

How Each Movement Changes the Next: An Experimental and Theoretical Study of Fast Adaptive Priors in Reaching

Timothy Verstynen and Philip N. Sabes

Department of Physiology, Keck Center for Integrative Neuroscience, University of California, San Francisco, California 94143

Most voluntary actions rely on neural circuits that map sensory cues onto appropriate motor responses. One might expect that for everyday movements, like reaching, this mapping would remain stable over time, at least in the absence of error feedback. Here we describe a simple and novel psychophysical phenomenon in which recent experience shapes the statistical properties of reaching, independent of any movement errors. Specifically, when recent movements are made to targets near a particular location subsequent movements to that location become less variable, but at the cost of increased bias for reaches to other targets. This process exhibits the variance–bias tradeoff that is a hallmark of Bayesian estimation. We provide evidence that this process reflects a fast, trial-by-trial learning of the prior distribution of targets. We also show that these results may reflect an emergent property of associative learning in neural circuits. We demonstrate that adding Hebbian (associative) learning to a model network for reach planning leads to a continuous modification of network connections that biases network dynamics toward activity patterns associated with recent inputs. This learning process quantitatively captures the key results of our experimental data in human subjects, including the effect that recent experience has on the variance–bias tradeoff. This network also provides a good approximation of a normative Bayesian estimator. These observations illustrate how associative learning can incorporate recent experience into ongoing computations in a statistically principled way.

Introduction

Experience has a profound effect on nearly all human behavior. For goal-directed behaviors such as reaching, the neural circuits that map sensory cues onto motor output continuously adapt to maintain behavioral stability in the face of external perturbations or internal noise (Shadmehr and Mussa-Ivaldi, 1994; Thoroughman and Shadmehr, 2000; Scheidt et al., 2001; Smith et al., 2006; Cheng and Sabes, 2007; Kluzik et al., 2008; Diedrichsen et al., 2010; Huang et al., 2011). There is evidence that these adaptive processes follow Bayesian statistical principles (Körding and Wolpert, 2004; Slijper et al., 2009; Wei and Körding, 2010), so that behavior is guided by both current sensory signals and a prior expectation of those signals derived from experience. Similar observations have been made for a variety of perceptual and cognitive processes (Weiss et al., 2002; Kersten et al., 2004; Miyazaki et al., 2005; Knill, 2007; Sato et al., 2007; Lages and Heron, 2008; Lu et al., 2008), suggesting that Bayesian principles may capture a general property of neural computation. Despite this wealth of experimental evidence, it is not well understood how neural circuits could learn such priors from recent experience.

To better understand how behavior is shaped by recent actions, we investigated a novel form of experience-dependent learning using visually guided reaching. We found that the sensorimotor system appears to maintain a prior expectation for movement planning that is continually updated based on the sequence of recent reaches. This phenomenon is consistent, at the qualitative level, with adaptive Bayesian estimation. We next explored our hypothesis that the activity of a sensorimotor network could itself create and maintain such priors via Hebbian learning. Specifically, we propose a model in which ongoing activity continuously modifies the structure of synaptic connections within a competitive neural network. These changes bias the network dynamics toward recent activity patterns, effectively creating a prior on the network computations. We show that this simple model accurately captures the results from several novel behavioral experiments, suggesting a potential mechanism for adaptive Bayesian estimation.

Materials and Methods

Experimental procedures

A total of 24 healthy, right-handed participants were tested (10 female, age range: 18–32 years). Subjects were paid for their participation and were naive to the purpose of the experiment. All the experimental procedures were approved by the University of California, San Francisco Human Research Protection Program. Subjects performed a series of trials in which they reached to visual targets in a virtual feedback apparatus (Fig. 1A) (Sober and Sabes, 2003). At the beginning of each trial, subjects placed the tip of their right index finger at a central start location positioned ~29 cm in front of the midline of their chest. We used an arrow-field paradigm to guide their finger to the start location without providing visual information about absolute position (Sober and Sabes, 2005). After the start location was reached and after a variable delay (500–1500 ms), the target appeared 12 cm from the start location (un-

Received Dec. 14, 2010; revised April 15, 2011; accepted May 17, 2011.

Author contributions: T.V. and P.N.S. designed research; T.V. performed research; T.V. and P.N.S. analyzed data; T.V. and P.N.S. wrote the paper.

This work was supported by NIH Grant P50 MH077970 (Conte Center) and the Swartz Foundation. We thank Charles Biddle-Snead for help with the network models.

The authors declare no competing financial interests.

Correspondence should be addressed to Philip N. Sabes, Department of Physiology, University of California, San Francisco, 513 Parnassus Avenue, San Francisco, CA 94143-0444. E-mail: sabes@phy.ucsf.edu.

DOI:10.1523/JNEUROSCI.6525-10.2011

Copyright © 2011 the authors 0270-6474/11/3110050-10\$15.00/0

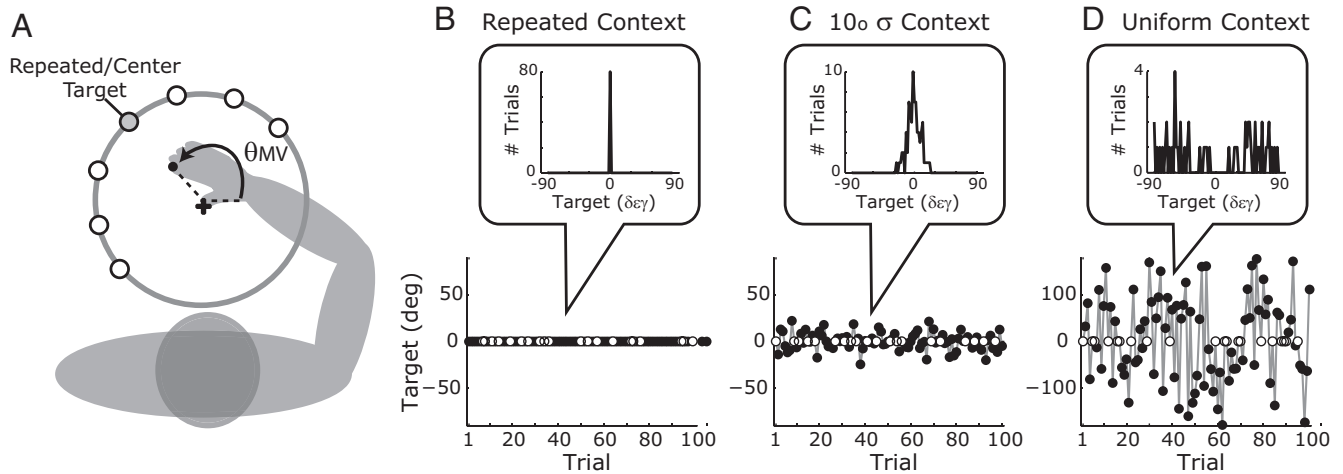


Figure 1. The reaching task. **A**, Subjects reached to visual targets with virtual visual feedback of the index fingertip (black dot) available ~ 100 ms after movement onset. For experiment 1, the central gray target is both the probe target and the center of the context-target distributions. For experiment 2, all seven probe targets were used and the center target was located either at 150° , as shown, or at 60° (randomized across subjects). The initial movement direction, θ_{MV} , was determined 100 ms after movement onset. **B–D**, Example trial blocks for three context conditions in experiment 1 (black circles, context trials; white circles, probe trials). Insets, Context target histograms.

filled circle, 15 mm in radius) and a “go tone” was played. Subjects were instructed to move as soon as possible and reach to put their finger in the center of the target circle. Once the finger had moved a quarter of the distance to the target, continuous feedback of finger position was displayed (filled white circle, 5 mm radius). Trials were terminated when the finger remained still in the target for 200 ms. At the end of each reach, participants received feedback in the form of a bonus score designed to encourage subjects to execute a quick, single, and accurate reach. The score was based on both reaction time and the distance between the target center and the location where the finger first decelerated below 25 mm/s. No bonus was given and a warning message appeared when the peak tangential velocity was <650 or >950 mm/s.

