

# MC2019 Part 1

## Optimal estimation: Maximum likelihood and Bayesian estimation.

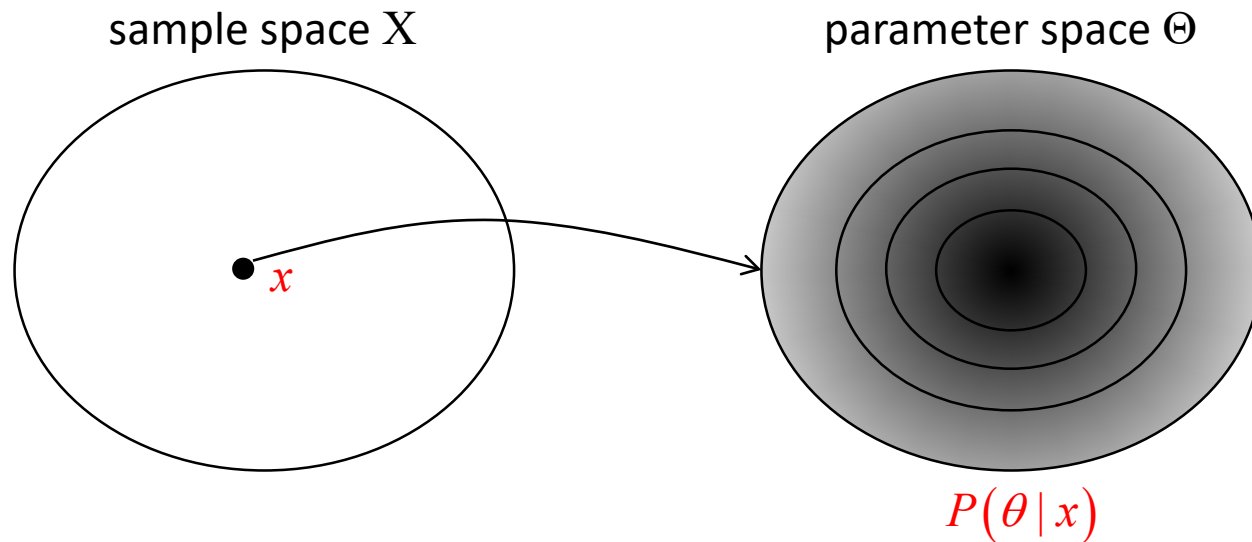
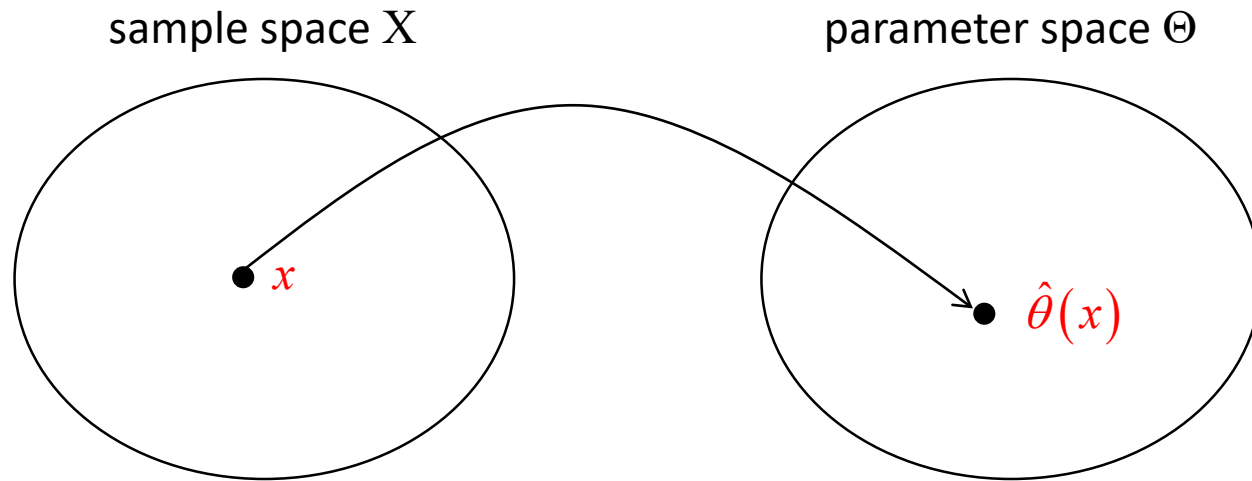
Hirokazu Tanaka

Estimation theory provides a theoretical framework for inference.

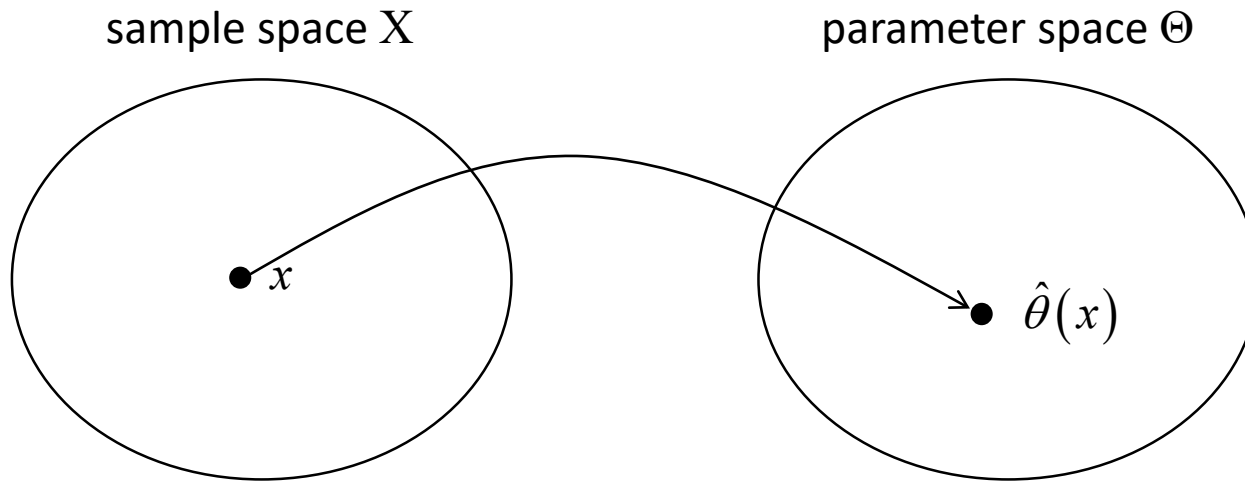
In this lecture, we will learn:

- Maximum-likelihood (ML) estimation
- Cramer-Rao's lower bound
- Bayesian estimation
- Wiener filter
- Kalman filter
- Kalman smoother
- Causal inference
- Information integration
- Postdiction

# Classical estimation (point) and Bayesian estimation (density).



# Maximum-likelihood (ML) estimation.

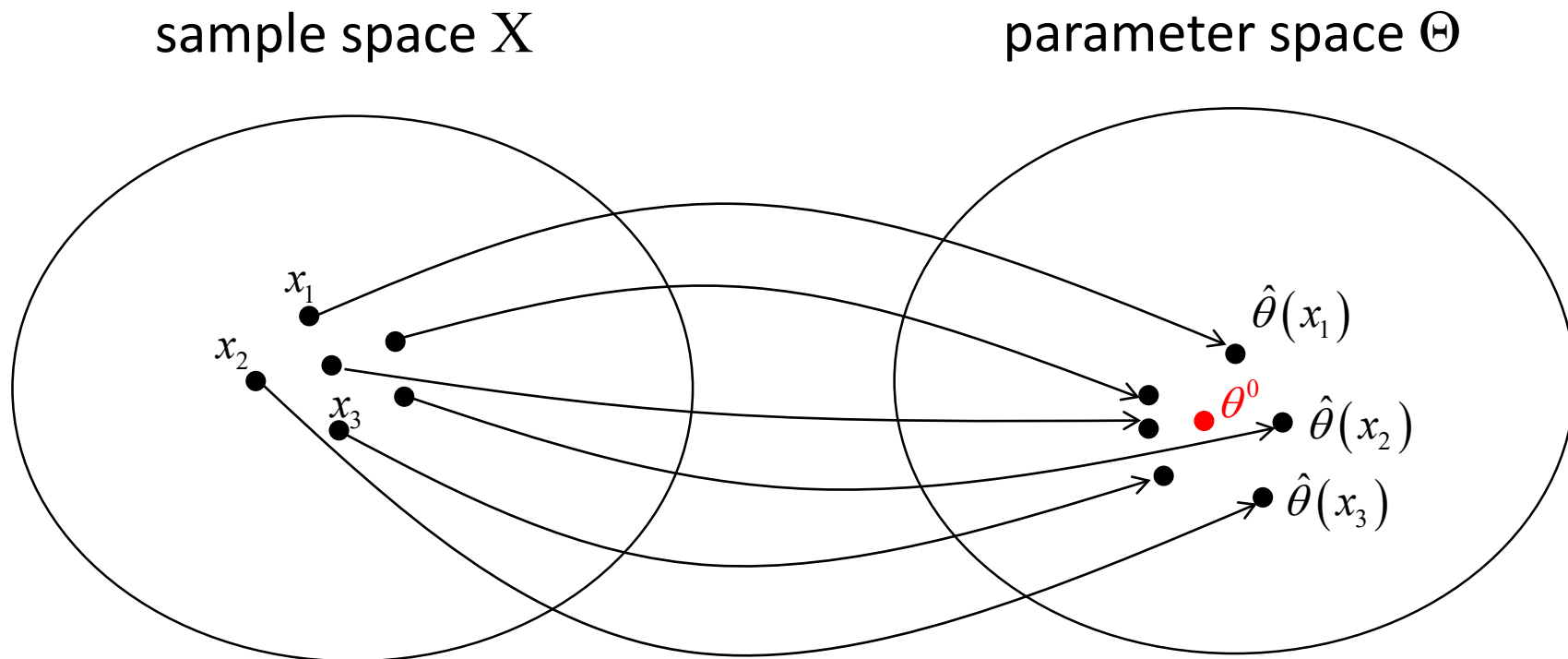


$$\begin{aligned}\hat{\theta}_{\text{ML}}(x) &= \arg \max_{\theta} P(x | \theta) \\ &= \arg \max_{\theta} \log P(x | \theta)\end{aligned}$$

ML is ...

- Asymptotically unbiased (i.e., approaches to a true value).
- Asymptotically efficient (i.e., achieves minimum variance, CRLB).

# Maximum-likelihood (ML) estimation.



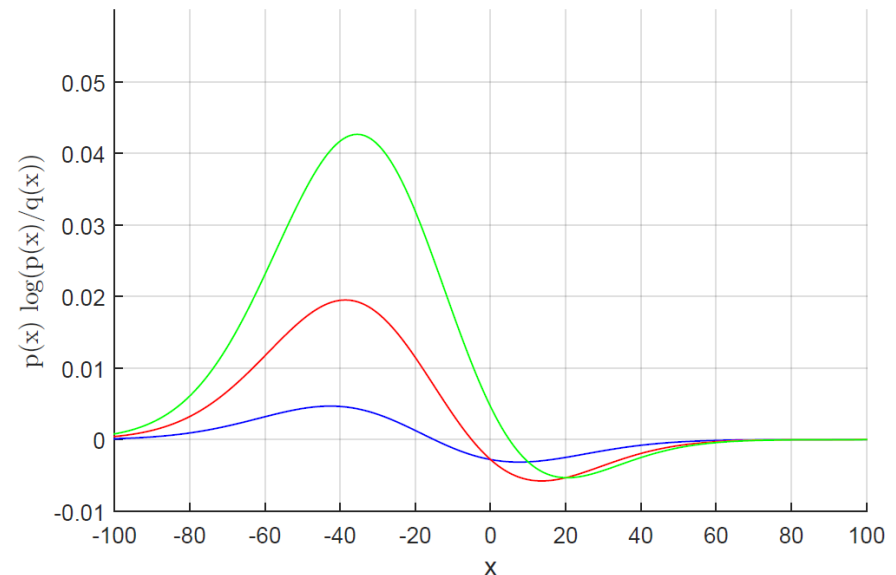
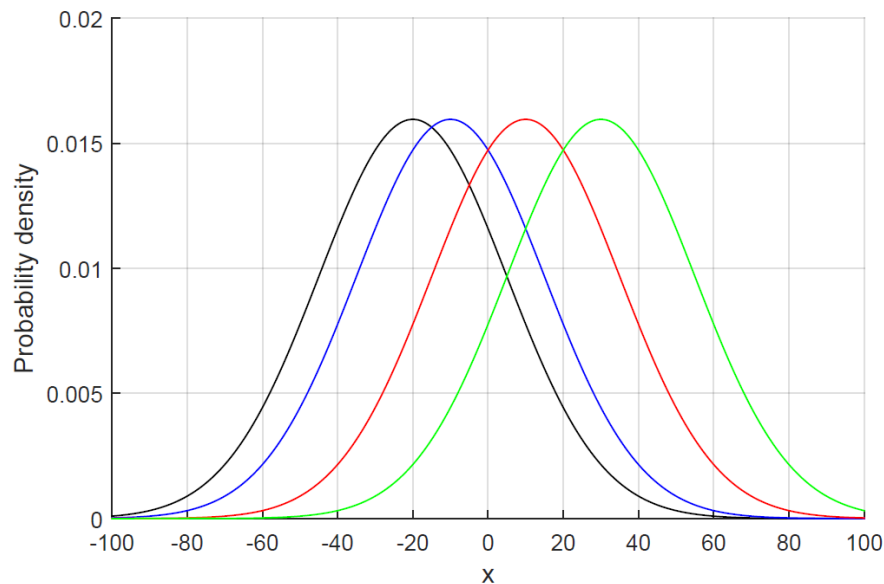
# Kullback-Leibler divergence: a distance between two density functions.

KL divergence for discrete variable

$$D_{\text{KL}} [q; p] = \sum_i p_i \log \frac{p_i}{q_i} \geq 0$$

KL divergence for continuous variable

$$D_{\text{KL}} [q; p] = \int dx p(x) \log \frac{p(x)}{q(x)} \geq 0$$



# ML as a minimization of KL divergence.

KL divergence between true density  $p(x)$  and parametrized density  $q(x|\theta)$ :

$$\begin{aligned} D_{\text{KL}}[p(x); q(x|\theta)] &= \int dx p(x) \log \frac{p(x)}{q(x|\theta)} \\ &= E[\log p(x)] - E[\log q(x|\theta)] \end{aligned}$$

|                               |                                                                                   |                 |                       |
|-------------------------------|-----------------------------------------------------------------------------------|-----------------|-----------------------|
| Minimization of KL divergence | $D_{\text{KL}}[p(x); q(x \theta)]$                                                |                 |                       |
|                               |  |                 |                       |
|                               |                                                                                   | Maximization of | $E[\log q(x \theta)]$ |

Sampling approximation:

$$E[\log q(x|\theta)] \simeq \frac{1}{N} \sum_{i=1}^N \log q(x_i|\theta)$$

# ML as a minimization of KL divergence.

Sampling approximation:

$$\begin{aligned} \mathbb{E} \left[ \log q(x | \hat{\theta}) \right] - \frac{1}{N} \sum_{i=1}^N \log q(x_i | \hat{\theta}) &\approx -(\hat{\theta} - \theta^0)^T \mathbb{E} \left[ \frac{\partial^2}{\partial \theta \partial \theta^T} \log q(x | \hat{\theta}) \right] (\hat{\theta} - \theta^0) \\ &= (\hat{\theta} - \theta^0)^T I(\hat{\theta}) (\hat{\theta} - \theta^0) \end{aligned}$$

Fisher information:

$$I(\hat{\theta}) \equiv \mathbb{E} \left[ \frac{\partial \log q(x | \hat{\theta})}{\partial \theta} \frac{\partial \log q(x | \hat{\theta})}{\partial \theta^T} \right] = \mathbb{E} \left[ -\frac{\partial^2}{\partial \theta \partial \theta^T} \log q(x | \hat{\theta}) \right]$$

In the limit of large samples (infinite  $N$ ), the ML estimator is unbiased and efficient.

$$\hat{\theta} \sim \mathcal{N} \left( \theta^0, \frac{1}{N} I^{-1}(\hat{\theta}) \right)$$

ML is ...

- Asymptotically unbiased (i.e., approaches to a true value).
- Asymptotically efficient (i.e., achieves minimum variance, CRLB).



# Cramer-Rao lower bound: scalar parameter case.

Suppose a random variable  $X \sim p(x; \theta)$ , where  $\theta$  is fixed but unknown. Assume that  $p(x; \theta)$  satisfies the “regularity” condition:

$$\mathbb{E} \left[ \frac{\partial}{\partial \theta} \log p(x | \theta) \right] = 0,$$

where the expectation is with respect to  $p(x; \theta)$ . Then the variance of any unbiased estimator  $\hat{\theta}$  satisfies

$$\text{Var}[\hat{\theta}] \geq \frac{1}{I(\theta)}$$

Fisher information:

$$I(\theta) \equiv \mathbb{E} \left[ \left( \frac{\partial \log p(x | \theta)}{\partial \theta} \right)^2 \right] = \mathbb{E} \left[ -\frac{\partial^2 \log p(x | \theta)}{\partial \theta^2} \right]$$

## Cramer-Rao lower bound: vector case.

Suppose a random variable  $X \sim p(x|\boldsymbol{\theta})$ , where  $\boldsymbol{\theta}$  is fixed but unknown. Assume that  $p(x|\boldsymbol{\theta})$  satisfies the “regularity” condition:

$$\mathbb{E}\left[\frac{\partial}{\partial \boldsymbol{\theta}} \log p(x|\boldsymbol{\theta})\right] = 0,$$

where the expectation is with respect to  $p(x;\theta)$ . Then the variance of any unbiased estimator  $\hat{\boldsymbol{\theta}}$  satisfies

$$\text{Cov}[\hat{\boldsymbol{\theta}}] \geq \mathbf{I}^{-1}(\boldsymbol{\theta})$$

Fisher information matrix:

$$\{\mathbf{I}(\boldsymbol{\theta})\}_{ij} \equiv \mathbb{E}\left[\frac{\partial \log p(x|\boldsymbol{\theta})}{\partial \theta_i} \frac{\partial \log p(x|\boldsymbol{\theta})}{\partial \theta_j}\right] = \mathbb{E}\left[-\frac{\partial^2 \log p(x|\boldsymbol{\theta})}{\partial \theta_i \partial \theta_j}\right]$$

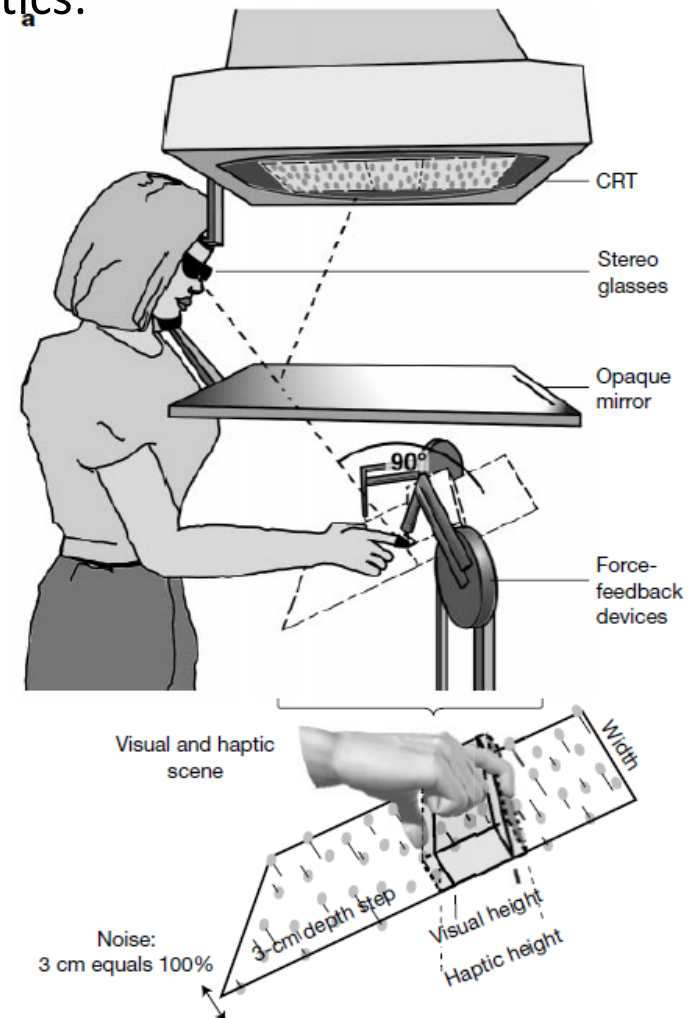
# Width estimation: optimal integration of vision and haptics.

