

An emerging view of neural geometry in motor cortex supports high-performance decoding

Reviewed Preprint

v2 • November 27, 2024

Revised by authors

Reviewed Preprint

v1 • October 10, 2023

Sean M Perkins, Elom A Amematsro, John P Cunningham, Qi Wang, Mark M Churchland 

Department of Biomedical Engineering, Columbia University, New York, USA • Zuckerman Institute, Columbia University, New York, USA • Department of Neuroscience, Columbia University Medical Center, New York, USA • Department of Statistics, Columbia University, New York, USA • Center for Theoretical Neuroscience, Columbia University Medical Center, New York, USA • Grossman Center for the Statistics of Mind, Columbia University, New York, USA • Kavli Institute for Brain Science, Columbia University Medical Center, New York, USA

 https://en.wikipedia.org/wiki/Open_access Copyright information

eLife Assessment

This paper presents a new method called MINT that is effective at BCI-style decoding tasks. The authors show **convincing** evidence to support their claims regarding how MINT is a new method that produces excellent decoding performance relative to the state-of-the-art. This work is **important** and will be of broad interest to neuroscientists and neuroengineers.

<https://doi.org/10.7554/eLife.89421.2.sa3>

Abstract

Decoders for brain-computer interfaces (BCIs) assume constraints on neural activity, chosen to reflect scientific beliefs while yielding tractable computations. Recent scientific advances suggest that the true constraints on neural activity, especially its geometry, may be quite different from those assumed by most decoders. We designed a decoder, MINT, to embrace statistical constraints that are potentially more appropriate. If those constraints are accurate, MINT should outperform standard methods that explicitly make different assumptions. Additionally, MINT should be competitive with expressive machine learning methods that can implicitly learn constraints from data. MINT performed well across tasks, suggesting its assumptions are well-matched to the data. MINT outperformed other interpretable methods in every comparison we made. MINT outperformed expressive machine learning methods in 37 of 42 comparisons. MINT's computations are simple, scale favorably with increasing neuron counts, and yield interpretable quantities such as data likelihoods. MINT's performance and simplicity suggest it may be a strong candidate for many BCI applications.

Introduction

Brain-computer interfaces (BCIs) seek to estimate target variables, in real time, from observations of neural spiking activity. Target variables may correspond to prosthetic motion [1–6], muscle activity [7–10], cursor control [11–17], navigation [18–20], speech [21–26], handwriting [27], or cognitive states [28–32]. Alternatively, one may estimate the state of the brain, a latent variable, for clinical monitoring [33], data visualization [34], or closed-loop neurally contingent experiments [35]. Because neural spiking is noisy, target variables must be estimated from patterns of spikes that were unobserved during training. Doing so requires assumptions.

Figure 1a illustrates a set of common assumptions. In this view, the primary objects are neural firing rates (which determine the probability of spiking) and a manifold [36–39] that constrains the possible values of the ‘neural state’: a vector containing the rate of every neuron. While that manifold is presumed to be nonlinear, a linear approximation – i.e. a subspace – may be acceptable for practical purposes [38]. Within this subspace, there exist neural dimensions where activity correlates with target variables such as hand velocity, providing a basis for estimating those variables. Generalization (e.g. [13, 14, 40]) is thought to rely on two features. First, if distributions of behavioral variables overlap across tasks, neural-state distributions will also overlap [39]. Second, out-of-distribution generalization can leverage the reliable correlation between behavioral variables and neural activity.

Decoding methods frequently assume many or all of these features. For example, the classic population vector [11, 41, 42] assumes a two- or three-dimensional manifold (for reaches in two or three physical dimensions) and leverages neural dimensions where activity correlates with hand velocity. Similarly, nearly all ‘interpretable’ methods (those that make explicit assumptions rather than learning them from training data) assume neural dimensions where activity correlates – perhaps incidentally, but at least usefully – with target variables. Such methods often assume additional constraints on neural activity related to structure in behavior [16, 17, 43–49]. E.g. a Kalman filter can embody the assumption that velocity and position are dynamically linked [13, 50], while Sani et al. [51] leveraged the assumption that only a subspace within a low-dimensional neural state is relevant to behavior.

Although the perspective in **Figure 1a** has been useful, its assumptions may be true only locally [52] and perhaps not even then. Under an alternative perspective (**Fig. 1b**), the primary objects are neural factors (which determine each neuron’s instantaneous probability of spiking) and a flow-field (gray arrows) governing factor-state trajectories. Within the space of potential factor states, the flow-field ensures that few are visited. The resulting manifold is complex, and may or may not be locally flat in any useful sense. For example, when cycling at different speeds, the manifold is an approximately seven-dimensional tube (as in **Figure 1b**, green) whose long-axis corresponds to cycling speed [53]. Many locations in factor-space are never visited, including the void within the tube. The manifold is thus sparse – a minority of factor-states are observable – even though individual-neuron responses are rarely sparse in the traditional sense (most neurons show time-varying activity during most movements). It is presumed that different tasks (or subtasks) may require different computations, and thus employ different regions of factor space with different local flow-fields. This may result in a complex and extremely sparse manifold. Some factors may, in some regions of the manifold, correlate incidentally with external variables such as velocity. Yet this is not expected to be consistently true.

The assumptions in **Figure 1b** accord with the hypothesis of factor-level dynamics [54–56] that sculpt computation-specific neural geometries [53, 57–63]. As a simple example, DePasquale et al. [55] and colleagues constructed a network of spiking neurons to generate

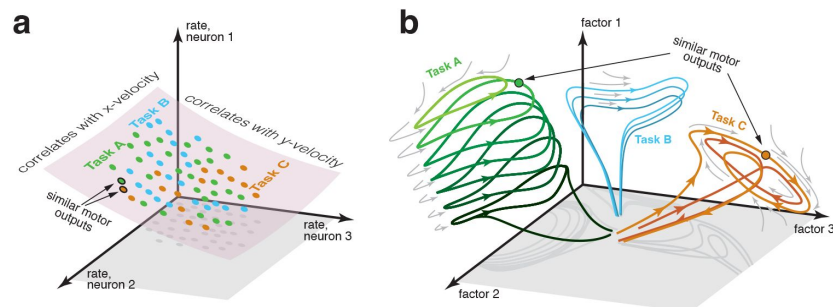


Figure 1.

Two perspectives on the structure of neural activity during motor tasks.

a) Illustration of a traditional perspective. Each neural state (colored dots) is an N-dimensional vector of firing rates. States are limited to a manifold that can be usefully approximated by a subspace. In this illustration, neural states are largely limited to a two-dimensional subspace within the full three-dimensional firing-rate space. Within that restricted subspace, there exist neural dimensions where neural activity correlates with to-be-decoded variables. BCI decoding leverages knowledge of the subspace and of those correlations. If two tasks have similar motor outputs at two particular moments, the corresponding neural states will also be similar, aiding generalization. **b)** An emerging perspective. Neural activity is summarized by neural factors (one per axis), with each neuron's firing rate being a simple function of the factors. Factors may be numerous (many dozens or even hundreds), resulting in a high-dimensional subspace of neural activity. Yet most of that space is unoccupied; factor-trajectories are heavily constrained by the underlying dynamics (gray arrows), such that most states are never visited. The order in which states are visited is also heavily sculpted by dynamics; e.g. trajectories rarely 'swim upstream'. The primary constraint is thus not a subspace, but a sparse manifold where activity largely flows in specific directions. Different tasks may often (but not always) employ different dynamics and thus different regions of the manifold. Because most aspects of neural trajectories fall in the null space of outgoing commands, there are no directions in neural state space that consistently correlate with kinematic variables, and distant states may correspond to similar motor outputs.

muscle activity during the cycling task [61]. Network dynamics produced a limit cycle in 12-dimensional factor space, with any deviations being swiftly corrected by the flow-field. Thus, although the data occupy a sizable subspace, the manifold consists of only those states near the limit cycle. When the network was trained to both cycle and reach, it did so using task-specific dynamics in different neural dimensions, resulting in a sparse, complex manifold with task-specific sub-regions.

Figure 1b also accords with analyses that focused on capturing trajectories within higher-dimensional subspaces [64–67], with the assumption of highly nonlinear mappings between neural activity and behavioral variables [68, 69], and with the finding that locally linear relationships between activity and behavior are prominent during some tasks but not others [20]. It agrees with the proposal that factors (though less numerous than neurons) are plentiful [70–74], with different tasks and task-epochs often using different factors [20, 55, 75–79] (the ‘flexibility-via-subspace’ hypothesis [56]). The assumption of a strong flow-field agrees with empirical limitations on BCI-generated trajectories [80, 81], with the ability of training to ‘open up’ previously unused degrees of freedom [82], and with the finding that decoding is aided by assuming neural-state dynamics [83, 84]. More broadly, low trajectory tangling [53, 61] – a feature of motor cortex activity in most tasks – constrains neural trajectories in ways that imply many of the features in **Figure 1b** and thus argues for approximating the manifold using collections of trajectories rather than a subspace [66, 85].

There is thus a potential mismatch between the structure of the data and the assumptions made by traditional interpretable decoders. This may be one reason why interpretable methods are often outperformed by expressive machine-learning methods [27, 86, 87]. Expressive methods [21, 23–27, 84, 86–93] may be able to implicitly learn, during training, many of the constraints illustrated in **Figure 1b**. Motivated by these considerations, we constructed a novel algorithm, Mesh of Idealized Neural Trajectories (MINT), that explicitly leverages those constraints. MINT takes a trajectory-centric view, where a complicated manifold is approximated using previously observed trajectories and interpolations between them. MINT abandons any notion of neural dimensions that reliably correlate with behavioral variables. Instead, MINT creates a direct correspondence between neural and behavioral trajectories, allowing it to capture highly nonlinear relationships. These relationships can be task-specific if the data so argue. One might have expected that mathematical tractability would be compromised by embracing the unusual assumptions in **Figure 1b**. Yet the requisite computations are surprisingly simple. It is also straightforward to decode a variety of behavioral variables – even those that do not correlate with neural activity – and to estimate the neural state.

MINT’s performance confirms that there are gains to be made by building decoders whose assumptions match a different, possibly more accurate view of population activity. At the same time, our results suggest fundamental limits on decoder generalization. Under the assumptions in **Figure 1b**, it will sometimes be difficult or impossible for decoders to generalize to not-yet-seen tasks. We found that this was true regardless of whether one uses MINT or a more traditional method. This finding has implications regarding when and how generalization should be attempted.

Results

We applied MINT to a total of nine datasets, recorded from primates performing a variety of motor, sensory, and cognitive tasks. Each dataset consisted of simultaneous neural recordings and behavioral variables. We used MINT to decode, based on spiking observations, behavioral variables and the neural state. A central goal was to compare MINT’s decoding performance with existing ‘interpretable’ methods (methods such as the Kalman filter that make explicit assumptions regarding data constraints) and with ‘expressive’ methods (e.g. neural networks that

can learn constraints from training data). Current interpretable methods make assumptions largely aligned with [Figure 1a](#). Thus, if [Figure 1b](#) better describes the data, MINT should consistently outperform other interpretable methods. Highly expressive methods can, given enough training data, implicitly learn a broad range of constraints, including those in [Figure 1a](#), [Figure 1b](#), or some other yet-to-be imagined scenario. This provides a strong test of the assumptions in [Figure 1b](#), because MINT should perform well if those assumptions are correct. Indeed, if [Figure 1b](#) is correct, MINT's performance should be similar to that of highly expressive methods, and may even be better in some cases because MINT can begin with good assumptions rather than having to learn them from data.

We also document properties of the data – directly and during decoding – that bear on the scientific question of whether [Figure 1a](#) or [1b](#) makes better assumptions. A comprehensive characterization of neural trajectory geometry is beyond the scope of this study. Yet, whenever possible, we examine features of neural trajectories relevant to MINT's assumptions. To do so, we begin by focusing on one dataset (MC_Cycle) recorded during a task in which a primate grasps a hand-pedal and moves it cyclically forward or backward to navigate a virtual environment.

Neural trajectories are locally stereotyped and sparsely distributed

During the cycling task, empirical neural trajectories are compatible with solutions found by neural networks trained to produce muscle commands as an output [61]. Those trajectories have features that tend to argue for the assumptions in [Figure 1b](#). Many of these features follow from the property of ‘low trajectory tangling’: during movement, it is never the case that very different factor-trajectory directions are associated with similar factor states. Trajectories are thus locally stereotyped: two trajectories that pass near one another are moving in similar directions. When trajectories travel in dissimilar directions, maintaining low tangling requires that they avoid one another, spreading out into additional dimensions and becoming more sparse in factor space. For example, the elliptical trajectories during forward and backward cycling ([Fig. 2a](#)) appear to overlap in the dominant two dimensions. Yet despite this appearance, the corresponding neural trajectories never come close to crossing. For example, at the apparent crossing-point indicated by the gray disk, forward and backward trajectories are well-separated in dimension 3 (all dimensions use the same scale). This is true at all apparent crossing points, leading to forward-cycling and backward-cycling trajectories that are well-separated. Additionally, both limit cycles contain a void in their center that is never occupied. One might expect that such voids would fill in as additional behaviors are considered. For example, perhaps the interior of the limit cycle becomes occupied when stationary, or at slower speeds? Yet when stationary, neural states are far from the limit-cycle center ([Fig. 2b](#)). Cycling at different speeds involves neural trajectories that spread out in additional ‘new’ dimensions, and thus remain sparse [53]. This illustrates a key difference between the notion of manifold in [Figure 1a](#) and [Figure 1b](#). In [Figure 1a](#), if two distant neural states are both likely to be observed, then the state halfway between them is also reasonably likely to be observed. In [Figure 1b](#), this will frequently not be true.

Neural and behavioral trajectories are non-isometric

A potential consequence of low-tangled trajectories is that similar behavioral outputs can be associated with distant neural states. We found that this was common. As a simple example, there exist many moments when two-dimensional velocity is matched between forward and backward cycling: e.g. at a 45° phase when cycling forward versus 225° when cycling backward. Yet forward and backward trajectories remain distant at all moments. One might suspect that this is simply because angular position is different, even though velocity is matched. Yet neural states can be quite different even when both position and velocity are matched. This point is illustrated ([Fig. 2b](#)) by the neural trajectory that unfolds when transitioning from stationary, to cycling forward seven times, to stationary again. This projection was chosen to highlight dimensions where the

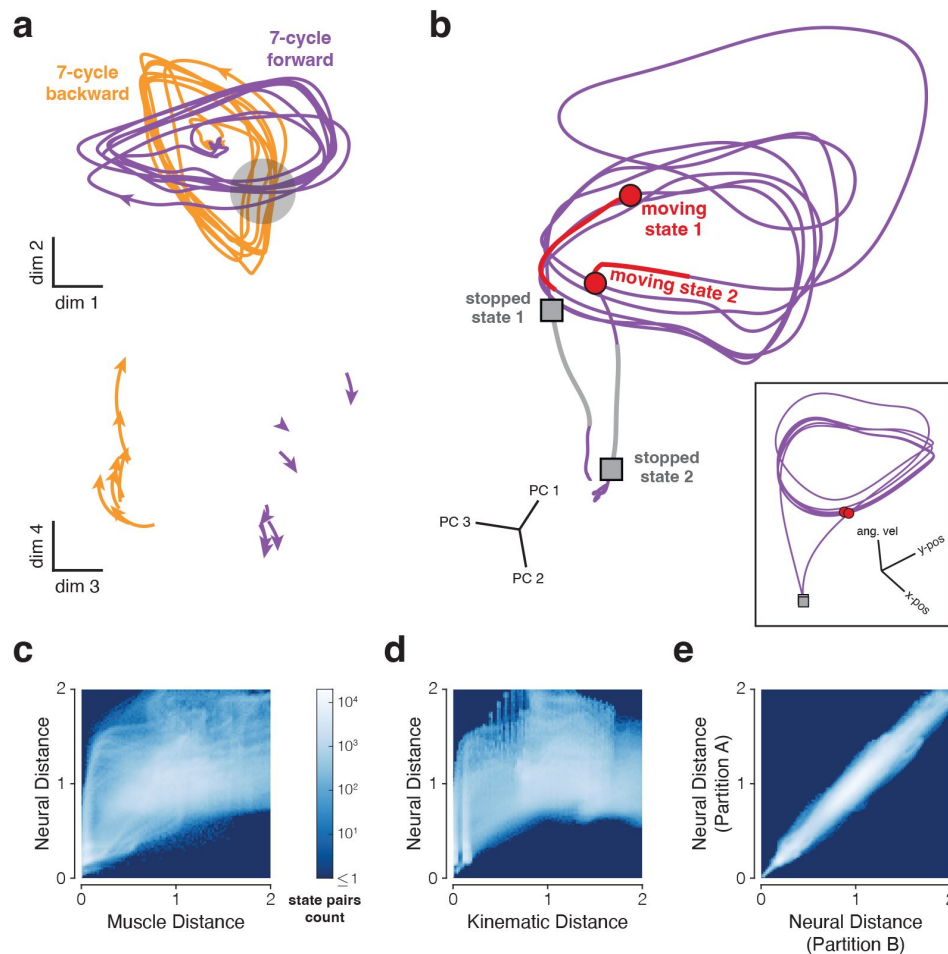


Figure 2.

Properties of neural trajectories in motor cortex, illustrated for data recorded during the cycling task (MC_Cycle dataset).

a) Low tangling implies separation between trajectories that might otherwise be close. Top. Neural trajectories for forward (purple) and backward (orange) cycling. Trajectories begin 600 ms before movement onset and end 600 ms after movement offset. Trajectories are based on trial-averaged firing rates, projected onto two dimensions. Dimensions were hand-selected to highlight apparent trajectory crossings while also capturing considerable variance (11.5%, comparable to the 11.6% captured by PCs 3 and 4). Gray region highlights one set of apparent crossings. Bottom. Trajectories during the restricted range of times in the gray region, but projected onto different dimensions. The same scale is used in top and bottom subpanels. **b)** Examples of well-separated neural states (main panel) corresponding to similar behavioral states (inset). Colored trajectory tails indicate the previous 150 ms of states. Data from 7-cycle forward condition. **c)** Joint distribution of pairwise distances for muscle and neural trajectories. Analysis considered all states across all conditions. For both muscle and neural trajectories, we computed all possible pairwise distances across all states. Each muscle state has a corresponding neural state, from the same time within the same condition. Thus, each pairwise muscle-state distance has a corresponding neural-state distance. The color of each pixel indicates how common it was to observe a particular combination of muscle-state distance and neural-state distance. Muscle trajectories are based on seven z-scored intramuscular EMG recordings. Correspondence between neural and muscle state pairs included a 50 ms lag to account for physiological latency. Results were not sensitive to the presence or size of this lag. Neural and muscle distances were normalized (separately) by average pairwise distance. **d)** Same analysis for neural and kinematic distances (based on phase and angular velocity). Correspondence between neural and kinematic state pairs included a 150 ms lag. Results were not sensitive to the presence or size of this lag. **e)** Control analysis to assess the impact of sampling error. If two sets of trajectories (e.g. neural and kinematic) are isometric and can be estimated perfectly, their joint distribution should fall along the diagonal. To estimate the impact of sampling error, we repeated the above analysis comparing neural distances across two data partitions, each containing 15-18 trials/condition.

neural trajectory resembles the behavioral trajectory (inset). Nevertheless, neural and behavioral trajectories are far from isometric. Red circles indicate two neural states where the corresponding positions and angular velocities are nearly identical: cycling at ~ 1.7 Hz with the hand approaching the cycle's bottom. Despite this behavioral match, the neural states are distant. Gray squares indicate neural states that are distant even though, in both cases, the hand is stopped at the cycle's bottom.

These observations suggest a many-to-one mapping from a neural manifold with one geometry to a behavioral manifold with a very different geometry. To quantify, we leveraged the fact that each neural distance (between a pair of neural states) has a corresponding behavioral distance, obtained for the same pair of conditions and times (allowing for a latency shift). If behavioral and neural geometries are similar, behavioral and neural distances should be strongly correlated. If geometries differ, a given behavioral distance could be associated with a broad range of neural distances. To establish a baseline that incorporates sampling error, we compared neural distances across two partitions. As expected, the joint distribution was strongly diagonal (**Fig. 2e**). In contrast, the joint distribution was non-diagonal when comparing neural versus muscle-population trajectories (**Fig. 2c**) and neural versus kinematic trajectories (**Fig. 2d**). Each behavioral distance (on the x-axis) was associated with a range of different neural distances (on the y-axis). This was true even when behavioral distances were small: for kinematic states with normalized distance < 0.1 , the corresponding neural states were as close as 0.02 and as far as 1.46. Thus, dissimilar neural states frequently correspond to similar behavioral states, as suggested by **Figure 1b** (see states indicated by arrows).

Importantly, the converse was not true. It was never the case that similar neural states corresponded to dissimilar behavioral states. Once past small values on the x-axis, the white regions in **Figure 2c,d** never extend all the way to the x-axis. For example, when kinematic distances were ~ 2 , neural distances were always at least 0.53. Thus, accurate decoding should be possible as long as the decoder can handle the many-to-one mapping of neural states to behavioral variables. In principle this might be accomplished linearly. Suppose neural trajectories mirrored behavioral trajectories but with additional non-behavior-encoding dimensions. A linear decoder could simply ignore those additional dimensions. Yet this will be suboptimal if the ignored dimensions capture the majority of response variance. The above observations thus argue that one may desire quite nonlinear decoding.

A potential concern regarding the analyses in **Figure 2c,d** is that they require explicit choices of behavioral variables: muscle population activity in **Figure 2c** and angular phase and velocity in **Figure 2d**. Perhaps these choices were misguided. Might neural and behavioral geometries become similar if one chooses 'the right' set of behavioral variables? This concern relates to the venerable search for movement parameters that are reliably encoded by motor cortex activity [70, 94–97]. If one chooses the wrong set of parameters (e.g. chooses muscle activity when one should have chosen joint angles) then of course neural and behavioral geometries will appear non-isometric. There are two reasons why this 'wrong parameter choice' explanation is unlikely to account for the results in **Figure 2c,d**. First, consider the implications of the left-hand side of **Figure 2d**. A small kinematic distance implies that angular position and velocity are nearly identical for the two moments being compared. Yet the corresponding pair of neural states can be quite distant. Under the concern above, this distance would be due to other encoded behavioral variables – perhaps joint angle and joint velocity – differing between those two moments. However, there are not enough degrees of freedom in this task to make this plausible. The shoulder remains at a fixed position (because the head is fixed) and the wrist has limited mobility due to the pedal design [61]. Thus, shoulder and elbow angles are almost completely determined by cycle phase. More generally, 'external variables' (positions, angles, and their derivatives) are unlikely to differ more than slightly when phase and angular velocity are matched. Muscle activity could be different because many muscles act on each joint, creating redundancy. However, as illustrated in **Figure 2c**, the key effect is just as clear when analyzing

muscle activity. Thus, the above concern seems unlikely even if it can't be ruled out entirely. A broader reason to doubt the 'wrong parameter choice' proposition is that it provides a vague explanation for a phenomenon that already has a straightforward explanation. A lack of isometry between the neural population response and behavior is expected when neural-trajectory tangling is low and output-null factors are plentiful [56, 61]. For example, in networks that generate muscle activity, neural and muscle-activity trajectories are far from isometric [53, 59, 61]. Given this straightforward explanation, and given repeated failures over decades to find the 'correct' parameters (muscle activity, movement direction, etc.) that create neural-behavior isometry, it seems reasonable to conclude that no such isometry exists.

We further explored the topic of isometry by considering pairs of distances. To do so, we chose two random neural states and computed their distance, yielding $d_{neural1}$. We repeated this process, yielding $d_{neural2}$. We then computed the corresponding pair of distances in muscle space ($d_{muscle1}$ and $d_{muscle2}$) and kinematic space (d_{kin1} and d_{kin2}). We considered cases where $d_{neural1}$ was meaningfully larger than (or smaller than) $d_{neural2}$, and asked whether the behavioral variables had the same relationship; e.g. was $d_{muscle1}$ also larger than $d_{muscle2}$? For kinematics, this relationship was weak: across 100,000 comparisons, the sign of $d_{kin1} - d_{kin2}$ agreed with $d_{neural1} - d_{neural2}$ only 67.3% of the time (with 50% being chance). The relationship was much stronger for muscles: the sign of $d_{muscle1} - d_{muscle2}$ agreed with $d_{neural1} - d_{neural2}$ 79.2% of the time, which is far more than expected by chance yet also far from what is expected given isometry (e.g. the sign agrees 99.7% of the time for the truly isometric control data in **Figure 2e**). Indeed there were multiple moments during this task when $d_{neural1}$ was much larger than $d_{neural2}$, yet $d_{muscle1}$ was smaller than $d_{muscle2}$. These observations are consistent with the proposal that neural trajectories resemble muscle trajectories in some dimensions, but with additional output-null dimensions that break the isometry [61].

Leveraging trajectories to estimate the manifold

The results above, along with other recent results, argue for the perspective in **Figure 1b**. In this view, the manifold of observable states is complicated and sparse. We thus designed MINT to approximate that manifold using the trajectories themselves, rather than their covariance matrix or corresponding subspace. Unlike a covariance matrix, neural trajectories indicate not only which states are likely, but also which state-derivatives are likely. If a neural state is near previously observed states, it should be moving in a similar direction. MINT leverages this directionality.

Training-set trajectories can take various forms, depending on what is convenient to collect. Most simply, training data might include one trajectory per condition, with each condition corresponding to a discrete movement. Alternatively, one might instead employ one long trajectory spanning many movements. Another option is to employ many sub-trajectories, each briefer than a whole movement. The goal is simply for training-set trajectories to act as a scaffolding, outlining the manifold that might be occupied during decoding and the directions in which decoded trajectories are likely to be traveling.

Because training-set trajectories are unlikely to sample the manifold with sufficient density (e.g. one may train using eight reach directions but wish to decode any direction), we designed MINT to interpolate during decoding. We use the term 'mesh' to describe the scaffolding created by the training-set trajectories and the interpolated states that arise at runtime. The term mesh is apt because, if MINT's assumptions are correct, interpolation will almost always be local. If so, the set of decodable states will resemble a mesh, created by line segments connecting nearby training-set trajectories. However, this mesh-like structure is not enforced by MINT's operations. Interpolation could, in principle, create state-distributions that depart from the assumption of a sparse manifold. For example, interpolation could fill in the center of the green tube in **Figure 1b**, resulting in a solid manifold rather than a mesh around its outer surface. However, this would occur only if spiking observations argued for it. As will be documented below, we find that

essentially all interpolation is local. This is a useful fact, because it implies that the estimated neural state (during decoding) will rarely be far from previously observed neural states for which the corresponding behavioral states are known. MINT finds those behavioral states using a direct (and highly nonlinear) mapping between neural and behavioral states. MINT then decodes behavior using local linear interpolation.

Although the mesh is formed of stereotyped trajectories, decoded trajectories can move along the mesh in non-stereotyped ways as long as they generally obey the flow-field implied by the training data. This flexibility supports many types of generalization, including generalization that is compositional in nature. Other types of generalization – e.g. from the green trajectories to the orange trajectories in **Figure 1b** – are unavailable when using MINT and are expected to be challenging for any method (as will be documented in a later section).

Training and decoding using MINT

Training MINT requires computing a library of neural trajectories (Ω) and a corresponding library of behavioral trajectories (Φ). For ease of subsequent computation, neural trajectories are expressed in firing-rate space. MINT is agnostic regarding how neural trajectories are found. For example, one can estimate factor-trajectories, then convert to firing rates via a rectifying nonlinearity [55, 84]. When training data contain repeated trials per condition, it is simplest to use trial-averaging to compute each neuron's rate. We use this approach to illustrate MINT's operations during a center-out reaching task. Neural and behavioral trajectories (**Fig. 3a**) were learned from a training set containing repeated reaches (trials) to each target location. Each moment in that training data corresponds to a condition (c) and an index (k) indicating progress through the movement. For example, $c = 3$ corresponds to the purple reach condition and $k = 240$ indicates a moment 240 ms into that movement. There were many trials for this condition, and thus many measurements of neural activity and behavior for $c = 3$ and $k = 240$. By temporally filtering spikes and averaging across trials, one obtains a neural state $\bar{x}_k^c$ containing each neuron's estimated rate. If desired, firing-rate estimation can be further improved by a variety of means (see Methods). A similar process yields a behavioral state $\bar{z}_k^c$, which may contain position, velocity, or any other variables of interest.

As expected given results above, neural and behavioral trajectories differ during reaching: neural trajectories are stacked loops while position trajectories are arranged radially (velocity trajectories are also radial but return to center as the reach ends). These different neural and behavioral geometries might seem like an impediment to decoding, or to argue for a different choice of target behavioral variables. Yet because neural and behavioral trajectories share the same indices c and k , decoding is straightforward regardless of which variables one wishes to decode.

The parameters of MINT are the libraries of neural trajectories (Ω) and behavioral trajectories (Φ). The trajectories in Ω form a scaffolding that outlines the expected neural manifold during decoding. Ω also specifies the directionality with which decoded states are expected to traverse that manifold. On long timescales, decoded trajectories may be very different from any library trajectory. Yet decoded trajectories will typically be composed of events resembling those from the library, reflecting the local stereotype described above. Prior work has utilized stereotyped behavioral trajectories [43, 47] or a mixture of condition-specific behavioral trajectory models that capture trial-by-trial movement variability [49]. For tractability, those methods assumed low-dimensional neural states that are roughly isometric to arm velocities. With MINT, neural trajectories can take any geometry, are learned empirically, and are related to behavioral trajectories via c and k rather than by an assumed geometric isometry.

During decoding, spike-based observations are binned and counted (**Fig. 3b**, $\Delta = 20$ ms for all analyses unless otherwise specified). Suppose we are at time-bin t' within a decoding session. Recent spiking observations are denoted by $\bar{S}_{t' - \tau' : t'}$, where τ' specifies the number of previous

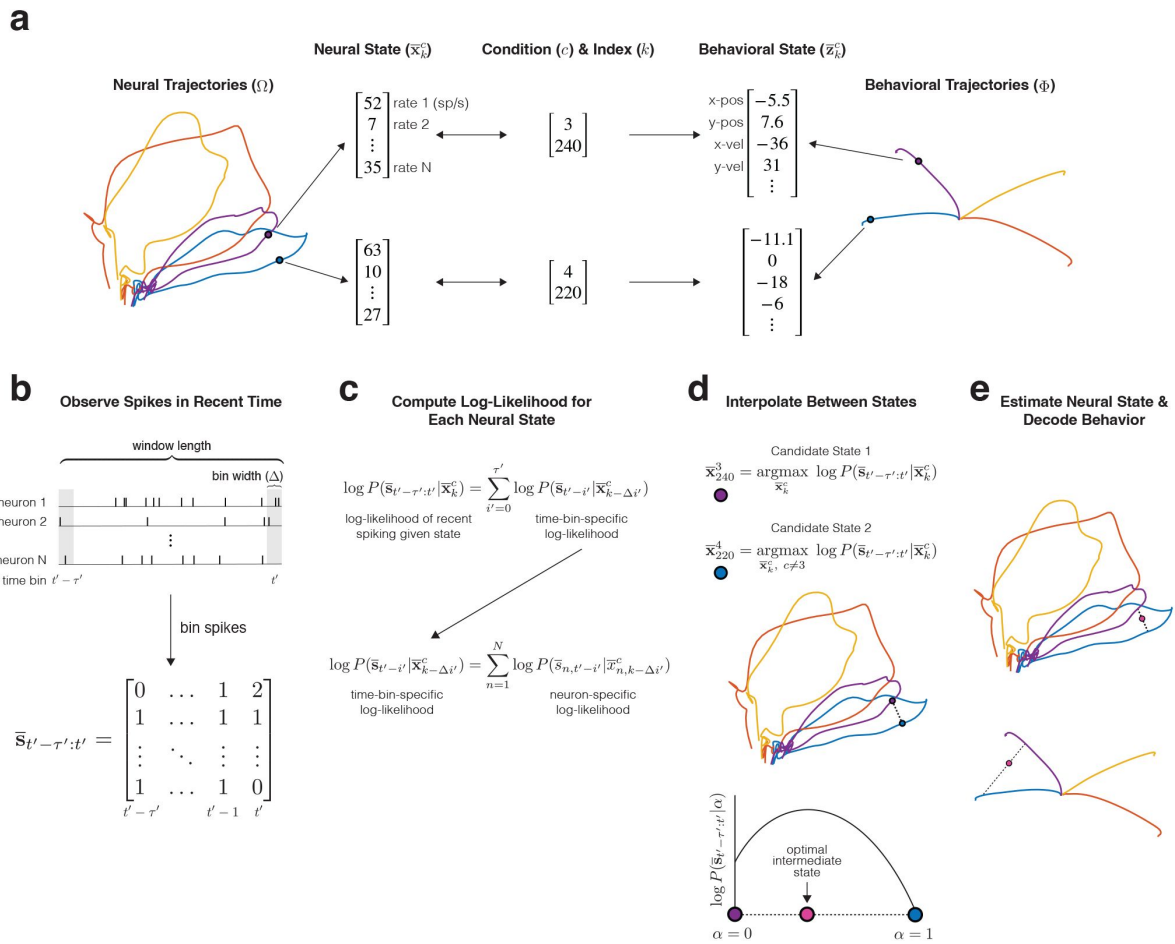


Figure 3.

Example training (top panel) and decoding (bottom panels) procedures for MINT, illustrated using four conditions from a reaching task.

a) Libraries of neural and behavioral trajectories are learned such that each neural state $\bar{x}_k^c$ corresponds to a behavioral state $\bar{z}_k^c$. **b)** Spiking observations are binned spike counts (20 ms bins for all main analyses). $\bar{s}_{t' - \tau': t'}^c$ contains the spike count of each neuron for the present bin, t' , and τ' bins in the past. **c)** At each time step during decoding, the log-likelihood of observing $\bar{s}_{t' - \tau': t'}^c$ computed for each and every state in the library of neural trajectories. Log-likelihoods decompose across neurons and time bins into Poisson log-likelihoods that can be queried from a precomputed lookup table. A recursive procedure (not depicted) further improves computational efficiency. **d)** Two candidate neural states (purple and blue) are identified. The first is the state within the library of trajectories that maximizes the log-likelihood of the observed spikes. The second state similarly maximizes that log-likelihood, with the restriction that the second state must not come from the same trajectory as the first (i.e. must be from a different condition). Interpolation identifies an intermediate state that maximizes log-likelihood. **e)** The optimal interpolation is applied to candidate neural states – yielding the final neural-state estimate – and their corresponding behavioral states – yielding decoded behavior. Despite utilizing binned spiking observations, neural and behavioral states can be updated at millisecond resolution (Methods).

time bins retained in memory. The foundation of MINT is the computation of the log-likelihood of the data, $\bar{\mathbf{S}}_{t'-\tau':t'}$, for a variety of neural states. The decoded state is chosen with the goal of maximizing that log-likelihood. Not all decoders involve data likelihood computations, but likelihoods are desirable when it is possible to obtain them. One typically wishes to decode the state that maximizes data likelihood, or maximizes a cost function that involves likelihoods across multiple states. Knowledge of likelihoods can also indicate the trustworthiness of decoding, as will be discussed in a subsequent section.

To decode, we first compute the log-likelihood that we would have observed $\bar{\mathbf{S}}_{t'-\tau':t'}$ for each neural state in the library (**Fig. 3c**). Importantly, this computation is performed for relatively few states. Even with 1000 samples for each of the neural trajectories in **Figure 3**, there are only 4000 possible neural states for which log-likelihoods must be computed (in practice it is fewer still, see Methods). This is far fewer than if one were to naively consider all possible neural states in a typical rate- or factor-based subspace. It thus becomes tractable to compute log-likelihoods using a Poisson observation model. A Poisson observation model is usually considered desirable, yet can pose tractability challenges for methods that utilize a continuous model of neural states. For example, when using a Kalman filter, one is often restricted to assuming a Gaussian observation model to maintain computational tractability.

The library, Ω , makes it simple to compute data likelihoods that take spike-history into account. Without the constraint of a flow-field, a great many past trajectories could have led to the present neural state. Computing the likelihood of $\bar{\mathbf{S}}_{t'-\tau':t'}$ given all these possible histories, may be impractical during real-time decoding. Yet given a strong flow-field, such histories will be similar on short timescales. For example, if we are presently 240 ms into the purple trajectory in **Figure 3a** (i.e. $c = 3$ and $k = 240$) then we know (approximately) the rate of every neuron over the last few hundred milliseconds. One can thus compute the log-likelihood of the observed pattern of spikes, over the last few hundred milliseconds, given that we are presently at $c = 3$ and $k = 240$. This computation is aided by the fact that, under a Poisson model, spike counts in non-overlapping bins are conditionally independent given knowledge of the underlying rate. Similarly, spike counts for different neurons are conditionally independent given knowledge of their rates. Thus, the log-likelihood of $\bar{\mathbf{S}}_{t'-\tau':t'}$, for a particular current neural state, is simply the sum of many individual log-likelihoods (one per neuron and time-bin). Each individual log-likelihood depends on only two numbers: the firing rate at that moment and the spike count in that bin. To simplify online computation, one can precompute the log-likelihood, under a Poisson model, for every plausible combination of rate and spike-count. For example, a lookup table of size 2001×21 is sufficient when considering rates that span 0-200 spikes/s in increments of 0.1 spikes/s, and considering 20 ms bins that contain at most 20 spikes (only one lookup table is ever needed, so long as its firing-rate range exceeds that of the most-active neuron at the most active moment in Ω). Now suppose we are observing a population of 200 neurons, with a 200 ms history divided into ten 20 ms bins. For each library state, the log-likelihood of the observed spike-counts is simply the sum of $200 \times 10 = 2000$ individual log-likelihoods, each retrieved from the lookup table. In practice, computation is even simpler because many terms can be reused from the last time bin using a recursive solution (Methods). This procedure is lightweight and amenable to real-time applications. The assumption of locally stereotyped trajectories also enables neural states (and decoded behaviors) to be updated between time bins. While waiting for the next set of spiking observations to arrive, MINT simply assumes that each neural state advances deterministically according to the flow-field.

