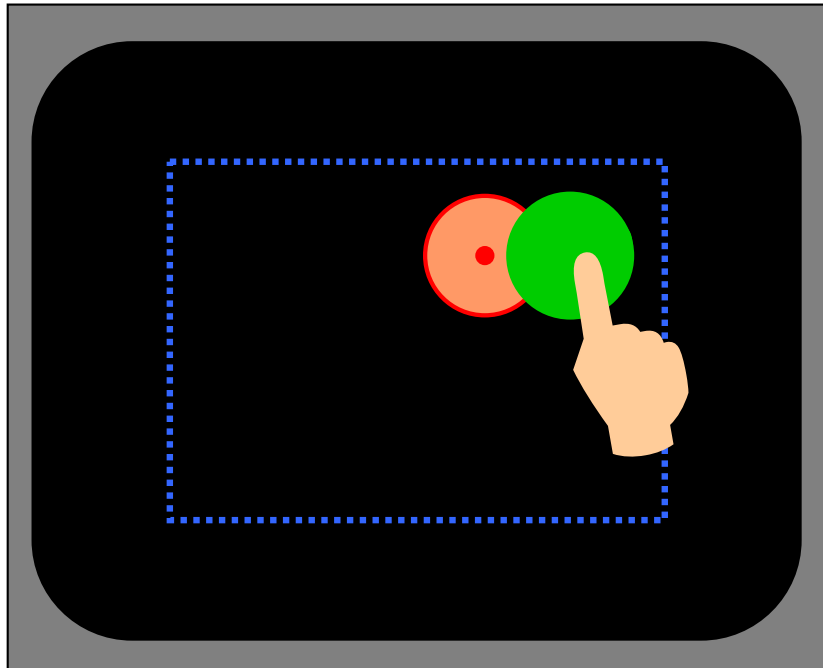
Target display (700 ms)



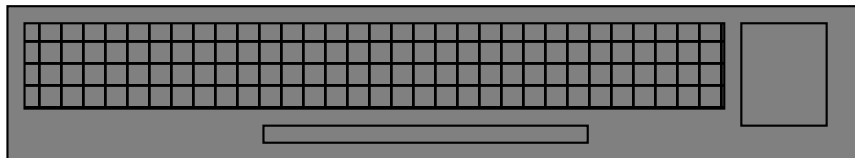
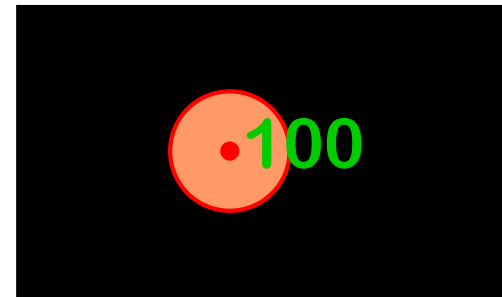
Experimental Task



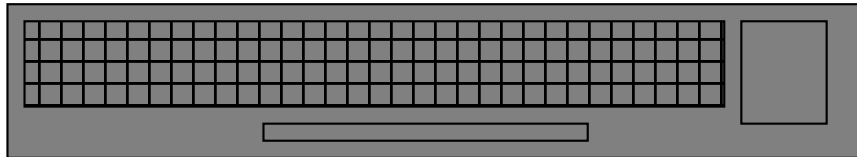
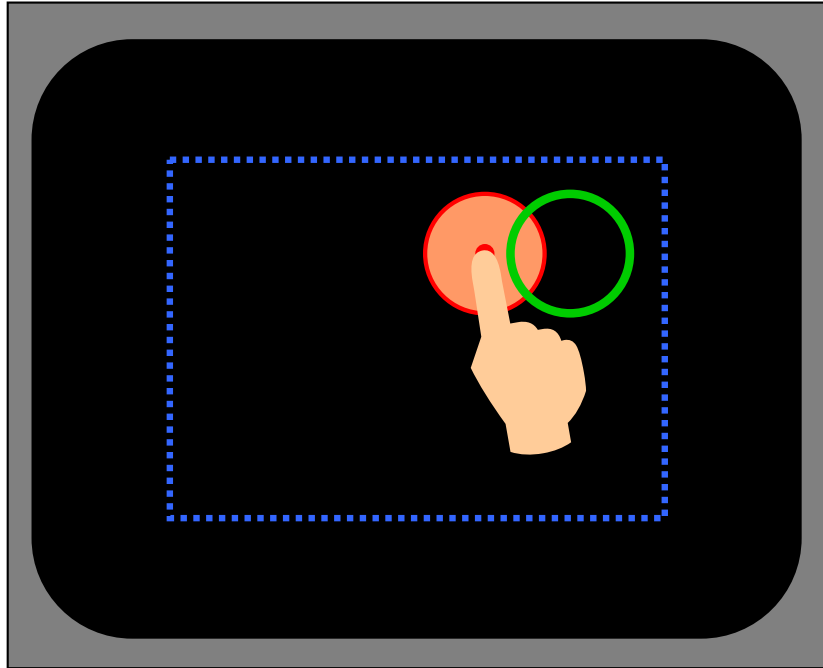
Experimental Task



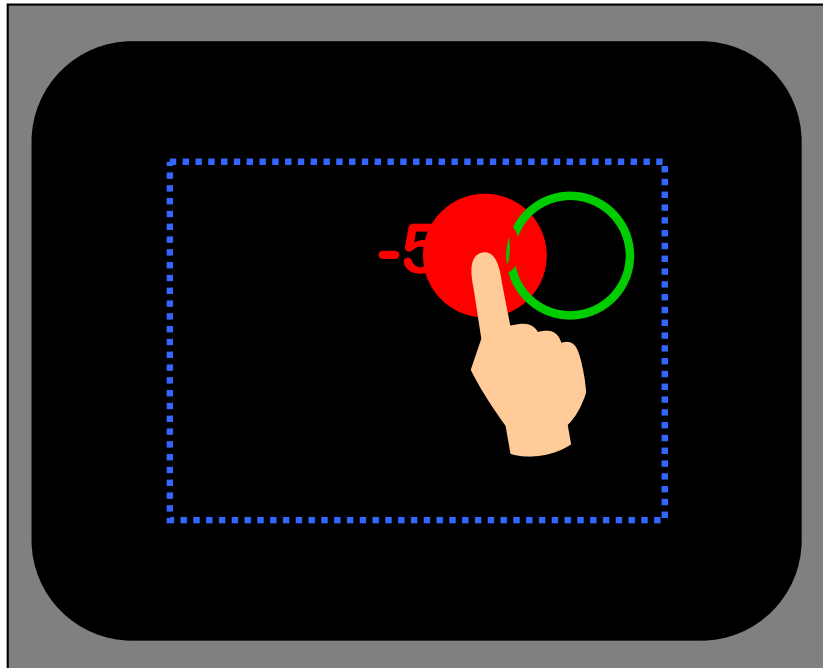
The green target is hit:
+100 points



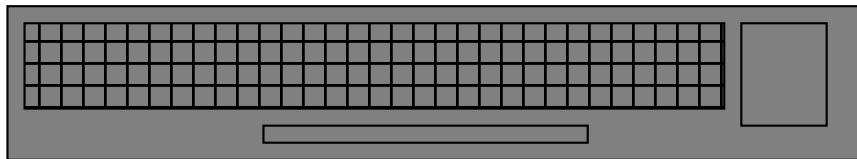
Experimental Task



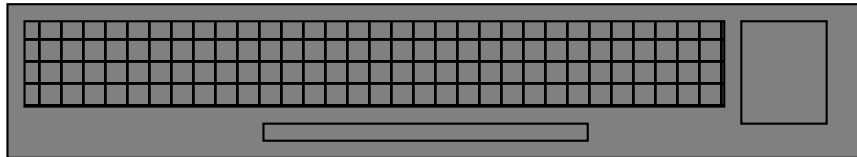
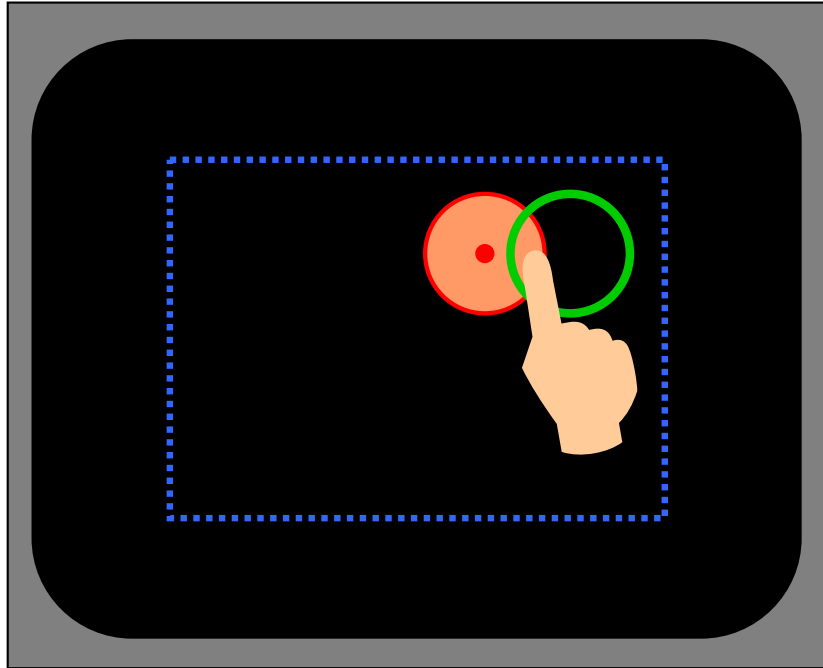
Experimental Task



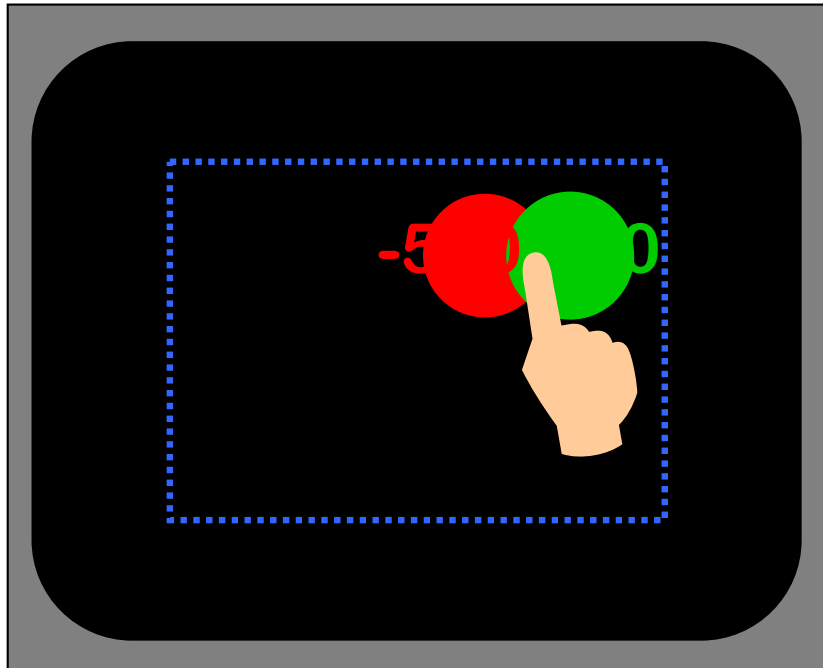
The red target is hit:
-500 points



Experimental Task



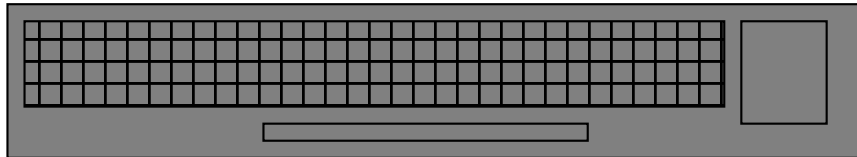
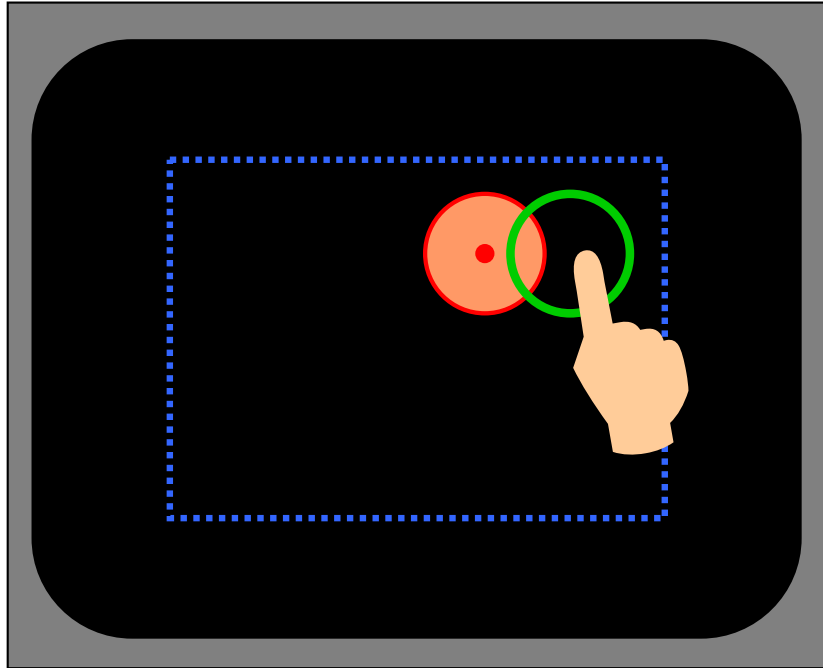
Experimental Task



Scores add if both targets are hit:

-500 100

Experimental Task

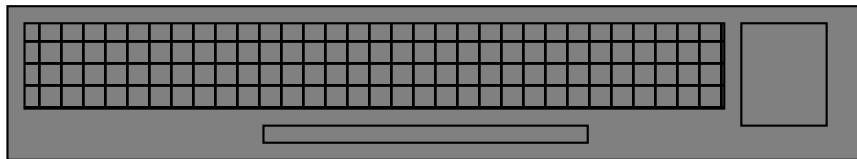


Experimental Task

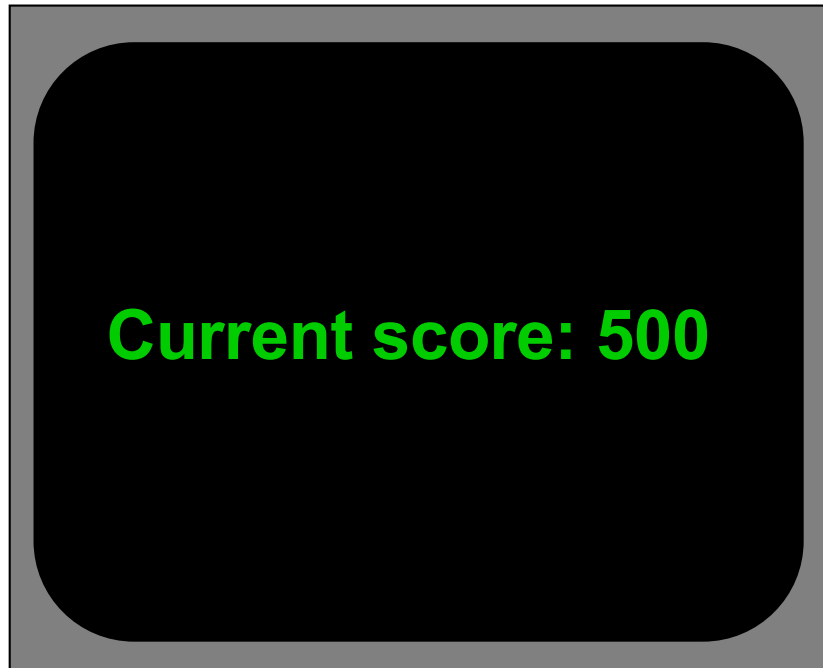


The screen is hit
later than 700 ms
after target display:
-700 points.

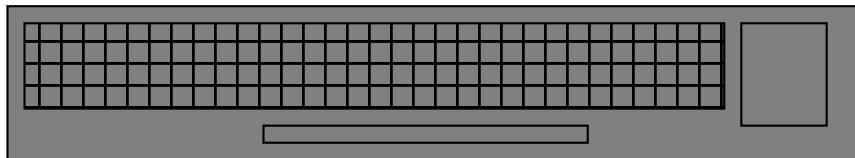
If you are on time but
Miss the targets, 0.



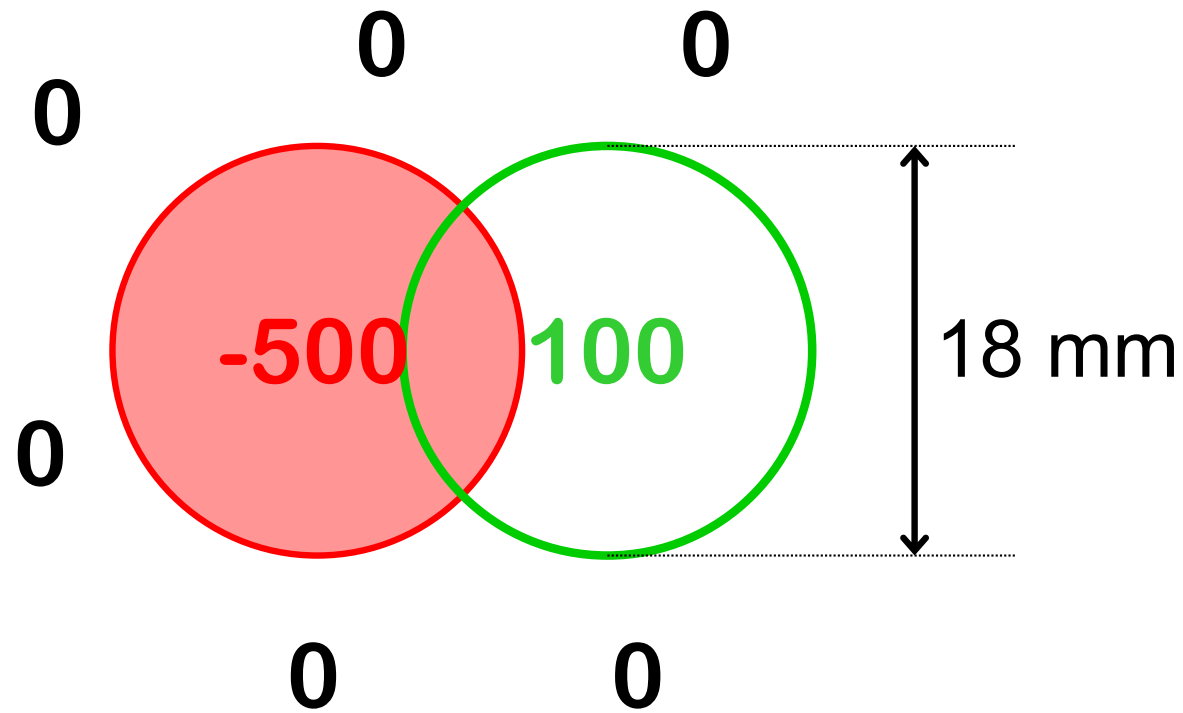
Experimental Task



End of trial

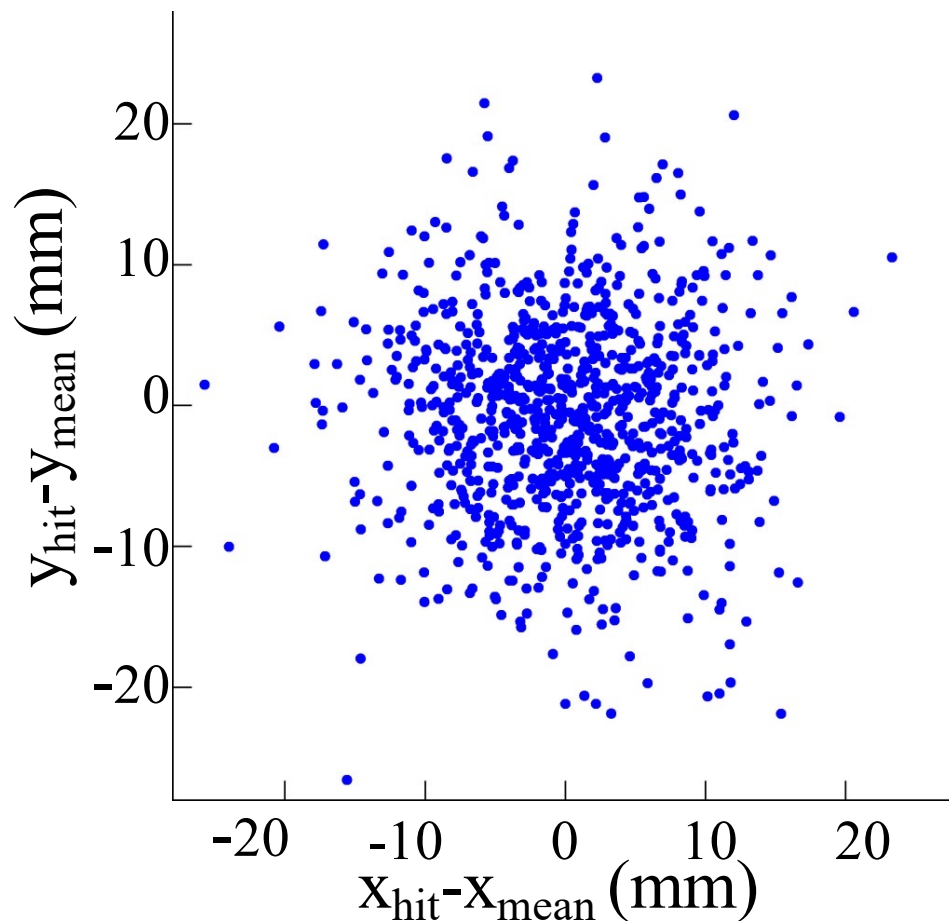


Choice among Movement Strategies



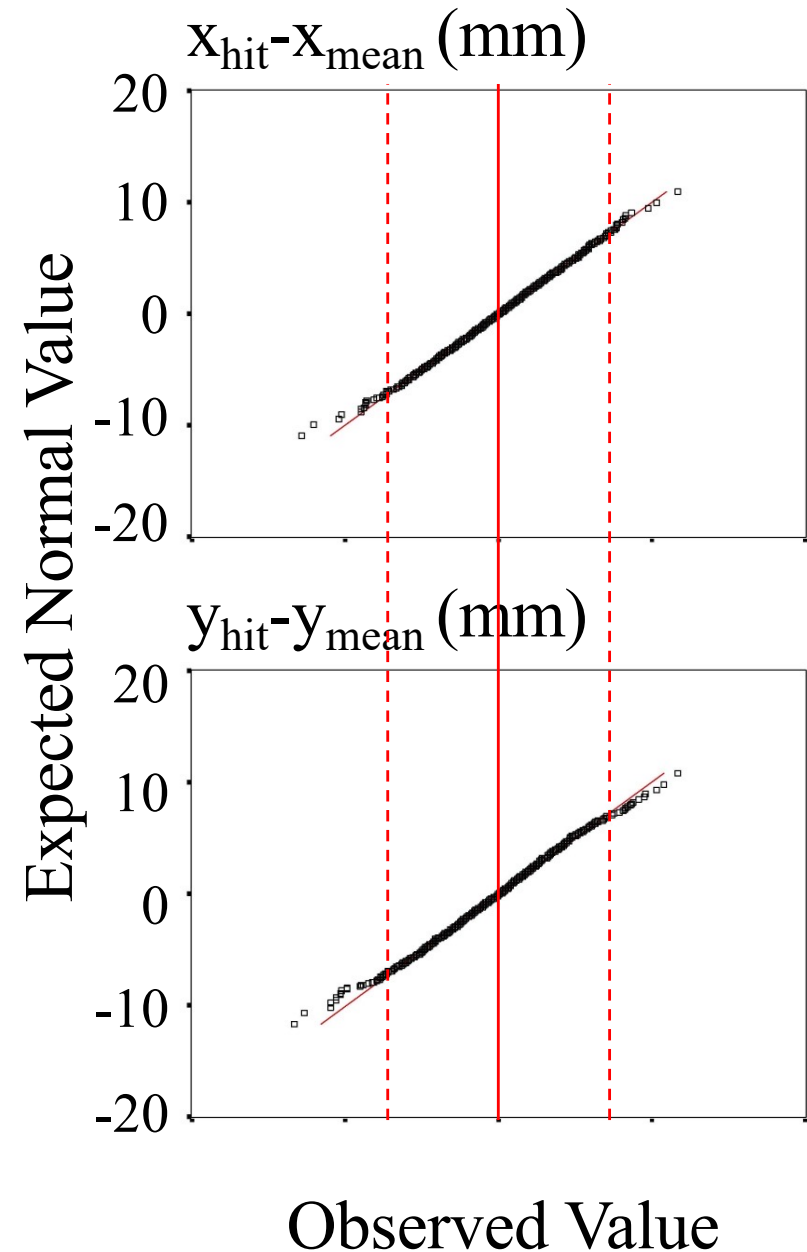
What should Paulina do?