Participants in experiment 1 ($n = 8$, three female) were tested on eight blocks of 110 trials in a single session. Each block began with 10 context trials, followed by a randomized sequence of 80 context and 20 probe trials. Probe trial targets were fixed at $\theta = 150^\circ$ (relative to rightward axis) for all trial blocks. For the context trials, target angles were selected from a different distribution in each trial block: the repeated-target condition with all trials at the probe-target location; a normal distribution of targets with standard deviation $\sigma_{\text{Target}} = 1^\circ, 2^\circ, 3^\circ, 5^\circ, 10^\circ$, or 15° ; or a uniform distribution of targets on the circle.

Participants in experiment 2 ($n = 8$, three female) were tested on six blocks of 90 trials in each of four experimental sessions. Each block began with 10 context trials, followed by two repetitions of each probe target (14 trials) randomized with 66 context trials. A single context-target distribution was used for each session: the repeated target condition, a normal distribution with $\sigma_{\text{Target}} = 7.5^\circ$ or 15° , or a uniform target distribution. Seven probe target locations were used in this experiment. These were defined with respect to the repeated target location: $\theta - \theta_{\text{repeat}} = 0^\circ, \pm 30^\circ, \pm 60^\circ$, or $\pm 90^\circ$. Half of the subjects were tested with $\theta_{\text{repeat}} = 150^\circ$ and half with $\theta_{\text{repeat}} = 60^\circ$.

Participants in experiment 3 ($n = 8$, four female) were tested on six blocks of 120 trials. Targets were presented in sequential order, starting at 0° and proceeding in 3° increments around the circle. Trial blocks had either clockwise or counterclockwise target rotation and there were three blocks per condition with order randomized.

Data analysis

Movement trajectories were obtained from the position of an infrared LED located on the right index fingertip. Generally, positional data were not smoothed before analysis. However on $\sim 5\%$ of trials there were missing data samples due to obstructed view of the fingertip LED. For these trials, the missing datapoints were interpolated using a cubic spline method (*spline* in Matlab). However, removing these trials appeared to have no qualitative affect on the final results.

The initial movement angle, θ_{MV} , was computed as the angle between fingertip position at the beginning of the reach and its position 100 ms after reach onset. This timing was chosen because it samples the movement before feedback-control mechanisms can affect reach dynamics. Since the variability of initial movement angles for a given target θ is always small compared with the full circle, circular statistics are not needed; therefore, the standard deviation of initial angles, σ_{MV} , was computed in the usual linear way; analysis with circular statistics yielded nearly identical results (data not shown). For each subject, Grubbs’ test (Grubbs, 1950) was used to identify and exclude outlier trials from analysis (0.11% of trials excluded, with no more than three trials total excluded per subject). To quantify the movement error in the direction of the center probe target, θ_{repeat} , we defined the bias (toward θ_{repeat}) in terms of the angular errors for each pair of probe targets located at the same distance δ from the center target position,

$$\text{Bias} = \frac{1}{N_{-\delta} + N_{+\delta}} \left(\sum_{t=1}^{N_{-\delta}} (\theta_{MV,t} - \theta_{-\delta}) - \sum_{t=1}^{N_{+\delta}} (\theta_{MV,t} - \theta_{+\delta}) \right). \quad (1)$$

In experiment 3, we estimated the mean angular error, $\theta_{MV} - \theta_{\text{target}}$, separately for blocks with clockwise and counterclockwise target rotations.

Normative Bayesian model

We quantified the predictions of a simple Bayesian estimator (see Fig. 3) by computing the MAP (maximum a posteriori) estimate for a normal prior and likelihood with parameters fit to the experimental data. Bayes’ rule states

$$\underbrace{p(\theta|x)}_{\text{posterior}} \propto \underbrace{p(x|\theta)}_{\text{likelihood}} \underbrace{p(\theta)}_{\text{prior}}. \quad (2)$$

The true target θ gives rise to a sensory signal x , and we model the likelihood of θ given x as a normal distribution centered on θ with variance $\sigma_{\text{Likelihood}}^2$. The prior is modeled as a normal distribution with mean $\bar{\theta}$ and variance σ_{Prior}^2 . In this case, the peak of the posterior distribution, i.e., the MAP estimate of θ , is given by

$$\theta_{\text{MAP}} = \frac{\sigma_{\text{Posterior}}^2}{\sigma_{\text{Prior}}^2} \bar{\theta} + \frac{\sigma_{\text{Posterior}}^2}{\sigma_{\text{Likelihood}}^2} x, \quad (3)$$

where $\sigma_{\text{Posterior}}^2 = (\sigma_{\text{Prior}}^2 + \sigma_{\text{Likelihood}}^2)^{-1}$. If we assume that the mean of the prior is known without uncertainty, then the variance of the MAP estimate is given by

$$\sigma_{\text{MAP}}^2 = \sigma_{\text{Likelihood}}^2 (\sigma_{\text{Prior}}^2 + \sigma_{\text{Likelihood}}^2)^{-2}. \quad (4)$$

To compare the normative Bayesian model with empirical data, we assume that the empirical initial angle is a direct readout of θ_{MAP} and that the prior distribution accurately reflects the distribution of context trial targets (i.e., $\bar{\theta} = \theta_{\text{repeat}}$ and $\sigma_{\text{Prior}}^2 = \sigma_{\text{Context}}^2$). The only remaining free parameter of the model is $\sigma_{\text{Likelihood}}^2$, which we fit to the empirical data using Equation 4. This was done separately for each participant in experiment 1 (maximum likelihood fit using *fmincon* in Matlab). With this procedure, we obtained a mean $\sigma_{\text{Likelihood}}^2 = 7.2^\circ$. We then predicted the mean bias in experiment 2, $\theta_{\text{MAP}} - \theta$, using

$$\text{mean bias} = \frac{\sigma_{\text{Posterior}}^2}{\sigma_{\text{Prior}}^2} (\bar{\theta} - \theta). \quad (5)$$

Note, however, that this model and the fit value of $\sigma_{\text{Likelihood}}^2$ are only meant to provide qualitative comparisons to the data, as the assumptions that go into the model, in particular that the prior variance matches the context variance, are unlikely to be correct, as described below.

Adaptive Bayesian model

The normative Bayes model described above does not address how priors are learned or the rate of this learning. To estimate subjects' learning rates, we also implemented a simple adaptive version of the Bayesian model. For each trial n , the MAP estimate of the target location is computed using Equation 3, where the mean and variance of the prior are iteratively updated, given the current input θ_n , using the update rules

$$\begin{aligned} \bar{\theta}_{n+1} &= (1 - \beta)\bar{\theta}_n + \beta\theta_n \\ \sigma_{\text{Prior}, n+1}^2 &= (1 - \beta)\sigma_{\text{Prior}, n}^2 + \beta(\bar{\theta}_n - \theta_n)^2 \end{aligned} \quad (6)$$

where $\beta \in [0, 1]$ is the learning rate. The free parameters β and $\sigma_{\text{Likelihood}}^2$ were fit to the experimental data, minimizing the square error between the per trial MAP estimates of target location and subjects' movement errors (*fmincon* in Matlab). These fits were performed separately for each subject and each experimental session in experiment 2 (excluding the uniform condition). Cases where the fitted learning rate for a particular subject and session was either the maximum or minimum allowed value (0.999 and 0.001, respectively) were excluded from subsequent analysis (four of 24 fits). With this procedure, the mean estimate for $\sigma_{\text{Likelihood}}^2$ was 10.0° (SD, 6.9°), slightly larger than that found with the normative Bayesian model.