To decode stereotyped trajectories, one could simply obtain the maximum-likelihood neural state from the library, then render a behavioral decode based on the behavioral state with the same values of c and k . This would be appropriate for applications in which conditions are categorical, such as typing or handwriting. Yet in most cases we wish for the trajectory library to serve not as an exhaustive set of possible states, but as a scaffolding for the mesh of possible states. MINT's operations are thus designed to estimate any neural trajectory – and any corresponding

behavioral trajectory – that moves along the mesh in a manner generally consistent with the trajectories in Ω . To illustrate, suppose the subject made a reach that split the difference between the purple and blue reach conditions in **Figure 3a**. MINT identifies two candidate states with high likelihoods, along different neural trajectories (**Fig. 3d**) in the library. MINT then interpolates between these high-likelihood states, and determines whether there exists a higher-likelihood state between them. MINT does so by considering a line segment of neural states parameterized by α , ranging from 0 (at the purple state) to 1 (at the blue state). Because the log-likelihood of the observed spikes is a concave function of α , the optimal neural state can be rapidly identified online using Newton’s method. The decoded behavioral output is then produced by interpolation between the corresponding two behavioral states in Φ , using the same α associated with the optimal neural state (**Fig. 3e**).

The value of α may change with time. For example, this could occur if the true behavioral trajectory began close to the purple trajectory but became progressively more similar to the blue trajectory. The library trajectories that MINT interpolates between can also change with time. Thus, interpolation allows considerable flexibility. Not only is one not ‘stuck’ on a trajectory from Φ , one is also not stuck on trajectories created by weighted averaging of trajectories in Φ . For example, if cycling speed increases, the decoded neural state could move steadily up a scaffolding like that illustrated in **Figure 1b** (green). In such cases, the decoded trajectory might be very different in duration from any of the library trajectories. Thus, one should not think of the library as a set of possible trajectories that are selected from, but rather as providing a mesh-like scaffolding that defines where future neural states are likely to live and the likely direction of their local motion. The decoded trajectory may differ considerably from any trajectory within Ω . Nevertheless, individual decoded states are (empirically) close to states within Ω , as will be discussed in a subsequent section. That proximity makes it reasonable to use the linear interpolation step when decoding behavior.

Behavioral decoding when the ground-truth is known

We begin by assessing MINT’s performance when applied to the activity of an artificial recurrent network of spiking neurons, trained to generate the empirical activity of the deltoid [55]. Doing so provides a situation where the ground truth is known: we know the network’s true output and that spiking is approximately Poisson, resembling that of cortical neurons [55]. The network performed two tasks – reaching and cycling – governed by different local flow-fields in different regions of factor-space (as proposed by the flexibility-via-subspace hypothesis [56]). Indeed, the neural dimensions occupied during cycling are orthogonal to those occupied during reaching (**Fig. 4b**). There are also moments when behavioral output – instantaneous muscle activity – is identical across tasks even though the underlying neural state is very different, in agreement with the empirical many-to-one mapping (**Fig. 2c,d**). These features provide a test of whether MINT can handle geometry similar to that proposed in **Figure 1b**.

Decoded muscle activity (**Fig. 4a**, green) was virtually identical to network output (black). This was true during reaches (R3, R4, R6) and bouts of cycling (C). Decoding seamlessly transitioned between tasks. Decoding R^2 was .968 over ~7.1 minutes of test trials based on ~4.4 minutes of training data. This demonstrates that MINT performs well when data match its motivating assumptions, and also that it may be particularly useful in situations where decoding must automatically switch amongst tasks.

We leveraged this simulated data to compare MINT with a biomimetic decoder whose operations (unlike MINT) mimic the manner in which the true output is produced. The network’s true output is simply a weighted sum of the spikes of 1200 neurons. We thus employed a linear readout (learned via ridge regression) leveraging all 1200 neurons. The resulting decode was very good: $R^2 = .982$ (and would have been unity if readout weights were inferred from connectivity rather than learned from activity). However, the decoding dimension – just like the actual output dimension – captured exceedingly little (~2%) of the total variance of network activity, a feature that is likely

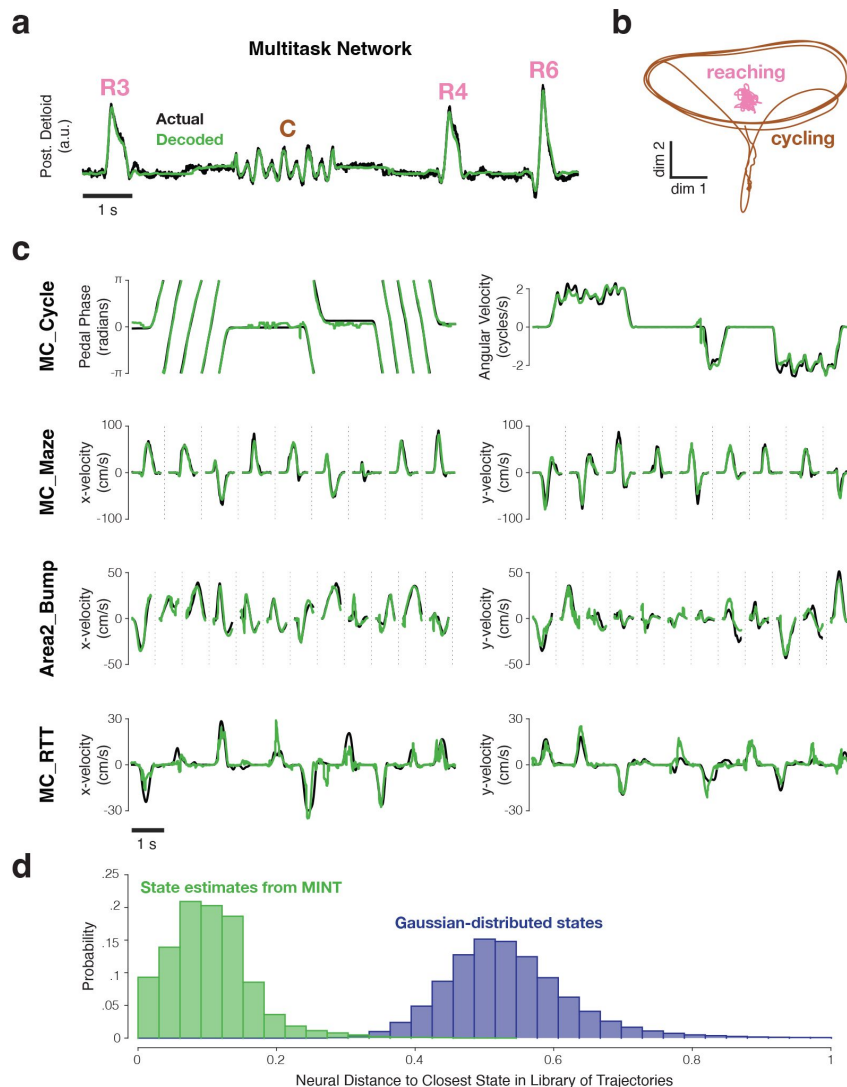


Figure 4.

Examples of behavioral decoding provided by MINT for one simulated dataset and four empirically recorded datasets.

All decoding is causal; only spikes from before the decoded moment are used. **a)** MINT was applied to spiking data from an artificial spiking network. That network was trained to generate posterior deltoid activity and to switch between reaching and cycling tasks. Based on spiking observations, MINT approximately decoded the true network output at each moment. 'R3', 'R4', and 'R6' indicate three different reach conditions. 'C' indicates a cycling bout. MINT used no explicit task-switching, but simply tracked neural trajectories across tasks as if they were conditions within a task. **b)** Illustration of the challenging nature, from a decoding perspective, of network trajectories. Trajectories are shown for two dimensions that are strongly occupied during cycling. Trajectories for the 8 reaching conditions (pink) are all nearly orthogonal to the trajectory for cycling (brown) and thus appear compressed in this projection. **c)** Decoded behavioral variables (green) compared to actual behavioral variables (black) across four empirical datasets. MC_Cycle and MC_RTT show 10 seconds of continuous decoding. MC_Maze and Area2_Bump show randomly selected trials, demarcated by vertical dashed lines. **d)** Distribution of distances between each decoded neural state and the nearest state in the library Ω (green), for the MC_Cycle dataset. To provide a reference, we drew neural states from a Gaussian distribution whose mean and covariance matched the library states, and computed distances to the library states (purple). This emulates what would occur if the manifold were defined by the data covariance and associated neural subspace. The green distribution mostly contains non-zero values, but is much closer to zero than the purple distribution. Thus, MINT rarely decodes a neural state that exactly matches a library state, but most decoded states are close to the scaffolding provided by the library states. Neural distances are normalized by the average pairwise distance between library states.

true of many real networks as well [61]. Biomimetic decoding may therefore suffer when not all neurons are recorded. Indeed, when only using 5% of neurons in the network, performance of the biomimetic decoder dropped from $R^2 = 0.982$ to $R^2 = 0.773$. In contrast, MINT's performance declined only slightly from $R^2 = 0.968$ to $R^2 = 0.967$. This illustrates that, even when biomimetic decoding is possible, it may be outperformed by methods that leverage all neural dimensions rather than just output dimensions. MINT can also decode variables – e.g. kinematics – that are not literally encoded by network activity. These observations call into question an often-implicit assumption: that the activity-to-behavior ‘model’ used during decoding should attempt to approximate the true activity-to-behavior mapping. This is presumably true if a great many neurons can be recorded and the true output of the neural population is the quantity one wishes to decode – yet that is often not the case in real-world BCI applications.

Behavioral decoding across multiple datasets

We investigated MINT's decoding performance across four datasets in which primates performed different motor tasks. All datasets included simultaneously collected spiking activity and behavioral measurements. The MC_Cycle dataset involved cycling a hand-held pedal and provided the motor cortex data in **Figure 2**. The MC_Maze dataset involved recordings from motor cortex during straight and curved reaches [98]. The Area2_Bump dataset involved recordings from somatosensory cortex (area 2) during active and passive center-out-reaches [99]. The MC_RTT dataset involved motor cortex recordings during point-to-point reaches with no condition structure [90]. All decoding was offline, to facilitate comparison across methods using matched data. Many of these datasets are suspected to involve trajectories with properties matching those in **Figure 1b**. However, unlike for the network above, this is an empirical inference and is not known to be true for all datasets. It was thus unclear, a priori, how well MINT would perform. All decoding was causal; i.e. based on spikes from before the decoded time. All decoding examples and performance quantifications utilized held-out test sets. Selection of MINT's (few) hyperparameters is documented in a later section.

For the MC_Cycle dataset (**Fig. 4c**, top row), MINT's decode tracked pedal phase and angular velocity (also see Supp. Video). When the pedal was stationary, erroneous deviations from non-zero angular velocity sometimes occurred (e.g. **Figure 4c**, just before the middle cycling bout) but were brief and rare: decoded angular speed was $<0.1\text{Hz}$ for $\sim 99\%$ of non-moving times. Transitions between rest and cycling occurred with near-zero latency. During movement, angular velocity was accurately decoded and angular phase was decoded with effectively no net drift over time. This is noteworthy because angular velocity on test trials never perfectly matched any of the trajectories in Φ . Thus, if decoding were restricted to a library trajectory, one would expect growing phase discrepancies. Yet decoded trajectories only need to locally (and approximately) follow the flow-field defined by the library trajectories. Based on incoming spiking observations, decoded trajectories speed up or slow down (within limits).

This decoding flexibility presumably relates to the fact that the decoded neural state is allowed to differ from the nearest state in Ω . To explore, we computed, for each moment in the MC_Cycle dataset, the distance between the decoded neural state and the closest library state. Most distances were non-zero (**Fig. 4d**, green distribution). At the same time, distances were much smaller than expected if decoded states obeyed a Gaussian with the empirical covariance (purple). These small distances are anticipated under the view that the empirical manifold (which produces the spiking observations that drive decoding) differs from the distribution defined by the data's covariance matrix.

For the MC_Maze dataset (**Fig. 4c**, second row), decoded velocity tracked actual reach velocity across a variety of straight and curved reaches (108 conditions). Decoded and actual velocities were highly correlated (**Fig. 4c**, middle row, $R_x = .963 \pm .001$, $R_y = .950 \pm .002$; ranges indicate standard errors, computed by resampling test trials with replacement). For comparison with the

other reach-like datasets below, we computed the mean absolute error (MAE) between decoded and actual velocity. Absolute errors were low ($MAE_x = 4.5 \pm 0.1$ cm/s, $MAE_y = 4.7 \pm 0.1$ cm/s), and constituted only $\sim 2\%$ of x- and y-velocity ranges (224.8 and 205.7 cm/s, respectively).

For the Area2_Bump dataset (**Fig. 4c**, third row), correlations with horizontal and vertical hand velocity were also high: $R_x = .939 \pm .007$, $R_y = .895 \pm .015$. Mean absolute errors were small ($MAE_x = 4.1 \pm 0.2$ cm/s, $MAE_y = 4.4 \pm 0.2$ cm/s) and similar to those for the MC_Maze dataset. In relative terms, errors were slightly larger for Area2_Bump versus MC_Maze because the range of velocities was smaller (95.8 and 96.8 cm/s for x- and y-velocity). Errors are thus slightly clearer in the traces plotted in **Figure 4c** (compare second and third rows).

Decoding was acceptable, but noticeably worse, for the MC_RTT dataset (**Fig. 4c**, bottom row). In part this can be attributed to this dataset's slower reaches (86.6 and 61.3 cm/s x- and y-velocity ranges). When analyzing only movement periods (reach speed > 1 cm/s), the median speed across all movement times was 38.1, 16.8, and 5.9 cm/s for MC_Maze, Area2_Bump, and MC_RTT, respectively. Thus, despite being modest in absolute terms ($MAE_x = 2.7 \pm 0.1$ cm/s, $MAE_y = 2.1 \pm 0.1$ cm/s), errors resulted in noticeably weakened correlations between decoder output and behavior: $R_x = .785 \pm .013$, $R_y = .843 \pm .013$. As will be discussed below, every decode method achieved its worst estimates of velocity for the MC_RTT dataset. In addition to the impact of slower reaches, MINT was likely impacted by training data that made it challenging to accurately estimate library trajectories. Due to the lack of repeated trials, MINT used AutoLFADS to estimate the neural state during training. In principle this should work well. In practice AutoLFADS may have been limited by having only ~ 10 minutes of training data. Because the random-target task involved more variable reaches, it may also have stressed the ability of all methods to generalize, perhaps for the reasons illustrated in **Figure 1b**. To aid interpolation amongst library trajectories (and possibly improve generalization), for this dataset we allowed MINT to consider multiple pairwise interpolations when finding the maximum-likelihood state. This aided performance, but only very modestly. We return to the issue of generalization in a later section.

Comparison to other decoders

We compared MINT with four other decode algorithms: the Kalman filter, Wiener filter, feedforward neural network, and recurrent neural network (GRU). The Kalman filter and Wiener filter are historically popular BCI algorithms [1, 50] that are computationally simple and make linear assumptions. The feedforward neural network and GRU are expressive nonlinear function approximators that have been proposed for neural decoding [87, 100] and can be implemented causally. At each decoding time, both the feedforward neural network and the GRU leverage recent spiking history. The primary difference between the two is whether that history is processed all at once (feedforward network) or sequentially (GRU).

Each method was evaluated on each of the four datasets and was asked to decode multiple behavioral variables on held-out test sets (**Fig. 5**). All decoders were provided with a trailing window of spiking observations, binned in 20 millisecond increments. The length of this trailing window was optimized separately for each decoder and dataset (and, like all hyperparameters, was chosen using validation data). For non-MINT decoders, hyperparameters were tuned, using Bayesian optimization, to maximize performance [101]. MINT's few hyperparameters were less extensively optimized. For example, window length was optimized just once for a given dataset for MINT, rather than separately for each set of behavioral variables as was done for most other decoders (excepting the Kalman filter). Minimal hyperparameter optimization embraces an unusual and useful aspect of MINT: there is no need to retrain if one wishes to decode a different behavioral variable. Once the neural state has been estimated, all behavioral variables are readily decodable. In principle, less-extensive optimization could have put MINT at a relative disadvantage. In practice, MINT is robust across reasonable hyperparameter choices (**Fig. S1**);

any disadvantage is thus likely very slight. By necessity and desire, all comparisons were made offline, enabling benchmarked performance across a variety of tasks and decoded variables, where each decoder had access to the exact same data and recording conditions.

There were 15 total ‘behavioral groups’ across the four datasets. For example, MC_Maze yielded two behavioral groups: position and velocity. We computed performance, for each group, by averaging across group members (e.g. across horizontal and vertical velocity). For every behavioral group, across all datasets, MINT’s performance improved upon the ‘interpretable’ methods: the Kalman and Wiener filters. MINT’s performance was typically comparable to, and often better than, that of the expressive GRU and feedforward neural network. MINT yielded the best performance, of all methods, for 11 of 15 behavioral groups. For 3 of the 4 behavioral groups where MINT’s performance was not the highest, the deficit relative to the best method was negligible: $\Delta R^2 < .004$. Thus, in only 1 of 15 cases was MINT at a disadvantage relative to another decoder.

When decoding velocity, MINT provided a small but noticeable improvement over the expressive methods for two datasets (MC_Maze and Area2_Bump), was modestly worse for one dataset (MC_RTT), and was involved in essentially a three-way tie for one dataset (MC_Cycle). For all datasets where position was a behavioral group, MINT and both expressive methods provided extremely similar and very accurate performance, considerably better than that of the Wiener and Kalman filter. MINT tended to display the largest improvements, relative to the next-best decoder, when decoding muscle-related variables (EMG, force, muscle velocity, and muscle length) and phase. This is noteworthy because MINT used the same hyperparameters in all cases and was not re-optimized or re-trained across variables (unlike the other methods). Thus, the same MINT operations that produce a velocity decode naturally produce a position decode and a muscle-length decode.

For each dataset, we computed overall performance by averaging across behavioral groups. MINT had the best overall performance for three of four datasets: MC_Cycle, MC_Maze, and Area2_Bump. For two datasets (MC_Cycle and MC_Maze) MINT considerably outperformed the Kalman and Wiener filter, and had a tiny advantage over the expressive feedforward network and GRU. For the third (Area2_Bump), MINT’s performance was noticeably better than all other methods.

The only dataset where MINT did not perform the best overall was the MC_RTT dataset, where it was outperformed by the feedforward network and GRU. As noted above, this may relate to the need for MINT to learn neural trajectories from training data that lacked repeated trials of the same movement (a design choice one might wish to avoid). Alternatively, the less-structured MC_RTT dataset may strain the capacity to generalize; all methods experienced a drop in velocity-decoding R^2 for this dataset compared to the others. MINT generalizes somewhat differently than other methods, and may have been at a modest disadvantage for this dataset. A strong version of this possibility is that perhaps the perspective in **Figure 1a** [↗](#) is correct, in which case MINT might struggle because it cannot use forms of generalization that are available to other methods (e.g. generalization based on neuron-velocity correlations). This strong version seems unlikely; MINT continued to significantly outperform the Wiener and Kalman filters, which make assumptions aligned with **Figure 1a** [↗](#).

MINT’s generally high level of decoding performance is supported by Bayesian computations: the estimation of data likelihoods and the selection of a decoded state that maximizes likelihood. A natural question is thus whether a simpler Bayesian decoder would have yielded similar results. We explored this possibility by testing a Naïve Bayes regression decoder [\[87\]](#) [↗](#) using the MC_Maze dataset. This decoder performed poorly, especially when decoding velocity ($R^2 = .688$ and $.093$ for hand position and velocity, respectively), indicating that the specific modeling assumptions that differentiate MINT from a naive Bayesian decoder are important drivers of

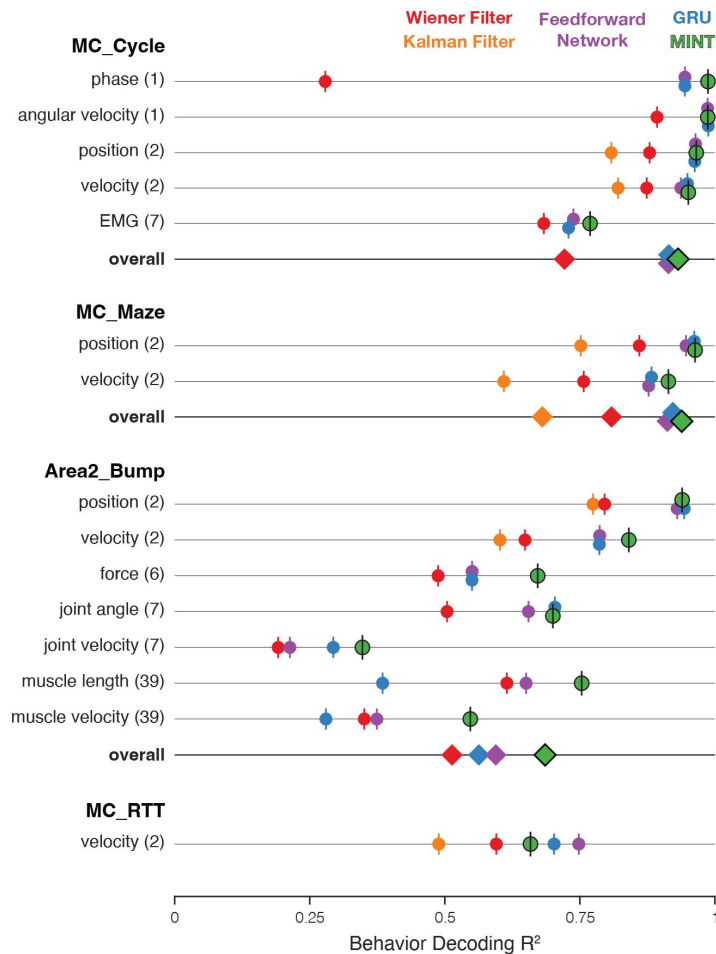


Figure 5.

Comparison of decoding performance, for MINT and four additional algorithms, for four datasets and multiple decoded variables.

On a given trial, MINT decodes all behavioral variables in a unified way based on the same inferred neural state. For non-MINT algorithms, separate decoders were trained for each behavioral group (with the exception of the Kalman filter, which used the same model to decode position and velocity). E.g. separate GRUs were trained to output position and velocity in MC_Maze. Parentheticals indicate the number of behavioral variables within a group. E.g. 'position (2)' has two components: x- and y-position. R^2 is averaged across behavioral variables within a group. 'Overall' plots performance averaged across all behavioral groups. R^2 for feedforward networks and GRUs are additionally averaged across runs for 10 random seeds. The Kalman filter is traditionally utilized for position- and velocity-based decoding and was therefore only used to predict these behavioral groups. Accordingly, the 'overall' category excludes the Kalman filter for datasets in which the Kalman filter did not contribute predictions for every behavioral group. Results are based on the following numbers of training / test trials: MC_Cycle (174 train, 99 test), MC_Maze (1721 train, 574 test), Area2_Bump (272 train, 92 test), MC_RTT (810 train, 268 test). Vertical offsets and vertical ticks are used to increase visibility of data when symbols are close due to similar values.

MINT's performance. We also investigated the possibility that MINT gained its performance advantage simply by having access to trial-averaged neural trajectories during training, while all other methods were trained on single-trial data. This difference arises from the fundamental requirements of the decoder architectures: MINT needs to estimate typical trajectories while other methods don't. Yet it might still be the case that other methods would benefit from including trial-averaged data in the training set, in addition to single-trial data. Alternatively, this might harm performance by creating a mismatch, between training and decoding, in the statistics of decoder inputs. We found that the latter was indeed the case: all non-MINT methods performed better when trained purely on single-trial data.

How challenging is generalization?

The contrasting assumptions in **Figure 1a** and **1b** yield different implications regarding generalization, illustrated in **Figure 6a** and **6b**. In both views, generalization is possible when a newly observed neural trajectory is 'composed' of previously observed elements. However, the nature of those elements is very different. Under the view in **Figure 6a**, those elements are neural states within the manifold. New movements involve novel neural trajectories, but those trajectories are composed of states within the previously observed manifold. Because those neural states are assumed to correlate with behavioral variables (e.g. velocity) those variables can be decoded during novel movements.

While this scenario would be advantageous, empirical observations question its assumptions. The constraint of low tangling implies that it is not possible to traverse previously observed states in entirely new ways. Indeed, even when encouraged to do so under BCI control, monkeys are incapable of producing new neural trajectories that conflict with the flow-field implied by previously seen trajectories [80, 81]. This suggests that newly observed trajectories will fall into two categories. First, new trajectories may be composed of trajectory segments that largely obey the flow-field implied by previously observed trajectories (**Fig. 6b**, green and blue traces). For example, the green trajectory climbs the scaffolding of previously observed single-speed trajectories, mimicking what might occur in a new situation where cycling speed increases steadily [53]. Second, newly observed trajectories may sometimes evolve in previously unobserved portions of state-space (**Fig. 6b**, orange trajectories) where the flow-field has never been estimated. This could occur if the internal computation must change considerably [102], a scenario that may be difficult to anticipate externally.

Under **Figure 6a**, MINT would be at a disadvantage relative to traditional decode methods. In particular, the standard Wiener filter should generalize well because it leverages neuron-behavior correlations and is not limited by any assumptions about a flow-field. In contrast, MINT might often generalize poorly because it cannot leverage neuron-behavior correlations and is constrained by previously observed trajectories. For example, if the green trajectory in **Figure 6a** is in the library, then MINT would be predisposed against decoding the blue trajectory during generalization because the two trajectories approach one another with opposing derivatives. The Wiener filter would have no such limitation.

If the scenario in **Figure 6b** holds, generalization will hinge less on the decoder itself and more on how new neural trajectories (colored traces) relate to the previously observed scaffolding (dashed gray traces). When newly observed trajectories hew close to that scaffolding (e.g. green and blue) any good decode method should generalize well. When newly observed trajectories are far from the scaffolding (e.g. orange), all methods will struggle. MINT will struggle because it has no scaffolding in this region. Other methods will struggle because previously observed neuron-behavior correlations are no longer valid.

We desired a dataset where the predictions of **Figure 6a** and **6b** are likely to be different. With the possible exception of MC_RTT, the datasets above lack this property; generalization to held-out trials is not particularly challenging and could be subserved by the features in **Figure**

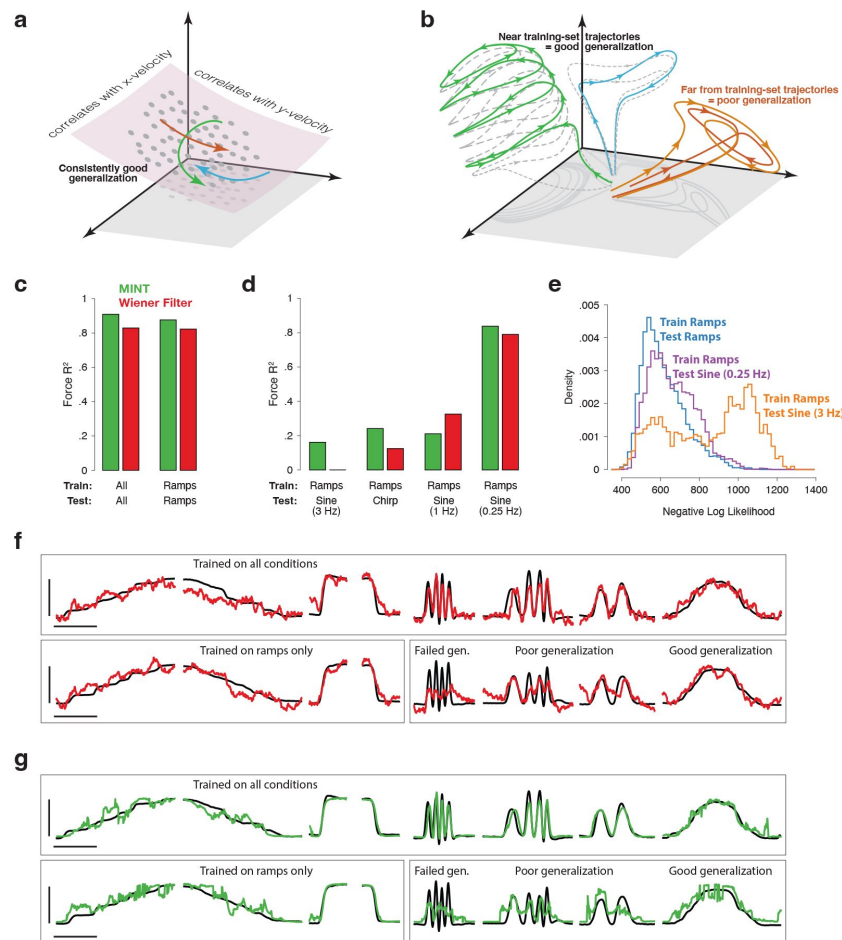


Figure 6.

Implications of neural geometry for decoder generalization, in principle and in practice.

a) Under a traditional perspective, training that explores the full range of to-be-decoded variables will also explore the relevant neural manifold, allowing future generalization. Gray dots indicate neural states observed during training. Colored traces indicate newly observed trajectories. Because these trajectories traverse the previously defined manifold, they can be decoded by leveraging the established correlations between manifold dimensions and kinematics. **b)** Under an emerging perspective, generalization is possible when new trajectories lie within a similar region of factor-space as training-set trajectories. Gray-dashed traces indicate training-set neural trajectories. Colored traces indicate newly observed trajectories. Green and blue trajectories largely follow the flow-field implied by training-set trajectories, allowing generalization. In contrast, the orange trajectory evolves far from any previously explored region, challenging generalization. This may occur even when behavioral outputs overlap with those observed during training. **c)** Basic decoding performance on the MC_PacMan task for MINT (green) and a Wiener filter (red). Test sets used held-out trials from the same conditions as training sets. Force-decoding was excellent both when considering all conditions and when considering only ramps (fast and slow, increasing and decreasing). Results are based on the following numbers of training and test trials: All (234 train, 128 test), Ramps (91 train, 57 test). **d)** Generalization when training employed the four ramp conditions and testing employed a sine or chirp. Results are based on 91 training trials (all cases) and the following numbers of test trials: Ramps (57), 0.25 Hz Sine (19), 1 Hz Sine (12), 3 Hz Sine (15), Chirp (25). **e)** Data likelihoods indicate strained generalization. We computed the log-likelihood of the spiking observations for each neural state decoded by MINT. Rightward values indicate that those observations were unlikely even for the state selected as most likely to have generated those observations. Training employed the four ramp conditions. Distributions are shown when test data employed ramps, a 0.25 Hz sine, or a 3 Hz sine. **f)** Examples of force decoding traces (for held-out trials) when decoding using the Wiener filter. Top row: all conditions were used during training. Bottom row: training employed the four ramp conditions. Horizontal and vertical scale bars correspond to 2 seconds and 16 Newtons, respectively. **g)** As above, but when decoding using MINT. Trials and scale bars are matched across **f** and **g**. Representative trials were selected as those that achieved median decoding performance by MINT within each condition.

6a or by decoding something like the green trajectory in **Figure 6b**. One way to challenge generalization, using existing data, is to restrict training to fewer conditions while ensuring that training still employs a broad range of values for each decoded variable. In the MC_Cycle dataset, the same set of x- and y-velocities occur during forward and backward cycling, just at different phases. If neural activity were dominated by correlations with velocity, traditional methods that leverage those correlations should generalize across cycling directions. Instead, we found that all tested methods failed to generalize when trained on one cycling direction and tested on the other. Most R^2 values were close to zero or even negative. The likely reason is that neural trajectories for forward and backward cycling occupy nearly orthogonal subspaces (**Fig. S2**), especially for this particular dataset [20]. These observations are consistent with the view in **Figure 6b**. Of course, this did not have to be the case; the empirical neural trajectories during cycling could have obeyed the hypothesis in **Figure 6a** and thus supported generalization.

Whether one finds the above test conclusive depends upon which behavioral variables one considers primary. If only velocity (or only position) matters, then generalization should have been easy under **Figure 6a** because the same range of velocities (and positions) was present in training and test sets. Yet one may have different views regarding which variables are central. One may thus suspect that generalization failed because those variables (or their combinations) were insufficiently sampled during training. To address this concern, we explored generalization using a dataset where it is simple to confirm that the relevant behavioral variables are fully explored during training. The MC_PacMan dataset involved Neuropixels-based recordings of motor cortex spiking activity and simultaneous measurements of isometric force applied to an immovable handle. All forces were generated in a forward direction (away from the body) and the arm did not move. Force is thus the central behavioral variable in this task; all other behavioral variables (including muscle activity and cursor height) mirror force or its derivative. Every condition explores a similar range of forces, yet involves very different temporal force profiles. Under the scenario in **Figure 6a**, generalization to new force profiles should be straightforward (for traditional methods) because neural states associated with each force level have already been observed. New force profiles are simply new trajectories through previously observed states. This scenario thus predicts that a simple traditional decoder, such as a Wiener filter, will generalize well but that MINT will not. In contrast, under the scenario in **Figure 6b**, some forms of generalization should be possible (for any method) while others will be intrinsically challenging (for any method).

To test these predictions, we trained both MINT and a Wiener filter to decode force. Training data included fast and slow increasing and decreasing ramps, and thus spanned the full range of forces used in this task. Generalization was tested using sinusoids (0.25, 2, and 3 Hz) and a slow-to-fast chirp. In agreement with **Figure 6b**, generalization performance was strongly condition-specific and only weakly decoder-specific (**Fig. 6d,f,g**). Both MINT and the Wiener filter generalized well to the 0.25 Hz sinusoid: performance was almost as good as when training included all conditions. Yet both decoders generalized poorly at higher frequencies, especially the 3 Hz sinusoid. This was true even though the training data contained swiftly increasing and decreasing forces.

These empirical results make little sense under **Figure 6a**, because training data should have explored the relevant manifold. However, they do make sense under **Figure 6b**. Under this perspective, neural activity is not dominated by correlations with behavioral variables, but is shaped by the flow-field performing the underlying computation. Two conditions can involve different flow-fields, even if their behavioral outputs span similar ranges. Generalization will be strained in such situations, which may be difficult to anticipate ‘from the outside’.

In instances where generalization is strained, and performance is poor overall, different decoders may experience performance drops of different, and perhaps difficult-to-predict, magnitudes. For example, MINT outperformed the Wiener filter when generalizing from ramps to the chirp, while the Wiener filter outperformed MINT when generalizing from ramps to the 1 Hz sine. This

observation does not suggest a meaningful advantage for either decoder. These advantages were idiosyncratic and, more importantly, overall performance was so poor that either decoder would have been essentially unusable without retraining that includes additional conditions. The MC_RTT task may also be an example of this effect: the two expressive methods outperformed MINT, but all decoders achieved their worst velocity-decode performance for this dataset (Fig. 5 [bottom](#)).

For these reasons, one may wish to know, during decoding, if generalization is becoming strained. Usefully, MINT computes the log-likelihood of spiking observations for its estimated neural state (which is the most-probable state it can find). A prediction of Figure 6b [is](#) that, when generalization is strained, these likelihoods should become unusually low because the true neural state (which generated the spiking observations) will often be far from any state that MINT can estimate. To test this prediction, we computed the distribution of negative log-likelihoods (Fig. 6e [is](#)) for all estimated neural states during the test of generalization described above. The distribution for 0.25 Hz sinusoids (purple) was similar to that for ramps (blue). This is consistent with the set of true neural states, during the 0.25 Hz sinusoids, being similar to states that are ‘reachable’ by MINT’s interpolation step. In contrast, the distribution for 3 Hz sinusoids (orange) showed a second mode of spiking observations that were particularly unlikely. This is expected if the true neural state, at these moments, is far from any state that can be reached by MINT’s interpolation (as would be true for the orange trace in Figure 6b [is](#)). Data log-likelihoods could thus potentially be used to indicate – without having to ask the participant – when decoding is strained because neural states are far from those observed during training. Training could then be modified to include additional conditions.