Distribution of movement end points

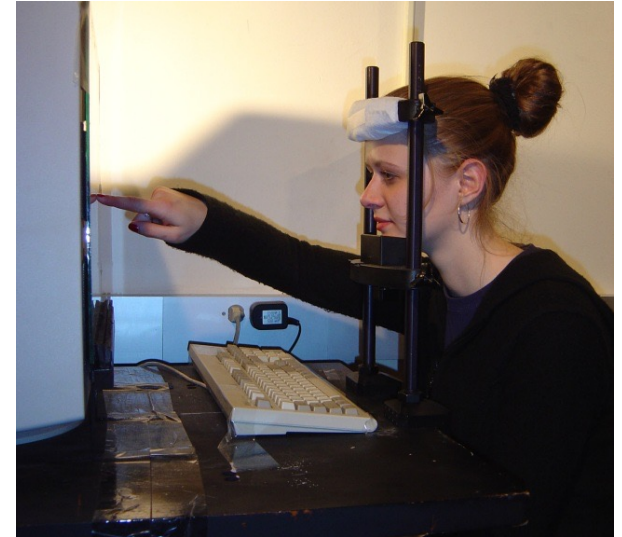
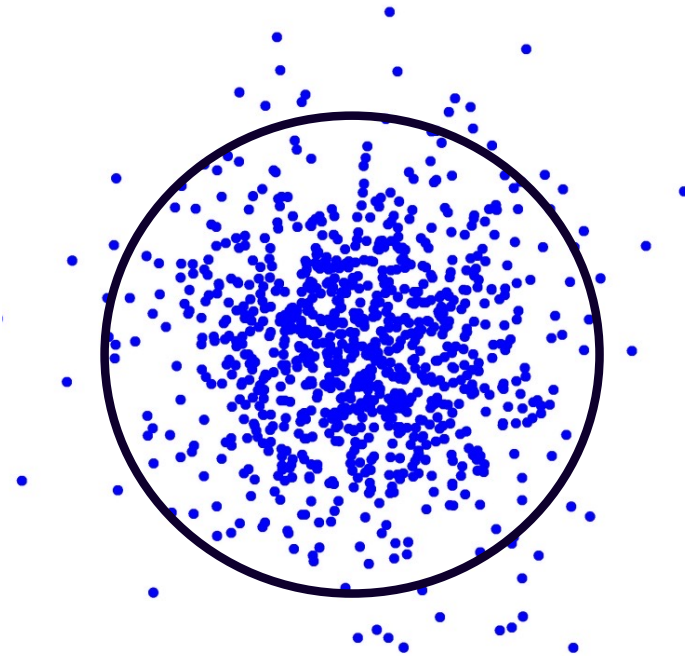


Subject S4, $\sigma = 3.62$ mm,
72x15 = 1080 end points

Q-Q Plot



If there were no red penalty circle

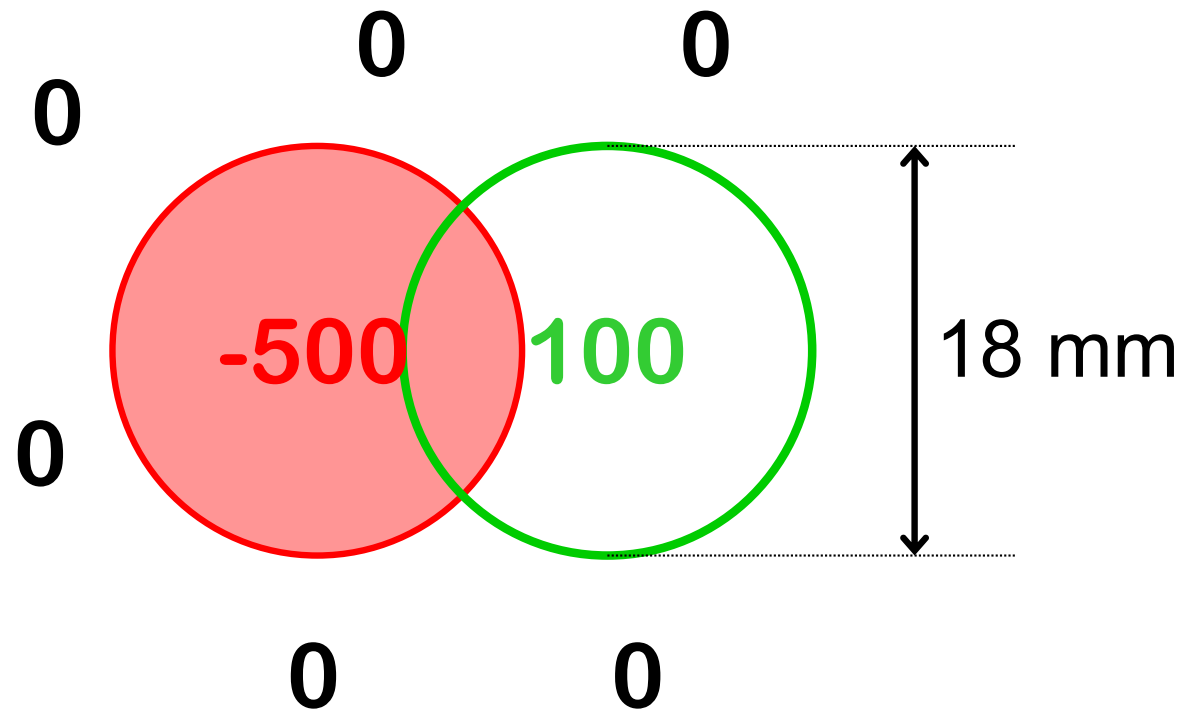


Aim for center

*Select perceptual-motor strategies that
minimize variance*

Harris & Wolpert (1998)

Choice among Movement Strategies

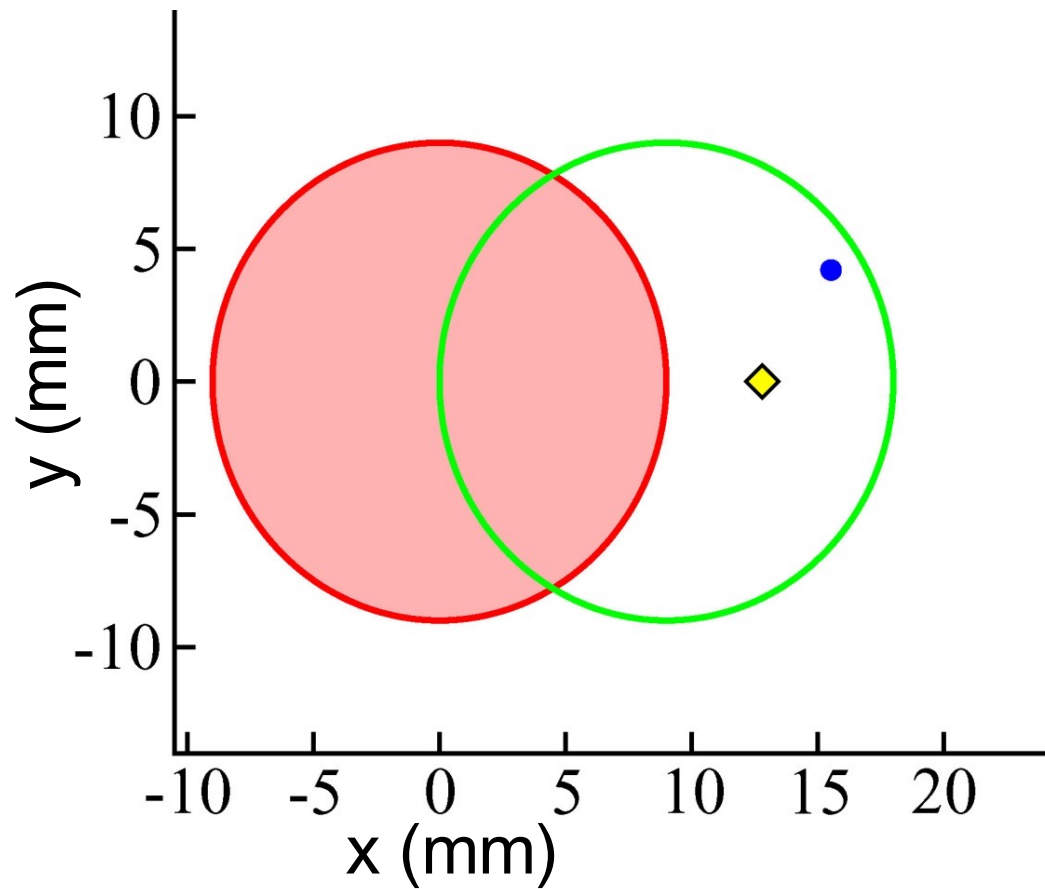


What should Paulina do?

Thought Experiment

● : -500

○ : 100 points (2.5 ¢)



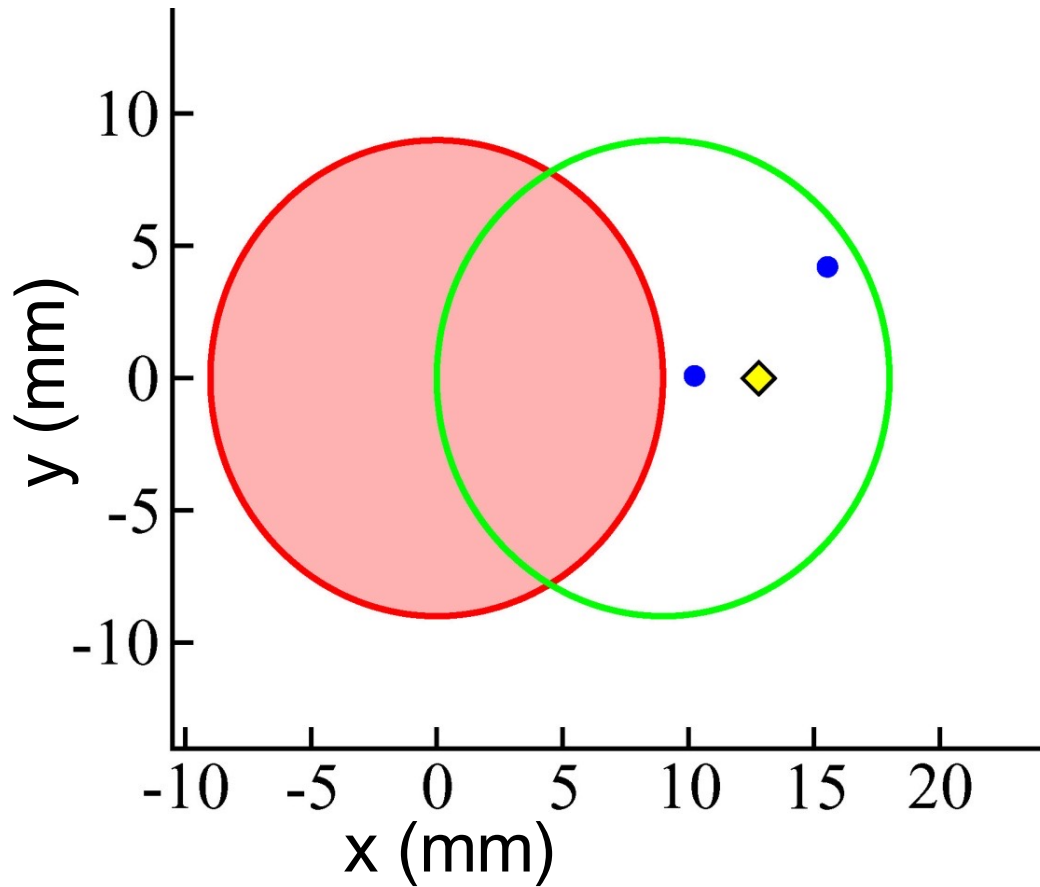
100 points

$$\sigma = 4.83 \text{ mm}$$

Thought Experiment

● : -500

○ : 100 points (2.5 ¢)



100 points

100 points

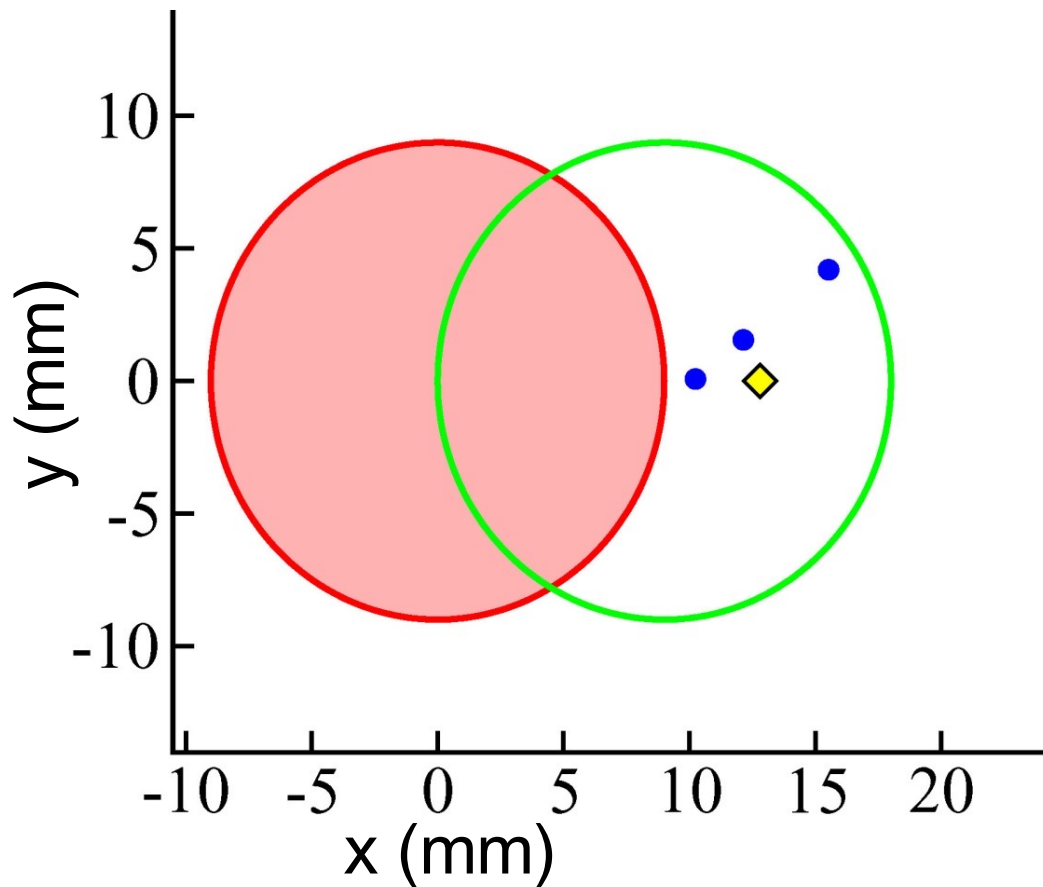
200 points

$\sigma = 4.83 \text{ mm}$

Thought Experiment

● : -500

○ : 100 points (2.5 ¢)



100 points

100 points

100 points

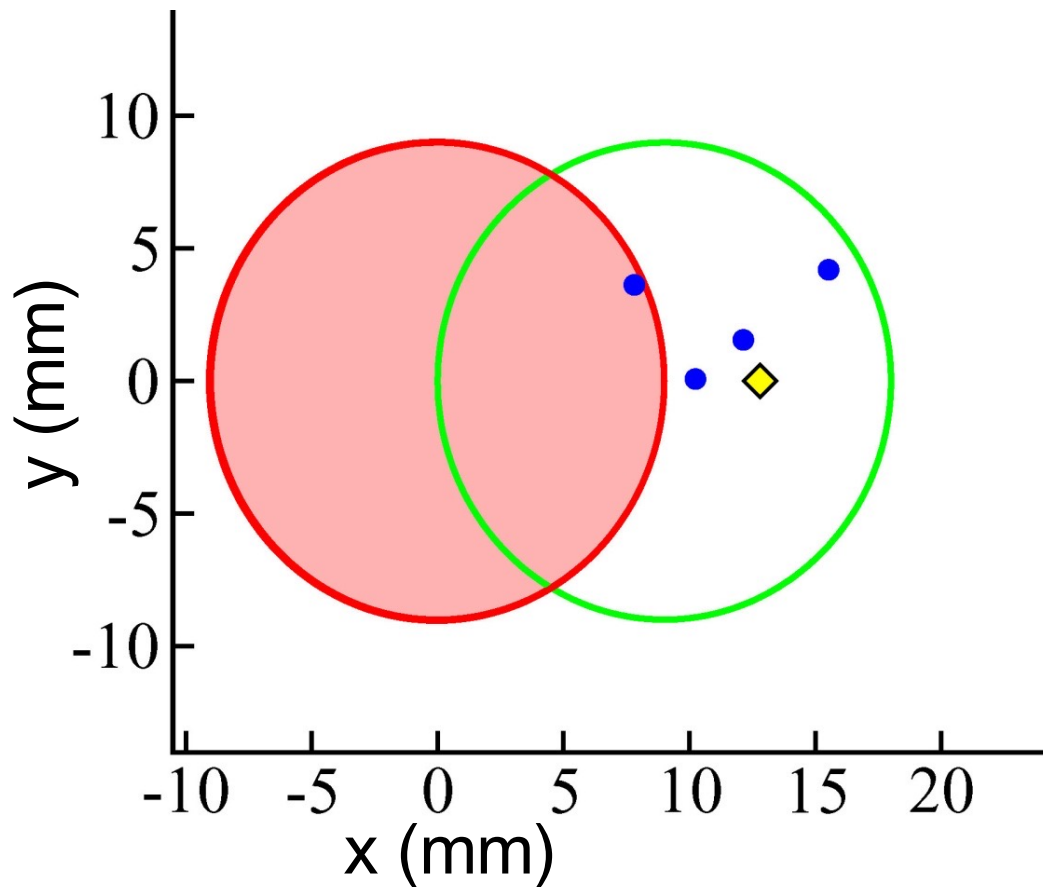
300 points

$\sigma = 4.83 \text{ mm}$

Thought Experiment

● : -500

○ : 100 points (2.5 ¢)



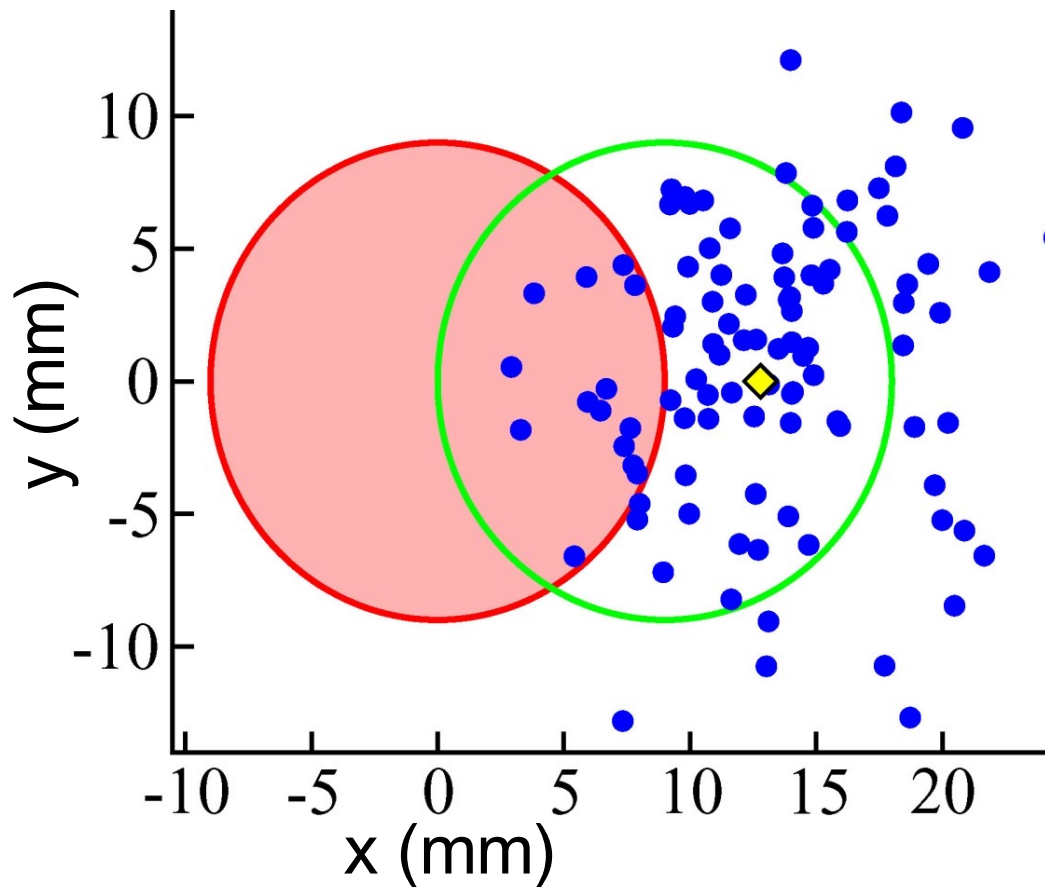
100 points
100 points
100 points
-400 points
-100 points

$\sigma = 4.83 \text{ mm}$

Thought Experiment

● : -500

○ : 100 points (2.5 ¢)