Adaptive network: model architecture

The network architecture and dynamics largely follow an approach used previously (Deneve et al., 1999, 2001, 2007; Pouget et al., 2002; Latham et al., 2003) and were designed to yield line-attractor dynamics. The network consists of a single layer of 180 neurons that receive activation from both exogenous inputs and lateral connections. On simulated trial n , the mean input activation to neuron i is a function of the difference between the target location $\theta(n)$ and the preferred direction of the neuron θ_i^* :

$$\bar{I}_i(n) = I_0 + \gamma \exp\left(-\ln(2)\left(\frac{(\theta(n) - \theta_i^*)}{\omega/2}\right)^2\right), \quad (7)$$

where ω determines the spread of input activation across the population [full-width at half-max (FWHM)], γ is the input gain, and I_0 is the baseline input rate. The preferred directions were evenly spaced at $\Delta = 2^\circ$ intervals. For network simulations, the dynamics of the network were iterated five times for each trial. The input activation for unit i on iteration t is given by $I_{i,t} = \bar{I}_i + \xi_{i,t}$, where $|x|_+$ is the positive part of x and the vector ξ_t of noise terms is a mean-zero multivariate Gaussian with covariance matrix

$$\Sigma_{ij} = \begin{cases} \sigma_i^2 = F\bar{I}_i & i=j \\ \sigma_{ij} = \sigma_i\sigma_j\rho\left(\frac{1}{2}\right)^{\frac{|\theta_i - \theta_j| - \Delta\theta}{\phi/2}} & i \neq j \end{cases} \quad (8)$$

The noise has a Fano factor F (ratio of variance to mean), a correlation coefficient ρ between nearest-neighbor units, and a correlation coefficient for other pairs that falls off with the distance between the two

Table 1. Parameter values for the adaptive network simulations

Initial weights	Input activation	Learning (Oja's rule)
$\omega_L = 38.8^\circ$	$\gamma = 13.4$ Hz	$\beta = 5.28 \times 10^{-7}$
	$\omega_i = 18.9^\circ$	$\alpha = 0.2$
	$I_0 = 5$ Hz	$\epsilon = 0.646$
	$F = 2$	
	$\phi = 5.49^\circ$	

Parameters were fit to the experimentally observed dependence of reach variance on context and probe target location (Fig. 7A, C).

preferred directions, with a FWHM of ϕ . With the parameter values used in our simulations (Table 1), we observed an average neuron–neuron correlation coefficient of 0.10 in the input activations.

In addition to the input activation, units received recurrent activation via lateral connections within the network. The initial connection strength between a pair of units is determined by the distance between the units' preferred directions in the circular space

$$W_{ij}^0 = \exp\left(-\ln(2)\left(\frac{\theta_i^* - \theta_j^*}{\omega_L/2}\right)^2\right), \quad (9)$$

where ω_L is the FWHM. On each iteration t , the total activation for unit i was then computed as the sum of the input and recurrent activations:

$$U_{i,t} = I_{i,t} + \sum_j W_{ij} X_{j,t-1}, \quad t = 1 \dots 5, \quad (10)$$

where $X_{i,t-1}$ are the firing rates on the previous iteration, with initial values $X_{i,0} = 0$. Total activation was normalized on each iteration to obtain the new firing rates:

$$X_{i,t} = \frac{U_{i,t}^2}{a + b \sum_i U_{i,t}^2}, \quad (11)$$

where the parameters a and b were set to the values 0.002 and 0.001, respectively.

After each iteration, the recurrent weights W_{ij} were updated using a normalized Hebbian learning rule (Oja, 1982):

$$W_{ij,t+1} = W_{ij,t} + \beta(X_{i,t} X_{j,t} - \alpha X_{i,t} X_{j,t} W_{ij,t}), \quad (12)$$

where β is the learning rate and α is the normalization parameter (Oja's α).

At the end of each trial, the estimated movement vector (θ_{MV}) was calculated using a population vector decode (Georgopoulos et al., 1988):

$$\theta_{\text{MV}} = \arctan\left(\frac{\sum_i X_{i,5} [\cos\theta_i^*, -\sin\theta_i^*]}{\sum_i X_{i,5} [\sin\theta_i^*, \cos\theta_i^*]}\right). \quad (13)$$

We chose to use five iterations per simulated trial because this decode converged after ~ 3 – 4 network iterations.

Adaptive network: model fitting

The network model parameters used in our simulations are shown in Table 1. These values were obtained by fitting the model to the context-dependent changes in movement variance observed experimentally. Specifically, we fit the model to the variance curves in Figures 2A (experiment 1) and 5A, B (experiment 2) using the full-experiment simulations (see below); the parameters were selected to maximize the sum of the R^2 values for each plot. For computational efficiency, this optimization was performed in three steps. First, using a reduced network with $N = 90$ neurons, we ran a large number of full-experiment simulations with random parameter values (~ 3000 runs). We then qualitatively determined the subset of the parameters that did not correlate well with the goodness of fit: the baseline firing rate I_0 , Fano factor F , and the Oja's α . In the second stage, we ran a large number of optimization runs for the remaining parameters on the $N = 90$ network with different initial conditions, using a generic nonlinear optimization routine (Matlab's *fmincon*).

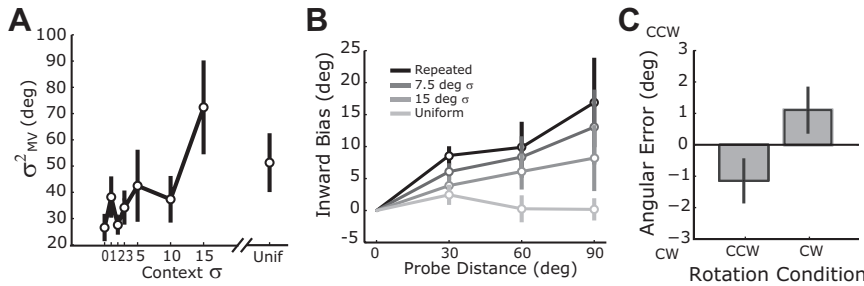


Figure 2. Experimental results. **A**, Reach variability of the initial movement direction in experiment 1. **B**, Reach bias in experiment 2, with positive values reflecting bias toward the repeated-target position. The value of zero bias at the repeated-target position (0°) is nominal, since bias is defined as the error toward that location; there was no significant change across context conditions in the average angular error at this target ($F_{(2,21)} = 1.13, p = 0.34$). **C**, Angular error in the CW and CCW trial blocks of experiment 3. Error bars represent SEs.

con). From all these runs, we selected the best $N = 90$ network. Finally, we increased the number of neurons from 90 to 180 and reoptimized the parameters from the best network, with the scale-dependent parameters τ and β adjusted inversely proportional to network size.

Adaptive network: simulations

Blockwise simulations. We first measured the effect of context target distribution by independently simulating blocks of trials with different input distributions. Each simulation consisted of a series of 100 training trials with the inputs $\theta(n)$ drawn from a given distribution followed by a series of probe trials in which learning was turned off ($\beta = 0$). All context distributions were centered at 0°, which we refer to as the “repeated target angle”. Probe inputs were presented from -140° to $+140^\circ$ in 20° intervals and each input was repeated 100 times. From each simulation, we computed the bias and trial-by-trial variance of the network output θ_{MV} for the probe targets. For each context distribution, we repeated this simulation 50 times and computed the average variance and bias.

Full-experiment simulations. For each behavioral experiment, we ran the network with the exact sequence of targets that each human subject experienced. The sequence included both context and probe trials. To allow for some unlearning between trial blocks, the network weights were decayed back toward their initial values, W_{ij}^0 (Eq. 12), after each block:

$$W_{ij} \leftarrow (1 - \epsilon) W_{ij} + \epsilon W_{ij}^0, \quad (14)$$

where the parameter ϵ determines the degree of unlearning. The network output was analyzed in the same manner as the experimental data.

Comparison of adaptive network and normative Bayesian models

Lastly, we quantified how closely the network model approximates the normative Bayesian model. Specifically, we set out to ask whether the bias and variance of the network output changes as predicted by the Bayesian model as a function of the context and likelihood variances. The network was trained using the blockwise protocol described above with four different context distributions and tested with a series of probe targets and with different input gains, γ (see Fig. 8C–F).

