

Robust Coding of Eye Position in Posterior Parietal Cortex despite Context-Dependent Tuning

Jamie R. McFadyen,¹ Barbara Heider,² Anushree N. Karkhanis,²  Shaun L. Cloherty,³ Fabian Muñoz,^{4,5} Ralph M. Siegel,² and Adam P. Morris^{1,6}

¹Neuroscience Program, Biomedicine Discovery Institute, Department of Physiology, Monash University, Clayton, VIC, 3800, Australia, ²Center for Molecular and Behavioral Neuroscience, Rutgers University, Newark, NJ, 07102, ³School of Engineering, RMIT University, Melbourne, VIC, 3001, Australia, ⁴Department of Neuroscience, Columbia University, New York, NY, 10027, ⁵Zuckerman Mind Brain Behavior Institute, Columbia University, New York, NY 10027, and ⁶Monash Data Futures Institute, Monash University, Clayton, VIC, 3800, Australia

Neurons in posterior parietal cortex (PPC) encode many aspects of the sensory world (e.g., scene structure), the posture of the body, and plans for action. For a downstream computation, however, only some of these dimensions are relevant; the rest are “nuisance variables” because their influence on neural activity changes with sensory and behavioral context, potentially corrupting the read-out of relevant information. Here we show that a key postural variable for vision (eye position) is represented robustly in male macaque PPC across a range of contexts, although the tuning of single neurons depended strongly on context. Contexts were defined by different stages of a visually guided reaching task, including (1) a visually sparse epoch, (2) a visually rich epoch, (3) a “go” epoch in which the reach was cued, and (4) during the reach itself. Eye position was constant within trials but varied across trials in a 3×3 grid spanning $24^\circ \times 24^\circ$. Using demixed principal component analysis of neural spike-counts, we found that the subspace of the population response encoding eye position is orthogonal to that encoding task context. Accordingly, a context-naïve (fixed-parameter) decoder was nevertheless able to estimate eye position reliably across contexts. Errors were small given the sample size ($\sim 1.78^\circ$) and would likely be even smaller with larger populations. Moreover, they were comparable to that of decoders that were optimized for each context. Our results suggest that population codes in PPC shield encoded signals from crosstalk to support robust sensorimotor transformations across contexts.

Key words: Bayesian; decode; eye position; gaze; orthogonal; PPC

Significance Statement

Neurons in posterior parietal cortex (PPC) which are sensitive to gaze direction are thought to play a key role in spatial perception and behavior (e.g., reaching, navigation), and provide a potential substrate for brain-controlled prosthetics. Many, however, change their tuning under different sensory and behavioral contexts, raising the prospect that they provide unreliable representations of egocentric space. Here, we analyze the structure of encoding dimensions for gaze direction and context in PPC during different stages of a visually guided reaching task. We use demixed dimensionality reduction and decoding techniques to show that the coding of gaze direction in PPC is mostly invariant to context. This suggests that PPC can provide reliable spatial information across sensory and behavioral contexts.

Received Mar. 30, 2021; revised Mar. 30, 2022; accepted Mar. 30, 2022.

Author contributions: J.R.M., B.H., A.N.K., and F.M. performed research; J.R.M., B.H., S.L.C., F.M., and A.P.M. contributed unpublished reagents/analytic tools; J.R.M. analyzed data; J.R.M. wrote the first draft of the paper; J.R.M. and A.P.M. edited the paper; J.R.M. wrote the paper; R.M.S. and A.P.M. designed research.

This work was supported by National Institute of Health Grants EY09223, T32NS081287, and T32NS081287; Whitehall Foundation; Charles and Johanna Busch Faculty Research Grant; Mind Science Foundation; National Partnership for Advanced Computational Infrastructure; and Army W911NF-09-1-0116. This work was also supported by National Health and Medical Research Council of Australia APP1083898 to A.P.M. The contents of the published material are solely the responsibility of the administering institution, a participating institution, or individual authors and do not reflect the views of the National Health and Medical Research Council. We thank Nirmala Ramalingam and Jasmine Siegel for assistance in data collection; and Nicholas Price for feedback on the manuscript. The late Ralph Siegel contributed to conceptual design and data collection, but not to the analyses or text presented in this paper (he passed in September 2011).

The authors declare no competing financial interests.

Correspondence should be addressed to Adam P. Morris at adam.morris@monash.edu.

<https://doi.org/10.1523/JNEUROSCI.0674-21.2022>

Copyright © 2022 the authors

Introduction

To reach for an object, the brain needs to combine visual information with internal postural signals, including the position of the eyes and the position of the hand (Soechting and Flanders, 1992). Thus, like most brain functions, the computations for reaching incorporate many variables that change over short time scales and that must be estimated from internal sources of information. This information is encoded in the rates or timing of spiking activity across a population of neurons (i.e., as a “population code”), and is pooled in downstream computations (Pitkow and Angelaki, 2017).

Spiking patterns, however, are “noisy”; they change across seemingly identical repeats of a stimulus or behavior. Some of

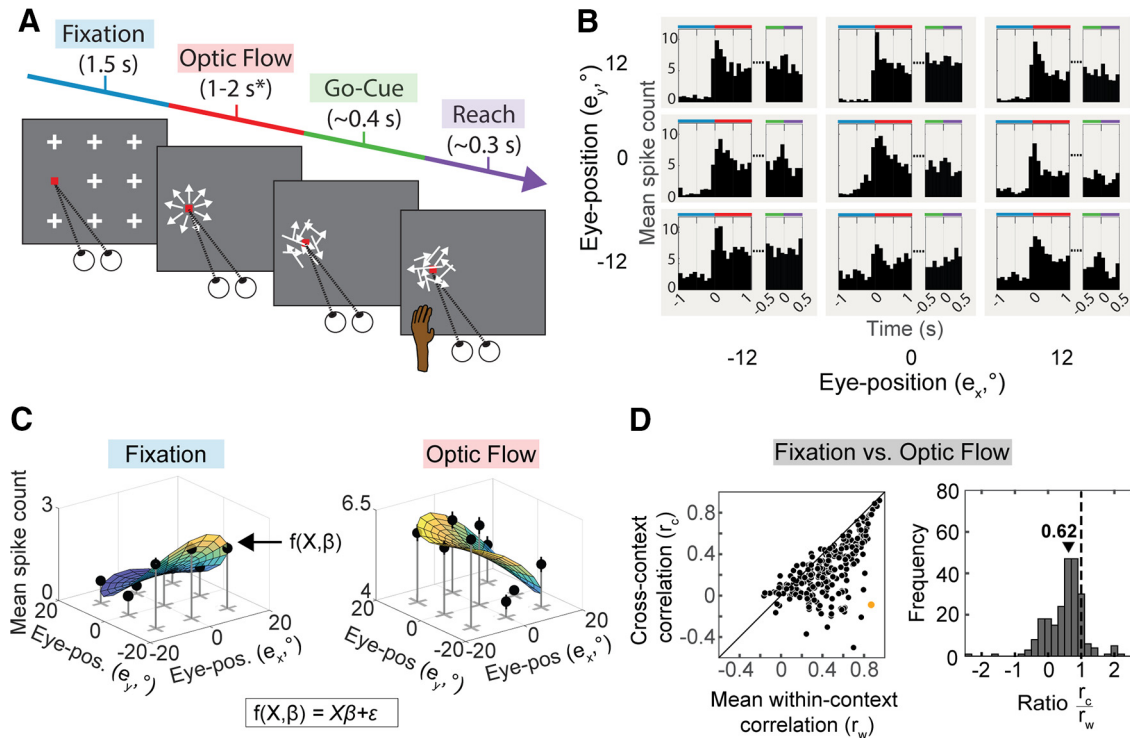


Figure 1. Single neurons in PPC have different tuning for eye position across behavioral and sensory contexts. **A**, Animals performed a visually guided reaching task that we partitioned into four analysis epochs (i.e., contexts), indicated by the colored labels (the context-color mapping shown here is maintained throughout the paper). Animals maintained fixation throughout each trial on a small dot at one of nine positions arranged in a 3×3 grid (shown as plus symbols but not visible on the actual display). Shortly after the onset of fixation, a cloud of moving dots appeared within a circular aperture (diameter = 12°) centered on the fixation point. The dots moved in a radial expansion pattern simulating optic flow. After a delay (1–2 s or 2–3 s, in separate sessions), the dots switched from expansion to random, incoherent motion. This was the “go” cue, signaling the animal to move its hand from a resting position near the torso toward the display and touch the dot cloud. While the animal performed the task, we recorded the activity of neurons in areas 7a or DP with single electrodes. **B**, Peristimulus time histograms (100 ms bins) for an example neuron from area DP (MFR06102), aligned to the onset of the Optic Flow or Reach context. Each subpanel corresponds to one position, arranged to match the grid of eye positions in the task. The neuron showed strong effects of context and eye position on spike-counts. **C**, The data from **B**, plotted as a 2D spatial tuning curve (i.e., EPF), separately for two task contexts. The data points are the mean spike-counts ($\pm$ SEM) for all 100 ms windows within each context and pooled over trials. The fitted regression surfaces (Eq. 3) provide a quantitative link between eye position and neural activity in each context and form the basis of our probabilistic decoder (see Fig. 2). The neuron shows a striking change in spatial tuning across contexts. Although this is an extreme case to illustrate the phenomenon (see **D**), all neurons that had significant spatial tuning changed the shape of their EPFs across contexts (Heider et al., 2010). **D**, The scatterplot (left) shows that eye position tuning was more similar when two independent datasets from the same context were compared (within-context correlation) than when the datasets came from different task contexts (cross-context correlation). The independent datasets were obtained by repeatedly subsampling the available spike-counts into two sets, from either the same or different contexts, and then calculating the means for the nine eye positions and correlating the two sets. Each point in the plot is the mean correlation across 5000 resampled sets for the Optic Flow and Fixation contexts. **B**, **C**, Amber circle represents the example neuron. Right, Ratios of cross- and within-context correlations across the population. Black triangle represents the population median. Dashed line indicates the expected value for no change in tuning across contexts. The values for five neurons (2% of total sample) were outside the plotted range. Figure 3 shows these ratios for all combinations of cross-context pairs.

this variability is stochastic, but it also includes contributions from experimentally uncontrolled internal (e.g., thoughts, intentions, arousal) and external (e.g., image background or lighting) variables that could, in principle, be predicted (Niell and Stryker, 2010; Masquelier, 2013; Pachitariu et al., 2015). Single neurons in area 7a of the posterior parietal cortex (PPC), for example, are sensitive to variables, including retinal stimulus location, eye position, spatial attention, reach trajectory, and vestibular input (Motter and Mountcastle, 1981; Andersen et al., 1985; MacKay, 1992; Quraishi et al., 2007; Avila et al., 2019). Hence, the context in which a downstream computation is performed consists of the joint state of all these variables.

A key question is how computations remain robust across contexts. Consider a population in which each neuron is sensitive to three variables (a , b , and c) and a downstream area that computes a function that incorporates estimates of only two (e.g., $f = \hat{a} + \hat{b}$). Here, changes in activity related to the value of c pose a threat to the computation; it is a nuisance variable and must be factored out. Information must therefore be encoded in a way that provides robust access to $\hat{a} + \hat{b}$

regardless of context. This is not achieved in single neurons; a scalar spike-count cannot separate these variables. Worse, many neurons have “mixed” selectivity, in which their tuning for one variable changes depending on the values of others (Quraishi et al., 2007; Heider et al., 2010; Rigotti et al., 2013; Kim et al., 2016).