100 points
100 points
100 points
-400 points
⋮

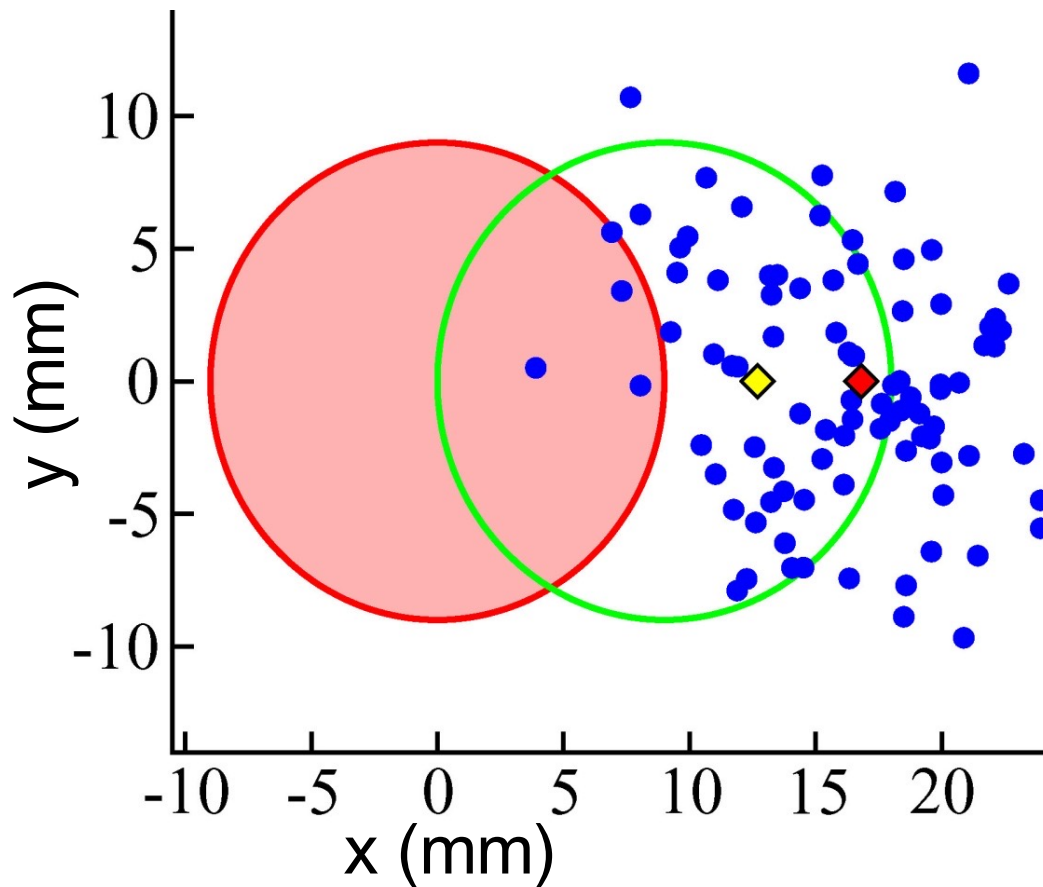
-0.3 pts. per trial

$\sigma = 4.83 \text{ mm}$

Thought Experiment

● : -500

○ : 100 points (2.5 ¢)



- 0.3 pts. per trial

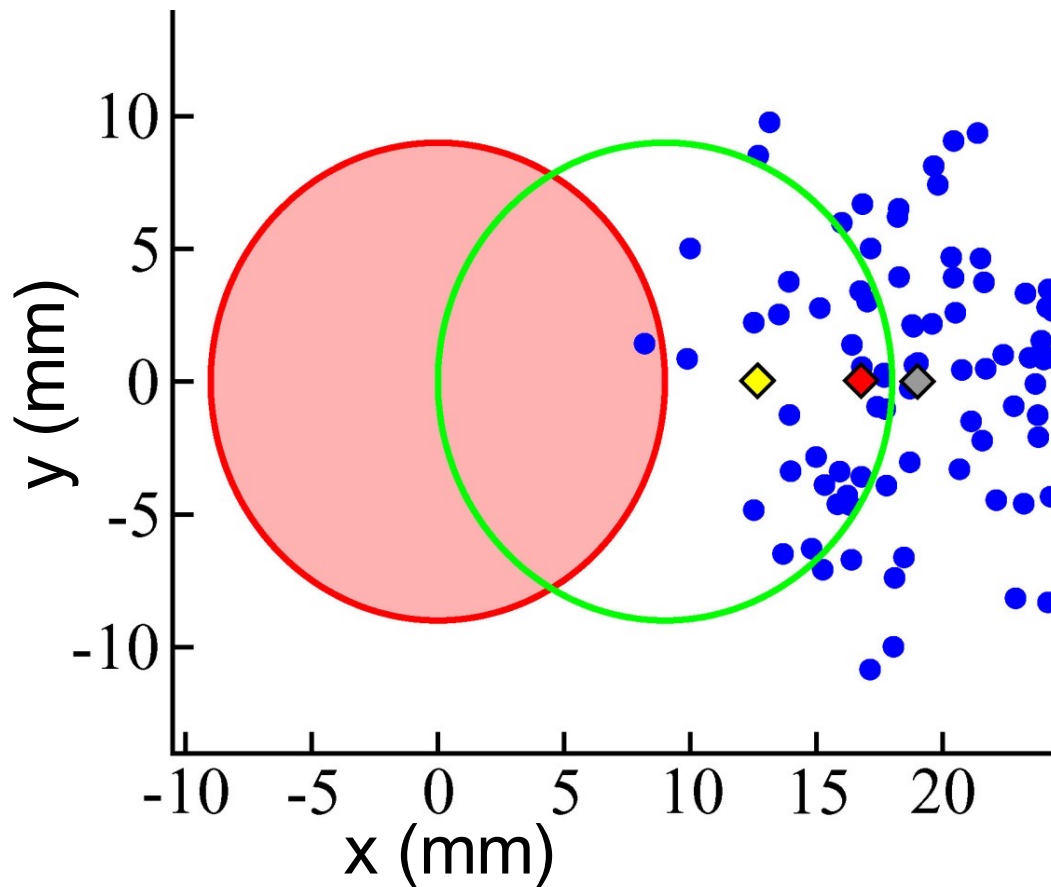
30.7 pts. per trial

$$\sigma = 4.83 \text{ mm}$$

Thought Experiment

● : -500

○ : 100 points (2.5 ¢)



- 0.3 pts. per trial

30.7 pts. per trial

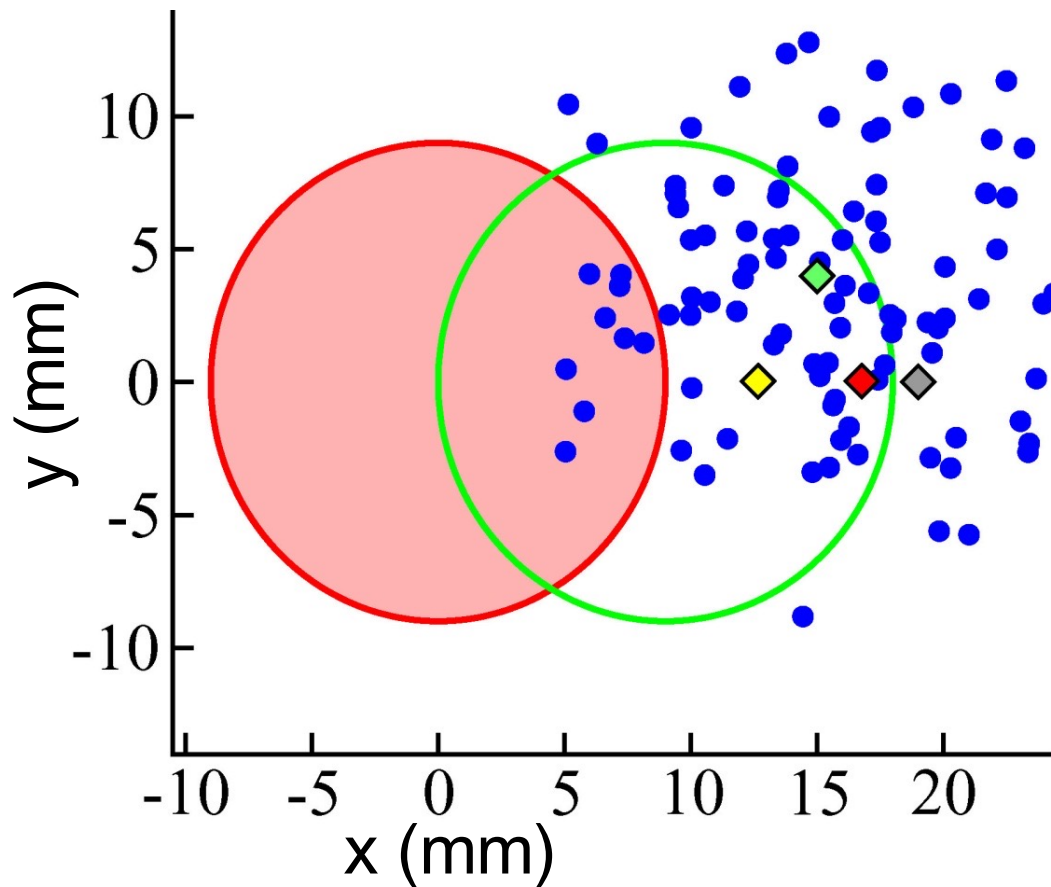
25.5 pts. per trial

$\sigma = 4.83 \text{ mm}$

Thought Experiment

● : -500

○ : 100 points (2.5 ¢)



- 0.3 pts. per trial

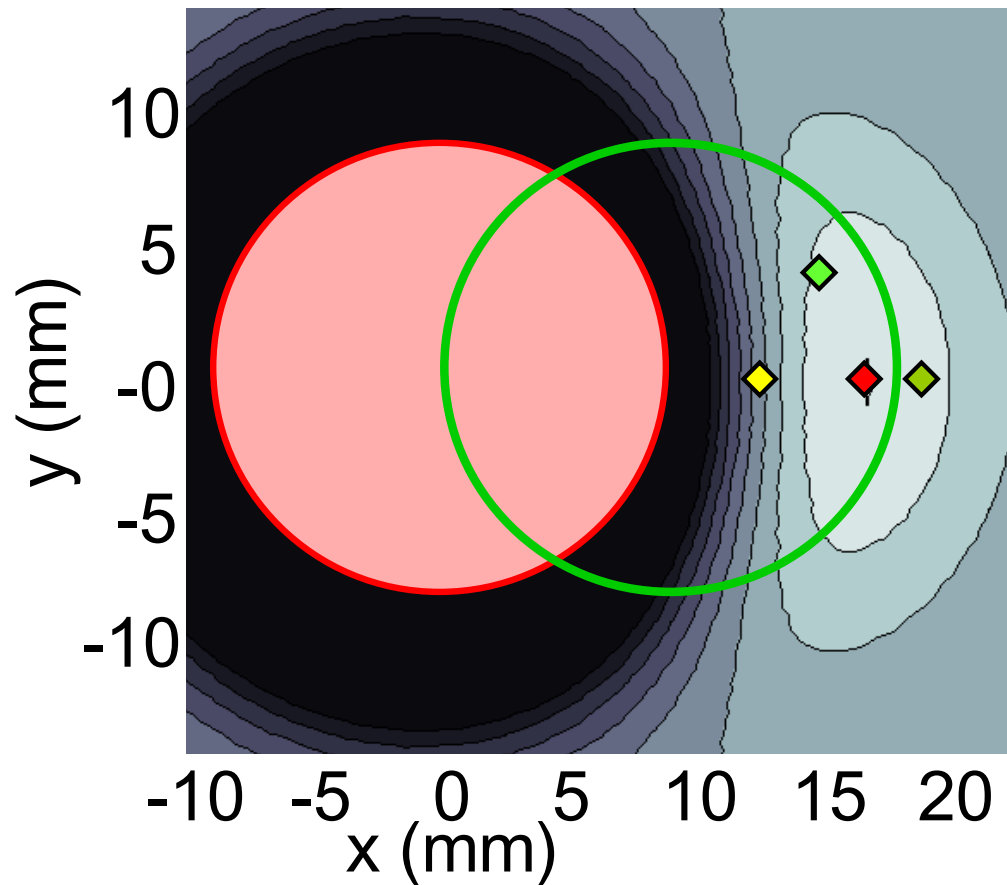
30.7 pts. per trial

25.5 pts. per trial

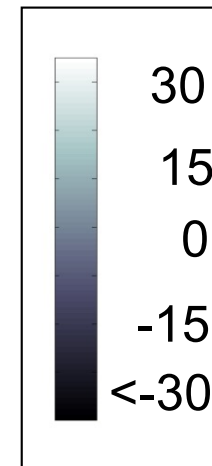
22.6 pts. per trial

$\sigma = 4.83 \text{ mm}$

Expected value as function of
mean movement end point (x,y):



points
per trial



target: 100
penalty: -500

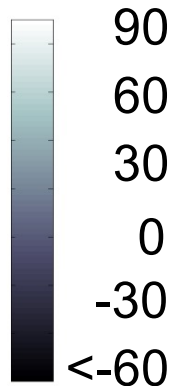
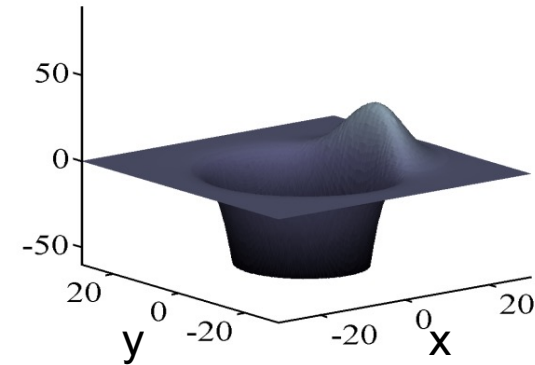
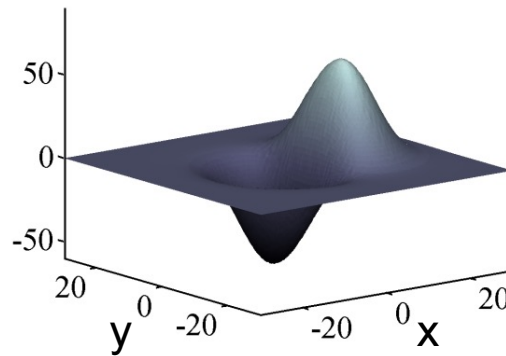
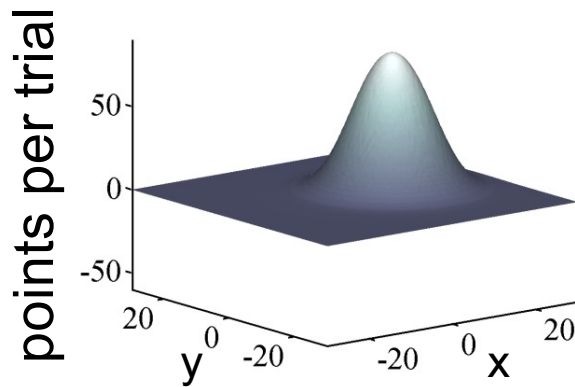
$\sigma = 4.83$ mm

Thought Experiment

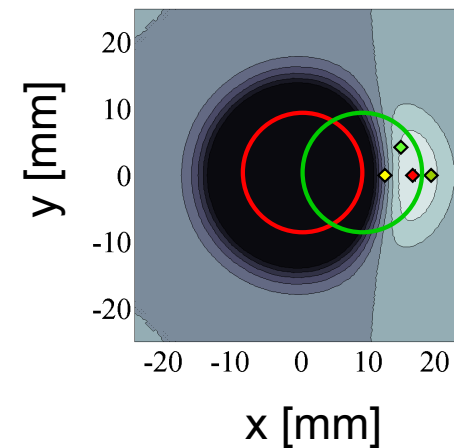
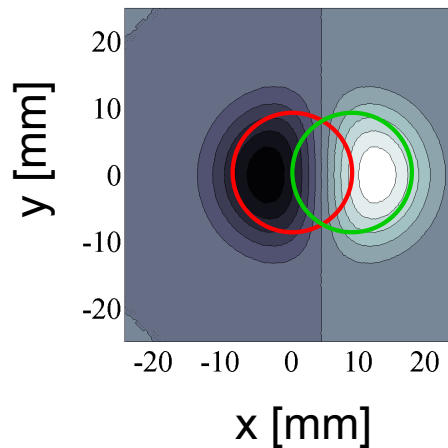
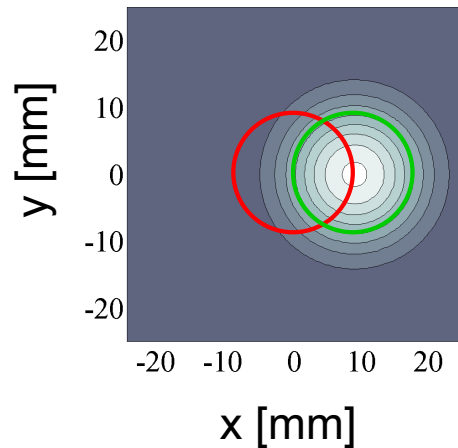
penalty: 0

penalty: 100

penalty: 500

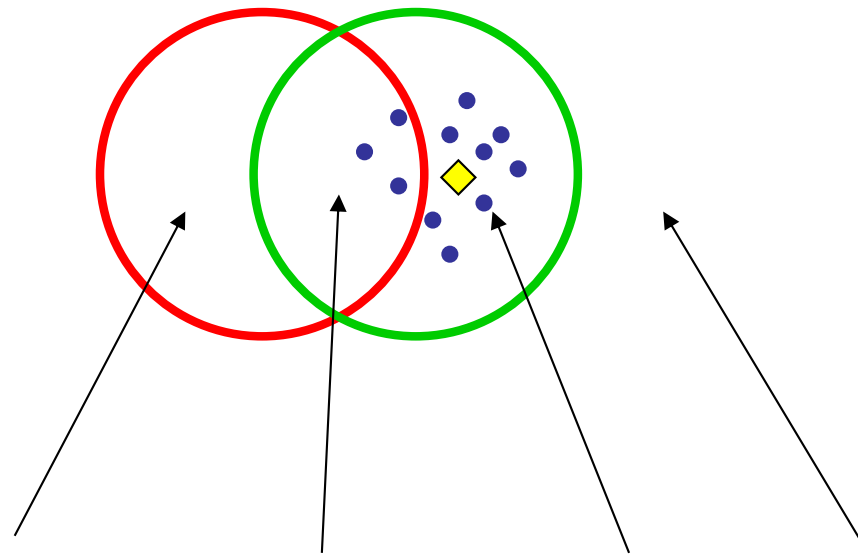


x, y: mean movement end point [mm]



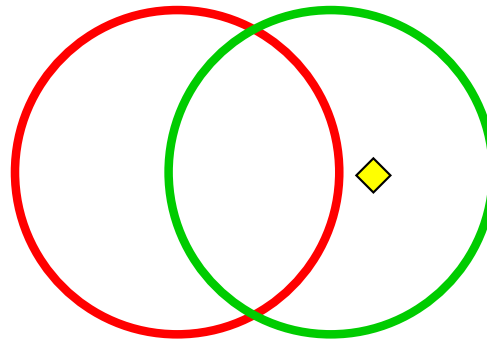
$$\sigma = 4.83 \text{ mm}$$

Movement plans as lotteries



$(p_1, -500; p_2, -400; p_3, 100; p_4, 0)$

Movement plans as lotteries



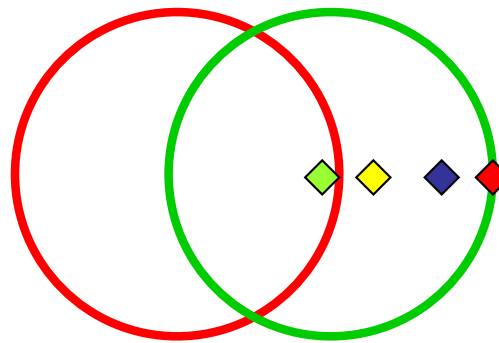
$$\sigma = 4 \text{ mm}$$

Lottery:

(1.3%, -500; 30.3%, -400; 60.9%, 100; 7.5%, 0)

Movement plans as lotteries

Optimal aim point: lottery with MEV



$\sigma = 4$ mm

(6.6%, -500; 52.3%, -400; 37.0%, 100; 4.0%, 0)

(1.3%, -500; 30.3%, -400; 60.9%, 100; 7.5%, 0)

(0%, -500; 4.6%, -400; 62.6%, 100; 32.8%, 0)

(0%, -500; 0.7%, -400; 37.6%, 100; 61.7%, 0)

Experiment

Reaching with Asymmetric Gain/Loss

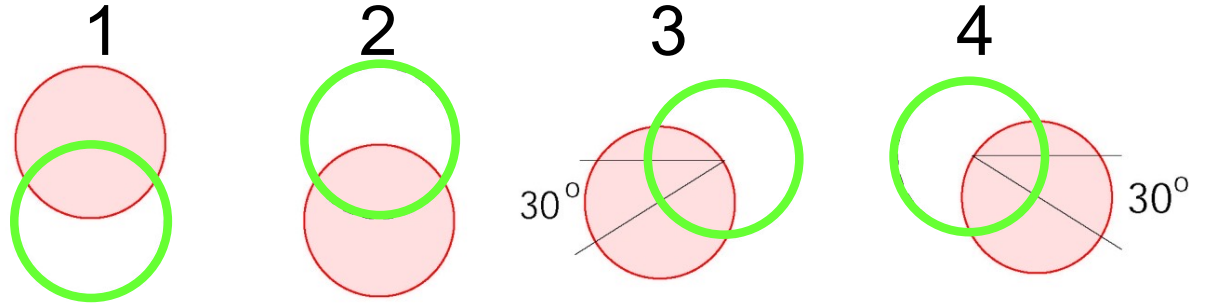


Julia Trommershäuser

Trommershäuser, Maloney, Landy (2003) JOSA A

Test of the model: Experiment 1

4 stimulus configurations:
(varied within block)



2 penalty conditions:
0 and -500 points (varied between blocks)

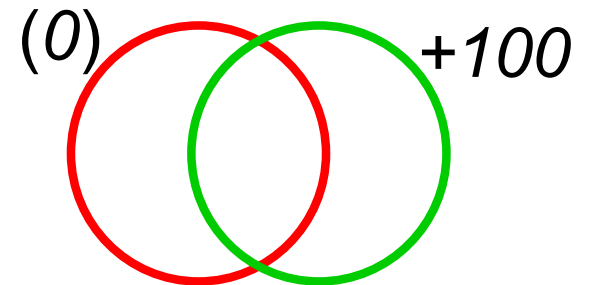
5 “practiced movers”

1 session of data collection: 360 trials
24 data points per condition

General Methods: Training

For all experiments:

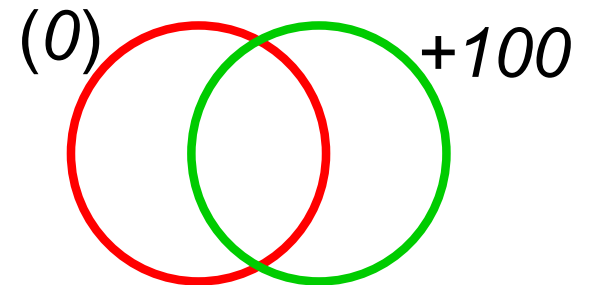
- *All subjects practice the task for 360 trials or more until their variance stabilizes.*
- *The timeout limit is gradually decreased to 700 ms during training.*
- *There are no **penalties** during training (the concept is never mentioned).*
- *We verify that each subject's movement variance has stabilized.*
- *They are told only to **make money**.*