Modeling & preprocessing choices

MINT uses a direct mapping from neural states to behavioral states to instantiate highly nonlinear decoding, then uses locally linear interpolation to support decoding of behavioral states not observed during training. To determine the impact of these choices on performance, we ran a full factorial analysis that always used neural state estimates generated by MINT, but compared the direct neural-to-behavioral-state mapping to a linear neural-to-behavioral-state mapping (Fig. S3a [is](#)), and assessed the impact of interpolation (Fig. S3b [is](#)). Both choices did indeed improve performance. We also assessed the impact of acausal versus causal decoding (Fig. S3c [is](#)). Causal decoding is important for real-time BCIs, and is thus employed for all our main decoding analyses. However, future applications could potentially introduce a small lag into online decoding, effectively allowing some ‘acausal decoding’ at the expense of longer-latency responses. As expected, acausal decoding provided modest performance increases, raising the possibility of an interesting tradeoff. Decoding can be zero-lag (causal), positive-lag (acausal), or even negative-lag. The latter would result in a very ‘snappy’ decoder that anticipated actions at the expense of some decoding accuracy. The lag hyperparameter wouldn’t need to be specified prior to training and could be freely adjusted at any time.

All datasets were curated to contain sorted spikes. Yet during online performance, decoding must typically rely on unsorted threshold crossings. Prior studies have found that moving from sorted spikes to threshold crossings produces a modest but tolerable drop in performance [103 [is](#)]. We similarly found that moving from sorted spikes to threshold crossings (selected to be from ‘good’ channels with reasonable signal-to-noise) produced only a small drop in performance (Fig. S3d [is](#)). The operations of MINT highlight why such robustness is expected. When spikes from two neurons are combined, the resulting ‘unit’ acts statistically as if it were a high firing-rate neuron. Spiking is still approximately Poisson, because the sum of Poisson processes is a Poisson process. Thus, MINT’s statistical assumptions remain appropriate.

Neural state estimation

MINT's success at behavioral decoding suggests underlying success in neural-state estimation. However, that would not necessarily have to be true. For example, the GRU also typically supported excellent behavioral decodes but doesn't provide neural state estimation. MINT does (by its very construction) but whether those estimates are accurate remains to be evaluated.

To evaluate neural-state estimates, we submitted firing-rate predictions to the Neural Latents Benchmark [104], an online benchmark for latent variable models that compares acausal neural-state estimates across methods and datasets. The four primary datasets for the benchmark are MC_Maze, Area2_Bump, MC_RTT, and an additional dataset, DMFC_RSG, that contains neural recordings from dorsomedial frontal cortex while a monkey performed a cognitive timing task [62]. Three secondary datasets (MC_Maze-L, MC_Maze-M, and MC_Maze-S) contain additional sessions of the maze task with progressively fewer training trials (500, 250, 100, respectively).

By submitting to the Neural Latents Benchmark, we can ask how MINT performs relative to a set of 'baseline' methods. One baseline method (smoothed spikes) is traditional and simple. Other baseline methods, Gaussian Process Factor Analysis (GPFA) [105] and Switching Linear Dynamical System (SLDS) [106], are modern yet well-established. Finally, two baseline methods are quite new and cutting-edge: Neural Data Transformers (NDT) [92] and AutoLFADS [107]. These are expected to provide a high bar by which to judge the performance of other approaches.

Although the neural state is a well-defined quantity in spiking models [55], experimental data does not typically provide a ground-truth neural state against which estimates can be compared (unlike for behavioral-state decoding). Yet one can define attributes that an accurate neural-state decode should possess. A critical attribute is prediction of spiking activity. This attribute is assessed via bits per spike, which is closely related to the log-likelihood of spiking observations given a set of firing rates (assuming Poisson spiking). MINT performed similarly to AutoLFADS and was consistently slightly better than NDT on the four primary datasets (Fig. 7a, left). MINT performed slightly better than AutoLFADS and NDT on the largest of the secondary datasets. That advantage grew as dataset-size shrank (Fig. 7a, right). MINT outperformed the other three baseline methods by modest-to-sizeable margins for all seven datasets.

MINT was not designed to maximize bits per spike – good performance in this regard simply results from the primary goal of estimating a neural state containing firing rates. This matters because there exist spiking features – e.g. synchrony beyond that due to correlated rates – that some methods may be able to leverage when predicting spike probabilities. Doing so is reasonable but separate from MINT's goal of estimating a rate-based neural state. It is thus worth considering other attributes expected of good neural-state estimates. The Neural Latents Benchmark includes two such attributes. First, if state estimation is accurate, trial-averaged neural state estimates should resemble empirical trial-averaged firing rates (computed for test trials), an attribute assessed by PSTH R^2 . MINT achieved higher PSTH R^2 values than every other method on every dataset (Fig. 7b). This is in some ways unsurprising: MINT estimates neural states that tend to resemble (at least locally) trajectories 'built' from training-set-derived rates, which presumably resemble test-set rates. Yet strong performance is not a trivial consequence of MINT's design. MINT does not 'select' whole library trajectories; PSTH R^2 will be high only if condition (c), index (k), and the interpolation parameter (α) are accurately estimated for most moments.

A second attribute is that, in sensorimotor areas during motor tasks, behavior should be decodable from the neural state. The benchmark thus measured R^2 for a linear decode of velocity from the neural state. MINT outperformed all baseline approaches (Fig. 7c) for all datasets. (MINT could of course have achieved even better decoding using its typical direct association, but that would have undermined the utility of this comparison).

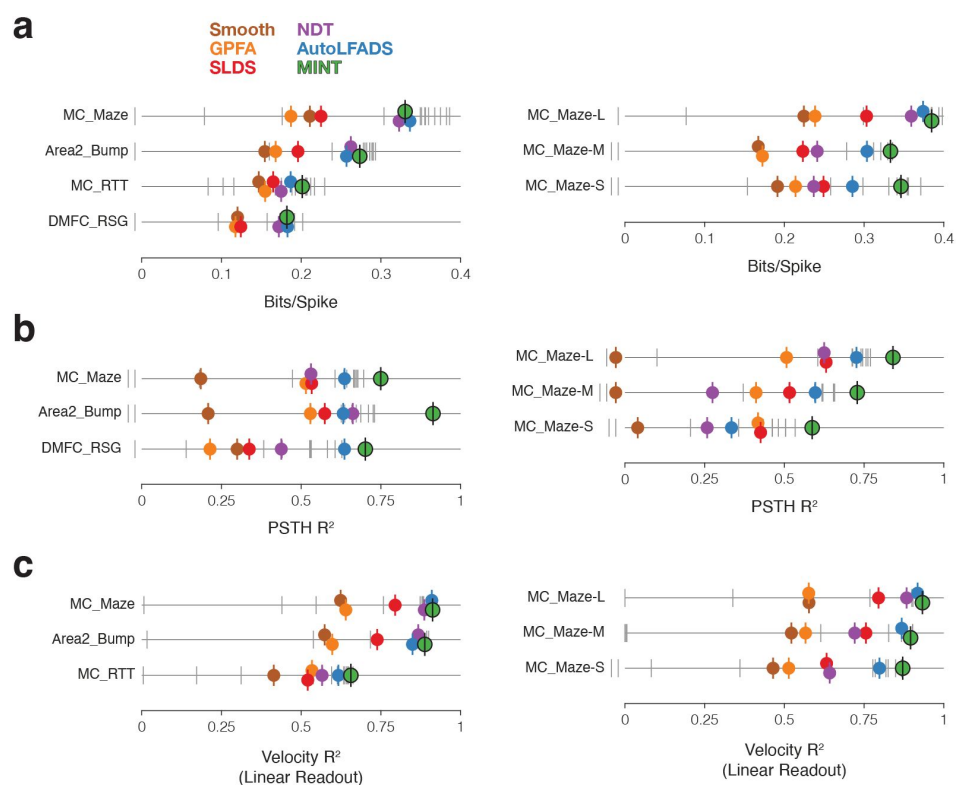


Figure 7.

Evaluation of neural state estimates for seven datasets.

a) Performance quantified using bits per spike. The benchmark's baseline methods have colored markers. All other submissions have gray markers. Vertical offsets and vertical ticks are used to increase visibility of data when symbols are close due to similar values. Results with negative values are designated by markers to the left of zero, but their locations don't reflect the magnitude of the negative values. Data from Neural Latents Benchmark (<https://neurallatents.github.io/>). **b)** Performance quantified using PSTH R^2 . **c)** Performance quantified using velocity R^2 , after velocity was decoded linearly from the neural state estimate. A linear decode is used simply as a way of evaluating the quality of neural state estimates, especially in dimensions relevant to behavior.

In addition to the five baseline methods, the Neural Latents Benchmark solicited submissions from other methods and many were submitted. These submissions spanned a broad range in terms of performance (gray vertical lines in [Fig. 7](#)), highlighting the challenging nature of this problem. Some submissions involved ensemble approaches, some of which outperformed both MINT and the baseline methods in terms of bits per spike. For the primary datasets, the highest bits per spike overall was achieved by an ensemble of SpatioTemporal Neural Data Transformers [108], achieving .386 bits per spike on the MC_Maze dataset, compared to .330 for MINT. For the smaller secondary datasets, the very best submissions provided minimal-to-modest improvement over MINT in terms of bits per spike ([Fig. 7a](#), right), and the majority performed less well.

The ability of some expressive neural networks (and some ensembles of these networks) to improve the bits per spike metric indicates there exists variability in the neural recordings that MINT does not capture. That variability could arise from behaviorally irrelevant fluctuations in the neural state, from synchrony based effects, or from instabilities in the neural recordings. These are ‘real’ things and it is thus valid for a method to leverage them to predict spiking when that is the goal. Yet as noted above, they may often be incidental to one’s goals in estimating the neural state, highlighting the need to consider additional attributes such as PSTH R^2 and velocity R^2 . For PSTH R^2 ([Fig. 7b](#)), MINT outperformed all other methods for all six datasets. For velocity R^2 ([Fig. 7c](#)), MINT had the best performance for three out of six datasets and was very close to the highest R^2 values for the remaining three datasets (the largest deficit was $\Delta R^2 = .012$).

In summary, MINT’s strong decode performance ([Fig. 5](#)) does indeed follow from good neural-state estimation, as one would hope. MINT provided estimates that were competitive with the best present methods. This includes existing well-documented methods, as well as a bevy of additional submissions. This implies that MINT could be used in situations where neural-state estimation is the goal, possibly including monitoring of the neural state for clinical purposes. MINT uses simple operations and is readily implemented causally – indeed this is typically how we would expect to see it used. Yet MINT’s performance is typically similar to – and often better than – that of complex methods that would likely have to be limited to offline situations.

Practical considerations

Training & execution times

For MINT, ‘training’ simply means computation of standard quantities (e.g. firing rates) rather than parameter optimization. MINT is thus typically very fast to train ([Table 1](#)), on the order of seconds using generic hardware (no GPUs). This speed reflects the simple operations involved in constructing the library of neural-state trajectories: filtering of spikes and averaging across trials. At the same time we stress that MINT is a method for leveraging a trajectory library, not a method for constructing it. One may sometimes wish to use alternatives to trial-averaging, either of necessity or because they improve trajectory estimates. For example, for the MC_RTT task we used AutoLFADS to infer the library. Training was consequently much slower (hours rather than seconds) because of the time taken to estimate rates. Training time could be reduced back to seconds using a different approach – grouping into pseudo-conditions and averaging – but performance was reduced. Thus, training will typically be very fast, but one may choose time-consuming methods when appropriate.

Execution times were well below the threshold for real-time applications ([Table 1](#)), even using a laptop computer. State estimation in MINT is $\mathcal{O}(NT_{total})$, where T_{total} is the total number of times (across conditions) in the library of neural trajectories. This is $\mathcal{O}(NC)$ when making the simplifying assumption that all conditions are of equal length. Memory requirements are similarly $\mathcal{O}(N + MC)$, where M is the number of behavioral variables (a value that will typically remain quite small). Thus, both execution times and memory requirements grow only linearly with neurons or

Data Set	Training Time (s)	Execution Time (ms)
Area2_Bump	4.8 ± 0.2	0.31 ± 0.03
MC_Cycle	20.3 ± 2.9	0.55 ± 0.03
MC_Maze	42.8 ± 0.6	2.26 ± 0.09
MC_Maze-L	10.8 ± 0.4	0.83 ± 0.02
MC_Maze-M	8.0 ± 0.4	0.80 ± 0.04
MC_Maze-S	5.8 ± 0.2	0.59 ± 0.07
MC_RTT	~ 3.2 hours	6.99 ± 0.28

Table 1.

MINT training times and average execution times (average time it took to decode a 20 ms bin of spiking observations).

To be appropriate for real-time applications, execution times will ideally be shorter than the bin width. Note that the 20 ms bin width does not prevent MINT from decoding every millisecond; MINT updates the inferred neural state and associated behavioral state between bins, using neural and behavioral trajectories that are sampled every millisecond. For all table entries (except MC_RTT training time), means and standard deviations were computed across 10 train/test runs for each dataset. Training times exclude loading datasets into memory and any hyperparameter optimization. Timing measurements taken on a Macbook Pro (on CPU) with 32GB RAM and a 2.3 GHz 8-Core Intel Core i9 processor. Training and execution code used for timing measurements was written in MATLAB (with the core recursion implemented as a MEX file). For MC_RTT, training involved running AutoLFADS twice (and averaging the resulting rates) to generate neural trajectories. This training procedure utilized 10 GPUs and took ~ 1.6 hours per run. For AutoLFADS, hyperparameter optimization and model fitting procedures are intertwined. Thus, the training time reported includes time spent optimizing hyperparameters.

with conditions. Additionally, much of the estimation procedure in MINT is condition-specific and therefore parallelizable. This makes MINT a good candidate for decoding in BCI applications that involve large neuron counts and/or large behavioral repertoires.

MINT performs well on small training sets

To investigate robustness to training with small trial-counts, we leveraged the fact that the four maze-reaching datasets – one primary (MC_Maze) and three secondary (MC_Maze-L, MC_Maze-M, and MC_Maze-S) – contained training sets of decreasing size. Because training involved computing trial-averaged rates, performance is expected to decline when fewer trials are available for averaging. For comparison, we chose the GRU and the feedforward neural network because of their generally good position and velocity decoding performance in the MC_Maze task. We assessed performance in terms of position and velocity R^2 (Fig. 8a). MINT always performed at least as well as the two expressive methods. This gap was often (though not always) larger for smaller training sets. More broadly, MINT was quite robust to small training sets: an approximately five-fold reduction in training data led to a decline in R^2 of only .037 (position) and .050 (velocity).

Performance when neurons are lost

It is desirable for a decoder to be robust to the unexpected loss of the ability to detect spikes from some neurons. Such loss might occur while decoding, without being immediately detected. Additionally, one desires robustness to a known loss of neurons / recording channels. For example, there may have been channels that were active one morning but are no longer active that afternoon. At least in principle, MINT makes it very easy to handle this second situation: there is no need to retrain the decoder, one simply ignores the lost neurons when computing likelihoods. This is in contrast to nearly all other methods, which require retraining because the loss of one neuron alters the optimal parameters associated with every other neuron.

To evaluate, we performed a neuron-dropping analysis using the MC_Maze-L dataset (Fig. 8b). For comparison, we also evaluated a Wiener filter. The Wiener filter provides a useful comparison because (like MINT) it is easy to retrain quickly, with reliable results, across many training sets (which was critical in the present situation). To simulate undetected neuron loss, we chose a random subset of neurons and removed all their spikes, without making any adjustment to the decoders. This simulates a neuron (or channel) falling silent. This was repeated many times with different random draws. To simulate known neuron loss, we did the same but allowed the decoders to adjust: by ignoring the lost neurons in the case of MINT, and by retraining all parameters in the case of the Wiener filter. We evaluated both position (Fig. 8b, top) and velocity (bottom) decoding.

For MINT, undetected loss of fewer than 20 (of 162) neurons caused essentially no performance deficit. Average position and velocity decoding R^2 remained above 95% of peak performance until fewer than 113 and 109 neurons remained, respectively. Large undetected losses – e.g. of half the neurons – caused major deficits in decoding. Yet decoding remained accurate if those losses were known, with the lost neurons ignored during decoding. MINT was extremely robust in this situation: MINT's average R^2 for position decoding remained above 95% of peak performance until fewer than 58 neurons (36%) remained. Results were similar for velocity decoding: performance was >95% of its peak until fewer than 62 neurons (38%) remained. The Wiener filter was also fairly robust, but less so than MINT. To remain above 95% of peak performance, the Wiener filter needed approximately twice as many remaining neurons: 118 for position and 126 for velocity.

These results illustrate that MINT is reasonably robust to undetected neuron loss. In practice, the loss of many neurons would likely be detected. In such cases, MINT is very robust. Additionally, because there is no need for any true retraining, any adjustment to known neuron-loss can occur

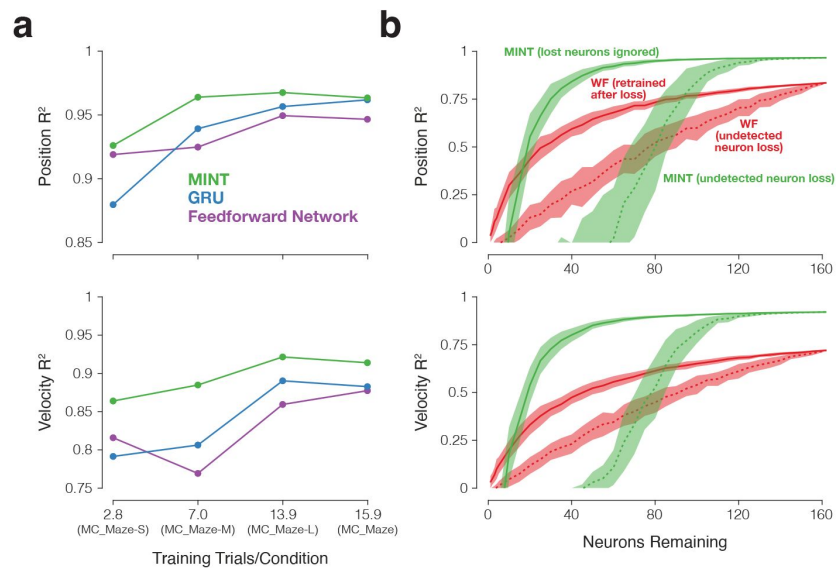


Figure 8.

Decoding robustness in the face of small training sets and reduced neuron counts.

a) R_2 values for MINT and two neural network decoders (GRU and the feedforward network) when decoding position and velocity from four maze datasets with progressively fewer training trials per condition. Results are based on the following numbers of training and testing trials: MC_Maze-S (75 train, 25 test), MC_Maze-M (188 train, 62 test), MC_Maze-L (375 train, 125 test), MC_Maze (1721 train, 574 test). MC_Maze contains 108 conditions and the other maze datasets each contain 27 conditions. **b)** Decoding performance in the face of undetected (dashed traces) and known (solid traces) loss of neurons. R_2 values are shown for MINT (green) and the Wiener filter (red) when decoding position and velocity from the MC_Maze-L dataset (same train/test trials as in **a**). To simulate undetected neuron loss, we chose k neurons (randomly, with k being the number of lost neurons out of 162 total) and eliminated their spikes without adjusting the decoders. This procedure was repeated 50 times for each value of k . Traces show the mean and standard deviation of the sampling distribution (equivalent to the standard error). To simulate known neuron loss, we followed the above procedure but altered the decoder to account for the known loss. For MINT, this simply involved ignoring the lost neurons when computing data likelihoods. In contrast, the Wiener filter had to be retrained using training data with those neurons eliminated (as would be true for most methods). This was done separately for every set of lost neurons.

on the fly, with no need to pause decoding.

Interpretability

One form of interpretability relates to whether an algorithm uses ‘black box’ optimization to learn a solution. For example, neural networks can be highly expressive and learn constraints directly from data. However, when constraints are known, one may prefer to rely more on explicit assumptions and less on expressivity. Interpretable methods, such as the Kalman filter, rely on explicit assumptions. Such assumptions may (if they match the data) improve performance. They may also provide other potential advantages. For example, the Kalman filter is a canonical ‘interpretable’ algorithm, with a probabilistic graphical model that makes an explicit Gaussian assumption regarding distributions of spike count measurements. This explicitness may be helpful in understanding successes and failures (e.g. the Gaussian assumption can break down during low-rate periods).

The performance of interpretable methods should tend to reflect the extent to which their assumptions match the statistics of the data. To investigate, we considered the performance of both the Kalman filter and MINT on the MC_Maze reaching dataset. The Kalman filter produces a relatively poor decode ($R^2 = .609$ for hand velocity, versus .914 for MINT). This doesn’t imply the Kalman filter is a poor algorithm, simply that its assumptions may not be well-aligned with the data. We confirmed this by generating synthetic data whose properties are closer to those assumed by the Kalman filter. The Kalman filter then performed quite competitively (**Fig. S4**). This result underscores that MINT does not outperform the Kalman filter because it is intrinsically a better method. Instead, MINT performs well because its assumptions are presumably a better match to the underlying properties of the data. If **Figure 1a** were the correct perspective, the Kalman filter would have been expected to perform at least as well as MINT. Under **Figure 1b** the Kalman filter is predicted to perform less well, which is indeed what we found.

MINT is also interpretable regarding why each prediction is made. Suppose that, at a given moment during decoding, MINT selects neural state $\mathbf{x}$ as the most likely state and renders its decode using the associated behavioral state $\mathbf{z}$. As analyzed in **Figure 6e** above, selection of $\mathbf{x}$ comes with an associated likelihood that conveys information regarding how well the spiking data matched that neural state. Likelihoods are known not only for this ‘selected’ state, but for all non-selected states on all trajectories in the library. These likelihoods can be analyzed to understand why the selected state was chosen over the others, how close the algorithm was to selecting a different state, and which neurons contributed most strongly to state selection. More generally, because of its probability-based approach, MINT is readily modifiable and extensible in a variety of ways (**Fig. S5**).

Discussion

Implications regarding neural geometry

Our results argue that the empirical structure of motor-cortex activity is closer to the hypothesis in **Figure 1b**. We began with the prediction that, under that hypothesis, MINT should consistently outperform existing interpretable methods such as the Kalman and Wiener filter. This was indeed the case: MINT always performed better than these two methods, often by sizable margins. In contrast, MINT and the Kalman filter performed comparably on simulated data that better approximated the assumptions in **Figure 1a**. Thus, MINT is not a ‘better’ algorithm – simply better aligned with the empirical properties of motor cortex data. This highlights an important caveat. Although MINT performs well when decoding from motor areas, its assumptions may be a poor match in other areas (e.g. the hippocampus). MINT performed well on two non-motor-cortex

datasets – Area2_Bump (S1) and DMFC_RSG (dorsomedial frontal cortex) – yet there will presumably be other brain areas and/or contexts where one would prefer a different method that makes assumptions appropriate for that area.

We also began with the prediction that, under the hypothesis in **Figure 1b**, MINT should be competitive with highly expressive approaches that can learn their constraints from data. MINT may even perform better in some cases because MINT can begin with appropriate constraints rather than having to learn them from training data. At the same time, MINT might sometimes be at a disadvantage if there are additional constraints or features that it does not take into account (a simple example would be drift in firing rates over time – one would desire a drift stabilization module to correct for this, as in **Fig. S5**). Empirically, decoding performance was often exceedingly similar for MINT and both expressive methods (the feedforward network and GRU). Consider velocity, which is probably the most commonly decoded parameter for motor BCIs. Across the four datasets in **Figure 5**, the average R^2 for the velocity decode was .841 (MINT), .830 (GRU), and .837 (feedforward network). These values are all considerably higher than for the Kalman filter (.630) and Wiener filter (.719). There were multiple cases where MINT performed noticeably better than either expressive method, and one case (the MC_RTT dataset) where MINT performed noticeably worse. These results are consistent with the set of expectations outlined above: MINT should be highly competitive with expressive methods, often performing modestly better and sometimes performing modestly worse.

The hypothesis that MINT's assumptions are well-aligned with the data is further supported by its performance during neural-state estimation. Of the well-established baseline methods for state-estimation, the highest performing were AutoLFADS and a Neural Data Transformer (NDT). Across a variety of datasets and performance metrics, MINT either performed similarly to these methods or performed noticeably better. MINT consistently outperformed simple lightweight state estimators, despite being a simple lightweight estimator itself. We find it noteworthy that MINT naturally performs very well as a neural-state estimator, despite being designed as a decoder of behavior. MINT does so with no need for modification. This relates to the observation that MINT can readily switch from decoding one parameter (e.g. cycling phase) to a very different one (muscle activity) with no need for additional training, an option not available with any other standard method.

Practical considerations for use of MINT as an online decoder

Across the four datasets we tested, the best decoders (MINT and the two expressive methods) were all very close in terms of performance. Given this, decoder-choice would likely be based on considerations beyond raw performance. Here MINT has many advantages. Because MINT's computations are so simple, it is typically fast to train and always very fast to run. Its interpretability supports principled extensions (discussed below) and the ability to investigate the source of any failures (as in **Figure 6e**). These considerations, combined with the fact that MINT often yielded the best behavioral decode of all methods, makes MINT a compelling candidate for online decoding in scientific or clinical applications. With that goal in mind, there exist three important practical considerations. First, some decode algorithms experience a performance drop when used online. One presumed reason is that, when decoding is imperfect, the participant alters their strategy which in turn alters the neural responses upon which decoding is based. Because MINT produces particularly accurate decoding, this effect may be minimized, but this cannot be known in advance. If a performance drop does indeed occur, one could adapt the known solution of retraining using data collected during online decoding [13]. Another presumed reason (for a gap between offline and online decoding) is that offline decoders can overfit the temporal structure in training data [109]. This concern is somewhat mitigated by MINT's use of a short spike-count history, but MINT may nevertheless benefit from data augmentation strategies such as including time-dilated versions of learned trajectories in the libraries.

A second practical consideration is that, for a participant with limited ability to move, training would involve asking subjects to observe and/or attempt movements. Both strategies engage neurons in motor cortex [110], but can also create uncertainty regarding the exact behavioral label for each moment (a challenge any decode method must face). MINT's decoding performance benefits from good estimates of Ω and Φ , and will decline if those estimates are poor or misaligned. Fortunately, methods exist for temporally aligning neural data based on spiking activity alone [111] and the performance of such methods should only improve as more neurons are simultaneously recorded. Such methods would allow computation of trial-averaged rates (or single-trial rates) that could then be aligned with a canonical movement trajectory (e.g. the usual bell-shaped velocity profile). Conveniently, any alterations to that alignment are trivial to accomplish, including at runtime. For example, to shift velocity earlier, one would simply shift MINT's library indices. More generally, MINT is amenable to a variety of methods – simple or sophisticated – for inferring the trajectory library. Possible future approaches include leveraging the typical neural geometry for a particular task. This could allow a participant-specific library to be learned using fewer trials. This might even be done in an unsupervised way (similar to Dyer et al. [112]): a trajectory library could be built without behavioral labels, then appropriately linked to a 'reference' behavioral library via prior knowledge of task-specific neural geometry. As a simple example, cycling across different speeds produces a tube-shaped manifold with different velocities at each end [53]. An empirically observed tube might be given behavioral labels based on this knowledge.

A final practical consideration concerns generalization. Because MINT uses a library of canonical trajectories, it is tempting to suppose that its decoded trajectories will simply be replicas of a limited set of stereotyped behaviors. By removing the interpolation step, MINT can indeed approximate this mode, which would be desirable in situations where the to-be-decoded behavior is itself stereotyped (e.g. producing keystrokes or handwritten characters). Yet one often wishes to decode behaviors beyond those observed in the training set. MINT can do so if the empirical neural trajectory mostly (and locally) obeys the flow-field implied by the library. This provides considerable, but not unlimited, flexibility. Might MINT thus be at a disadvantage, relative to other methods, with respect to generalization? This is predicted to be true under the hypothesis in **Figure 1a** and **6a**. However, it is not expected to be true under the hypothesis in **Figure 1b** and **6b**. Under that hypothesis, the forms of generalization that are unavailable to MINT will also be unavailable to standard methods. Consistent with that prediction, generalization success for the MC_PacMan dataset was condition-specific, not decoder-specific.

Multi-task decoding

The fact that generalization may often be challenging does not imply that decoding across multiple tasks is a challenge. Indeed MINT is ideally suited for this situation. For example, MINT successfully decoded motor output for both tasks performed by the network in **Figure 4a**, and across all conditions in the MC_PacMan dataset (when trained on all conditions, **Fig. 6c**). As long as trajectories relevant to both tasks are present in the library, 'task switching' occurs naturally and implicitly. Because MINT does not rely on correlations between neural activity and behavioral variables, task switching will not be impaired if those correlations change across tasks, something anticipated under **Figure 1b** and **6b**. As the number of situations decoders are expected to handle increases, this may be a useful property.

Future avenues

MINT's probability-based framework is amenable to principled extensions. Rather than simply decode the state with the highest data likelihood, one could incorporate priors. Priors could be fixed or guided by incoming information (e.g. an eye tracker). Decoding could also focus on utility functions rather than probabilities. This would respect the fact that a participant may value avoiding certain types of errors more than others. For example, erroneous transitions from stationary to moving may be jarring or unsafe, and could be dissuaded via an appropriately

designed utility function. These extensions are natural because MINT computes data likelihoods for a broad range of candidate states before choosing the maximum-likelihood state. Converting likelihoods to probabilities (based on a prior) or to an estimated cost can thus be done in a principled manner.

Acknowledgements

We thank Felix Pei, Chethan Pandarinath, and the rest of the Neural Latents Benchmark team for curating many of the datasets used in this paper and releasing them publicly. We thank those who originally collected those datasets: Rameed Chowdhury (Area2_Bump), Hansem Sohn (DMFC_RSG), Matt Kaufman and Mark Churchland (MC_Maze, MC_Maze-L,M,S), and Joseph O'Doherty (MC_RTT). We additionally thank Felix and Chethan for providing AutoLFADS rates for the MC_RTT dataset. We thank Karen Schroeder for collecting and preprocessing the MC_Cycle dataset in collaboration with the first and last authors. We thank Najja Marshall, Eric Trautmann and Andrew Zimnik for their central roles in designing the PacMan task and collecting data (with Elom Amematsro) using Neuropixels probes. We thank Yana Pavlova for expert animal care while collecting the MC_Cycle and MC_PacMan datasets. We thank Brian DePasquale for providing code for simulating the artificial multitask spiking network. We thank Andrew Zimnik and Eric Trautmann for helping the second author present this work at NCM 2022 when the first author had COVID. This work was supported by the Simons Foundation and the Grossman Center for the Statistics of Mind.

Additional information

Author contributions

SP and MC designed the decoder. JC provided input on how to conceptualize and formalize the decoder. SP wrote the software. SP and MC designed data analyses. SP performed analyses. QW provided guidance regarding analyses. EA supplied processed data for the MC_PacMan dataset and analysis advice. SP and MC wrote the paper. All authors contributed to editing.

Competing interests

SP, JC, and MC hold a patent pertaining to this work. The patent has been licensed to Blackrock Neurotech. The authors declare no additional competing interests.

Code availability

- MINT is available at <https://github.com/seanmperkins/mint> (MATLAB).
- Other decoders are available at <https://github.com/seanmperkins/bci-decoders> (Python).
- Multitask spiking network was simulated with code from DePasquale et al. [55] that is available at <https://github.com/briandepasquale/factor-based-spiking-nets> (MATLAB).

Data availability

All empirical datasets (except MC_Cycle and MC_PacMan) were curated by the Neural Latents Benchmark team [104] and made publicly available on DANDI (linked below). The MC_Cycle dataset is available upon request. Useful functions for loading the DANDI datasets are available at

https://github.com/neurallatents/nlb_tools .

- Area2_Bump: <https://dandiarchive.org/dandiset/000127> 
- DMFC_RSG: <https://dandiarchive.org/dandiset/000130> 
- MC_Maze: <https://dandiarchive.org/dandiset/000128> 
- MC_Maze-L: <https://dandiarchive.org/dandiset/000138> 
- MC_Maze-M: <https://dandiarchive.org/dandiset/000139> 
- MC_Maze-S: <https://dandiarchive.org/dandiset/000140> 
- MC_RTT: <https://dandiarchive.org/dandiset/000129> 

Methods

MINT (Part I): Estimating candidate states

Model of idealized trajectories

Let $\mathbf{s}_t \in \mathcal{S}^N$ where $\mathcal{S} = \{0, 1, \dots, S\}$ be the measured spiking activity from N neurons at time sample t . Spikes are associated with some underlying neural state $\mathbf{x}_t \in \mathbb{R}^N$ and a corresponding behavioral state $\mathbf{z}_t \in \mathbb{R}^M$ where M is the number of behavioral variables (e.g. x- and y-velocity of the hand). We bin spiking activity in time such that $\bar{\mathbf{s}}_{i'} = \sum_{i=0}^{\Delta-1} \mathbf{s}_{\Delta i' - i}$. Given recent spiking history $\bar{\mathbf{s}}_{t'-\tau':t'}$ (spike counts from the most recently completed bin $t' = \lfloor t/\Delta \rfloor$ and τ' previous bins), we are interested in inferring posterior distributions over neural states $P(\mathbf{x}_t | \bar{\mathbf{s}}_{t'-\tau':t'})$ and over behavioral states $P(\mathbf{z}_t | \bar{\mathbf{s}}_{t'-\tau':t'})$. We can then perform *maximum a posteriori* estimation to generate candidate estimates for both the neural state and behavior. We introduce a model of the form:

$$\mathbf{x}_t \sim \text{Unif}(\Omega) \quad (1)$$

$$\mathbf{x}_{t-i-1} = g(\mathbf{x}_{t-i}), \quad i \in \{0, 1, \dots, \tau-1\} \quad (2)$$

$$\bar{\mathbf{s}}_{t'-i'} \sim \text{Pois} \left(\frac{1}{\Delta} \sum_{i=0}^{\Delta-1} \mathbf{x}_{\Delta(t'-i')-i} \right), \quad i' \in \{0, 1, \dots, \tau'\} \quad (3)$$

$$\mathbf{z}_t = f(\mathbf{x}_t) \quad (4)$$

where $g: \Omega^+ \rightarrow \Omega^+$ is a state-transition lookup table, $f: \Omega \rightarrow \Phi$ is a lookup table that associates neural states with behavioral states, $T = t - \Delta t'$ is the number of samples elapsed since the last time bin completed, and $\tau = T + \Delta(\tau' + 1) - 1$ indicates the extent of the deterministic history for $\mathbf{x}_t$ (matching the extent of recent spiking history). This model has parameters Ω^+ and Φ . These parameters are described below, along with definitions for g and f .

Ω^+ is a set (library) of C neural trajectories, with each trajectory consisting of an ordered set of neural states. Each neural state in the library is notated as $\bar{\mathbf{x}}_k^c$, where c indexes trajectories and k indexes locations along each trajectory (e.g. $c = 2$ and $k = 100$ indicates the 100th neural state along the second trajectory in the library). Each trajectory is presumed to contain the sequence of neural states that are traversed for a particular behavior. Thus, k increasing along a neural trajectory corresponds to the passage of time during the execution of a behavior. Eq. (2)  and Eq. (3)  indicate that each $\mathbf{x}_t$ has a history of τ states that are responsible for generating recent spike count observations. However, the first τ states along each neural trajectory lack a complete state history.

Thus, we learn the full Ω^+ , but in Eq. (1) we only consider a reduced set of neural states, $\Omega = \{\bar{\mathbf{x}}_k^c \in \Omega^+ \mid k > \tau\}$, because only these states have the necessary state histories. The state-transition lookup table is defined as follows,

$$g(\bar{\mathbf{x}}_k^c) = \bar{\mathbf{x}}_{k-1}^c \quad (5)$$

simply indicating that each neural state in the library has a recent history determined by the stereotyped ordering of neural states within each neural trajectory.

Φ is a set (library) of C behavioral trajectories, where each trajectory consists of an ordered set of behavioral states. Each $\bar{\mathbf{x}}_k^c \in \Omega$ is associated with a behavioral state, $\bar{\mathbf{z}}_k^c$, via

$$f(\bar{\mathbf{x}}_k^c) = \bar{\mathbf{z}}_k^c \quad (6)$$

In other words, each neural state in the library is paired with a behavioral state for the same c and k . Once the libraries of neural and behavioral trajectories have been learned, MINT's parameters are fully learned. g will be determined by the ordering of states within each trajectory and f will be determined by each state's c and k indices. There are a variety of methods one can employ for learning these trajectories. The methods used in this paper are described in Learning idealized trajectories. These trajectories can often be learned by averaging neural and behavioral data across many repeated trials of the same movement, yielding $\bar{\mathbf{x}}_k^c$ and $\bar{\mathbf{z}}_k^c$ that are highly representative of the neural and behavioral states expected for a particular movement.

We assume $\bar{\mathbf{z}}_{k_1}^{c_1} = \bar{\mathbf{z}}_{k_2}^{c_2}$ if and only if $c_1 = c_2$ and $k_1 = k_2$. (We similarly assume $\bar{\mathbf{x}}_{k_1}^{c_1} = \bar{\mathbf{x}}_{k_2}^{c_2}$ if and only if $c_1 = c_2$ and $k_1 = k_2$.) This can be made trivially true by letting c and k be the first two behavioral variables. Thus, f is invertible. This trick is mathematically convenient for subsequent derivations—it does not change that multiple neural states can be associated with the same behavior as in Figure 2.