We recently reported a striking example of mixed tuning in area 7a and the dorsal prelunate area (DP) of macaque PPC during a visually guided reaching task in which the location of a fixation point was varied across trials (Heider et al., 2010; Karkhanis et al., 2014). Single neurons were strongly tuned to fixation position, consistent with previous reports in PPC (Andersen et al., 1985). This tuning, however, changed within trials as the context changed, reversing spatial selectivity in some cases (see Figs. 1B–D and 3). Many of the contextual nuisance variables were observable (e.g., changes in the visual stimulus and the animal’s arm position), whereas others were unobservable but presumed (e.g., changes in attention and motor planning) (Andersen et al., 1997, 1998; Gottlieb et al., 1998; Snyder et al., 1998; Quraishi et al., 2007).

Context-dependent tuning suggests that the population code for eye position could be corrupted by changes in context. This would be incompatible with the putative role of eye position signals in maintaining stable representations of visual space across eye movements (Zipser and Andersen, 1988; Morris et al., 2013, 2016; Morris and Krekelberg, 2019). Here, however, we reanalyzed data reported previously (Heider et al., 2010; Karkhanis et al., 2014) to show that eye position can be read out from PPC neurons robustly even if the decoder is blind to context. Using demixed dimensionality reduction (Kobak et al., 2016), we show that this invariance arises from the geometry of how population activity changes across contexts and eye positions. Specifically, context shifts population vectors along axes that are orthogonal to the subspace in which eye position is encoded, shielding the read out of information from crosstalk. The result mirrors analogous effects in other systems (Kaufman et al., 2014; Gallego et al., 2017; Perich et al., 2018).

Materials and Methods

This study is a new analysis of data reported in two previous papers (Heider et al., 2010; Karkhanis et al., 2014). Experimental methods are described only briefly. The same animals and behavioral task (referred therein as the EVAR task) were used in both studies; hence, here the recordings were combined into a single dataset. This included 115 neurons from Karkhanis et al. (2014) (80 in area 7a, 32 in DP, and 3 potentially within the boundaries of either 7a or DP) and 120 neurons from Heider et al. (2010) (67 in area 7a and 53 in DP).

Electrophysiology

Experiments were conducted in 2 male macaque monkeys (*Macaca mulatta*) and conformed to the National Institutes of Health's *Guide for the care and use of laboratory animals*. Protocols were approved by the Rutgers University Animal Care and Facilities Committee. Animals had imaging chambers (20 mm diameter) implanted over the PPC of the right hemisphere for previous studies (Siegel et al., 2003; Heider et al., 2005). Electrode penetrations were guided by anatomic landmarks visible through the transparent artificial dura. Area 7a was defined as the region between the intraparietal sulcus and the end of the superior temporal sulcus, while DP identified as the region between the superior temporal sulcus and the lunate sulcus (Heider et al., 2010). Extracellular action potentials of isolated neurons were recorded using single platinum-iridium microelectrodes (UEPSESG2N5G, FHC), inserted orthogonal to the cortical surface (mean depth of recordings = 1140 μ m, SD = 565, number of penetrations/sessions = 155). Only a single penetration was made per recording session. Neural activity was amplified (Model 1800 Microelectrode AC Amplifier; A-M. Systems), filtered using a 50/60 Hz noise eliminator Humbug (Auto-Mate Scientific) and a bandpass filter with a bandwidth ranging from 300 to 20,000 Hz, and digitized at 40 kHz. Spike sorting was performed offline using the Plexon Offline Sorter version 1.39 (Plexon).

Behavior

Animals were trained to perform a visually guided reaching task. The animal was placed in a custom primate chair and faced a touch-sensitive display (Elo TouchSystems), with the left hand resting on a touch-sensitive panel (Crist Instrument) or capacitive proximity sensor (IFM Efcator) near the torso. The chair design allowed free movement of the upper limbs, and the screen was positioned close enough to be touched with an extended arm (30 and 35 cm for each animal, respectively). Animals were head-fixed, allowing only changes in eye position when performing the fixation task. Gaze direction was monitored and recorded at 60 Hz using an infra-red video-based eye tracking system (ISCAN) and the NIMH Cortex system.

A fixation point (a small red square 0.8° in width) appeared on the screen at one of nine locations forming a 3 × 3 grid (Fig. 1A), each spaced 12° apart horizontally and vertically. The animal was required to maintain fixation on the target for 1.5 s, after which a random dot

pattern appeared, centered on the fixation point (12° in diameter). The dots moved in a radial expansion pattern, simulating optic flow, and were intended to evoke visually driven activity from the recorded neuron. The animal was required to maintain fixation for a random duration of 1–2 s (for neurons from Karkhanis et al., 2014) or 2–3 s (for neurons from Heider et al., 2010). At this point, the dot pattern switched from expansion to random motion, in which the direction of each dot was drawn randomly from 0 to 360. This change in motion pattern signaled that it was time to perform a reach toward the display at the target location. The animal was rewarded if it touched the display within the 12° optic flow target diameter of the fixation target point and held it for 1.5 s. Target position was randomly selected within each block, and an average of 9.37 trials (SD = 1.02) were recorded at each of the nine potential target positions.

Data analysis

All analyses were performed using MATLAB R2017a (The MathWorks). Each trial was separated into four contexts: Fixation, Optic Flow, Go Cue, and Reach. Fixation was the interval between –1000 and –100 ms relative to the onset of the random dot pattern. Optic Flow was the interval in which the dots moved in an expansion pattern (1000–3000 ms duration). Go Cue was the interval between the onset of the reach cue (random dot motion) and the initiation of the reach (mean = 438.8 ms, SD = 141.2 ms), and Reach was the interval between the release of the hand from its start position and its first contact with the display (mean = 297.98 ms, SD = 72.1 ms).

Recordings that included fewer than six trials for any of the nine eye positions were excluded. A small number of recordings (4) were also removed because of difficulty with eye tracking calibration. In those cases, the same experimental condition was repeated many times (20+) to collect a complete dataset; hence, the animal's behavior was different from other sessions. Only correctly completed trials were analyzed. Analyses were performed using 235 neurons that satisfied these criteria (58 from Animal MFR and 177 from Animal MJR).

Spikes were counted within nonoverlapping 100 ms bins throughout each trial. All bins that fell within the boundaries of a context were collected across trials to create a pool of spike-counts for each neuron for every combination of eye position and task context (we refer to a conjunction of eye position and context as an “experimental condition”). This resulted in an average number of independent spike-counts per eye position and neuron of 76.74 (SD = 5.64) for the Fixation context, 162.97 (SD = 43.86) for the Optic Flow context, 41.02 (SD = 8.49) for the Go Cue context, and 27.88 (SD = 4.63) for the Reach context.

To allow analysis of how eye position and context affect activity at the population level, we aggregated the spike-counts of the n neurons (which were recorded one at a time) into n -dimensional response vectors (i.e., pseudo-population codes). This was done by selecting a random spike-count, r_i , from each neuron (i) within a given condition and concatenating them into a population response vector, r_{pop} . These population vectors served as the test data for our decoder.

For the demixed dimensionality reduction and geometric analyses, these vectors were averaged within each neuron and condition to yield mean population vectors. That is, we quantified the multivariate (eye position and context) tuning of the population. Although this pseudo-population response approach precludes analysis of trial-by-trial variability, it is assumed to preserve the enduring tuning characteristics of neurons and therefore allows insight into the higher-order geometry of encoded information at the population level. For example, the collection of vectors for one of these variables (e.g., across the nine eye positions) traces out a surface in the population space along which that variable is encoded, which can then be compared with that in which another variable (i.e., context) is encoded. We used this approach, combined with demixed principal component analysis (dPCA), to assess the orthogonality or collinearity of eye position and context effects on population activity.

dPCA. Previously, we have reported context-dependent changes in the eye position tuning of these single 7a and DP neurons (Heider et al., 2010; Karkhanis et al., 2014). We used dPCA to determine the geometric structure of these context effects at the population level (see Results).

This technique produces low-dimensional representations of the population response (“subspaces”) similar to the more familiar principal component analysis (PCA). PCA, however, is an unsupervised technique, blind to task variables; dPCA instead identifies components that capture variance in the population response associated with specific task variables (here, eye position and context). It does so by marginalizing over irrelevant task variables (e.g., context, when identifying eye position subspaces) before the dimensionality reduction step, similar to the marginalization used in factorial ANOVA (for detailed methods, see Kobak et al., 2016). dPCA components can thus be considered to define the low-dimensional subspace encoding a variable of interest (though further information is usually available in other subspaces).

We ran dPCA on the full set of spiking data recorded across eye positions and contexts. As the number of output dPCA components is limited to being one less than the number of unique combinations of task variables (here, four contexts and nine eye positions), this yielded a matrix of weights that projects the 235-dimensional population responses into a 35-dimensional subspace (if retaining all components). The 35 dimensions can be sorted into 8 principal components that define an eye position subspace, 3 that define the context subspace, and the remaining 24 that define interaction dimensions. The technique also quantifies the variance explained by each of these components, allowing a comparison of relative effect size on the overall population activity (Kobak et al., 2016).

While the components identified by dPCA within a task variable subspace (e.g., eye position) are constrained to be orthogonal, as in traditional PCA, those from different subspaces (e.g., the first components from the eye position and context subspaces) are not subject to this constraint. This means that the geometric relationship between different subspaces is a property of the neural data (i.e., the brain), not the technique, and can be assessed analytically.

Of key interest is the extent to which the eye position and context subspaces are orthogonal, which can be quantified using the set of principal angles between their component dimensions. Subspaces that are close to 90° apart can be considered orthogonal and would suggest that an invariant, context-naïve read-out of eye position should be possible, despite changes in the tuning of single neurons across contexts. For collinear subspaces (i.e., principal angles closer to 0°), the effect of one variable (e.g., context) changes the population response in ways that are similar to changes caused by the other variable (e.g., eye position), confounding the two signals and likely leading to errors in downstream read-out (Averbeck and Lee, 2006; Moreno-Bote et al., 2014; Pitkow et al., 2015).

Principal angles between the eye and context subspaces were determined using the *subspace* function for MATLAB (Knyazev and Argentati, 2002; Knyazev, 2021) and then expressed in degrees. Because the eye and context subspaces estimated by dPCA consisted of eight and three dimensions respectively, their geometric relationship was quantified using three principal angles.

As a point of reference, we compared the observed principal angles with null distributions; that is, with the angles expected if subspaces of equivalent dimensionality are positioned randomly within the 35-dimensional population subspace. We generated 1000 pairs of random subspaces by assigning random coefficients (from −1 to 1) to each axis to define an arbitrary vector in 35-dimensional space. The full set of dimensions orthogonal to the randomly generated vector were found using MATLAB’s *null* function, and these dimensions were randomly sampled without replacement until the new random subspace matched the dimensionality of the eye subspace. This process was repeated to generate a second random subspace matching the dimensionality of the context subspace. We then calculated the principal angles between the pair of subspaces, repeating this process 1000 times. The principal angles for the subspaces produced by dPCA were then compared with the 5th percentile of these null distributions to determine whether they departed significantly from that expected by chance.

Probabilistic decoding: universal decoder. If the eye position and context subspaces of the population response are orthogonal, it should be possible to find a decoding algorithm that accurately estimates eye position from spike-counts using a fixed read-out strategy, regardless of

context and associated changes in tuning. Accordingly, we developed a “universal decoder” that maps observed spike-counts to quasi-continuous, 2D estimates of eye position in degrees of visual angle (Fig. 2). Decoder accuracy was calculated according to the Euclidean distance of its estimates to the true eye position.