General Methods: Training

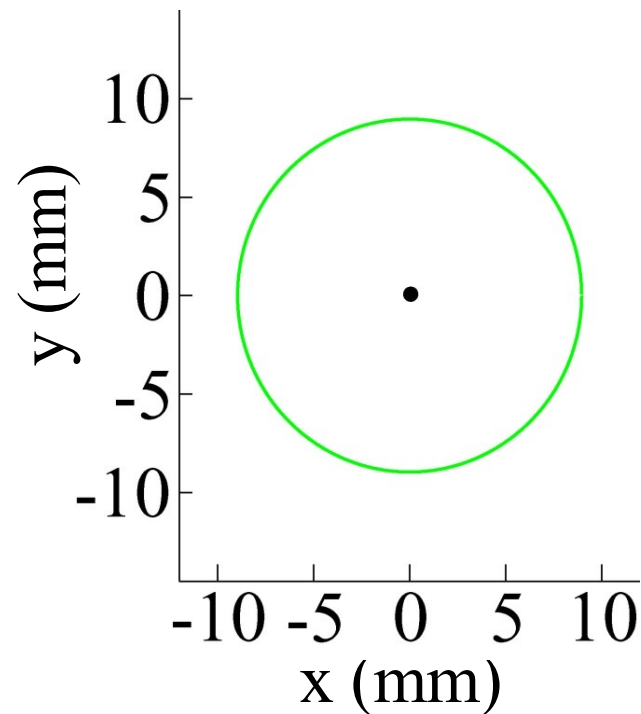
For all experiments:

- *All subjects practice the task for 360 trials or more until their variance stabilizes.*
- *The timeout limit is gradually decreased to 700 ms during training.*
- *There are no **penalties** during training (the concept is never mentioned).*
- *We verify that each subject's movement variance has stabilized.*
- *They are told only to **make money**.*



Results: Experiment 1

Model prediction:

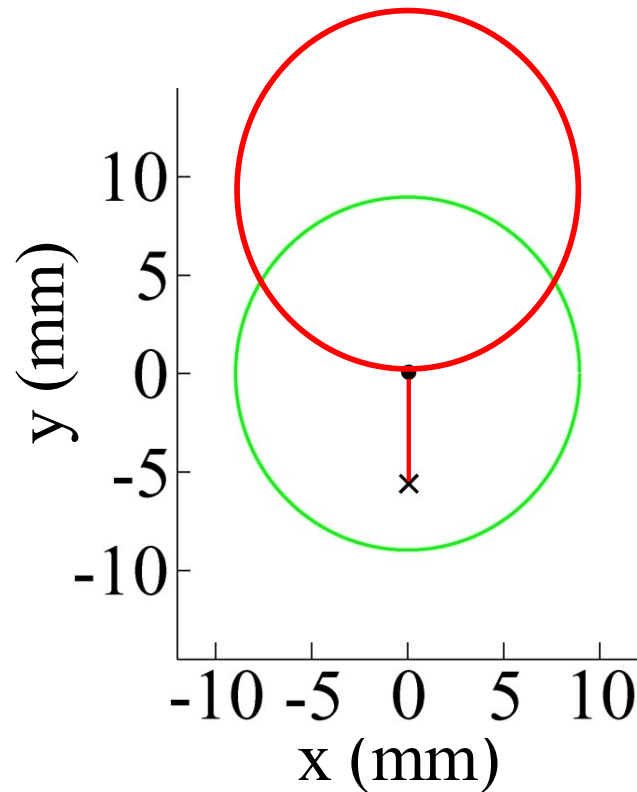


- model, penalty = 0

Subject S5, $\sigma = 2.99$ mm

Results: Experiment 1

Model prediction: configuration 1

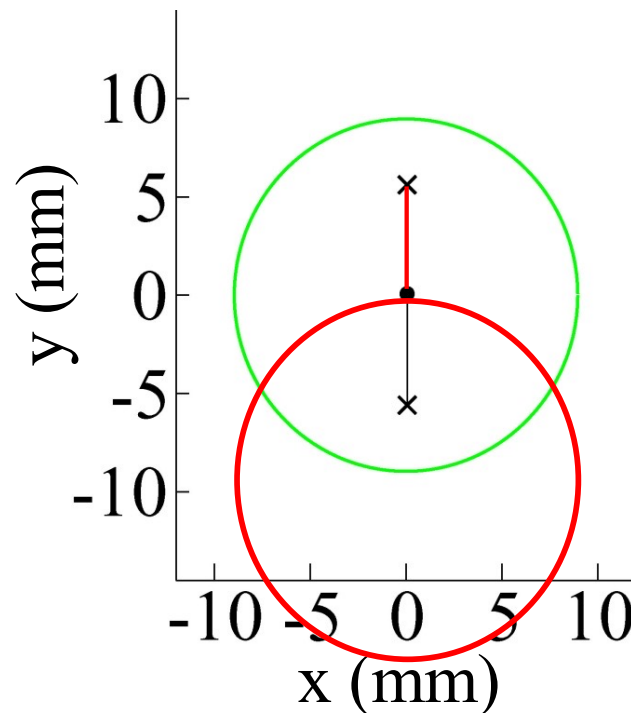


- model, penalty = 0
- × model, **penalty = 500**

Subject S5, $\sigma = 2.99$ mm

Results: Experiment 1

Model prediction: configuration 2

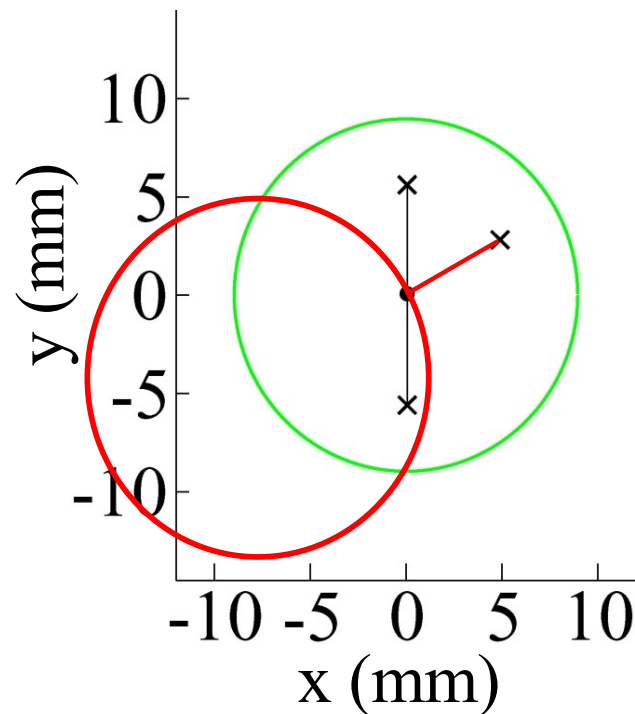


- model, penalty = 0
- × model, **penalty = 500**

Subject S5, $\sigma = 2.99$ mm

Results: Experiment 1

Model prediction: configuration 3

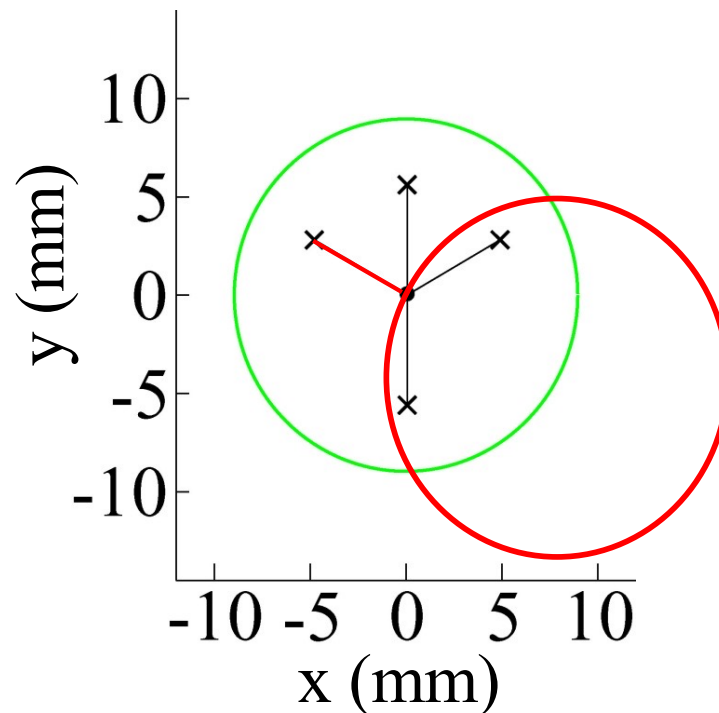


- model, penalty = 0
- × model, **penalty = 500**

Subject S5, $\sigma = 2.99$ mm

Results: Experiment 1

Model prediction: configuration 4

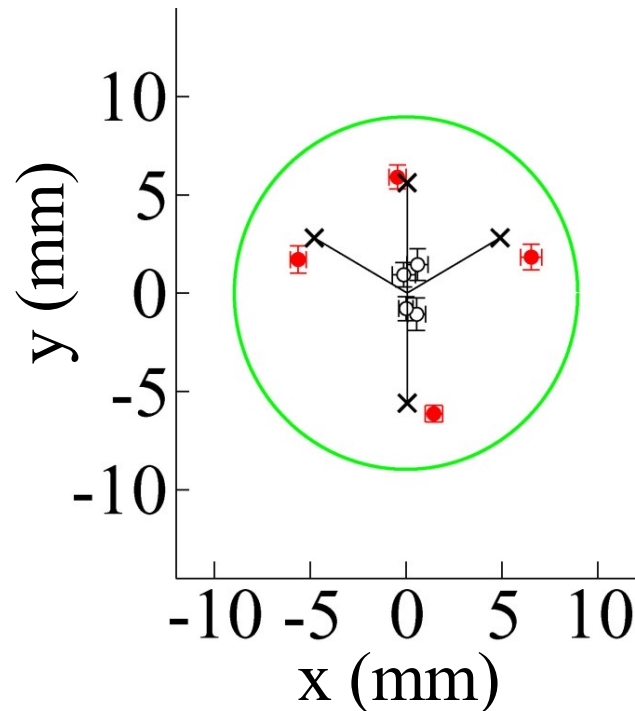


- model, penalty = 0
- × model, **penalty = 500**

Subject S5, $\sigma = 2.99$ mm

Results: Experiment 1

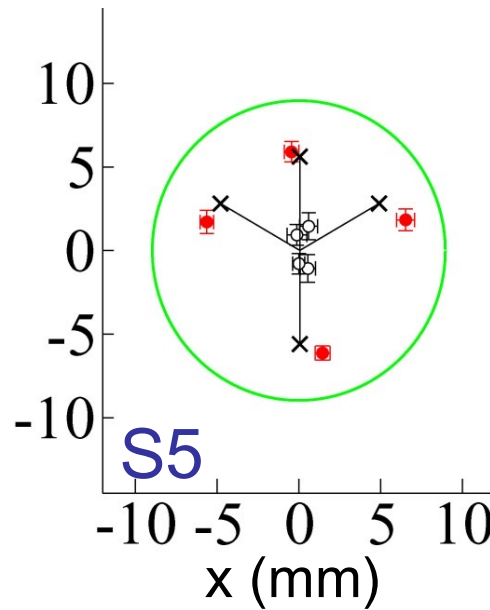
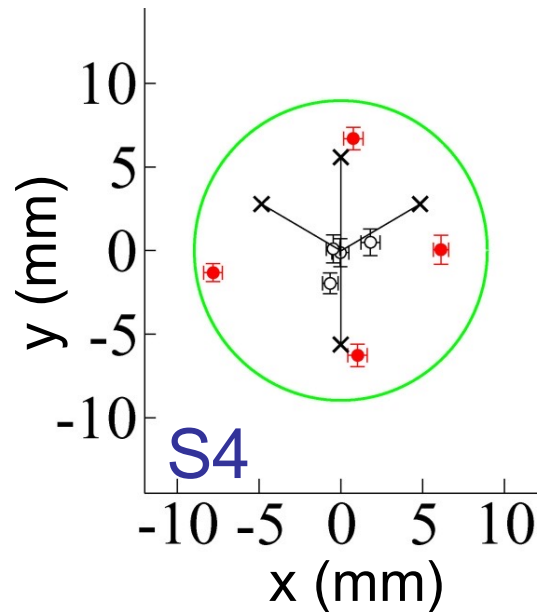
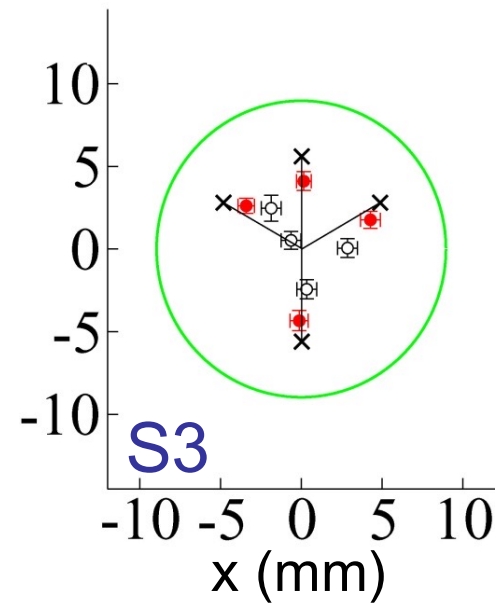
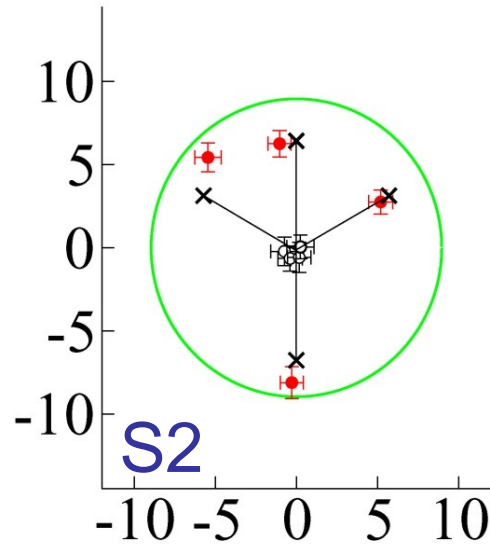
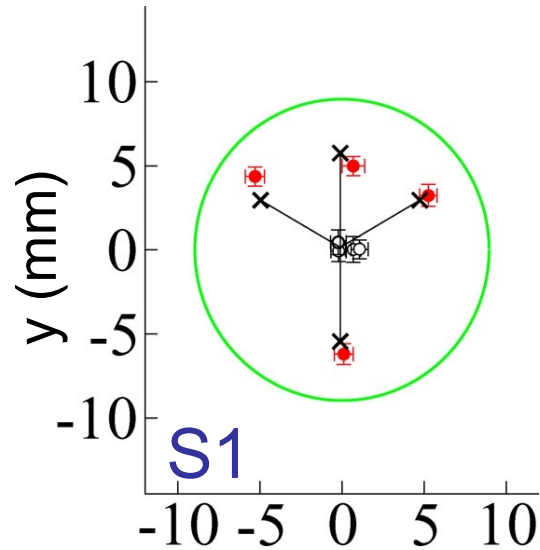
Comparison with experiment