Estimating candidate states

The prior in Eq. (1) ensures that $P(\mathbf{x}_t = \mathbf{x} \mid \bar{\mathbf{s}}_{t'-\tau':t'}) = 0$ if $\mathbf{x} \notin \Omega$. Every element of Ω can be written as $\bar{\mathbf{x}}_k^c$ for some c and k , and the prior in Eq. (1) is uniform over the states in Ω . Eq. (2) and Eq. (5) indicate that each $\bar{\mathbf{x}}_k^c$ has only one possible recent history, a history that is specified by the stereotyped ordering of neural states in the c -th trajectory. Thus, the posterior probabilities over neural states in Ω do not need to marginalize over multiple potential state histories. Rather, the probability associated with each $\bar{\mathbf{x}}_k^c$ is proportional to the likelihood of recently observed spikes, $\bar{\mathbf{s}}_{t'-i':t'}$ given the unique history of neural states associated with $\bar{\mathbf{x}}_k^c$. Thus, the posterior probabilities over neural states in Ω can be written

$$P(\mathbf{x}_t = \bar{\mathbf{x}}_k^c \mid \bar{\mathbf{s}}_{t'-\tau':t'}) \propto \prod_{i'=0}^{\tau'} \prod_{n=1}^N h(\bar{s}_{n,t'-i'}, \lambda_{n,k-\delta-\Delta i'}^c) \quad (7)$$

$$\lambda_{n,j}^c = \frac{1}{\Delta} \sum_{i=0}^{\Delta-1} \bar{x}_{n,j-i}^c \quad (8)$$

where $h(s, \lambda) = \frac{\lambda^s e^{-\lambda}}{s!}$ (i.e. Poisson likelihoods, as dictated by Eq. (3)), and $\bar{s}_{n,i'}$ and $\bar{x}_n^c$ are the n -th components of $\bar{\mathbf{s}}_{i'}$ and $\bar{\mathbf{x}}_i^c$, respectively. Recall that T is the number of time samples elapsed since the completion of the most recent time bin. A normalizing constant can convert Eq.

(7) into full posterior probabilities. Computing this constant is straightforward because Ω is finite.

Eq. (4) and Eq. (6) allow the posterior over behavioral states to be written in terms of the posterior over neural states. $P(\mathbf{z}_t = \mathbf{z} | \bar{\mathbf{s}}_{t'-\tau':t'}) = 0$ if $\mathbf{z} \notin \Phi$, but for behaviors in Φ the posterior probabilities are

$$P(\mathbf{z}_t = \bar{\mathbf{z}}_k^c | \bar{\mathbf{s}}_{t'-\tau':t'}) = P(\mathbf{x}_t = f^{-1}(\bar{\mathbf{z}}_k^c) | \bar{\mathbf{s}}_{t'-\tau':t'}) \quad (9)$$

we can perform *maximum a posteriori* estimation on the log-transformed neural state posterior and read out the behavioral estimate directly.

$$\hat{\mathbf{x}}_t = \underset{\bar{\mathbf{x}}_k^c \in \Omega}{\operatorname{argmax}} \sum_{i'=0}^{\tau'} \sum_{n=1}^N \log h(\bar{s}_{n,t'-i'}, \lambda_{n,k-\delta-\Delta i'}^c) \quad (10)$$

$$\hat{\mathbf{z}}_t = f(\hat{\mathbf{x}}_t) \quad (11)$$

Lookup table of log-likelihoods

One could proceed with querying Eq. (10) for all $\bar{\mathbf{x}}_k^c \in \Omega$, selecting the most probable neural state, then reading out the associated behavioral state with Eq. (11). However, one may often wish to speed up this process, e.g. to deploy in a real-time application. In this section (and the following two sections), we describe a procedure for approximating Eq. (10) with considerably reduced computational burden.

Notice that the log-likelihood term in Eq. (10) only varies as a function of spike count and rate. Spike counts are finite and firing rates have limited dynamic ranges, which can be discretized, so it is possible to precompute and store these log-likelihoods in a lookup table in memory. Suppose the dynamic range of rates is given by $[\lambda_{\min}, \lambda_{\max}]$ and we sample this range with $V + 1$ values such that $\tilde{\lambda}(v) = \lambda_{\min} + \frac{v}{V}(\lambda_{\max} - \lambda_{\min})$ for $v \in \mathcal{V}$ where $\mathcal{V} = \{0, 1, \dots, V\}$. Every rate $\lambda_{n,j}^c$ can now be approximated by $\tilde{\lambda}(v_{n,j}^c)$ for some $v_{n,j}^c \in \mathcal{V}$ that minimizes the approximation error. If we define a lookup table with entries $L(s, v) = \log h(s, \tilde{\lambda}(v))$, then Eq. (10) can be rewritten

$$\hat{\mathbf{x}}_t = \underset{\bar{\mathbf{x}}_k^c \in \Omega}{\operatorname{argmax}} \sum_{i'=0}^{\tau'} \sum_{n=1}^N L(\bar{s}_{n,t'-i'}, v_{n,k-\delta-\Delta i'}^c) \quad (12)$$

Thus, during training we can compute L for all s and v . Similarly, we can compute v^c for all c, n , and j . This formulation ensures the costly process of computing log-likelihoods only needs to be performed once, during training.

Recursive solution

The estimation procedure in Eq. (12) can be made faster still by exploiting recursive structure. Consider the following definitions.

$$q_{t,k}^c = \sum_{i'=0}^{\tau'} r_{t'-i',k-\delta-\Delta i'}^c \quad (13)$$

$$r_{j',j}^c = \begin{cases} \sum_{n=1}^N L(\bar{s}_{n,j'}, v_{n,j}^c) & j' > 0, j > 0 \\ 0 & \text{otherwise} \end{cases} \quad (14)$$

This expression can be generated recursively,

$$q_{t,k}^c = \begin{cases} q_{t-\Delta,k-\Delta}^c + r_{t',k}^c - r_{t'-\tau'-1,k-\tau-1}^c, & \delta = 0 \\ q_{t-1,k-1}^c, & \delta > 0 \end{cases} \quad (15)$$

with initial conditions $q_{t,0}^c = q_{0,k}^c = 0$. Given that $q_{t,k}^c$ is simply shifted by one index relative to the previous time step when $T > 0$, the state estimate can be written

$$\hat{\mathbf{x}}_t = \begin{cases} \underset{\bar{\mathbf{x}}_k^c \in \Omega}{\operatorname{argmax}} q_{t,k}^c, & \delta = 0 \\ \bar{\mathbf{x}}_{\hat{k}+\delta}^{\hat{c}}, & \delta > 0 \end{cases} \quad (16)$$

where $\hat{c}$ and $\hat{k}$ are the indices such that $\hat{\mathbf{x}}_{t-\delta} = \bar{\mathbf{x}}_{\hat{k}}^{\hat{c}}$. Note that $\hat{\mathbf{x}}_t$ cannot be computed until $t > \tau$ (when sufficient spiking history has been collected).

Querying fewer states

The approach described in [Eq. \(16\)](#) is efficient insofar as it renders a new estimate at each time t while only requiring substantial computations every Δ time samples. Nevertheless, when $T = 0$, the recursion requires computing new $r_{t',k}^c$ and $r_{t'-\tau'-1,k-\tau-1}^c$ for every element in Ω^+ to ensure that $q_{t,k}^c$ is defined for every $\bar{\mathbf{x}}_k^c \in \Omega$. However, we typically assume some smoothness in neural trajectories over time. Thus, for reasonably small Δ , we can approximate the most likely neural state when $T = 0$ by only considering a subset of neural states defined by

$$\tilde{\Omega} = \left\{ \bar{\mathbf{x}}_k^c \in \Omega \mid \frac{k}{\Delta} \in \mathbb{Z} \right\} \quad (17)$$

which is equivalent to downsampling the neural trajectories in Ω with Δ spacing between samples. Letting $k' = \frac{k}{\Delta}$, this choice yields a new expression for the neural state estimate,

$$\hat{\mathbf{x}}_t = \begin{cases} \underset{\bar{\mathbf{x}}_k^c \in \tilde{\Omega}}{\operatorname{argmax}} \tilde{q}_{t',k'}^c, & \delta = 0 \\ \bar{\mathbf{x}}_{\hat{k}+\delta}^{\hat{c}}, & \delta > 0 \end{cases} \quad (18)$$

where $\tilde{q}_{t',k'}^c = q_{\Delta t', \Delta k'}^c$. The recursion can now be performed over fewer states.

$$\tilde{q}_{t',k'}^c = \tilde{q}_{t'-1,k'-1}^c + r_{t',\Delta k'}^c - r_{t'-\tau'-1,\Delta k'-\tau-1}^c \quad (19)$$

This simplification reduces memory and execution time. Furthermore, by using interpolation (described in the next section) we can still estimate neural states outside of $\tilde{\Omega}$.

Generating estimates at higher-than-bin-width resolution

[Eq. \(18\)](#) describes a procedure for generating new neural state estimates at every time sample. When $T = 0$, the estimate is updated using new spiking observations. When $T > 0$, the estimate is updated using the model of stereotyped trajectories, which specifies which neural state is likely to come next even in the absence of new spiking observations. However, to implement these equations in real time requires that the most time-consuming computations (those performed when $T = 0$) be performed within one time sample (e.g. 1 ms). MINT is very fast and in many cases

will be capable of generating estimates in under a millisecond (Table 1). When this isn't possible, a small decoding lag can be introduced (e.g. 5 ms) such that the computations required when $T = 0$ can begin prior to needing the result.

MINT (Part II): Interpolation

The algorithm described above leverages a library of neural trajectories composed of discrete sequences of neural states typically observed for a discrete set of conditions. These trajectories are presumed to sample an underlying neural manifold that is fundamentally continuous. The library of neural trajectories will more finely sample that manifold if one uses a large C such that small variations on a behavior each get their own neural and behavioral trajectories. Additionally, recall from Querying fewer states that, while decoding, only a subset of neural states (spaced apart by Δ) are considered. The larger Δ is, the more likely it is that the neural state estimate will make a small error in time (estimating a neural state that is slightly ahead or slightly behind the true state along a trajectory). Thus, neural state estimation can be more accurate in time and reflect more behavioral variability if one uses a small Δ and a large C . However, the computational burden of the algorithm grows as Δ shrinks or C grows. (The training data requirement also increases with C). This creates a potential trade-off between performance and computational burden. However, MINT side-steps this trade-off by using linear interpolation between states.

When the neural manifold is presumed to smoothly vary between neural states (across time or conditions) and there is also smooth variation between the corresponding behavioral states, it is unnecessary to use an overly small Δ and/or large C . Instead, one can identify candidate states along a small set of neural trajectories that coarsely sample the manifold, and then interpolate between those candidate states to find an intermediate state that is more consistent with the observed spikes. This two-step procedure (identify candidate states and then interpolate) allows the neural state estimate (and corresponding behavioral state estimate) to be accurate in time and capture variability between conditions (effectively as though a smaller Δ and larger C had been used) while preserving the computational benefits of a reasonable Δ (e.g. 20 ms) and small C .

Model of interpolated states

Suppose we identify two neural states, $\bar{\mathbf{x}}_{k_1}^{c_1}$ and $\bar{\mathbf{x}}_{k_2}^{c_2}$, and we'd like to consider the possibility that the actual neural state is somewhere between them. We can use a model of the following form:

$$\alpha \sim \text{Unif}(0, 1) \quad (20)$$

$$\mathbf{x}_{t-i} = (1 - \alpha)\bar{\mathbf{x}}_{k_1}^{c_1} + \alpha\bar{\mathbf{x}}_{k_2}^{c_2}, \quad i \in \{0, 1, \dots, \tau\} \quad (21)$$

$$\bar{\mathbf{s}}_{t'-i'} \sim \text{Pois} \left(\frac{1}{\Delta} \sum_{i=0}^{\Delta-1} \mathbf{x}_{\Delta(t'-i')-i} \right), \quad i' \in \{0, 1, \dots, \tau'\} \quad (22)$$

$$\mathbf{z}_t = (1 - \alpha)\bar{\mathbf{z}}_{k_1}^{c_1} + \alpha\bar{\mathbf{z}}_{k_2}^{c_2} \quad (23)$$

This model assumes that as the neural state smoothly varies from $\bar{\mathbf{x}}_{k_1}^{c_1}$ to $\bar{\mathbf{x}}_{k_2}^{c_2}$, the corresponding behavior smoothly varies from $\bar{\mathbf{z}}_{k_1}^{c_1}$ to $\bar{\mathbf{z}}_{k_2}^{c_2}$.

Interpolating between states

Explicitly evaluating the probabilities associated with all $\mathbf{x}_t$ in this model would be computationally inefficient. Given that the prior on α is uniform over the range $[0, 1]$, we can simply select the α that maximizes the log-likelihood of the observed spikes. The log-likelihood of

the observed spikes is given by the equation:

$$\check{q}_t(\alpha) = \log P(\bar{s}_{t',-\tau':t'}|\alpha) = \sum_{i'=0}^{\tau'} \sum_{n=1}^N \log h\left(\bar{s}_{n,t'-i'}, (1-\alpha)\lambda_{n,k_1-\delta-\Delta i'}^{c_1} + \alpha\lambda_{n,k_2-\delta-\Delta i'}^{c_2}\right) \quad (24)$$

Differentiating $\check{q}_t$ twice with respect to α yields

$$\frac{d^2\check{q}_t}{d\alpha^2} = - \sum_{i'=0}^{\tau'} \sum_{n=1}^N \bar{s}_{n,t'-i'} \left(\frac{\lambda_{n,k_2-\delta-\Delta i'}^{c_2} - \lambda_{n,k_1-\delta-\Delta i'}^{c_1}}{(1-\alpha)\lambda_{n,k_1-\delta-\Delta i'}^{c_1} + \alpha\lambda_{n,k_2-\delta-\Delta i'}^{c_2}} \right)^2 \quad (25)$$

Notice that $\frac{d^2\check{q}_t}{d\alpha^2} \leq 0$ always, i.e. $\check{q}_t$ is a concave function of α . Thus, we can rapidly compute $\hat{\alpha}$ using Newton's method and then estimate $\hat{\mathbf{x}}_t$ and $\hat{\mathbf{z}}_t$ via [Eq. \(21\)](#) and [Eq. \(23\)](#). This procedure will yield the same $\hat{\alpha}$ for all t associated with the same t' . Thus, the interpolation only needs to be performed once per time bin, when $T = 0$. In this paper, we used the following stopping criteria for Newton's method: 1) the estimate of α is within .01 of the estimate at the previous iteration, 2) the estimate of α saturates at 0 or 1, or 3) optimization runs for 10 iterations.

Interpolating across indices

[Eq. \(17\)](#) restricts the neural states under consideration to be spaced apart by Δ indices. Thus, a straightforward use of the interpolation scheme described above is to interpolate between nearby neural states along the same trajectory to restore the precision in the estimate that was lost by this restriction. First, a candidate neural state, which we'll denote $\bar{\mathbf{x}}_{k_1}^{c_1}$, is identified using [Eq. \(18\)](#).

Then, a second candidate neural state is identified, which we'll denote $\bar{\mathbf{x}}_{k_2}^{c_1}$.

This second state is whichever of the two adjacent states (either $\bar{\mathbf{x}}_{k_1+\Delta}^{c_1}$ or $\bar{\mathbf{x}}_{k_1-\Delta}^{c_1}$) has a larger log-likelihood of the observed spikes. Then, interpolation between $\bar{\mathbf{x}}_{k_1}^{c_1}$ and $\bar{\mathbf{x}}_{k_2}^{c_1}$ is performed as in [Interpolating between states](#) to yield neural and behavioral state estimates $\hat{\mathbf{x}}_t$ and $\hat{\mathbf{z}}_t$. This form of interpolation is less about generalization and more about enabling [Eq. \(17\)](#) to improve computational efficiency without loss of precision in the neural state estimates along each trajectory.

Interpolating across conditions

To interpolate across conditions, we also identify candidate neural states. The first candidate neural state, $\bar{\mathbf{x}}_{k_1}^{c_1}$, is again identified using [Eq. \(18\)](#). The second candidate neural state, $\bar{\mathbf{x}}_{k_2}^{c_2}$, is identified as the neural state that maximizes the log-likelihood of the observed spikes but is not a state along the same trajectory as $\bar{\mathbf{x}}_{k_1}^{c_1}$ (i.e. $c_2 \neq c_1$). Next, both $\bar{\mathbf{x}}_{k_1}^{c_1}$ and $\bar{\mathbf{x}}_{k_2}^{c_2}$ are each interpolated across indices with their adjacent states to yield improved candidate neural and behavioral state estimates. Then, an additional interpolation is performed between the improved candidate state estimates to yield the condition-interpolated neural and behavioral state estimates $\hat{\mathbf{x}}_t$ and $\hat{\mathbf{z}}_t$. This allows for generalization between similar behaviors.

Other forms of interpolation

The exact manner in which interpolation is applied during MINT may vary by task or dataset. Often, interpolation across indices and conditions will be sufficient. However, in this section we'll review some alternative interpolation implementations that were used for various analyses throughout this paper.

First, a dataset like MC_RTT that lacks explicit condition structure will not necessarily have one trajectory per behavior. We used AutoLFADS to learn neural trajectories for MC_RTT, which would have yielded a single neural trajectory for the entire training set were it not for recording gaps that broke up that trajectory. Thus, different portions of the same trajectory frequently corresponded to similar movements that it made sense to consider interpolations between. We therefore modified interpolation such that the second candidate state could be on the same trajectory as the first candidate state, but it had to be at least 1000 milliseconds away from that first state. Additionally, in a task with many degrees of behavioral freedom, there may be many trajectory fragments near a candidate neural state that could be worthwhile to interpolate to. Thus, for the MC_RTT dataset we expanded to multiple candidate neural states (6 for decoding analyses, 5 for the neural state estimation analysis; these values were optimized as hyperparameters). Interpolation occurred between all pairs of candidate neural states. Then, we proceeded to only use the interpolation that yielded the highest spike count log-likelihoods. These modifications ensured that utilizing long trajectories spanning multiple behaviors in a reaching task with many degrees of behavioral freedom would not limit the ability of interpolation to generalize between similar behaviors.

Second, in DMFC_RSG we computed multiple libraries of neural trajectories, each library corresponding to a different section of the training data in which the overall firing rate statistics differed due to recording instabilities throughout the session. For this dataset, we wished to expand interpolation to occur across indices, conditions, *and* libraries. We simply extended the approach for interpolating across conditions one step further. The first candidate neural state was selected as the most likely neural state across all libraries. The second candidate neural state was selected as the most likely neural state in any library except the one in which the first candidate state resided. For each of these two candidate states, within-library interpolation across indices and conditions proceeded to improve the estimates. Then, a final interpolation occurred between the best estimate in the first library with the best estimate in the second library. This approach enabled MINT to robustly estimate the neural state even when fed test data from various sections of the dataset in which different recording instabilities were present.

Third, in MC_Cycle we decoded a circular variable: phase. To accommodate this circularity, phase interpolations were modified such that they always occurred across the acute angle between the two phases. For example, if two candidate states corresponding to phases of 10 degrees and 350 degrees were identified, interpolation occurred across the 20 degree difference rather than the 340 degree difference.

Datasets

We analyzed 9 empirical datasets (Area2_Bump, DMFC_RSG, MC_Cycle, MC_Maze, MC_Maze-L, MC_Maze-M, MC_Maze-S, MC_RTT, MC_PacMan), several simulated maze datasets, and an artificial multitask spiking network. All empirical datasets contained spike times from offline or online sorted units. Detailed descriptions of all empirical datasets (except MC_Cycle and MC_PacMan) can be found in Pei et al. [104]. Here, we briefly describe each of the datasets.

Area2_Bump

This dataset contains neural recordings (96-electrode Utah array) from Brodmann's area 2 of S1 (a proprioceptive area) while a monkey used a manipulandum to make center-out-reaches to one of eight targets. On some trials, the monkey volitionally reached to the targets ('active' trials). On other trials, the manipulandum perturbed the monkey's arm out toward one of the targets ('passive' trials). Thus, right after movement onset the 'active' trials involved predictable sensory feedback (the monkey planned the movement) and the 'passive' trials involved unexpected sensory feedback (the monkey was not planning to move). Between the eight targets and two trial types ('active' and 'passive'), there were a total of 16 conditions. The behavioral data for this dataset includes: hand position (x- and y-components), hand velocity (x- and y-components), forces

and torques applied to the manipulandum (x- y- and z-forces; x- y- and z-torques), 7 joint angles, 7 joint velocities, 39 muscle lengths, and 39 muscle velocities. Position data was zeroed at movement onset for each trial to ensure position data was relative to a pre-movement baseline. More information on the task and data can be found in Chowdhury, Glaser, and Miller [99].

DMFC_RSG

This dataset contains neural recordings (three Plexon probes) from dorsomedial frontal cortex while a monkey performed a time-interval reproduction task (aka Ready-Set-Go). In this task, the monkey received two visual cues, 'Ready' and 'Set', separated in time by a sample interval. The monkey was required to estimate the sample interval and report this estimate by performing an action ('Go') that followed the 'Set' cue by a production interval. The monkey was rewarded depending on how close the production interval was to the sample interval on a given trial. The sampling interval was selected on each trial from one of two prior distributions. The short prior distribution contained sampling intervals of 480, 560, 640, 720, and 800 ms. The long prior distribution contained sampling intervals of 800, 900, 1000, 1100, and 1200 ms. Trials were selected from prior distributions in blocks (i.e. many consecutive trials were drawn from the same prior distribution). Depending on the trial, the action could be reported in one of four ways: a joystick movement or an eye saccade to either the left or right. Between the two prior distributions, five sampling intervals per prior distribution, and four actions, there were a total of 40 conditions. More information on the task can be found in Sohn et al. [62].

MC_Cycle

This dataset contains neural recordings (two 96-electrode Utah arrays) from motor cortex (M1 and PMd) while a monkey performed a cycling task to navigate a virtual environment. Offline sorting of spike waveforms (using Kilosort [113]) yielded single-neuron and high-quality multi-neuron isolations. The dataset also includes threshold crossings that were detected online. On each trial, the monkey grasped a hand-pedal and moved it cyclically forward or backward to navigate the virtual environment. The color of the virtual landscape cued the monkey as to whether forward or backward pedaling would advance the avatar in the virtual environment ('green' landscape: forward pedaling advanced the avatar; 'tan' landscape: backward pedaling advanced the avatar). Pedaling distance was cued by a virtual target that appeared a fixed distance away in the virtual environment. The number of complete cycles of pedaling required to reach the target was either 1, 2, 4, or 7 cycles. Between the two pedaling directions and four pedaling distances, there were a total of 8 conditions. The behavioral data for this dataset includes: pedal phase, angular velocity of the pedal, x- and y- pedal position, x- and y- pedal velocity, and 7 intramuscular EMG recordings. The dataset includes both sorted spikes and threshold crossings. More information on the task can be found in Schroeder et al. [20].

MC_Maze

This dataset contains neural recordings (two 96-electrode Utah arrays) from motor cortex (M1 and PMd) while a monkey performed straight and curved reaches. The monkey was presented with a maze on a vertical screen with a cursor that tracked the motion of the monkey's hand. Targets appeared on the screen and the monkey had to make reaches that would move the cursor onto the target. Virtual barriers were presented that the cursor could not pass through, requiring curved reaches that avoided the barriers. Each maze configuration had three variants. The first variant had a single target with no barriers (straight reach). The second variant had a single target with a barrier (curved reach). The third variant had a target with a barrier (curved reach) as well as two unreachable distractor targets elsewhere in the maze. There were 36 maze configurations for a total of 108 conditions. The behavioral data for this dataset includes x- and y-components of hand position and velocity. Position data was zeroed at movement onset for each trial to ensure position data was relative to a pre-movement baseline. More information on the task can be found in Churchland et al. [98].

MC_Maze-(L,M,S)

These three datasets were collected from the same monkey as in MC_Maze performing the same task, but on different days with different conditions. The same arrays were used to collect neural recordings (though different numbers of sorted units were identified in each session) and the same behavioral data was collected. These datasets differ from MC_Maze primarily in that they have fewer conditions (nine maze configurations with three variants, for a total of 27 conditions) and fewer trials.

MC_RTT

This dataset contains neural recordings (96-electrode Utah array) from motor cortex while a monkey made reaches in a horizontal plane to targets in an 8×8 grid. The task had no explicit trial structure nor did it require that individual movements be repeated throughout the session. Rather, a target would appear in one of the 64 grid locations and the monkey would reach to that target. Then, a new target would appear at a different random grid location and the monkey would reach there. This process continued throughout the session, leading to a collection of reaches that varied in terms of reach direction, reach distance, and reach speed. Despite lacking meaningful trial structure, the session was nevertheless broken up into 600 ms ‘trial’ segments (hence the trial count listed in [Table 2](#)). The behavioral data for this dataset includes x- and y-components of finger position and velocity. Due to the structure of the training data, neural trajectories were learned in long stretches spanning multiple movements. Thus, the zeroing of position data at movement onset that was employed for other datasets would have led to discontinuities in the behavioral trajectories for this task. We could have decoded position with no zeroing – however, similar movements (with similar corresponding neural activity) could have very different starting and ending positions in this dataset. Thus, we decided to focus decoding analyses on velocity only. More information on the task can be found in Makin et al. [\[90\]](#).

MC_PacMan

This dataset contains neural recordings (128 electrodes recorded simultaneously from a passive pre-production version of the Neuropixels 1.0-NHP probe) from motor cortex while a monkey generated forward isometric force against an immovable handle. The monkey was presented with a PacMan icon on a screen. PacMan’s vertical position corresponded to the force applied to the handle. A path of scrolling dots moved horizontally across the screen and the monkey was rewarded when PacMan intercepted the dots. The pattern of scrolling dots was used to cue various dynamic force profiles. The task contained eight conditions that all spanned approximately the same range of forces: four ramping force profiles (fast and slow, ascending and descending), three sinusoidal force profiles (0.25 Hz, 1 Hz, 3 Hz), and a chirp profile in which the frequency steadily increased throughout the trial from 0-3 Hz. More information on the task can be found in Marshall et al. [\[74\]](#).

Maze task simulations

Several synthetic datasets were generated based on the maze task. These simulations were based on the MC_Maze dataset and matched the neuron, condition, and trial counts from MC_Maze. In all cases, the actual behavior from MC_Maze was used — spiking activity was the only component that was simulated.

To simulate firing rates, we first z-scored hand position, velocity, and acceleration (x- and y-components). Then, we simulated rates for each neuron as a random weighted sum of these z-scored kinematics with a constant offset. The weights and offsets were scaled such that the means

Dataset	Neurons	Conditions	Trials
Area2_Bump	65	16	364
DMFC_RSG	54	40	1006
MC_Cycle	112	8	273
MC_Maze	182	108	2295
MC_Maze-L	162	27	500
MC_Maze-M	152	27	250
MC_Maze-S	142	27	100
MC_RTT	130	N/A	1080
MC_PacMan	128	8	362
Maze Simulations	182	108	2295
Multitask Network	1200	9	270

Table 2.

Neuron, condition, and trial counts for each dataset. For some datasets, there are additional trials that are excluded from these counts. These trials are excluded because they were only usable for Neural Latents Benchmark submissions due to hidden behavioral data and partially hidden spiking data (see Neural Latents Benchmark).

and standard deviations of firing rates for the simulated neurons matched the means and standard deviations of the empirical trial-averaged rates in MC_Maze. In some simulations, we increased simulated firing rates by a factor of 5 or 10.

Simulated spiking activity was generated from simulated rates in one of two ways. In one simulation, spiking activity was simulated as a Poisson process, which corresponds to exponential-interval spiking. In other simulations, spiking activity was simulated with gamma-interval spiking. More specifically, we let the interval between spikes be drawn from a gamma distribution with a shape parameter $\alpha = 2$ and a rate parameter $\beta = 2\lambda$, where λ is the simulated firing rate. For reference, exponential-interval spiking corresponds to a gamma distribution with $\alpha = 1$ and $\beta = \lambda$. The overall rate of spiking was matched regardless of whether exponential-interval or gamma-interval spiking was simulated, but spiking occurred more regularly given the same rate for the gamma-interval simulations.

Multitask spiking network

A simulated multitask dataset was generated using the method described in DePasquale et al. [55]. A spiking network was trained to generate posterior deltoid activity during reaching (eight reach conditions) and cycling (one cycling condition). In the network, there were 39 latent factors related to reaching and 12 latent factors related to cycling. The factors, and therefore the neural trajectories, for reaching and cycling were roughly orthogonal to one another. The network simulated 500 consecutive trials. A 135-trial stretch of this simulated dataset was designated as the test set (~7.1 minutes of data). The training set was constructed by randomly selecting 15 trials per condition (135 training trials total) from the remaining trials.

Learning idealized trajectories

When training data includes repeated trials of the same behavior, neural and behavioral trajectories can be learned via standard trial-averaging. This procedure begins with filtering spikes temporally (e.g. with a Gaussian filter) to yield single-trial rates. Then, rates and behavioral variables are extracted on each trial relative to behaviorally relevant trial events (e.g. movement onset). Depending on the structure of the task, time may need to be warped to ensure that task epochs unfold on the same timescale across different trials of the same condition. Finally, the single-trial rates are averaged across trials of the same condition to yield firing rates. The vector of firing rates at each time within each condition defines a neural state and the collection of neural states within each condition defines a neural trajectory. Single-trial behavioral variables are similarly averaged across trials to yield behavioral states and trajectories.

In the sections below, we describe how this standard trial-averaging procedure was applied to the datasets in this paper. We describe an additional trial-smoothing step that can be applied prior to averaging that often yields a small performance boost for MINT. We also describe a procedure for smoothing across conditions and neurons after averaging that can improve performance. All trajectories were learned at millisecond resolution.

There are two datasets for which non-standard trajectory-learning procedures were utilized. First, the training set for MC_RTT does not contain repeated trials of the same behavior and therefore was not amenable to the trial-averaging procedure. Trajectories were instead learned using AutoLFADS, as described in its own section below. Second, DMFC_RSG contained substantial recording instabilities. Thus, although this dataset still used trial-averaging, an additional procedure was developed for learning multiple libraries of neural trajectories corresponding to the same behaviors at different portions of the session in which different recording instabilities were present. This procedure is also described in its own section below.

Filtering, extracting, and warping data on each trial

First, spiking activity for each neuron on each trial was temporally filtered with a Gaussian to yield single-trial rates. **Table 3** reports the Gaussian standard deviations σ (in milliseconds) used for each dataset. The optimal σ for a dataset may depend on neuron count, typical spike rate, number of trials, and the timescale with which rates change in any given task or brain area. Given these factors vary across datasets, the optimal σ is expected to vary as well. Thus, σ was treated as a hyperparameter and optimized per-dataset.

Next, single-trial rates and behavioral data were extracted on each trial relative to key trial events. For example, in the maze datasets, data was extracted beginning 500 ms prior to movement onset and ending 700 ms after movement onset. These extraction boundaries (listed in **Table 3**) determine when the neural trajectories begin and end. These boundaries should be set to ensure that the behaviors one is interested in decoding are represented in the library of trajectories. It is also important to have extra data at the beginning of each trajectory because the first τ indices on each trajectory cannot be selected as candidate states (due to lack of sufficient state history).

In the maze datasets, Area2_Bump, and MC_PacMan, there was no need to warp time – all movements of the same condition unfolded over similar timescales relative to movement onset. However, in the cycling dataset this was not the case. The duration between movement onset and offset was variable across trials of the same condition. Thus, data in MC_Cycle was time-warped (separately for each condition) using the procedure described in Russo et al. [61]. In the DMFC_RSG dataset, the duration between the ‘set’ cue and ‘go’ action was variable across trials. Thus, uniform time-warping to match the average duration for each condition was used. More generally, one may find the time-warping technique in Williams et al. [111] to be a useful tool. It is important that warping take place *after* filtering spikes, because warping raw spiking activity would change the neurons’ rates.

The result of filtering, extracting, and warping is that, for each condition c , single-trial rates can be formatted into a tensor $\mathbf{X}^c \in \mathbb{R}^{N \times K_c \times R}$ where N is the number of neurons, K_c is the number of extracted time points on each trial, and R is the number of trials. Similarly, behavioral data can be formatted into tensors $\mathbf{Z}^c \in \mathbb{R}^{M \times K_c \times R}$ where M is the number of behavioral variables.

Averaging across trials

It is straightforward to average each $\mathbf{X}^c$ across trials yielding matrices $\bar{\mathbf{X}}^c \in \mathbb{R}^{N \times K_c}$. We refer to this as Type I averaging. However, we found that a modified trial-averaging procedure (Type II) led to slightly improved neural state estimation and decoding in most cases. First, all rates ($\mathbf{X}^c$ for all c) are mean-centered and soft-normalized as described in Russo et al. [61]. The offset and scaling factor for this step are computed based on Type I averaged rates. Then, the mean-centered and soft-normalized rates are formatted into matrices $\mathbf{X}_{in}^c \in \mathbb{R}^{R \times N K_c}$. PCA is used to reduce the trial-dimensionality of the rates to 1 for each c via

$$\mathbf{X}_{out}^c = \mathbf{W}_c \mathbf{W}_c^\top \mathbf{X}_{in}^c \quad (26)$$

where $\mathbf{W}_c \in \mathbb{R}^R$ is a column-vector corresponding to the first principal component of $\mathbf{X}_{in}^c$. Then, we reverse the soft-normalization and mean-centering operations on $\mathbf{X}_{out}^c$ and average across trials to get $\bar{\mathbf{X}}^c$ (Type II). **Table 3** indicates for each dataset which trial-averaging procedure was used.

Dataset	Trajectory start	Trajectory end	σ	Trial Averaging	D_{neural}	$D_{condition}$
Area2_Bump	$t_{move} - 350$	$t_{move} + 750$	25	Type II	Full	Full
DMFC_RSG	$t_{go} - 1950$	$t_{go} + 750$	55	Type I	49	17
MC_Cycle	$t_{move} - 600$	$t_{stop} + 600$	30	Type I	Full	N/A
MC_Maze	$t_{move} - 500$	$t_{move} + 700$	30	Type II	Full	21
MC_Maze-L	$t_{move} - 500$	$t_{move} + 700$	35	Type II	Full	16
MC_Maze-M	$t_{move} - 500$	$t_{move} + 700$	60	Type II	Full	20
MC_Maze-S	$t_{move} - 500$	$t_{move} + 700$	85	Type II	120	10
MC_PacMan	$t_{prof_start} - 1000$	$t_{prof_end} + 1000$	35	Type II	Full	Full
Maze Simulations	$t_{move} - 500$	$t_{move} + 700$	30	Type II	Full	21
Multitask Network	(C) $t_{move} - 1000$	(C) $t_{stop} + 1000$	15	Type I	Full	N/A
	(R) $t_{move} - 800$	(R) $t_{move} + 1000$				

Table 3.

Hyperparameters for learning neural trajectories via standard trial-averaging. σ is the standard deviation of the Gaussian kernel used to temporally filter spikes. Trial-averaging Type I and Type II procedures are described in Averaging across trials. D_{neural} and $D_{condition}$ are the neural- and condition-dimensionalities described in Smoothing across neurons and/or conditions. 'Full' means that no dimensionality reduction was performed. Condition smoothing could not be performed for MC_Cycle or the multitask network because different conditions in these datasets are of different lengths (i.e. K_c is not the same for all c). (C) and (R) refer to the cycling and reaching trajectories, respectively, in the multitask network. t_{move} , t_{stop} , and t_{go} correspond to movement onset, movement offset, and the 'go' time in the ready-set-go task, respectively. In MC_PacMan, each force profile is padded with static target forces at the beginning and end. t_{prof_start} and t_{prof_end} mark the beginning and end of the non-padded force profile on each trial.

We average each $\mathbf{Z}^c$ across trials yielding matrices $\bar{\mathbf{Z}}^c \in \mathbb{R}^{M \times K_c}$. Typically, no modified trial-averaging procedure is required for behavioral data. However, there are cases where a behavioral variable should not be averaged without some pre- and/or post-processing. For example, in MC_Cycle, phase is a circular variable and therefore can't be linearly averaged. To accommodate this, we unwrapped phase, averaged across trials, then re-wrapped it.

Smoothing across neurons and/or conditions

To additionally improve trial-averaged rate estimates we can use dimensionality reduction to smooth across neurons and/or conditions. First, trial-averaged rates are mean-centered and soft-normalized. Then, they are concatenated across conditions into a matrix $\bar{\mathbf{X}}_{in} \in \mathbb{R}^{N \times K}$ where $K = \sum_c K_c$. PCA is used to reduce the neural-dimensionality of the rates via

$$\bar{\mathbf{X}}_{out} = \mathbf{W} \mathbf{W}^T \bar{\mathbf{X}}_{in} \quad (27)$$

where $\mathbf{W} \in \mathbb{R}^{N \times D_{neural}}$ is the top D_{neural} principal components of $\bar{\mathbf{X}}_{in}$. Then, $\bar{\mathbf{X}}_{out}$ is reformatted into a new $\bar{\mathbf{X}}_{in} \in \mathbb{R}^{C \times N K_c}$ where K_c is the same for all c (this form of smoothing cannot be applied when this is not the case). PCA is used to reduce the condition-dimensionality of the rates via

$$\bar{\mathbf{X}}_{out} = \mathbf{W} \mathbf{W}^T \bar{\mathbf{X}}_{in} \quad (28)$$

where $\mathbf{W} \in \mathbb{R}^{C \times D_{condition}}$ is the top $D_{condition}$ principal components of $\bar{\mathbf{X}}_{in}$. Then, we reverse the soft-normalization and mean-centering operations on $\bar{\mathbf{X}}_{out}$, rectify to ensure non-negative rates, and reformat the smoothed rates back into matrices $\tilde{\mathbf{X}}^c \in \mathbb{R}^{N \times K_c}$. Not all datasets benefited from neuron smoothing and/or condition smoothing. The smoothing dimensionalities used for each dataset are reported in [Table 3](#).