We first equate output of the network to the MAP estimate of the normative Bayesian model that it effectively implements. Thus, the output variance of the network for a given gain γ and context, $\sigma_{\text{Output}|\gamma, \sigma_{\text{Context}}}^2$, is a measure of the σ_{MAP}^2 for those parameters. The effective likelihood variance for gain γ , $\sigma_{\text{Likelihood}|\gamma}^2$, was thus determined by measuring the output variability of the network in the absence of a prior, i.e., with the network trained in the uniform context. The effective variance of the prior for a given training context, $\sigma_{\text{Prior}|\sigma_{\text{Context}}}^2$, was then computed from Equation 4 using the network simulations with gain $\gamma = 60$. In summary, the procedure for determining the effective parameters is given by the following two expressions:

$$\begin{aligned} \sigma_{\text{Likelihood}|\gamma}^2 &= \sigma_{\text{Output}|\gamma, \sigma_{\text{Context}=0}}^2 \\ \sigma_{\text{Prior}|\sigma_{\text{Context}}}^2 &= (\sigma_{\text{Output}|\gamma=60, \sigma_{\text{Context}}}^{-1} \sigma_{\text{Likelihood}|\gamma}^{-1} \sigma_{\text{Likelihood}|\gamma}^2)^{-1} \end{aligned} \quad (15)$$

We used these parameters to predict the variance and bias of the network for other combinations of gain and context variance (see Fig. 8C–F).

Results

How past actions shape future behavior

Experience-dependent changes in reach planning were evaluated using the well studied paradigm of center-out reaching to visual targets that were arrayed radially about a fixed starting point (Fig. 1A). Across blocks of trials, we manipulated the statistics of recent movements by varying the probability distributions of the target angles for the majority trials within a block, i.e., for the context trials (90 per block in experiment 1, 76 in experiment 2). If recent experience shapes the sensorimotor map, then we expected that changing the context target distribution should affect both the precision and accuracy reach planning. We evaluated such changes using a fixed set of probe targets that were randomly interleaved with the context trials (20 per block in experiment 1, 14 in experiment 2). Our principle measures of reach planning are the mean of the initial reach direction (θ_{MV} , for movement vector) (Fig. 1A) and its standard deviation across trials (σ_{MV}). These measures are made 100 ms after movement onset, before feedback can affect the action (Desmurget and Grafton, 2000).

In the first experiment, we measured how the variance of initial reach directions to a single probe target (Fig. 1A, gray circle) depended on the spread of context targets about the probe location. Eight different target distributions were used, ranging from the repeated target condition, with a single target location for all probe and context reaches (Fig. 1B), to the uniform condition, with context trial targets selected uniformly about the circle (Fig. 1D). The remaining conditions had normally distributed target angles, with the mean at the probe target and with different standard deviations (Fig. 1C). We found that the variance of reaches to the probe target changed across conditions (repeated measures, $F_{(7,49)} = 4.33, p < 0.001$), with less variable contexts generally leading to less variable probe reaches (Fig. 2A). These data show that the repetition of similar reaching movements improves performance on those actions, i.e., “practice makes perfect” on a short timescale.

In a second experiment, we measured how the distribution of context targets affects the mean movement error (bias) at an array of probe targets (Fig. 1A, white circles). When context reaches are all made to a single target, in this case the center probe target (Fig. 1A, gray circle), movements to the other probe targets are strongly biased inwards toward the center position compared with trial blocks with uniformly distributed context targets (Fig. 2B). This bias is stronger for probe locations further from the center target position; this effect attenuates as the distribution of context targets becomes more variable (repeated measures context $\times$ target interaction: $F_{(9,63)} = 5.20, p < 0.001$). These results show that the reduced movement variability for repeated target directions comes at the cost of increased movement bias for other target directions.

Experience-dependent changes, not cognitive strategy

We will argue below that the experience-dependent changes in reaching shown in Figure 2 are the result of an automatic learning process. However, it is also possible that these effects reflect a high-level strategy, where subjects simply predict future target

positions given recent history. To distinguish these possibilities, we conducted a third experiment in which future targets are predictable from recent reaches yet different from them. Within each block of 120 trials, targets were presented sequentially around the circle in either the clockwise (CW) or counterclockwise (CCW) direction, stepping in 3° increments. Subjects were aware of this pattern and could have predicted the location of the next target, resulting in little or no direction-dependent errors. In contrast, if future movements are always biased toward recently presented targets, then subjects should demonstrate a CW bias during CCW trial blocks and vice versa. Indeed, significant direction-dependent errors were observed ($t_{(7)} = 3.28$, $p < 0.025$) (Fig. 2C), confirming that these experience-dependent effects are not the result of predictive, cognitive strategies.

Bayesian interpretation

The tradeoff we observed between variance and bias can be qualitatively understood within a Bayesian framework for reach planning. In this framework, sensory signals, x , give rise to a likelihood function for the current target position, $L(\theta, x)$. Following Bayes' rule, this likelihood is combined with a prior expectation of the target, taking the form of a probability distribution $p(\theta)$, to yield a posterior distribution $p(\theta|x)$. The peak of $p(\theta|x)$ is the MAP estimate of the target location and is used to select the appropriate motor response (see Materials and Methods, above) (Fig. 3A). If $p(\theta)$ is an adaptive prior, then the bias and variance of reaching will reflect recent movement statistics. For example, when repeated movements are made to the center probe target, there is an increasing expectation that future movements will also be made in that direction; i.e., the prior probability distribution tightens about the center target. A tighter prior decreases the variance of the MAP estimate, but also biases it toward the center target. The width of the prior distribution modulates the degree of these effects (Fig. 3B,C).

The variance and bias effects of the simple Bayesian model follow the same trends as those observed experimentally; however, the detailed shapes of these curves are not the same (compare Fig. 3B,C with Fig. 2A,B). More sophisticated Bayesian estimation models can capture some of these details. For example, a robust Bayesian model (Körding and Wolpert, 2004; Knill, 2007) would predict that bias scales less than linearly with target distance, as seen in our data (Fig. 2B). However, we show later that the details of Figure 2 can be largely explained as the result of an iterative learning process acting on the specific order of trials used in our experiments.

Trial-by-trial learning rates

If the variance–bias tradeoff in Figure 2 arises from a process of adaptive Bayesian estimation, then we should be able to see the effects of learning evolve over time. Figure 4 illustrates how the average reach bias evolves for the $\pm 90^\circ$ probe targets over the course of an experimental session in experiment 2. We did not

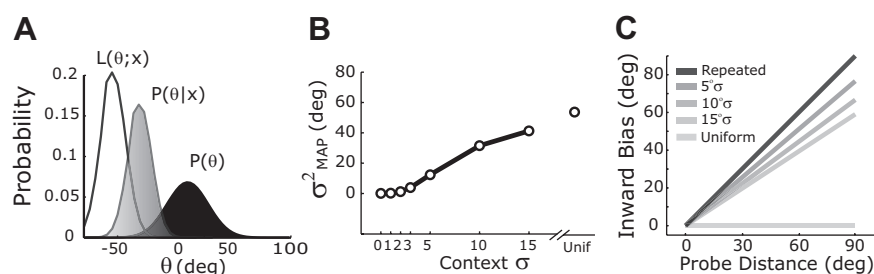


Figure 3. Variance–bias tradeoff in the normative Bayesian model. **A**, The components of the normative Bayesian model. **B**, Change in MAP estimator variance as a function of context variance. **C**, Output bias of the MAP estimator as a function of input location and context variance.

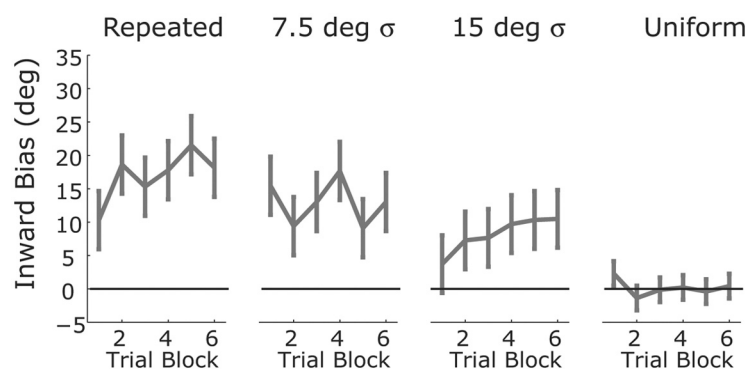


Figure 4. Evolution of reach bias across trial blocks. Data show the mean ($\pm$ SE) across subjects of the bias at the $\pm 90^\circ$ probe targets for each trial block within a session, separately for each context.

observe a large increase in bias across the session, although there is a trend for a slight increase within the first trial block, at least for the repeated target and 15° contexts. While these data might seem inconsistent with an adaptive Bayesian model, it is important to note that before the first $\pm 90^\circ$ probe target of a trial block occurred, a minimum of 10 context trials and an average of >22 context trials had already taken place. If learning occurs on relatively fast timescale, then it would have already approached the asymptotic value for that context by the time of the first probe trials.