- exp., penalty = 0
- exp., penalty = 500
- × model, penalty = 500

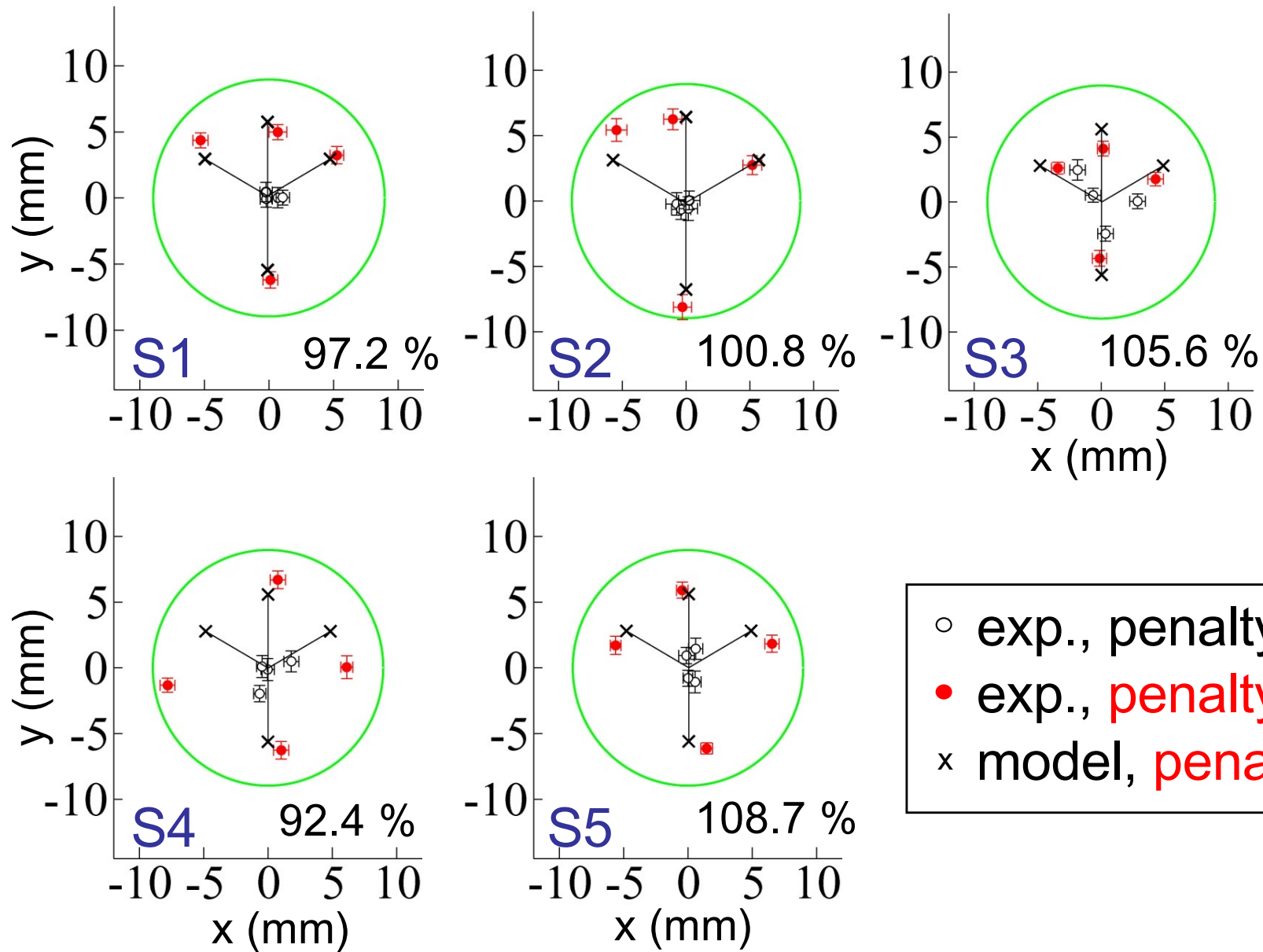
Subject S5, $\sigma = 2.99$ mm

Results: Experiment 1

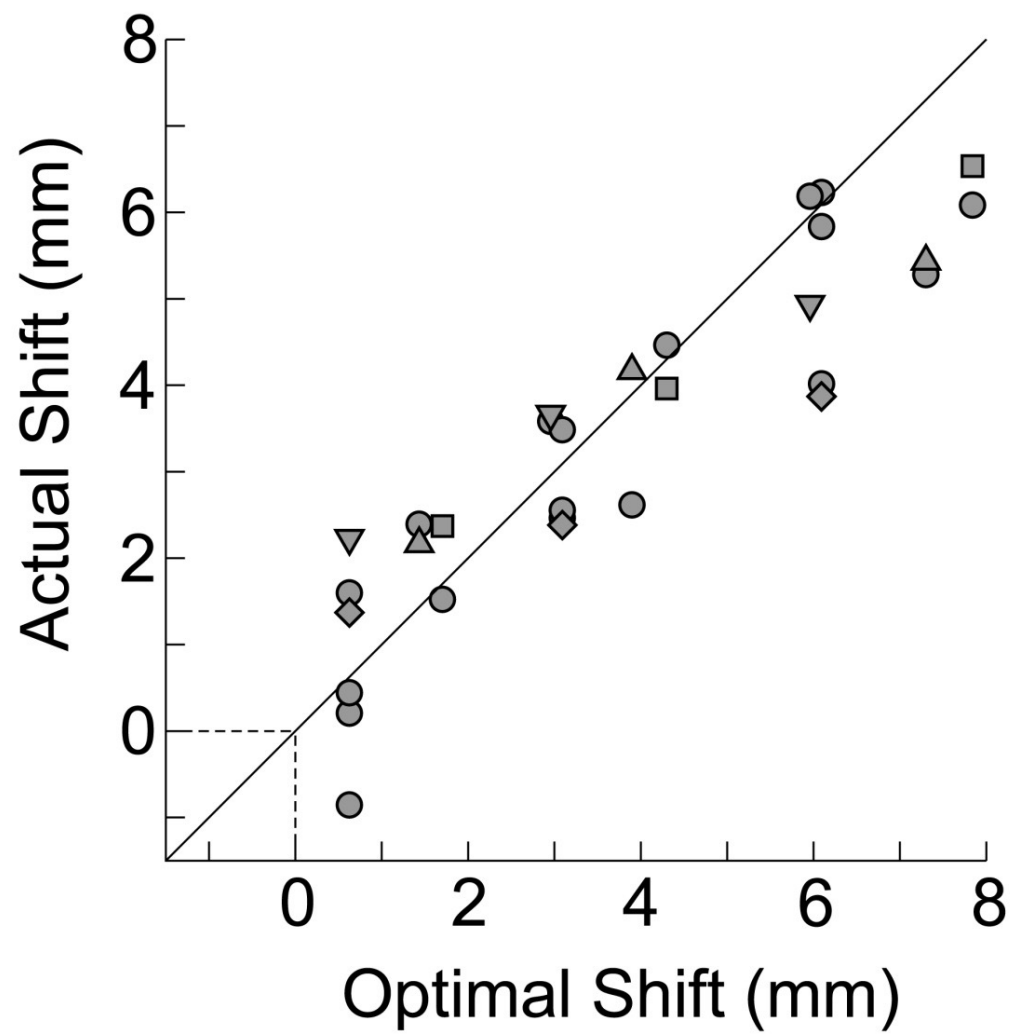


- exp., penalty = 0
- exp., penalty = 500
- × model, penalty = 500

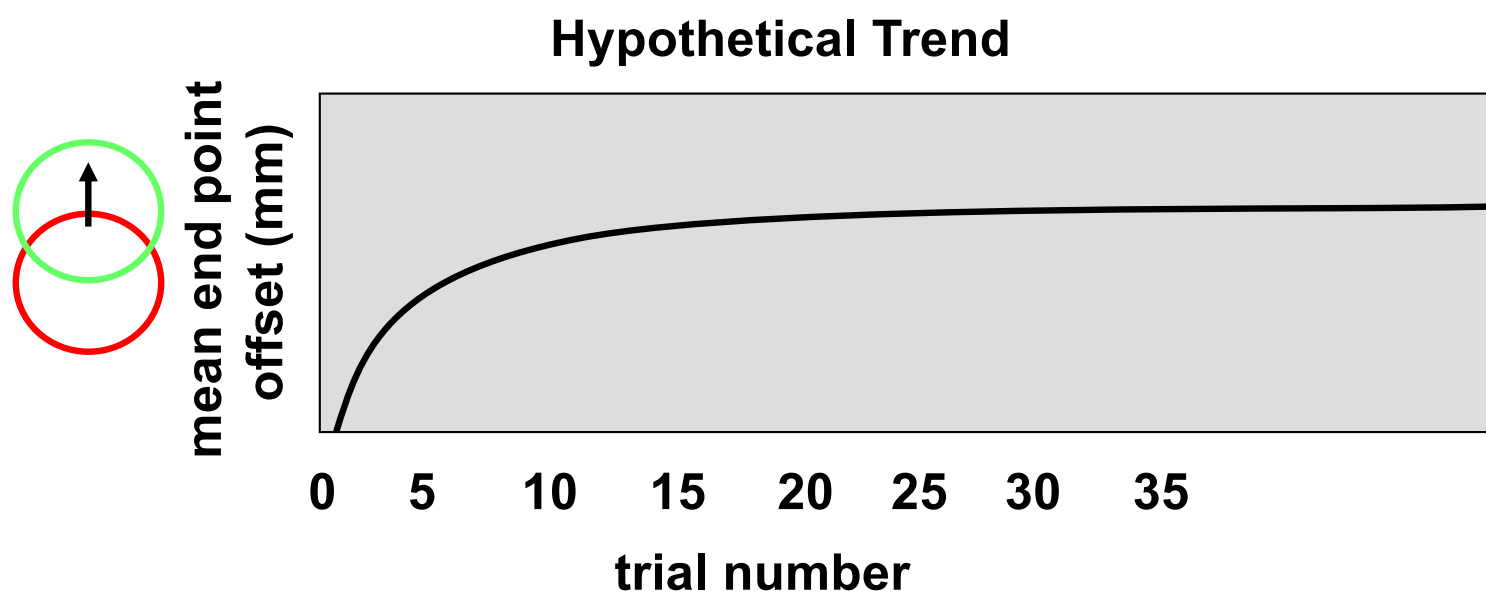
Results: Experiment 1



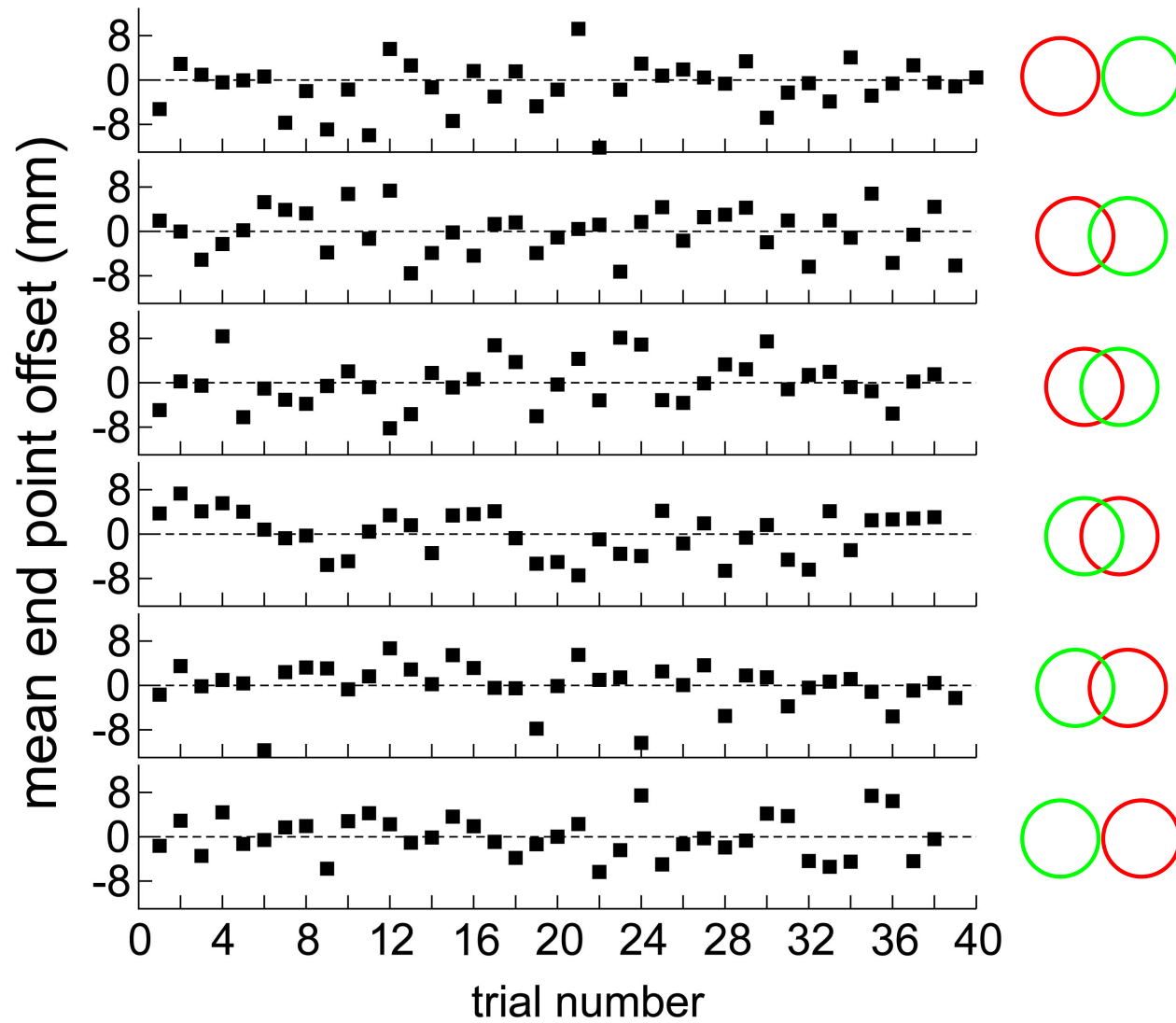
C



b

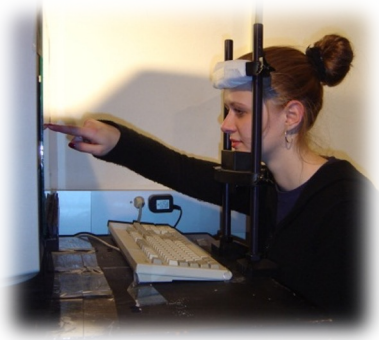


a



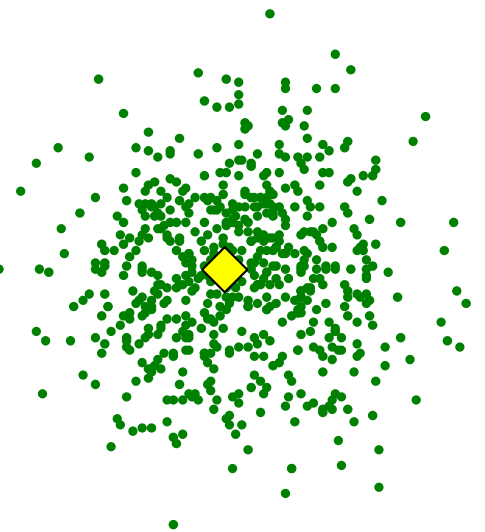
*No sequential effects (or learning).
Seem to know correct offset.*

Action selection

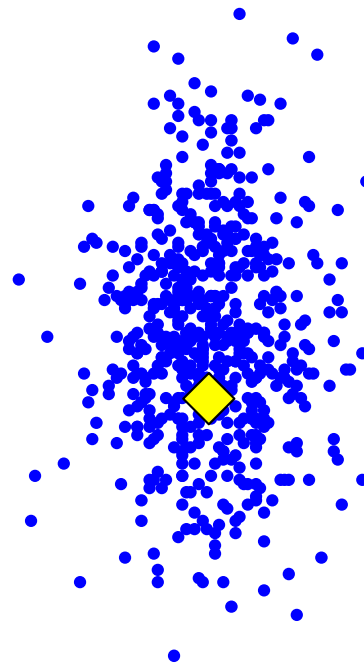
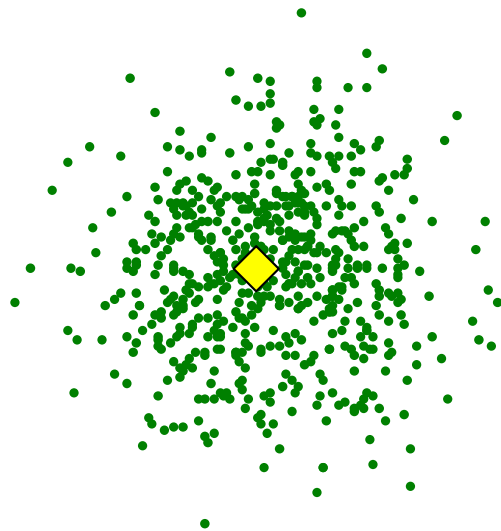
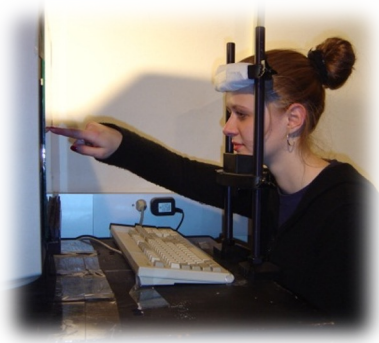


Our performance is close to optimal in the reaching and timing tasks considered so far.

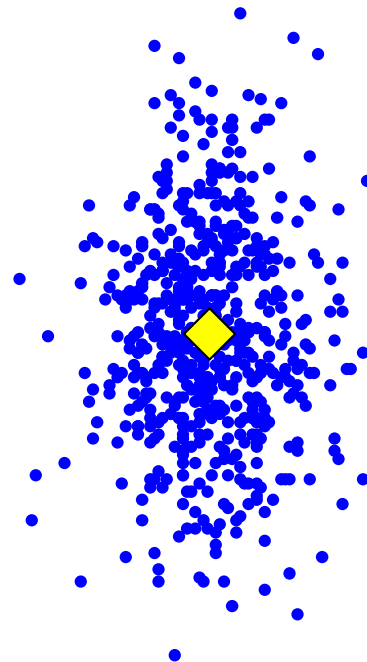
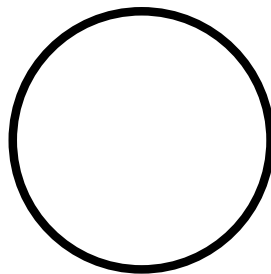
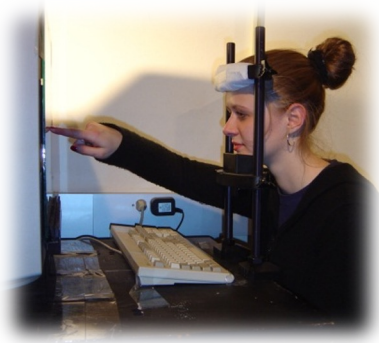
A thought problem: what if reaching error were *not* isotropic (round)?



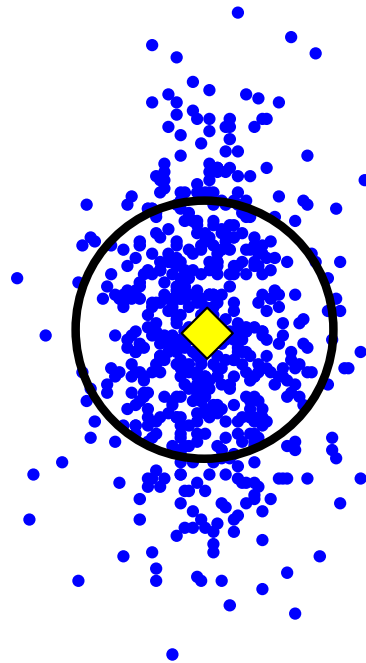
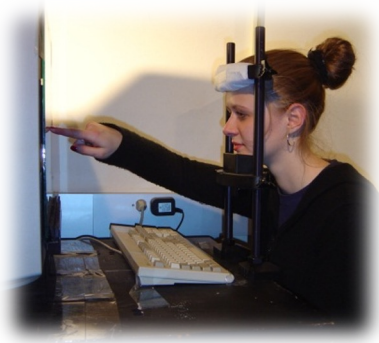
Action selection



Action selection



Action selection

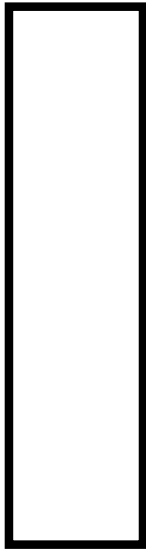


*Anisotropy doesn't affect
the optimal aim point in this
simple task.*

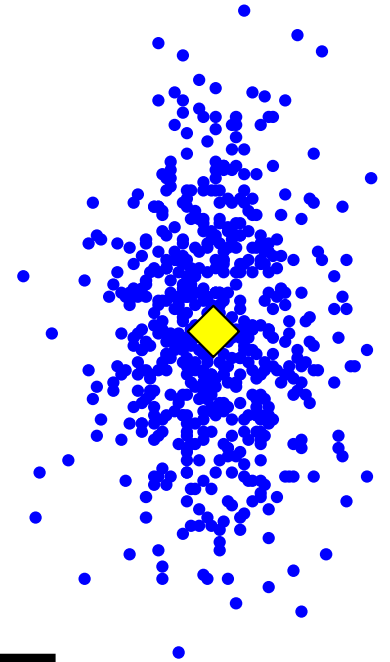
Add a gain function...

\$100

Which target would you like to try?

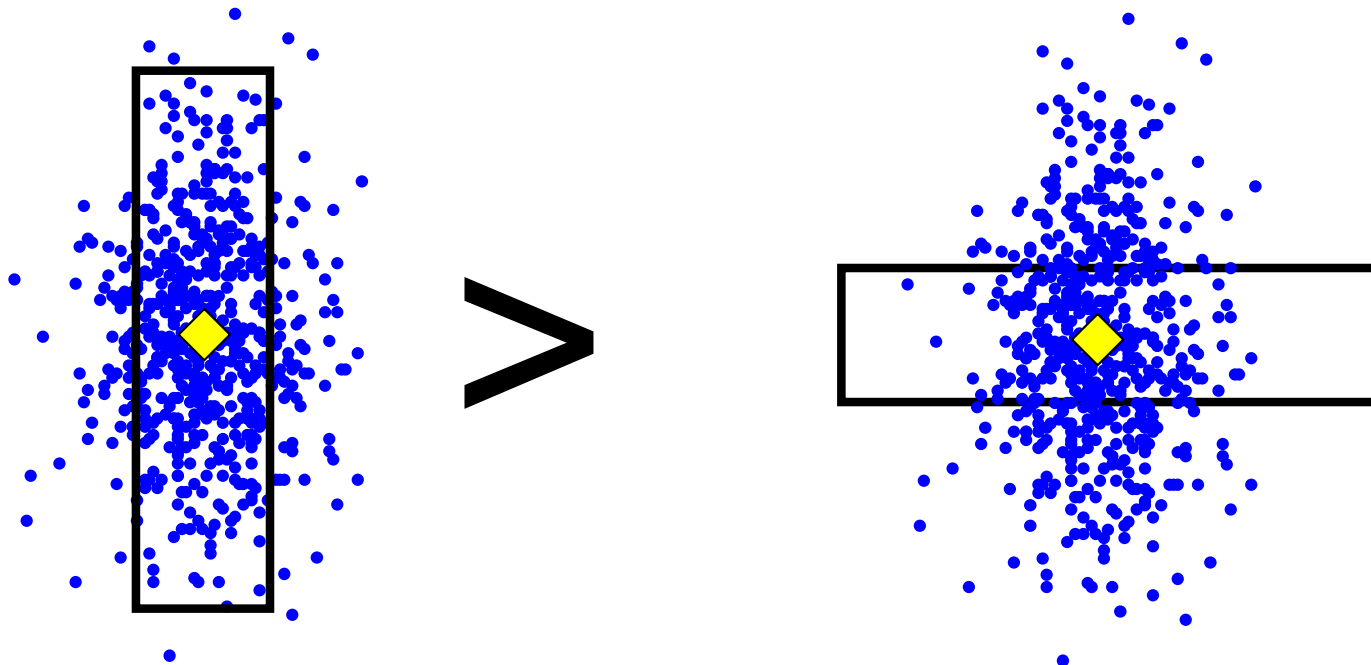


?



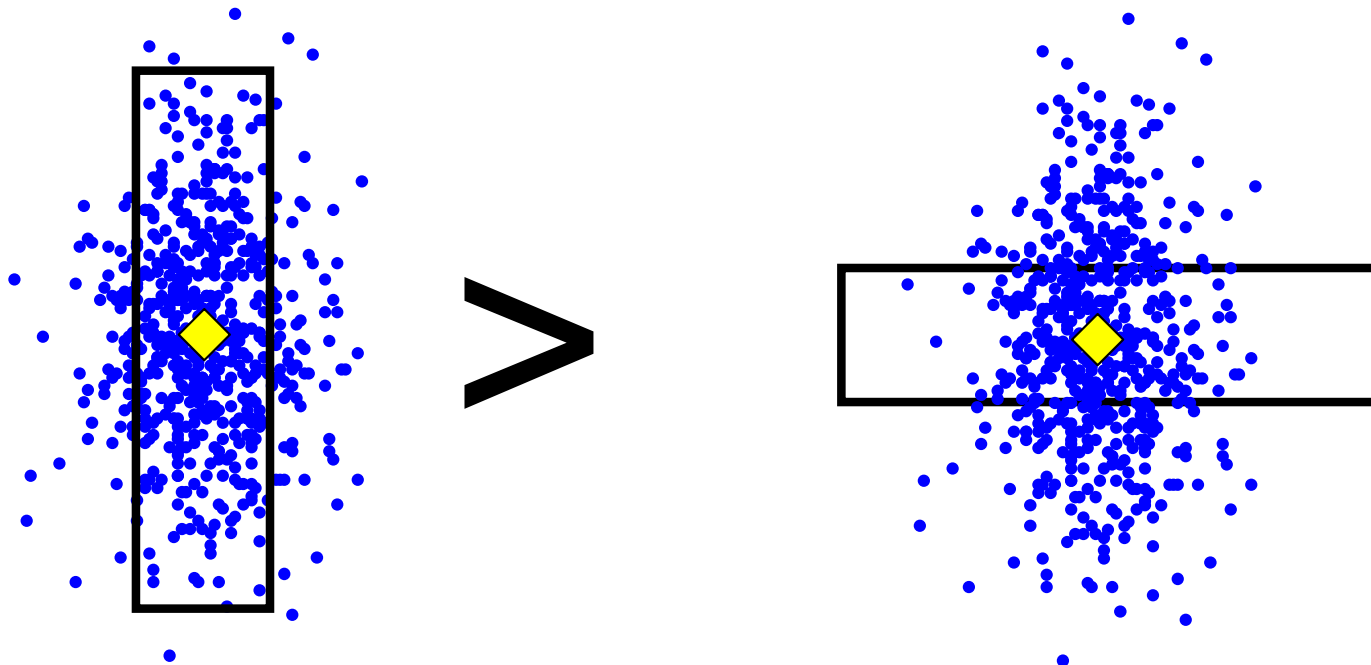
\$100

Which target would you like to try?



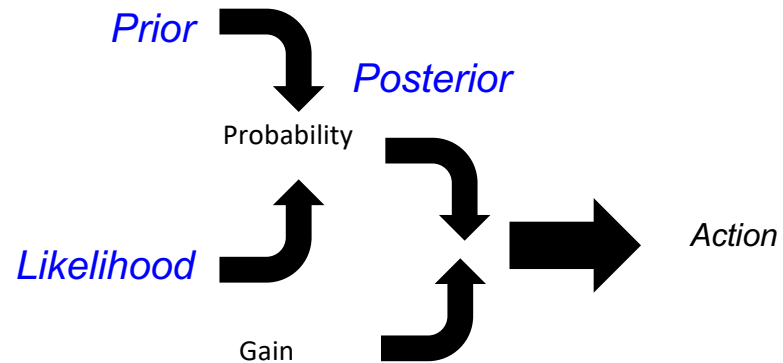
\$100

Which target would you like to try?