After completing the previous steps (some combination of smoothing rates over time, trials, neurons, and/or conditions), there will be matrices of firing rates $\tilde{\mathbf{X}}^c$ and behavioral data $\bar{\mathbf{Z}}^c$ for each c . The neural and behavioral states comprising the idealized neural and behavioral trajectories are then directly defined by these matrices.

$$\bar{\mathbf{x}}_k^c = \tilde{\mathbf{X}}_{:,k}^c \quad (29)$$

$$\bar{\mathbf{z}}_k^c = \bar{\mathbf{Z}}_{:,k}^c \quad (30)$$

Learning trajectories for MC_RTT

Neural trajectories were learned for MC_RTT using AutoLFADS [107]. Following the procedure described in Keshtkaran et al. [107], spiking activity in the training data was formatted into 600 ms segments with spike counts binned every 5 ms. Each segment overlapped with the previous segment by 200 ms and with the subsequent segment by 200 ms. AutoLFADS was then run twice on this formatted data (using 80/20 training and validation splits), yielding two sets of rate estimates that we then averaged across. Each AutoLFADS run yields slightly different results. Thus, we reasoned that averaging across runs would improve performance by reducing the impact of idiosyncratic rate variability that doesn't show up consistently across runs. We chose not to average more than two runs due to the computational burden this would accrue for training.

After averaging across runs, the rates from each segment were stitched back together (via the weighted average described in Keshtkaran et al. [107]) into long neural trajectories spanning many movements. The trajectories were then upsampled to 1 kHz via linear interpolation. If there had been no recording gaps, this procedure would have yielded a single neural trajectory. For the decoding analyses, there were 2 recording gaps and therefore 3 neural trajectories. For the neural state estimation analysis (which utilized a larger training set), there were 3 recording gaps and therefore 4 neural trajectories. In all cases, each trajectory contained ~ 2.7 minutes of data. Although one might prefer trajectory boundaries to begin and end at behaviorally relevant moments (e.g. a stationary state), rather than at recording gaps, the exact boundary points are unlikely to be consequential for trajectories of this length that span multiple movements. If MINT estimates a state near the end of a long trajectory, its estimate will simply jump to another likely state on a different trajectory (or earlier along the same trajectory) in subsequent moments. Clipping the end of each trajectory to an earlier behaviorally-relevant moment would only remove potentially useful states from the libraries.

When running AutoLFADS for the neural state estimation analysis, one of the runs returned a highly oscillatory set of rates (~ 25-40 Hz oscillations) that did not by eye resemble the output of the other AutoLFADS runs. Although we confirmed this set of rates performed similarly for neural state estimation (.200 bits/spike using the oscillatory rates alone, compared to .201 bits/spike when averaging across two non-oscillatory runs), we nevertheless chose to discard this run and average across two non-oscillatory AutoLFADS runs because the oscillatory solution was not one that AutoLFADS consistently returns for this particular training set.

See Other forms of interpolation for details on how these long neural trajectories are used during interpolation to improve neural state estimation and decoding.

Learning trajectories for DMFC_RSG

Learning neural trajectories for the DMFC_RSG dataset involved two main challenges. First, we needed to accommodate that each trial contained multiple trial events to align to. Second, we desired a procedure for creating multiple libraries of neural trajectories corresponding to different portions of the session in which different recording instabilities were present. The solutions we developed for each of these challenges are described below.

Each trial contained five main trial events: the ‘fixation’ time when the monkey began looking at the center of the screen, the ‘target onset’ time when a white target appeared in one of two locations on the screen, the ‘ready’ cue, the ‘set’ cue, and the ‘go’ time at which the monkey initiated a joystick movement or saccade (depending on the condition). Thus, when learning neural trajectories, we needed to compute rates for five trial epochs: fixation to target onset, target onset to ready cue, ready cue to set cue, set cue to go time, and go time to the end of the trial. The first three epochs could not be averaged across trials in the standard way (because the durations varied from one trial to the next), but it was also not appropriate to warp them. Warping should be performed when the monkey is executing a computation or behavior at a variable speed. In this case, the monkey was simply waiting for the next trial event without knowledge of when it would arrive. Thus, we chose to align to the initial event in each epoch, average across trials (ignoring missing data associated with the variable epoch durations), and then trim the averaged rates to the median epoch duration. In this task, certain conditions are not disambiguated for the monkey until the set cue arrives (e.g. the monkey doesn’t know if they are in an 800 ms or 900 ms interval trial until the set cue arrives). Thus, averages for a given condition in these first three epochs included all trials from the given condition plus all trials from other conditions that could not be distinguished from the given condition at this stage of the trial. In the set cue to go time epoch, rates were uniformly warped to match the average duration within each condition and then averaged across trials. All data after the go time were aligned to the go time and averaged. After

computing these epoch-specific rate averages, the rates were concatenated across epochs to yield neural trajectories for each condition. They were then trimmed according to the trajectory start and end times listed in **Table 3**.

With access to a very large dataset, computing multiple libraries of neural trajectories for different recording conditions is straightforward: simply break the training data up in time into different sections and separately compute a library of trajectories for each. However, given limited training data, we sought to implement a similar strategy while not limiting the trajectories in each section to be based only on the small amount of spiking data available in that section. Thus, we first learned a library of neural trajectories based on the entire session (following the procedure described above) and smoothed across neurons and conditions to improve this estimate (yielding $\tilde{\mathbf{X}}^c$ for each condition c). Then, we broke up the session into 6 consecutive sections and proceeded to learn a complete set of neural trajectories for each section with no smoothing across neurons or conditions (yielding $\bar{\mathbf{Y}}^{c,d}$ for each condition c and section d). Finally, we learned a linear transformation of $\tilde{\mathbf{X}}^c$ for each section d that would better match the transformed firing rates for each section to those in $\bar{\mathbf{Y}}^{c,d}$ while preserving some of the neural geometry learned at the level of the session. This occurred via the equation

$$\tilde{\mathbf{Y}}^{c,d} = (W^{(d)}\Lambda + I)\tilde{\mathbf{X}}^c + b^{(d)} \quad (31)$$

where $W^{(d)} \in \mathbb{R}^{N \times N}$, $b^{(d)} \in \mathbb{R}^N$, and $\Lambda \in \mathbb{R}^{N \times N}$. Here, Λ is simply a diagonal matrix whose action on $\tilde{\mathbf{X}}^c$ is to soft-normalize. The values in Λ are learned from $\tilde{\mathbf{X}}^c$ via the same soft-normalization procedure described in Russo et al. [61] and used previously in this paper. The linear transformation was formulated in Eq. (31) to ensure that all regularization when learning the weights only applied to the deviation between the session-wide rates and the section-specific transformed rates. $W^{(d)}$ and $b^{(d)}$ were learned via an L2-regularized weighted least squares regression,

$$W^{(d)} = Q^{(d)}SP^T(PSP^T + \lambda I)^{-1} \quad (32)$$

$$b^{(d)} = \mu_1^{(d)} - W^{(d)}\mu_2 \quad (33)$$

$$Q^{(d)} = \bar{\mathbf{Y}}^d - \tilde{\mathbf{X}} - \mu_1^{(d)} \quad (34)$$

$$P = \Lambda\tilde{\mathbf{X}} - \mu_2 \quad (35)$$

$$\bar{\mathbf{Y}}^d = \begin{bmatrix} \bar{\mathbf{Y}}^{1,d} & \bar{\mathbf{Y}}^{2,d} & \dots & \bar{\mathbf{Y}}^{C,d} \end{bmatrix} \quad (36)$$

$$\tilde{\mathbf{X}} = \begin{bmatrix} \tilde{\mathbf{X}}^1 & \tilde{\mathbf{X}}^2 & \dots & \tilde{\mathbf{X}}^C \end{bmatrix} \quad (37)$$

where S is a diagonal matrix of observation weights, $\mu_1^{(d)}$ is the mean of $\bar{\mathbf{Y}}^d - \tilde{\mathbf{X}}$ across columns, and μ_2 is the mean of $\Lambda\tilde{\mathbf{X}}$ across columns. In Eq. (34) and Eq. (35), the subtractions of $\mu_1^{(d)}$ and μ_2 are applied to each column of the matrices they subtract from. The diagonal entries of S weight how much each observation should matter when learning $W^{(d)}$. Given that the epoch between the set cue and the go time is when the monkey is performing an internal computation (keeping track of elapsed time internally) and this epoch is a large fraction of the evaluated trial period in the Neural Latents Benchmark, we decided $W^{(b)}$ should prioritize fitting this epoch well. Thus, diagonal entries of S corresponding to samples in the set-go epoch were given a weight $s_{\text{set go}}$ whereas all other diagonal entries of S were set to 1. Both the set-go epoch weight and the L2 regularization term were treated as hyperparameters and optimized on a validation set, yielding

$s_{set-go} = 4$ and $\lambda = 100$. To summarize, $W^{(d)}$ transforms the mean-centered, soft-normalized session-wide rates into a set of rate residuals that can be added to the session-wide rates to better match the firing rate statistics in each section of the data. The result of this procedure is a complete library of neural trajectories for each section of the training data d .

$$\bar{\mathbf{x}}_k^{c,d} = \tilde{\mathbf{Y}}_{:,k}^{c,d} \quad (38)$$

See Other forms of interpolation for details on how these multiple libraries are used during interpolation to improve the neural state estimate. There are no behavioral trajectories for this dataset (it was only used for neural state estimation analyses).

Visualizing neural trajectories

Neural trajectories have dimensionality matching the number of recorded neurons (or multi-units), but it is often useful to visualize them in a 2D state space. Throughout this paper, PCA is used to project neural trajectories into a low-dimensional subspace. When this done, firing rates are preprocessed with mean-centering and soft-normalization (as in Russo et al. [61]). In some cases, it is useful to rotate the data within the top PCs to find a perspective that highlights certain properties of the trajectories. When this is done, we report the neural variance captured by the plotted dimensions (along with the neural variance captured by the top PCs) to contextualize the scale of the neural dimensions. Percent neural variance captured was computed in the standard way, as described in Schroeder et al. [20].

Distance analyses

In **Figure 2c-e**, several distance metrics were used to analyze the MC_Cycle dataset. Neural distance is defined as the Euclidean distance between two neural states in the full-dimensional space. Muscle distance is defined by z-scoring the ‘muscle state’ (portion of behavioral state corresponding to muscle activity), then computing the Euclidean distance between two normalized muscle states. The muscle state consisted of seven intramuscular EMG recordings (long head of the biceps, anterior deltoid, lateral deltoid, posterior deltoid, trapezius, lateral tricep, and long head of the triceps). Kinematic distance is defined using the phases and angular velocities associated with two behavioral states. Both phase and angular velocity are z-scored (for phase, this utilizes the circular standard deviation). Then, the phase distance (d_{ph}) is computed as the circular distance between the two normalized phases. The angular velocity distance (d_{av}) is computed as the absolute difference between the two normalized angular velocities. Then, the kinematic distance is computed as $d_{kin} = \sqrt{d_{ph}^2 + d_{av}^2}$. All trajectories (neural and behavioral) were binned at 10 ms resolution prior to computing pairwise distances to keep the analysis computationally manageable. In **Figure 2e**, data was partitioned into two trial sets (‘Partition A’ and ‘Partition B’) for a control analysis. To keep the number of trials used to compute trajectories matched across all analyses in **Figure 2c-e**, only trials from ‘Partition A’ were utilized in the neural distance vs. muscle distance and neural distance vs. kinematic distance analyses. When plotting pairwise distances, all distances were normalized (separately for each plotting axis) such that 1 corresponded to the average pairwise distance.

The neural distances reported in **Figure 4d** are on the same scale as those in **Figure 2c-e**, i.e. they are Euclidean distances between neural states normalized by the average pairwise distance between neural states in the library of trajectories.

Decoding analyses

All decoding analyses utilized the train-test trial splits listed in **Table 4** (with the exception of generalization analyses that used subsets of these trials). Neural and behavioral trajectories were learned from the training set as described in Learning idealized trajectories. Then, MINT was

provided with spiking activity on test trials with which to decode behavior. The trial epochs over which performance was evaluated were set to match the evaluation epochs from the Neural Latents Benchmark (for datasets that were used in the benchmark) and are listed in **Table 4** [↗](#).

MINT used a bin size of $\Delta = 20$ ms. The amount of spiking history provided to MINT at each decoding time step varied by dataset and is listed in **Table 4** [↗](#). Decoding was always causal (i.e. only utilizing spiking history prior to the decoded moment), with one exception described below in MINT variants. MINT utilized interpolation for all datasets.

The lookup table of log-likelihoods utilized a minimum rate, λ_{min} , corresponding to 1 spike/second. Log-likelihoods in the table were also clipped so as not to drop below $\log(10^{-6})$. These two choices regularize decoding such that a spurious spike from a neuron whose rate is close to 0 cannot dominate the decision of which neural state is most probable.

Decoding R^2

Decoded behavioral variables were compared to ground truth behavioral variables over the evaluation epoch of all test trials. Performance was reported as the coefficient of determination (R^2) for the decoded behavioral variables, averaged across behavioral variables within a behavioral group. For example, R^2 for x-velocity and y-velocity was averaged and reported as ‘velocity R^2 ’. When computing R^2 for phase, circular distances and circular means were substituted for traditional distances and means in the R^2 definition. For the neural networks, the training/testing procedure was repeated 10 times with different random seeds. For most behavioral variables, there was very little variability in performance across repetitions. However, there were a few outliers for which variability was larger. Reported performance for each behavioral group is the average performance across the 10 repetitions to ensure results were not sensitive to any specific random initialization of each network. Decoding R^2 was always computed at 5 ms resolution (set to match the resolution used by the Neural Latents Benchmark).

Decoder comparison

MINT was primarily compared to four other decode algorithms: Kalman filter, Wiener filter, feedforward neural network, and a recurrent neural network (GRU). An additional comparison to a Naïve Bayes regression decoder was made for the MC_Maze dataset. All decoders were provided with the same training/testing data and used the same evaluation epochs. Complete details on their implementations are provided in the Appendix.

Neuron dropping

The neuron dropping analyses in **Figure 8** [↗](#) were performed on the MC_Maze-L dataset, which has 162 neurons. The analysis of decoding performance for known neuron loss consisted of training and testing the decoders with reduced sets of neurons. The analysis of performance for undetected neuron loss consisted of training and testing the decoder with all 162 neurons, but in the test set a subset of those neurons were artificially set to never spike. The analyses were performed for the following neuron counts: [1, 2, 3, 4, 5, 10, 15,..., 150, 155, 160, 162]. At each neuron count, the known loss and undetected loss analyses were each repeated 50 times (with different neurons randomly dropped at each repetition). Linear interpolation between these neuron counts was then applied to generate R^2 values for every neuron count between 1 and 162.

MINT variants

The analysis in **Figure S3a-c** [↗](#) required training and testing several variations on MINT for Area2_Bump, MC_Cycle, MC_Maze, and MC_RTT. These variations are described here.

Dataset	Training trials	Test trials	Evaluation start	Evaluation end	Window length (ms)
Area2_Bump	272	92	$t_{move} - 100$	$t_{move} + 500$	240
MC_Cycle	174	99	$t_{move} - 250$	$t_{stop} + 250$	200
MC_Maze	1721	574	$t_{move} - 250$	$t_{move} + 450$	300
MC_Maze-L	375	125	$t_{move} - 250$	$t_{move} + 450$	300
MC_Maze-M	188	62	$t_{move} - 250$	$t_{move} + 450$	300
MC_Maze-S	75	25	$t_{move} - 250$	$t_{move} + 450$	340
MC_RTT	810	268	t_{start}	$t_{start} + 599$	480
MC_PacMan	234	128	$t_{prof_start} - 500$	$t_{prof_end} + 1000$	300
Maze Simulations	1721	574	$t_{move} - 250$	$t_{move} + 450$	300
Multitask Network	135	135	t_{start} (first trial)	t_{end} (last trial)	300

Table 4.

Details for decoding analyses. The number of training and testing trials for each dataset are provided along with the evaluation period over which performance was computed. Generalization analyses used subsets of these training and testing trials. The window lengths refer to the amount of spiking history MINT used for decoding (e.g. when $\tau' = 14$ and $\Delta = 20$, the window length is $(\tau' + 1)\Delta = 300$ ms). t_{move} corresponds to movement onset, t_{stop} corresponds to movement offset, t_{start} refers to the beginning of a trial, and t_{end} refers to the end of a trial. There is no defined condition structure for MC_RTT to use for defining trial boundaries. Thus, each trial is simply a 600 ms segment of data, with no alignment to movement. Although 270 of these segments were available for testing, the first 2 segments lacked sufficient spiking history for all decoders to be evaluated and were therefore excluded, leaving 268 test trials. In MC_PacMan, each force profile is padded with static target forces at the beginning and end. t_{prof_start} and t_{prof_end} mark the beginning and end of the non-padded force profile on each trial. For the multitask network, performance was evaluated on a continuous stretch of 135 trials spanning ~ 7.1 minutes.

The first variation determines how behavior is decoded from the neural state estimate. The ‘Direct MINT Readout’ utilizes the direct association between neural and behavioral states that is standard for MINT. The ‘Linear MINT Readout’ begins by estimating the neural state using MINT, but behavior is then decoded as a weighted sum of the rates comprising the neural state estimate. The weights are learned from training data using ridge regression, with the L2 regularization term in the regression optimized via a grid search with 5-fold cross validation. To improve the quality of the fit, a lag between neural and behavioral states was included (lags set to match the values used in Pei et al. [104]; MC_Cycle set to 100 ms).

The second variation determined whether interpolation was used. When interpolation was used, the same methodology used for **Figure 5** analyses was followed (6 candidate states for MC_RTT; 2 candidate states for other datasets). When interpolation was not used, the neural state was selected from the library of neural trajectories as the state that maximized the log-likelihood of the observed spikes.

The third variation determined whether decoding occurred causally or acausally. When decoding causally, the spiking observations all came prior to the decoded moment. When decoding acausally, the spiking observations were centered on the decoding moment. For the analysis in **Figure S3a-c**, the extent of spiking observations (window length) was optimized separately for causal vs. acausal decoding to give each variant the opportunity to perform as well as possible. The causal window lengths are reported in **Table 4**. The acausal window lengths were 560 ms for Area2_Bump, 400 ms for MC_Cycle, 580 ms for MC_Maze, and 660 ms for MC_RTT.

SNR criteria for threshold crossings

When decoding based on threshold crossings in **Figure S3d**, a signal-to-noise ratio (SNR) was computed for each electrode channel. Then, only electrode channels with $\text{SNR} > 2$ were utilized for decoding. The SNR was computed as the ratio of the range of firing rates (determined from trial-averaged rates) and the standard deviation of the firing rate residuals. The firing rate residuals on each trial were computed as the difference between the filtered spikes and the corresponding trial-averaged rates. If the standard deviation was less than 5 spikes/second, it was set to 5 spikes/second in the SNR computation. This ensured that channels with low variability in the residuals due to saturation at small or large firing rates were rejected. For example, an electrode channel that almost never spikes will have very low variability in the residuals simply because the firing rate is typically 0.

Training non-MINT decoders with trial-averaged data

In the Comparison to other decoders section, it is reported that non-MINT decoders were trained with access to trial-averaged data and this did not improve decoding performance. This result comes from an analysis of the MC_Maze dataset. The training set was augmented to include two copies of each trial: the original trial, and an augmented trial in which spiking activity was replaced with trial-averaged rates for that trial’s condition (appropriately normalized such that binned rates matched the scale of binned spike counts).

Thus, during training, each method was exposed both to the original trials (whose neural data match the statistics of neural data expected during testing) and the augmented trials that, in theory, might benefit decoding by exposing the decoder to denoised rates (i.e. providing each method with access to the same trial-averaged data that MINT leverages). In practice, decoding performance declined for each of the non-MINT decoders when trained this way, presumably due to the mismatch it created between the statistics of inputs during training and testing.

Neural Latents Benchmark

All neural state estimation results in **Figure 7** came from the Neural Latents Benchmark (<https://neurallatents.github.io/>), a neural latent variable modeling competition released through the 2021 NeurIPS Datasets & Benchmarks Track [104]. For each of 7 datasets, the benchmark provided training data containing simultaneous spiking activity and behavioral measurements. The neurons in the training data were divided into two sets: held-in neurons and held-out neurons. For the test data, the spiking activity of the held-in neurons was provided, but the spiking activity of the held-out neurons and the behavioral data was kept private by the benchmark curators. These splits are provided in **Table 5**. The evaluation epochs used are the same as those used in **Table 4**. (For DMFC_RSG, the evaluation epoch was the 1500 ms leading up to the ‘go’ time.) In the test data, the held-in spiking activity was only provided for time points within the evaluation epochs. For each dataset, submissions consisted of rate estimates for every trial in the training and test sets (held-in rates and held-out rates). The benchmark was fundamentally acausal: rate estimates at a given moment could leverage spiking activity from the entire trial. The data in **Figure 7** reflect all benchmark submissions up through Feb. 24th, 2023.

MINT submissions

MINT first learned idealized neural trajectories during training as described in Learning idealized trajectories. Then, the neural trajectories were partitioned by neurons into two libraries: a held-in library of trajectories and a held-out library of trajectories. The held-in library contained neural states for the held-in neurons only. The held-out library similarly contained neural states for the held-out neurons only. MINT was run using the held-in library of trajectories to estimate a neural state in the held-in neural state space. A direct association between held-in neural states and held-out neural states (the same direct association typically used to map neural states to behavioral states) was then used generate rate estimates for the held-out neurons. Rate estimates were saturated such that they could not be lower than 0.1 spikes/second. Interpolation was utilized across indices and conditions (with modifications for DMFC_RSG and MC_RTT as described in Other forms of interpolation).

The window length used for acausal neural state estimation on each dataset is provided in **Table 5**. Note that test set data was not provided outside of the evaluation epochs. Thus, despite estimating rates acausally, the spiking observations often could not be centered on the estimated time. For example, when estimating the neural state for Area2_Bump, the window length was 500 ms, but the held-in spiking activity only spanned a 600 ms period. Thus, only neural states in the 250-350 ms portion of this 600 ms period could be estimated based on 500 ms of centered spiking activity. The neural states outside of this 250-350 ms zone were estimated by either propagating the earliest estimate backward or the latest estimate forward. This propagation is possible because each neural state estimate either occurs on a trajectory with a unique past and future or is an interpolation between trajectories that each have a unique past and future (i.e. the interpolation parameters are frozen and the states being interpolated are propagated backward or forward along their trajectories).

In addition to requiring test set rate estimates, the benchmark required training set rate estimates (these were needed for the velocity R^2 metric, see Evaluation metrics). Each moment in each trial of the training data corresponded to a particular condition c and a time within the execution of that condition k (with the exception of MC_RTT, which will be discussed in a moment). Thus, training rate estimates were simply constructed by assigning each moment within each training trial a vector of rates $\bar{x}_k^c$ for the corresponding c and k . For MC_RTT, the training rate estimates were simply the single-trial rates learned via AutoLFADS.

Table 5.

Details for neural state estimation results. Note that the training trial counts match the total number of trials reported in Table 2. This reflects that the Neural Latents Benchmark utilized an additional set of test trials not reflected in the Table 2 trial counts. The test trials used for this analysis have ground truth behavior hidden by the benchmark creators and are therefore only suitable for this analysis.

Dataset	Training trials	Test trials	Held-in neurons	Held-out neurons	Window length (ms)
Area2_Bump	364	98	49	16	500
DMFC_RSG	1006	283	40	14	1500
MC_Maze	2295	574	137	45	500
MC_Maze-L	500	100	122	40	500
MC_Maze-M	250	100	114	38	500
MC_Maze-S	100	100	107	35	500
MC_RTT	1080	271	98	32	500

Table 6.

Hyperparameters used for the Wiener filter. The L2 regularization term λ was optimized in the range [0, 2000], with the optimized values rounded to the closest multiple of 10. Window lengths were optimized (in 20 ms increments) in the range [200, 600] for Area2_Bump, [200, 1000] for MC_Cycle, [200, 1200] for MC_RTT, [200, 500] for MC_PacMan, and [200, 700] for MC_Maze, MC_Maze-L, MC_Maze-M, and MC_Maze-S. These ranges were determined by the structure of each dataset (e.g. Area2_Bump couldn't look back more than 600 ms from the beginning of the evaluation epoch without entering the previous trial). Window lengths are directly related to κ via Δ (e.g. $\kappa = 14$ would correspond to a window length of $\Delta(\kappa + 1) = 20(14 + 1) = 300$ ms.)

Dataset	Behavioral Group	Window Length (ms)	L2 Regularization (λ)
Area2_Bump	Position	560	350
	Velocity	320	320
	Force	600	900
	Joint Angles	300	250
	Joint Velocities	400	1410
	Muscle Lengths	380	0
	Muscle Velocities	460	940
MC_Cycle	Phase	1000	1730
	Angular Velocity	960	1470
	Position	920	1710
	Velocity	420	600
	EMG	560	1990
MC_Maze	Position	660	530
	Velocity	540	210
MC_Maze-L	Position	540	510
	Velocity	540	440
MC_Maze-M	Position	660	440
	Velocity	480	310
MC_Maze-S	Position	680	130
	Velocity	420	270
MC_RTT	Velocity	1180	1610
MC_PacMan	Force	480	470

Table 7.

Hyperparameters used for the Kalman filter. The lag (in increments of 20 ms time bins) between neural activity and behavior was optimized in the range [2, 8], corresponding to 40-160 ms, for all datasets except Area2_Bump. For Area2_Bump the lag was not optimized and was simply set to 0 due to the fact that, in a sensory area, movement precedes sensory feedback. Given that $\mathbf{x}_k$ aggregates spikes across the whole time bin, but $\mathbf{y}_k$ corresponds to the behavioral variables at the end of the time bin, the effective lag is actually half a bin (10 ms) longer — i.e. the effective range of lags considered for the non-sensory datasets was 50-170 ms.

Dataset	Behavioral Group	Lag (bins)
Area2_Bump	Position & Velocity	0
MC_Cycle	Position & Velocity	2
MC_Maze	Position & Velocity	4
MC_Maze-L	Position & Velocity	4
MC_Maze-M	Position & Velocity	4
MC_Maze-S	Position & Velocity	4
MC_RTT	Position & Velocity	2

Table 8.

Hyperparameters used for the feedforward neural network. The number of hidden layers (L) was optimized in the range [1, 15]. The number of units per hidden layer (D) was optimized in the range [50, 1000], with the optimized values rounded to the closest multiple of 10. The dropout rate was optimized in the range [0, 0.5] and the number of training epochs was optimized in the range [2, 100]. Window lengths were optimized (in 20 ms increments) in the range [200, 600] for Area2_Bump, [200, 1000] for MC_Cycle, [200, 1200] for MC_RTT, and [200, 700] for MC_Maze, MC_Maze-L, MC_Maze-M, and MC_Maze-S.

Dataset	Behavioral Group	Window Length (ms)	Hidden Layers (L)	Units/Layer (D)	Dropout Rate	Epochs
Area2_Bump	Position	440	4	690	0.20	58
	Velocity	520	5	260	0	62
	Force	240	4	560	0.01	54
	Joint Angles	320	7	480	0.15	59
	Joint Velocities	460	5	660	0.12	36
	Muscle Lengths	240	10	840	0.19	81
	Muscle Velocities	500	6	480	0.02	24
MC_Cycle	Phase	740	8	180	0.11	24
	Angular Velocity	700	8	710	0.01	43
	Position	360	9	470	0	45
	Velocity	460	5	330	0.25	55
	EMG	840	5	970	0.07	43
MC_Maze	Position	680	7	340	0.15	65
	Velocity	700	14	760	0	89
MC_Maze-L	Position	380	6	160	0.05	44
	Velocity	600	8	380	0.13	43
MC_Maze-M	Position	520	8	860	0.26	79
	Velocity	420	2	360	0.09	59
MC_Maze-S	Position	580	8	340	0.02	74
	Velocity	520	4	210	0.07	94
MC_RTT	Velocity	1040	2	700	0.15	17

Dataset	Behavioral Group	Window Length (ms)	Units (D)	Dropout Rate	Epochs
Area2_Bump	Position	560	700	0.06	12
	Velocity	340	820	0.32	3
	Force	380	380	0.27	26
	Joint Angles	480	630	0.31	8
	Joint Velocities	260	390	0.27	9
	Muscle Lengths	460	990	0.34	45
	Muscle Velocities	420	870	0.19	47
MC_Cycle	Phase	460	580	0.18	49
	Angular Velocity	480	390	0.12	27
	Position	920	760	0.06	42
	Velocity	520	170	0.43	45
	EMG	840	250	0.40	36
MC_Maze	Position	640	740	0.27	4
	Velocity	520	800	0.34	11
MC_Maze-L	Position	620	640	0.40	49
	Velocity	580	420	0.36	31
MC_Maze-M	Position	620	820	0.41	48
	Velocity	660	840	0.29	30
MC_Maze-S	Position	460	310	0.39	27
	Velocity	680	500	0.44	45
MC_RTT	Velocity	540	890	0.40	8

Table 9.

Hyperparameters used for the GRU network. The number of units (D) was optimized in the range [50, 1000], with the optimized values rounded to the closest multiple of 10. The dropout rate was optimized in the range [0, 0.5] and the number of training epochs was optimized in the range [2, 50]. Window lengths were optimized (in 20 ms increments) in the range [200, 600] for Area2_Bump, [200, 1000] for MC_Cycle, [200, 1200] for MC_RTT, and [200, 700] for MC_Maze, MC_Maze-L, MC_Maze-M, and MC_Maze-S.

Evaluation metrics

The evaluation metrics used in this paper are briefly described below. Detailed explanations can be found in Pei et al. [104 [↗](#)].

Bits per spike was used to assess how well the rate estimates for the held-out neurons on the test data matched the actual held-out spiking activity. The metric is computed by first computing the log-likelihood of the observed spikes, given rate estimates, across all neurons. Then, the log-likelihood that would have been obtained by letting the rate estimates be the neurons' mean rates is subtracted off. Finally, the metric is normalized. Positive values indicate that the held-out rate estimates are more predictive of held-out spiking activity than the mean rates.

PSTH R^2 was computed by collecting all rate estimates on the test set, sorting them by condition, and averaging rate estimates across trials of the same condition. Then the R^2 was computed between these rate-estimate-derived PSTHs and the empirical PSTHs (computed by smoothing spikes and averaging across trials within conditions). Empirical PSTHs were computed using trials from both the training and test sets. A separate R^2 value was generated for each neuron — the reported values are the average R^2 across all neurons. This metric could not be computed for MC_RTT due to lack of condition structure.

Velocity R^2 was computed by regressing lagged x- and y-velocity of the hand against estimated rates in the training data (ridge regression). Then, the linear mapping learned via regression is applied to estimated rates in the test data to generate estimated x- and y-velocity of the hand. R^2 is computed between estimated velocities and actual velocities on the test data and averaged across the x- and y-components. This metric was not computed for DMFC_RSG because the evaluated portion of the trials did not involve motion of the monkey's arm.

Hyperparameter optimization

MINT has very few hyperparameters, all of which can be readily set by hand. These hyperparameters typically relate straightforwardly to properties of the task or data, and performance is robust to their exact values (Fig. S1 [↗](#)). For example, we did not optimize bin size (Δ). Rather, we let $\Delta = 20$ ms for all analyses, reasoning that this value would reduce computation time while remaining considerably shorter than the timescale over which rates change. MINT's only other relevant hyperparameters are window length (duration of spiking history considered at each moment) and the number of candidate states to use for interpolation. The number of candidate states was simply set to two for all datasets except MC_RTT. For MC_RTT, the nature of the training data argued for considering more candidate states during the interpolation stage. Thus, the number of states was optimized using a grid search with 10-fold cross validation on the training set (using velocity decoding R^2 as an objective function to maximize). Window length was also optimized via this grid search procedure, separately for each dataset. In some cases (e.g. the maze datasets), we still chose to standardize window length across datasets because optimization yielded similar values. Because neural state estimation was acausal and decoding analyses were causal, window length was also optimized separately across these analyses. For the simulated datasets, we chose not to optimize window length at all – the maze simulations and multitask spiking network were simply set to use the same window length as MC_Maze.

This approach to hyperparameter selection contrasts with our approach for the non-MINT decoders. As described in the Appendix, the non-MINT decoders utilized a more sophisticated technique to select hyperparameters: Bayesian optimization [101 [↗](#)]. Additionally, the non-MINT decoders were allowed the flexibility of using a different window length for each behavioral group. The choice not to provide MINT with the same flexibility was not arbitrary. Rather, this

choice emphasizes a particularly useful aspect of MINT: any relevant behavioral variable can be read out from the same neural state estimate. There is no need to use separate training or decoding procedures for different behavioral variables.

The methods we used for learning neural trajectories to use with MINT also contained hyperparameters (e.g. **Table 3** [↗](#)). For all datasets involved in the Neural Latents Benchmark (except MC_RTT), these hyperparameters were optimized using the grid search procedure (using bits per spike as an objective function to maximize). These hyperparameters were then re-used for decoding analyses with no additional optimization. For MC_RTT, we utilized AutoLFADS to learn neural trajectories, which contains a built-in procedure for optimizing hyperparameters. MC_Cycle, MC_PacMan, the maze simulations, and the multitask network were not part of the Neural Latents Benchmark and therefore needed their trajectory-learning hyperparameters set in a different way. For MC_Cycle and MC_PacMan, these hyperparameters were optimized using the grid search procedure with velocity and force decoding R^2 , respectively, as the objective functions. The maze simulations simply used the same hyperparameters as MC_Maze. For the multitask network, the trajectory learning hyperparameters were set very conservatively by hand to only perform very light temporal smoothing on spikes along with standard trial-averaging.

The example decoding results in **Fig. 4** [↗](#) and the Supp. Video used the same hyperparameters that were used in generating the quantitative decoding results in **Fig. 5** [↗](#).

Supplementary Figures

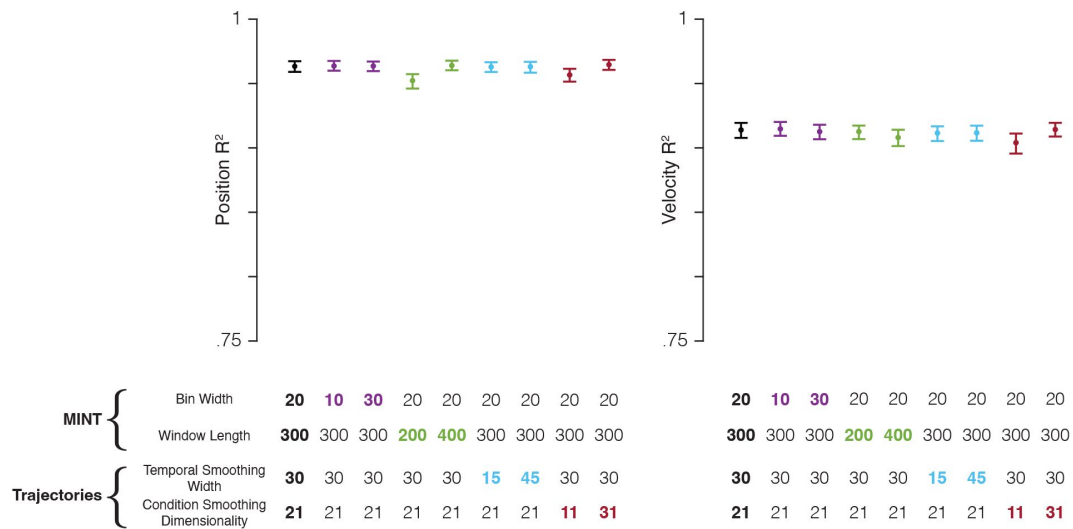


Figure S1.

MINT's decoding performance is robust to the choice of hyperparameters.