To measure the effective learning rates in our experiments, we used a simple incremental learning algorithm to model an adaptive Bayesian estimator. After each trial, the algorithm updates its estimate of the prior distribution, specifically the mean $\bar{\theta}$ and the variance σ^2_{Prior} , based on that trial's target location (Eq. 6) (see Materials and Methods, above). The model includes two free parameters, the likelihood variance $\sigma^2_{\text{Likelihood}}$ and a learning rate β that determines the weight given to the last trial in update algorithm. In the limit of $\beta = 0$, the system is not adaptive and the initial conditions for $\bar{\theta}$ and σ^2_{Prior} are used for all trials. In the limit of $\beta = 1$, the estimates have no memory, e.g., the mean of the prior is set to the target on the previous trial. We fit this model separately to the data from each subject and session in experiment 2 (excluding sessions with the uniform context), minimizing the sum-squared prediction error for movement angle. The best-fit learning rates show a large degree of heterogeneity across subjects, with a median value of 0.25 and a positive skew (SD = 0.27, mean = 0.29). Still, learning was generally fast. For example, with the median learning rate, the estimated mean of the prior would reach 66% of its asymptotic value within four trials. The presence of such fast learning rates explains why we see little change in the measured bias across a session (Fig. 4).

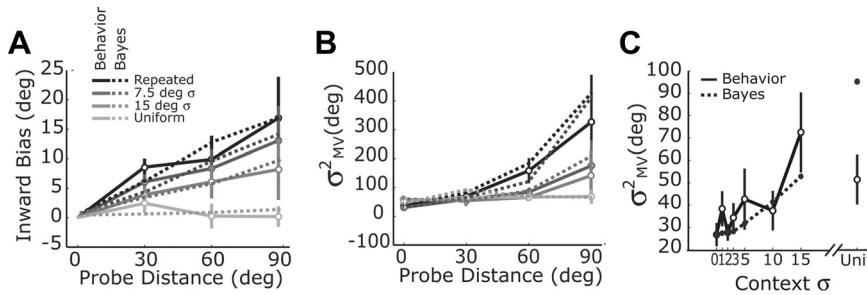


Figure 5. Comparison of adaptive Bayesian model and experimental results. **A, B**, Reach bias (**A**) and reach variance (**B**) as a function of probe target angle and context target distribution for experiment 2. Dashed lines show mean results from the iterative Bayesian model fit to each experimental session in experiment 2. Solid lines show the observed data (mean $\pm$ SE). **C**, Reach variance for experiment 1. Dashed lines show predictions from the adaptive Bayesian model with the mean per session parameters used in **A** and **B**. Behavioral data in **A** and **C** are replotted from Figure 2.

The effect of trial-by-trial learning on movement bias and variance

The adaptive Bayesian model provides a much better account of the mean bias data than the simple normative model, predicting a lower magnitude for the bias and better capturing the dependence of the bias on probe distance (compare Fig. 5A to Fig. 3C). This difference is not due to the difference in fit likelihoods for the two models (mean $\sigma = 7.2^\circ$ for the normative model and $\sigma = 10.0^\circ$ for the adaptive model), since the larger mean variance used in the adaptive model would only increase the magnitude of the bias. Rather, the improvement in fit is due to the effects of trial-by-trial learning and the actual sequence of targets experienced by the subjects. In particular, the presence of the probe trials prevents the prior distribution from converging to the context target distribution.

The influence of trial-by-trial learning can also be seen in the variability of movements. A notable feature of the experimental data is that the reach variability in experiment 2 depends significantly on the interaction between the context and the target location (repeated measures context $\times$ target, $F_{(9,63)} = 2.26$, $p = 0.029$) (Fig. 5B). In the context of the normative Bayesian model, this result is unexpected, since the model predicts that reach variance should decrease monotonically with context variance, independent of target location (Eq. 4; Fig. 3B). We observed this predicted trend at the repeated target location for both experiments 1 (Fig. 2A) and 2. However, the opposite trend was seen for probe targets further from the repeated target location, i.e., an increase in reach variability was observed as the context distribution narrows (Fig. 5B, solid lines). The increased variability is not simply due a large change in bias during the early phase of learning (Fig. 4), as the patterns are qualitatively unchanged when the data are analyzed separately for each of the six trial blocks in the session (data not shown).

These results can be understood as the result of the trial-by-trial learning process. Under the normative Bayesian model (Eq. 3), the MAP estimate of target location is given by

$$\theta_{\text{MAP}} = \bar{\theta} - \frac{\sigma_{\text{Likelihood}}^{-2}}{\sigma_{\text{Prior}}^{-2} + \sigma_{\text{Likelihood}}^{-2}} (\bar{\theta} - x), \quad (16)$$

where $\bar{\theta}$ is the mean of the prior and x is the target estimate based on the sensory input alone. Learning-dependent fluctuations in $\bar{\theta}$ and $\sigma_{\text{Prior}}^{-2}$ would give rise to additional trial-by-trial variability in θ_{MAP} beyond that predicted by the normative Bayesian model. If learning occurs on a short timescale, then such fluctuations are expected even after learning nominally asymptotes for a given

context. The second term in Equation 16 is linear in the target distance $\bar{\theta} - x$, and so one can show that the additional output variance due to fluctuations in $\sigma_{\text{Prior}}^{-2}$ should increase quadratically with the target distance. Specifically, under the normative Bayesian model, the additional output variance due to trial-by-trial fluctuations in the prior is given by

$$\text{Var}(\theta_{\text{MAP}}) = k_1 \text{Var}(\bar{\theta}) + k_2 \text{Var}(\sigma_{\text{Prior}}^{-2}) (\bar{\theta} - \theta)^2, \quad (17)$$

where θ reflects the true value of the target (i.e., the mean value of x) and the parameters k_1 and k_2 only depend on the current values of the likelihood and prior variances. As predicted, the exper-

imental data show an approximately quadratic increase in movement variance with target location (Fig. 5B, solid lines). Furthermore, the adaptive Bayesian model reproduces the variance effects seen in our data from experiment 2 (Fig. 5B, dashed lines).

The adaptive Bayesian model also provides a much better account of movement variance at the center target than the simple normative model (compare Fig. 5C with Fig. 3B). While the dependence on context is comparable for the two models, only the adaptive Bayesian model provides an accurate estimate of the lower bound on movement variance. Specifically, the simple normative model predicts no residual variability in the repeated target case, while the adaptive Bayesian model accurately predicts this value. Again, this difference is due to the effects of trial-by-trial learning. While the nominal target is at the same location on every trial, sensory noise injects trial-by-trial variability into the parameters of the prior distribution and prevents the variance of the prior from converging to zero.

Together, these results suggest that the context-dependent changes in reach variance and bias are indeed the effect of a trial-by-trial learning process. These findings also illustrate that the process of learning can itself be responsible for a large portion of the trial-by-trial variability observed in sensorimotor tasks (Cheng and Sabes, 2006, 2007), even in a case such as this where the apparent goal of learning is the reduction of movement variability.

Adaptive network model

While the adaptive Bayesian model captures many of the key features of the experimental data, it does not address the underlying mechanism for learning. In particular, we are interested in discovering candidate neural mechanisms that can link normative models of learning to the neural circuits that control behavior. Here we propose a parsimonious approach to adaptive Bayesian estimation within cortical sensorimotor circuits, similar to that proposed in previous studies (Wu et al., 2002, 2003; Wu and Amari, 2005).