Our decoder was based on the maximum-likelihood estimation approach used previously in PPC and elsewhere (Morris et al., 2013; Morris and Krekelberg, 2019). This decoder infers eye position from spiking activity via a process of probabilistic inference and is built on a set of generative models (one for each neuron) that capture conditional spike likelihood functions. Given a pseudo-population response vector, the decoder computes a probability map indicating the relative likelihood of different eye positions. The most likely eye position, according to the spike-count models, is then used as the decoder’s estimate. We used leave-one-out cross-validation for all decoding analyses.

The generative model for each neuron represented the likelihood of observing particular responses (e.g., $r = 2$) as a probability mass function over spike-counts (R) at arbitrary positions in 2D eye position space, $p(R|e_x, e_y)$. Previously, we have assumed a Poisson form for these functions (i.e., that spike-counts are random samples from a Poisson distribution), for which the mean and variance are equal. Here, however, our spike-count distributions consisted of samples pooled across the four contexts, and hence had greater (super-Poisson) variability, because the means tended to change across contexts. Indeed, Fano factors (the ratio of the variance to the mean) were >1 on average for the pooled dataset used for the universal decoder (mean = 1.38, SD = 0.53).

To capture this extra variability, we instead assumed a negative binomial form for spike-count distributions (Goris et al., 2014). The negative binomial distribution is defined by two parameters, a and b , which together define distributions for spike-counts in which the mean and variance are dissociable (unlike Poisson) and with variance that is greater than (but not less than) the mean [Definitions of a and b vary across texts and software packages; the equations below (Eqs. 1, 2, 4, and 5) are correct under the definitions and equations used in MATLAB’s *nbpdf* (version R2017a)]. a and b were estimated independently for each neuron at each of the nine measured eye positions, $[e_x, e_y]$, using the observed spike-count distributions and MATLAB’s *mle* function. To make our decoder “universal,” these distributions consisted of spike-count samples pooled across all four contexts.

We then extrapolated these local parameters to a global generative model that could return spike-count probabilities at arbitrary positions in 2D eye position space. To do so, we used a regression approach similar to that classically used to study eye position fields (EPFs, also known as gain-fields) in PPC, in which the mean firing rate is modeled as a second-order polynomial function of eye position (Eq. 3) (Morris et al., 2013; Morris and Krekelberg, 2019). Here, however, our aim was to also capture potential changes in spike-count variance across eye positions. Accordingly, we reparameterized the fitted a and b parameters for each eye position to the corresponding values for the negative binomial mean (μ) and variance (σ^2) using the following:

$$\mu = \frac{a(1-b)}{b} \quad (1)$$

$$\sigma^2 = \frac{a(1-b)}{b^2} \quad (2)$$

and then fitted a regression surface separately to these two moments of the negative binomial distribution across eye positions using the same function form but with independent parameters as follows:

$$v = f(e_x, e_y, \beta) = \beta_1 e_x + \beta_2 e_y + \beta_3 e_x^2 + \beta_4 e_y^2 + \beta_5 e_x e_y + \beta_6 \quad (3)$$

where e_x and e_y are the 2D eye positions, β_{1-6} are idiosyncratic tuning coefficients for the neuron, and v is the fitted moment ($\hat{\mu}$ or $\hat{\sigma}^2$). [For our dataset, fitting Eq. 3 to means and variances provided better and

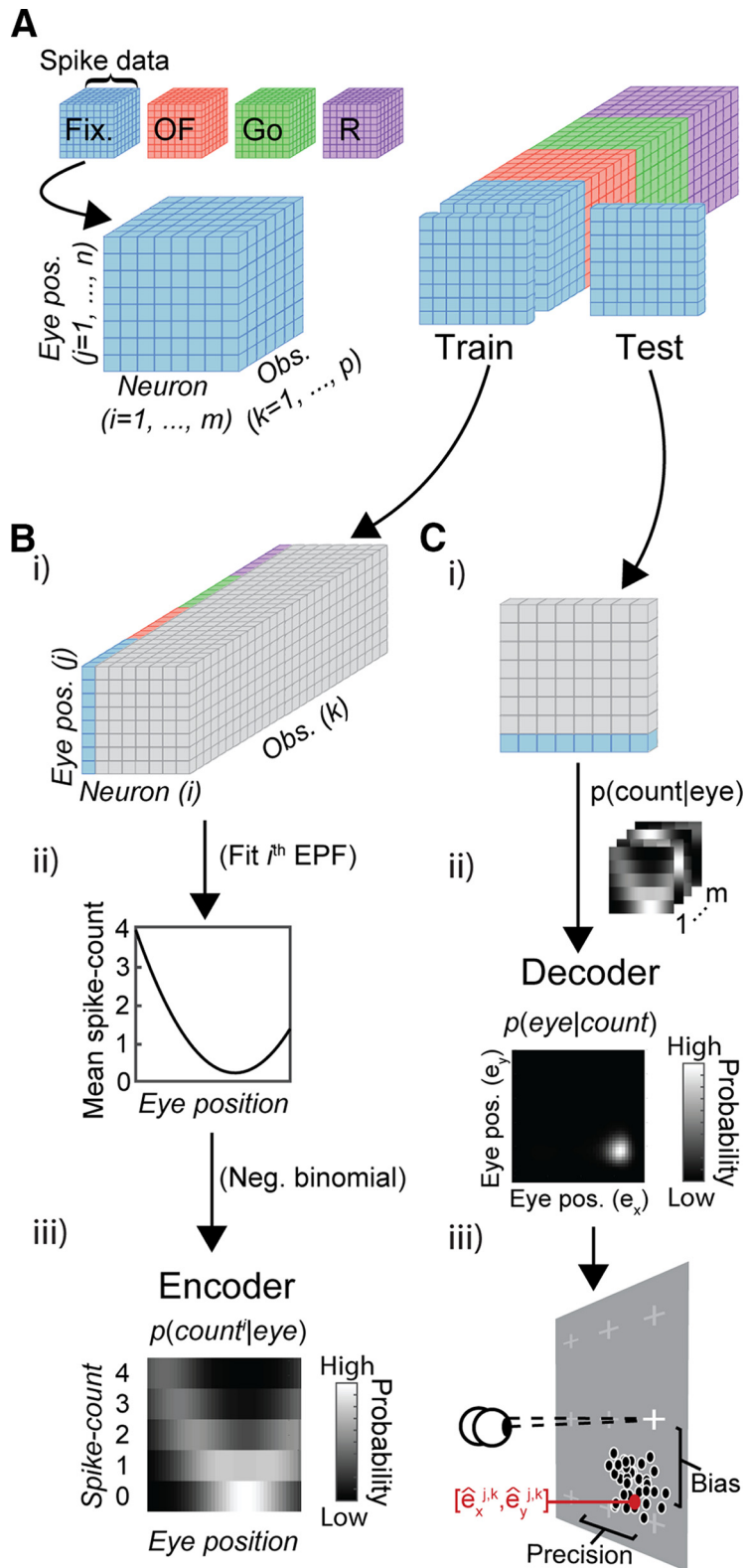


Figure 2. A maximum likelihood decoding approach for testing the robustness of eye position coding across task contexts. **A**, Spiking data for each context, depicted here as 3D matrices, were separated into a training set (consisting of spike-counts for each neuron collected during every eye position and context) and a testing set (consisting of only one set of spike-counts per neuron and eye position, collected from a single nominated context). Context names are abbreviated as Fix (Fixation), OF (Optic Flow), Go (Go Cue), and R (Reach). Training and testing sets were independent. **Bi**, “Training” the decoder was achieved by estimating a probabilistic generative model for spike-counts, separately for each neuron (e.g., using the colored slice of the “train” matrix for the i -th unit). **Bii**, As a first step, we fitted an EPF (Eq. 3) to the nine mean spike-counts (one for each eye position). The plot shows an EPF for a hypothetical neuron and for only a single dimension of eye position (in practice, EPF surfaces were 2D). A corresponding EPF was also fitted to the spike-count variance values

more robust solutions than fitting to a and b directly. It is also more consistent with traditional approaches to modeling EPFs].

The generative model built from Equations 1–3 and the negative binomial assumption fully specify a neuron’s spiking statistics across a continuous 2D eye position space and is referred to as a probabilistic EPF (pEPF) (Morris et al., 2013; Morris and Krekelberg, 2019). To compute the likelihood function at a given eye position, $p(R|e_x, e_y)$, we evaluated Equation 3 and converted the mean ($\hat{\mu}$) and variance ($\hat{\sigma}^2$) values back to the corresponding native parameters of the negative binomial as follows:

$$\hat{b} = \frac{\hat{\mu}}{\hat{\sigma}^2} \quad (4)$$

$$\hat{a} = \frac{\hat{\mu}\hat{b}}{(1 - \hat{b})} \quad (5)$$

The obtained values were then used to compute the spike-count likelihood function for this eye position as follows:

$$\hat{p}(R|e_x, e_y) = \text{Neg. Binomial}[\hat{a}, \hat{b}] \quad (6)$$

The process described above (Eqs. 1–6) was done separately for each neuron. In some cases, however, some care was required to tame the values of a and b . The negative binomial variance is always greater than the mean, and hence leads to fit errors for samples in which the variance is too low (which we defined as $<1.0001 \times r$). Accordingly, in those cases, instead of Equations 1 and 2, the mean spike-count, r , was estimated directly from the spike-counts and the variance was set to $1.0001 \times r$.

To decode eye position, given a pseudo-population response vector, r_{pop} , we first calculated the likelihood of the given spike-counts at each eye position, $\hat{p}(r|e_x, e_y)$ for each neuron, and then converted those values into an *a posteriori* density across eye position space, $p(E_x, E_y|r_{pop})$. This was done using Bayes’ rule (assuming a uniform prior across eye positions) to get $p(E_x, E_y|r)$ for each

(not shown). **Biii**, Assuming a negative binomial form for spike-count distributions, the mean and variance EPFs parameterized an “encoder” model linking eye positions to conditional spike-count probabilities. **C**, Testing the decoder. **Ci**, We sampled a vector of spike-counts (one for each unit in the population) for a given eye position from the test data (e.g., blue-shaded segment of “test” matrix). **Cii**, We then calculated the probability of this observation at each eye position in a 2D quasi-continuous grid under the assumed encoder models. The position with the highest probability was chosen as the decoder’s estimate of eye position for that observation (e.g., the red dot shown in the diagram in **Ciii**). **Ciii**, Performance across trials was quantified by the bias (absolute distance of the centroid [median] estimate to the true position) and precision (span of the central 50% of estimates, averaged over the nine eye positions (diagram represents hypothetical estimates for a single eye position)).

neuron, taking the log of those values, and then summing them pointwise across neurons (an operation that assumes independence across neurons).

Finally, the eye position associated with the maximum *a posteriori* probability (MAP) was used as the decoder's final estimate of eye position as follows:

$$\text{MAP}[\hat{e}_x, \hat{e}_y] = \underset{e_x, e_y}{\operatorname{argmax}} p(E_x, E_y | r_{\text{pop}}) \quad (7)$$

Although the pEPFs are continuous, in practice, these calculations were performed over a 2D grid of eye positions from -20° to 20° along horizontal and vertical dimensions with a resolution of 1° . This process was repeated 500 times for each of the nine experimental eye positions, using independent samples (spike-counts) to estimate the pEPFs of each neuron (i.e., train the decoder) and to test its performance by decoding via the $\text{MAP}[\hat{e}_x, \hat{e}_y]$. This yielded 500 decoded (predicted) eye positions for each of the nine true eye positions and for each of the four tested task contexts.

The Euclidean distance between the 2D centroid (median) of these predictions and the true position was used as a measure of decoder bias (accuracy) for each condition. The scatter of the predictions, defined as the span of the central 50%, was used as a measure of precision. Performance of the decoder was then assessed by comparing decoder biases across the different tested contexts.