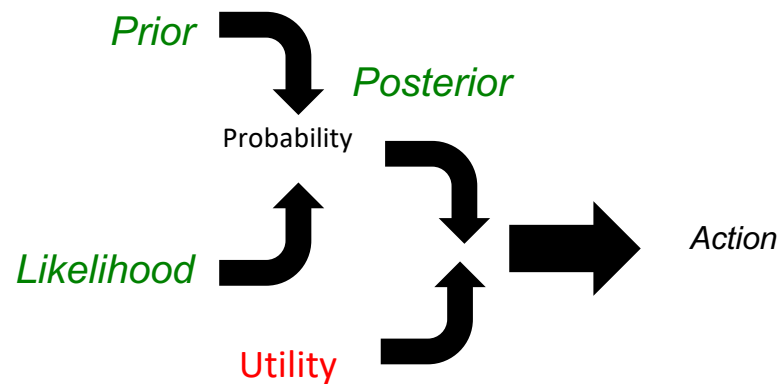
Bayesian

Objective Distributions



We *measure* these

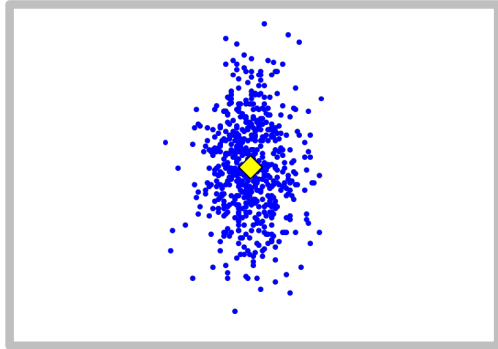
Internal Representations



We *measure* these!

von Neumann & Morgenstern

Representing motor uncertainty



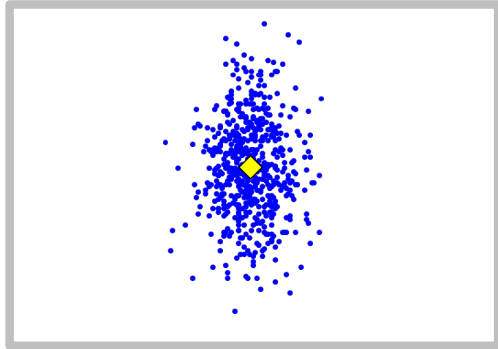
True distribution



Subject's representation

*Probability density function
pdf*

Representing motor uncertainty



True distribution



Subject's representation

*Probability density function
pdf*

Movement planning: near optimal

Trommershäuser, Maloney, & Landy (2003), Spat Vis

Trommershäuser, Maloney, & Landy (2003), J Opt Soc Am A

Körding & Wolpert (2004), Nature

Najemnik & Geisler (2005), Nature

Trommershäuser, Gepshtein, Maloney, Landy, & Banks (2005), J Neurosci

Trommershäuser, Mattis, Landy, & Maloney (2006), Exp Brain Res

Trommershäuser, Landy, & Maloney (2006), Psych Sci

Battaglia & Schrater (2007), J Neurosci

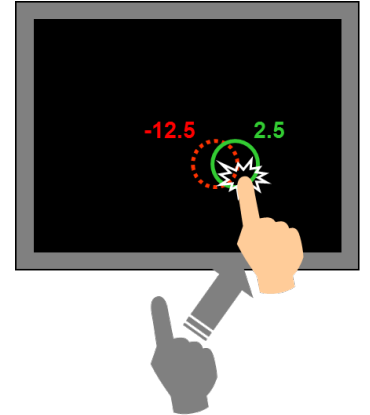
Dean, Wu, & Maloney (2007), JOV

Hudson, Maloney, & Landy (2008), PLoS Comp Biol

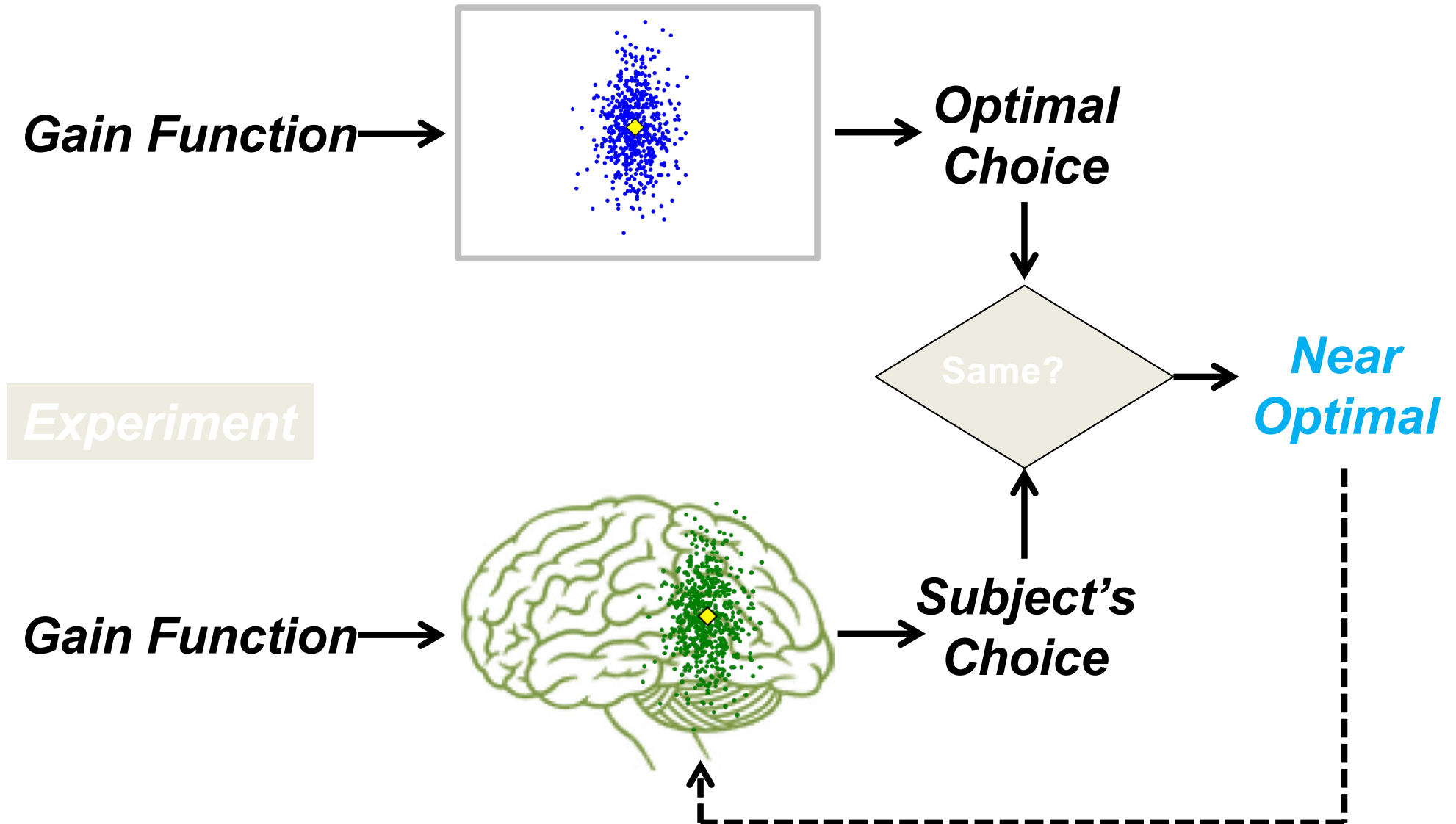
Faisal & Wolpert (2009), J Neurophysiol

Wei & Körding (2010), Front Comput Neurosci

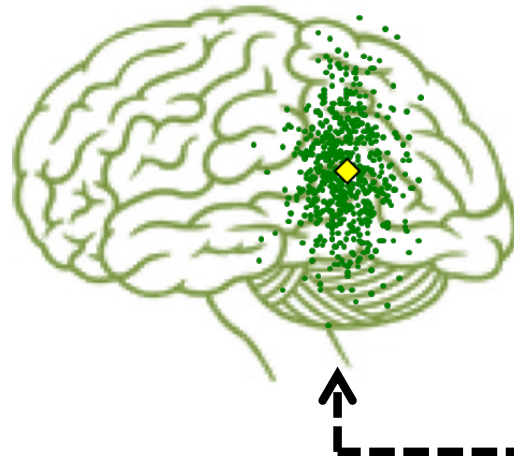
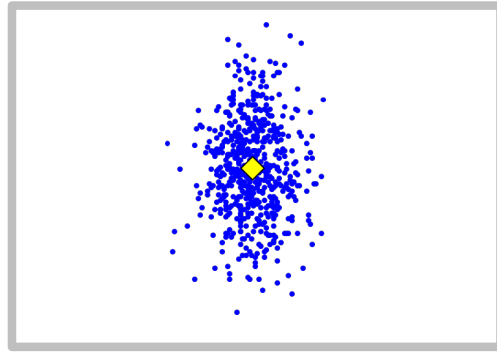
.....



Bayesian computation



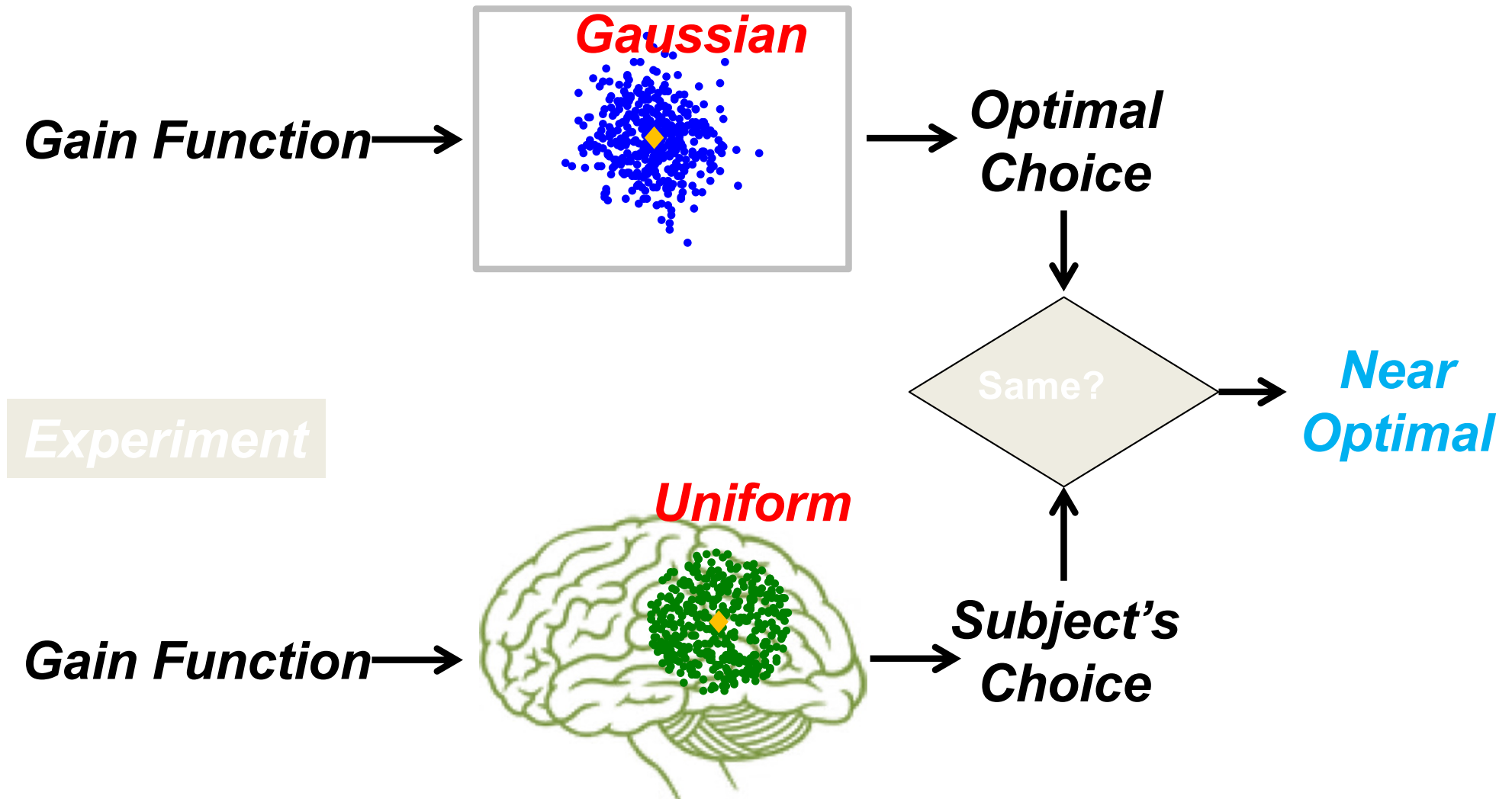
Bayesian computation



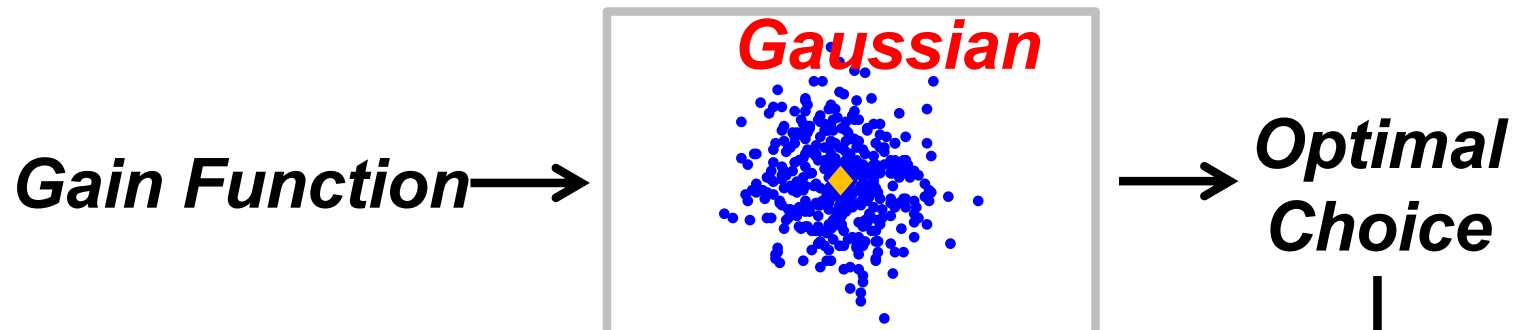
*Near
Optimal*



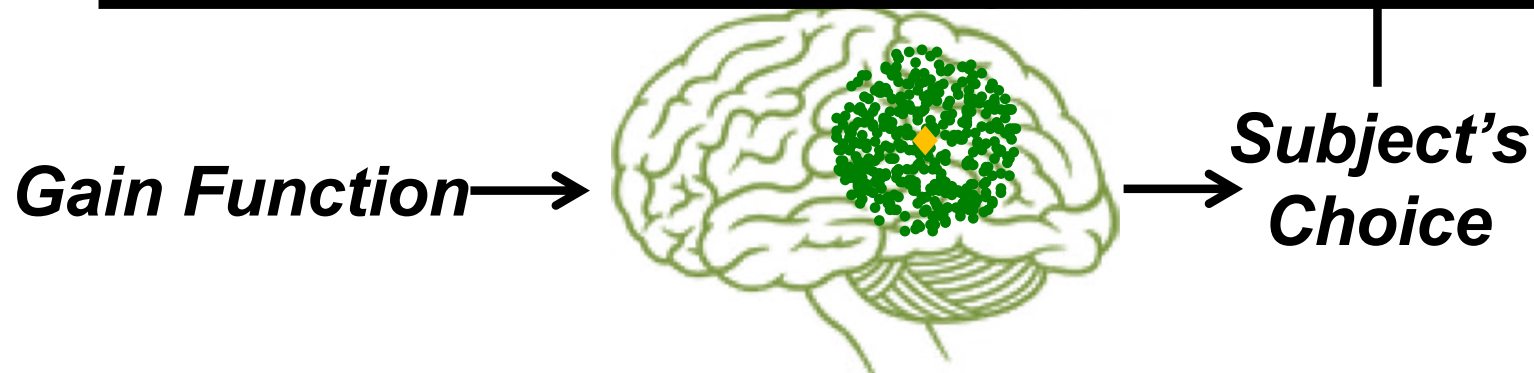
Bayesian computation



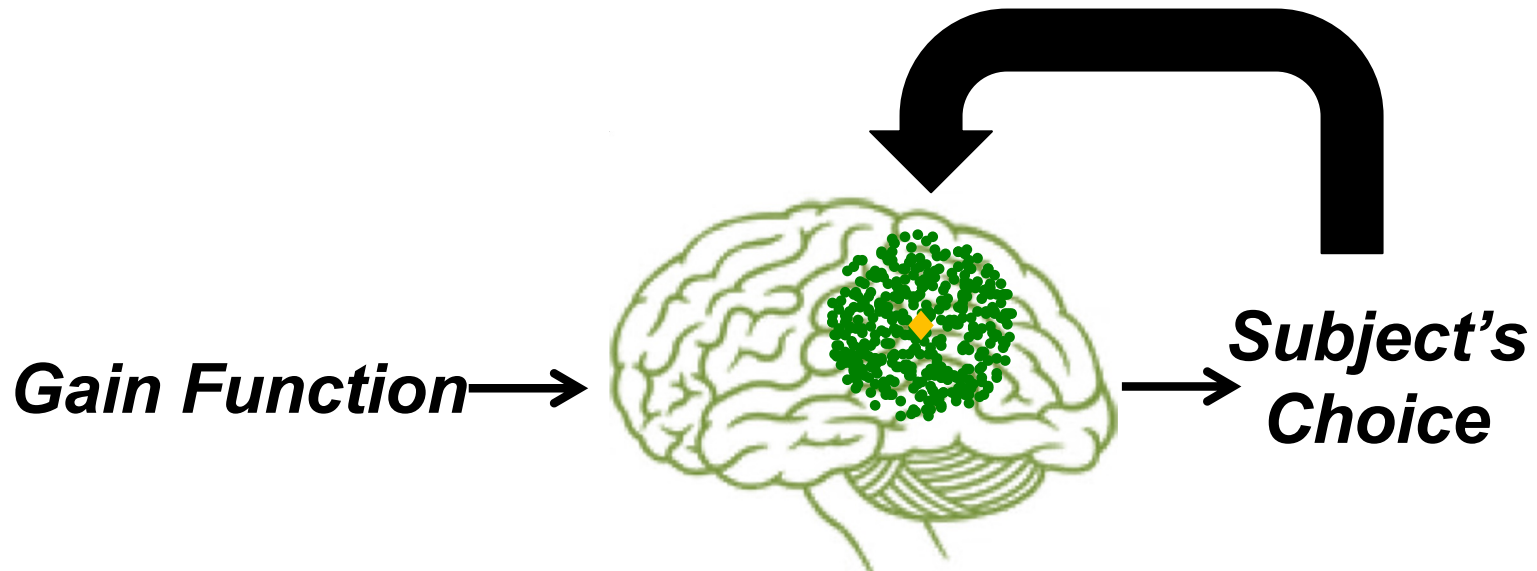
Bayesian computation



Many tasks may simply be insensitive to systematic deviations in the internal model of uncertainty



The inverse problem

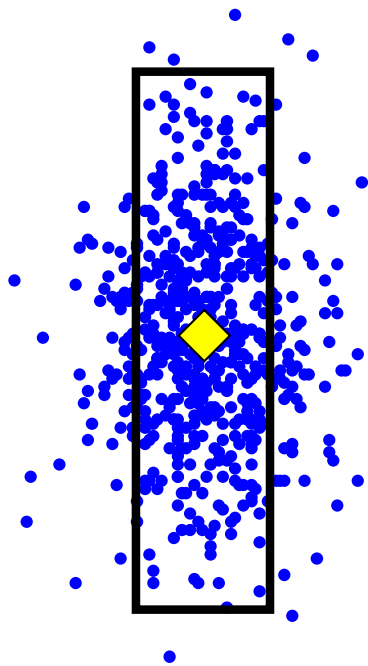


Inferring people's internal models of uncertainty based on their choices

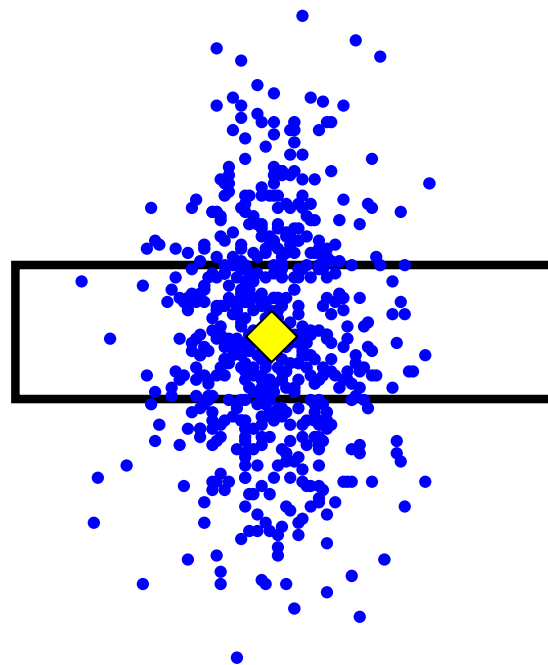
Choice Task

\$100

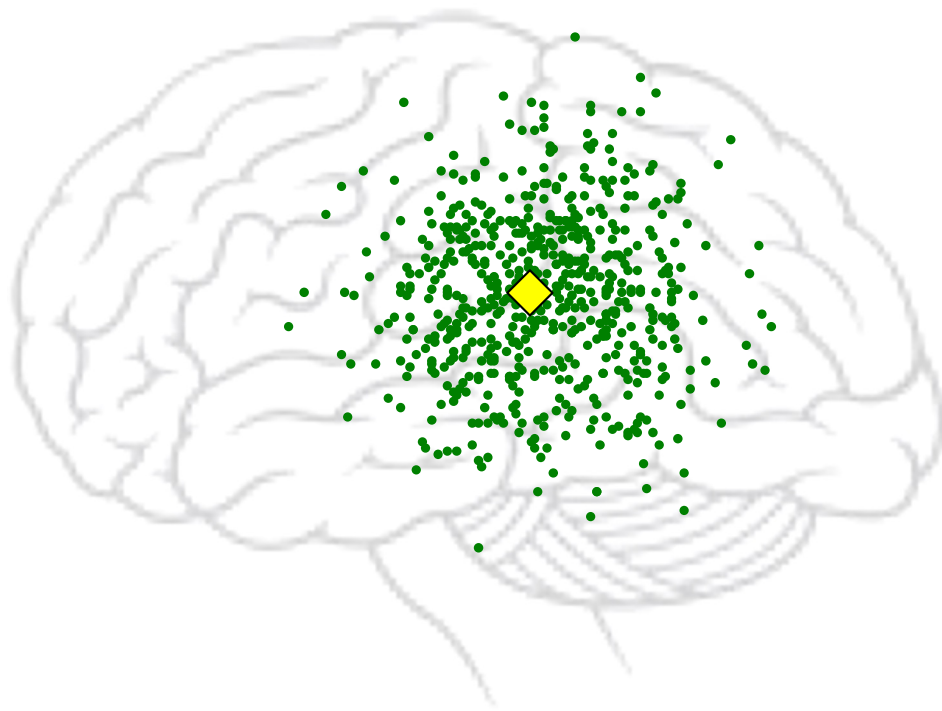
Which target would you like to try?



>



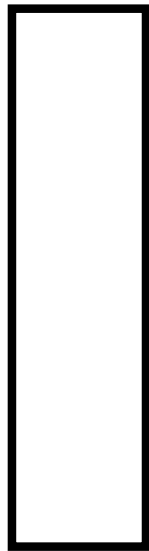
Bayesian computation



Choice task

\$100

Which target would you like to try?



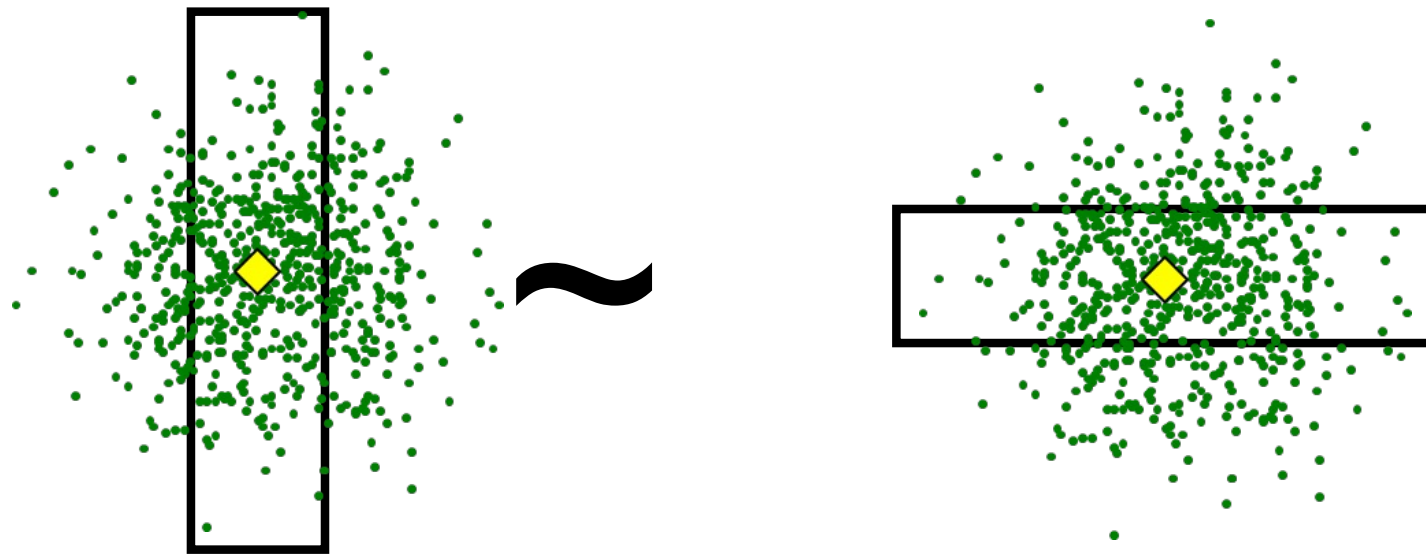
?