MINT was run on the MC_Maze dataset with systematic perturbations to two of MINT's hyperparameters: bin width (ms) and window length (ms). Bin width is the size of the bin in which spikes are counted: 20 ms for all analyses in the main figures. Window length is the length, in milliseconds, of spiking history that is considered. For all main analyses of the MC_Maze dataset, this was 300 ms, i.e. MINT considered the spike count in the present bin and in 14 previous bins. Perturbations were also made to two hyperparameters related to learning neural trajectories: temporal smoothing width (standard deviation of Gaussian filter applied to spikes) and condition-smoothing dimensionality (see Methods). These two hyperparameters describe how aggressively the trial-averaged data are smoothed (across time and conditions, respectively) when learning rates. Baseline decoding performance (*black circles*) was computed using the same hyperparameters that were used with the MC_Maze dataset in the analyses from [Figure 4](#) and [Figure 5](#). Then, decoding performance was computed by perturbing each of the four hyperparameters twice (*colored circles*): once to ~50% of the hyperparameter's baseline value and once to ~150%. Trials were bootstrapped (1000 resamples) to generate 95% confidence intervals (*error bars*). Perturbations of hyperparameters had little impact on performance. Altering bin width had essentially no impact, nor did altering temporal smoothing. Shortening window length had a negative impact, presumably because MINT had to estimate the neural state using fewer observations. However, the drop in performance was minimal: R^2 dropped by .011 for position decoding only. Reducing the number of dimensions used for across-condition smoothing, and consequently over-smoothing the data, had a negative impact on both position and velocity decoding. Yet again this was small: e.g. velocity R^2 dropped by .010. These results demonstrate that MINT can achieve high performance using hyperparameter values that span a large range. Thus, they do not need to be meticulously optimized to ensure good performance. In general, optimization may not be needed at all, as MINT's hyperparameters can often be set based on first principles. For example, in this study, bin width was never optimized either for specific datasets or in general. We chose to always count spikes in 20 ms bins (except in the perturbations shown here) because this duration is long enough to reduce computation time yet short relative to the timescales over which rates change. Additionally, window length can be optimized (as we did for decoding analyses), but it could also simply be chosen to roughly match the timescale over which past behavior tends to predict future behavior. Temporal smoothing of trajectories when building the library can simply use the same values commonly used when analyzing such data. For example, in prior studies, we have used smoothing kernels of width 20 to 30 ms when computing trial-averaged rates, and these values also support excellent decoding. Condition smoothing is optional and need not be applied at all, but may be useful if one wishes to record fewer trials for more conditions. For example, rather than record 15 trials for 8 reach directions, one might wish to record 5 trials for 24 conditions, then use condition smoothing to reduce sampling error.

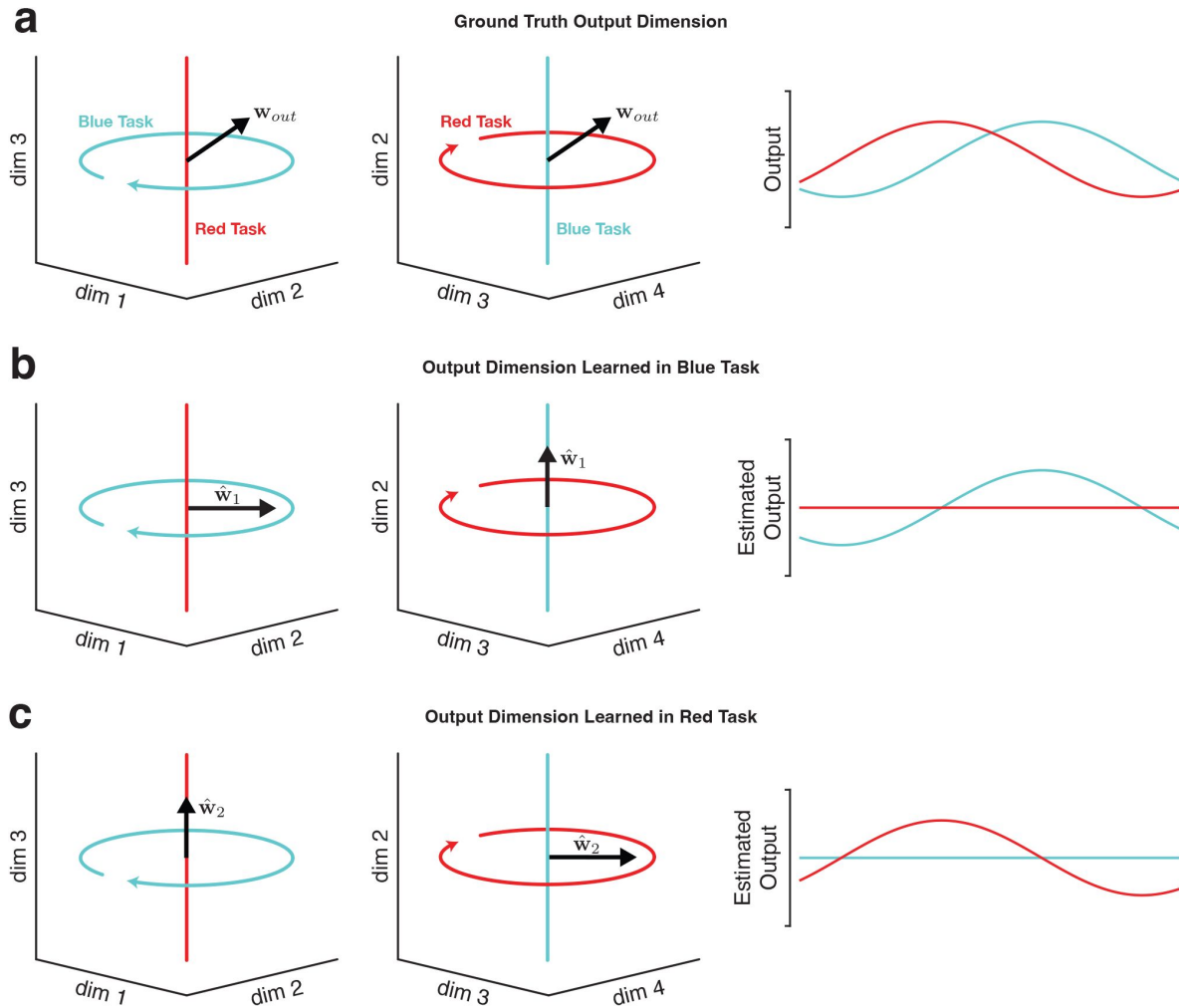


Figure S2.

Illustration of why decoding will typically fail to generalize across tasks when neural trajectories occupy orthogonal subspaces.

In theory, learning the output-potent dimensions in motor cortex would be an effective strategy for biomimetic decoding that generalizes to novel tasks. In practice, it is difficult (and often impossible) to statistically infer these dimensions without observing the subject's full behavioral repertoire (at which point generalization is no longer needed because all the relevant behaviors were directly observed). **a**) In this toy example, two neural trajectories occupy fully orthogonal subspaces. In the 'blue task', the neural trajectory occupies dimensions 1 and 2. In the 'red task', the neural trajectory occupies dimensions 3 and 4. Trajectories in both subspaces have a non-zero projection onto $\mathbf{w}_{out}$, enabling neural activity from each task to drive the appropriate output. The output at right is simply the projection of the blue-task and red-task trajectories onto $\mathbf{w}_{out}$. **b**) Illustration of the difficulty of inferring $\mathbf{w}_{out}$ from only one task. In this example, $\mathbf{w}_{out}$ is learned using data from the blue task, by linearly regressing the observed output against the neural trajectory for the blue task. The resulting estimate of $\hat{\mathbf{w}}_1$ correctly translates neural activity in the blue task into the appropriate time-varying output, but fails to generalize to the red task. This failure to generalize is a straightforward consequence of the fact that $\hat{\mathbf{w}}_1$ was learned from data that didn't explore dimensions 3 and 4. Note that this same phenomenon would occur if $\hat{\mathbf{w}}_1$ were estimated by regressing intended output versus neural activity (as might occur in a paralyzed patient) **c**) Estimating the output-potent dimension based only on the red task yields an estimate $\hat{\mathbf{w}}_2$ that fails to generalize to the blue task. This phenomenon is illustrated here for a linear readout, but would apply to most nonlinear methods as well, unless some other form of knowledge can allow interpretation of neural trajectories in previously unseen dimensions.

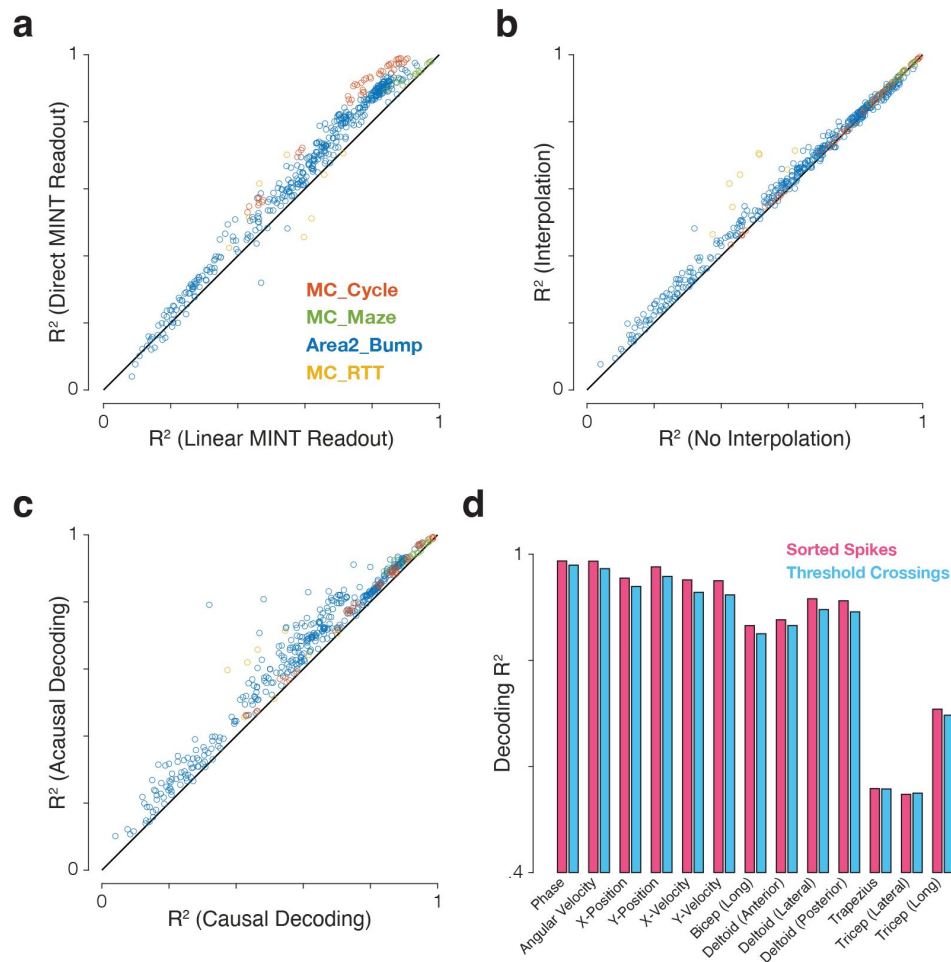


Figure S3.

Impact of different modeling and preprocessing choices on performance of MINT.

For most applications we anticipate MINT will employ direct decoding that leverages the correspondence between neural and behavioral trajectories, but one could also choose to linearly decode behavior from the estimated neural state. For most applications, we anticipate MINT will use interpolation amongst candidate neural states on neighboring trajectories, but one could also restrict decoding to states within the trajectory library. For real-time applications we anticipate MINT will be run causally, but acausal decoding (using both past and future spiking observations) could be used offline or even online by introducing a lag. We anticipate MINT may be used both in situations where spike events have been sorted based on neuron identity, and situations where decoding simply uses channel-specific unsorted threshold crossings. Panels **a-c** explore the first three choices. Performance was quantified for all 8 combinations of: direct MINT readout vs. linear MINT readout, interpolation vs. no interpolation, causal decoding vs. acausal decoding. This was done for 121 behavioral variables across 4 datasets for a total of 964 R^2 values. The 'phase' behavioral variable in MC_Cycle was excluded from 'linear MINT readout' variants because its circularity makes it a trivially poor target for linear decoding. **a)** MINT's direct neural-to-behavioral-state association outperforms a linear readout based on MINT's estimated neural state. Performance was significantly higher using the direct readout ($\Delta R^2 = .061 \pm .002$ SE; $p < .001$, paired t-test). Note that the linear readout still benefits from the ability of MINT to estimate the neural state using all neural dimensions, not just those that correlate with kinematics. **b)** Decoding with interpolation significantly outperformed decoding without interpolation ($\Delta R^2 = .018 \pm .001$ SE; $p < .001$, paired t-test). **c)** Running acausally significantly improved performance relative to causal decoding ($\Delta R^2 = .051 \pm .002$ SE; $p < .001$, paired t-test). Although causal decoding is required for real-time applications, this result suggests that (when tolerable) introducing a small decoding lag could improve performance. For example, a decoder using 200 ms of spiking activity could introduce a 50 ms lag such that the decode at time t is rendered by generating the best estimate of the behavioral state at time $t - 50$ using spiking data from $t - 200$ through t . **d)** Decoding performance for 13 behavioral variables in the MC_Cycle dataset when sorted spikes were used (112 neurons, pink) versus 'good' threshold crossings from electrodes for which the signal-to-noise ratio of the firing rates exceeded a threshold (93 electrodes, SNR > 2, cyan). The loss in performance when using threshold crossings was small ($\Delta R^2 = -.014 \pm .002$ SE). SE refers to standard error of the mean.

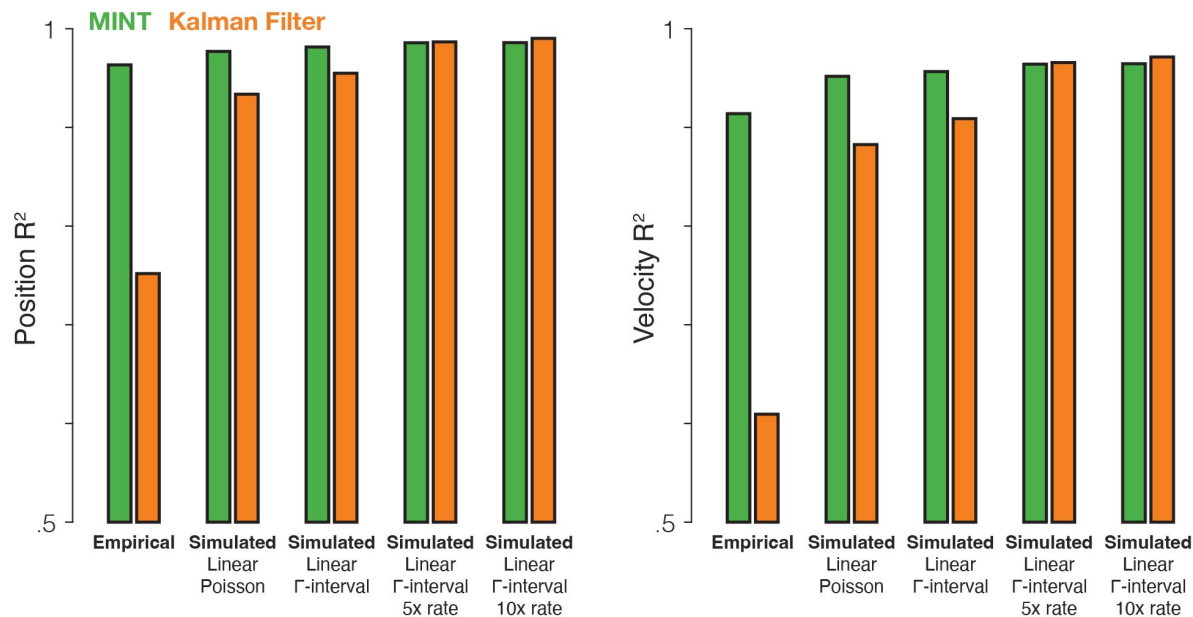


Figure S4.

The Kalman filter's relative performance would improve if neural data had different statistical properties.

We compared MINT to the Kalman filter across one empirical dataset (MC_Maze) and four simulated datasets (same behavior as MC_Maze, but simulated spikes). Simulated firing rates were linear functions of hand position, velocity, and acceleration (as is assumed by the Kalman filter). The means and standard deviations (across time and conditions) of the simulated firing rates were matched to actual neural data (with additional rate scaling in two cases). The Kalman filter assumes that observation noise is stationary and Gaussian. Although spiking variability cannot be Gaussian (spike counts must be positive integers), spiking variability can be made more stationary by letting that variability depend less on rate. Thus, although the first simulation generated spikes via a Poisson process, subsequent simulations utilized gamma-interval spiking (gamma distribution with $\alpha = 2$ and $\beta = 2\lambda$, where λ is firing rate). Gamma-interval spiking variability of this form is closer to stationary at higher rates. Thus, the third and fourth simulations scaled up firing rates to further push spiking variability into a more stationary regime at the expense of highly unrealistic rates (in the fourth simulation, a firing rate briefly exceeded 1800 Hz). Overall, as the simulated neural data better accorded with the assumptions of the Kalman filter, decoding performance for the Kalman filter improved. Interestingly, MINT continued to perform well even on the simulated data, likely because MINT can exploit a linear relationship between neural and behavioral states when the data argue for it and higher rates benefit both algorithms. These results demonstrate that algorithms like MINT and the Kalman filter are not intrinsically good or bad. Rather, they are best suited to data that match their assumptions. When simulated data approximate the assumptions of the Kalman filter, both methods perform similarly. However, MINT shows much higher performance for the empirical data, suggesting that its assumptions are a better match for the statistical properties of the data.

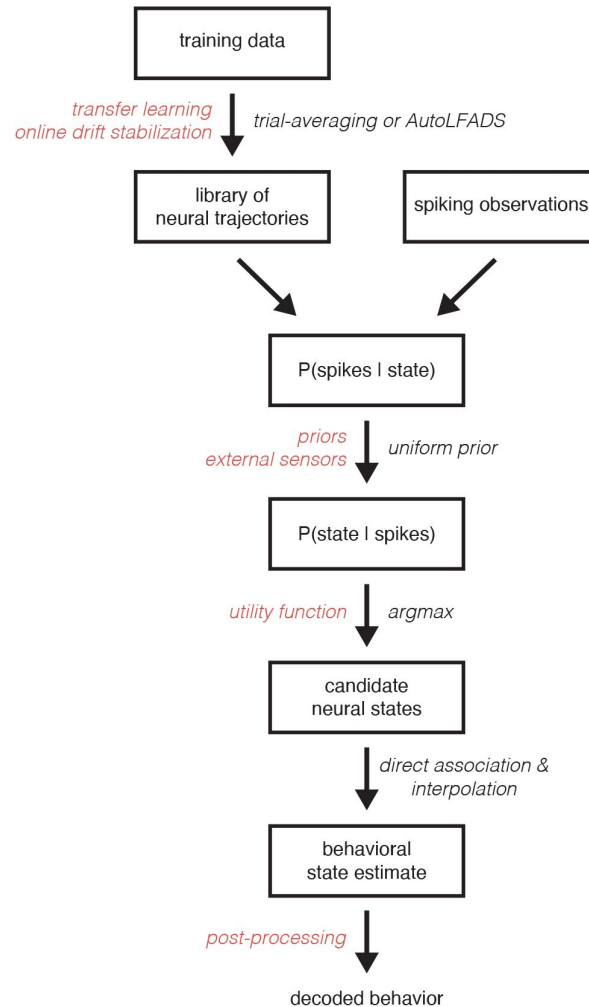


Figure S5.

MINT is a modular algorithm for which a variety of modifications and extensions exist.

The flowchart illustrates the standard MINT algorithm in *black* and lists potential changes to the algorithm in *red*. For example, the library of neural trajectories is typically learned via standard trial-averaging or a single-trial rate estimation technique like AutoLFADS. However, transfer learning could be utilized to learn the library of trajectories based on trajectories from a different subject or session. The trajectories could also be modified online while MINT is running to reflect changes in spiking activity that relate to recording instabilities. A potential modification to the method also occurs when likelihoods are converted into posterior probabilities. Typically, we assume a uniform prior over states in the library. However, that prior could be set to reflect the relative frequency of different behaviors and could even incorporate time-varying information from external sensors (e.g. state of a prosthetic limb, eye tracking, etc.) that include information about how probable each behavior is at a given moment. Another potential extension of MINT occurs at the stage where candidate neural states are selected. Typically, these states are selected based solely on spike count likelihoods. However, one could use a utility function that reflects a user's values (e.g. the user may dislike some decoding mistakes more than others) in conjunction with the likelihoods to maximize expected utility. Lastly, the behavioral estimate that MINT returns could be post-processed (e.g. temporally smoothed). MINT's modularity largely derives from the fact that the library of neural trajectories is finite. This assumption enables posterior probabilities to be directly computed for each state, rather than analytically derived. Thus, choices like how to learn the library of trajectories, which observation model to use (e.g. Poisson vs. generalized Poisson), and which (if any) state priors to use, can be made independently from one another. These choices all interact to impact performance. Yet they will not interact to impact tractability, as would have been the case if one analytically derived a continuous posterior probability distribution.

Appendix Decoder implementations

All decoders described in this section were provided with the same training and testing data as MINT. When decoding, spiking activity was binned every 20 ms for each neuron (with the exception of the Naïve Bayes regression decoder, which benefitted from a larger bin size). Each decoder was given access to recent binned spike counts to use for decoding. The following sections will: 1) provide an overview of the notation and definitions needed to understand the subsequent equations (notation *will not* in general match the notation used for MINT), 2) describe the hyperparameter optimization technique used for all decoders in this appendix, and 3) describe each decoder's implementation in detail.

Definitions

Time Bins

Each decoder received binned spiking activity and decoded behavior at time-bin resolution. These time bins are indexed by k . When higher-resolution estimates were required (e.g. for evaluating decoding performance), the decoded variables were simply upsampled with a zero-order-hold.

Observations

The vector of spike counts at time bin k is denoted by the column vector $\mathbf{x}_k \in \mathbb{R}^N$ where N is the number of neurons. When decoding, each decoder was given access to the current time bin of spike counts and κ previous time bins of spike counts. κ was not a fixed value and varied depending on the method, dataset, and behavioral group being decoded.

Target Variables

The vector of behavioral variables at time bin k is denoted by the column vector $\mathbf{y}_k \in \mathbb{R}^M$ where M is the number of behavioral variables. This corresponds to the values of the behavioral variables at the end of the time bin.

Decoded Variables

The vector of decoded behavioral variables at time bin k is denoted $\hat{\mathbf{y}}_k$.

Training Data

Observations and target variables at time bin i of trial j in the training data are denoted $\mathbf{x}_i^{(j)}$ and $\mathbf{y}_i^{(j)}$.

Hyperparameter Optimization

The hyperparameters for each decoder were set using Bayesian optimization [101]. The training set was split into a reduced training set (80% of the trials from the full training set) and a validation set (the remaining 20% of trials). The method is provided with a range of hyperparameter values to search. The method then seeks to learn a set of hyperparameters that, when trained on the reduced training set, will lead to maximal decoding R^2 on the validation set. The optimization occurs via an iterative process that involves exploring the space of hyperparameters and exploiting knowledge of how certain hyperparameter sets performed in previous iterations. In all cases, we set the method to initially perform 10 iterations of random exploration of hyperparameter space, followed by 10 iterations of Bayesian optimization. The

optimization was performed separately for each behavioral group (e.g. position vs. velocity) within each dataset. In the decoder-specific sections below, the exact hyperparameter values learned for each dataset and behavioral group are reported along with the hyperparameter ranges that were searched.

Wiener Filter

The Wiener filter [1] uses a model

$$\mathbf{y}_k = b + \sum_{i=0}^{\kappa} W_i \mathbf{x}_{k-i} + \epsilon_k \quad (\text{S1})$$

where $b \in \mathbb{R}^M$, $W_i \in \mathbb{R}^{M \times N}$, and the residuals are $\epsilon_k \in \mathbb{R}^M$. Decoding occurs via the equation

$$\hat{\mathbf{y}}_k = W \begin{bmatrix} \mathbf{x}_k \\ \mathbf{x}_{k-1} \\ \vdots \\ \mathbf{x}_{k-\kappa} \\ 1 \end{bmatrix} \quad (\text{S2})$$

for some matrix W that minimizes the squared error of the residuals in the training data.

Parameter fitting

For each trial j in the training data, we construct matrices

$$\mathbf{X}^{(j)} = \begin{bmatrix} \mathbf{x}_{\kappa+1}^{(j)} & \mathbf{x}_{\kappa+2}^{(j)} & \cdots & \mathbf{x}_{T_j}^{(j)} \\ \mathbf{x}_{\kappa}^{(j)} & \mathbf{x}_{\kappa+1}^{(j)} & \cdots & \mathbf{x}_{T_j-1}^{(j)} \\ \vdots & \vdots & \ddots & \vdots \\ \mathbf{x}_1^{(j)} & \mathbf{x}_2^{(j)} & \cdots & \mathbf{x}_{T_j-\kappa}^{(j)} \\ 1 & 1 & \cdots & 1 \end{bmatrix} \quad (\text{S3})$$

$$\mathbf{Y}^{(j)} = \begin{bmatrix} \mathbf{y}_{\kappa+1}^{(j)} & \mathbf{y}_{\kappa+2}^{(j)} & \cdots & \mathbf{y}_{T_j}^{(j)} \end{bmatrix} \quad (\text{S4})$$

where T_j is the number of time bins in trial j . Then, we concatenate across J trials to get

$$\mathbf{X} = [\mathbf{X}^{(1)} \quad \mathbf{X}^{(2)} \quad \cdots \quad \mathbf{X}^{(J)}] \quad (\text{S5})$$

$$\mathbf{Y} = [\mathbf{Y}^{(1)} \quad \mathbf{Y}^{(2)} \quad \cdots \quad \mathbf{Y}^{(J)}] \quad (\text{S6})$$

W can then be fit via a regularized regression,

$$W = \mathbf{YX}^T (\mathbf{XX}^T + \lambda I)^{-1} \quad (\text{S7})$$

Kalman filter

The Kalman filter [114] uses a model

$$\mathbf{y}_k = A\mathbf{y}_{k-1} + \mathbf{q}_k \quad (\text{S8})$$

$$\mathbf{x}_k = C\mathbf{y}_k + \mathbf{r}_k \quad (\text{S9})$$

where $\mathbf{q}_k \in \mathcal{N}(0, Q)$ and $\mathbf{r}_k \in \mathcal{N}(0, R)$. Decoding occurs via the standard update equations:

$$\hat{\mathbf{y}}_k^- = A\hat{\mathbf{y}}_{k-1} \quad (\text{S10})$$

$$\hat{\mathbf{y}}_k = \hat{\mathbf{y}}_k^- + K_k(\mathbf{x}_k - C\hat{\mathbf{y}}_k^-) \quad (\text{S11})$$

where

$$P_k^- = AP_{k-1}A^\top + Q \quad (\text{S12})$$

$$K_k = P_k^- C^\top (CP_k^- C^\top + R)^{-1} \quad (\text{S13})$$

$$P_k = P_k^- - K_k CP_k^- \quad (\text{S14})$$

The state vector $\mathbf{y}_k$ includes 7 elements: position, velocity, and acceleration (x- and y-components of each) as well as a constant 1. The 1 was included to allow firing rates to have a constant offset. Acceleration was included as in [50] to improve the position and velocity estimates.

Parameter fitting

For each trial j in the training set, we construct matrices

$$\mathbf{X}^{(j)} = \begin{bmatrix} \mathbf{x}_1^{(j)} & \mathbf{x}_2^{(j)} & \cdots & \mathbf{x}_{T_j - lag}^{(j)} \end{bmatrix} \quad (\text{S15})$$

$$\mathbf{Y}^{(j)} = \begin{bmatrix} \mathbf{y}_{1+lag}^{(j)} & \mathbf{y}_{2+lag}^{(j)} & \cdots & \mathbf{y}_{T_j}^{(j)} \end{bmatrix} \quad (\text{S16})$$

$$\mathbf{Y}_1^{(j)} = \begin{bmatrix} \mathbf{y}_{1+lag}^{(j)} & \mathbf{y}_{2+lag}^{(j)} & \cdots & \mathbf{y}_{T_j-1}^{(j)} \end{bmatrix} \quad (\text{S17})$$

$$\mathbf{Y}_2^{(j)} = \begin{bmatrix} \mathbf{y}_{2+lag}^{(j)} & \mathbf{y}_{3+lag}^{(j)} & \cdots & \mathbf{y}_{T_j}^{(j)} \end{bmatrix} \quad (\text{S18})$$

where T_j is the number of time bins in trial j . Then, we concatenate across J trials to get

$$\mathbf{X} = [\mathbf{X}^{(1)} \quad \mathbf{X}^{(2)} \quad \cdots \quad \mathbf{X}^{(J)}] \quad (\text{S19})$$

$$\mathbf{Y} = [\mathbf{Y}^{(1)} \quad \mathbf{Y}^{(2)} \quad \cdots \quad \mathbf{Y}^{(J)}] \quad (\text{S20})$$

$$\mathbf{Y}_1 = [\mathbf{Y}_1^{(1)} \quad \mathbf{Y}_1^{(2)} \quad \cdots \quad \mathbf{Y}_1^{(J)}] \quad (\text{S21})$$

$$\mathbf{Y}_2 = [\mathbf{Y}_2^{(1)} \quad \mathbf{Y}_2^{(2)} \quad \cdots \quad \mathbf{Y}_2^{(J)}] \quad (\text{S22})$$

The parameters can then be fit as follows:

[eq

$$A = \mathbf{Y}_2 \mathbf{Y}_1^\top (\mathbf{Y}_1 \mathbf{Y}_1^\top)^{-1} \quad (\text{S23})$$

$$C = \mathbf{X} \mathbf{Y}^\top (\mathbf{Y} \mathbf{Y}^\top)^{-1} \quad (\text{S24})$$

$$Q = \frac{(\mathbf{Y}_2 - A \mathbf{Y}_1)(\mathbf{Y}_2 - A \mathbf{Y}_1)^\top}{T_1} \quad (\text{S25})$$

$$R = \frac{(\mathbf{X} - C \mathbf{Y})(\mathbf{X} - C \mathbf{Y})^\top}{T} \quad (\text{S26})$$

$$\hat{\mathbf{y}}_1^- = \frac{1}{J} \sum_{j=1}^J \mathbf{y}_{1+lag}^{(j)} \quad (\text{S27})$$

$$P_1^- = \frac{1}{J} \sum_{j=1}^J (\mathbf{y}_{1+lag}^{(j)} - \hat{\mathbf{y}}_1^-)(\mathbf{y}_{1+lag}^{(j)} - \hat{\mathbf{y}}_1^-)^\top \quad (\text{S28})$$

where T_1 and T are the number of columns in $\mathbf{Y}_1$ and $\mathbf{Y}$, respectively.

Feedforward Neural Network

The feedforward neural network [115] we use takes spike count observations (current and previous time bins) and flattens them into a long vector,

$$\mathbf{f}_k = [\mathbf{x}_k^\top \quad \mathbf{x}_{k-1}^\top \quad \dots \quad \mathbf{x}_{k-\kappa}^\top]^\top \quad (\text{S29})$$

centers and normalizes that vector,

$$\tilde{\mathbf{f}}_k = (\mathbf{f}_k - \mu) \oslash \sigma \quad (\text{S30})$$

and then feeds it through multiple hidden network layers,

$$\mathbf{h}_k^1 = \text{ReLU}(W_1 \tilde{\mathbf{f}}_k + b_1) \quad (\text{S31})$$

$$\mathbf{h}_k^l = \text{ReLU}(W_l \mathbf{h}_k^{l-1} + b_l), \quad l \in \{2, \dots, L\} \quad (\text{S32})$$

before finally reading out behavior,

$$\hat{\mathbf{y}}_k = W_{out} \mathbf{h}_k^L + b_{out} \quad (\text{S33})$$

where $\mathbf{h}_k^l \in \mathbb{R}^D$, D is the number of units per hidden layer, and L is the number of hidden layers.

Parameter fitting

For each trial j in the training set, we create flattened observations for all $i > \kappa$ (sufficient spike count history doesn't exist for $i \leq \kappa$). We then let μ and σ be the element-wise mean and standard deviation of $\mathbf{f}^{(j)}$ across all observations in the training set (i.e. across all time bins in all J trials for

which $i > \kappa$). The parameters W_{out} , b_{out} , W_l and b_l for $l \in \{1, \dots, L\}$ are then learned by training the network with the Adam optimization routine [116] and a mean-squared error loss function. Training utilized dropout on the outputs of each hidden layer.

GRU

The GRU neural network [117] we use takes spike count observations (current and previous time bins), centers and normalizes those observations,

$$\mathbf{f}_i^{(j)} = \left[(\mathbf{x}_i^{(j)})^\top \quad (\mathbf{x}_{i-1}^{(j)})^\top \quad \dots \quad (\mathbf{x}_{i-\kappa}^{(j)})^\top \right]^\top \quad (\text{S34})$$

and then feeds the observations sequentially into the network (initializing with $\mathbf{h}_{k-\kappa-1} = \mathbf{0}$),

$$\tilde{\mathbf{x}}_i = (\mathbf{x}_i - \mu) \oslash \sigma, \quad i \in \{k - \kappa, \dots, k\} \quad (\text{S35})$$

$$\mathbf{z}_i = \text{sigmoid}(W_z \tilde{\mathbf{x}}_i + U_z \mathbf{h}_{i-1} + b_z) \quad (\text{S36})$$

$$\mathbf{r}_i = \text{sigmoid}(W_r \tilde{\mathbf{x}}_i + U_r \mathbf{h}_{i-1} + b_r) \quad (\text{S37})$$

$$\hat{\mathbf{h}}_i = \tanh(W_h \tilde{\mathbf{x}}_i + U_h (\mathbf{r}_i \odot \mathbf{h}_{i-1}) + b_h) \quad (\text{S38})$$

$$\mathbf{h}_i = \mathbf{z}_i \odot \mathbf{h}_{i-1} + (1 - \mathbf{z}_i) \odot \hat{\mathbf{h}}_i \quad (\text{S39})$$

ultimately reading out behavior from the final state,

$$\hat{\mathbf{y}}_k = W_{out} \mathbf{h}_k + b_{out} \quad (\text{S40})$$

where $\mathbf{z}_i \in \mathbb{R}^D$, $\mathbf{r}_i \in \mathbb{R}^D$, $\hat{\mathbf{h}}_i \in \mathbb{R}^D$, $\mathbf{h}_i \in \mathbb{R}^D$, and D is the number of GRU units. The GRU hidden states do not persist from one decoding time bin to the next. Rather, at each decoding time bin, the GRU is re-initialized with $\mathbf{h}_{k-\kappa-1} = \mathbf{0}$ and run sequentially over recent history to generate an estimate of the behavior at the current time bin.

Parameter fitting

μ and σ (both column vectors of length N) are computed as the mean and standard deviation, respectively, of the observed spike counts for each neuron in the training set. The parameters W_z , U_z , b_z , W_r , U_r , b_r , W_h , U_h , b_h , W_{out} and b_{out} are then learned by training the network with the RMSProp optimization routine [118] and a mean-squared error loss function. Training utilized dropout both on the linear transformation of inputs and on the linear transformation of the recurrent state.

Naïve Bayes regression

The Naïve Bayes regression decoder [87, 119] is fully described with code provided in Glaser et al. [87]. The hyperparameters for this decoder include the density of the kinematic grid, bin size, and the number of previous bins available to the decoder. This decoder was only evaluated for the MC_Maze dataset. The best performance was achieved for a 50×50 grid of kinematics. For both position and velocity, hyperparameter optimization indicated use of only a single bin of spike counts (220 ms and 320 ms bin sizes for position and velocity, respectively).