Within-context decoder. The universal decoder is trained on data from all contexts and so is unlikely to embody the best possible read-out strategy for samples from any given context, even if it performs well. We therefore compared the universal decoder with decoders trained and tested on the same context. MLE is an optimal decoding approach provided that the assumed generative model matches the true data-generating mechanism. This should be approximately true when data from the same context are used for both training and testing, save for errors introduced by simplifying assumptions in our analysis (e.g., negative binomial statistics, the simple model form in Eq. 3, or nonstationarities within a context) or sampling error (i.e., limited spike-counts per neuron). To obtain this approximate upper bound, we trained and tested a decoder within each context, again using leave-one-out cross validation. That is, all but one spike-count from each eye position was used to estimate a pEPF for each neuron and then held-out population response vectors were decoded. This was repeated 500 times for each eye position.

Cross-context decoder. Although the universal decoder is likely not the optimal decoder for any particular context, it is perhaps closer to that of some contexts than others. If the aim was to rapidly acquire a dataset to find the universal decoder (e.g., for a brain-machine interface), which behavioral context would most efficiently get you there? To address this question, we trained a new series of decoders on spike-counts from one context and tested them on data from each of the other three contexts (using all 12 unique train-test permutations of the four task contexts). As it was not necessary to hold out samples for cross-validation, all spike-counts in the training context were used to estimate the pEPF of each neuron. Again, this was performed using 500 unique pseudo-population responses to produce 500 estimates for each of the nine true eye positions.

Decoding significance. To compare performance across decoding paradigms, we assessed significance by bootstrapping the spike-count samples recorded for each neuron. Each decoding analysis (universal, within-context, and cross context) was repeated 100 times, with each repeat using an identical bootstrapped sample across decoding paradigms. The same number of spike-counts per eye position and context condition was maintained throughout for each neuron. The universal and within-context decoders were trained and tested on every context, whereas the cross-context decoder was trained on either Fixation or Optic Flow for each analysis and tested using all available contexts. The Fixation and Optic Flow decoders were selected as they represented the worst and best performing cross-context decoders, respectively. Two measures of bias were considered significantly different if there was no overlap between the central 95% of their bootstrapped distributions.

Decoder performance with varying population sizes. If eye position and context were truly encoded orthogonally, any observed bias in the

universal decoder would represent only sampling error, and hence would in principle reduce to zero as the population size increases. Although our sample is tiny compared with that of the brain, we estimated the effects of sample size on decoder performance by subsampling our data. This allows us to estimate an asymptotic value for bias with large populations.

By resampling our neurons with replacement, we ran the universal decoder using only n neurons, where values for n were 20, 40, 80, 120, 160, 200, or 235 neurons, for all four task contexts. These analyses were repeated 100 times, each time using a new random set of sampled neurons. We calculated the mean bias across eye positions for each decoder and then averaged those values across the 100 repeats for each population size, resulting in seven values of mean decoder performance. These values were fitted with a power function as follows:

$$f(n, \varphi, \omega, c) = \varphi n^\omega + c \quad (8)$$

The fitting process was weighted to evenly prioritize fitting the bias of the resampled data at $n = 235$ (a suboptimal estimate of the true value because of a reduction of diversity across neurons) and the observed value for our original sample of 235 unique neurons (our best estimate of the true value for this population size). The fitted function was subsequently evaluated at an n of 10,000 to predict the bias of our decoder given a much larger recorded sample.

Quantification of context effects on eye position tuning. Changes in eye position tuning across contexts for these PPC neurons were analyzed in detail in our previous report (Heider et al., 2010). Here we also quantified the changes in tuning of individual neurons using correlation analysis (see Figs. 1D and 3). The mean spike-count in each context was calculated for each of the nine eye positions for each neuron. For neurons with EPFs, these mean values should show reliable variation across eye positions and hence replicate across independent datasets. We used subsampling to measure the correlation in tuning between independent sets of spike-counts taken from within the same context (within-context correlation). The spike-counts for each eye position were randomly split into two sets and averaged. Tuning similarity was measured using Pearson correlation between the two sets of averaged values. This was repeated 5000 times to yield a mean within-context correlation.

The same approach was used to compare tuning across contexts by correlating a set from one context with that of another. However, because the four contexts had different numbers of spike-counts (i.e., time points), and sample size itself impacts on absolute correlation measures, it was necessary to match sample size for a given within- versus cross-comparison. Sample size (per neuron, per eye position) was thus set to half of the number of samples available in the context with the fewest samples, to allow multiple, independent datasets for both within- and cross-context comparisons.

The ratio of cross- to within-correlations was used as a scale-invariant, model-free metric of how much tuning changed across contexts.

Results

We (re)analyzed the extracellular spiking activity of 235 neurons in areas 7a and DP of macaque PPC during a visually guided reaching task (Fig. 1A). The task spanned four experimentally defined epochs ("contexts"), and the direction of gaze was constant within each trial but varied across trials. For the results reported here, spikes were counted in 100-ms-wide bins across each trial, separately for each neuron, and pooled into separate datasets for each context.

We showed previously that the tuning of these neurons to eye position changed across the different contexts of the task (Heider et al., 2010; Karkhanis et al., 2014). Figure 1B, C shows an example of this context effect for a neuron in which the change in tuning was large. When the animal fixated a point on an otherwise blank display, the neuron responded maximally when the eyes were directed downward; in contrast, when driven by a visual

motion stimulus, the neuron responded maximally when the eyes were directed leftward. The spatial tuning for eye position was thus conditional on task context, as reported in our earlier work.

Although the focus of this work is on how these context effects are structured in population codes, we first summarize these context effects in single neurons. We compared the similarity of spatial tuning between contexts (“cross-context”) with that of a baseline in which independent samples were selected from within the same context (“within-context”). To do so, we calculated the correlation between two independent samples of nine mean spike-counts (one for each eye position in the 3×3 grid) and expressed the context effect as the ratio between cross- and within-context correlations. A value of 1 thus indicates no context effect and values <1 indicate different tuning across contexts.

Figures 1D and 3 show the distributions of context effects across the sample of neurons and across all pairs of contexts. Most pairs of contexts showed large changes in tuning on average, particularly those that included the Fixation context (mean of 0.49 across pairs, indicating a halving of correlation compared with within-context). More detailed analysis of single neuron tuning is provided in our previous reports (Heider et al., 2010; Karkhanis et al., 2014).

Robust population coding of eye position

At first glance, these population-wide changes in tuning might appear to compromise downstream computations that incorporate this representation of eye position. A naive read-out, blind to these changes, would have to adopt a fixed mapping of spike-counts onto eye positions and hence could make large estimation errors. For example, a decoder that relied only on the example neuron in Figure 1C to infer eye position would confuse the horizontal and vertical meridian as context changes because of the rotated axis of maximal modulation. If all neurons changed in that same way, they would be of no benefit to such a decoder and it would be doomed to fail. If instead different neurons contribute different errors (in direction and magnitude), and the errors are uncorrelated, they increasingly cancel out as the population size gets larger (much like an SEM reduces with $\sqrt{N}$ for a Gaussian random variable) (Kim et al., 2016; Roth and Rust, 2019). That is, the tuning changes act as a source of noise (albeit one that is reproducible) and, like trial-by-trial variability, can be overcome through population coding as long they are structured in a way that is not information-limiting (Averbeck and Lee, 2006; Moreno-Bote et al., 2014; Pitkow et al., 2015).

To illustrate how context effects might alter the representation of eye position at the population level, it is useful to consider the instantaneous population activity for N neurons as a vector in an N -dimensional space. As eye position changes, and context is held fixed, the spatial tuning of single neurons gives rise to a high-dimensional population tuning surface (or manifold). Similarly, as context changes, and eye position is held fixed, the context effects shift the population vector along another surface. The geometry of these two coding subspaces determines whether

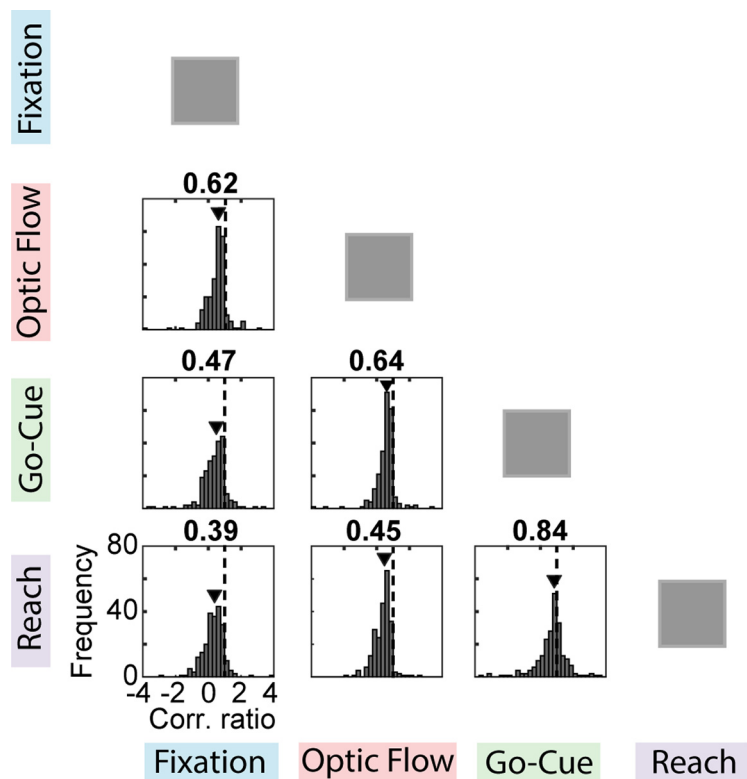


Figure 3. Eye position tuning of single neurons changes across different sensory and behavioral contexts. Each panel is plotted in the same format as in Figure 1D. All pairs of contexts are compared. Inverted triangle markers represent the median of each distribution. Dotted lines indicate a ratio of 1 for clarity. A small number of values ($<6\%$ across all cases) fell outside the plotted range of the histograms.

context effects are tolerable or catastrophic for the downstream readout of eye position. If these subspaces are collinear, changes in context cannot be easily discriminated from changes in eye position. Alternatively, if these subspaces are orthogonal, eye position and context could be interpreted accurately by fixed, independent read-outs.

Figure 4 shows a diagram example for a minimal population of three hypothetical neurons. Here, context changes tuning in a way that, when viewed at the population level, nevertheless provides access to an accurate linear read-out of eye position (shown as a projection onto a downstream readout dimension). Context thus changes population activity in a way that is orthogonal to the effects of eye position.