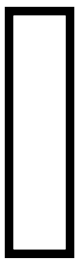
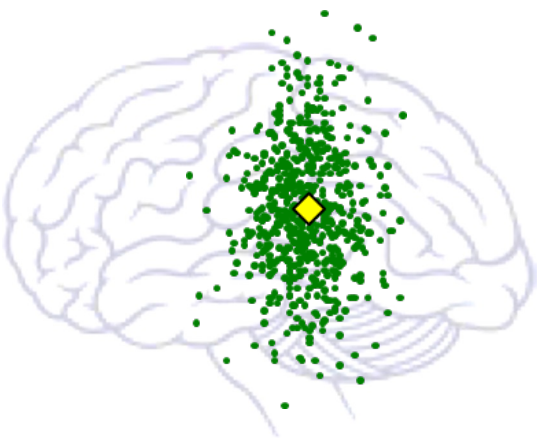


Bayesian computation

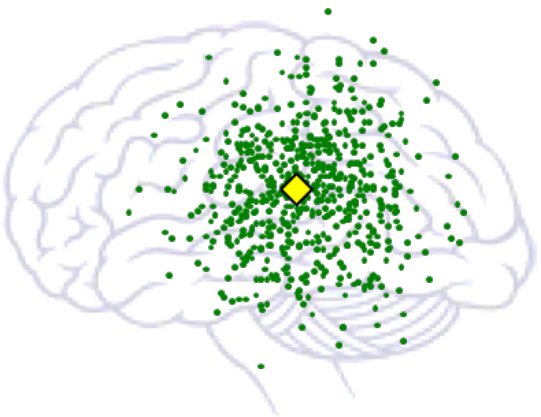
\$100

Which target would you like to try?

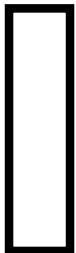
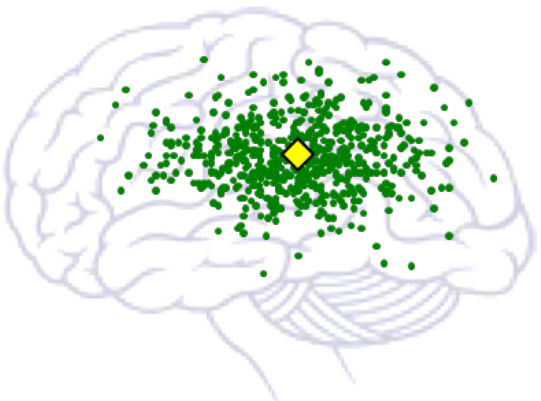




$>$

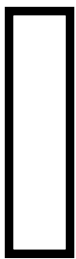
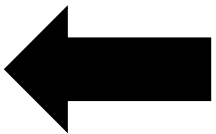
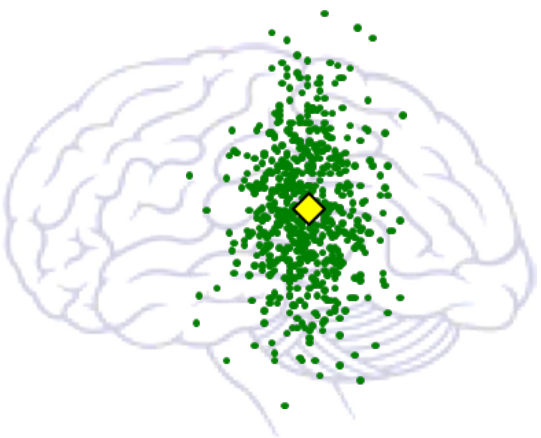


$\approx$

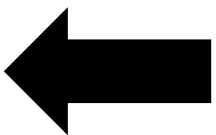
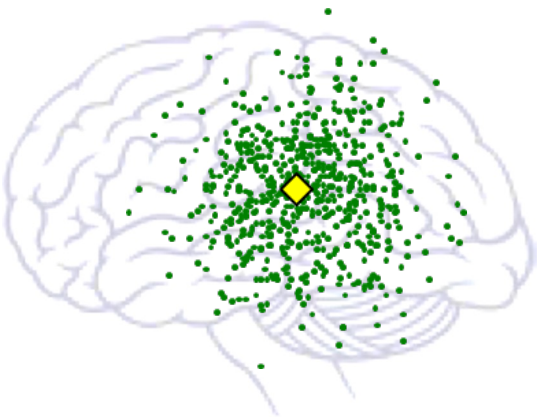


$>$

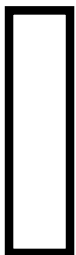
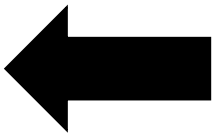
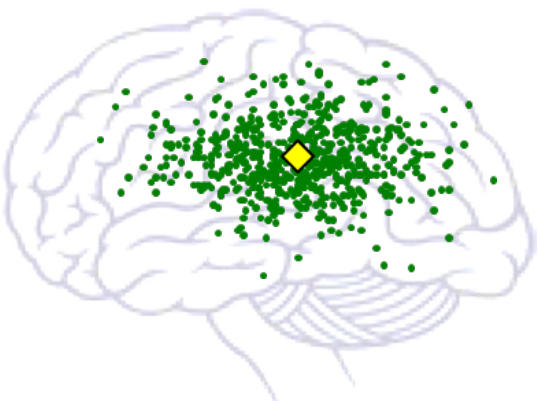




$>$



$\sim$



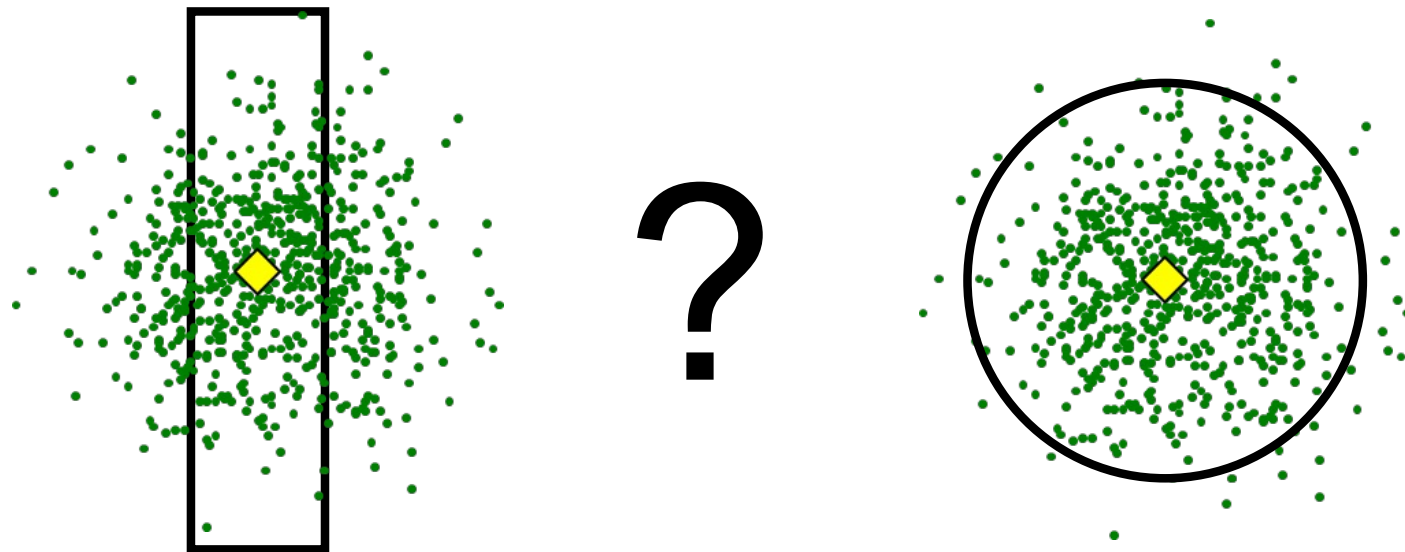
$>$



Experiment

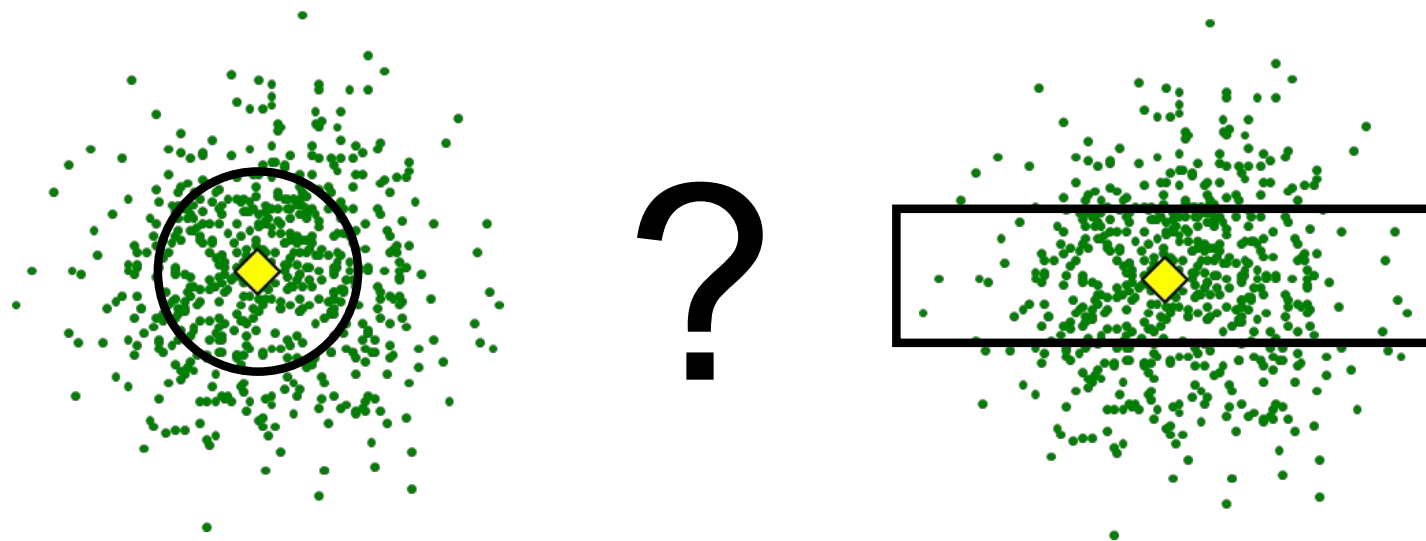
Zhang, Daw, Maloney (2013) PLoS CB

Measuring subjects' choices between targets of varying shapes and sizes allows us to infer their internal model of motor uncertainty



Riesz-Fischer theorem, see Maloney & Mamassian, 2009

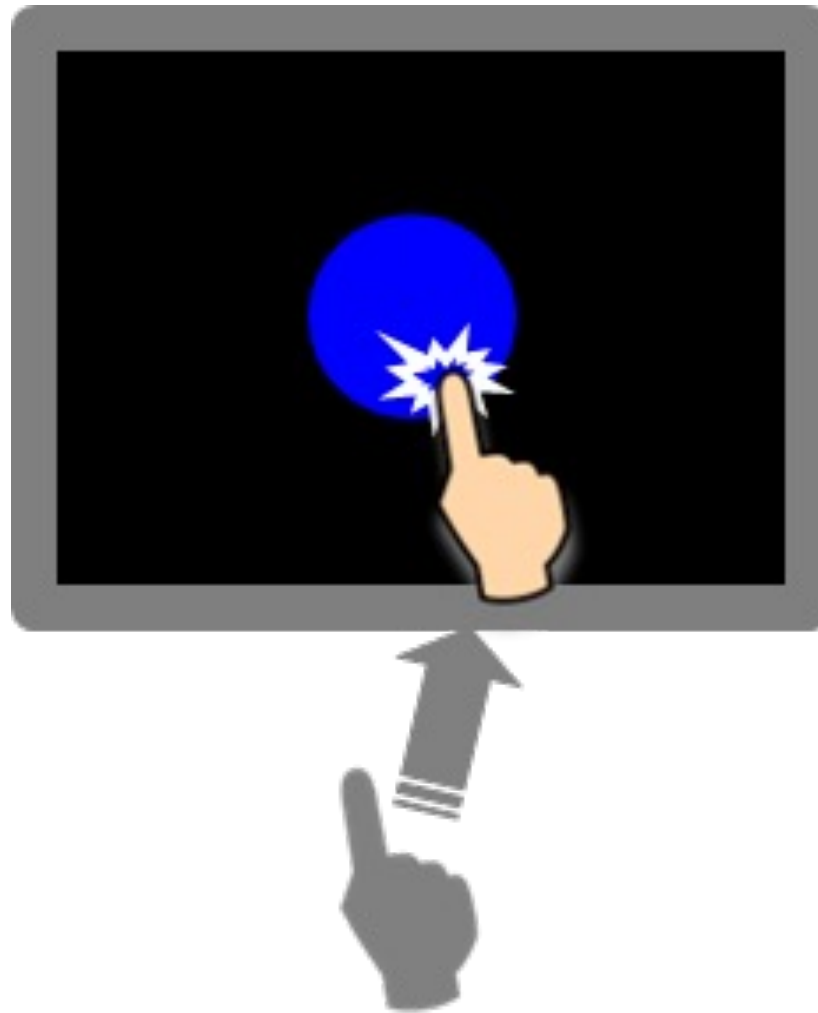
Measuring subjects' choices between targets of varying shapes and sizes allows us to infer their internal model of motor uncertainty



Riesz-Fischer theorem, see Maloney & Mamassian, 2009

Zhang et al (2013)

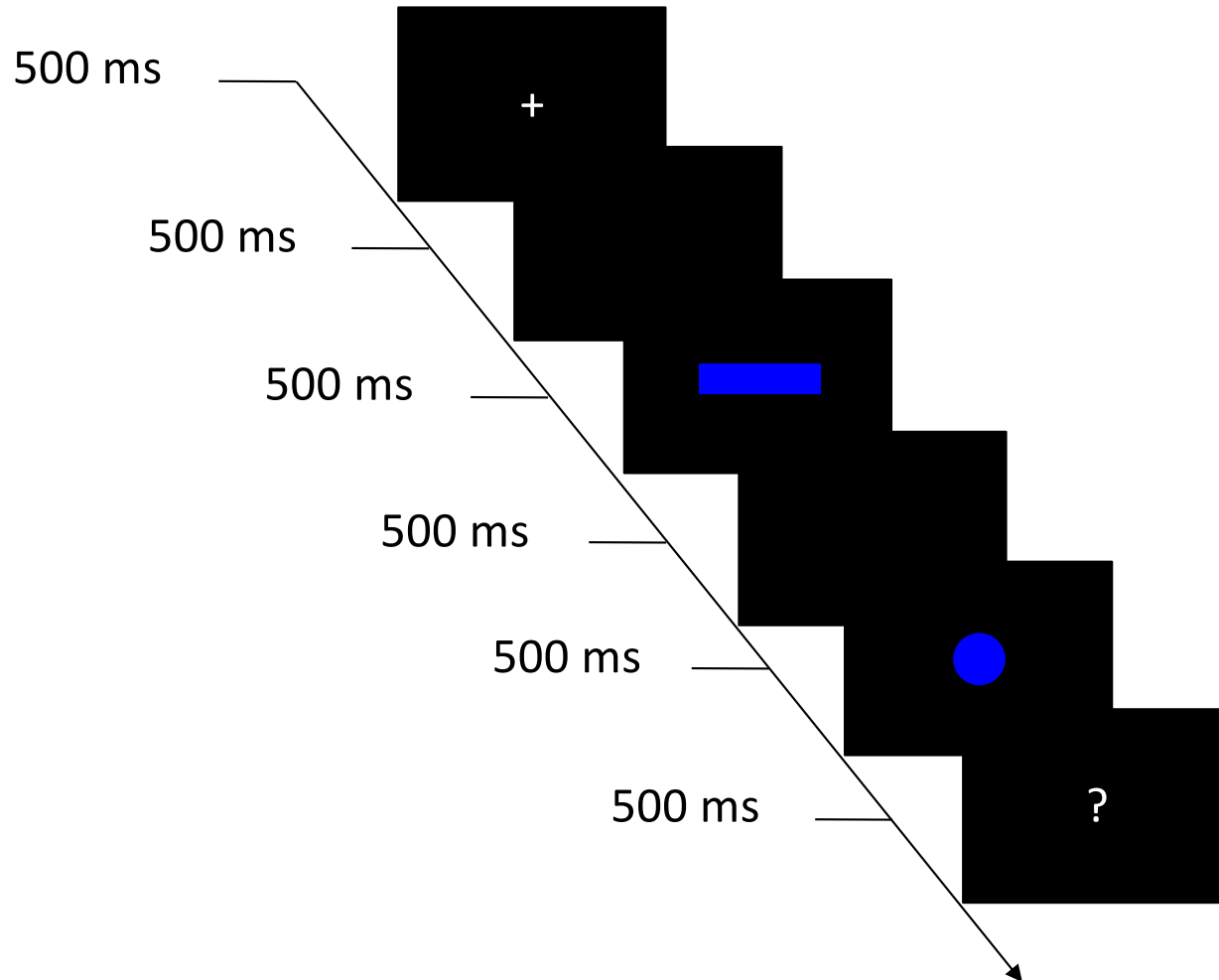
Procedure



Touch the target within
400 msec

300 trials

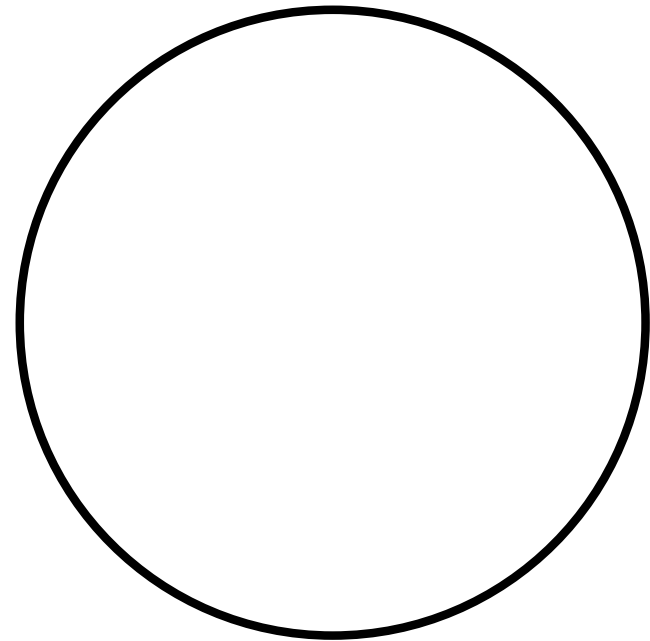
Which is easier to hit? 1st or 2nd?



Which is easier to hit?



?



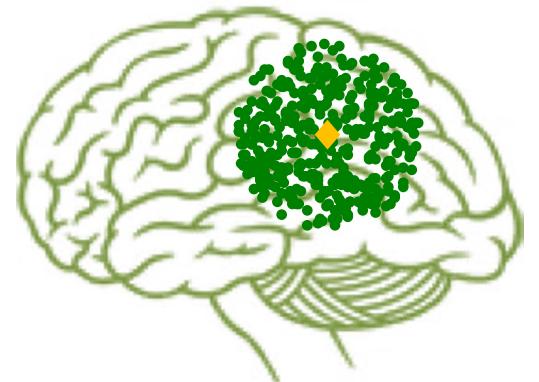
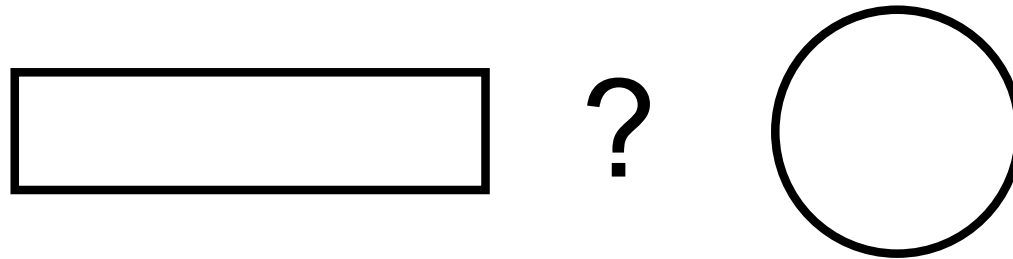
Definitely the circle

Which is easier to hit?

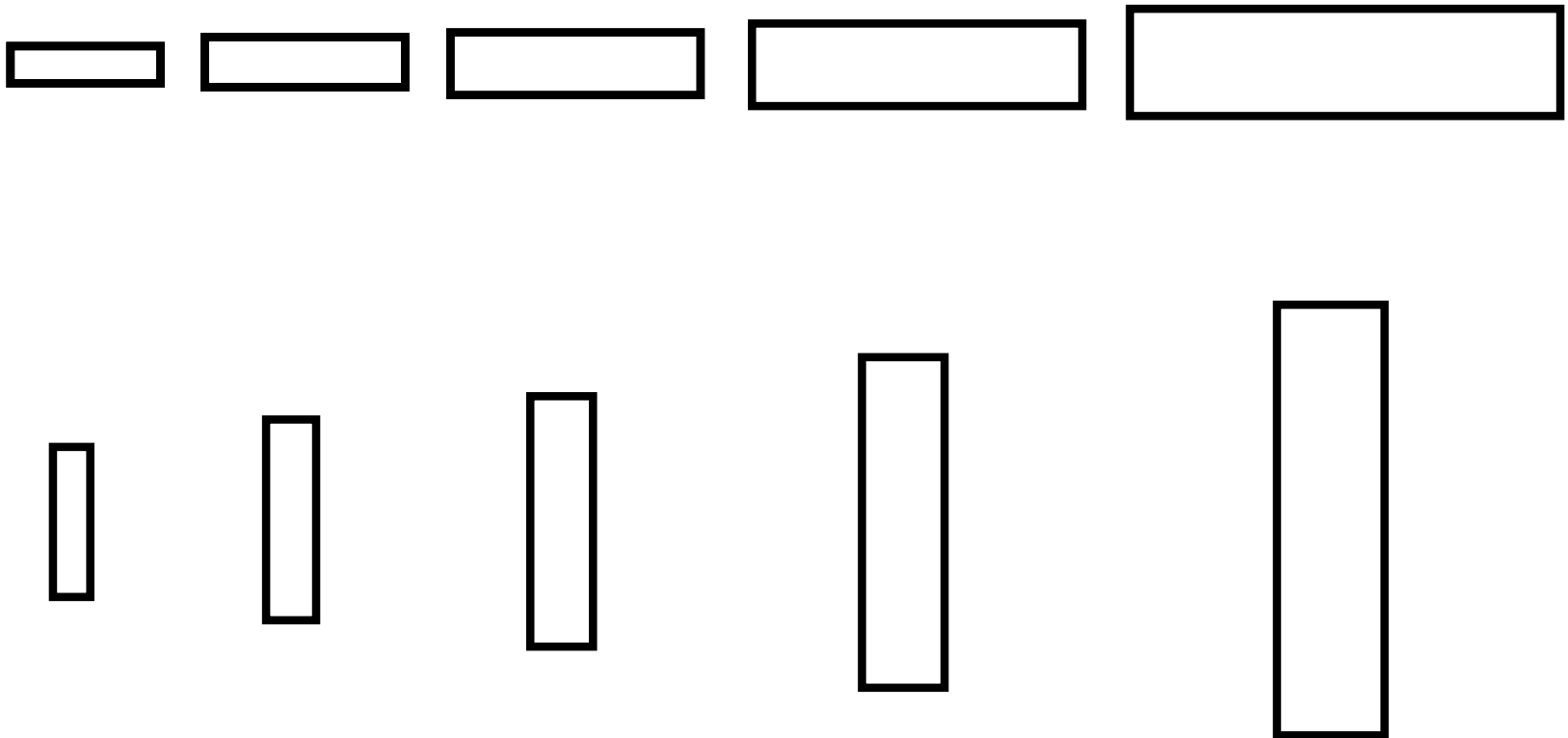


Definitely the rectangle

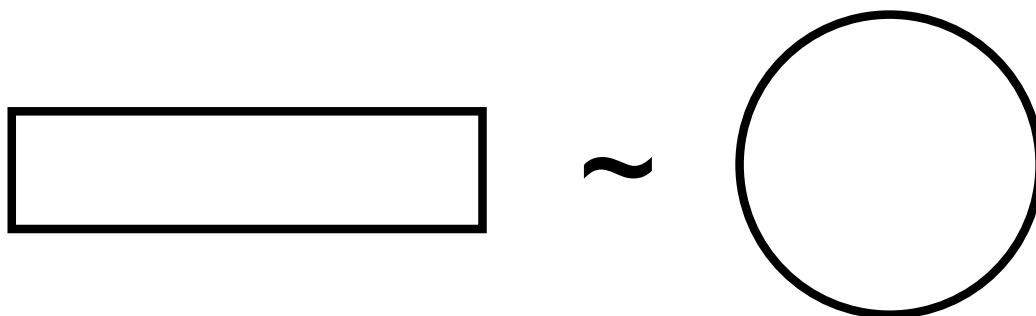
Which is easier to hit?