Consider a recurrently connected network of neurons (Fig. 6A) with dynamics that allow it to efficiently extract information from noisy inputs (Pouget et al., 1998; Deneve et al., 1999, 2001; Latham et al., 2003). On each simulated trial, neurons receive input activation that is determined by the current target angle, the neuronal tuning curves (mean activation vs target angle), and correlated noise (Wu et al., 2002), features that are consistent

with physiological observations of sensorimotor cortex (Burnod et al., 1999; Georgopoulos et al., 1986). The pattern of activity across the network is driven by both the input activation and recurrent activation between neurons with similar tuning curves. At the end of each trial (five iterations of the network dynamics), the planned movement direction is read out from the pattern of activity across the network using a population vector decoder (Georgopoulos et al., 1986, 1988).

In order for the network to learn from experience, a normalized Hebbian learning rule (Hebb, 1949; Oja, 1982) is applied to the recurrent connections so that the changes in connectivity between any two units reflects the trial-by-trial correlations in their firing rates. On every trial, this learning rule acts to strengthen connections that give rise to the pattern of activity associated with the current movement direction, slightly biasing the dynamics of the network toward that pattern. We expected that repeated presentations of a narrow range of targets would strengthen the associated patterns of activity, thereby creating an effective prior on subsequent trials.

We first tested this idea by examining whether the variance and bias of the network change with the statistics of recent experience in a manner similar to that observed experimentally. For each simulated trial block, the network was initialized with a weight structure that has been shown to yield nearly optimal outputs for the case of a flat prior, i.e., that approximates maximum likelihood estimation (Pouget et al., 1998; Latham et al., 2003), and with network parameters fit to the experimental data (see Materials and Methods, above). The network was then trained on a set of context trials with the input target angle drawn from one of the distributions used experimentally (see Blockwise simulations, Materials and Methods, above). After training, learning was turned off and a series of simulated probe trials was performed to measure the trial-by-trial variance and bias of the network output.

As expected, the distribution of training targets affects both the variance and bias of the network output. After training with repeated inputs to the same target, i.e., 0° , the network's output variability is greatly reduced compared with the case of training with uniform targets, while intermediate training distributions lead to intermediate output variability (Fig. 6B). Training with repeated inputs also resulted in a marked bias toward the repeated target (Fig. 6C). These effects become smaller with wider training target distributions.

Simulated experiments with the adaptive network model

While the block-wise network simulations in Figure 6 show the same variance and bias trends as the behavioral data, they do not account for the finer details of the data. This is perhaps to be expected. In the real experiment, learning is never turned off, so the probe trials themselves contribute to the learned effects. Also, the influences of learning can carry over between blocks, either because learning is slow compared with the timescale of a block or because of the presence of multiple timescales of sensorimotor learning (Avillac et al., 2005; Smith et al., 2006; Körding et al., 2007). This makes the results dependent on the specific ordering of trials and conditions. This effect was particularly important for the variance experiment, because the distribution of context targets differed across blocks within the same experimental session.

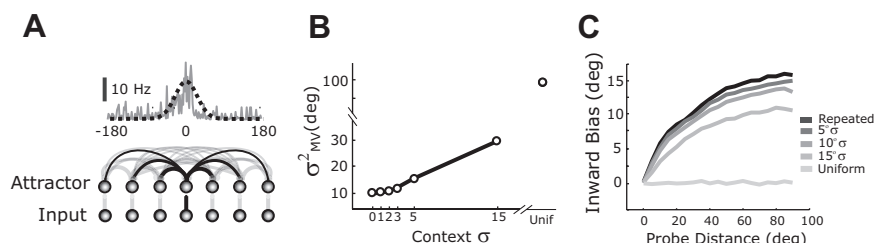


Figure 6. Variance–bias tradeoff in the adaptive network model. **A**, The network model. Top, Mean input activation (black dashed line) and a single example input (gray lines) for a target at $\theta = 0^\circ$. Bottom, Recurrent connections reflect network topography. **B**, Change in network output variance as a function of context variance. **C**, Bias in network output as a function of input location and context variance.

Thus, information learned in one block could carry over and influence performance in future blocks. We therefore conducted network simulations with the exact sequence of trials that subjects experienced in each session. Because subjects were allowed to rest between trial blocks and learned information may have been lost during this delay (Körding et al., 2007), we also allowed for some unlearning between simulated blocks, with the network weights partially decaying back to initial values.

With these full-experiment simulations, we found that the network was able to reproduce the human psychophysical data with good accuracy. The model matches much of the apparent noise in variance-by-context effects in experiment 1 (Fig. 7A), suggesting that these features can be explained by the specific sequence of trials and blocks used in our experiment. For example, the relatively low variance in the uniform context condition appears to be due to the fact that this condition was often presented in one of the last two blocks of the session, when the cumulative effects of learning were greatest. The model also naturally captures the pattern of context- and target-dependent variability that we observed in experiment 2 (Fig. 7C). Finally, even though the network parameters were fit on the reach variance data, the network model is able to predict the pattern of reach biases that were observed in both experiments 2 (Fig. 7B) and 3 (Fig. 7D).

The good match between the data and the adaptive network model for the variance in experiment 2 (Fig. 7C) suggests that the effective learning rate of the network model is also in the same range as that observed experimentally. We explored this issue further by looking at how the within-session changes in network bias compare with the changes in bias observed in experiment 2 (Fig. 7E). The network provides a good match to the data for wider context distributions. However, as the context variance becomes smaller, there are more pronounced learning-dependent increases in bias across trial block for the adaptive network model than for the experimental data, particularly during the repeated target condition. This difference is offset by a larger bias in the initial trial block for the experimental data, suggesting that the effective learning rate in the model is slower than that observed experimentally. As noted in the Discussion below, we hypothesize that this difference results from the fact that the network does not include a sufficiently strong stabilizing force that would cause learning to asymptote.

Adaptive Bayesian priors emerge from Hebbian learning

We have shown that the network model can accurately emulate the experience-dependent changes in reach variance and reach bias that we observed behaviorally. In the network, these changes arise from modifications to the recurrent connections, represented by the matrix of connection weights, with element (i, j) representing the connection strength from unit j to unit i (Fig.

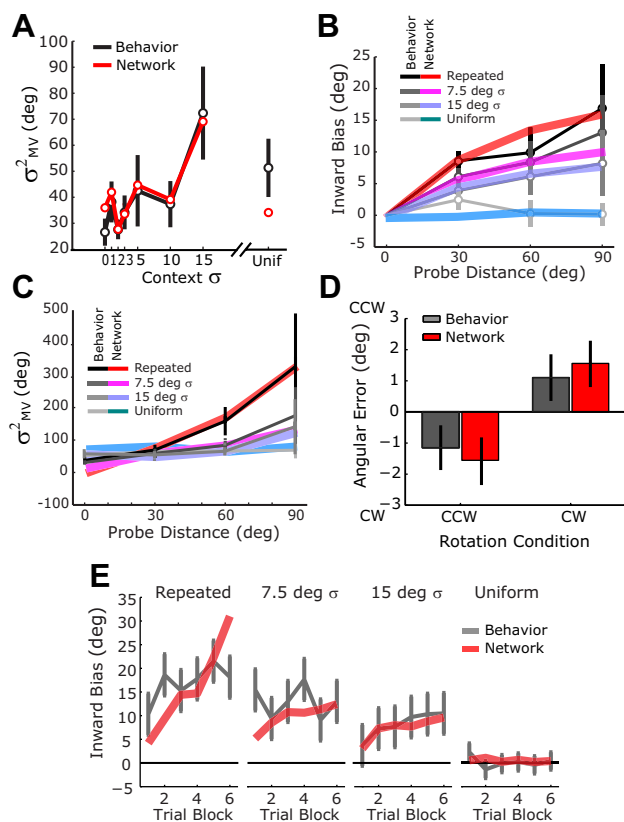


Figure 7. Comparison of network behavior (colored lines) and experimental results (gray lines) when the network is simulated with the same trials sequences that subjects experienced in experiments 1–3. **A**, Reach variance in experiment 1. **B**, Reach bias in experiment 2. **C**, Reach variance in experiment 2. **D**, Reach bias in experiment 3. **E**, Evolution of reach bias across trial blocks in experiment 2. Behavioral data are replotted from Figures 2 and 4. The same network parameters were used in all plots and were fitted to the data in **A** and **C**.