ML estimation of object width from vision and haptics:

$x_1$ : visually perceived width,  $\sigma_1$ : visual uncertainty  
 $x_2$ : haptically perceived width  $\sigma_2$ : haptic uncertainty

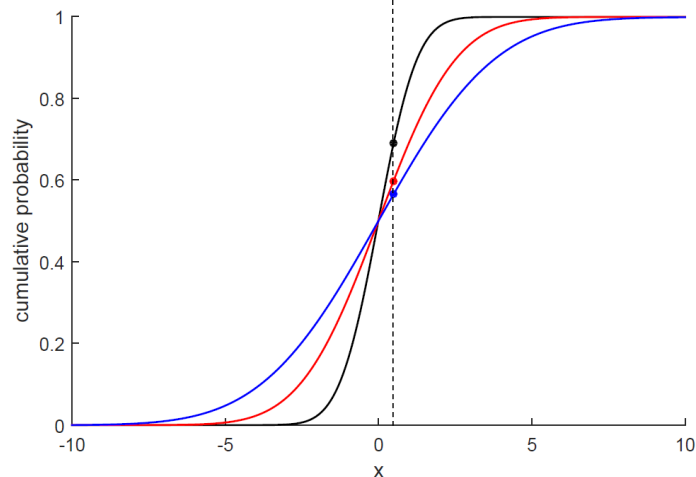
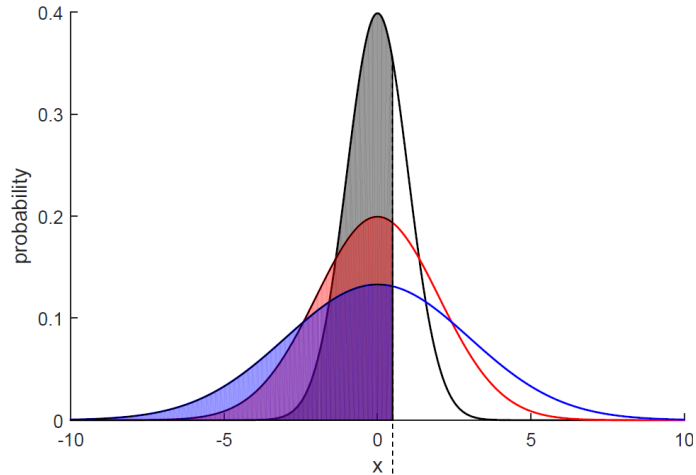
$$p(x_V, x_H | \mu) = p(x_V | \mu) p(x_H | \mu)$$
$$\propto \exp\left(-\frac{(x_V - \mu)^2}{2\sigma_V^2} - \frac{(x_H - \mu)^2}{2\sigma_H^2}\right)$$
$$\propto \exp\left(-\frac{(\hat{\mu} - \mu)^2}{2\sigma^2}\right)$$

$$\hat{\mu} = \frac{\sigma_H^2}{\sigma_V^2 + \sigma_H^2} x_V + \frac{\sigma_V^2}{\sigma_V^2 + \sigma_H^2} x_H$$
$$\frac{1}{\sigma^2} = \frac{1}{\sigma_V^2} + \frac{1}{\sigma_H^2}$$

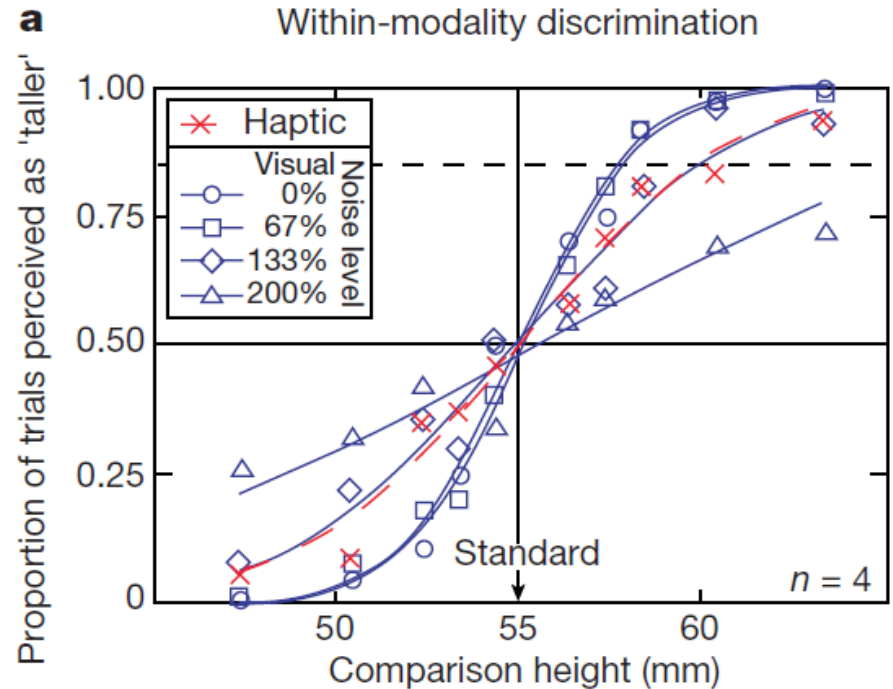


# Width estimation: optimal integration of vision and haptics.

## Density functions (sensory uncertainty)

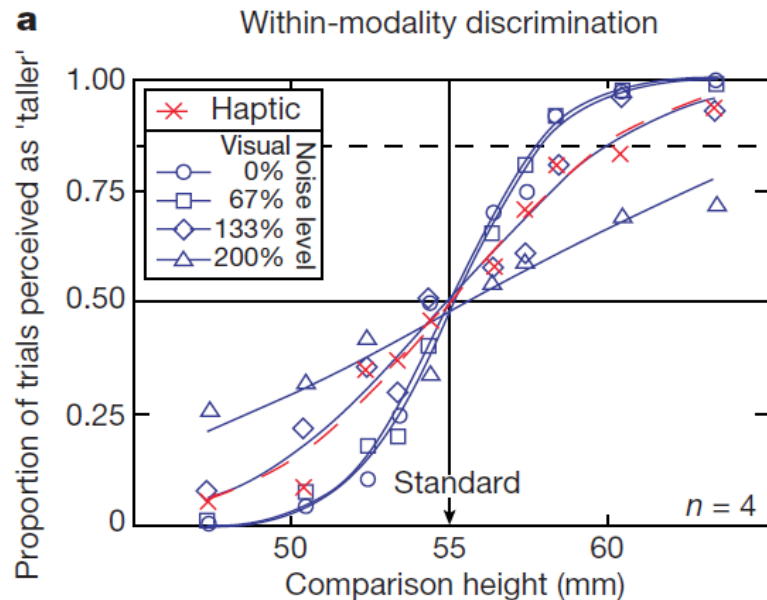


## Distribution functions (human response)

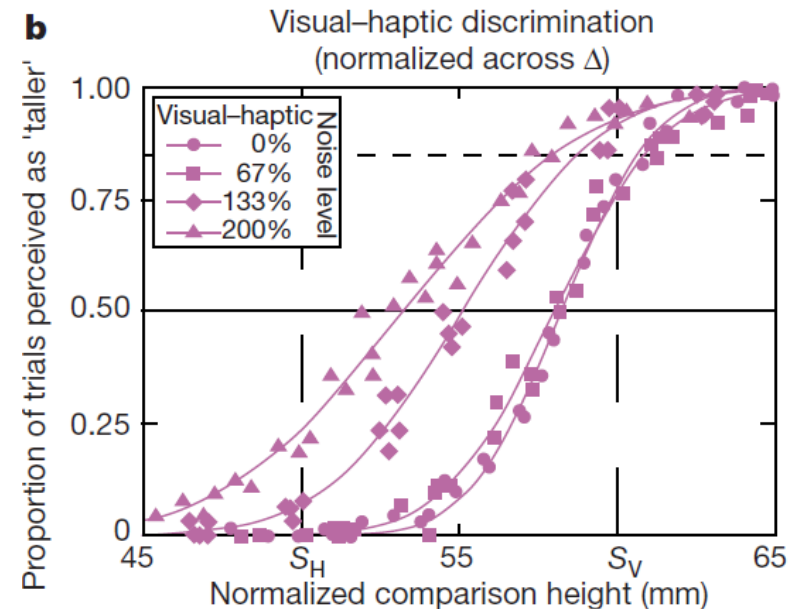


By examining human forced choice, the sensory density function can be estimated.

# Width estimation: optimal integration of vision and haptics.



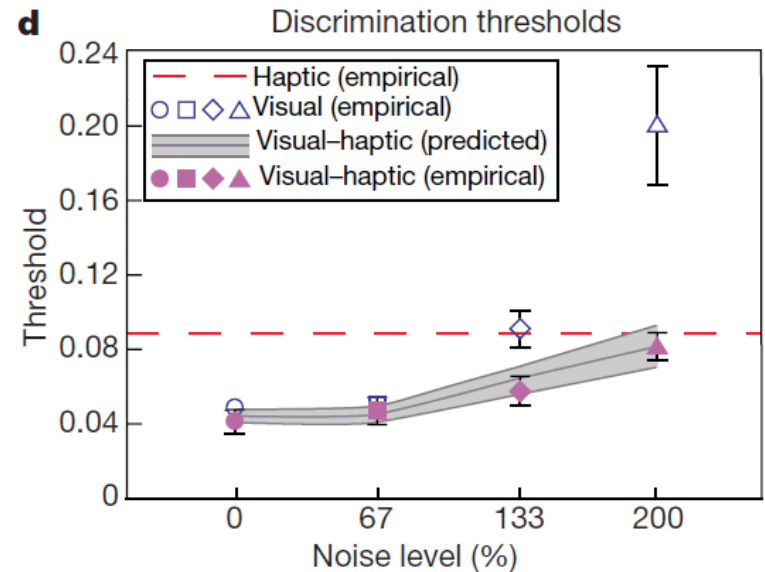
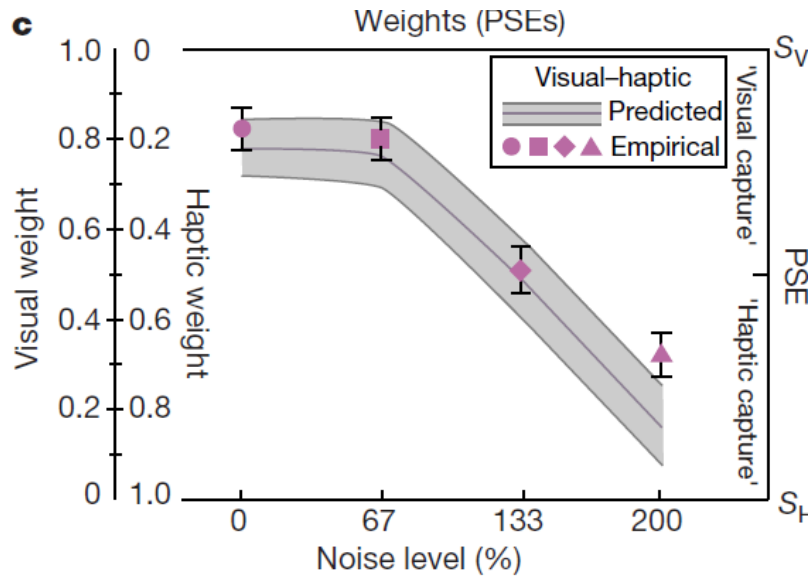
mean ( $x_V$ ) and variance ( $\sigma_V^2$ ) of vision  
mean ( $x_H$ ) and variance ( $\sigma_H^2$ ) of haptics



mean ( $x$ ) and variance ( $\sigma_V$ ) of  
multi-modal integration

Human performance in visual-haptic discrimination (RIGHT) is predicted by maximum-likelihood integration of those in visual and haptic discrimination measured separately (LEFT).

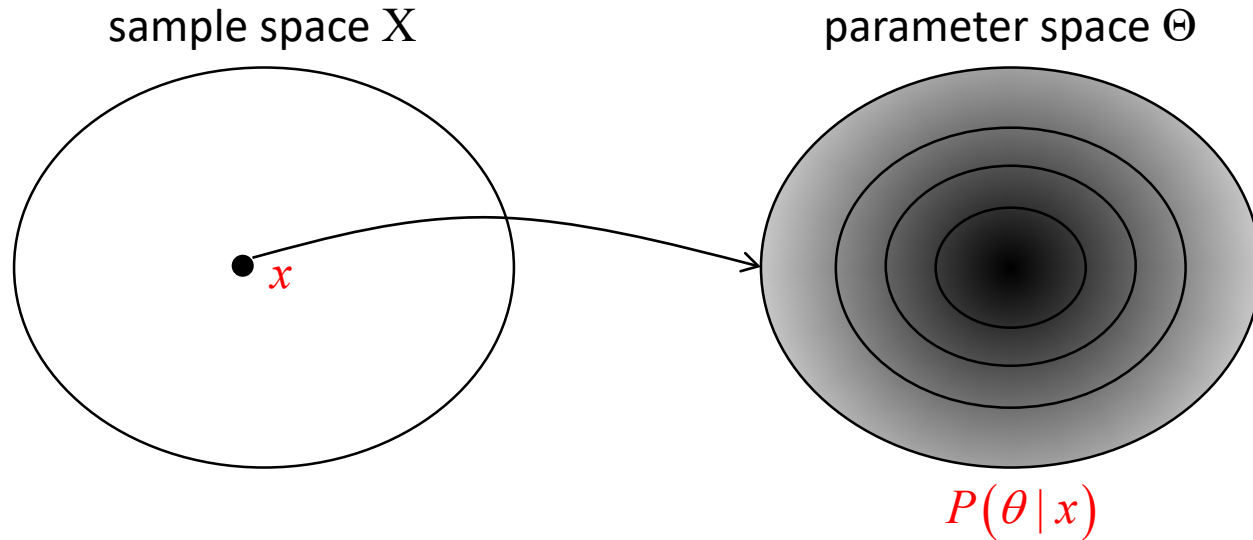
# Width estimation: optimal integration of vision and haptics.



$$x = \frac{\sigma_H^2}{\sigma_V^2 + \sigma_H^2} x_1 + \frac{\sigma_V^2}{\sigma_V^2 + \sigma_H^2} x_H$$

$$\frac{1}{\sigma^2} = \frac{1}{\sigma_V^2} + \frac{1}{\sigma_H^2}$$

# Bayes' theorem: mapping observation to probability density.



$$P(\theta | x) = \frac{P(x | \theta) P(\theta)}{P(x)}$$

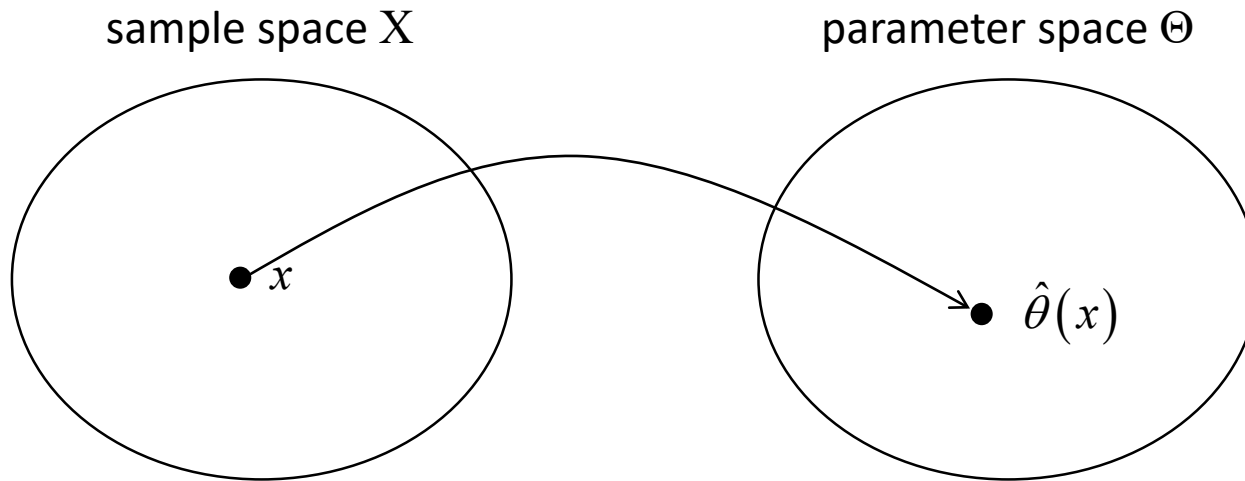
$P(\theta | x)$ : Posterior probability

$P(x | \theta)$ : Likelihood

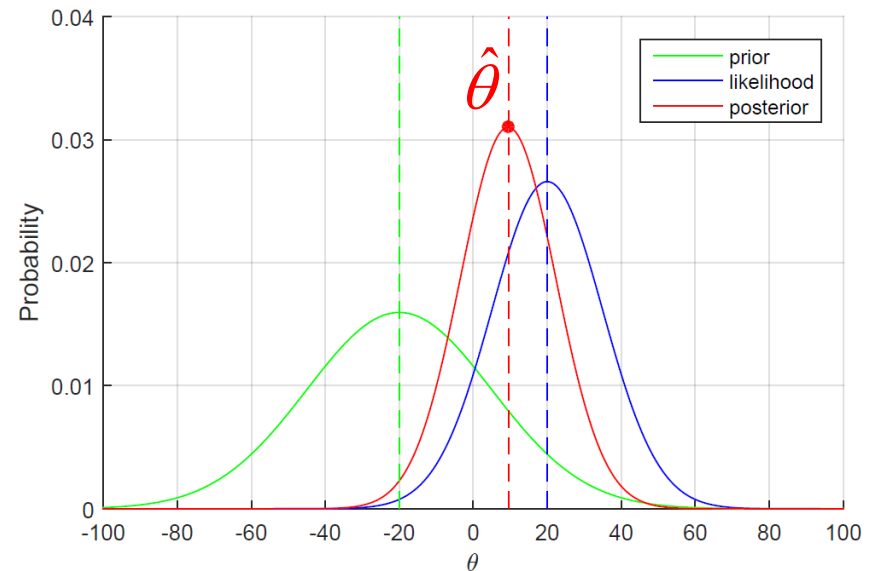
$P(\theta)$ : Prior probability

$P(x)$ : Marginal probability

# Bayesian estimation: Maximum A Posteriori (MAP).

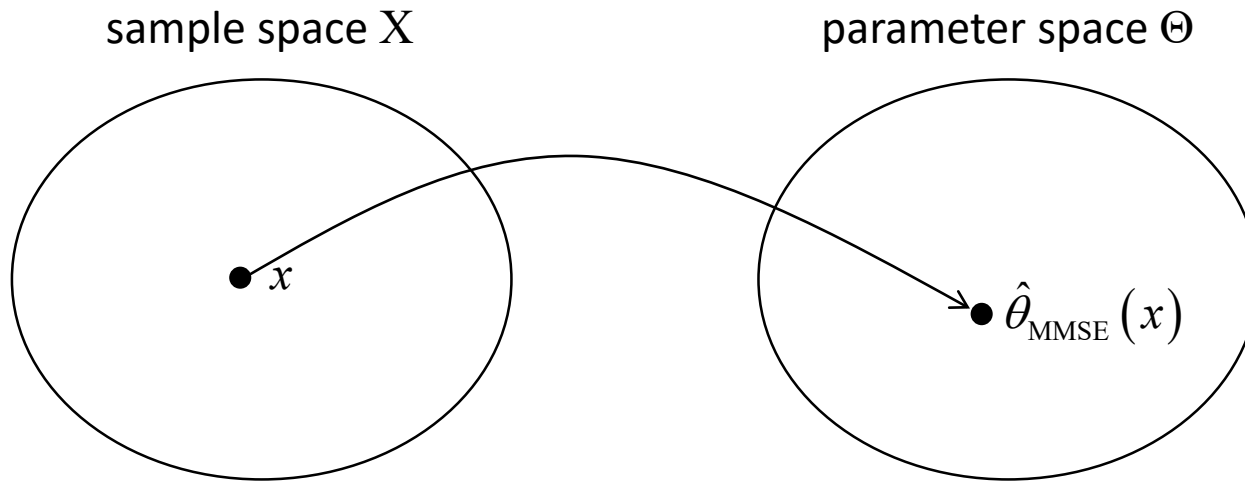


$$\begin{aligned}\hat{\theta}(x) &= \arg \max_{\theta} P(\theta | x) \\ &= \arg \max_{\theta} \frac{P(x | \theta) P(\theta)}{P(x)} \\ &= \arg \max_{\theta} P(x | \theta) P(\theta)\end{aligned}$$