Demixed principal component analysis

To assess whether this might be the case in parietal cortex neurons, we investigated the population-level activity of the 235 neurons recorded in areas 7a and DP of macaque PPC using dPCA (Fig. 5A). This analysis provides a low-dimensional view of population tuning properties and operates on mean spike-counts rather than individual spikes. Accordingly, the focus here is on “signal” variability associated with experimental variables and not “noise” variability observed trial-by-trial (Averbeck and Lee, 2006; Moreno-Bote et al., 2014; Pitkow et al., 2015). Indeed, these neurons were not recorded simultaneously, so our analysis assumes that tuning remains stable over time (Clopath et al., 2017; Dhawale et al., 2017).

dPCA decomposed the population responses across conditions into 35 linear components that captured variability related to the eye (8 components), context (3 components), or their

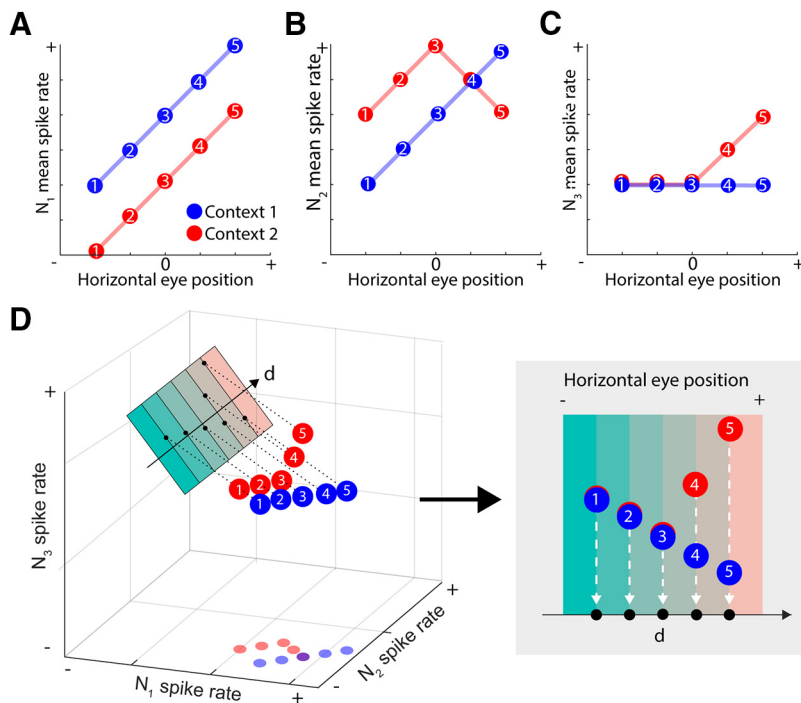


Figure 4. Robust decoding of eye position from a hypothetical population with orthogonal influences of eye position and context. **A–C**, Multivariate tuning curves for three neurons (N_1 , N_2 , and N_3), each with eye position tuning that is context-dependent in different ways. **D**, The data from **A–C** plotted as 3D “population response vectors” (only vector endpoints shown). The response vectors can be projected onto a fixed axis, d , to read out eye position in a way that is robust to changes in context. Each colored band represents a position along the horizontal eye position axis. Context shifts response vectors in ways that are orthogonal (invisible) to the read-out axis, and so this “nuisance variability” is filtered out during the projection.

interaction (24 components). Considered together as groups, the eye position subspace accounted for $\sim 35\%$ of the total variance in the population response captured by dPCA, with the remainder split between the context ($\sim 40\%$) and interaction ($\sim 25\%$) subspaces. Before assessing the orthogonality of these subspaces, we first provide insight into the tuning properties captured in these population-wide components.

Figure 5B, C shows the population state (i.e., mean response vector) for each eye position and context projected onto individual dPCA components. The context components (Fig. 5B) accounted for some well-established properties of single PPC neurons, but here, they are viewed as population-level phenomena. The first (capturing 26.8% of the variance) accounted for increases in spiking in the presence of strong visual stimulation (i.e., comparing Fixation and Optic Flow) and during reach preparation and execution (i.e., during Go and Reach contexts). The second and third context components (7.6% and 3.6% of the total variance) further shaped tuning across the four contexts.

Critically, eye position had almost no effect on population activity when projected onto these context components. As dPCA does not constrain context and eye components to be orthogonal (unlike traditional PCA), this is not a necessary outcome of the dPCA analysis and would not be the case if context and eye position had similar effects on population activity, that is, if context and eye position subspaces were collinear.

This apparent orthogonality is revealed further when the data are projected onto the eye position subspace (Fig. 5C). Again, a classic property of PPC neurons is demonstrated (spatial tuning for eye position), but here as population-wide components. The first and second eye position components (containing 13.6% and

5.7% of the total variance) captured sensitivity to changes in eye elevation and azimuth, respectively. Both showed monotonic effects of eye position on spiking activity and little effect of context. This again suggests that information about eye position is encoded in the population in a way that is orthogonal to the bulk of the context effects. As for the context subspace, further components shaped the population tuning to eye position, including nonmonotonic effects (not shown).

Finally, Figure 5D shows the contributions of the first two interaction (context-eye) components (2.8% and 1.7%) on population activity. The first represents a change in population eye position tuning that occurs in the presence of strong visual stimulation (i.e., from Fixation to Optic Flow contexts), as in the example neuron shown in Figure 1C. The second describes further interactions between eye position and context during reach planning and execution.

The geometry of eye position and context subspaces

The central question of our dPCA analysis is whether the eye and context linear subspaces are oriented orthogonally. To quantify orthogonality, we calculated the principal angles between the eye and context subspaces (see Materials and Methods). These angles have a possible range of 0 (fully collinear) to 90 (orthogonal). To assess statistical significance, we compared the observed angles to a null distribution of angles expected by chance (by sampling random pairs of comparable subspaces, see Materials and Methods).

Figure 6 shows the principal angles between the empirically derived eye and context subspaces. All observed angles were close to the maximum possible value of 90. Comparing observed angles with the range of angles expected by chance (see Materials and Methods) revealed that none significantly departed from orthogonality ($p > 0.05$).

Together, the dPCA and principal angle analyses suggest that 7a and DP neurons support a population-level subspace in which eye position is encoded orthogonally to the effects of changes in task context. In turn, this suggests that a decoder with a fixed interpretation of population activity might be able to estimate eye position reliably despite the effects of context on the spatial tuning of single neurons.

Decoding eye position from neurons in PPC

We assessed whether a statistical decoder with a fixed read-out strategy (i.e., a “universal decoder”) is capable of estimating eye position from PPC activity accurately across experimental contexts. The maximum-likelihood decoder transformed a population vector of spike-counts into a 2D estimate of the eye’s azimuth and elevation. The decoded spike-counts were individual observations (not averaged over trials) and collected over a biologically relevant timescale (100 ms windows). We used three variants of this approach that differed solely in the data used to train and test the decoder. The “universal decoder” was trained

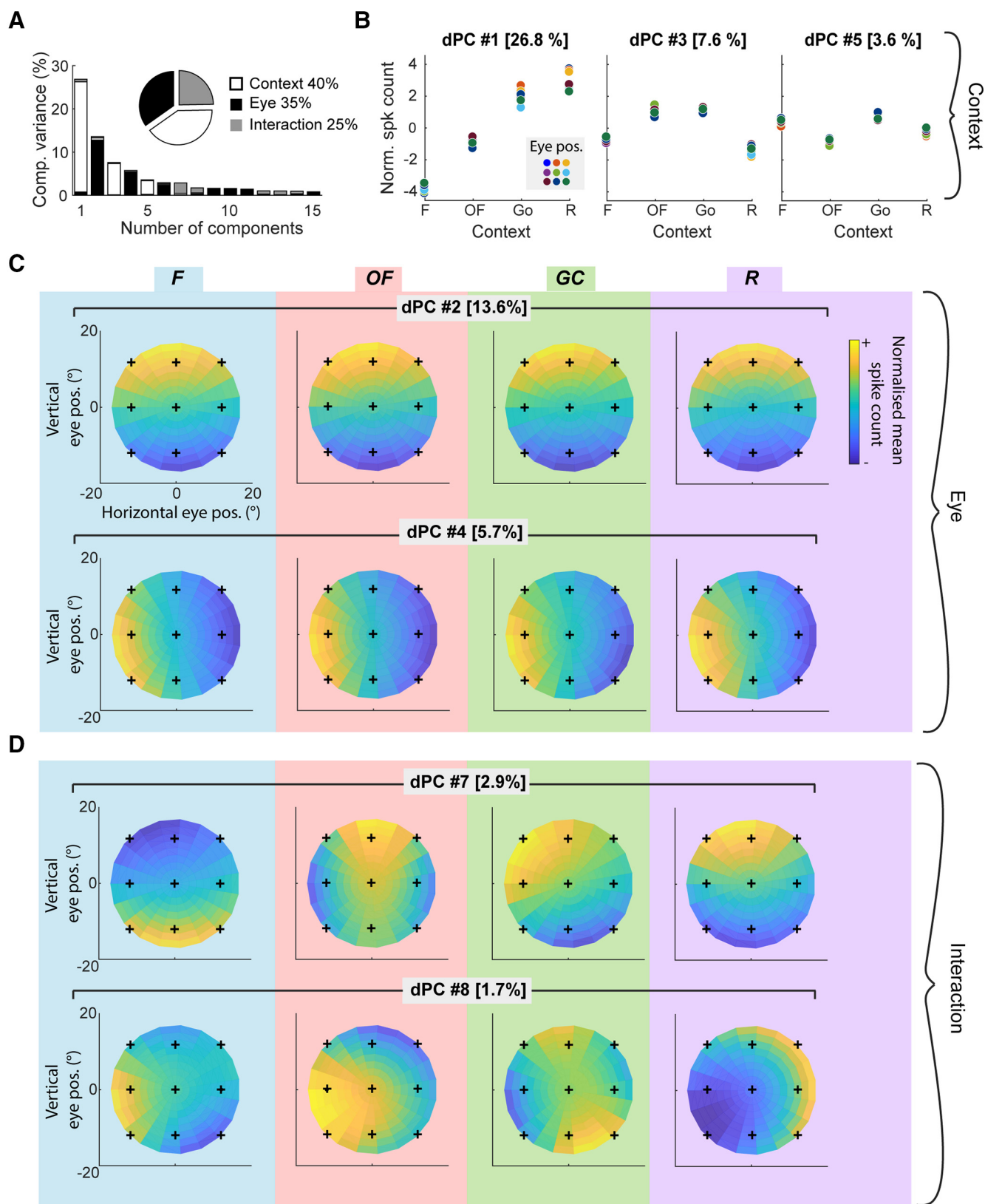


Figure 5. dPCA was used to identify encoding subspaces for eye position, context, and their interaction. **A**, Variance associated with changes in eye position, context, or their interaction captured by each of the first 15 dimensions identified by dPCA (components). Components are generally well separated; however, the first context component (dPC #1) also captures a small amount of eye position- and interaction-related variance (evident also in **B**). The percentages of the total variance explained by all the components associated with eye position, context, or interactions are expressed in the inset pie chart. **B**, The three context components explaining the most variance, plotted as individual traces for each eye position. Components are labeled with percentages of the total variance explained. dPC #1 shows increased activity approaching the reach context, with a concurrent divergence in the responses for each eye position. The other context components do not share a similar divergence across eye positions. **C**, The dPCs explaining the greatest proportion of variance for the eye marginalizations. Each is represented as a flat EPF plotted individually for each context to display any interaction effects. Black crosses represent the eye positions used in the behavioral task. Both eye position components remain

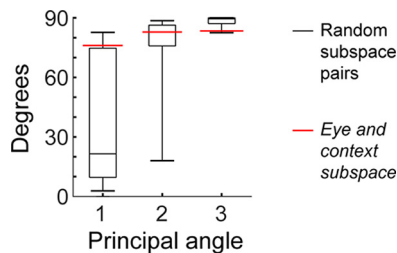


Figure 6. The observed eye and context dPCA subspaces are near-orthogonal. Principal angles between the eye and context subspaces (red lines) compared with a null distribution (box plots with the box spanning the central 50% and whiskers to the 5th and 95th percentiles). The eye subspace consisted of dPCA components 2, 4, 6, 9, 10, 11, 15, and 17. The context subspace consisted of components 1, 3, and 5. Horizontal black lines indicate the median of the distribution (this is obscured by the red marker for principal angle 2 and is compressed at the top of the distribution for principal angle 3). Because the context subspace contains three dimensions, its alignment relative to the eight-dimensional eye position encoding subspace can be captured using three principal angles. Magnitudes of all three angles revealed a high level of orthogonality between the two subspaces (76.05°, 82.89°, and 83.41°).

on data from all contexts and then used to decode samples from each of the four contexts separately (Fig. 2A). The “within-context” decoder was trained and tested using independent datasets from within the same context (see Fig. 9A). Finally, the “cross-context” decoder was trained on one context and then tested on each of the remaining three contexts (see Fig. 10A). Respectively, these decoder variants allowed us to address three specific questions: First, is there a universal, fixed read-out strategy that can accurately estimate eye position from spiking activity regardless of context? If so, how much worse (if at all) is this universal decoder compared with the optimal read-out for each context individually? Last, which of the context-optimized decoders performs best as an approximation to the universal decoder?