8A). In the repeated-target condition, the activity pattern representing the repeated 0° target is reinforced by Hebbian learning on every trial, causing an increase in the connection strengths between units with preferred directions near this target (Fig. 8B). Similar, but attenuated, changes are apparent when the network is trained with broader target distributions centered at 0° . The enhancement in recurrent connectivity around the 0° unit effectively shifts the energy landscape of the network, deepening the basin of attraction in that region.

These learning-related changes in recurrent connections, and the resulting changes in network dynamics, alter the variance and bias of the network output. We characterized these effects with a range of input gains following training with different target distributions and then compared the network performance to the predictions of a matched normative Bayesian model. The matched value of $\sigma^2_{\text{Likelihood}}$ for each input gain was determined from the output variance of the network trained on the uniform context condition (Fig. 8C, gray curve). The matched values of σ^2_{Prior} were then determined from the network outputs in the 60 Hz input condition (Fig. 8C, vertical gray bar). Details of these computations are given in the Materials and Methods, above.

Figure 8C shows the network output variance as a function of training context and testing input gain. At lower input gains, the network and Bayesian models diverge sharply. This difference arises from the fact that in the normative Bayesian model, the prior distribution is known, and so there is no noise in estimating its mean (see Materials and Methods, above). Thus, when the

input gain is low, $\sigma^2_{\text{Likelihood}}$ is high and the MAP estimate is dominated by the noise-free prior. In contrast, the network never achieves a noise-free prior because there is persistent and stochastic baseline input (see Materials and Methods). As a result, the network exhibits monotonically increasing output variance as the input gain is reduced. If, however, we reduce the level of baseline noise (from 5 Hz in the network simulation to 1 Hz), then the network model provides a very close approximation to the matched normative Bayesian model (Fig. 8D). Furthermore, the same Bayesian model provides a close match to the network bias across training context and input gain (Fig. 8E), even though the parameters of the Bayesian model were determined using only the network variance. These simulations show that the adaptive network model can provide a close approximation to an ideal Bayesian estimator.

Discussion

We have shown that visually guided reaching exhibits a novel form of experience-dependent learning where the statistics of recent movements affect future actions. This learning produces a variance–bias tradeoff that is qualitatively consistent with the process of Bayesian estimation. In particular, we show that repeated performance of movements with similar goal parameters (i.e., target locations) improves the precision of subsequent actions with these goals, resulting in a short-term version of “practice makes perfect.” However, this advantage comes at the cost of reduced accuracy for movements with dissimilar goals.

Others have reported evidence for Bayesian processes in sensorimotor control, primarily reflected as changes in movement bias (Körding and Wolpert, 2004; Miyazaki et al., 2005; Körding et al., 2007) or learning rates (Huang and Shadmehr, 2009; Wei and Körding, 2010). Here we show that such Bayesian affects can be observed even in the absence of the feedback perturbations used in previous studies (Körding and Wolpert, 2004; Wei and Körding, 2010). We show that changes in both movement bias and variance are consistent with the presence of an adaptive prior. In particular, an iteratively adaptive Bayesian model can capture many of the key features of our experimental data.

One limitation of the experimental design used here is that we cannot determine the stage at which these bias and variance effects occur, e.g., target selection, movement vector estimation, or both (Sober and Sabes, 2003). A recent report by Diedrehsen et al. (2010) shows a potentially related form of use-dependent learning (see also Huang et al., 2011). In one of their experiments, passive movements were made to one side or another of an elongated target. This resulted in a bias toward the same side of the target during subsequent voluntary movements. Because their training movements were passive and the target was not precisely defined, it seems likely that the effect is movement-dependent, not target-dependent. To the extent that the same mechanism is at play in our study, our effect is also likely to be at least partly movement-dependent.

With the adaptive Bayesian model, we have argued that these effects are the result of a trial-by-trial learning process. With this model, we were able to estimate the rate of learning and illustrate features of the bias and variance effects that likely result from the presence of trial-by-trial learning. We then showed that Hebbian learning in a simple network model makes a plausible candidate for the mechanism underlying this adaptive estimation process, providing a quantitative match to the experience-dependent changes in movement variance and bias that we observed experimentally (Fig. 7A–D). It is important to note that the adaptive Bayesian model and the network model serve very different pur-

poses in this paper. Therefore, they were fit differently to the data and have different numbers of free parameters. Thus, the goodness-of-fit of the two models cannot be directly compared and we do not claim that one provides a better account of the data than the other. An important point of comparison, however, is the effective learning rates of the two models. The learning rates estimated from adaptive Bayesian model are fast, consistent with the evolution of the reach bias across blocks in experiment 2. In contrast, the best-fit network model appears to have a slower effective learning rate (Fig. 7E). This slower rate is most likely due to the fact that the model does not include a sufficiently strong stabilizing force that causes learning to asymptote in the way empirical learning does (Fig. 4). This may simply result from the choice of learning rule: Oja's rule guarantees stability for the square norm of the weights converging onto a cell, but only in the asymptotic limit. A detailed study of weight normalization in this model and its effects on trial-by-trial learning are left to future work.

Finally, we have shown that the network model can implement a close approximation to a normative Bayesian estimator. Other network models have been proposed for how Bayesian priors can be incorporated into cortical networks (Pouget et al., 2003; Wu et al., 2003; Deneve and Pouget, 2004; Wu and Amari, 2005; Ma et al., 2006). In many of these models, priors are explicitly represented by a separate set of units whose inputs act as an additional source of information. The adaptive network approach offers two advantages. First, the prior naturally emerges within the network connections, removing the need for another set of units to encode this information. Second, Hebbian learning provides a simple mechanism by which these priors can be learned. A similar approach has been previously explored by Wu and colleagues (Wu et al., 2003; Wu and Amari, 2005). They showed that Hebbian learning in a very similar network model acts to smooth stochastic input signals over time, reducing output variance when the input is stationary at the cost of slow reaction to changing inputs (i.e., increased bias). This smoothing approximates a form of Bayesian estimation, i.e., a simple Kalman filter. We have extended these results, showing that a recurrent network with Hebbian learning approximates Bayesian estimation across a range of input statistics and on behaviorally relevant timescales. Furthermore, we showed that the network output provides a good quantitative match to psychophysical data. Overall, these results suggest that prior experience may be incorporated into sensorimotor planning via Hebbian learning and as a natural byproduct of the ongoing dynamics of the underlying neural circuits.

NOTE

Supplemental material for this article is available at <http://arxiv.org/abs/1106.2977>. The supplement presents an analysis of the effects of recent inputs on the steady-state dynamics of the network model, providing additional insight into why the network approximates Bayesian estimation. This material has not been peer reviewed.