References

- [1] Carmena Jose M., et al. (2003) **Learning to Control a Brain–Machine Interface for Reaching and Grasping by Primates** *PLoS Biology* **1** <https://doi.org/10.1371/journal.pbio.0000042>
- [2] Hochberg Leigh R., et al. (2006) **Neuronal ensemble control of prosthetic devices by a human with tetraplegia** *Nature* **442**:164–171 <https://doi.org/10.1038/nature04970>
- [3] Velliste Meel, et al. (2008) **Cortical control of a prosthetic arm for self-feeding** *Nature* **453**:1098–1101 <https://doi.org/10.1038/nature06996>
- [4] Hochberg Leigh R., et al. (2012) **Reach and grasp by people with tetraplegia using a neurally controlled robotic arm** *Nature* **485**:372–375 <https://doi.org/10.1038/nature11076>
- [5] Collinger Jennifer L., et al. (2013) **High-performance neuroprosthetic control by an individual with tetraplegia** *The Lancet* **381**:557–564 [https://doi.org/10.1016/s0140-6736\(12\)61816-9](https://doi.org/10.1016/s0140-6736(12)61816-9)
- [6] Wodlinger B, et al. (2015) **Ten-dimensional anthropomorphic arm control in a human brain-machine interface: difficulties, solutions, and limitations** *Journal of Neural Engineering* **12** <https://doi.org/10.1088/1741-2560/12/1/016011>
- [7] Moritz Chet T., Perlmutter Steve I., Fetz Eberhard E. (2008) **Direct control of paralysed muscles by cortical neurons** *Nature* **456**:639–642 <https://doi.org/10.1038/nature07418>
- [8] Ethier C., et al. (2012) **Restoration of grasp following paralysis through brain-controlled stimulation of muscles** *Nature* **485**:368–371 <https://doi.org/10.1038/nature10987>
- [9] Bouton Chad E., et al. (2016) **Restoring cortical control of functional movement in a human with quadriplegia** *Nature* **533**:247–250 <https://doi.org/10.1038/nature17435>
- [10] Ajiboye A. Bolu, et al. (2017) **Restoration of reaching and grasping movements through brain-controlled muscle stimulation in a person with tetraplegia: a proof-of-concept demonstration** *The Lancet* **389**:1821–1830 [https://doi.org/10.1016/s0140-6736\(17\)30601-3](https://doi.org/10.1016/s0140-6736(17)30601-3)
- [11] Taylor Dawn M., Tillery Stephen I. Helms, Schwartz Andrew B. (2002) **Direct Cortical Control of 3D Neuroprosthetic Devices** *Science*
- [12] Serruya Mijail D., et al. (2002) **Instant neural control of a movement signal** *Nature* **416**:141–142 <https://doi.org/10.1038/416141a>
- [13] Gilja Vikash, et al. (2012) **A high-performance neural prosthesis enabled by control algorithm design** *Nature Neuroscience* **15**:1752–1757 <https://doi.org/10.1038/nn.3265>
- [14] Jarosiewicz Beata, et al. (2015) **Virtual typing by people with tetraplegia using a self-calibrating intracortical brain-computer interface** *Science Translational Medicine* **7** <https://doi.org/10.1126/scitranslmed.aac7328>

- [15] Nuyujukian Paul, et al. (2015) **A High-Performance Keyboard Neural Prosthesis Enabled by Task Optimization** *IEEE Transactions on Biomedical Engineering* **62**:21–29 <https://doi.org/10.1109/tbme.2014.2354697>
- [16] Shanechi Maryam M., et al. (2017) **Rapid control and feedback rates enhance neuroprosthetic control** *Nature Communications* **8** <https://doi.org/10.1038/ncomms13825>
- [17] Pandarinath Chethan, et al. (2017) **High performance communication by people with paralysis using an intracortical brain-computer interface** *eLife* **6** <https://doi.org/10.7554/elife.18554>
- [18] Libedinsky Camilo, et al. (2016) **Independent Mobility Achieved through a Wireless Brain-Machine Interface** *PLOS ONE* **11** <https://doi.org/10.1371/journal.pone.0165773>
- [19] Rajangam Sankaranarayani, et al. (2016) **Wireless Cortical Brain-Machine Interface for Whole-Body Navigation in Primates** *Scientific Reports* **6** <https://doi.org/10.1038/srep22170>
- [20] Schroeder Karen E, et al. (2022) **Cortical Control of Virtual Self-Motion Using Task-Specific Subspaces** *The Journal of Neuroscience* **42**:220–239 <https://doi.org/10.1523/jneurosci.2687-20.2021>
- [21] Anumanchipalli Gopala K., Chartier Josh, Chang Edward F. (2019) **Speech synthesis from neural decoding of spoken sentences** *Nature* **568**:493–498 <https://doi.org/10.1038/s41586-019-1119-1>
- [22] Wilson Guy H., et al. (2020) **Decoding spoken English from intracortical electrode arrays in dorsal precentral gyrus** *Journal of Neural Engineering* **17** <https://doi.org/10.1088/1741-2552/abbfef>
- [23] Willett Francis R., et al. (2023) **A high-performance speech neuroprosthesis** *Nature* **620**:1031–1036 <https://doi.org/10.1038/s41586-023-06377-x>
- [24] Metzger Sean L., et al. (2023) **A high-performance neuroprosthesis for speech decoding and avatar control** *Nature* **620**:1037–1046 <https://doi.org/10.1038/s41586-023-06443-4>
- [25] Wairagkar Maitreyee, et al. (2023) **Synthesizing Speech by Decoding Intracortical Neural Activity from Dorsal Motor Cortex** *2023 11th International IEEE/EMBS Conference on Neural Engineering (NER)* 0:1–4 <https://doi.org/10.1109/ner52421.2023.10123880>
- [26] Card Nicholas S., et al. (2024) **An accurate and rapidly calibrating speech neuroprosthesis** *medRxiv* <https://doi.org/10.1101/2023.12.26.23300110>
- [27] Willett Francis R., et al. (2021) **High-performance brain-to-text communication via handwriting** *Nature* **593**:249–254 <https://doi.org/10.1038/s41586-021-03506-2>
- [28] Musallam S., et al. (2004) **Cognitive Control Signals for Neural Prosthetics** *Science* **305**:258–262 <https://doi.org/10.1126/science.1097938>
- [29] Wallis Joni D. (2018) **Decoding Cognitive Processes from Neural Ensembles** *Trends in Cognitive Sciences* **22**:1091–1102 <https://doi.org/10.1016/j.tics.2018.09.002>
- [30] Sani Omid G., et al. (2018) **Mood variations decoded from multi-site intracranial human brain activity** *Nature Biotechnology* **36**:954–961 <https://doi.org/10.1038/nbt.4200>

- [31] Yousefi Ali, et al. (2019) **Decoding Hidden Cognitive States From Behavior and Physiology Using a Bayesian Approach** *Neural Computation* **31**:1751–1788 https://doi.org/10.1162/neco_a_01196
- [32] Provenza Nicole R., et al. (2019) **Decoding task engagement from distributed network electrophysiology in humans** *Journal of Neural Engineering* **16** <https://doi.org/10.1088/1741-2552/ab2c58>
- [33] Chari Aswin, et al. (2020) **Microelectrode recordings in human epilepsy: A case for clinical translation?** *Brain Communications* **2** <https://doi.org/10.1093/braincomms/fcaa082>
- [34] Cowley Benjamin R., et al. (2013) **DataHigh: graphical user interface for visualizing and interacting with high-dimensional neural activity** *Journal of Neural Engineering* **10** <https://doi.org/10.1088/1741-2560/10/6/066012>
- [35] Peixoto Diogo, et al. (2021) **Decoding and perturbing decision states in real time** *Nature* **591**:604–609 <https://doi.org/10.1038/s41586-020-03181-9>
- [36] Sadtler Patrick T., et al. (2014) **Neural constraints on learning** *Nature* **512**:423–426 <https://doi.org/10.1038/nature13665>
- [37] Golub Matthew D., et al. (2018) **Learning by neural reassociation** *Nature Neuroscience* **21**:607–616 <https://doi.org/10.1038/s41593-018-0095-3>
- [38] Gallego Juan A., et al. (2017) **Neural Manifolds for the Control of Movement** *Neuron* **94**:978–984 <https://doi.org/10.1016/j.neuron.2017.05.025>
- [39] Gallego Juan A., et al. (2018) **Cortical population activity within a preserved neural manifold underlies multiple motor behaviors** *Nature Communications* **9** <https://doi.org/10.1038/s41467-018-06560-z>
- [40] Weiss Jeffrey M, et al. (2019) **Demonstration of a portable intracortical brain-computer interface** *Brain-Computer Interfaces* **6**:106–117
- [41] Georgopoulos Apostolos P., Schwartz Andrew B., Kettner Ronald E. (1986) **Neuronal Population Coding of Movement Direction** *Science* **233**:1416–1419 <https://doi.org/10.1126/science.3749885>
- [42] Schwartz Andrew B. (1994) **Direct Cortical Representation of Drawing** *Science* **265**:540–542 <https://doi.org/10.1126/science.8036499>
- [43] Kemere Caleb T., et al. (2002) **Decoding of Plan and Peri-Movement Neural Signals in Prosthetic Systems** *IEEE Workshop on Signal Processing Systems* :276–283 <https://doi.org/10.1109/sips.2002.1049722>
- [44] Kao Jonathan C., et al. (2017) **A High-Performance Neural Prosthesis Incorporating Discrete State Selection With Hidden Markov Models** *IEEE Transactions on Biomedical Engineering* **64**:935–945 <https://doi.org/10.1109/tbme.2016.2582691>
- [45] Brockwell A. E., Rojas A. L., Kass R. E. (2004) **Recursive Bayesian Decoding of Motor Cortical Signals by Particle Filtering** *Journal of Neurophysiology* **91**:1899–1907 <https://doi.org/10.1152/jn.00438.2003>

- [46] Kemere Caleb, Sahanil Maneesh, Meng Teresa (2003) **Robust Neural Decoding of Reaching Movements for Prosthetic Systems** *Proceedings of the 25th Annual International Conference of the IEEE Engineering in Medicine and Biology Society (IEEE Cat. No. 03CH37439)* **3**:2079–2082 <https://doi.org/10.1109/iembs.2003.1280146>
- [47] Kemere C., Shenoy K. V., Meng T. H. (2004) **Model-based neural decoding of reaching movements: a maximum likelihood approach** *IEEE Transactions on Biomedical Engineering* **51**:925–932 <https://doi.org/10.1109/tbme.2004.826675>
- [48] Kemere Caleb, et al. (2004) **Model-Based Decoding of Reaching Movements for Prosthetic Systems** *The 26th Annual International Conference of the IEEE Engineering in Medicine and Biology Society* **2**:4524–4528 <https://doi.org/10.1109/iembs.2004.1404256>
- [49] Yu Byron M., et al. (2007) **Mixture of Trajectory Models for Neural Decoding of Goal-Directed Movements** *Journal of Neurophysiology* **97**:3763–3780 <https://doi.org/10.1152/jn.00482.2006>
- [50] Wu W., et al. (2003) **Neural Decoding of Cursor Motion Using a Kalman Filter** *Advances in Neural Information Processing Systems 15 (NIPS 2002)* **15**
- [51] Sani Omid G., et al. (2021) **Modeling behaviorally relevant neural dynamics enabled by preferential subspace identification** *Nature Neuroscience* **24**:140–149 <https://doi.org/10.1038/s41593-020-00733-0>
- [52] Fortunato Cátia, et al. (2024) **Nonlinear manifolds underlie neural population activity during behaviour** *bioRxiv* <https://doi.org/10.1101/2023.07.18.549575>
- [53] Saxena Shreya, et al. (2022) **Motor cortex activity across movement speeds is predicted by network-level strategies for generating muscle activity** *eLife* **11** <https://doi.org/10.7554/elife.67620.sa0>
- [54] Vyas Saurabh, et al. (2020) **Computation Through Neural Population Dynamics** *Annual Review of Neuroscience* **43**:249–275 <https://doi.org/10.1146/annurev-neuro-092619-094115>
- [55] DePasquale Brian, et al. (2023) **The centrality of population-level factors to network computation is demonstrated by a versatile approach for training spiking networks** *Neuron* <https://doi.org/10.1016/j.neuron.2022.12.007>
- [56] Churchland Mark M., Shenoy Krishna V. (2024) **Preparatory activity and the expansive null-space** *Nature Reviews Neuroscience* **25**:213–236 <https://doi.org/10.1038/s41583-024-00796-z>
- [57] Churchland Mark M, et al. (2012) **Neural population dynamics during reaching** *Nature* **487**:51–56 <https://doi.org/10.1038/nature11129>
- [58] Mante Valerio, et al. (2013) **Context-dependent computation by recurrent dynamics in prefrontal cortex** *Nature* **503**:78–84 <https://doi.org/10.1038/nature12742>
- [59] Sussillo David, et al. (2015) **A neural network that finds a naturalistic solution for the production of muscle activity** *Nature Neuroscience* **18**:1025–1033 <https://doi.org/10.1038/nn.4042>
- [60] Remington Evan D., et al. (2018) **Flexible Sensorimotor Computations through Rapid Reconfiguration of Cortical Dynamics** *Neuron* **98**:1005–1019 <https://doi.org/10.1016/j.neuron.2018.05.020>

- [61] Russo Abigail A., et al. (2018) **Motor Cortex Embeds Muscle-like Commands in an Untangled Population Response** *Neuron* **97**:953–966 <https://doi.org/10.1016/j.neuron.2018.01.004>
- [62] Sohn Hansem, et al. (2019) **Bayesian Computation through Cortical Latent Dynamics** *Neuron* **103**:934–947 <https://doi.org/10.1016/j.neuron.2019.06.012>
- [63] Russo Abigail A., et al. (2020) **Neural Trajectories in the Supplementary Motor Area and Motor Cortex Exhibit Distinct Geometries, Compatible with Different Classes of Computation** *Neuron* **107**:745–758 <https://doi.org/10.1016/j.neuron.2020.05.020>
- [64] Stopfer Mark, Jayaraman Vivek, Laurent Gilles (2003) **Intensity versus Identity Coding in an Olfactory System** *Neuron* **39**:991–1004 <https://doi.org/10.1016/j.neuron.2003.08.011>
- [65] Mishne Gal, et al. (2016) **Hierarchical Coupled-Geometry Analysis for Neuronal Structure and Activity Pattern Discovery** *IEEE Journal of Selected Topics in Signal Processing* **10**:1238–1253 <https://doi.org/10.1109/jstsp.2016.2602061>
- [66] Goudar Vishwa, Buonomano Dean V (2018) **Encoding sensory and motor patterns as time-invariant trajectories in recurrent neural networks** *eLife* **7** <https://doi.org/10.7554/elife.31134>
- [67] Chaudhuri Rishidev, et al. (2019) **The intrinsic attractor manifold and population dynamics of a canonical cognitive circuit across waking and sleep** *Nature Neuroscience* **22**:1512–1520 <https://doi.org/10.1038/s41593-019-0460-x>
- [68] Zhou Ding, Wei Xue-Xin (2020) **Learning identifiable and interpretable latent models of high-dimensional neural activity using pi-VAE** *arXiv* <https://doi.org/10.48550/arxiv.2011.04798>
- [69] Schneider Steffen, Lee Jin Hwa, Mathis Mackenzie Weygandt (2023) **Learnable latent embeddings for joint behavioural and neural analysis** *Nature* **617**:360–368 <https://doi.org/10.1038/s41586-023-06031-6>
- [70] Churchland Mark M., Shenoy Krishna V. (2007) **Temporal Complexity and Heterogeneity of Single-Neuron Activity in Premotor and Motor Cortex** *Journal of Neurophysiology* **97**:4235–4257 <https://doi.org/10.1152/jn.00095.2007>
- [71] Seely Jeffrey S., et al. (2016) **Tensor Analysis Reveals Distinct Population Structure that Parallels the Different Computational Roles of Areas M1 and V1** *PLoS Computational Biology* **12** <https://doi.org/10.1371/journal.pcbi.1005164>
- [72] Gao Peiran, Ganguli Surya (2015) **On simplicity and complexity in the brave new world of large-scale neuroscience** *Current Opinion in Neurobiology* **32**:148–155 <https://doi.org/10.1016/j.conb.2015.04.003>
- [73] Gao Peiran, et al. (2017) **A theory of multineuronal dimensionality, dynamics and measurement** *bioRxiv* <https://doi.org/10.1101/214262>
- [74] Marshall Najja J., et al. (2022) **Flexible neural control of motor units** *Nature Neuroscience* :1–13 <https://doi.org/10.1038/s41593-022-01165-8>
- [75] Elsayed Gamaleldin F., et al. (2016) **Reorganization between preparatory and movement population responses in motor cortex** *Nature Communications* **7** <https://doi.org/10.1038/ncomms13239>

- [76] Miri Andrew, et al. (2017) **Behaviorally Selective Engagement of Short-Latency Effector Pathways by Motor Cortex** *Neuron* **95**:683–696 <https://doi.org/10.1016/j.neuron.2017.06.042>
- [77] Warriner Claire L., et al. (2022) **Motor cortical influence relies on task-specific activity covariation** *Cell Reports* **40** <https://doi.org/10.1016/j.celrep.2022.111427>
- [78] Xing David, Truccolo Wilson, Borton David A. (2022) **Emergence of Distinct Neural Subspaces in Motor Cortical Dynamics during Volitional Adjustments of Ongoing Locomotion** *The Journal of Neuroscience* **42**:9142–9157 <https://doi.org/10.1523/jneurosci.0746-22.2022>
- [79] Ames Katherine Cora, Churchland Mark M. (2019) **Motor cortex signals for each arm are mixed across hemispheres and neurons yet partitioned within the population response** *eLife* **8** <https://doi.org/10.7554/elife.46159>
- [80] Athalye Vivek R., et al. (2023) **Invariant neural dynamics drive commands to control different movements** *Current Biology* **33**:2962–2976 <https://doi.org/10.1016/j.cub.2023.06.027>
- [81] Oby Emily R., et al. (2024) **Dynamical constraints on neural population activity** *bioRxiv* <https://doi.org/10.1101/2024.01.03.573543>
- [82] Oby Emily R., et al. (2019) **New neural activity patterns emerge with long-term learning** *Proceedings of the National Academy of Sciences* **116**:15210–15215 <https://doi.org/10.1073/pnas.1820296116>
- [83] Kao Jonathan C., et al. (2015) **Single-trial dynamics of motor cortex and their applications to brain-machine interfaces** *Nature Communications* **6** <https://doi.org/10.1038/ncomms8759>
- [84] Pandarinath Chethan, et al. (2018) **Inferring single-trial neural population dynamics using sequential auto-encoders** *Nature Methods* **15**:805–815 <https://doi.org/10.1038/s41592-018-0109-9>
- [85] Brennan Connor, et al. (2023) **One dimensional approximations of neuronal dynamics reveal computational strategy** *PLOS Computational Biology* **19** <https://doi.org/10.1371/journal.pcbi.1010784>
- [86] Schwemmer Michael A., et al. (2018) **Meeting brain–computer interface user performance expectations using a deep neural network decoding framework** *Nature Medicine* **24**:1669–1676 <https://doi.org/10.1038/s41591-018-0171-y>
- [87] Glaser Joshua I., et al., ENEURO.0506–19.2020 (2020) **Machine learning for neural decoding** *eNeuro* **7** <https://doi.org/10.1523/eneuro.0506-19.2020>
- [88] Sussillo David, et al. (2012) **A recurrent neural network for closed-loop intracortical brain–machine interface decoders** *Journal of Neural Engineering* **9** <https://doi.org/10.1088/1741-2560/9/2/026027>
- [89] Sussillo David, et al. (2016) **Making brain–machine interfaces robust to future neural variability** *Nature Communications* **7** <https://doi.org/10.1038/ncomms13749>
- [90] Makin Joseph G., et al. (2018) **Superior arm-movement decoding from cortex with a new, unsupervised-learning algorithm** *Journal of Neural Engineering* **15** <https://doi.org/10.1088/1741-2552/aa9e95>

- [91] Tseng Po-He, et al. (2019) **Decoding Movements from Cortical Ensemble Activity Using a Long Short-Term Memory Recurrent Network** *Neural Computation* **31**:1085–1113 https://doi.org/10.1162/neco_a_01189
- [92] Ye Joel, Pandarinath Chethan (2021) **Representation learning for neural population activity with Neural Data Transformers** *Neurons, Behavior, Data Analysis, and Theory* <https://doi.org/10.51628/001c.27358>
- [93] Ahmadi Nur, Constandinou Timothy G., Bouganis Christos-Savvas (2021) **Robust and accurate decoding of hand kinematics from entire spiking activity using deep learning** *Journal of Neural Engineering* **18** <https://doi.org/10.1088/1741-2552/abde8a>
- [94] Fetz Eberhard E. (1992) **Are movement parameters recognizably coded in the activity of single neurons?** *Behavioral and Brain Sciences* **15**:679–690
- [95] Todorov Emanuel (2000) **Direct cortical control of muscle activation in voluntary arm movements: a model** *Nature Neuroscience* **3**:391–398 <https://doi.org/10.1038/73964>
- [96] Scott Stephen H. (2008) **Inconvenient Truths about neural processing in primary motor cortex** *The Journal of Physiology* **586**:1217–1224 <https://doi.org/10.1113/jphysiol.2007.146068>
- [97] Reimer Jacob, Hatsopoulos Nicholas G (2009) **Progress in Motor Control, A Multidisciplinary Perspective** *Advances in Experimental Medicine and Biology* **629**:243–259 https://doi.org/10.1007/978-0-387-77064-2_12
- [98] Churchland Mark M., et al. (2010) **Cortical Preparatory Activity: Representation of Movement or First Cog in a Dynamical Machine?** *Neuron* **68** <https://doi.org/10.1016/j.neuron.2010.09.015>
- [99] Chowdhury Raeed H, Glaser Joshua I, Miller Lee E (2020) **Area 2 of primary somatosensory cortex encodes kinematics of the whole arm** *eLife* **9** <https://doi.org/10.7554/elife.48198>
- [100] Wessberg Johan, et al. (2000) **Real-time prediction of hand trajectory by ensembles of cortical neurons in primates** *Nature* **408**:361–365 <https://doi.org/10.1038/35042582>
- [101] Snoek Jasper, Larochelle Hugo, Adams Ryan P. (2012) **Practical Bayesian Optimization of Machine Learning Algorithms** *Advances in Neural Information Processing Systems 25 (NIPS 2012)* **25**
- [102] Trautmann Eric, et al. (2022) **Motor Cortex Isolates Skill-Specific Dynamics in a Context Switching Task** *Computational and Systems Neuroscience (COSYNE) Abstracts*
- [103] Christie Breanne P, et al. (2015) **Comparison of spike sorting and thresholding of voltage waveforms for intracortical brain-machine interface performance** *Journal of Neural Engineering* **12** <https://doi.org/10.1088/1741-2560/12/1/016009>
- [104] Pei Felix, et al. (2021) **Neural Latents Benchmark '21: Evaluating latent variable models of neural population activity** *arXiv* <https://doi.org/10.48550/arxiv.2109.04463>
- [105] Yu Byron M., et al. (2009) **Gaussian-process factor analysis for low-dimensional single-trial analysis of neural population activity** *Advances in Neural Information Processing Systems 21 (NIPS 2008)* **21**

- [106] Fox Emily, et al. (2008) **Bayesian Nonparametric Inference of Switching Dynamic Linear Models** *EEE Transactions on Signal Processing* **59**
- [107] Keshtkaran Mohammad Reza, et al. (2022) **A large-scale neural network training framework for generalized estimation of single-trial population dynamics** *Nature Methods* :1–6 <https://doi.org/10.1038/s41592-022-01675-0>
- [108] Shlizerman Trung Le and Eli (2022) **STNDT: Modeling Neural Population Activity with a Spatiotemporal Transformer** *36th Conference on Neural Information Processing Systems (NeurIPS 2022)*. **35** <https://doi.org/10.48550/arxiv.2206.04727>
- [109] Deo Darrel R., et al. (2023) **Translating deep learning to neuroprosthetic control** *bioRxiv* <https://doi.org/10.1101/2023.04.21.537581>
- [110] Vargas-Irwin Carlos E., et al. (2018) **Watch, Imagine, Attempt: Motor Cortex Single-Unit Activity Reveals Context-Dependent Movement Encoding in Humans With Tetraplegia** *Frontiers in Human Neuroscience* **12** <https://doi.org/10.3389/fnhum.2018.00450>
- [111] Williams Alex H., et al. (2020) **Discovering Precise Temporal Patterns in Large-Scale Neural Recordings through Robust and Interpretable Time Warping** *Neuron* **105**:246–259 <https://doi.org/10.1016/j.neuron.2019.10.020>
- [112] Dyer Eva L., et al. (2017) **A cryptography-based approach for movement decoding** *Nature Biomedical Engineering* **1**:967–976 <https://doi.org/10.1038/s41551-017-0169-7>
- [113] Pachitariu Marius, et al. (2024) **Spike sorting with Kilosort4** *Nature Methods* **21**:914–921 <https://doi.org/10.1038/s41592-024-02232-7>
- [114] Kalman R. E. (1960) **A New Approach to Linear Filtering and Prediction Problems** *Journal of Basic Engineering* **82**:35–45 <https://doi.org/10.1115/1.3662552>
- [115] Rosenblatt Frank (1961) **Principles of neurodynamics. Perceptrons and the theory of brain mechanisms. Tech. rep.** Buffalo, NY: Cornell Aeronautical Laboratory
- [116] Kingma Diederik P, Ba Jimmy (2014) **Adam: A Method for Stochastic Optimization** *arXiv* <https://doi.org/10.48550/arxiv.1412.6980>
- [117] Cho Kyunghyun, et al. (2014) **Learning Phrase Representations using RNN Encoder-Decoder for Statistical Machine Translation** *arXiv* <https://doi.org/10.48550/arxiv.1406.1078>
- [118] Tieleman Tijmen, Hinton Geoffrey (2012) **Lecture 6.5-rmsprop: Divide the gradient by a running average of its recent magnitude** *COURSERA: Neural networks for machine learning*
- [119] Zhang Kechen, et al. (1998) **Interpreting Neuronal Population Activity by Reconstruction: Unified Framework With Application to Hippocampal Place Cells** *Journal of Neurophysiology* **79**:1017–1044 <https://doi.org/10.1152/jn.1998.79.2.1017>

Author information

Sean M Perkins

Department of Biomedical Engineering, Columbia University, New York, USA, Zuckerman Institute, Columbia University, New York, USA
ORCID iD: [0000-0001-9456-4648](https://orcid.org/0000-0001-9456-4648)

Elom A Amematsro

Zuckerman Institute, Columbia University, New York, USA, Department of Neuroscience, Columbia University Medical Center, New York, USA
ORCID iD: [0000-0003-4843-4513](https://orcid.org/0000-0003-4843-4513)

John P Cunningham

Zuckerman Institute, Columbia University, New York, USA, Department of Statistics, Columbia University, New York, USA, Center for Theoretical Neuroscience, Columbia University Medical Center, New York, USA, Grossman Center for the Statistics of Mind, Columbia University, New York, USA

Qi Wang

Department of Biomedical Engineering, Columbia University, New York, USA
ORCID iD: [0000-0001-8656-1439](https://orcid.org/0000-0001-8656-1439)

Mark M Churchland

Zuckerman Institute, Columbia University, New York, USA, Department of Neuroscience, Columbia University Medical Center, New York, USA, Grossman Center for the Statistics of Mind, Columbia University, New York, USA, Kavli Institute for Brain Science, Columbia University Medical Center, New York, USA
ORCID iD: [0000-0001-9123-6526](https://orcid.org/0000-0001-9123-6526)

For correspondence: mc3502@columbia.edu

Editors

Reviewing Editor

Mackenzie Mathis

École Polytechnique Fédérale de Lausanne, Genève, Switzerland

Senior Editor

Joshua Gold

University of Pennsylvania, Philadelphia, United States of America

Reviewer #1 (Public Review):

Summary:

This paper presents an innovative decoding approach for brain-computer interfaces (BCIs), introducing a new method named MINT. The authors develop a trajectory-centric approach to decode behaviors across several different datasets, including eight empirical datasets from the Neural Latents Benchmark. Overall, the paper is well written and their method shows impressive performance compared to more traditional decoding approaches that use a simpler approach. While there are some concerns (see below), the paper's strengths, particularly its emphasis on a trajectory-centric approach and the simplicity of MINT, provide a compelling contribution to the field.

Strengths:

The adoption of a trajectory-centric approach that utilizes statistical constraints presents a substantial shift in methodology, potentially revolutionizing the way BCIs interpret and predict neural behaviour. This is one of the strongest aspects of the paper.

The thorough evaluation of the method across various datasets serves as an assurance that the superior performance of MINT is not a result of overfitting. The comparative simplicity of the method in contrast to many neural network approaches is refreshing and should facilitate broader applicability.

Weaknesses:

Scope: Despite the impressive performance of MINT across multiple datasets, it seems predominantly applicable to M1/S1 data. Only one of the eight empirical datasets comes from an area outside the motor/somatosensory cortex. It would be beneficial if the authors could expand further on how the method might perform with other brain regions that do not exhibit low tangling or do not have a clear trial structure (e.g. decoding of position or head direction from hippocampus)

When comparing methods, the neural trajectories of MINT are based on averaged trials, while the comparison methods are trained on single trials. An additional analysis might help in disentangling the effect of the trial averaging. For this, the authors could average the input across trials for all decoders, establishing a baseline for averaged trials. Note that inference should still be done on single trials. Performance can then be visualized across different values of N , which denotes the number of averaged trials used for training.

Comments on revisions:

I have looked at the responses and they are thorough and answer all of my questions.

<https://doi.org/10.7554/eLife.89421.2.sa2>

Reviewer #2 (Public Review):

Summary:

The goal of this paper is to present a new method, termed MINT, for decoding behavioral states from neural spiking data. MINT is a statistical method which, in addition to outputting a decoded behavioral state, also provides soft information regarding the likelihood of that behavioral state based on the neural data. The innovation in this approach is neural states are assumed to come from sparsely distributed neural trajectories with low tangling, meaning that neural trajectories (time sequences of neural states) are sparse in the high-dimensional space of neural spiking activity and that two dissimilar neural trajectories tend to correspond to dissimilar behavioral trajectories. The authors support these assumptions through analysis of previously collected data, and then validate the performance of their method by comparing it to a suite of alternative approaches. The authors attribute the typically improved decoding performance by MINT to its assumptions being more faithfully aligned to the properties of neural spiking data relative to assumptions made by the alternatives.

Strengths:

The paper did an excellent job critically evaluating common assumptions made by neural analytical methods, such as neural state being low-dimensional relative to the number of recorded neurons. The authors made strong arguments, supported by evidence and literature, for potentially high-dimensional neural states and thus the need for approaches that do not rely on an assumption of low dimensionality.

The paper was thorough in considering multiple datasets across a variety of behaviors, as well as existing decoding methods, to benchmark the MINT approach. This provided a valuable comparison to validate the method. The authors also provided nice intuition regarding why MINT may offer performance improvement in some cases and in which instances MINT may not perform as well.

In addition to providing a philosophical discussion as to the advantages of MINT and benchmarking against alternatives, the authors also provided a detailed description of practical considerations. This included training time, amount of training data, robustness to data loss or changes in the data, and interpretability. These considerations not only provided objective evaluation of practical aspects but also provided insights to the flexibility and robustness of the method as they relate back to the underlying assumptions and construction of the approach.

Impact:

This work is motivated by brain-computer interfaces applications, which it will surely impact in terms of neural decoder design. However, this work is also broadly impactful for neuroscientific analysis to relate neural spiking activity to observable behavioral features. Thus, MINT will likely impact neuroscience research generally. The methods are made publicly available, and the datasets used are all in public repositories, which facilitates adoption and validation of this method within the greater scientific community.

<https://doi.org/10.7554/eLife.89421.2.sa1>

Author response:

The following is the authors' response to the original reviews.

Summary of reviewers' comments and our revisions:

We thank the reviewers for their thoughtful feedback. This feedback has motivated multiple revisions and additions that, in our view, have greatly improved the manuscript. This is especially true with regard to a major goal of this study: clearly defining existing scientific perspectives and delineating their decoding implications. In addition to building on this conceptual goal, we have expanded existing analyses and have added a new analysis of generalization using a newly collected dataset. We expect the manuscript will be of very broad interest, both to those interested in BCI development and to those interested in fundamental properties of neural population activity and its relationship with behavior.

Importantly, all reviewers were convinced that MINT provided excellent performance, when benchmarked against existing methods, across a broad range of standard tasks:

“their method shows impressive performance compared to more traditional decoding approaches” (R1)

“The paper was thorough in considering multiple datasets across a variety of behaviors, as well as existing decoding methods, to benchmark the MINT approach. This provided a valuable comparison to validate the method.” (R2)

“The fact that performance on stereotyped tasks is high is interesting and informative...” (R3)

This is important. It is challenging to design a decoder that performs consistently across multiple domains and across multiple situations (including both decoding and neural state estimation). MINT does so. MINT consistently outperformed existing lightweight ‘interpretable’ decoders, despite being a lightweight interpretable decoder itself. MINT was

very competitive with expressive machine-learning methods, yet has advantages in flexibility and simplicity that more ‘brute force’ methods do not. We made a great many comparisons, and MINT was consistently a strong performer. Of the many comparisons we made, there was only one where MINT was at a modest disadvantage, and it was for a dataset where all methods performed poorly. No other method we tested was as consistent. For example, although the GRU and the feedforward network were often competitive with MINT (and better than MINT in the one case mentioned above), there were multiple other situations where they performed less well and a few situations where they performed poorly. Moreover, no other existing decoder naturally estimates the neural state while also readily decoding, without retraining, a broad range of behavioral variables.

R1 and R2 were very positive about the broader impacts of the study. They stressed its impact both on decoder design, and on how our field thinks, scientifically, about the population response in motor areas:

“This paper presents an innovative decoding approach for brain-computer interfaces” (R1)

“presents a substantial shift in methodology, potentially revolutionizing the way BCIs interpret and predict neural behaviour” (R1)

“the paper's strengths, particularly its emphasis on a trajectory-centric approach and the simplicity of MINT, provide a compelling contribution to the field” (R1)

“The authors made strong arguments, supported by evidence and literature, for potentially high-dimensional neural states and thus the need for approaches that do not rely on an assumption of low dimensionality” (R2)

“This work is motivated by brain-computer interfaces applications, which it will surely impact in terms of neural decoder design.” (R2)

“this work is also broadly impactful for neuroscientific analysis... Thus, MINT will likely impact neuroscience research generally.” (R2)

We agree with these assessments, and have made multiple revisions to further play into these strengths. As one example, the addition of Figure 1b (and 6b) makes this the first study, to our knowledge, to fully and concretely illustrate this emerging scientific perspective and its decoding implications. This is important, because multiple observations convince us that the field is likely to move away from the traditional perspective in Figure 1a, and towards that in Figure 1b. We also agree with the handful of weaknesses R1 and R2 noted. The manuscript has been revised accordingly. The major weakness noted by R1 was the need to be explicit regarding when we suspect MINT would (and wouldn't) work well in other brain areas. In non-motor areas, the structure of the data may be poorly matched with MINT's assumptions. We agree that this is likely to be true, and thus agree with the importance of clarifying this topic for the reader. The revision now does so. R1 also wished to know whether existing methods might benefit from including trial-averaged data during training, something we now explore and document (see detailed responses below). R2 noted two weaknesses: 1) The need to better support (with expanded analysis) the statement that neural and behavioral trajectories are non-isometric, and 2) The need to more rigorously define the ‘mesh’. We agree entirely with both suggestions, and the revision has been strengthened by following them (see detailed responses below).

R3 also saw strengths to the work, stating that:

“This paper is well-structured and its main idea is clear.”

“The fact that performance on stereotyped tasks is high is interesting and informative, showing that these stereotyped tasks create stereotyped neural trajectories.”

“The task-specific comparisons include various measures and a variety of common decoding approaches, which is a strength.”

However, R3 also expressed two sizable concerns. The first is that MINT might have onerous memory requirements. The manuscript now clarifies that MINT has modest memory requirements. These do not scale unfavorably as the reviewer was concerned they might. The second concern is that MINT is:

“essentially a table-lookup rather than a model.”

Although we don’t agree, the concern makes sense and may be shared by many readers, especially those who take a particular scientific perspective. Pondering this concern thus gave us the opportunity to modify the manuscript in ways that support its broader impact. Our revisions had two goals: 1) clarify the ways in which MINT is far more flexible than a lookup-table, and 2) better describe the dominant scientific perspectives and their decoding implications.

The heart of R3’s concern is the opinion that MINT is an effective but unprincipled hack suitable for situations where movements are reasonably stereotyped. Of course, many tasks involve stereotyped movements (e.g. handwriting characters), so MINT would still be useful. Nevertheless, if MINT is not principled, other decode methods would often be preferable because they could (unlike MINT in R3’s opinion) gain flexibility by leveraging an accurate model. Most of R3’s comments flow from this fundamental concern:

“This is again due to MINT being a lookup table with a library of stereotyped trajectories rather than a model.”

“MINT models task-dependent neural trajectories, so the trained decoder is very task-dependent and cannot generalize to other tasks.”

“Unlike MINT, these works can achieve generalization because they model the neural subspace and its association to movement.”

“given that MINT tabulates task-specific trajectories, it will not generalize to tasks that are not seen in the training data even when these tasks cover the exact same space (e.g., the same 2D computer screen and associated neural space).”

“For proper training, the training data should explore the whole movement space and the associated neural space, but this does not mean all kinds of tasks performed in that space must be included in the training set (something MINT likely needs while modeling-based approaches do not).”

The manuscript has been revised to clarify that MINT is considerably more flexible than a lookup table, even though a lookup table is used as a first step. Yet, on its own, this does not fully address R3’s concern. The quotes above highlight that R3 is making a standard assumption in our field: that there exists a “movement space and associated neural space”. Under this perspective, one should, as R3 argues fully explore the movement space. This would perforce fully explore the associated neural subspace. One can then “model the neural subspace and its association to movement”. MINT does not use a model of this type, and thus (from R3’s perspective) does not appear to use a model at all. A major goal of our study is to question this traditional perspective. We have thus added a new figure to highlight the contrast between the traditional (Figure 1a) and new (Figure 1b) scientific perspectives, and to clarify their decoding implications.

While we favor the new perspective (Figure 1b), we concede that R3 may not share our view. This is fine. Part of the reason we believe this study is timely, and will be broadly read, is that it raises a topic of emerging interest where there is definitely room for debate. If we are

misguided – i.e. if Figure 1a is the correct perspective – then many of R3’s concerns would be on target: MINT could still be useful, but traditional methods that make the traditional assumptions in Figure 1a would often be preferable. However, if the emerging perspective in Figure 1b is more accurate, then MINT’s assumptions would be better aligned with the data than those of traditional methods, making it a more (not less) principled choice.

Our study provides new evidence in support of Figure 1b, while also synthesizing existing evidence from other recent studies. In addition to Figure 2, the new analysis of generalization further supports Figure 1b. Also supporting Figure 1b is the analysis in which MINT’s decoding advantage, over a traditional decoder, disappears when simulated data approximate the traditional perspective in Figure 1a.

That said, we agree that the present study cannot fully resolve whether Figure 1a or 1b is more accurate. Doing so will take multiple studies with different approaches (indeed we are currently preparing other manuscripts on this topic). Yet we still have an informed scientific opinion, derived from past, present and yet-to-be-published observations. Our opinion is that Figure 1b is the more accurate perspective. This possibility makes it reasonable to explore the potential virtues of a decoding method whose assumptions are well-aligned with that perspective. MINT is such a method. As expected under Figure 1b, MINT outperforms traditional interpretable decoders in every single case we studied.