Universal decoding

We used the universal decoder to estimate eye position from testing data from each of the four contexts. Figure 7 shows the median estimated eye position (accuracy) and scatter (precision) for each of the nine eye positions, plotted separately for each task context. In all cases, the decoded eye position closely matched the true eye position. A global measure of bias was calculated as the mean distance of the nine points to their respective true eye positions. The universal decoder achieved biases of 1.89°, 1.68°, 1.43°, and 2.11° when tested on the Fixation, Optic Flow, Go Cue, and Reach contexts, respectively, an average bias of 1.78°. As a point of comparison, the bias expected from simply guessing the central eye position every time (the optimal *a priori* strategy) is 12.88°. Bias scores for all decoding methods used in this study are shown in Figure 11. As an alternative measure of accuracy, the “grid” of the median estimates produced by the universal decoder accounted for >97% of the variance in eye position for all tested contexts.

The accurate performance of the universal decoder for all tested contexts shows that eye position information can be

accurately interpreted from any of the tested sensory or behavioral contexts, overcoming large changes in the tuning of single neurons. There remained, however, a small amount of bias in the estimates produced by the universal decoder. With these results alone, it is unclear whether these biases were the result of a mismatch between the assumed generative models in the decoder and the true data-generating mechanism because of the effects of context, or if they result from inevitable sampling errors that arise with small experimental datasets. To disambiguate these possibilities, we assessed the effect of varying the population size to infer how bias would diminish with the larger populations available to the brain.

Decoding with different population sizes

If the universal read-out for eye position were truly orthogonal to the effects of context, it would be expected that any observed bias would simply be resolved by the inclusion of more neurons in our decoding analyses. On the other hand, if the bias arose from a systematic collinearity in the tuning changes, no amount of population pooling would eliminate the bias completely (according to the same computational principles as so-called “information-limiting” correlations in trial-by-trial variability) (Averbeck and Lee, 2006; Moreno-Bote et al., 2014; Pitkow et al., 2015; Roth and Rust, 2019). To assess asymptotic bias levels, we repeated the universal decoding analyses for subsampled populations of different sizes (from 20 to 235 neurons) and fitted the observed bias values for each context with power functions (R^2 of the fits exceeded 0.99 for all cases; Fig. 8).

Using the fitted functions, we evaluated the bias expected for a population of 10,000 neurons. The predicted bias (averaged over contexts) was 0.75, equivalent to a 56% reduction in the bias observed with our original sample of 235 neurons. The Optic Flow context benefitted least from increased population sizes, showing a 39% reduction in bias to 0.65° (38.91%), while the Reach context benefitted most with a reduction of 70.6%.

The performance of the universal decoder for the original sample was consistently better than the mean of the resampled distributions of equal population size (235 neurons). This is expected because sampling with replacement means that many of the neurons are included more than once in each resampled set and the overall diversity of population is reduced; in other words, their individual contribution to bias is magnified rather than nulled by a statistically independent counterpart (i.e., context and eye position subspaces become more collinear) (Kim et al., 2016; Roth and Rust, 2019). In this way, it is clear that context effects are tantamount to a source of noise (albeit reproducible for a given context) that may be overcome through pooling or averaging.

Within-context decoding

Given the large effects of context on eye position tuning in single neurons, the universal decoder is unlikely to be statistically optimal for any given context. Intuitively, a decoder that is trained and tested within each context alone should outperform the universal decoder because the assumed generative models should better match the true data-generating mechanisms (i.e., the true tuning and variance of neurons). Within-context decoding thus provides a useful estimate of the expected upper bound of decoder performance and was performed for each of the four contexts in our experiment.

Figure 9B shows the accuracy and precision of estimates for each of the nine true eye positions, formatted as in Figure 7. Like

unchanged across contexts. The first eye component reflects tuning to changes in eye elevation, the second to changes in eye azimuth. **D**, The first interaction component reveals a change in neural responsiveness that occurs at the onset of the optic flow stimulus, that stabilizes during reach planning. The second interaction component reveals a change in response that occurs during reach planning and continues to change until reach execution.

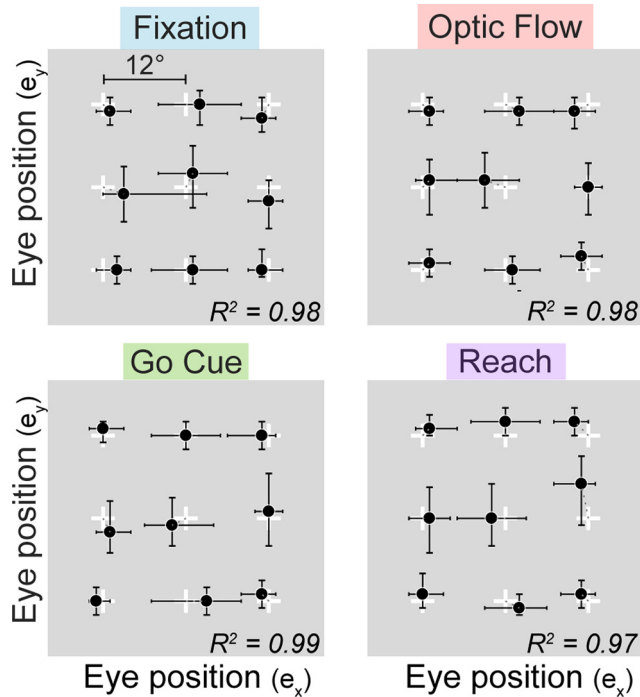


Figure 7. An accurate representation of eye position is available via a fixed read-out of population activity in PPC, regardless of context. The universal decoding approach shown in Figure 2 was used to generate metric estimates of eye position in 2D oculomotor space. Each plot represent the median decoded position (black circles) for each target fixation position (crosses) for a decoder that was trained on spike-counts from all contexts and tested on spike-counts from the labeled context. Error bars capture the precision of decoder estimates from sample to sample, defined as the range of the central 50% of the distribution across 500 pseudo-population responses. In all contexts, the “grid” of decoded positions accounted for >97% of the variance in the true grid of fixation positions.

the universal decoder, decoded eye positions closely matched the true eye position for all contexts. The within-context decoder, however, showed larger biases than the universal decoder, with values of 2.12°, 1.66°, 2.31°, and 2.84° for the Fixation, Optic Flow, Go Cue, and Reach contexts, respectively. The median eye position estimates accounted for >96% of the variance in eye position in all cases.

This potentially counterintuitive finding is likely explained by the effects of sampling error; within-context decoding meant fewer spike-count samples were available during training to infer the generative model for each neuron compared with the universal decoder that pooled over all contexts. The decoder parameters were thus more likely to be mis-estimated, even if in the limit within-context decoding would yield an optimal decoder.

Cross-context decoding

Finally, we used a “cross-context” variant of our approach to establish whether a decoder trained on any one context might effectively approximate the universal decoder and perform well when tested on the other contexts. This would be useful for future potential applications of such a decoder, such as a brain-machine interface, where data acquisition is challenging to obtain and efficient estimation of decoder parameters is paramount. We evaluated performance of cross-context decoders across all train-test pairwise combinations of the four contexts (Fig. 10B).

Cross-context decoding performance was measured in two ways: first by taking the average of the biases found for each of the three tested contexts, and second by considering the worst

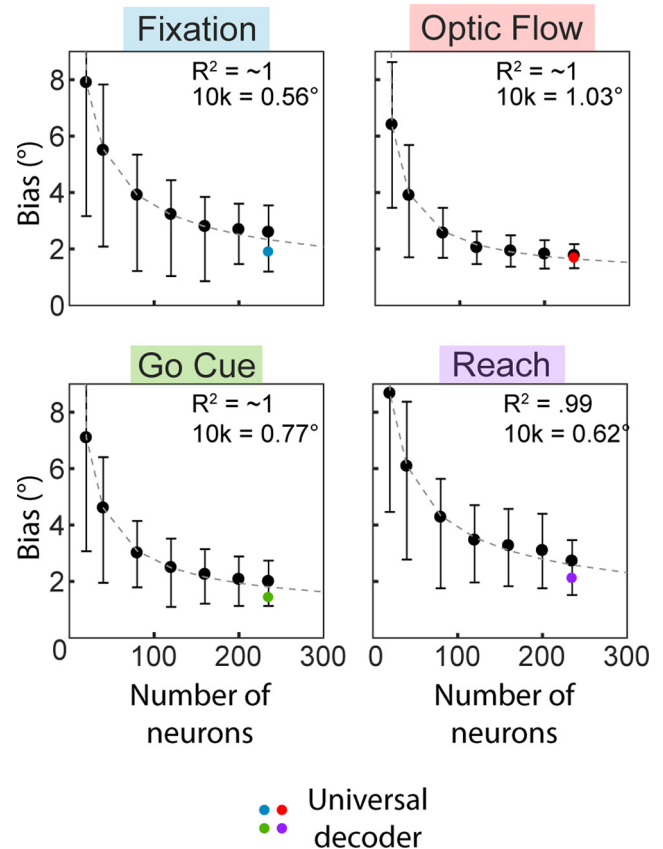


Figure 8. Mis-estimation of eye position reduces with increased population size because of increasingly diverse effects of context on eye position tuning. Universal decoding was performed using subsampled population sizes from 20 to 235 neurons, sampled with replacement from the total 235 recorded neurons. Resampling was repeated 100 times for each data point in the plot. The means of the resulting bias distributions (black circles) were fit by a power function (dashed lines; Eq. 8), which was evaluated to predict the bias anticipated if the decoder had access to 10,000 neurons. Error bars indicate 2.5th to 97.5th percentile ranges. Colored points represent the true universal decoding result. In all cases, the bias of the true universal decoder was less than the mean of the resampled distribution at 235 neurons. The power function was weighted to equally account for both the resampled and original data point.

performing test context. The average biases achieved across contexts was 5.54°, 3.42°, 4.48°, and 6.96° for decoders trained on the Fixation, Optic Flow, Go Cue, and Reach contexts, respectively. All were less accurate than the universal decoder’s average performance of 1.78° across contexts and the within-context average of 2.23° (Fig. 11).

Considering the worst performing test context of each cross-context decoder, we again found that no single context contains a representation of eye position that supports estimation of a decoder that approximates the performance of the universal decoder. The decoder that was trained on the visually sparse Fixation context performed worst overall, and particularly when tested on the Go Cue context with a bias of 6.04° (4.61° greater than the equivalent universal decoder result).