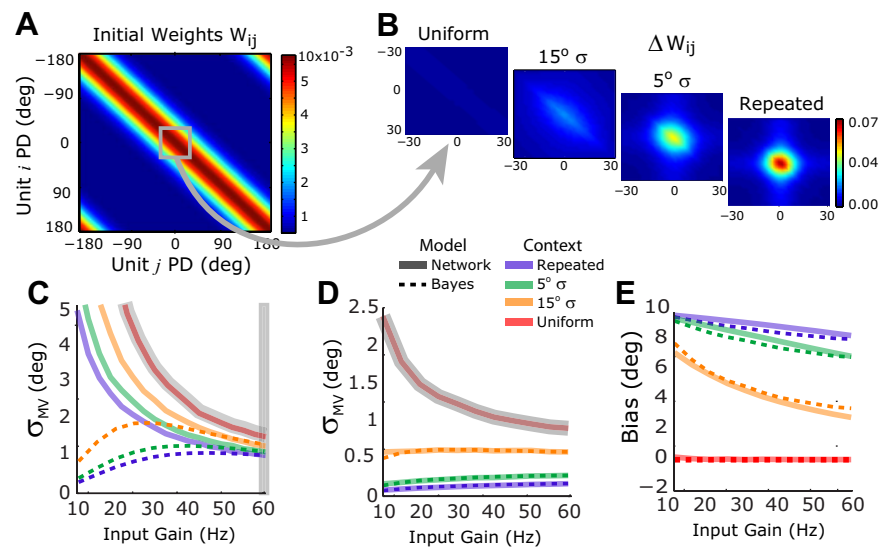


Figure 8. Adaptive network model approximates Bayesian estimation. **A**, Matrix of initial recurrent weights before training, with units arranged topographically by preferred direction (PD). 0° represents the repeated target angle. **B**, Changes in the central portion of the weight matrix following 100 training trials with different context target distributions. **C–E**, Comparison of network output and a matched normative Bayesian model as a function of context variance during training and input gain during testing. **C**, Solid lines, Network output variance. Dashed lines, Predictions from a matched Bayesian model with likelihoods determined after training with the uniform context variance (gray curve) and priors determined by the network results with 60 Hz input gain (gray vertical bar); the Bayesian model necessarily matches the network model for those data points. **D**, Same plot as in **C** but using a network with network baseline firing rates lowered to 1 Hz. Priors for this network were also estimated with the 60 Hz input gain (not shown). **E**, Network output bias (solid lines) for $\theta = 10^\circ$ compared with the prediction of the matched Bayesian model from **D** (dashed lines).

References

- Avillac M, Deneve S, Olivier E, Pouget A, Duhamel JR (2005) Reference frames for representing visual and tactile locations in parietal cortex. *Nat Neurosci* 8:941–949.
- Burnod Y, Baraduc P, Battaglia-Mayer A, Guigon E, Koechlin E, Ferraina S, Lacquaniti F, Caminiti R (1999) Parieto-frontal coding of reaching: an integrated framework. *Exp Brain Res* 129:325–346.
- Cheng S, Sabes PN (2006) Modeling sensorimotor learning with linear dynamical systems. *Neural Comput* 18:760–793.
- Cheng S, Sabes PN (2007) Calibration of visually guided reaching is driven by error-corrective learning and internal dynamics. *J Neurophysiol* 97:3057–3069.
- Deneve S, Pouget A (2004) Bayesian multisensory integration and cross-modal spatial links. *J Physiol Paris* 98:249–258.
- Deneve S, Latham PE, Pouget A (1999) Reading population codes: a neural implementation of ideal observers. *Nat Neurosci* 2:740–745.
- Deneve S, Latham PE, Pouget A (2001) Efficient computation and cue integration with noisy population codes. *Nat Neurosci* 4:826–831.
- Deneve S, Duhamel JR, Pouget A (2007) Optimal sensorimotor integration in recurrent cortical networks: a neural implementation of Kalman filters. *J Neurosci* 27:5744–5756.
- Desmurget M, Grafton S (2000) Forward modeling allows feedback control for fast reaching movements. *Trends Cogn Sci* 4:423–431.
- Diedrichsen J, White O, Newman D, Lally N (2010) Use-dependent and error-based learning of motor behaviors. *J Neurosci* 30:5159–5166.
- Georgopoulos AP, Schwartz AB, Kettner RE (1986) Neuronal population coding of movement direction. *Science* 233:1416–1419.
- Georgopoulos AP, Kettner RE, Schwartz AB (1988) Primate motor cortex and free arm movements to visual targets in three-dimensional space. II. Coding of the direction of movement by a neuronal population. *J Neurosci* 8:2928–2937.
- Grubbs FE (1950) Sample criteria for testing outlying observations. *Ann Math Stat* 21:27–58.
- Hebb DO (1949) *The organization of behavior: a neurophysiological theory*. New York: Wiley.
- Huang VS, Shadmehr R (2009) Persistence of motor memories reflects statistics of the learning event. *J Neurophysiol* 102:931–940.
- Huang VS, Haith A, Mazzoni P, Krakauer J (2011) Rethinking motor learn-

- ing and savings in adaptation paradigms: model-free memory for successful actions combines with internal models. *Neuron* 70:787–801.
- Kersten D, Mamassian P, Yuille A (2004) Object perception as Bayesian inference. *Annu Rev Psychol* 55:271–304.
- Kluzik J, Diedrichsen J, Shadmehr R, Bastian AJ (2008) Reach adaptation: what determines whether we learn an internal model of the tool or adapt the model of our arm? *J Neurophysiol* 100:1455–1464.
- Knill DC (2007) Robust cue integration: a Bayesian model and evidence from cue-conflict studies with stereoscopic and figure cues to slant. *J Vis* 7:5.1–5.24.
- Körding KP, Wolpert DM (2004) Bayesian integration in sensorimotor learning. *Nature* 427:244–247.
- Körding KP, Tenenbaum JB, Shadmehr R (2007) The dynamics of memory as a consequence of optimal adaptation to a changing body. *Nat Neurosci* 10:779–786.
- Lages M, Heron S (2008) Motion and disparity processing informs Bayesian 3D motion estimation. *Proc Natl Acad Sci U S A* 105:E117.
- Latham PE, Deneve S, Pouget A (2003) Optimal computation with attractor networks. *J Physiol Paris* 97:683–694.
- Lu H, Yuille AL, Liljeholm M, Cheng PW, Holyoak KJ (2008) Bayesian generic priors for causal learning. *Psychol Rev* 115:955–984.
- Ma WJ, Beck JM, Latham PE, Pouget A (2006) Bayesian inference with probabilistic population codes. *Nat Neurosci* 9:1432–1438.
- Miyazaki M, Nozaki D, Nakajima Y (2005) Testing Bayesian models of human coincidence timing. *J Neurophysiol* 94:395–399.
- Oja E (1982) A simplified neuron model as a principal component analyzer. *J Math Biol* 15:267–273.
- Pouget A, Zhang K, Deneve S, Latham PE (1998) Statistically efficient estimation using population coding. *Neural Comput* 10:373–401.
- Pouget A, Deneve S, Duhamel JR (2002) A computational perspective on the neural basis of multisensory spatial representations. *Nat Rev Neurosci* 3:741–747.
- Pouget A, Dayan P, Zemel RS (2003) Inference and computation with population codes. *Annu Rev Neurosci* 26:381–410.
- Sato Y, Toyoizumi T, Aihara K (2007) Bayesian inference explains perception of unity and ventriloquism aftereffect: identification of common sources of audiovisual stimuli. *Neural Comput* 19:3335–3355.
- Scheidt RA, Dingwell JB, Mussa-Ivaldi FA (2001) Learning to move amid uncertainty. *J Neurophysiol* 86:971–985.
- Shadmehr R, Mussa-Ivaldi FA (1994) Adaptive representation of dynamics during learning of a motor task. *J Neurosci* 14:3208–3224.
- Slijper H, Richter J, Over E, Smeets J, Frens M (2009) Statistics predict kinematics of hand movements during everyday activity. *J Mot Behav* 41:3–9.
- Smith MA, Ghazizadeh A, Shadmehr R (2006) Interacting adaptive processes with different timescales underlie short-term motor learning. *PLoS Biol* 4:e179.
- Sober SJ, Sabes PN (2003) Multisensory integration during motor planning. *J Neurosci* 23:6982–6992.
- Sober SJ, Sabes PN (2005) Flexible strategies for sensory integration during motor planning. *Nat Neurosci* 8:490–497.
- Thoroughman KA, Shadmehr R (2000) Learning of action through adaptive combination of motor primitives. *Nature* 407:742–747.
- Wei K, Körding K (2010) Uncertainty of feedback and state estimation determines the speed of motor adaptation. *Front Comput Neurosci* 4:11.
- Weiss Y, Simoncelli EP, Adelson EH (2002) Motion illusions as optimal percepts. *Nat Neurosci* 5:598–604.
- Wu S, Amari S (2005) Computing with continuous attractors: stability and online aspects. *Neural Comput* 17:2215–2239.
- Wu S, Amari S, Nakahara H (2002) Population coding and decoding in a neural field: a computational study. *Neural Comput* 14:999–1026.
- Wu S, Chen D, Niranjana M, Amari S (2003) Sequential Bayesian decoding with a population of neurons. *Neural Comput* 15:993–1012.