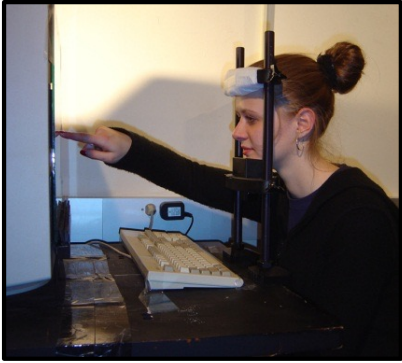


10 possible rectangles



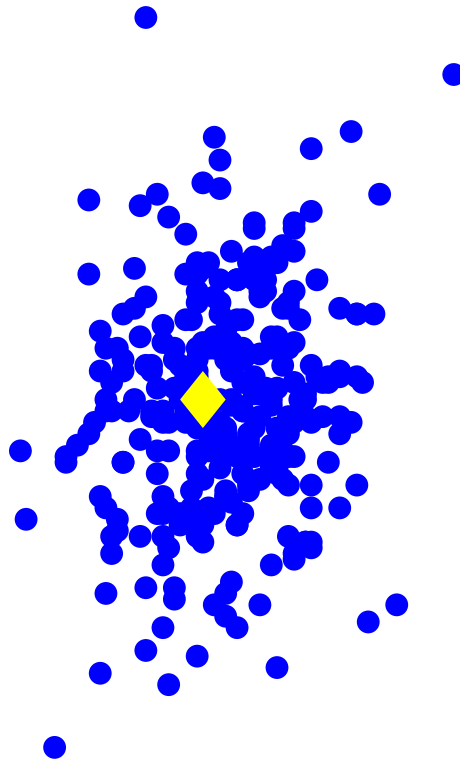
Radius of the circle was adjusted by adaptive procedures
(staircase)

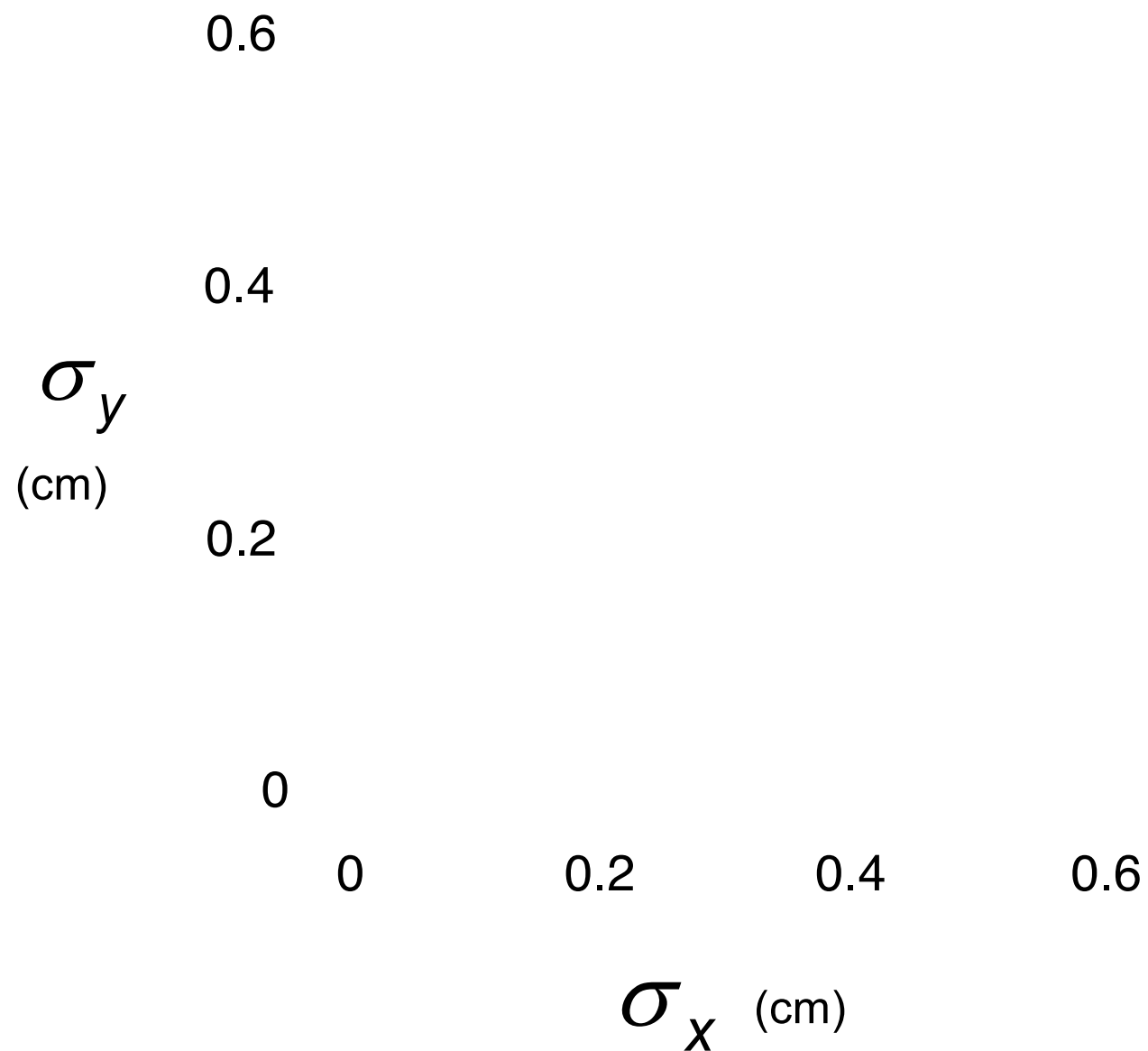


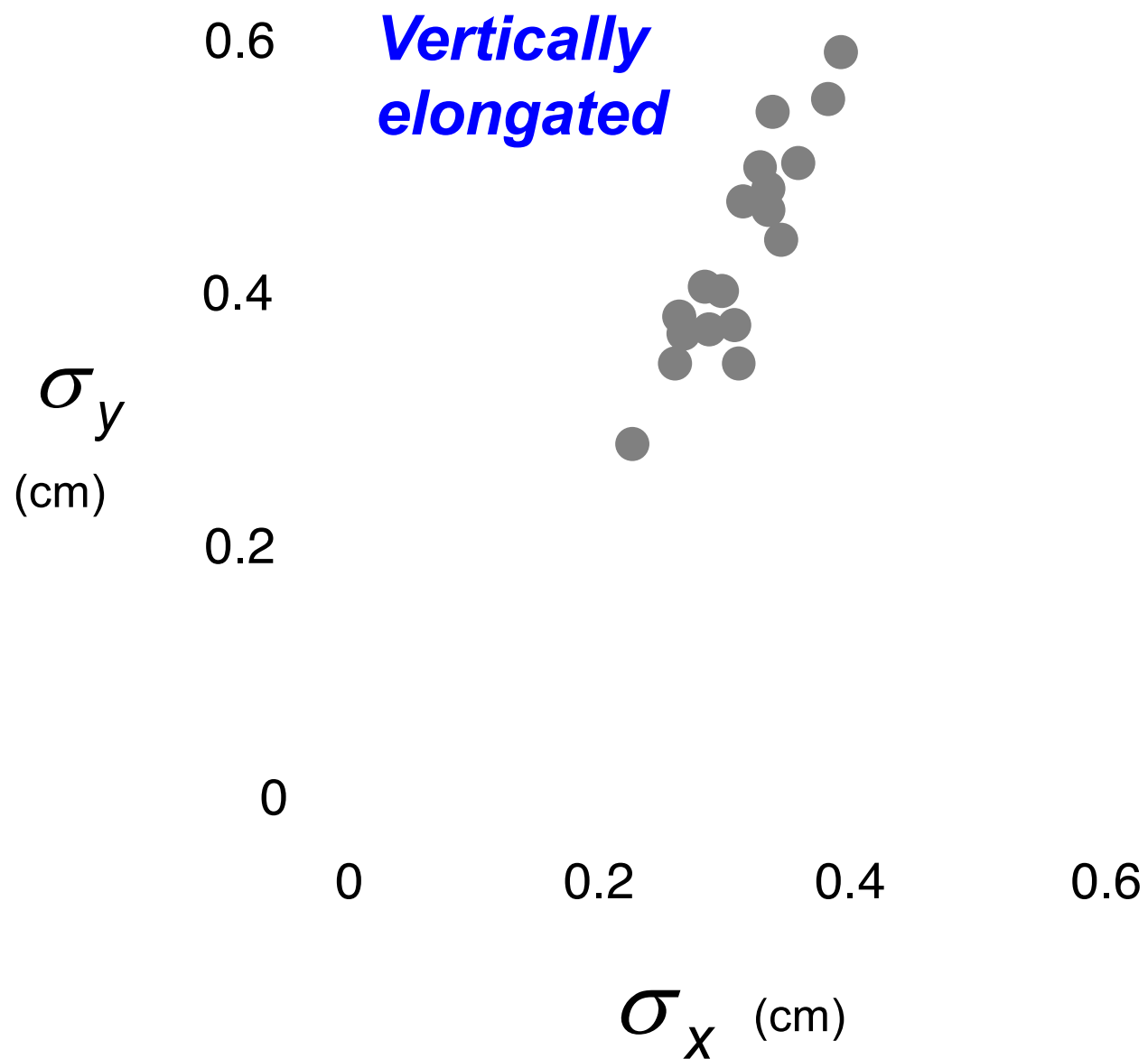


Results: True error distribution

***Vertically elongated,
bivariate Gaussian***

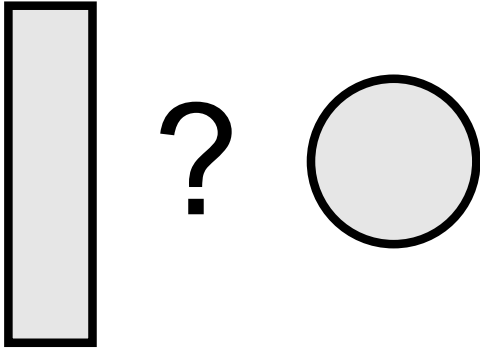






Median

$$\sigma_y / \sigma_x = 1.44$$



Results: Subjects' internal model of their own error distribution

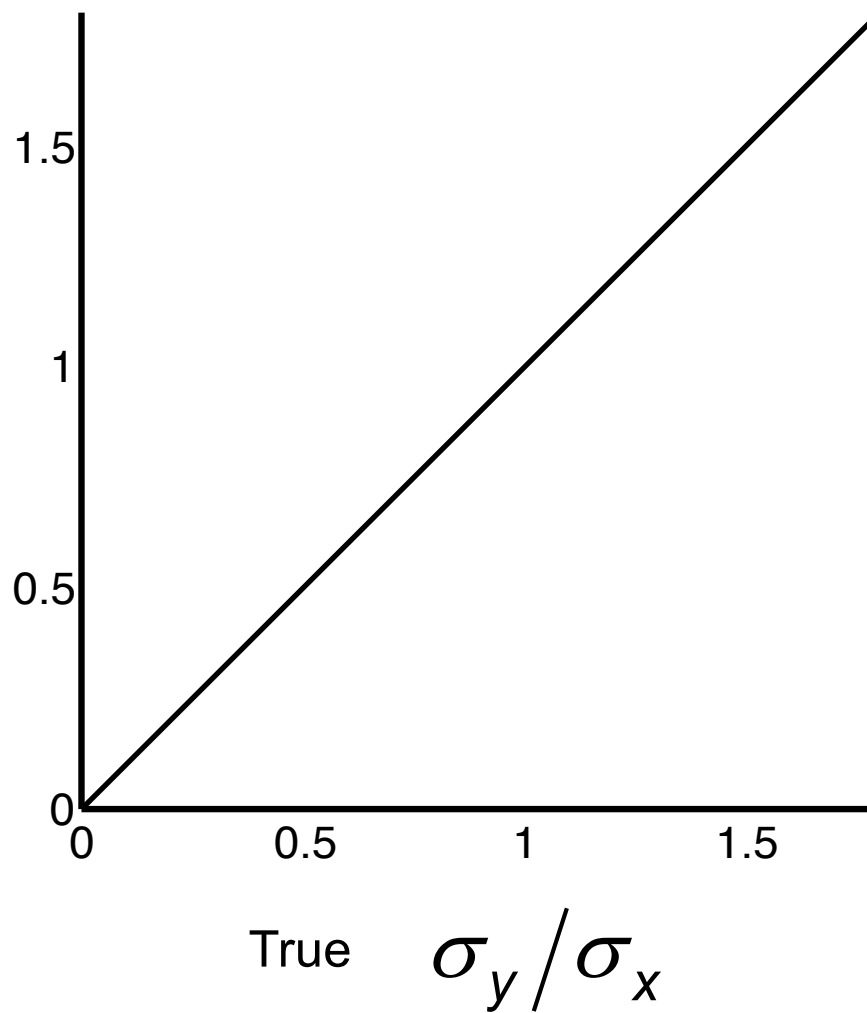


Each subject's internal model was fitted to the subject's choices as a bivariate Gaussian distribution with two free parameters:

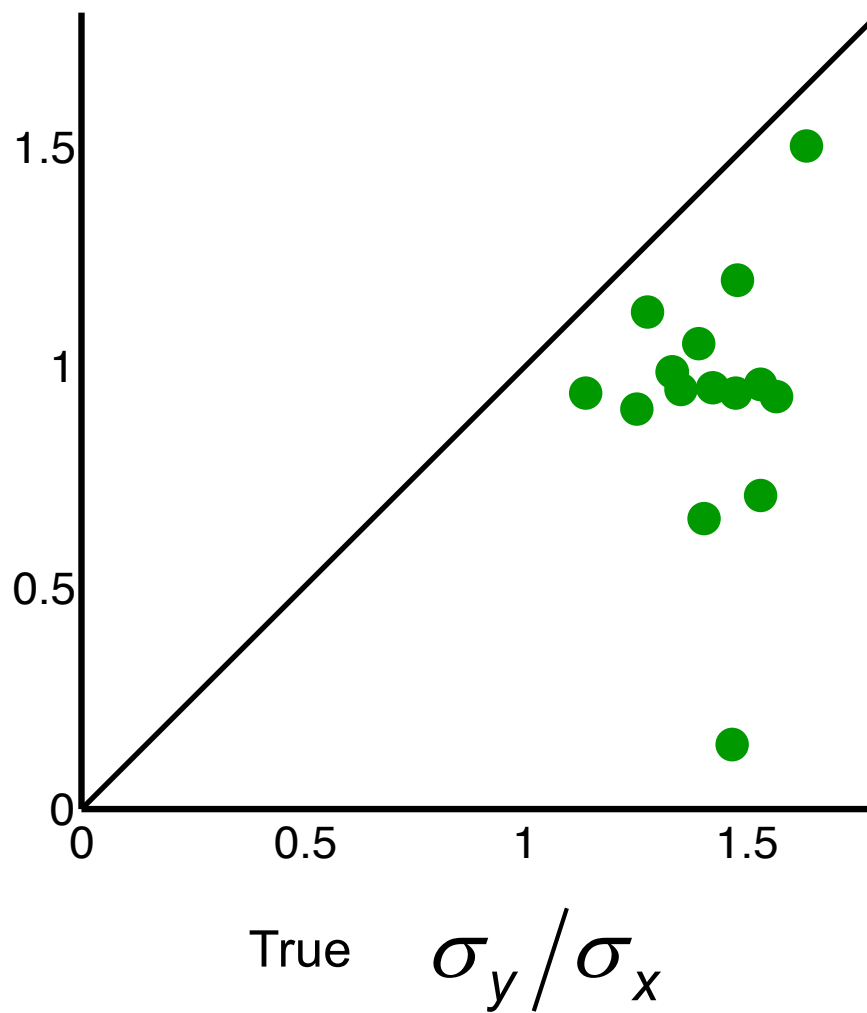
$$\sigma_x' \quad \text{and} \quad \sigma_y'$$

Maximum likelihood estimates

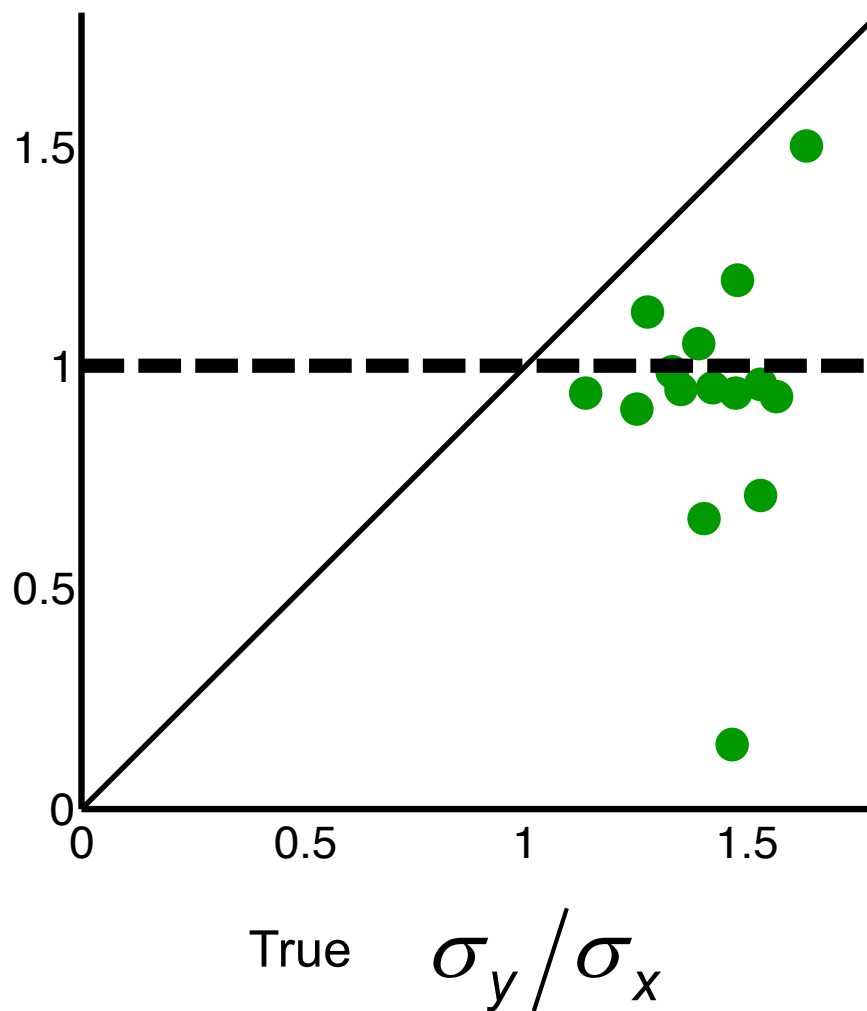
$\sigma_y' / \sigma_x^{\text{in}'}$
Subjects' Model



$\sigma_y' / \sigma_x^{\text{in}'}$
Subjects' Model

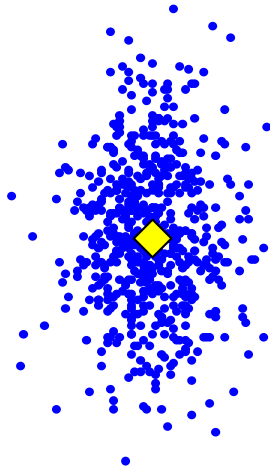


$\sigma_y' / \sigma_x^{\text{in}'}$
Subjects' Model

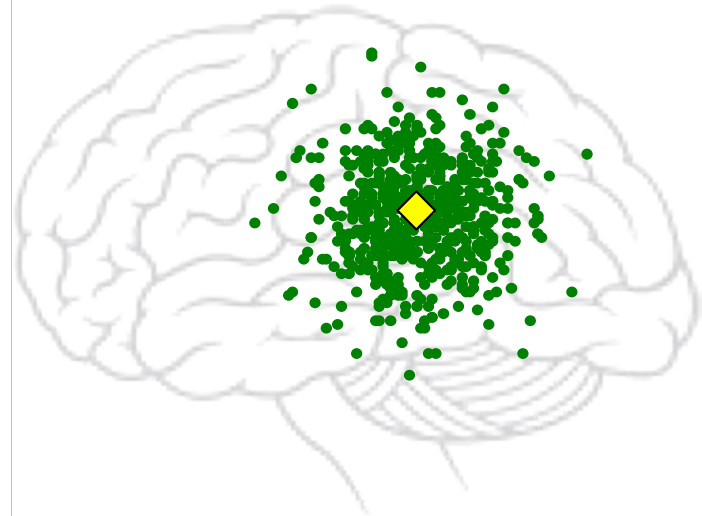


Summary

true distribution



subject's model



Further Readings

Zhang, H. , Ren, X. J. & Maloney, L. T. (2020), The bounded rationality of probability distortion. *Proceedings of the National Academy of Sciences, USA*, 1-11. <http://www.pnas.org/cgi/doi/10.1073/pnas.1922401117>

Zhang, H., Daw, N. & Maloney, L. T. (2013), Testing whether humans have an accurate model of their own motor uncertainty in a speeded reaching task. *PLoS Computational Biology*, 9(5):e1003080, 1-11.

Hudson, T. E., Wolfe, U. & Maloney, L. T. (2012), Speeded reaching movements around invisible obstacles. *PLoS Computational Biology*, 8(9), e1002676, 1-9.

Maloney, L. T. & Mamassian, P. (2009), Bayesian decision theory as a model of visual perception: Testing Bayesian transfer. *Visual Neuroscience*, 26, 147-155.

Thank you!