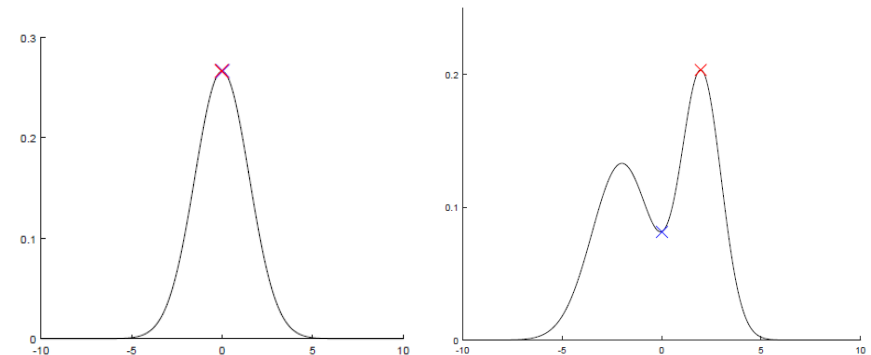


# Bayesian estimation: Mean minimum-squared error estimator.



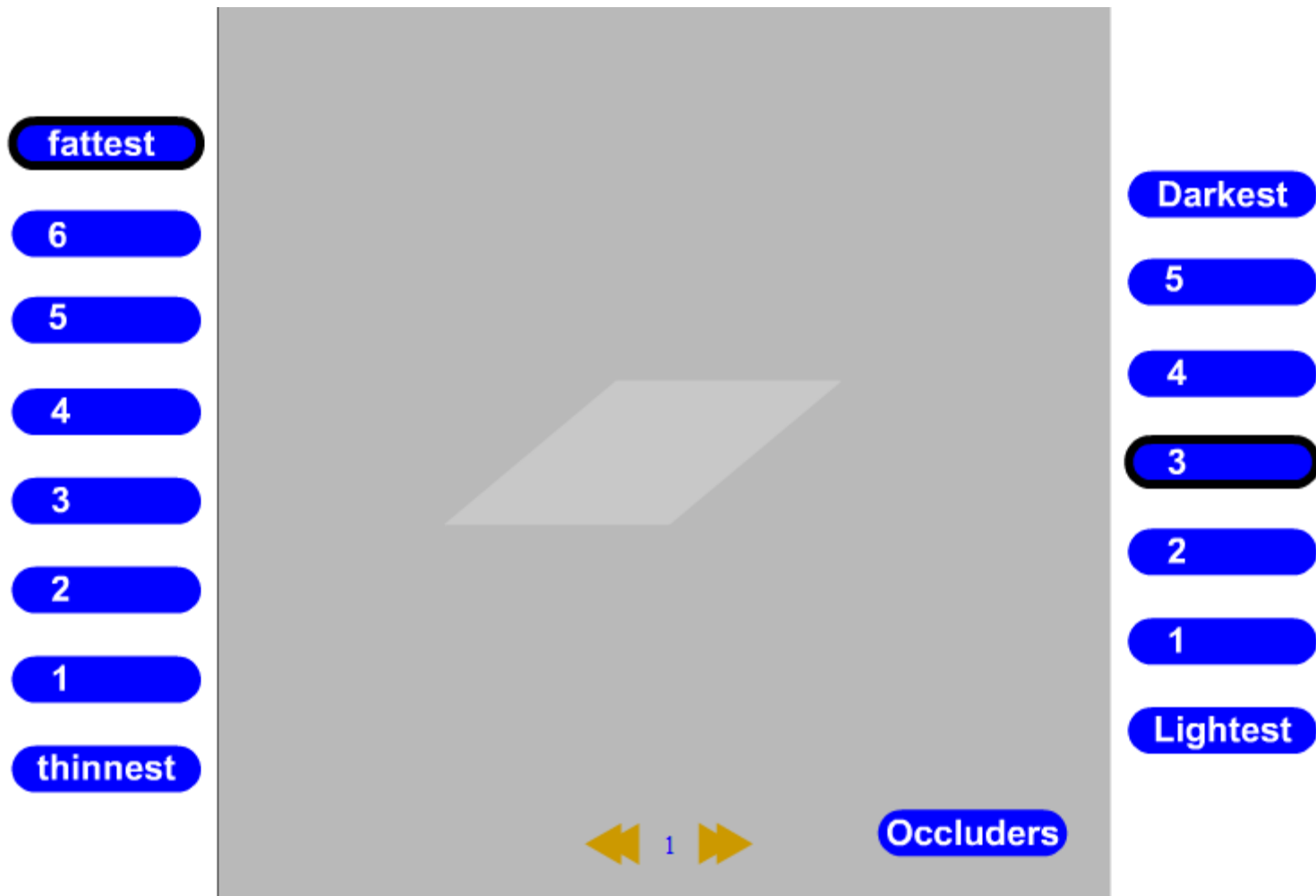
$$\begin{aligned}\hat{\theta}_{\text{MMSE}}(x) &= \arg \min_{\hat{\theta}} E_{\theta} \left[ (\theta - \hat{\theta})^2 \middle| x \right] \\ &= \arg \min_{\hat{\theta}} \int (\theta - \hat{\theta})^2 p(\theta|x) d\theta\end{aligned}$$

×: MAP, ×: MMSE



$$\hat{\theta}_{\text{MMSE}}(x) = \int \theta p(\theta|x) d\theta = E_{\theta}[\theta|x]$$

# Motion illusion as optimal percepts.



Use the left blue buttons to change the size/shape of the moving figure; use the right blue buttons to change the color of the background. Use the yellow buttons to control the speed.

# Motion illusion as optimal percepts.

Prior density: most of objects have small velocity (not moving).

$$P(v_x, v_y) \propto \exp \left\{ -\frac{1}{2\sigma_v^2} (v_x^2 + v_y^2) \right\}$$

Likelihood: Probability of observed image with given velocity  $(v_x, v_y)$ .

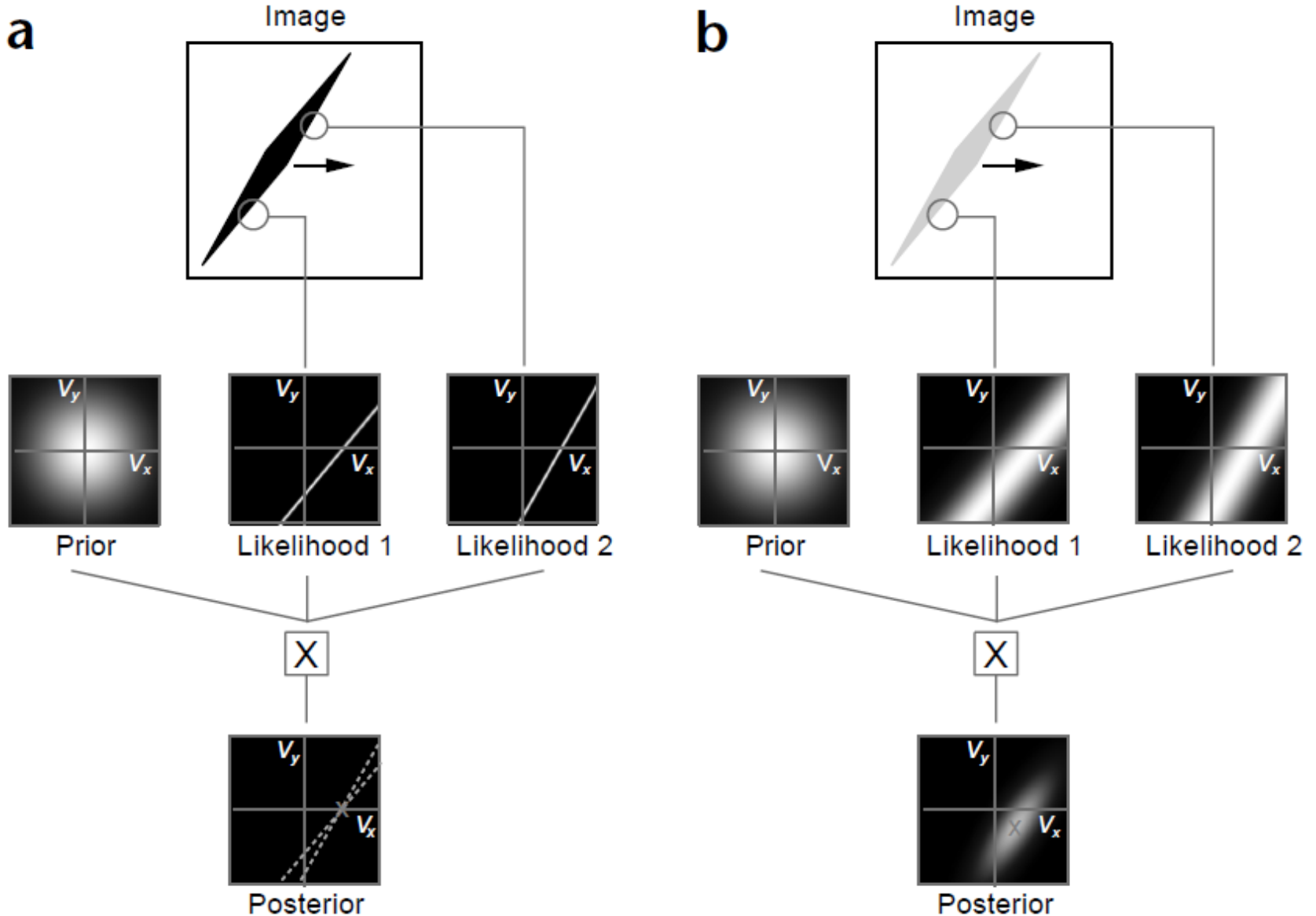
Lucas-Kanade image model:  $I(x + v_x \Delta t, y + v_y \Delta t, t + \Delta t) \approx I(x, y, t)$

$$P(I(x_i, y_i, t) | v_x, v_y) \propto \exp \left\{ -\frac{1}{2\sigma^2} \left( \frac{\partial I(x_i, y_i, t)}{\partial x} v_x + \frac{\partial I(x_i, y_i, t)}{\partial y} v_y + \frac{\partial I(x_i, y_i, t)}{\partial t} \right)^2 \right\}$$

Posterior density

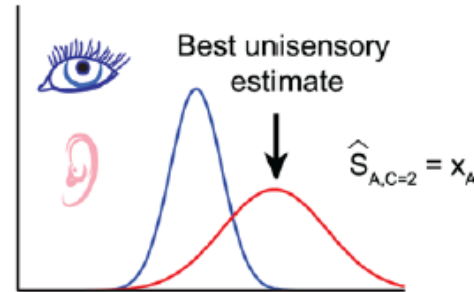
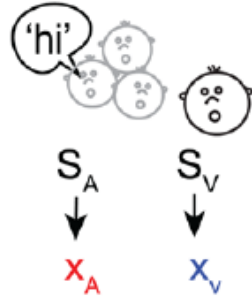
$$P(v_x, v_y | I(x_i, y_i, t)) \propto \exp \left\{ -\frac{1}{2\sigma_v^2} (v_x^2 + v_y^2) - \frac{1}{2\sigma^2} \sum_i \left( \frac{\partial I(x_i, y_i, t)}{\partial x} v_x + \frac{\partial I(x_i, y_i, t)}{\partial y} v_y + \frac{\partial I(x_i, y_i, t)}{\partial t} \right)^2 \right\}$$

# Motion illusion as optimal percepts.

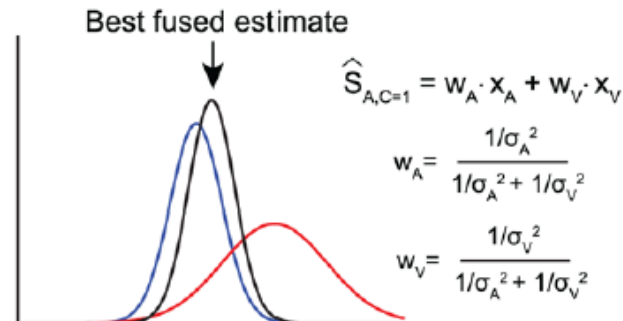
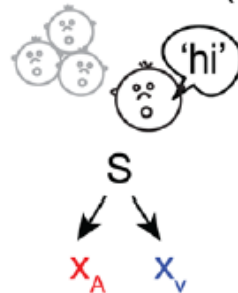


# Causal inference of sensory information.

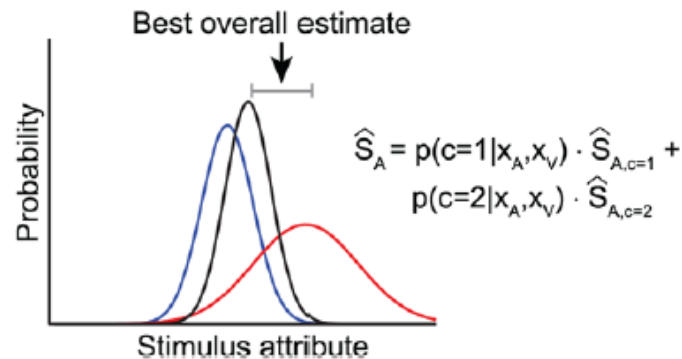
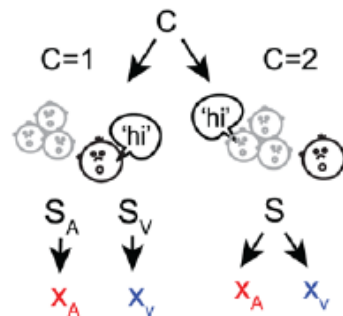
## A Separate sources (C=2)



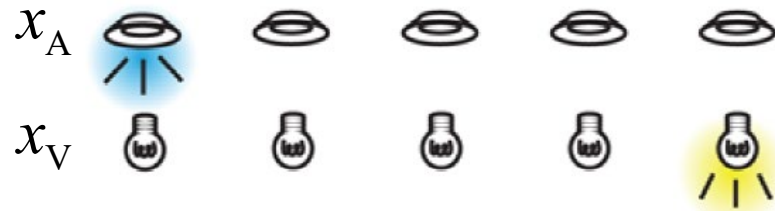
## B Common source (C=1)



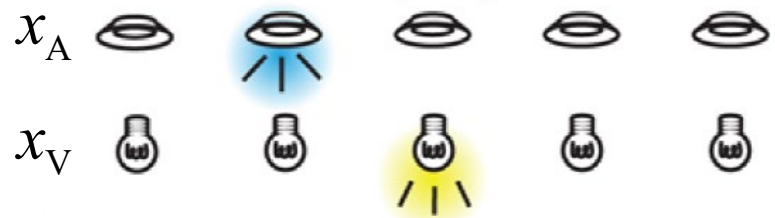
## C Causal inference



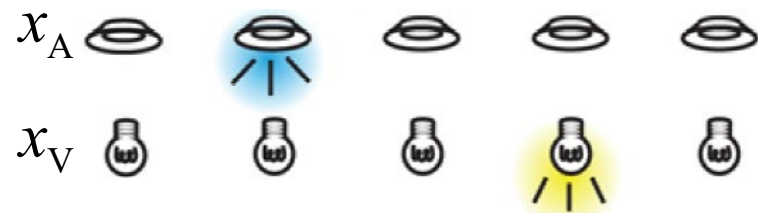
# Causal inference of sensory information.



Given visual ( $x_V$ ) and auditory ( $x_A$ ) observations...



Q1: Do they belong to a single object ( $C=1$ ) or come from two different objects ( $C=2$ )?



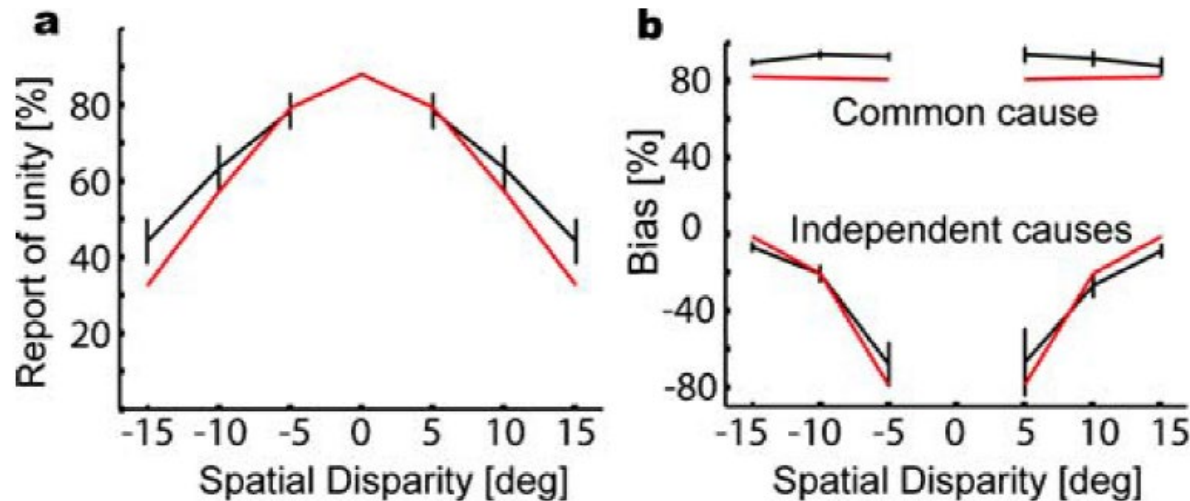
Q2: What are the most likely positions for the visual and the auditory objects?

# Causal inference of sensory information.

Bayes' theorem:

$$P(C = 1 | x_A, x_V) = \frac{P(x_A, x_V | C = 1)P(C = 1)}{P(x_A, x_V | C = 1)P(C = 1) + P(x_A, x_V | C = 2)P(C = 2)}$$

$$P(C = 2 | x_A, x_V) = \frac{P(x_A, x_V | C = 2)P(C = 2)}{P(x_A, x_V | C = 1)P(C = 1) + P(x_A, x_V | C = 2)P(C = 2)}$$

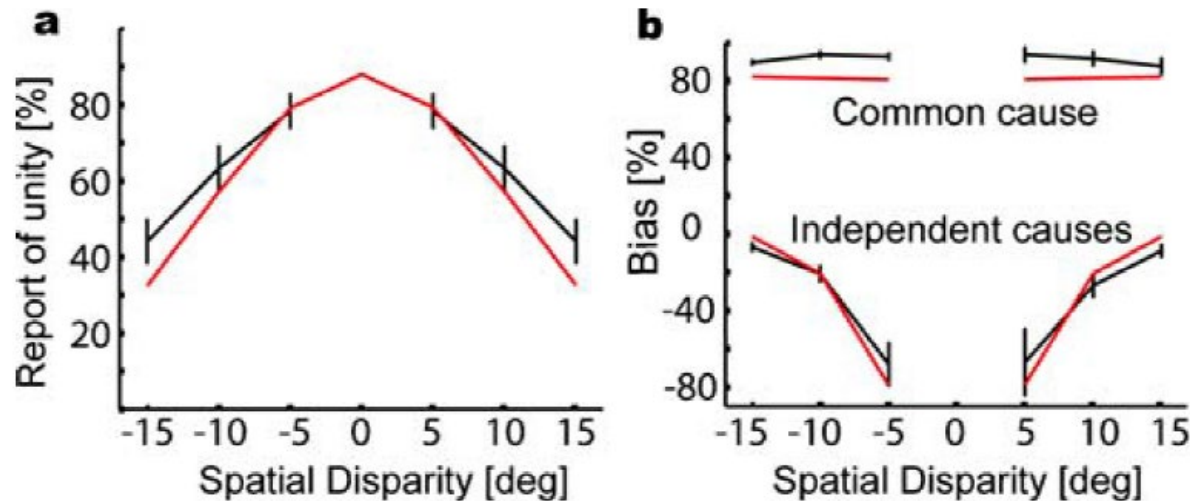
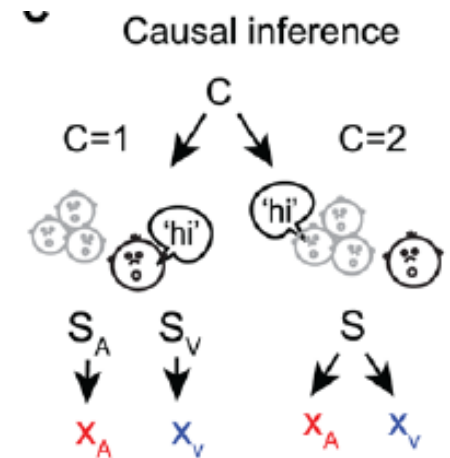


# Causal inference of sensory information.

Model averaging.:

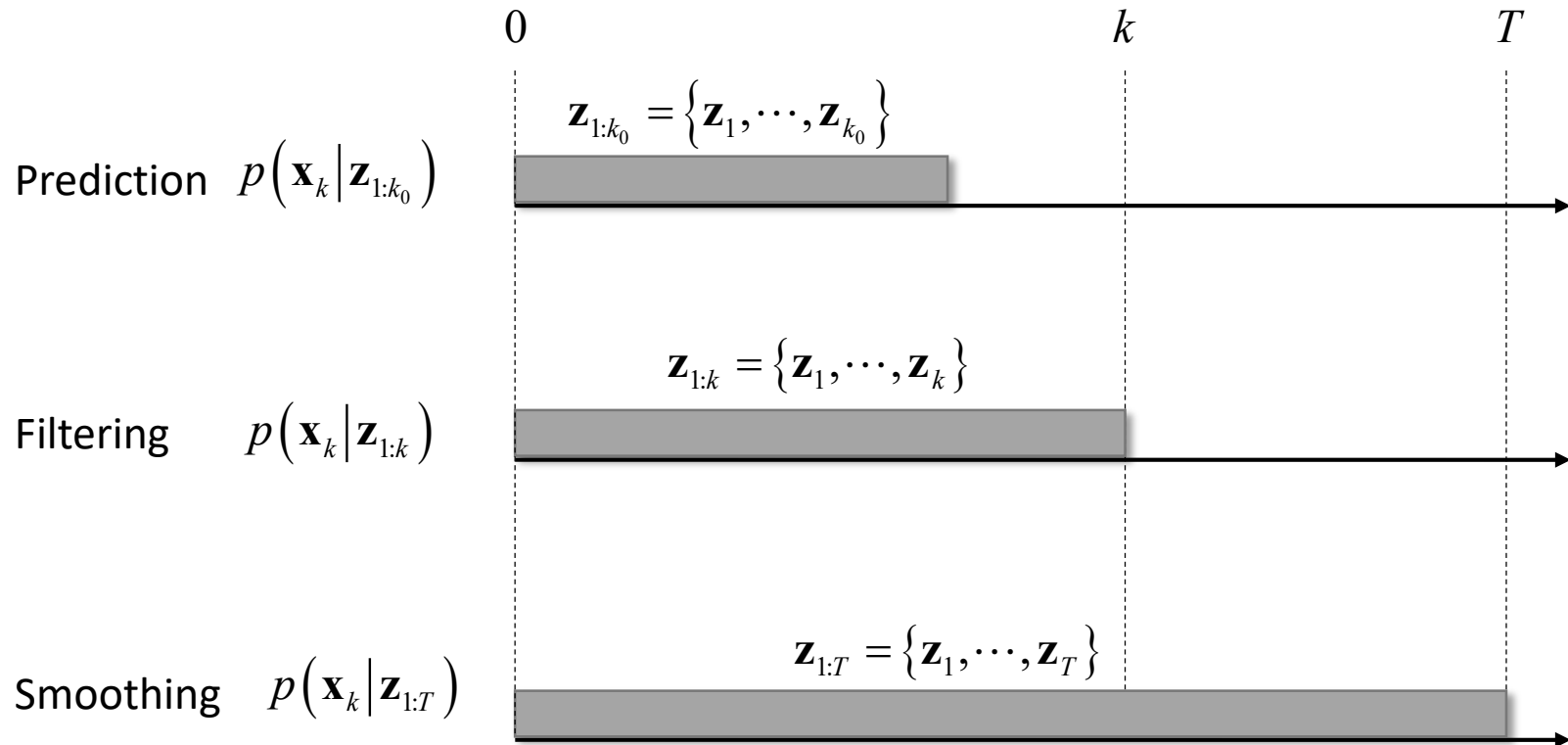
$$\hat{s}_A = \hat{s}_{A,C=1} \times P(C=1 | x_A, x_V) + \hat{s}_{A,C=2} \times P(C=2 | x_A, x_V)$$

$$\hat{s}_V = \hat{s}_{V,C=1} \times P(C=1 | x_A, x_V) + \hat{s}_{V,C=2} \times P(C=2 | x_A, x_V)$$





# Prediction, filtering and smoothing of a dynamic system.



# Estimation of dynamic systems: Kalman filter.

## Stochastic state-space model

$$\mathbf{x}_{k+1} = \mathbf{A}\mathbf{x}_k + \mathbf{w}_k$$

$$\mathbf{z}_k = \mathbf{C}\mathbf{x}_k + \mathbf{v}_k$$

Problem: Given a sequence of measurements up to step  $k$ ,

$$\mathbf{z}_{1:k} = \{\mathbf{z}_1, \mathbf{z}_2, \dots, \mathbf{z}_k\},$$

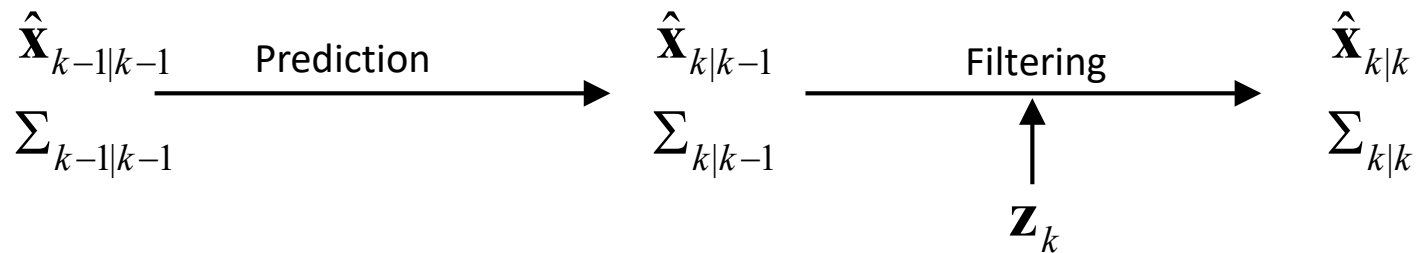
then estimate the conditional mean

$$\hat{\mathbf{x}}_{k|k} = E[\mathbf{x}_k | \mathbf{z}_{1:k}],$$

and the conditional covariance of the state vector at step  $k$ ,

$$\Sigma_{k|k} = \text{Cov}[\mathbf{x}_k | \mathbf{z}_{1:k}] = E\left[(\mathbf{x}_k - \hat{\mathbf{x}}_{k:k})(\mathbf{x}_k - \hat{\mathbf{x}}_{k:k})^T | \mathbf{z}_{1:k}\right],$$

Kalman filter optimally integrates prediction and measurement.



Prediction

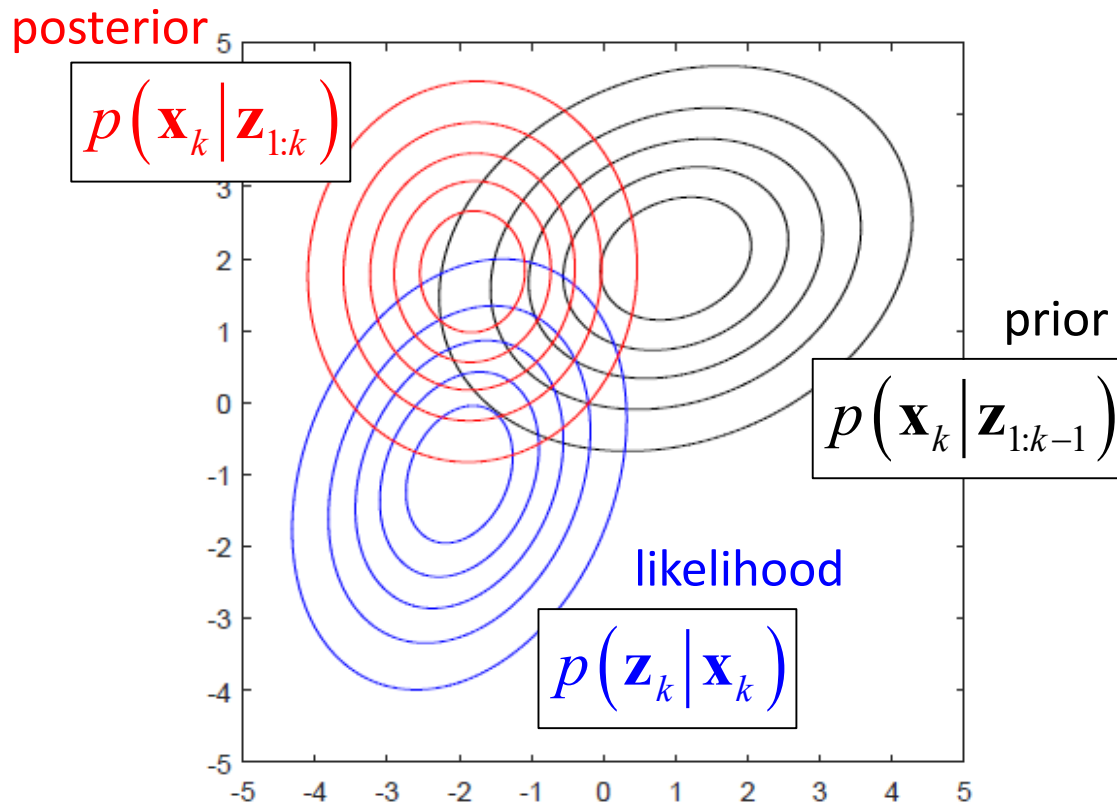
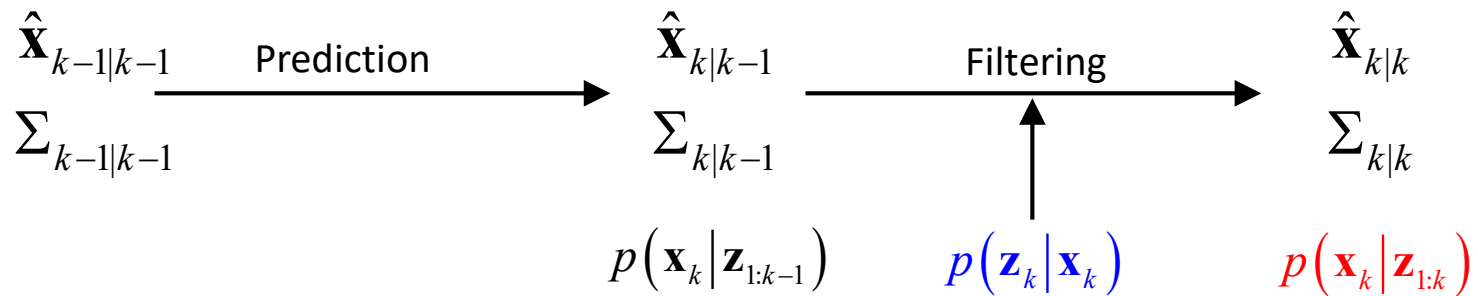
$$\begin{aligned}\hat{\mathbf{x}}_{k|k-1} &= \mathbf{A}\hat{\mathbf{x}}_{k-1|k-1} \\ \Sigma_{k|k-1} &= \mathbf{A}\Sigma_{k-1|k-1}\mathbf{A}^T + \mathbf{Q}\end{aligned}$$

Filtering

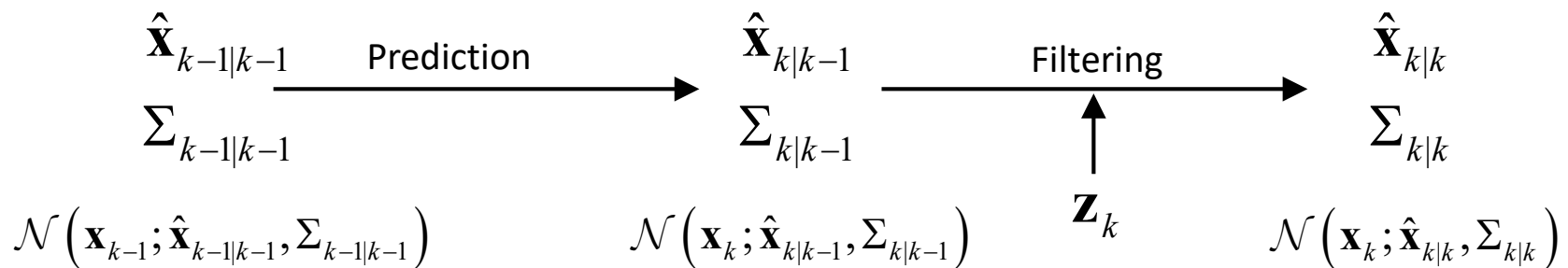
$$\begin{aligned}\hat{\mathbf{x}}_{k|k} &= \hat{\mathbf{x}}_{k|k-1} + \mathbf{K}_k (\mathbf{z}_k - \mathbf{C}\hat{\mathbf{x}}_{k|k-1}) \\ \Sigma_{k|k} &= (\mathbf{I} - \mathbf{K}_k \mathbf{C})\Sigma_{k|k-1}\end{aligned}$$

Kalman gain  $\mathbf{K}_k = \Sigma_{k|k-1} \mathbf{C}^T (\mathbf{C}\Sigma_{k|k-1} \mathbf{C}^T + \mathbf{R})^{-1}$

Kalman filter optimally integrates prediction and measurement.



# Kalman filter: Derivation.



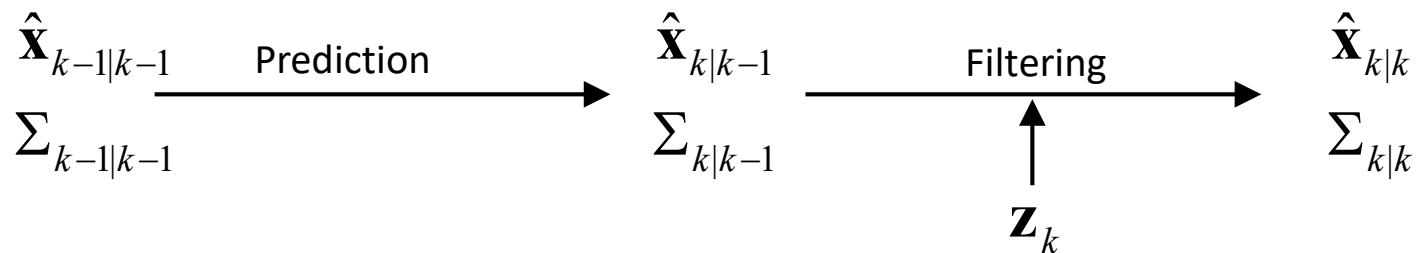
## Prediction

$$\begin{aligned} p(\mathbf{x}_k | \mathbf{z}_{1:k-1}) &= \int d\mathbf{x}_{k-1} p(\mathbf{x}_k | \mathbf{x}_{k-1}) p(\mathbf{x}_{k-1} | \mathbf{z}_{1:k-1}) \\ &= \int d\mathbf{x}_{k-1} \mathcal{N}(\mathbf{x}_k; \mathbf{A}\mathbf{x}_{k-1}, \mathbf{Q}) \mathcal{N}(\mathbf{x}_{k-1}; \hat{\mathbf{x}}_{k-1|k-1}, \Sigma_{k-1|k-1}) \\ &= \mathcal{N}(\mathbf{x}_k; \hat{\mathbf{x}}_{k|k-1}, \Sigma_{k|k-1}) \end{aligned}$$

## Filtering

$$p(\mathbf{x}_k | \mathbf{z}_{1:k}) = \frac{p(\mathbf{z}_k | \mathbf{x}_k) p(\mathbf{x}_k | \mathbf{z}_{1:k-1})}{p(\mathbf{z}_k | \mathbf{z}_{1:k-1})} = \mathcal{N}(\mathbf{x}_k; \hat{\mathbf{x}}_{k|k}, \Sigma_{k|k})$$

# Kalman filter: Derivation.



Lemma: When a joint density function of two variables  $\mathbf{x}$  and  $\mathbf{y}$  are given as

$$\begin{pmatrix} \mathbf{x} \\ \mathbf{y} \end{pmatrix} \sim \mathcal{N} \left( \begin{pmatrix} \boldsymbol{\mu}_x \\ \boldsymbol{\mu}_y \end{pmatrix}, \begin{pmatrix} \Sigma_{xx} & \Sigma_{xy} \\ \Sigma_{yx} & \Sigma_{yy} \end{pmatrix} \right),$$

then conditional probability densities are given as follows:

$$\mathbf{x} | \mathbf{y} \sim \mathcal{N} \left( \boldsymbol{\mu}_x + \Sigma_{xy} \Sigma_{yy}^{-1} (\mathbf{y} - \boldsymbol{\mu}_y), \Sigma_{xx} - \Sigma_{xy} \Sigma_{yy}^{-1} \Sigma_{yx} \right)$$

$$\mathbf{y} | \mathbf{x} \sim \mathcal{N} \left( \boldsymbol{\mu}_y + \Sigma_{yx} \Sigma_{xx}^{-1} (\mathbf{x} - \boldsymbol{\mu}_x), \Sigma_{yy} - \Sigma_{yx} \Sigma_{xx}^{-1} \Sigma_{xy} \right)$$

# Matlab code of Kalman filter and smoother (1/5)

```
function [x, z, w, v] = sim_statespace(m, N)

x = zeros(m.dx, N);
z = zeros(m.dz, N);

w = m.sQ*randn(m.dx, N);
v = m.sR*randn(m.dz, N);

x(:,1) = m.x0 + sqrtm(m.P0)*randn(m.dx, 1);
z(:,1) = m.C*x(:,1) + v(1);

for n=1:N-1
    x(:,n+1) = m.A*x(:,n) + w(:,n);
    z(:,n+1) = m.C*x(:,n+1) + v(:,n+1);
end
```

# Matlab code of Kalman filter and smoother (3/5)

```
function [xp, xf, Pp, Pf, K] = kalmanfiter(m, z)

N = size(z,2);
xp = zeros(m.dx, N); xf = zeros(m.dx, N);
Pp = zeros(m.dx, m.dx, N); Pf = zeros(m.dx, m.dx, N);
K = zeros(m.dx, m.dz, N);

% initialize:
xp(:,1) = m.x0; Pp(:, :, 1) = m.P0;

K(:, :, 1) = Pp(:, :, 1)*m.C'/(m.C*Pp(:, :, 1)*m.C'+m.R);
xf(:,1) = xp(:,1) + K(:, :, 1)*(z(:,1)-m.C*xp(:,1));
Pf(:, :, 1) = (eye(size(m.x0,1))-K(:, :, 1)*m.C)*Pp(:, :, 1);

% main loop for kalman filter:
for n=1:N-1
    % prediction step:
    xp(:,n+1) = m.A*xf(:,n);
    Pp(:, :, n+1) = m.A*Pf(:, :, n)*m.A'+m.Q;

    % filtering step:
    K(:, :, n+1) = Pp(:, :, n+1)*m.C'/(m.C*Pp(:, :, n+1)*m.C'+m.R);
    xf(:,n+1) = xp(:,n+1) + K(:, :, n+1)*(z(:,n+1)-m.C*xp(:,n+1));
    Pf(:, :, n+1) = (eye(size(m.x0,1))-K(:, :, n+1)*m.C)*Pp(:, :, n+1);
end
```



# Matlab code of Kalman filter and smoother (4/5)

```
function [xs, Ps, G] = kalman smoother(m)

N = size(m.xp,2);
xs = zeros(m.dx, N);
Ps = zeros(m.dx, m.dx, N);
G = zeros(m.dx, m.dx, N);

% terminal conditions:
xs(:,N) = m.xf(:,N);
Ps(:, :, N) = m.Pf(:, :, N);

for n=N-1:(-1):1
    G(:, :, n) = m.Pf(:, :, n)*m.A'/m.Pp(:, :, n+1);
    xs(:, n) = m.xf(:, n) + G(:, :, n)*(xs(:, n+1)-m.xp(:, n+1));
    Ps(:, :, n) = m.Pf(:, :, n) + G(:, :, n)*(Ps(:, :, n+1)-
m.Pp(:, :, n+1))*G(:, :, n)';
end
```