Figure 11 summarizes and compares the performance of the three decoder variants used in this study. The universal decoder showed the most accurate performance, significantly improving on the within-context decoder for the Go Cue and Reach contexts (showing equivalent performance for the others) and significantly improving on the results of the best cross-context decoder (Optic Flow-trained; Fig. 11). Together, the results show that eye position can be interpreted from the activity of neurons

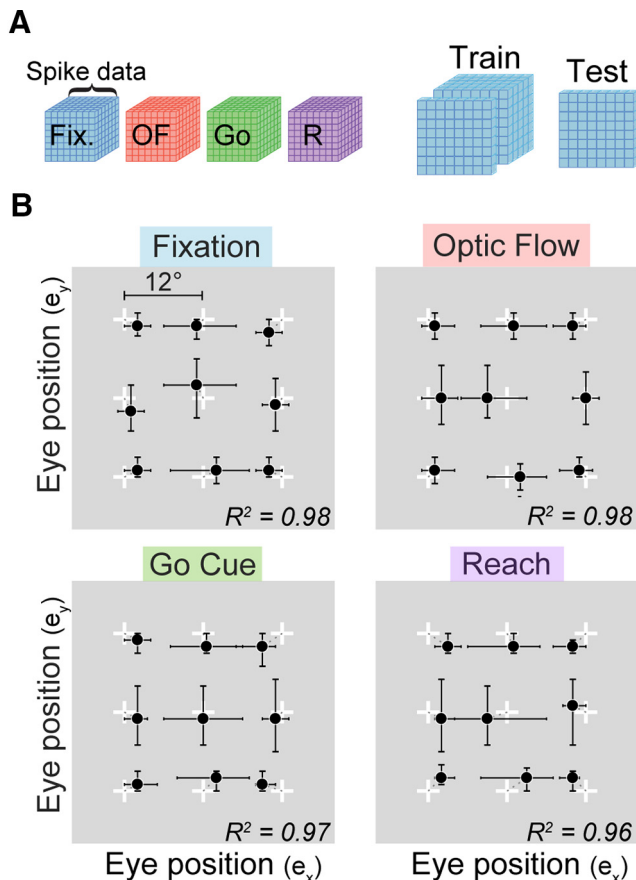


Figure 9. Performance of decoders that were optimized for each experimental context. **A**, Eye position was estimated using a “within-context” variant of our approach in which the decoder was trained and tested on spike-counts from within the same experimental context (Fixation is shown as an example). **B**, Performance of the within-context decoders, plotted as in Figure 7.

in areas 7a and DP using a fixed, context-invariant read-out strategy.

Discussion

Our results show that metric, 2D estimates of eye position can be decoded accurately from neural activity in PPC (areas 7a and DP) throughout four different contexts of a visually guided reaching task. This was achieved using a fixed, context-invariant readout of spike-counts by a probabilistic “universal decoder.” The task contexts included changes in the visual stimulus, a delay epoch that required vigilance, an emerging intention to move the arm toward the target, and movement of the arm. These task variables not only influence neural responses in these cortical areas but also cause widespread changes in the tuning of neurons for eye position (Heider et al., 2010). Thus, it is perhaps surprising that the information represented at the population level is so robust, even when extracted by a naive, fixed decoder.

Using a supervised dimensionality reduction approach (dPCA), however, we demonstrated that the result can be understood by the geometry of the subspaces of the population response in which eye position and context are (respectively) encoded. We showed that context shifts the population response vectors along axes that are approximately orthogonal to those that encode eye position. This allows the universal decoder to “view” the population activity in a way that retains variability because of changes in eye position but is blind to variability

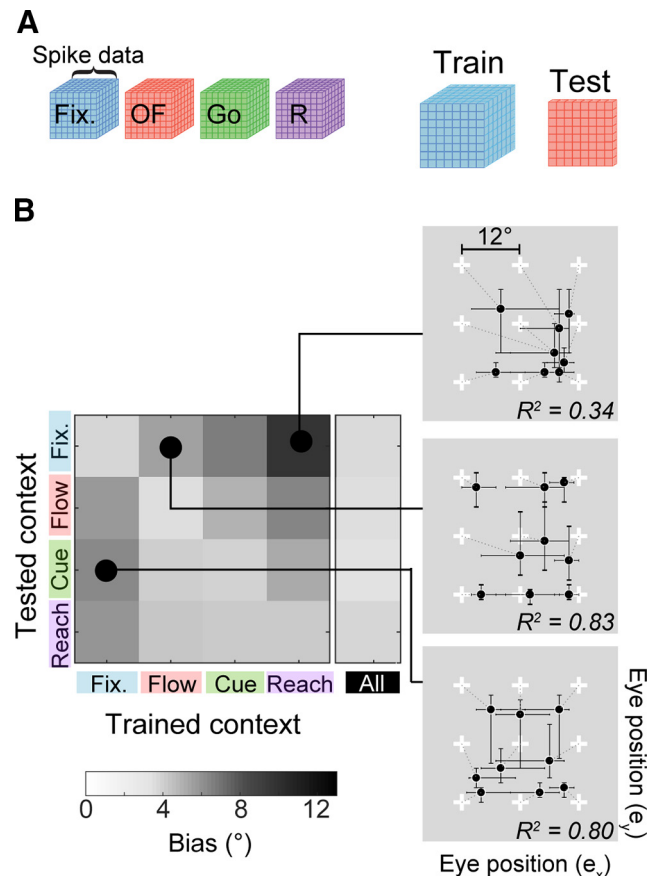


Figure 10. No single sensory or behavioral context can be used to yield an accurate universal decoder. **A**, Eye position was estimated using a “cross-context” variant of our approach in which the decoder was trained on spike-counts from one context and then tested on spike-counts from each of the other three contexts (in the example, Fixation and Optic Flow data were used for training and testing, respectively). **B**, Bias (mean constant error) scores for all combinations of train and test datasets. Pixels along the leading diagonal correspond to the within-context decoders shown in Figure 9. The discontinuous column (ALL) corresponds to the universal decoder performance shown in Figure 7. Off-diagonal pixels show the performance of suboptimal, cross-context decoders. Breakout panels show the full “grid” of decoded eye positions for the tested contexts that showed the greatest bias for three of the different cross-context decoders, plotted as in Figure 7.

arising from changes in context (Churchland et al., 2012; Kaufman et al., 2014; Elsayed et al., 2016; Gallego et al., 2017; Roth and Rust, 2019). In this sense, the effects of the irrelevant “nuisance variable” (context) act as a source of noise that can be filtered out given a sufficiently large population.

The most accurate estimates of eye position were produced using our universal decoder, which was trained on spiking data from all experimental contexts. With an average bias of 1.78°, performance was similar to, or better than, those produced using decoders that were optimized for each context independently, despite the large qualitative differences in single neuron tuning. Although we expect that context-optimized decoders would, by definition, outperform the universal decoder in the limit, for real-world experiments like ours, the universal decoder has the advantage of having a larger training set on which to estimate its internal parameters (i.e., the parameters of the negative binomial distributions and regression equations used to quantify and extrapolate the effect of eye position on single neuron spiking behavior).

We consider the observed bias of the universal decoder low given our small sample size (235 neurons) and expect it to be

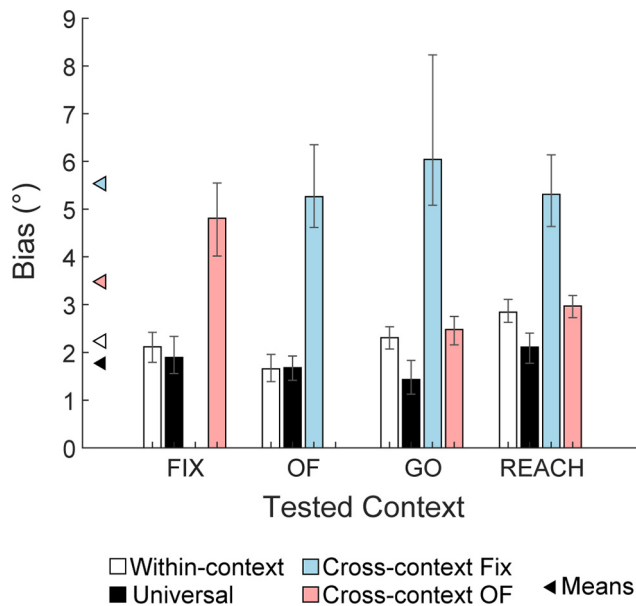


Figure 11. The universal decoder outperforms within- and cross-context decoders and has good absolute accuracy. The plot shows the average bias for the universal, within-context, and cross-context decoders. Only cross-context decoders trained on the Fixation (worst-performing) and Optic Flow (best-performing) contexts are shown. Error bars indicate the 2.5th to 97.5th percentile range of 100 bootstrapped datasets (see Materials and Methods). The universal decoder performed better than other decoders for all test contexts, except for the within-context decoder when tested on Fixation and Optic Flow, with which it was comparable.

even lower for larger population sizes. Our analysis of varying population size through subsampling our data suggests that asymptotic accuracy is likely well below 1°, although additional experiments with larger samples are needed to confirm this prediction. Just as noise can be averaged out through pooling large populations, adding neurons increases the diversity of the population; and if encoding subspaces are orthogonal, as our dPCA analysis suggests, errors should approach zero.

The results suggest that PPC is home to a representation of eye position that is robust to changes in task context and likely of sufficient accuracy to support goal-directed behaviors (reaching, navigation, etc.) and stable spatial vision across changes in eye position (White and Snyder, 2007; Morris et al., 2012, 2013, 2016; Graf and Andersen, 2014; Bremmer et al., 2016; Morris and Krekelberg, 2019). This supports the purported functional roles of EPFs (or “gain-fields”) in perception and behavior (Andersen et al., 1985, 1993; Zipser and Andersen, 1988; Salinas and Abbott, 1996, 2001; Pouget et al., 2002). It also suggests that areas DP and 7a could be suitable targets for a brain–machine interface for spatial vision and assisted behavior (Aflalo et al., 2015; Bremmer et al., 2016; Filippini et al., 2020).

This result was far from guaranteed. If eye position tuning changed across contexts in different, systematic ways, the universal decoder would fail. For example, if a change in context tended to rotate or invert the EPFs in a common direction (e.g., clockwise) or along a common axis (e.g., flipping horizontally), the change in population response would resemble a true shift in eye position and the decoder would consistently make the same types of error. In the geometric framework, such changes in activity because of context would work along an axis that is collinear with those that encode eye position and would thus bias the decoder. In this sense, changes in tuning can affect computations in similar ways to the effects of “information-limiting” correlations in trial-by-trial variability, which cannot be simply

averaged out through pooling larger populations (Averbeck and Lee, 2006; Moreno-Bote et al., 2014; Pitkow et al., 2015; Roth and Rust, 2019).

Rapid estimation of a universal decoder?

A simple algorithm that can estimate eye position or the location of visual objects in peripersonal space robustly would be of great value for brain–machine interfaces. A challenge in that area is to collect enough data to estimate the parameters of the decoding algorithm, so efficient methods are important. We assessed whether any one of the four task contexts allowed estimation of a decoder that performed as well as the universal decoder. In other words, in which context does the PPC population activity most reflect the “universal code”?

Our cross-context decoding approach, in which we trained on one context and tested on each other context, revealed that no single context is sufficient to yield a universal decoder: decoding performance was considerably worse in all cases. This suggests that the nonstationary statistical properties of neurons as context changes need to be incorporated into the decoding model. We achieved this in part by assuming a negative binomial distribution of spike-counts, which captured the increased variability associated with tuning changes. Further, it suggests that training must incorporate exposure to highly varied environments and tasks, most likely more than the four used in the present work.

Multiplexed population codes

Our results are consistent with previous reports of invariant, multiplexed population codes in cortex; as examples, V1 activity supports the read-out of eye position regardless of whether the eyes are stationary, jumping from one point to another, or pursuing a moving target (Morris and Krekelberg, 2019); PPC activity supports the read-out of past, current, and future eye positions simultaneously (i.e., time-lagged variants of the same signal) (Morris et al., 2016); MSTd activity supports the read-out of heading direction and object motion (Kim et al., 2016; Sasaki et al., 2017); and M1 and PMd neurons encode preparatory motor signals in a way that prevents premature motor execution (Churchland et al., 2012; Kaufman et al., 2014; Elsayed et al., 2016; Gallego et al., 2017, 2020).