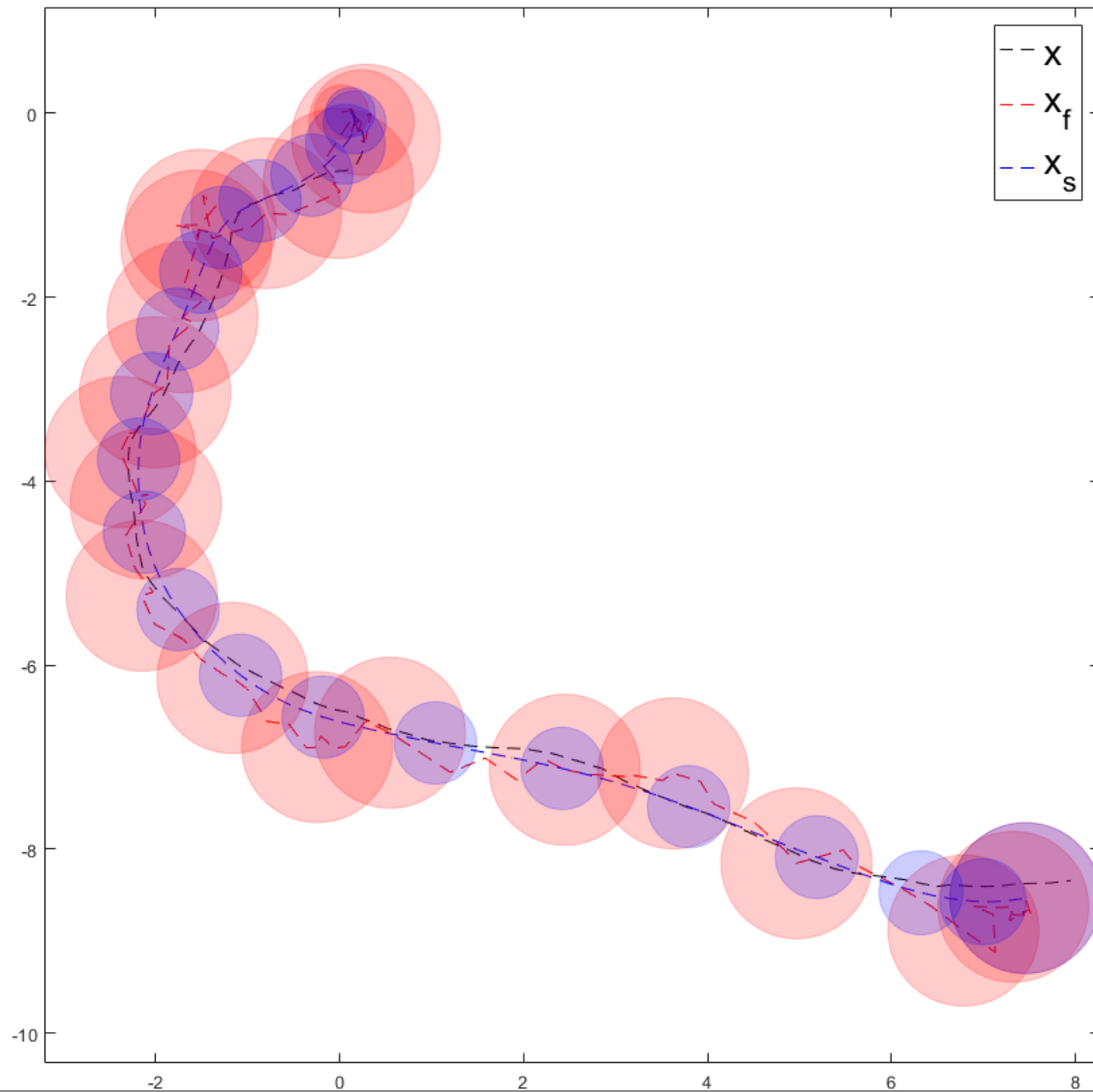
# Kalman filter

```
dt = 0.1;
q1 = 1.;
q2 = 1.;
sigma1 = 0.5;
sigma2 = 0.5;

model.A = [1 0 dt 0; 0 1 0 dt; 0 0 1 0; 0 0 0 1];
model.C = [1 0 0 0; 0 1 0 0];
model.Q = [q1*dt^3/3 0 q1*dt^2/2 0; 0 q2*dt^3/3 0
           q2*dt^2/2;
           q1*dt^2/2 0 q1*dt 0; 0 q2*dt^2/2 0 q2*dt];
model.sQ = sqrtm(model.Q);
model.R = [sigma1^2 0; 0 sigma2^2];
model.sR = sqrtm(model.R);
model.x0 = [0 0 0 0]';
model.P0 = diag([0.1^2 0.1^2 0.01^2 0.01^2]);
model.dx = size(model.A,1);
model.dz = size(model.C,1);
```

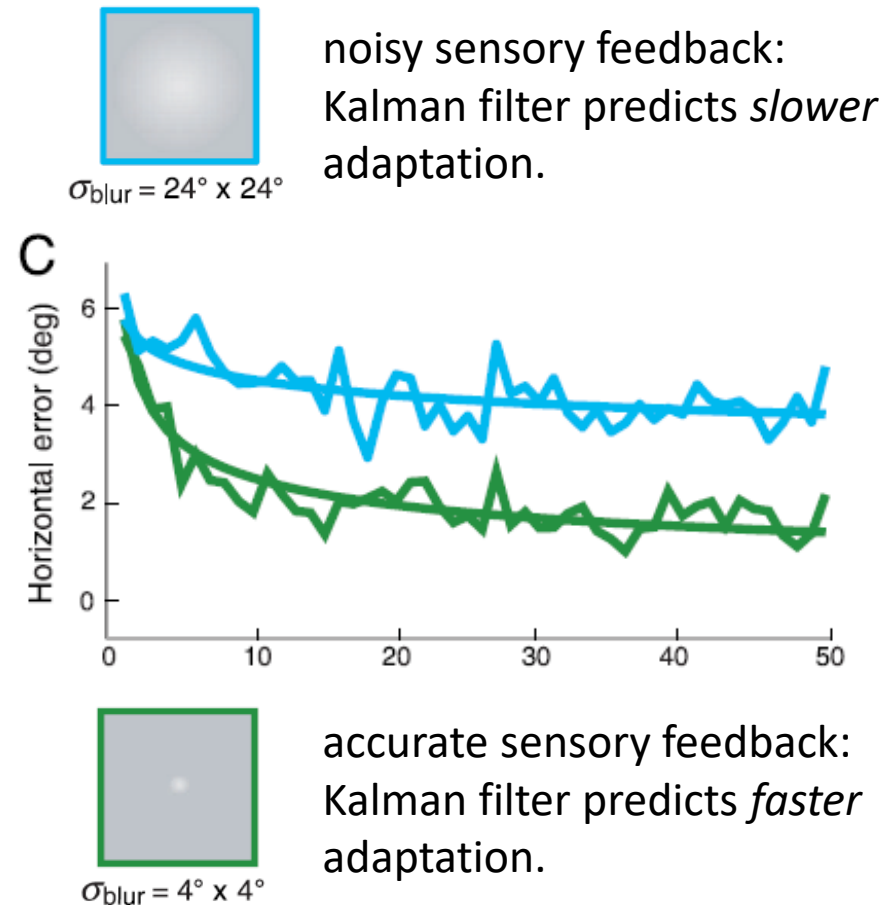
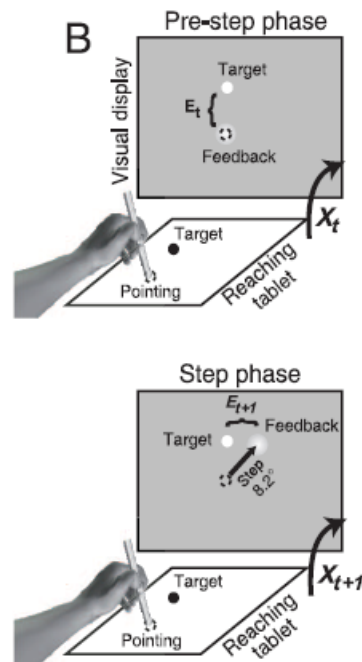
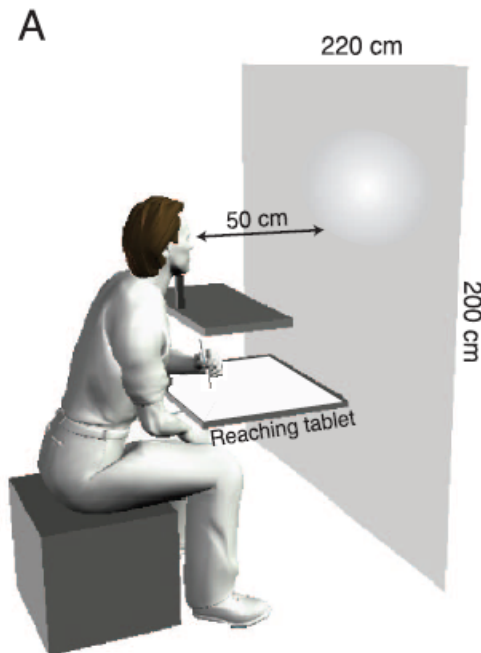
Parameters taken from: Example 4.3, “Bayesian Filtering and Smoothing” by Sarkka

# Matlab code of Kalman filter and smoother (4/5)

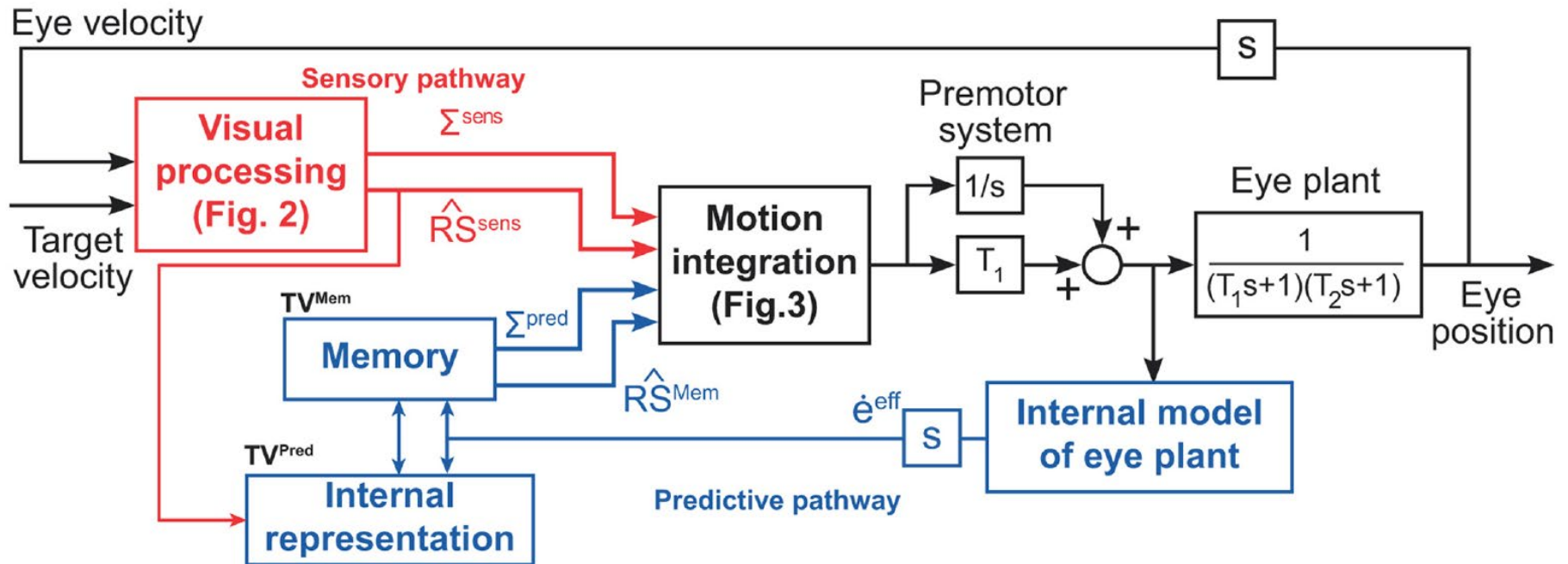


# Kalman filter explains smooth eye pursuit movements.

$$\hat{\mathbf{x}}_{k|k} = \hat{\mathbf{x}}_{k|k-1} + \mathbf{K}_k \left( \mathbf{z}_k - \mathbf{C}\hat{\mathbf{x}}_{k|k-1} \right)$$



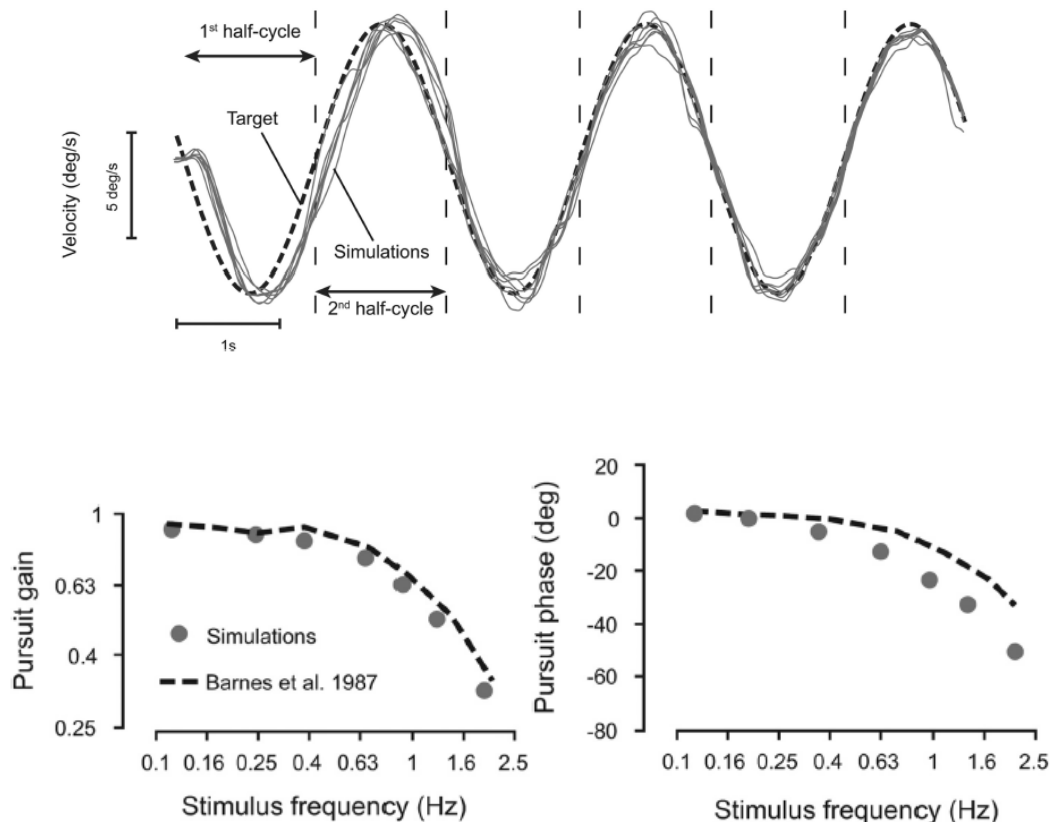
# Kalman filter explains smooth eye pursuit movements.



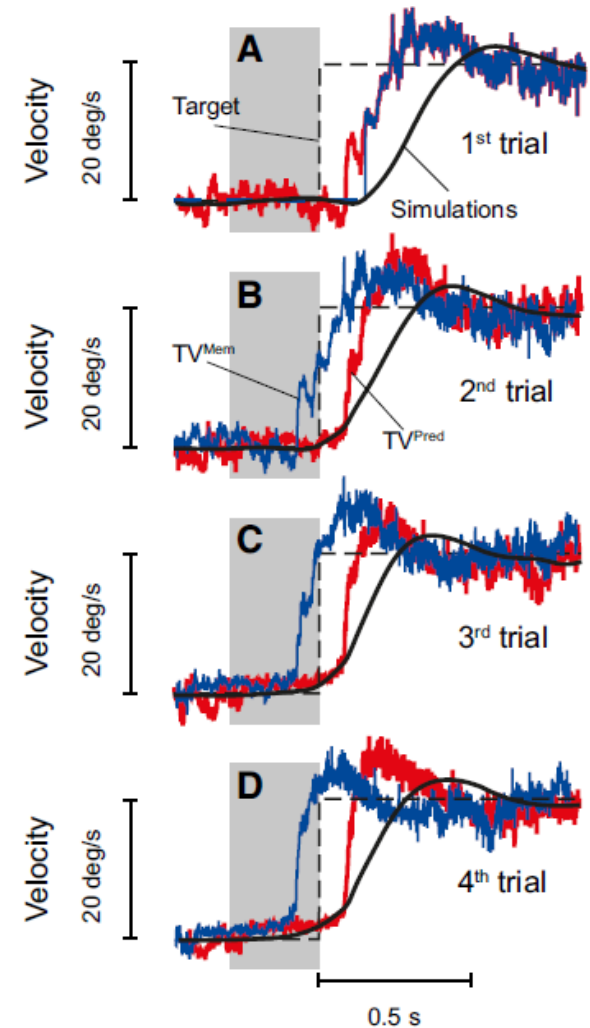
Kalman filter #1 for visual processing of retinal slip.  
Kalman filter #2 for remembered target dynamics.

# Kalman filter explains smooth eye pursuit movements.

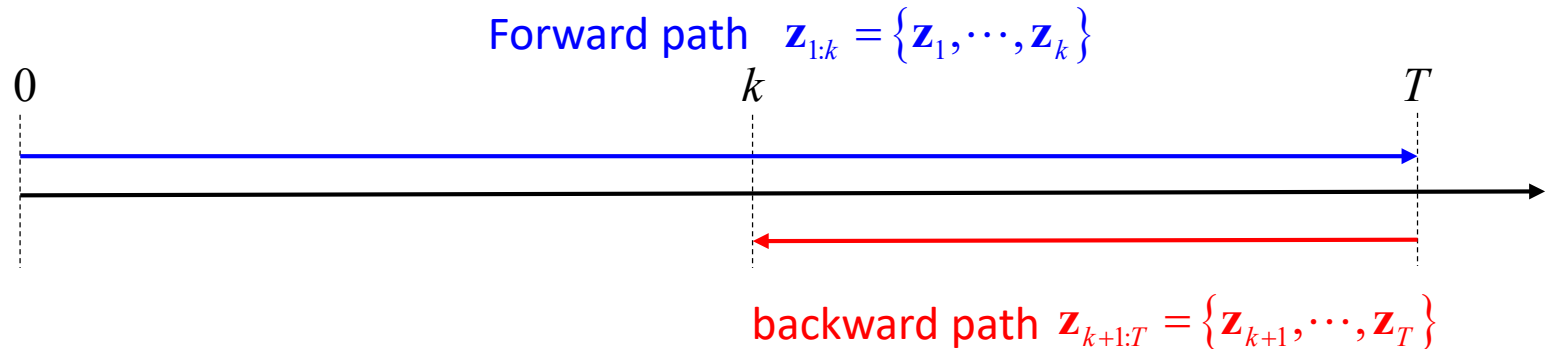
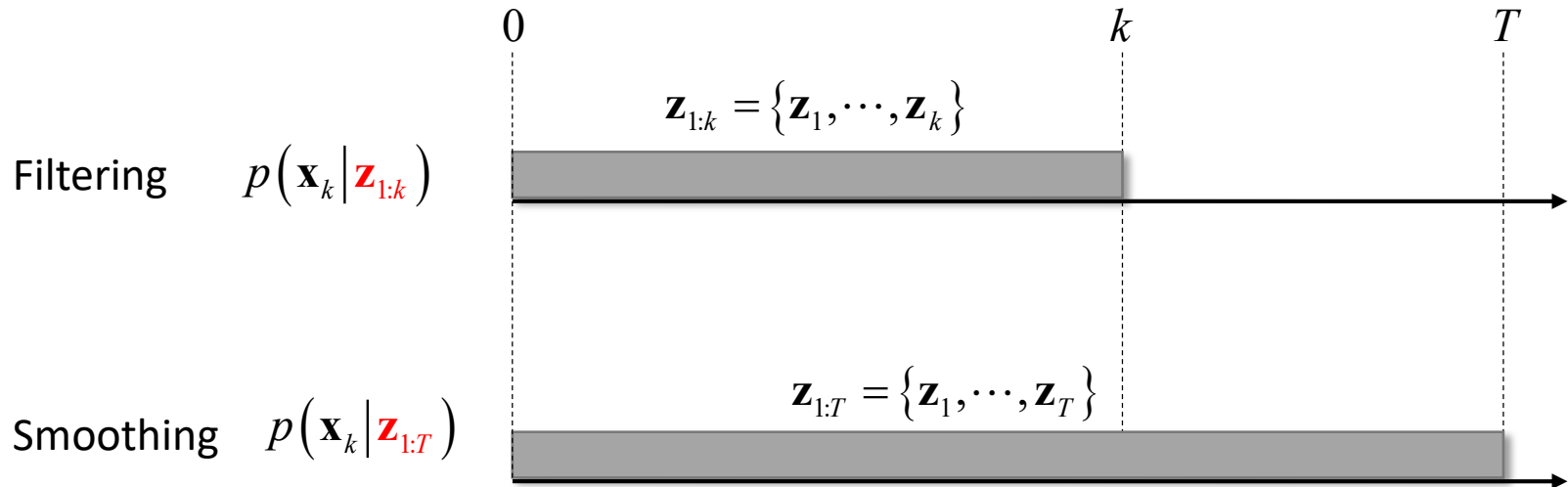
## Sinusoidal target motion



## Anticipatory smooth pursuit



Smoothing consists of forward and backward paths.



# Kalman smoother algorithm: RTS smoother.

Forward path: Using the standard algorithm of Kalman filter, compute the predicted mean and covariance,

$$\left( \hat{\mathbf{x}}_{k+1|k}, \Sigma_{k+1|k} \right) \quad k = 0, \dots, T-1,$$

and the filtered mean and covariance,

$$\left( \hat{\mathbf{x}}_{k|k}, \Sigma_{k|k} \right) \quad k = 0, \dots, T.$$

Backward path: Then, the smoothed mean and covariance,

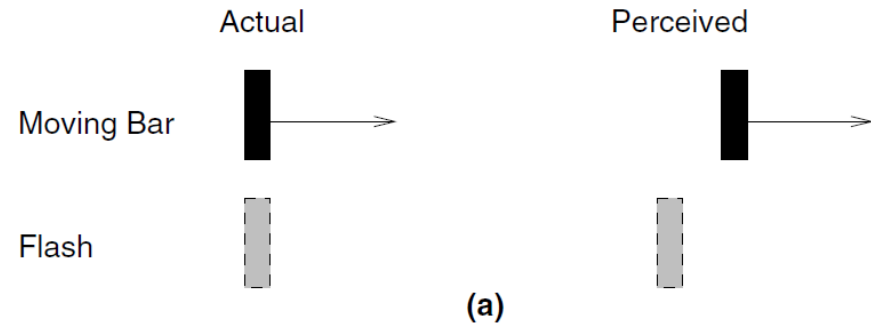
$$\left( \hat{\mathbf{x}}_{k|T}, \Sigma_{k|T} \right) \quad k = 0, \dots, T.$$

are solved backward in time as follows:

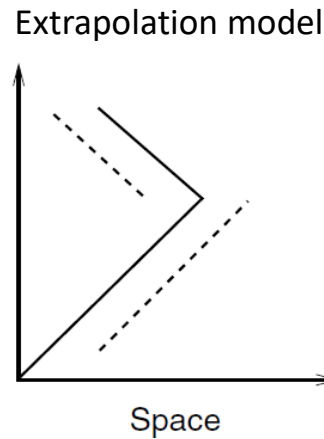
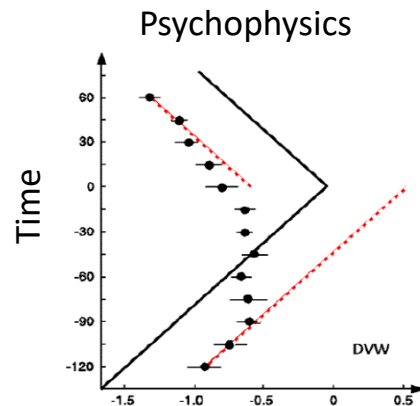
$$\begin{aligned} \hat{\mathbf{x}}_{k|T} &= \hat{\mathbf{x}}_{k|k} + \mathbf{G}_k \left( \hat{\mathbf{x}}_{k+1|T} - \hat{\mathbf{x}}_{k+1|k} \right) \\ \Sigma_{k|T} &= \Sigma_{k|k} + \mathbf{G}_k \left( \Sigma_{k+1|T} - \Sigma_{k+1|k} \right) \mathbf{G}_k^T \end{aligned} \quad \text{where} \quad \mathbf{G}_k = \Sigma_{k|k} \mathbf{A}^T \left( \Sigma_{k+1|k} \right)^{-1}$$



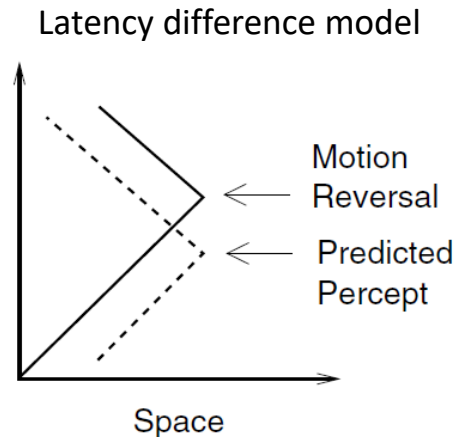
# Flash-lag effect explained by Kalman smoother.



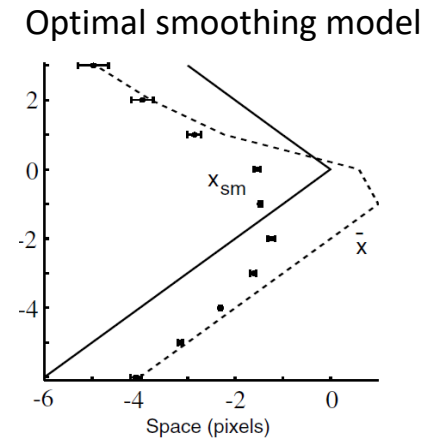
<http://visionlab.harvard.edu/Members/Alumni/David/flash-lag.htm>



Nijhawan (1994)



Whitney & Murakami (1998)



Rao et al. (2001)

# Attractor neural networks for ML estimation.

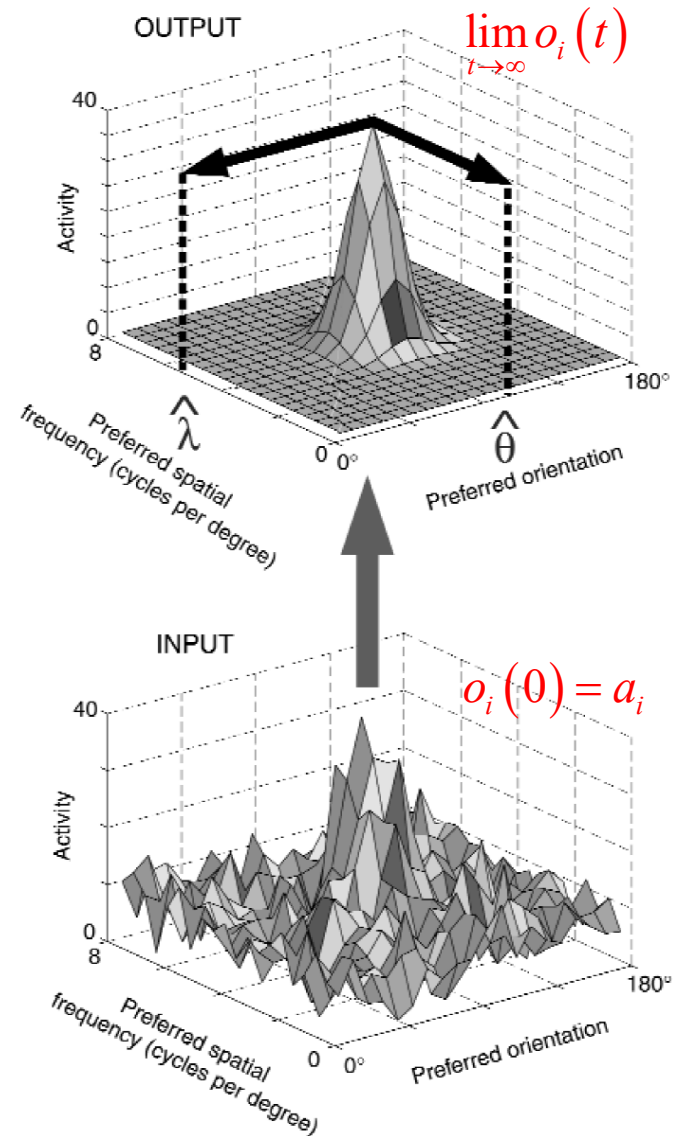
## Dynamics of recurrent neural network

### Local averaging

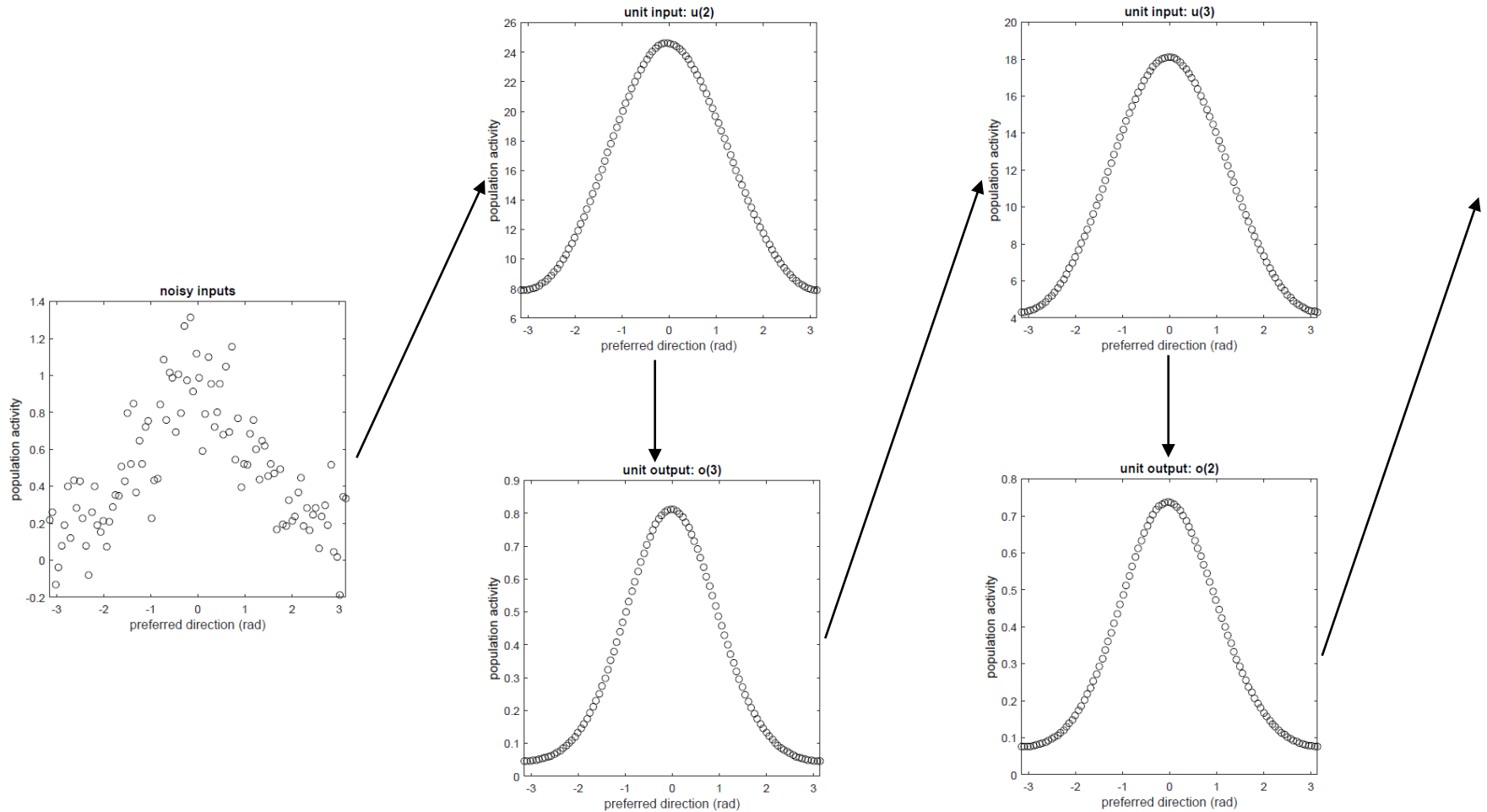
$$u_i(t+1) = \sum_j w_{ij} o_j(t)$$

### Divisive normalization (“winner-takes-all”)

$$o_i(t+1) = \frac{u_i^2(t+1)}{1 + \mu \sum_j u_j^2(t+1)}$$



# Attractor neural networks for ML estimation.



# Neural network algorithm for maximum likelihood.

Probabilistic neural responses to stimulus  $s$ :

$$n_i \sim \text{Poisson}(f_i(s))$$

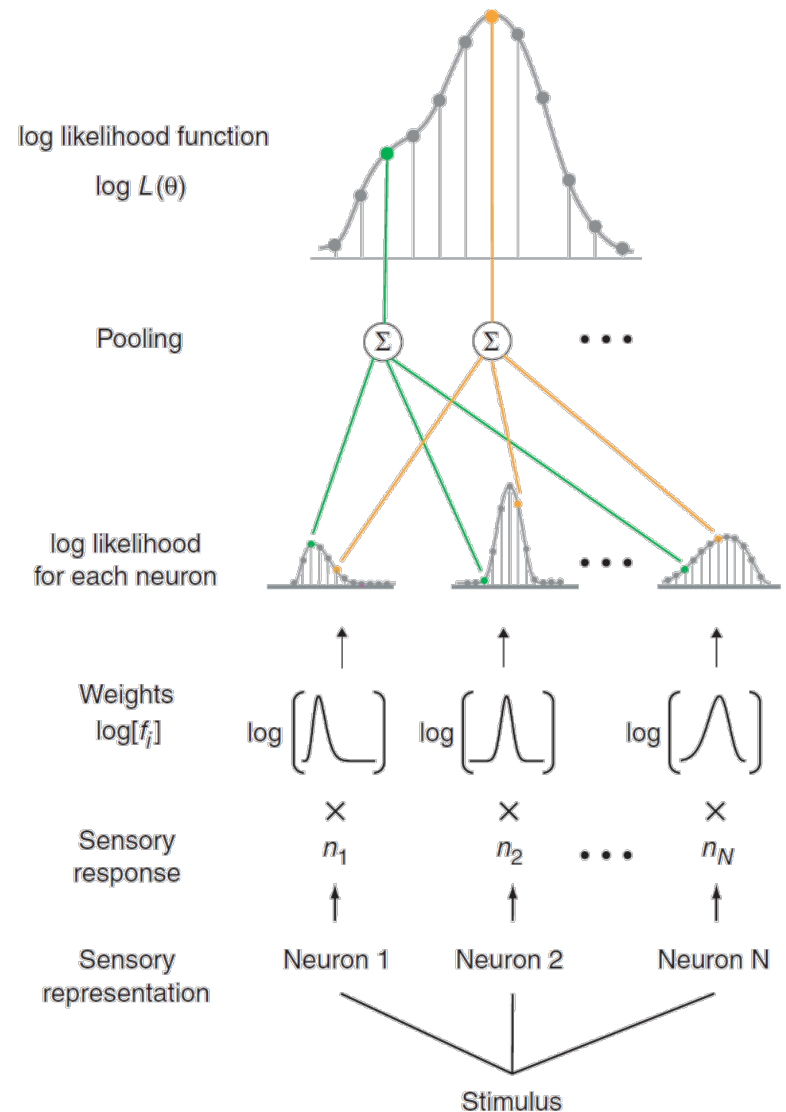
$$P(r_i | s) = \frac{e^{-f_i(s)} f_i(s)^{n_i}}{n_i!}$$

Likelihood function:

$$P(\mathbf{n} | s) = \prod_{i=1}^N P(n_i | s) = \prod_{i=1}^N \frac{e^{-f_i(s)} f_i(s)^{n_i}}{n_i!}$$

Log likelihood function:

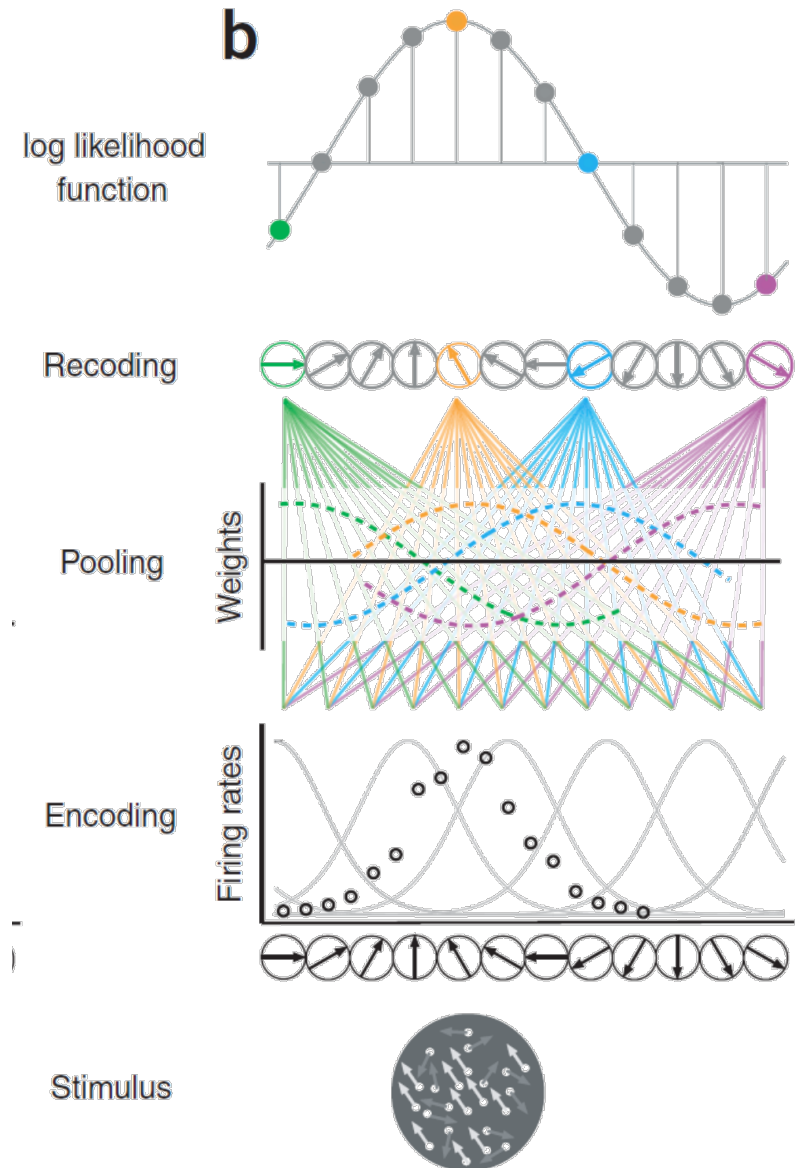
$$\begin{aligned} \log P(\mathbf{n} | s) &= \sum_{i=1}^N \log P(n_i | s) \\ &= -\sum_{i=1}^N f_i(s) + \sum_{i=1}^N n_i \log f_i(s) + \sum_{i=1}^N n_i! \end{aligned}$$



# Neural network algorithm for maximum likelihood.

$$\begin{aligned}\log P(\mathbf{r} | s) &= \sum_{i=1}^N \log P(r_i | s) \\ &= -\sum_{i=1}^N f_i(s) + \sum_{i=1}^N r_i \log f_i(s) + \sum_{i=1}^N r_i!\end{aligned}$$

$$\begin{aligned}\text{Maximize} \quad & \sum_{i=1}^N r_i \log f_i(s) \\ \text{under constraint} \quad & \sum_{i=1}^N f_i(s) = \text{constant}\end{aligned}$$



# Neural implementation of multimodal integration.

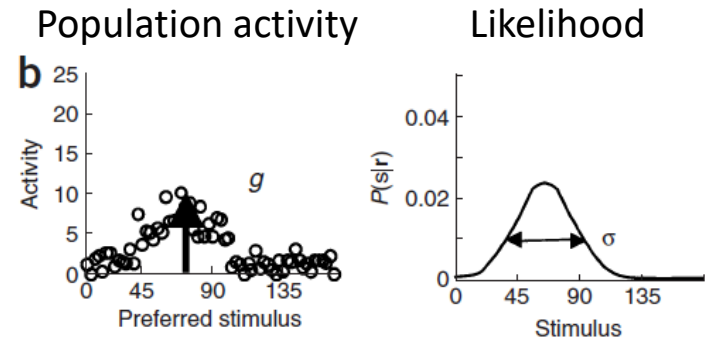
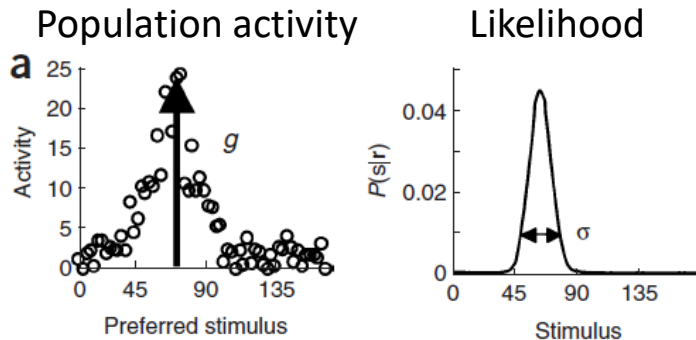
If tuning function is Gaussian with mean  $s_i$  and variance  $\sigma_0^2$ ,

$$f_i(s) = \mathcal{N}(s_i, \sigma_0^2) = \frac{1}{\sqrt{2\pi\sigma_0^2}} \exp\left(-\frac{(s - s_i)^2}{2\sigma_0^2}\right),$$

then, the likelihood function is also Gaussian:

$$\log P(\mathbf{n} | s) = \sum_{i=1}^N n_i \log f_i(s) + \text{const.} = -\sum_{i=1}^N \frac{n_i}{2\sigma_0^2} (s - s_i)^2 + \text{const.}$$

$$\therefore P(\mathbf{n} | s) = \mathcal{N}(s; \mu, \sigma^2) \quad \text{where} \quad \mu = \frac{\sum_{i=1}^N n_i s_i}{\sum_{i=1}^N n_i}, \quad \sigma^2 = \frac{\sigma_0^2}{\sum_{i=1}^N n_i}$$