A stable representation of egocentric space

Eye position information is central to how the brain constructs an internal model of the outside world for perception and action. For vision, it is needed to map the positions of objects in the retinal image, which change with every eye movement, onto their true positions in the world and relative to the body. Thus, visual information only really becomes useful for guiding actions (e.g., reaching, locomotion) when it is combined with information about the direction of gaze (i.e., head and eye position).

For audition, eye position information is needed to shift gaze toward a target that is heard but currently out of view (or to use auditory information to improve localization) (Witten and Knudsen, 2005). Sound locations are encoded relative to the head initially, but eye movements require their locations relative to the fovea (i.e., eye-centered coordinates). Eye position necessarily provides the mapping between these two spatial frames of reference (Soechting and Flanders, 1992). Some neurons in PPC (e.g., the lateral intraparietal area) are sensitive to auditory targets in eye-centered coordinates, evidencing a coordinate transformation built from eye position signals (Mullette-Gillman et al., 2005).

Our results suggest a computational basis for the robustness of these fundamental sensorimotor transformations to changing sensory and behavioral contexts in everyday life: eye position is encoded along dimensions of PPC population responses (at least in areas 7a and DP) that are orthogonal to those that encode other latent variables (e.g., reach plans, changes in visual input, hand/arm posture). This is analogous to findings in premotor and primary motor cortices, where signals for motor preparation and execution are represented through a common population of neurons but encoded orthogonally (Churchland et al., 2012; Kaufman et al., 2014; Gallego et al., 2017; Perich et al., 2018).

Orthogonal encoding thus appears to be a general strategy for multiplexing sensorimotor signals in neural population codes. This would support invariant read-out of information downstream for action and perception through a fixed pattern of synaptic weights. Similarly, it would allow targeted feedback (e.g., attention, efference copy, etc.) through comparable incoming projections that selectively modulate individual signals. In sum, our results point to a stable and context-invariant representation of eye position in primate PPC and provide a striking example of the robustness of population codes, even when single neurons provide highly unreliable information.

References

- Aflalo T, Kellis S, Klaes C, Lee B, Shi Y, Pejsa K, Shanfield K, Hayes-Jackson S, Aisen M, Heck C, Liu C, Andersen RA (2015) Decoding motor imagery from the posterior parietal cortex of a tetraplegic human. *Science* 348:906–910.
- Andersen RA, Essick GK, Siegel RM (1985) Encoding of spatial location by posterior parietal neurons. *Science* 230:456–458.
- Andersen RA, Snyder LH, Li CS, Stricane B (1993) Coordinate transformations in the representation of spatial information. *Curr Opin Neurobiol* 3:171–176.
- Andersen RA, Snyder LH, Bradley DC, Xing J (1997) Multimodal representation of space in the posterior parietal cortex and its use in planning movements. *Annu Rev Neurosci* 20:303–330.
- Andersen RA, Snyder LH, Batista AP, Buneo CA, Cohen YE (1998) Posterior parietal areas specialized for eye movements (LIP) and reach (PRR) using a common coordinate frame. *Novartis Found Symp* 218:109–122.
- Averbeck BB, Lee D (2006) Effects of noise correlations on information encoding and decoding. *J Neurophysiol* 95:3633–3644.
- Avila E, Lakshminarasimhan KJ, DeAngelis GC, Angelaki DE (2019) Visual and vestibular selectivity for self-motion in macaque posterior parietal area 7a. *Cereb Cortex* 29:3932–3947.
- Bremmer F, Kaminiarz A, Klingenhoefer S, Churan J (2016) Decoding target distance and saccade amplitude from population activity in the macaque lateral intraparietal area (LIP). *Front Integr Neurosci* 10:30–47.
- Churchland MM, Cunningham JP, Kaufman MT, Foster JD, Nuyujukian P, Ryu SI, Shenoy KV, Shenoy KV (2012) Neural population dynamics during reaching. *Nature* 487:51–56.
- Clopath C, Bonhoeffer T, Hübner M, Rose T (2017) Variance and invariance of neuronal long-term representations. *Philos Trans R Soc Lond B Biol Sci* 372:20160161.
- Dhawale AK, Poddar R, Wolff SB, Normand VA, Kopelowitz E, Ölveczky BP (2017) Automated long-term recording and analysis of neural activity in behaving animals. *Elife* 6:e27702.
- Elsayed GF, Lara AH, Kaufman MT, Churchland MM, Cunningham JP (2016) Reorganization between preparatory and movement population responses in motor cortex. *Nat Commun* 7:13239.
- Filippini M, Morris AP, Breveglieri R, Hadjimitsakis K, Fattori P (2020) Decoding of standard and non-standard visuomotor associations from parietal cortex. *J Neural Eng* 17:046027.
- Gallego JA, Perich MG, Miller LE, Solla SA (2017) Neural manifolds for the control of movement. *Neuron* 94:978–984.
- Gallego JA, Perich MG, Chowdhury RH, Solla SA, Miller LE (2020) Long-term stability of cortical population dynamics underlying consistent behavior. *Nat Neurosci* 23:260–270.
- Goris R, Movshon J, Simoncelli E (2014) Partitioning neuronal variability. *Nat Neurosci* 17:858–865.
- Gottlieb JP, Kusunoki M, Goldberg ME (1998) The representation of visual salience in monkey parietal cortex. *Nature* 391:481–484.
- Graf AB, Andersen RA (2014) Inferring eye position from populations of lateral intraparietal neurons. *Elife* 3:2813.
- Heider B, Jandó G, Siegel RM (2005) Functional architecture of retinotopy in visual association cortex of behaving monkey. *Cereb Cortex* 15:460–478.
- Heider B, Karnik A, Ramalingam N, Siegel RM (2010) Neural representation during visually guided reaching in macaque posterior parietal cortex. *J Neurophysiol* 104:3494–3509.
- Karkhanis AN, Heider B, Silva FM, Siegel RM (2014) Spatial effects of shifting prisms on properties of posterior parietal cortex neurons. *J Physiol* 592:3625–3646.
- Kaufman MT, Churchland MM, Ryu SI, Shenoy KV (2014) Cortical activity in the null space: permitting preparation without movement. *Nat Neurosci* 17:440–448.
- Kim HR, Pitkow X, Angelaki DE, DeAngelis GC (2016) A simple approach to ignoring irrelevant variables by population decoding based on multisensory neurons. *J Neurophysiol* 116:1449–1467.
- Knyazev AV (2021) subspace.m, MATLAB Central File Exchange. Accessed March 24, 2021. Available at <https://www.mathworks.com/matlabcentral/fileexchange/55-subspace-m>.
- Knyazev AV, Argentati ME (2002) Principal angles between subspaces in an A-based scalar product: algorithms and perturbation estimates. *SIAM J Sci Comput* 23:2008–2040.
- Kobak D, Brendel W, Constantinidis C, Feierstein CE, Kepecs A, Mainen ZF, Qi XL, Romo R, Uchida N, Machens CK (2016) Demixed principal component analysis of neural population data. *Elife* 5:e10989.
- MacKay WA (1992) Properties of reach-related neuronal activity in cortical area 7A. *J Neurophysiol* 67:1335–1345.
- Masquelier T (2013) Neural variability, or lack thereof. *Front Comput Neurosci* 7:7.
- Moreno-Bote R, Beck J, Kanitscheider I, Pitkow X, Latham P, Pouget A (2014) Information-limiting correlations. *Nat Neurosci* 17:1410–1417.
- Morris AP, Krekelberg B (2019) A stable visual world in primate primary visual cortex. *Curr Biol* 29:1471–1480.
- Morris AP, Kubischik M, Hoffmann KP, Krekelberg B, Bremmer F (2012) Dynamics of eye position signals in the dorsal visual system. *Curr Biol* 22:173–179.
- Morris AP, Bremmer F, Krekelberg B (2013) Eye position signals in the dorsal visual system are accurate and precise on short timescales. *J Neurosci* 33:12395–12406.
- Morris AP, Bremmer F, Krekelberg B (2016) The dorsal visual system predicts future and remembers past eye position. *Front Syst Neurosci* 10:9.
- Motter BC, Mountcastle VB (1981) The functional properties of the light-sensitive neurons of the posterior parietal cortex studied in waking monkeys: foveal sparing and opponent vector organization. *J Neurosci* 1:3–26.
- Mullette-Gillman OA, Cohen YE, Groh JM (2005) Eye-centered, head-centered, and complex coding of visual and auditory targets in the intraparietal sulcus. *J Neurophysiol* 94:2331–2352.
- Niell CM, Stryker MP (2010) Modulation of visual responses by behavioral state in mouse visual cortex. *Neuron* 65:472–479.
- Pachitariu M, Lyamzin DR, Sahani M, Lesica NA (2015) State-dependent population coding in primary auditory cortex. *J Neurosci* 35:2058–2073.
- Perich MG, Gallego JA, Miller LE (2018) A neural population mechanism for rapid learning. *Neuron* 100:964–976.
- Pitkow X, Angelaki DE (2017) Inference in the brain: statistics flowing in redundant population codes. *Neuron* 94:943–953.
- Pitkow X, Liu S, Angelaki DE, DeAngelis GC, Pouget A (2015) How can single sensory neurons predict behavior? *Neuron* 87:411–423.
- Pouget A, Deneve S, Duhamel JR (2002) A computational perspective on the neural basis of multisensory spatial representations. *Nat Rev Neurosci* 3:741–747.

- Quraishi S, Heider B, Siegel RM (2007) Attentional modulation of receptive field structure in area 7a of the behaving monkey. *Cereb Cortex* 17:1841–1857.
- Rigotti M, Barak O, Warden MR, Wang XJ, Daw ND, Miller EK, Fusi S (2013) The importance of mixed selectivity in complex cognitive tasks. *Nature* 497:585–590.
- Roth N, Rust NC (2019) Rethinking assumptions about how trial and nuisance variability impact neural task performance in a fast processing regime. *J Neurophysiol* 121:115–130.
- Salinas E, Abbott LF (1996) A model of multiplicative neural responses in parietal cortex. *Proc Natl Acad Sci USA* 93:11956–11961.
- Salinas E, Abbott LF (2001) Coordinate transformations in the visual system: how to generate gain fields and what to compute with them. *Prog Brain Res* 130:175–190.
- Sasaki R, Angelaki DE, DeAngelis GC (2017) Dissociation of self-motion and object motion by linear population decoding that approximates marginalization. *J Neurosci* 37:11204–11219.
- Siegel RM, Raffi M, Phinney RE, Turner JA, Jandó G (2003) Functional architecture of eye position gain fields in visual association cortex of behaving monkey. *J Neurophysiol* 90:1279–1294.
- Snyder LH, Grieve KL, Brotchie P, Andersen RA (1998) Separate body- and world-referenced representations of visual space in parietal cortex. *Nature* 394:887–891.
- Soechting JF, Flanders M (1992) Moving in three-dimensional space: frames of reference, vectors, and coordinate systems. *Annu Rev Neurosci* 15:167–191.
- White RL, Snyder LH (2007) Spatial constancy and the brain: insights from neural networks. *Philos Trans R Soc Lond B Biol Sci* 362:375–382.
- Witten IB, Knudsen EI (2005) Why seeing is believing: merging auditory and visual worlds. *Neuron* 48:489–496.
- Zipser D, Andersen RA (1988) A back-propagation programmed network that simulates response properties of a subset of posterior parietal neurons. *Nature* 331:679–684.