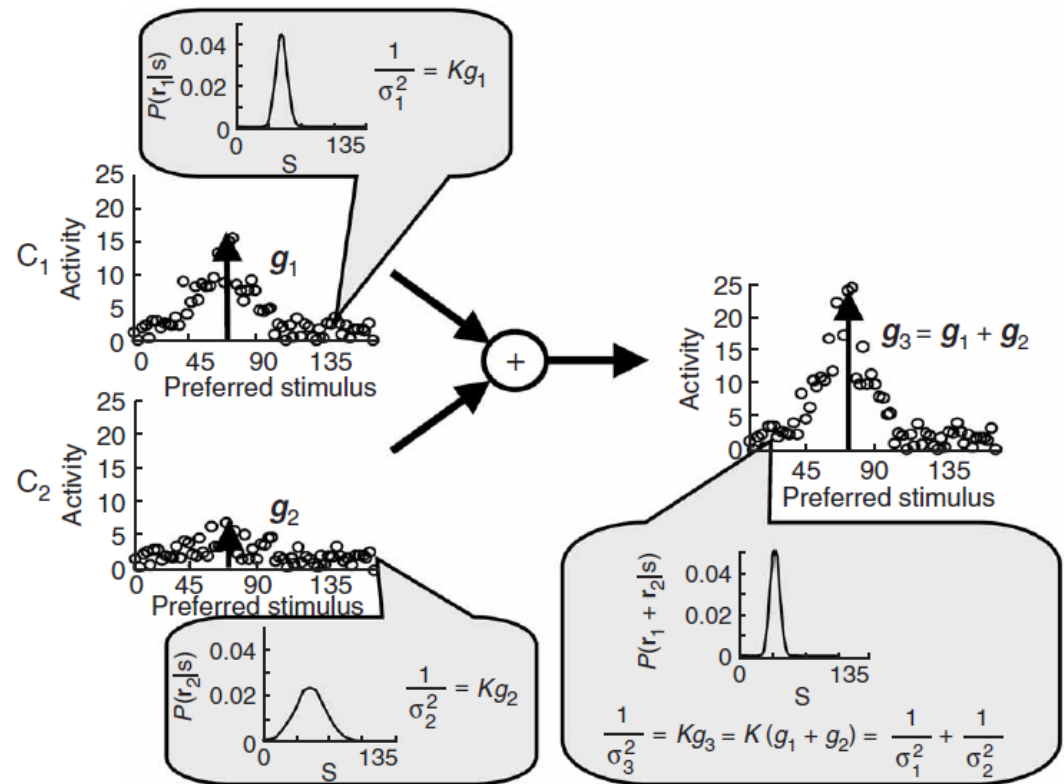


# Neural implementation of multimodal integration.

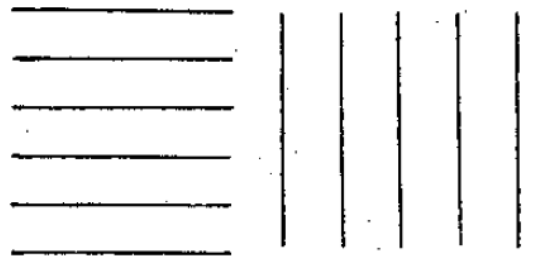
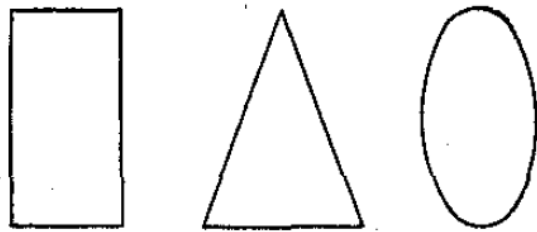
$$\begin{aligned}
 P(\mathbf{n}^{(3)} | s) &= P(\mathbf{n}^{(1)} + \mathbf{n}^{(2)} | s) \\
 &= P(\mathbf{n}^{(1)} | s) P(\mathbf{n}^{(2)} | s) \\
 &= \mathcal{N}(\mu_1, \sigma_1^2) \mathcal{N}(\mu_2, \sigma_2^2) \\
 &= \mathcal{N}(\mu_3, \sigma_3^2)
 \end{aligned}$$

$$\mu_3 = \frac{\sigma_2^2}{\sigma_1^2 + \sigma_2^2} x_1 + \frac{\sigma_1^2}{\sigma_1^2 + \sigma_2^2} x_2$$

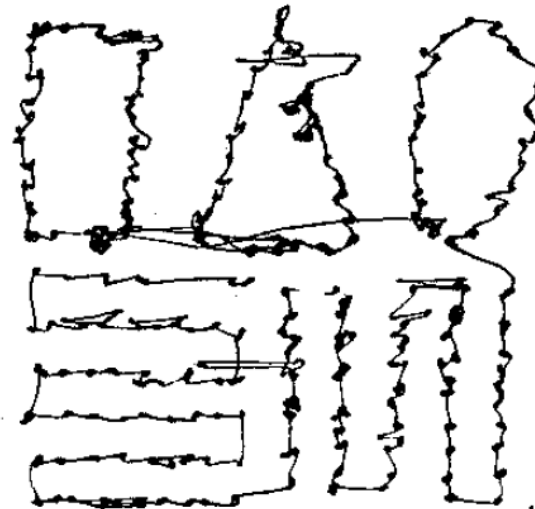
$$\frac{1}{\sigma_3^2} = \frac{1}{\sigma_1^2} + \frac{1}{\sigma_2^2}$$



# Eye movements are goal oriented.

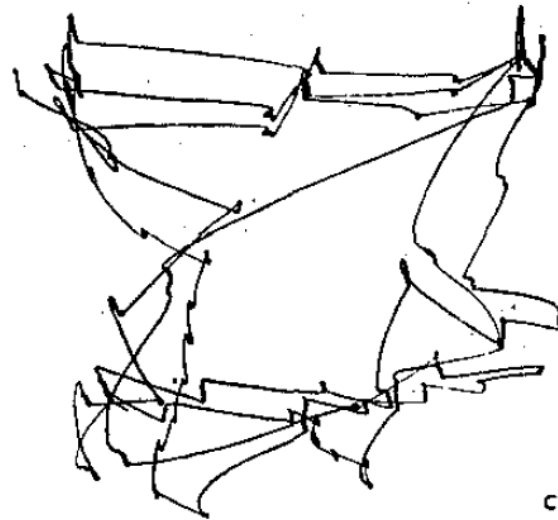


a

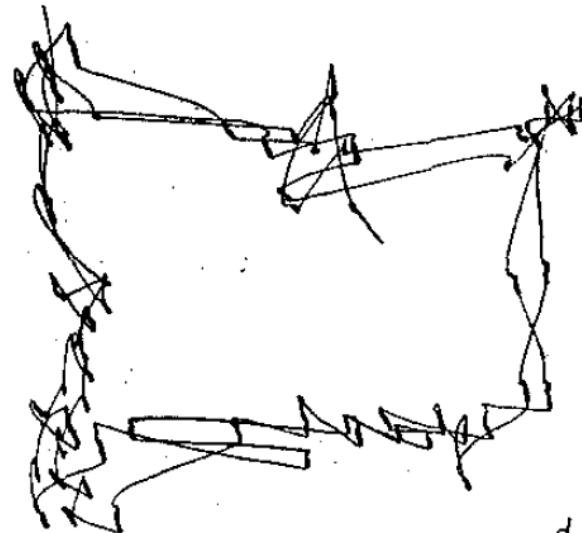


"Trace the lines smoothly."

b



c



d

"Count the number of straight lines"

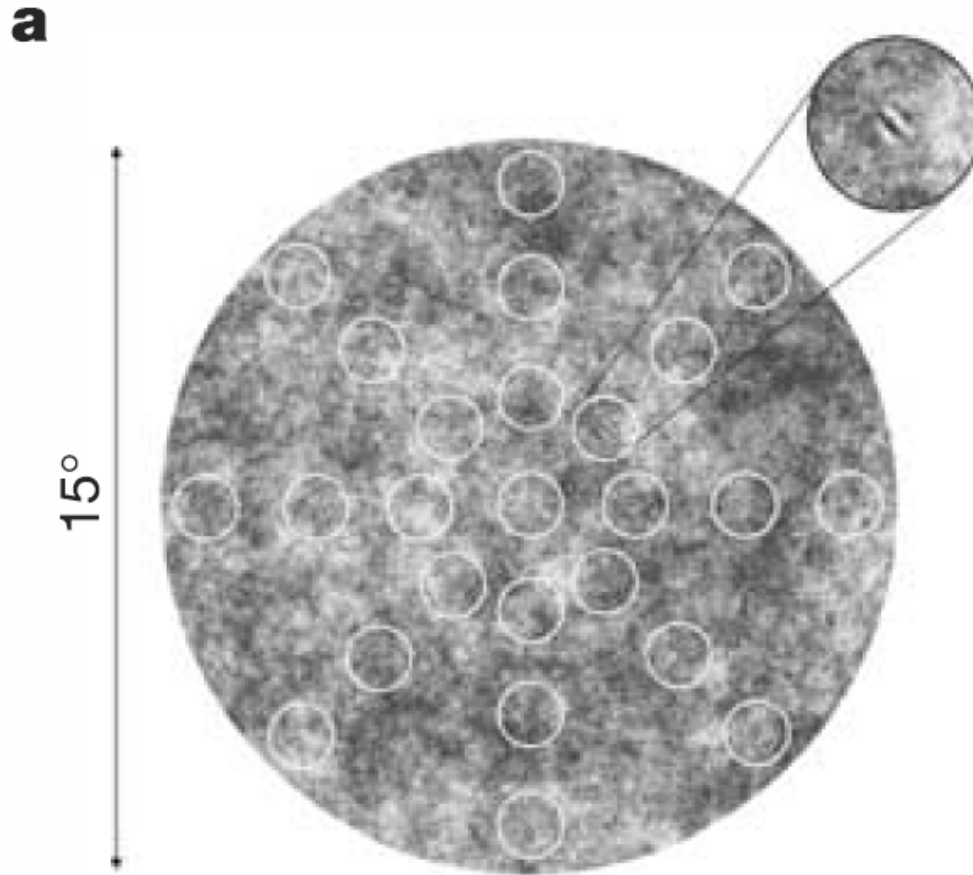
Original figures

Free viewing  
(no instruction)



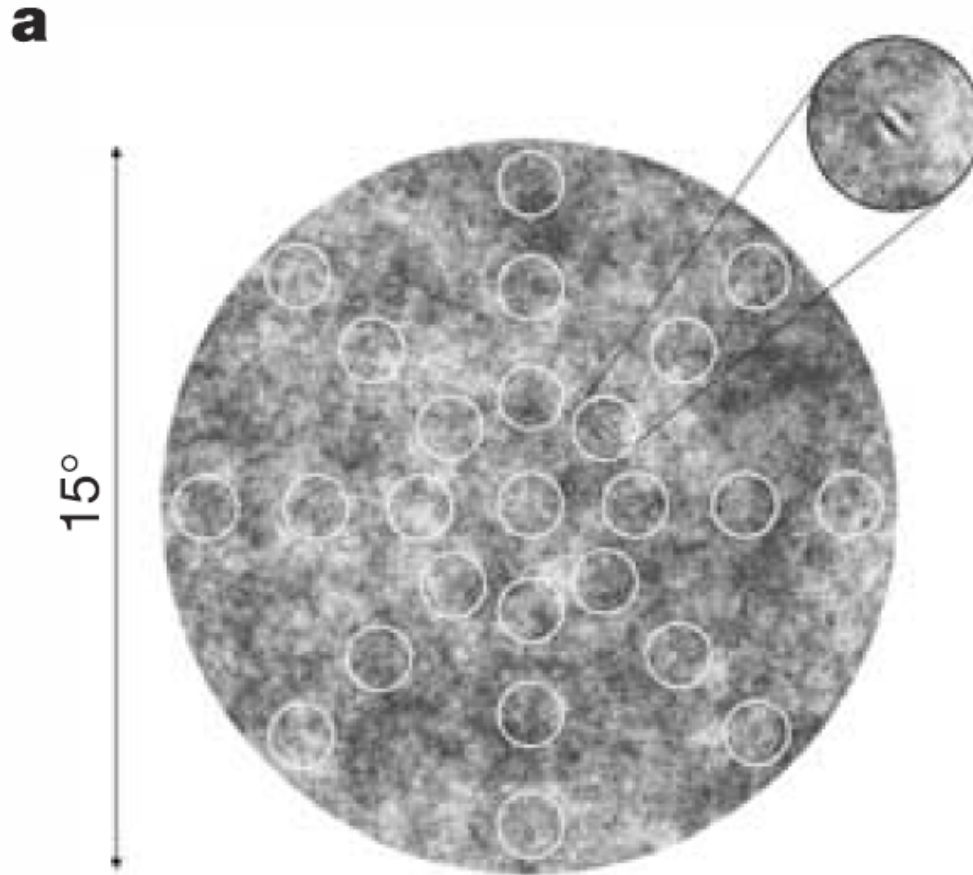
# Eye movements as optimal decision making

“Look for the Gabor patch in the pink-noise background.”



# Eye movements as optimal decision making

“Look for the Gabor patch in the pink-noise background.”



# Summary

- Inference from noisy measurements is formulated as a statistical inference problem.
- Statistical inference is categorized into classical point estimation and Bayesian probability estimation.
- Human performance in psychophysical experiments matches the predictions of optimal estimation.
- The computation of optimal inference can be implemented by a few neural mechanisms.

# References

- Ernst, M. O., & Banks, M. S. (2002). Humans integrate visual and haptic information in a statistically optimal fashion. *Nature*, 415(6870), 429-433.
- Weiss, Y., Simoncelli, E. P., & Adelson, E. H. (2002). Motion illusions as optimal percepts. *Nature neuroscience*, 5(6), 598-604.
- Kording, K. P., & Wolpert, D. M. (2004). Bayesian integration in sensorimotor learning. *Nature*, 427(6971), 244-247.
- Kording, K. P., Beierholm, U., Ma, W. J., Quartz, S., Tenenbaum, J. B., & Shams, L. (2007). Causal inference in multisensory perception.
- Kayser, C., & Shams, L. (2015). Multisensory causal inference in the brain. *PLoS biology*, 13(2).
- Burge, J., Ernst, M. O., & Banks, M. S. (2008). The statistical determinants of adaptation rate in human reaching. *Journal of Vision*, 8(4), 20.
- de Xivry, J. J. O., Coppe, S., Blohm, G., & Lefevre, P. (2013). Kalman filtering naturally accounts for visually guided and predictive smooth pursuit dynamics. *The Journal of Neuroscience*, 33(44), 17301-17313.
- Rao, R. P., Eagleman, D. M., & Sejnowski, T. J. (2001). Optimal smoothing in visual motion perception. *Neural Computation*, 13(6), 1243-1253.
- Deneve, S., Latham, P. E., & Pouget, A. (1999). Reading population codes: a neural implementation of ideal observers. *Nature neuroscience*, 2(8), 740-745.
- Jazayeri, M., & Movshon, J. A. (2006). Optimal representation of sensory information by neural populations. *Nature neuroscience*, 9(5), 690-696.
- Ma, W. J., Beck, J. M., Latham, P. E., & Pouget, A. (2006). Bayesian inference with probabilistic population codes. *Nature neuroscience*, 9(11), 1432-1438.

## Exercise

- Write a Matlab code for Kalman filter and Kalman smoothing, and examine the difference.
- Psychophysics: Smooth eye pursuit experiment using GazeParser (<http://gazeparser.sourceforge.net/>).

# Appendix A: Derivation of Kalman filter equations (1)

- Strategy 1.
  - Compute the mean and average.
  - Intuitive but a bit heuristic.
- Strategy 2.
  - Use Bayes' formula.
  - Complete the square.
- Strategy 3.
  - Use Bayes' formula.
  - Use the formula for conditional normal probability.

# Appendix A: Derivation of Kalman filter equations (1)

- Prediction step:

$$\hat{\mathbf{x}}_{k+1|k} = \mathbb{E}[\mathbf{x}_{k+1} \mid \mathbf{z}_{1:k}] = \mathbb{E}[\mathbf{A}\mathbf{x}_k + \mathbf{w}_k \mid \mathbf{z}_{1:k}] = \mathbf{A}\mathbf{x}_{k|k}$$

$$\Sigma_{k+1|k} = \text{Var}[\mathbf{x}_{k+1} \mid \mathbf{z}_{1:k}]$$

$$= \mathbb{E}\left[\left(\mathbf{x}_{k+1} - \hat{\mathbf{x}}_{k+1|k}\right)\left(\mathbf{x}_{k+1} - \hat{\mathbf{x}}_{k+1|k}\right)^T \mid \mathbf{z}_{1:k}\right]$$

$$= \mathbb{E}\left[\left(\mathbf{A}\left(\mathbf{x}_k - \hat{\mathbf{x}}_{k|k}\right) + \mathbf{w}_k\right)\left(\mathbf{A}\left(\mathbf{x}_k - \hat{\mathbf{x}}_{k|k}\right) + \mathbf{w}_k\right)^T \mid \mathbf{z}_{1:k}\right]$$

$$= \mathbf{A}\mathbb{E}\left[\left(\mathbf{x}_k - \hat{\mathbf{x}}_{k|k}\right)\left(\mathbf{x}_k - \hat{\mathbf{x}}_{k|k}\right)^T \mid \mathbf{z}_{1:k}\right]\mathbf{A}^T + \mathbb{E}\left[\mathbf{w}_k\mathbf{w}_k^T \mid \mathbf{z}_{1:k}\right]$$

$$= \mathbf{A}\Sigma_{k|k}\mathbf{A}^T + \mathbf{Q}$$

$$\hat{\mathbf{x}}_{k+1|k} = \mathbf{A}\mathbf{x}_{k|k}$$

$$\Sigma_{k+1|k} = \mathbf{A}\Sigma_{k|k}\mathbf{A}^T + \mathbf{Q}$$

# Appendix A: Derivation of Kalman filter equations (1)

- Filter step:

$$\hat{\mathbf{x}}_{k+1|k+1} = \mathbb{E}[\mathbf{x}_{k+1} \mid \mathbf{z}_{1:k}] = \hat{\mathbf{x}}_{k+1|k} + \mathbf{K}_{k+1} (\mathbf{z}_{k+1} - \mathbf{C}\hat{\mathbf{x}}_{k+1|k})$$

$$\begin{aligned}\Sigma_{k+1|k+1} &= \text{Var}[\mathbf{x}_{k+1} \mid \mathbf{z}_{1:k+1}] \\ &= \mathbb{E}\left[\left(\mathbf{x}_{k+1} - \hat{\mathbf{x}}_{k+1|k+1}\right)\left(\mathbf{x}_{k+1} - \hat{\mathbf{x}}_{k+1|k+1}\right)^T \mid \mathbf{z}_{1:k+1}\right] \\ &= \dots \\ &= (\mathbf{I} - \mathbf{K}_{k+1}\mathbf{C})\Sigma_{k+1|k}(\mathbf{I} - \mathbf{K}_{k+1}\mathbf{C})^T + \mathbf{K}_{k+1}\mathbf{R}\mathbf{K}_{k+1}^T\end{aligned}$$

$$0 = \frac{1}{2} \frac{\partial}{\partial \mathbf{K}_{k+1}} \text{tr} \Sigma_{k+1|k+1} = -(\mathbf{I} - \mathbf{K}_{k+1}\mathbf{C})\Sigma_{k+1|k}\mathbf{C}^T + \mathbf{K}_{k+1}\mathbf{R}$$

$$\mathbf{K}_{k+1} = \Sigma_{k+1|k}\mathbf{C}^T (\mathbf{C}\Sigma_{k+1|k}\mathbf{C}^T + \mathbf{R})^{-1}$$



# Appendix A: Derivation of Kalman filter equations (1)

- Note on matrix derivatives.

$$\begin{aligned}\frac{\partial}{\partial K_{ij}} \text{tr}(\mathbf{KAK}^T) &= \frac{\partial}{\partial K_{ij}} \sum_{k,l,m} K_{kl} A_{lm} K_{mk}^T = \frac{\partial}{\partial K_{ij}} \sum_{k,l,m} K_{kl} A_{lm} K_{km} \\ &= \sum_{k,l,m} (\delta_{i,k} \delta_{j,l} A_{lm} K_{km} + K_{kl} A_{lm} \delta_{i,k} \delta_{m,j}) = \sum_m A_{jm} K_{im} + \sum_l K_{il} A_{lj} = (\mathbf{KA}^T)_{ij} + (\mathbf{KA})_{ij} = (2\mathbf{KA})_{ij}\end{aligned}$$

$$\frac{\partial}{\partial \mathbf{K}} \text{tr}(\mathbf{KAK}^T) = 2\mathbf{KA}$$

# Appendix A: Derivation of Kalman filter equations (1)

$$\hat{\mathbf{x}}_{k+1|k+1} = \left( \mathbf{\Sigma}_{k+1|k}^{-1} + \mathbf{C}^T \mathbf{R}^{-1} \mathbf{C} \right)^{-1} \left( \mathbf{\Sigma}_{k+1|k}^{-1} \hat{\mathbf{x}}_{k+1|k} + \mathbf{C}^T \mathbf{R}^{-1} \mathbf{z}_{k+1} \right)$$

$$\mathbf{\Sigma}_{k+1|k+1} = \left( \mathbf{\Sigma}_{k+1|k}^{-1} + \mathbf{C}^T \mathbf{R}^{-1} \mathbf{C} \right)^{-1}$$

$$\left( \mathbf{\Sigma}_{k+1|k}^{-1} + \mathbf{C}^T \mathbf{R}^{-1} \mathbf{C} \right)^{-1} \mathbf{\Sigma}_{k+1|k}^{-1} = \mathbf{I} - \mathbf{\Sigma}_{k+1|k} \mathbf{C}^T \left( \mathbf{R} + \mathbf{C} \mathbf{\Sigma}_{k+1|k} \mathbf{C}^T \right)^{-1} \mathbf{C} = \mathbf{I} - \mathbf{K}_{k+1} \mathbf{C}$$

$$\left( \mathbf{\Sigma}_{k+1|k}^{-1} + \mathbf{C}^T \mathbf{R}^{-1} \mathbf{C} \right)^{-1} \mathbf{C}^T \mathbf{R}^{-1} = \mathbf{\Sigma}_{k+1|k} \mathbf{C}^T \left( \mathbf{R} + \mathbf{C} \mathbf{\Sigma}_{k+1|k} \mathbf{C}^T \right)^{-1} = \mathbf{K}_{k+1}$$

$$\begin{aligned} \mathbf{\Sigma}_{k+1|k+1} &= \left( \mathbf{\Sigma}_{k+1|k}^{-1} + \mathbf{C}^T \mathbf{R}^{-1} \mathbf{C} \right)^{-1} \\ &= \mathbf{\Sigma}_{k+1|k} - \mathbf{\Sigma}_{k+1|k} \mathbf{C}^T \left( \mathbf{R} + \mathbf{C} \mathbf{\Sigma}_{k+1|k} \mathbf{C}^T \right)^{-1} \mathbf{C} \mathbf{\Sigma}_{k+1|k} \\ &= (\mathbf{I} - \mathbf{K}_{k+1} \mathbf{C}) \mathbf{\Sigma}_{k+1|k} \end{aligned}$$

# Appendix A: Derivation of Kalman filter equations (1)

- “Weighted average” expression:

$$\hat{\mathbf{x}}_{k+1|k+1} = \left( \boldsymbol{\Sigma}_{k+1|k}^{-1} + \mathbf{C}^T \mathbf{R}^{-1} \mathbf{C} \right)^{-1} \left( \boldsymbol{\Sigma}_{k+1|k}^{-1} \hat{\mathbf{x}}_{k+1|k} + \mathbf{C}^T \mathbf{R}^{-1} \mathbf{z}_{k+1} \right)$$
$$\boldsymbol{\Sigma}_{k+1|k+1} = \left( \boldsymbol{\Sigma}_{k+1|k}^{-1} + \mathbf{C}^T \mathbf{R}^{-1} \mathbf{C} \right)^{-1}$$

- “Innovation-correction” expression:

$$\hat{\mathbf{x}}_{k+1|k+1} = \hat{\mathbf{x}}_{k+1|k} + \mathbf{K}_{k+1} \left( \mathbf{z}_{k+1} - \mathbf{C} \hat{\mathbf{x}}_{k+1|k} \right)$$
$$\boldsymbol{\Sigma}_{k+1|k+1} = \left( \mathbf{I} - \mathbf{K}_{k+1} \mathbf{C} \right) \boldsymbol{\Sigma}_{k+1|k}$$
$$\mathbf{K}_{k+1} = \boldsymbol{\Sigma}_{k+1|k} \mathbf{C}^T \left( \mathbf{C} \boldsymbol{\Sigma}_{k+1|k} \mathbf{C}^T + \mathbf{R} \right)^{-1}$$

# Matrix identities

- For matrices  $\mathbf{A}$ ,  $\mathbf{B}$ ,  $\mathbf{C}$  and  $\mathbf{D}$ :

$$\mathbf{A} \in \mathbb{R}^{n \times n}, \mathbf{B} \in \mathbb{R}^{n \times m}, \mathbf{C} \in \mathbb{R}^{m \times m}, \mathbf{D} \in \mathbb{R}^{m \times n},$$

$$(\mathbf{A} + \mathbf{BCD})^{-1} = \mathbf{A}^{-1} - \mathbf{A}^{-1}\mathbf{B}(\mathbf{C}^{-1} + \mathbf{DA}^{-1}\mathbf{B})^{-1}\mathbf{DA}^{-1}$$

$$\mathbf{A}^{-1}\mathbf{B}(\mathbf{A} + \mathbf{BCD})^{-1} = (\mathbf{C}^{-1} + \mathbf{DA}^{-1}\mathbf{B})^{-1}\mathbf{DA}^{-1}$$

- Rank-1 modification

$$(\mathbf{A} + \mathbf{uv}^T)^{-1} = \mathbf{A}^{-1} - \frac{\mathbf{A}^{-1}\mathbf{uv}^T\mathbf{A}^{-1}}{1 + \mathbf{v}^T\mathbf{A}^{-1}\mathbf{u}}$$

## Appendix B: Kalman filter in continuous-time dynamics

- Let us start with a discrete-time state-space model:

$$\mathbf{x}_{k+1} = \mathbf{A}\mathbf{x}_k + \mathbf{w}_k$$

$$\mathbf{z}_k = \mathbf{C}\mathbf{x}_k + \mathbf{v}_k$$

If the matrix  $\mathbf{A}$  is expanded for a small time step  $dt$ , the prediction update reads as:

$$\hat{\mathbf{x}}_{k+1|k} = (\mathbf{I} + \bar{\mathbf{A}}dt) \hat{\mathbf{x}}_{k|k}$$

$$\Sigma_{k+1|k} = (\mathbf{I} + \bar{\mathbf{A}}dt) \Sigma_{k|k} (\mathbf{I} + \bar{\mathbf{A}}dt)^T + \mathbf{Q} = \Sigma_{k|k} + (\bar{\mathbf{A}}\Sigma_{k|k} + \Sigma_{k|k}\bar{\mathbf{A}}^T)dt + \mathbf{Q}$$

From the covariance update, the matrix  $\mathbf{Q}$  should be of the order of  $dt$ :

$$\mathbf{Q} = \bar{\mathbf{Q}}dt$$

then, the covariance update now becomes:

$$\Sigma_{k+1|k} = \Sigma_{k|k} + (\bar{\mathbf{A}}\Sigma_{k|k} + \Sigma_{k|k}\bar{\mathbf{A}}^T + \bar{\mathbf{Q}})dt$$

# Appendix B: Kalman filter in continuous-time dynamics

- Next, consider the filter step:

$$\hat{\mathbf{x}}_{k+1|k+1} = \hat{\mathbf{x}}_{k+1|k} + \mathbf{K}_{k+1} \left( \mathbf{z}_{k+1} - \mathbf{C}\hat{\mathbf{x}}_{k+1|k} \right)$$

$$\begin{aligned} \Sigma_{k+1|k+1} &= (\mathbf{I} - \mathbf{K}_{k+1}\mathbf{C})\Sigma_{k+1|k}(\mathbf{I} - \mathbf{K}_{k+1}\mathbf{C})^T + \mathbf{K}_{k+1}\mathbf{R}\mathbf{K}_{k+1}^T \\ &= \underbrace{\Sigma_{k+1|k}}_{\mathcal{O}(1)} - \underbrace{\mathbf{K}_{k+1}\mathbf{C}\Sigma_{k+1|k}}_{\mathcal{O}(dt)} - \underbrace{\Sigma_{k+1|k}\mathbf{C}^T\mathbf{K}_{k+1}^T}_{\mathcal{O}(dt)} + \underbrace{\mathbf{K}_{k+1}\mathbf{C}\Sigma_{k+1|k}\mathbf{C}^T\mathbf{K}_{k+1}^T}_{\mathcal{O}(dt^2)} + \underbrace{\mathbf{K}_{k+1}\mathbf{R}\mathbf{K}_{k+1}^T}_{\mathcal{O}(dt)} \end{aligned}$$

Hence, it is appropriate to define:

$$\mathbf{R} = \frac{\bar{\mathbf{R}}}{dt}$$

In this limit, the Kalman gain is of the order of  $dt$ :

$$\mathbf{K}_{k+1} = \Sigma_{k+1|k}\mathbf{C}^T \left( \mathbf{C}\Sigma_{k+1|k}\mathbf{C}^T + \frac{\bar{\mathbf{R}}}{dt} \right)^{-1} = \Sigma_{k+1|k}\mathbf{C}^T\bar{\mathbf{R}}^{-1}dt + \mathcal{O}(dt^2)$$

## Appendix B: Kalman filter in continuous-time dynamics

- The prediction equations up to the order of  $dt$  are:

$$\hat{\mathbf{x}}_{k+1|k} = (\mathbf{I} + \bar{\mathbf{A}}dt) \hat{\mathbf{x}}_{k|k}$$

$$\Sigma_{k+1|k} = \Sigma_{k|k} + (\bar{\mathbf{A}}\Sigma_{k|k} + \Sigma_{k|k}\bar{\mathbf{A}}^T + \bar{\mathbf{Q}})dt$$

and the filter equations up to the order of  $dt$  are:

$$\hat{\mathbf{x}}_{k+1|k+1} = \hat{\mathbf{x}}_{k+1|k} + \mathbf{K}_{k+1} (\mathbf{z}_{k+1} - \mathbf{C}\hat{\mathbf{x}}_{k+1|k})$$

$$\Sigma_{k+1|k+1} = \Sigma_{k+1|k} - \Sigma_{k+1|k} \mathbf{C}^T \bar{\mathbf{R}}^{-1} \mathbf{C} \Sigma_{k+1|k} dt$$

- Combining these equations and eliminating the prediction variables, the update equations for the posterior mean and covariance are:

$$\hat{\mathbf{x}}_{k+1|k+1} = \hat{\mathbf{x}}_{k|k} + \left\{ \bar{\mathbf{A}}\hat{\mathbf{x}}_{k|k} + \Sigma_{k+1|k} \mathbf{C}^T \bar{\mathbf{R}}^{-1} (\mathbf{z}_{k+1} - \mathbf{C}\hat{\mathbf{x}}_{k+1|k}) \right\} dt$$

$$\Sigma_{k+1|k+1} = \Sigma_{k|k} + \left( \bar{\mathbf{A}}\Sigma_{k|k} + \Sigma_{k|k}\bar{\mathbf{A}}^T + \bar{\mathbf{Q}} - \Sigma_{k|k} \mathbf{C}^T \bar{\mathbf{R}}^{-1} \mathbf{C} \Sigma_{k|k} \right) dt$$

## Appendix B: Kalman filter in continuous-time dynamics

- By taking the limit of infinitesimal  $dt$

$$\hat{\mathbf{x}}_{k+1|k+1} = \hat{\mathbf{x}}_{k|k} + \left\{ \bar{\mathbf{A}}\mathbf{x}_{k|k} + \Sigma_{k+1|k} \mathbf{C}^T \bar{\mathbf{R}}^{-1} \left( \mathbf{z}_{k+1} - \mathbf{C}\hat{\mathbf{x}}_{k+1|k} \right) \right\} dt$$

$$\Sigma_{k+1|k+1} = \Sigma_{k|k} + \left( \bar{\mathbf{A}}\Sigma_{k|k} + \Sigma_{k|k} \bar{\mathbf{A}}^T + \bar{\mathbf{Q}} - \Sigma_{k|k} \mathbf{C}^T \bar{\mathbf{R}}^{-1} \mathbf{C} \Sigma_{k|k} \right) dt$$

and introducing continuous-time variables as

$$\hat{\mathbf{x}}(t) \equiv \hat{\mathbf{x}}_{k|k}, \quad \Sigma(t) \equiv \Sigma_{k|k}, \quad \mathbf{z}(t) \equiv \mathbf{z}_k,$$

the equations of Kalman filter for a continuous-time system are obtained as:

$$\dot{\hat{\mathbf{x}}}(t) = \bar{\mathbf{A}}\hat{\mathbf{x}}(t) + \Sigma(t) \mathbf{C}^T \bar{\mathbf{R}}^{-1} \left( \mathbf{z}(t) - \mathbf{C}\hat{\mathbf{x}}(t) \right)$$

$$\dot{\Sigma}(t) = \bar{\mathbf{A}}\Sigma(t) + \Sigma(t) \bar{\mathbf{A}}^T + \bar{\mathbf{Q}} - \Sigma(t) \mathbf{C}^T \bar{\mathbf{R}}^{-1} \mathbf{C} \Sigma(t)$$



# Matrix formulas

- Matrix inversion lemma

$$(\mathbf{A} + \mathbf{BCD})^{-1} = \mathbf{A}^{-1} - \mathbf{A}^{-1}\mathbf{B}(\mathbf{C}^{-1} + \mathbf{DA}^{-1}\mathbf{B})^{-1}\mathbf{DA}^{-1}$$

$$\begin{aligned} & (\mathbf{A} + \mathbf{BCD})(\mathbf{A} + \mathbf{BCD})^{-1} \\ &= (\mathbf{A} + \mathbf{BCD}) \left[ \mathbf{A}^{-1} - \mathbf{A}^{-1}\mathbf{B}(\mathbf{C}^{-1} + \mathbf{DA}^{-1}\mathbf{B})^{-1}\mathbf{DA}^{-1} \right] \\ &= \mathbf{I} + \mathbf{BCDA}^{-1} - (\mathbf{I} + \mathbf{BCDA}^{-1})\mathbf{B}(\mathbf{C}^{-1} + \mathbf{DA}^{-1}\mathbf{B})^{-1}\mathbf{DA}^{-1} \\ &= \mathbf{I} + \mathbf{BCDA}^{-1} - \mathbf{B}(\mathbf{I} + \mathbf{CDA}^{-1}\mathbf{B})(\mathbf{C}^{-1} + \mathbf{DA}^{-1}\mathbf{B})^{-1}\mathbf{DA}^{-1} \\ &= \mathbf{I} + \mathbf{BCDA}^{-1} - \mathbf{BC}(\mathbf{C}^{-1} + \mathbf{DA}^{-1}\mathbf{B})(\mathbf{C}^{-1} + \mathbf{DA}^{-1}\mathbf{B})^{-1}\mathbf{DA}^{-1} \\ &= \mathbf{I} + \mathbf{BCDA}^{-1} - \mathbf{BCDA}^{-1} \\ &= \mathbf{I} \end{aligned}$$

$$(\mathbf{I} + \mathbf{BCDA}^{-1})\mathbf{B} = \mathbf{B}(\mathbf{I} + \mathbf{CDA}^{-1}\mathbf{B})$$

$$(\mathbf{I} + \mathbf{CDA}^{-1}\mathbf{B}) = \mathbf{C}(\mathbf{C}^{-1} + \mathbf{DA}^{-1}\mathbf{B})$$

# Matrix formulas

- Matrix inversion lemma

$$(\mathbf{A} + \mathbf{BCD})^{-1} \mathbf{BC} = \mathbf{A}^{-1} \mathbf{B} (\mathbf{C}^{-1} + \mathbf{DA}^{-1} \mathbf{B})^{-1}$$

$$\begin{aligned} & (\mathbf{A} + \mathbf{BCD})^{-1} \mathbf{BC} \\ &= \mathbf{A}^{-1} (\mathbf{I} + \mathbf{BCDA}^{-1})^{-1} \mathbf{BC} \\ &= \mathbf{A}^{-1} \mathbf{B} (\mathbf{I} + \mathbf{CDA}^{-1} \mathbf{B})^{-1} \mathbf{C} \\ &= \mathbf{A}^{-1} \mathbf{B} (\mathbf{C}^{-1} + \mathbf{DA}^{-1} \mathbf{B})^{-1} \end{aligned}$$

$$(\mathbf{I} + \mathbf{BCDA}^{-1}) \mathbf{B} = \mathbf{B} (\mathbf{I} + \mathbf{CDA}^{-1} \mathbf{B})$$

$$(\mathbf{I} + \mathbf{CDA}^{-1} \mathbf{B}) = \mathbf{C} (\mathbf{C}^{-1} + \mathbf{DA}^{-1} \mathbf{B})$$

# Recursive least squares

- Least squares problem:

$$\sum_{k=1}^n (y_k - \mathbf{w}^T \mathbf{x}_k)^2$$

$$\hat{\mathbf{w}}_n = \arg \min_{\mathbf{w}} \sum_{k=1}^n (y_k - \mathbf{w}^T \mathbf{x}_k)^2 = \left( \sum_{k=1}^n \mathbf{x}_k \mathbf{x}_k^T \right)^{-1} \left( \sum_{k=1}^n y_k \mathbf{x}_k^T \right)$$

- We would like a recursive formula for  $n+1$  when the optimal weight vector up to  $n$  is available.

$$\hat{\mathbf{w}}_{n+1} = \arg \min_{\mathbf{w}} \sum_{k=1}^{n+1} (y_k - \mathbf{w}^T \mathbf{x}_k)^2 = \left( \sum_{k=1}^{n+1} \mathbf{x}_k \mathbf{x}_k^T \right)^{-1} \left( \sum_{k=1}^{n+1} y_k \mathbf{x}_k^T \right)$$

# Recursive least squares

- Least squares problem:

$$\hat{\mathbf{w}}_{n+1} = \left( \sum_{k=1}^n \mathbf{x}_k \mathbf{x}_k^T + \mathbf{x}_{n+1} \mathbf{x}_{n+1}^T \right)^{-1} \left( \sum_{k=1}^n y_k \mathbf{x}_k^T + y_{n+1} \mathbf{x}_{n+1}^T \right)$$

$$\begin{aligned} \hat{\mathbf{w}}_{n+1} &= \left( \mathbf{P}_n + \mathbf{x}_{n+1} \mathbf{x}_{n+1}^T \right)^{-1} \left( \sum_{k=1}^n y_k \mathbf{x}_k + y_{n+1} \mathbf{x}_{n+1} \right) \\ &= \left( \mathbf{P}_n + \mathbf{x}_{n+1} \mathbf{x}_{n+1}^T \right)^{-1} \sum_{k=1}^n y_k \mathbf{x}_k + \left( \mathbf{P}_n + \mathbf{x}_{n+1} \mathbf{x}_{n+1}^T \right)^{-1} \mathbf{x}_{n+1} y_{n+1} \\ &= \left( \mathbf{P}_n^{-1} - \frac{\mathbf{P}_n^{-1} \mathbf{x}_{n+1} \mathbf{x}_{n+1}^T \mathbf{P}_n^{-1}}{1 + \mathbf{x}_{n+1}^T \mathbf{P}_n^{-1} \mathbf{x}_{n+1}} \right) \sum_{k=1}^n y_k \mathbf{x}_k + \frac{\mathbf{P}_n^{-1} \mathbf{x}_{n+1}}{1 + \mathbf{x}_{n+1}^T \mathbf{P}_n^{-1} \mathbf{x}_{n+1}} y_{n+1} \\ &= \mathbf{P}_n^{-1} \sum_{k=1}^n y_k \mathbf{x}_k^T + \frac{\mathbf{P}_n^{-1} \mathbf{x}_{n+1}}{1 + \mathbf{x}_{n+1}^T \mathbf{P}_n^{-1} \mathbf{x}_{n+1}} \left( y_{n+1} - \hat{\mathbf{w}}_n^T \mathbf{x}_{n+1} \right) \\ &= \hat{\mathbf{w}}_{n+1} + \mathbf{K}_{n+1} \left( y_{n+1} - \hat{\mathbf{w}}_n^T \mathbf{x}_{n+1} \right) \end{aligned}$$

# Recursive least squares

- Least squares problem:

$$\mathbf{P}_1 = \mathbf{x}_1 \mathbf{x}_1^T$$

$$\mathbf{P}_{n+1}^{-1} = \mathbf{P}_n^{-1} - \frac{\mathbf{P}_n^{-1} \mathbf{x}_{n+1} \mathbf{x}_{n+1}^T \mathbf{P}_n^{-1}}{1 + \mathbf{x}_{n+1}^T \mathbf{P}_n^{-1} \mathbf{x}_{n+1}} = \left( \mathbf{I} - \mathbf{K}_{n+1} \mathbf{x}_{n+1}^T \right) \mathbf{P}_n^{-1}$$

$$\hat{\mathbf{w}}_{n+1} = \hat{\mathbf{w}}_n + \mathbf{K}_{n+1} \left( y_{n+1} - \hat{\mathbf{w}}_n^T \mathbf{x}_{n+1} \right)$$

$$\mathbf{K}_{n+1} = \frac{\mathbf{P}_n^{-1} \mathbf{x}_{n+1}}{1 + \mathbf{x}_{n+1}^T \mathbf{P}_n^{-1} \mathbf{x}_{n+1}}$$