As noted above, we have added a new generalization-focused analysis (Figure 6) based on a newly collected dataset. We did so because R3’s comments highlight a deep point: which scientific perspective one takes has strong implications regarding decoder generalization. These implications are now illustrated in the new Figure 6a and 6b. Under Figure 6a, it is possible, as R3 suggests, to explore “the whole movement space and associated neural space” during training. However, under Figure 6b, expectations are very different. Generalization will be ‘easy’ when new trajectories are near the training-set trajectories. In this case, MINT should generalize well as should other methods. In contrast, generalization will be ‘hard’ when new neural trajectories have novel shapes and occupy previously unseen regions / dimensions. In this case, all current methods, including MINT, are likely to fail. R3 points out that traditional decoders have sometimes generalized well to new tasks (e.g. from center-out to ‘pinball’) when cursor movements occur in the same physical workspace. These findings could be taken to support Figure 6a, but are equally consistent with ‘easy’ generalization in Figure 6b. To explore this topic, the new analysis in Figure 6c-g considers conditions that are intended to span the range from easy to hard. Results are consistent with the predictions of Figure 6b.

We believe the manuscript has been significantly improved by these additions. The revisions help the manuscript achieve its twin goals: 1) introduce a novel class of decoder that performs very well despite being very simple, and 2) describe properties of motor-cortex activity that will matter for decoders of all varieties.

Reviewer #1:

Summary:

This paper presents an innovative decoding approach for brain-computer interfaces (BCIs), introducing a new method named MINT. The authors develop a trajectory-centric approach to decode behaviors across several different datasets, including eight empirical datasets from the Neural Latents Benchmark. Overall, the paper is well written and their method shows impressive performance compared to more traditional decoding approaches that use a simpler approach. While there are some concerns (see below), the paper’s strengths, particularly its emphasis on a trajectory-centric approach and the simplicity of MINT, provide a compelling contribution to the field.

We thank the reviewer for these comments. We share their enthusiasm for the trajectory-centric approach, and we are in complete agreement that this perspective has both scientific and decoding implications. The revision expands upon these strengths.

Strengths:

The adoption of a trajectory-centric approach that utilizes statistical constraints presents a substantial shift in methodology, potentially revolutionizing the way BCIs interpret and predict neural behaviour. This is one of the strongest aspects of the paper.

Again, thank you. We also expect the trajectory-centric perspective to have a broad impact, given its relevance to both decoding and to thinking about manifolds.

The thorough evaluation of the method across various datasets serves as an assurance that the superior performance of MINT is not a result of overfitting. The comparative simplicity of the method in contrast to many neural network approaches is refreshing and should facilitate broader applicability.

Thank you. We were similarly pleased to see such a simple method perform so well. We also agree that, while neural-network approaches will always be important, it is desirable to also possess simple ‘interpretable’ alternatives.

Weaknesses:

Comment 1) Scope: Despite the impressive performance of MINT across multiple datasets, it seems predominantly applicable to M1/S1 data. Only one of the eight empirical datasets comes from an area outside the motor/somatosensory cortex. It would be beneficial if the authors could expand further on how the method might perform with other brain regions that do not exhibit low tangling or do not have a clear trial structure (e.g. decoding of position or head direction from hippocampus)

We agree entirely. Population activity in many brain areas (especially outside the motor system) presumably will often not have the properties upon which MINT’s assumptions are built. This doesn’t necessarily mean that MINT would perform badly. Using simulated data, we have found that MINT can perform surprisingly well even when some of its assumptions are violated. Yet at the same time, when MINT’s assumptions don’t apply, one would likely prefer to use other methods. This is, after all, one of the broader themes of the present study: it is beneficial to match decoding assumptions to empirical properties. We have thus added a section on this topic early in the Discussion:

“In contrast, MINT and the Kalman filter performed comparably on simulated data that better approximated the assumptions in Figure 1a. Thus, MINT is not a ‘better’ algorithm – simply better aligned with the empirical properties of motor cortex data. This highlights an important caveat. Although MINT performs well when decoding from motor areas, its assumptions may be a poor match in other areas (e.g. the hippocampus). MINT performed well on two non-motor-cortex datasets – Area2_Bump (S1) and DMFC_RSG (dorsomedial frontal cortex) – yet there will presumably be other brain areas and/or contexts where one would prefer a different method that makes assumptions appropriate for that area.”

Comment 2) When comparing methods, the neural trajectories of MINT are based on averaged trials, while the comparison methods are trained on single trials. An additional analysis might help in disentangling the effect of the trial averaging. For this, the authors could average the input across trials for all decoders, establishing a baseline for averaged trials. Note that inference should still be done on single trials. Performance can

then be visualized across different values of N , which denotes the number of averaged trials used for training.

We explored this question and found that the non-MINT decoders are harmed, not helped, by the inclusion of trial-averaged responses in the training set. This is presumably because the statistics of trial-averaged responses don't resemble what will be observed during decoding. This statistical mismatch, between training and decoding, hurts most methods. It doesn't hurt MINT, because MINT doesn't 'train' in the normal way. It simply needs to know rates, and trial-averaging is a natural way to obtain them. To describe the new analysis, we have added the following to the text.

"We also investigated the possibility that MINT gained its performance advantage simply by having access to trial-averaged neural trajectories during training, while all other methods were trained on single-trial data. This difference arises from the fundamental requirements of the decoder architectures: MINT needs to estimate typical trajectories while other methods don't. Yet it might still be the case that other methods would benefit from including trial-averaged data in the training set, in addition to single-trial data. Alternatively, this might harm performance by creating a mismatch, between training and decoding, in the statistics of decoder inputs. We found that the latter was indeed the case: all non-MINT methods performed better when trained purely on single-trial data."

Reviewer #2:

Summary:

The goal of this paper is to present a new method, termed MINT, for decoding behavioral states from neural spiking data. MINT is a statistical method which, in addition to outputting a decoded behavioral state, also provides soft information regarding the likelihood of that behavioral state based on the neural data. The innovation in this approach is neural states are assumed to come from sparsely distributed neural trajectories with low tangling, meaning that neural trajectories (time sequences of neural states) are sparse in the high-dimensional space of neural spiking activity and that two dissimilar neural trajectories tend to correspond to dissimilar behavioral trajectories. The authors support these assumptions through analysis of previously collected data, and then validate the performance of their method by comparing it to a suite of alternative approaches. The authors attribute the typically improved decoding performance by MINT to its assumptions being more faithfully aligned to the properties of neural spiking data relative to assumptions made by the alternatives.

We thank the reviewer for this accurate summary, and for highlighting the subtle but important fact that MINT provides information regarding likelihoods. The revision includes a new analysis (Figure 6e) illustrating one potential way to leverage knowledge of likelihoods.

Strengths:

The paper did an excellent job critically evaluating common assumptions made by neural analytical methods, such as neural state being low-dimensional relative to the number of recorded neurons. The authors made strong arguments, supported by evidence and literature, for potentially high-dimensional neural states and thus the need for approaches that do not rely on an assumption of low dimensionality.

Thank you. We also hope that the shift in perspective is the most important contribution of the study. This shift matters both scientifically and for decoder design. The revision expands on this strength. The scientific alternatives are now more clearly and concretely illustrated

(especially see Figure 1a,b and Figure 6a,b). We also further explore their decoding implications with new data (Figure 6c-g).

The paper was thorough in considering multiple datasets across a variety of behaviors, as well as existing decoding methods, to benchmark the MINT approach. This provided a valuable comparison to validate the method. The authors also provided nice intuition regarding why MINT may offer performance improvement in some cases and in which instances MINT may not perform as well.

Thank you. We were pleased to be able to provide comparisons across so many datasets (we are grateful to the Neural Latents Benchmark for making this possible).

In addition to providing a philosophical discussion as to the advantages of MINT and benchmarking against alternatives, the authors also provided a detailed description of practical considerations. This included training time, amount of training data, robustness to data loss or changes in the data, and interpretability. These considerations not only provided objective evaluation of practical aspects but also provided insights to the flexibility and robustness of the method as they relate back to the underlying assumptions and construction of the approach.

Thank you. We are glad that these sections were appreciated. MINT's simplicity and interpretability are indeed helpful in multiple ways, and afford opportunities for interesting future extensions. One potential benefit of interpretability is now explored in the newly added Figure 6e.

Impact:

This work is motivated by brain-computer interfaces applications, which it will surely impact in terms of neural decoder design. However, this work is also broadly impactful for neuroscientific analysis to relate neural spiking activity to observable behavioral features. Thus, MINT will likely impact neuroscience research generally. The methods are made publicly available, and the datasets used are all in public repositories, which facilitates adoption and validation of this method within the greater scientific community.

Again, thank you. We have similar hopes for this study.

Weaknesses (1 & 2 are related, and we have switched their order in addressing them):

Comment 2) With regards to the idea of neural and behavioral trajectories having different geometries, this is dependent on what behavioral variables are selected. In the example for Fig 2a, the behavior is reach position. The geometry of the behavioral trajectory of interest would look different if instead the behavior of interest was reach velocity. The paper would be strengthened by acknowledgement that geometries of trajectories are shaped by extrinsic choices rather than (or as much as they are) intrinsic properties of the data.

We agree. Indeed, we almost added a section to the original manuscript on this exact topic. We have now done so:

“A potential concern regarding the analyses in Figure 2c,d is that they require explicit choices of behavioral variables: muscle population activity in Figure 2c and angular phase and velocity in Figure 2d. Perhaps these choices were misguided. Might neural and behavioral geometries become similar if one chooses ‘the right’ set of behavioral variables? This concern relates to the venerable search for movement parameters that are reliably encoded by motor

cortex activity [69, 92–95]. If one chooses the wrong set of parameters (e.g. chooses muscle activity when one should have chosen joint angles) then of course neural and behavioral geometries will appear non-isometric. There are two reasons why this ‘wrong parameter choice’ explanation is unlikely to account for the results in Figure 2c,d. First, consider the implications of the left-hand side of Figure 2d. A small kinematic distance implies that angular position and velocity are nearly identical for the two moments being compared. Yet the corresponding pair of neural states can be quite distant. Under the concern above, this distance would be due to other encoded behavioral variables – perhaps joint angle and joint velocity – differing between those two moments. However, there are not enough degrees of freedom in this task to make this plausible. The shoulder remains at a fixed position (because the head is fixed) and the wrist has limited mobility due to the pedal design [60]. Thus, shoulder and elbow angles are almost completely determined by cycle phase. More generally, ‘external variables’ (positions, angles, and their derivatives) are unlikely to differ more than slightly when phase and angular velocity are matched. Muscle activity could be different because many muscles act on each joint, creating redundancy. However, as illustrated in Figure 2c, the key effect is just as clear when analyzing muscle activity. Thus, the above concern seems unlikely even if it can’t be ruled out entirely. A broader reason to doubt the ‘wrong parameter choice’ proposition is that it provides a vague explanation for a phenomenon that already has a straightforward explanation. A lack of isometry between the neural population response and behavior is expected when neural-trajectory tangling is low and output-null factors are plentiful [55, 60]. For example, in networks that generate muscle activity, neural and muscle-activity trajectories are far from isometric [52, 58, 60]. Given this straightforward explanation, and given repeated failures over decades to find the ‘correct’ parameters (muscle activity, movement direction, etc.) that create neural-behavior isometry, it seems reasonable to conclude that no such isometry exists.”

Comment 1) The authors posit that neural and behavioral trajectories are non-isometric. To support this point, they look at distances between neural states and distances between the corresponding behavioral states, in order to demonstrate that there are differences in these distances in each respective space. This supports the idea that neural states and behavioral states are non-isometric but does not directly address their point. In order to say the trajectories are non-isometric, it would be better to look at pairs of distances between corresponding trajectories in each space.

We like this idea and have added such an analysis. To be clear, we like the original analysis too: isometry predicts that neural and behavioral distances (for corresponding pairs of points) should be strongly correlated, and that small behavioral distances should not be associated with large neural distances. These predictions are not true, providing a strong argument against isometry. However, we also like the reviewer’s suggestion, and have added such an analysis. It makes the same larger point, and also reveals some additional facts (e.g. it reveals that muscle-geometry is more related to neural-geometry than is kinematic-geometry). The new analysis is described in the following section:

“We further explored the topic of isometry by considering pairs of distances. To do so, we chose two random neural states and computed their distance, yielding d_{neural1} . We repeated this process, yielding d_{neural2} . We then computed the corresponding pair of distances in muscle space (d_{muscle1} and d_{muscle2}) and kinematic space (d_{kin1} and d_{kin2}). We considered cases where d_{neural1} was meaningfully larger than (or smaller than) d_{neural2} , and asked whether the behavioral variables had the same relationship; e.g. was d_{muscle1} also larger than d_{muscle2} ? For kinematics, this relationship was weak: across 100,000 comparisons, the sign of $d_{\text{kin1}} - d_{\text{kin2}}$ agreed with $d_{\text{neural1}} - d_{\text{neural2}}$ only 67.3% of the time (with 50% being chance). The relationship was much stronger for muscles: the sign of $d_{\text{muscle1}} - d_{\text{muscle2}}$ agreed with $d_{\text{neural1}} - d_{\text{neural2}}$ 79.2% of the time, which is far more than expected by chance yet also far from what is expected given isometry (e.g. the sign

agrees 99.7% of the time for the truly isometric control data in Figure 2e). Indeed there were multiple moments during this task when $d_{neural1}$ was much larger than $d_{neural2}$, yet $d_{muscle1}$ was smaller than $d_{muscle2}$. These observations are consistent with the proposal that neural trajectories resemble muscle trajectories in some dimensions, but with additional output-null dimensions that break the isometry [60].”

Comment 3) The approach is built up on the idea of creating a "mesh" structure of possible states. In the body of the paper the definition of the mesh was not entirely clear and I could not find in the methods a more rigorous explicit definition. Since the mesh is integral to the approach, the paper would be improved with more description of this component.

This is a fair criticism. Although MINTs actual operations were well-documented, how those operations mapped onto the term ‘mesh’ was, we agree, a bit vague. The definition of the mesh is a bit subtle because it only emerges during decoding rather than being precomputed. This is part of what gives MINT much more flexibility than a lookup table. We have added the following to the manuscript.

“We use the term ‘mesh’ to describe the scaffolding created by the training-set trajectories and the interpolated states that arise at runtime. The term mesh is apt because, if MINT’s assumptions are correct, interpolation will almost always be local. If so, the set of decodable states will resemble a mesh, created by line segments connecting nearby training-set trajectories. However, this mesh-like structure is not enforced by MINT’s operations.

Interpolation could, in principle, create state-distributions that depart from the assumption of a sparse manifold. For example, interpolation could fill in the center of the green tube in Figure 1b, resulting in a solid manifold rather than a mesh around its outer surface. However, this would occur only if spiking observations argued for it. As will be documented below, we find that essentially all interpolation is local”

We have also added Figure 4d. This new analysis documents the fact that decoded states are near trainingset trajectories, which is why the term ‘mesh’ is appropriate.

Reviewer #3:

Summary:

This manuscript develops a new method termed MINT for decoding of behavior. The method is essentially a table-lookup rather than a model. Within a given stereotyped task, MINT tabulates averaged firing rate trajectories of neurons (neural states) and corresponding averaged behavioral trajectories as stereotypes to construct a library. For a test trial with a realized neural trajectory, it then finds the closest neural trajectory to it in the table and declares the associated behavior trajectory in the table as the decoded behavior. The method can also interpolate between these tabulated trajectories. The authors mention that the method is based on three key assumptions: (1) Neural states may not be embedded in a lowdimensional subspace, but rather in a high-dimensional space. (2) Neural trajectories are sparsely distributed under different behavioral conditions. (3) These neural states traverse trajectories in a stereotyped order.

The authors conducted multiple analyses to validate MINT, demonstrating its decoding of behavioral trajectories in simulations and datasets (Figures 3, 4). The main behavior decoding comparison is shown in Figure 4. In stereotyped tasks, decoding performance is comparable (M_Cycle, MC_Maze) or better (Area 2_Bump) than other linear/nonlinear algorithms

(Figure 4). However, MINT underperforms for the MC_RTT task, which is less stereotyped (Figure 4).

This paper is well-structured and its main idea is clear. The fact that performance on stereotyped tasks is high is interesting and informative, showing that these stereotyped tasks create stereotyped neural trajectories. The task-specific comparisons include various measures and a variety of common decoding approaches, which is a strength. However, I have several major concerns. I believe several of the conclusions in the paper, which are also emphasized in the abstract, are not accurate or supported, especially about generalization, computational scalability, and utility for BCIs. MINT is essentially a table-lookup algorithm based on stereotyped task-dependent trajectories and involves the tabulation of extensive data to build a vast library without modeling. These aspects will limit MINT's utility for real-world BCIs and tasks. These properties will also limit MINT's generalizability from task to task, which is important for BCIs and thus is commonly demonstrated in BCI experiments with other decoders without any retraining. Furthermore, MINT's computational and memory requirements can be prohibitive it seems. Finally, as MINT is based on tabulating data without learning models of data, I am unclear how it will be useful in basic investigations of neural computations. I expand on these concerns below.

We thank the reviewer for pointing out weaknesses in our framing and presentation. The comments above made us realize that we needed to 1) better document the ways in which MINT is far more flexible than a lookup-table, and 2) better explain the competing scientific perspectives at play. R3's comments also motivated us to add an additional analysis of generalization. In our view the manuscript is greatly improved by these additions. Specifically, these additions directly support the broader impact that we hope the study will have.

For simplicity and readability, we first group and summarize R3's main concerns in order to better address them. (These main concerns are all raised above, in addition to recurring in the specific comments below. Responses to each individual specific comment are provided after these summaries.)

(1) R3 raises concerns about 'computational scalability.' The concern is that "MINT's computational and memory requirements can be prohibitive." This point was expanded upon in a specific comment, reproduced below:

I also find the statement in the abstract and paper that "computations are simple, scalable" to be inaccurate. The authors state that MINT's computational cost is $O(NC)$ only, but it seems this is achieved at a high memory cost as well as computational cost in training. The process is described in section "Lookup table of log-likelihoods" on line [978-990]. The idea is to precompute the log-likelihoods for any combination of all neurons with discretization $\times$ all delay/history segments $\times$ all conditions and to build a large lookup table for decoding. Basically, the computational cost of precomputing this table is $O(V^{\{N\tau\}} \times TC)$ and the table requires a memory of $O(V^{\{N\tau\}})$, where V is the number of discretization points for the neural firing rates, N is the number of neurons, τ is the history length, T is the trial length, and C is the number of conditions. This is a very large burden, especially the $V^{\{N\tau\}}$ term. This cost is currently not mentioned in the manuscript and should be clarified in the main text. Accordingly, computation claims should be modified including in the abstract.

The revised manuscript clarifies that our statement (that computations are simple and scalable) is absolutely accurate. There is no need to compute, or store, a massive lookup table. There are three tables: two of modest size and one that is tiny. This is now better explained:

“Thus, the log-likelihood of $\bar{s}_{t'-\tau':t'}$ for a particular current neural state, is simply the sum of many individual log-likelihoods (one per neuron and time-bin). Each individual log-likelihood depends on only two numbers: the firing rate at that moment and the spike count in that bin. To simplify online computation, one can precompute the log-likelihood, under a Poisson model, for every plausible combination of rate and spike-count. For example, a lookup table of size 2001×21 is sufficient when considering rates that span 0-200 spikes/s in increments of 0.1 spikes/s, and considering 20 ms bins that contain at most 20 spikes (only one lookup table is ever needed, so long as its firing-rate range exceeds that of the most-active neuron at the most active moment in Ω). Now suppose we are observing a population of 200 neurons, with a 200 ms history divided into ten 20 ms bins. For each library state, the log-likelihood of the observed spike-counts is simply the sum of $200 \times 10 = 2000$ individual loglikelihoods, each retrieved from the lookup table. In practice, computation is even simpler because many terms can be reused from the last time bin using a recursive solution (Methods). This procedure is lightweight and amenable to real-time applications.”

In summary, the first table simply needs to contain the firing rate of each neuron, for each condition, and each time in that condition. This table consumes relatively little memory. Assuming 100 one-second-long conditions (rates sampled every 20 ms) and 200 neurons, the table would contain $100 \times 50 \times 200 = 1,000,000$ numbers. These numbers are typically stored as 16-bit integers (because rates are quantized), which amounts to about 2 MB. This is modest, given that most computers have (at least) tens of GB of RAM. A second table would contain the values for each behavioral variable, for each condition, and each time in that condition. This table might contain behavioral variables at a finer resolution (e.g. every millisecond) to enable decoding to update in between 20 ms bins (1 ms granularity is not needed for most BCI applications, but is the resolution used in this study). The number of behavioral variables of interest for a particular BCI application is likely to be small, often 1-2, but let's assume for this example it is 10 (e.g. x-, y-, and z-position, velocity, and acceleration of a limb, plus one other variable). This table would thus contain $100 \times 1000 \times 10 = 1,000,000$ floating point numbers, i.e. an 8 MB table. The third table is used to store the probability of s spikes being observed given a particular quantized firing rate (e.g. it may contain probabilities associated with firing rates ranging from 0 – 200 spikes/s in 0.1 spikes/s increments). This table is not necessary, but saves some computation time by precomputing numbers that will be used repeatedly. This is a very small table (typically $\sim 2000 \times 20$, i.e. 320 KB). It does not need to be repeated for different neurons or conditions, because Poisson probabilities depend on only rate and count.

(2) R3 raises a concern that MINT “is essentially a table-lookup rather than a model.” R3 states that MINT

“is essentially a table-lookup algorithm based on stereotyped task-dependent trajectories and involves the tabulation of extensive data to build a vast library without modeling.”

and that,

“as MINT is based on tabulating data without learning models of data, I am unclear how it will be useful in basic investigations of neural computations.”

This concern is central to most subsequent concerns. The manuscript has been heavily revised to address it. The revisions clarify that MINT is much more flexible than a lookup table, even though MINT uses a lookup table as its first step. Because R3's concern is intertwined with one's scientific assumptions, we have also added the new Figure 1 to explicitly illustrate the two key scientific perspectives and their decoding implications.

Under the perspective in Figure 1a, R3 would be correct in saying that there exist traditional interpretable decoders (e.g. a Kalman filter) whose assumptions better model the data. Under this perspective, MINT might still be an excellent choice in many cases, but other methods would be expected to gain the advantage when situations demand more flexibility. This is R3's central concern, and essentially all other concerns flow from it. It makes sense that R3 has this concern, because their comments repeatedly stress a foundational assumption of the perspective in Figure 1a: the assumption of a fixed lowdimensional neural subspace where activity has a reliable relationship to behavior that can be modeled and leveraged during decoding. The phrases below accord with that view:

"Unlike MINT, these works can achieve generalization because they model the neural subspace and its association to movement."

"it will not generalize... even when these tasks cover the exact same space (e.g., the same 2D computer screen and associated neural space)."

"For proper training, the training data should explore the whole movement space and the associated neural space"

"I also believe the authors should clarify the logic behind developing MINT better. From a scientific standpoint, we seek to gain insights into neural computations by making various assumptions and building models that parsimoniously describe the vast amount of neural data rather than simply tabulating the data. For instance, low-dimensional assumptions have led to the development of numerous dimensionality reduction algorithms and these models have led to important interpretations about the underlying dynamics"

Thus, R3 prefers a model that 1) assumes a low-dimensional subspace that is fixed across tasks and 2) assumes a consistent 'association' between neural activity and kinematics. Because R3 believes this is the correct model of the data, they believe that decoders should leverage it. Traditional interpretable method do, and MINT doesn't, which is why they find MINT to be unprincipled. This is a reasonable view, but it is not our view. We have heavily revised the manuscript to clarify that a major goal of our study is to explore the implications of a different, less-traditional scientific perspective.

The new Figure 1a illustrates the traditional perspective. Under this perspective, one would agree with R3's claim that other methods have the opportunity to model the data better. For example, suppose there exists a consistent neural subspace – conserved across tasks – where three neural dimensions encode 3D hand position and three additional neural dimensions encode 3D hand velocity. A traditional method such as a Kalman filter would be a very appropriate choice to model these aspects of the data.

Figure 1b illustrates the alternative scientific perspective. This perspective arises from recent, present, and to-be-published observations. MINT's assumptions are well-aligned with this perspective. In contrast, the assumptions of traditional methods (e.g. the Kalman filter) are not well-aligned with the properties of the data under this perspective. This does not mean traditional methods are not useful. Yet under Figure 1b, it is traditional methods, such as the Kalman filter, that lack an accurate model of the data. Of course, the reviewer may disagree with our scientific perspective. We would certainly concede that there is room for debate. However, we find the evidence for Figure 1b to be sufficiently strong that it is worth exploring the utility of methods that align with this scientific perspective. MINT is such a method. As we document, it performs very well.

Thus, in our view, MINT is quite principled because its assumptions are well aligned with the data. It is true that the features of the data that MINT models are a bit different from those that are traditionally modeled. For example, R3 is quite correct that MINT does not attempt to use a biomimetic model of the true transformation from neural activity, to muscle activity,

and thence to kinematics. We see this as a strength, and the manuscript has been revised accordingly (see paragraph beginning with “We leveraged this simulated data to compare MINT with a biomimetic decoder”).

(3) R3 raises concerns that MINT cannot generalize. This was a major concern of R3 and is intimately related to concern #2 above. The concern is that, if MINT is “essentially a lookup table” that simply selects pre-defined trajectories, then MINT will not be able to generalize. R3 is quite correct that MINT generalizes rather differently than existing methods. Whether this is good or bad depends on one’s scientific perspective. Under Figure 1a, MINT’s generalization would indeed be limiting because other methods could achieve greater flexibility. Under Figure 1b, all methods will have serious limits regarding generalization. Thus, MINT’s method for generalizing may approximate the best one can presently do. To address this concern, we have made three major changes, numbered i-iii below:

i) Large sections of the manuscript have been restructured to underscore the ways in which MINT can generalize. A major goal was to counter the impression, stated by R3 above, that:

“for a test trial with a realized neural trajectory, [MINT] then finds the closest neural trajectory to it in the table and declares the associated behavior trajectory in the table as the decoded behavior”.

This description is a reasonable way to initially understand how MINT works, and we concede that we may have over-used this intuition. Unfortunately, it can leave the misimpression that MINT decodes by selecting whole trajectories, each corresponding to ‘a behavior’. This can happen, but it needn’t and typically doesn’t. As an example, consider the cycling task. Suppose that the library consists of stereotyped trajectories, each four cycles long, at five fixed speeds from 0.5-2.5 Hz. If the spiking observations argued for it, MINT could decode something close to one of these five stereotyped trajectories. Yet it needn’t. Decoded trajectories will typically resemble library trajectories locally, but may be very different globally. For example, a decoded trajectory could be thirty cycles long (or two, or five hundred) perhaps speeding up and slowing down multiple times across those cycles.

Thus, the library of trajectories shouldn’t be thought of as specifying a limited set of whole movements that can be ‘selected from’. Rather, trajectories define a scaffolding that outlines where the neural state is likely to live and how it is likely to be changing over time. When we introduce the idea of library trajectories, we are now careful to stress that they don’t function as a set from which one trajectory is ‘declared’ to be the right one:

“We thus designed MINT to approximate that manifold using the trajectories themselves, rather than their covariance matrix or corresponding subspace. Unlike a covariance matrix, neural trajectories indicate not only which states are likely, but also which state-derivatives are likely. If a neural state is near previously observed states, it should be moving in a similar direction. MINT leverages this directionality.

Training-set trajectories can take various forms, depending on what is convenient to collect. Most simply, training data might include one trajectory per condition, with each condition corresponding to a discrete movement. Alternatively, one might instead employ one long trajectory spanning many movements. Another option is to employ many sub-trajectories, each briefer than a whole movement. The goal is simply for training-set trajectories to act as a scaffolding, outlining the manifold that might be occupied during decoding and the directions in which decoded trajectories are likely to be traveling.”

Later in that same section we stress that decoded trajectories can move along the ‘mesh’ in nonstereotyped ways:

“Although the mesh is formed of stereotyped trajectories, decoded trajectories can move along the mesh in non-stereotyped ways as long as they generally obey the flow-field implied by the training data. This flexibility supports many types of generalization, including generalization that is compositional in nature. Other types of generalization – e.g. from the green trajectories to the orange trajectories in Figure 1b – are unavailable when using MINT and are expected to be challenging for any method (as will be documented in a later section).”

The section “Training and decoding using MINT” has been revised to clarify the ways in which interpolation is flexible, allowing decoded movements to be globally very different from any library trajectory.

“To decode stereotyped trajectories, one could simply obtain the maximum-likelihood neural state from the library, then render a behavioral decode based on the behavioral state with the same values of c and k . This would be appropriate for applications in which conditions are categorical, such as typing or handwriting. Yet in most cases we wish for the trajectory library to serve not as an exhaustive set of possible states, but as a scaffolding for the mesh of possible states. MINT’s operations are thus designed to estimate any neural trajectory – and any corresponding behavioral trajectory – that moves along the mesh in a manner generally consistent with the trajectories in Ω .”

“...interpolation allows considerable flexibility. Not only is one not ‘stuck’ on a trajectory from Φ , one is also not stuck on trajectories created by weighted averaging of trajectories in Φ . For example, if cycling speed increases, the decoded neural state could move steadily up a scaffolding like that illustrated in Figure 1b (green). In such cases, the decoded trajectory might be very different in duration from any of the library trajectories. Thus, one should not think of the library as a set of possible trajectories that are selected from, but rather as providing a mesh-like scaffolding that defines where future neural states are likely to live and the likely direction of their local motion. The decoded trajectory may differ considerably from any trajectory within Ω .”

This flexibility is indeed used during movement. One empirical example is described in detail:

“During movement... angular phase was decoded with effectively no net drift over time. This is noteworthy because angular velocity on test trials never perfectly matched any of the trajectories in Φ . Thus, if decoding were restricted to a library trajectory, one would expect growing phase discrepancies. Yet decoded trajectories only need to locally (and approximately) follow the flow-field defined by the library trajectories. Based on incoming spiking observations, decoded trajectories speed up or slow down (within limits).

This decoding flexibility presumably relates to the fact that the decoded neural state is allowed to differ from the nearest state in Ω . To explore... [the text goes on to describe the new analysis in Figure 4d, which shows that the decoded state is typically not on any trajectory, though it is typically close to a trajectory].”

Thus, MINT’s operations allow considerable flexibility, including generalization that is compositional in nature. Yet R3 is still correct that there are other forms of generalization that are unavailable to MINT. This is now stressed at multiple points in the revision. However, under the perspective in Figure 1b, these forms of generalization are unavailable to any current method. Hence we made a second major change in response to this concern... ii) We explicitly illustrate how the structure of the data determines when generalization is or isn’t possible. The new Figure 1a,b introduces the two perspectives, and the new Figure 6a,b lays out their implications for generalization. Under the perspective in Figure 6a, the reviewer is quite right: other methods can generalize in ways that MINT cannot. Under the perspective in Figure 6b, expectations are very different. Those expectations make testable predictions.

Hence the third major change... iii) We have added an analysis of generalization, using a newly collected dataset. This dataset was collected using Neuropixels Probes during our Pac-Man force-tracking task. This dataset was chosen because it is unusually well-suited to distinguishing the predictions in Figure 6a versus Figure 6b. Finding a dataset that can do so is not simple. Consider R3's point that training data should "explore the whole movement space and the associated neural space". The physical simplicity of the Pac-Man task makes it unusually easy to confirm that the behavioral workspace has been fully explored. Importantly, under Figure 6b, this does not mean that the neural workspace has been fully explored, which is exactly what we wish to test when testing generalization. We do so, and compare MINT with a Wiener filter. A Wiener filter is an ideal comparison because it is simple, performs very well on this task, and should be able to generalize well under Figure 1a. Additionally, the Wiener filter (unlike the Kalman Filter) doesn't leverage the assumption that neural activity reflects the derivative of force. This matters because we find that neural activity does not reflect $d\text{force}/dt$ in this task. The Wiener filter is thus the most natural choice of the interpretable methods whose assumptions match Figure 1a.

The new analysis is described in Figure 6c-g and accompanying text. Results are consistent with the predictions of Figure 6b. We are pleased to have been motivated to add this analysis for two reasons. First, it provides an additional way of evaluating the predictions of the two competing scientific perspectives that are at the heart of our study. Second, this analysis illustrates an underappreciated way in which generalization is likely to be challenging for any decode method. It can be tempting to think that the main challenge regarding generalization is to fully explore the relevant behavioral space. This makes sense if a behavioral space has "an associated neural space". However, we are increasingly of the opinion that it doesn't. Different tasks often involve different neural subspaces, even when behavioral subspaces overlap. We have even seen situations where motor output is identical but neural subspaces are quite different. These facts are relevant to any decoder, something highlighted in the revised Introduction:

"MINT's performance confirms that there are gains to be made by building decoders whose assumptions match a different, possibly more accurate view of population activity. At the same time, our results suggest fundamental limits on decoder generalization. Under the assumptions in Figure 1b, it will sometimes be difficult or impossible for decoders to generalize to not-yet-seen tasks. We found that this was true regardless of whether one uses MINT or a more traditional method. This finding has implications regarding when and how generalization should be attempted."

We have also added an analysis (Figure 6e) illustrating how MINT's ability to compute likelihoods can be useful in detecting situations that may strain generalization (for any method). MINT is unusual in being able to compute and use likelihoods in this way.

Detailed responses to R3: we reproduce each of R3's specific concerns below, but concentrate our responses on issues not already covered above.

Main comments:

Comment 1. MINT does not generalize to different tasks, which is a main limitation for BCI utility compared with prior BCI decoders that have shown this generalizability as I review below. Specifically, given that MINT tabulates task-specific trajectories, it will not generalize to tasks that are not seen in the training data even when these tasks cover the exact same space (e.g., the same 2D computer screen and associated neural space).

First, the authors provide a section on generalization, which is inaccurate because it mixes up two fundamentally different concepts: 1) collecting informative training data and 2) generalizing from task to task. The former is critical for any algorithm, but it does not imply the latter. For example, removing one direction of cycling from the training set as the authors

do here is an example of generating poor training data because the two behavioral (and neural) directions are non-overlapping and/or orthogonal while being in the same space. As such, it is fully expected that all methods will fail. For proper training, the training data should explore the whole movement space and the associated neural space, but this does not mean all kinds of tasks performed in that space must be included in the training set (something MINT likely needs while modeling-based approaches do not). Many BCI studies have indeed shown this generalization ability using a model. For example, in Weiss et al. 2019, center-out reaching tasks are used for training and then the same trained decoder is used for typing on a keyboard or drawing on the 2D screen. In Gilja et al. 2012, training is on a center-out task but the same trained decoder generalizes to a completely different pinball task (hit four consecutive targets) and tasks requiring the avoidance of obstacles and curved movements. There are many more BCI studies, such as Jarosiewicz et al. 2015 that also show generalization to complex realworld tasks not included in the training set. Unlike MINT, these works can achieve generalization because they model the neural subspace and its association to movement. On the contrary, MINT models task-dependent neural trajectories, so the trained decoder is very task-dependent and cannot generalize to other tasks. So, unlike these prior BCIs methods, MINT will likely actually need to include every task in its library, which is not practical.

I suggest the authors remove claims of generalization and modify their arguments throughout the text and abstract. The generalization section needs to be substantially edited to clarify the above points. Please also provide the BCI citations and discuss the above limitation of MINT for BCIs.

As discussed above, R3's concerns are accurate under the view in Figure 1a (and the corresponding Figure 6a). Under this view, a method such as that in Gilja et al. or Jarosiewicz et al. can find the correct subspace, model the correct neuron-behavior correlations, and generalize to any task that uses "the same 2D computer screen and associated neural space", just as the reviewer argues. Under Figure 1b things are quite different.

This topic – and the changes we have made to address it – is covered at length above. Here we simply want to highlight an empirical finding: sometimes two tasks use the same neural subspace and sometimes they don't. We have seen both in recent data, and it is can be very non-obvious which will occur based just on behavior. It does not simply relate to whether one is using the same physical workspace. We have even seen situations where the patterns of muscle activity in two tasks are nearly identical, but the neural subspaces are fairly different. When a new task uses a new subspace, neither of the methods noted above (Gilja nor Jarosiewicz) will generalize (nor will MINT). Generalizing to a new subspace is basically impossible without some yet-to-be-invented approach. On the other hand, there are many other pairs of tasks (center-out-reaching versus some other 2D cursor control) where subspaces are likely to be similar, especially if the frequency content of the behavior is similar (in our recent experience this is often critical). When subspaces are shared, most methods will generalize, and that is presumably why generalization worked well in the studies noted above.

Although MINT can also generalize in such circumstances, R3 is correct that, under the perspective in Figure 1a, MINT will be more limited than other methods. This is now carefully illustrated in Figure 6a. In this traditional perspective, MINT will fail to generalize in cases where new trajectories are near previously observed states, yet move in very different ways from library trajectories. The reason we don't view this as a shortcoming is that we expect it to occur rarely (else tangling would be high). We thus anticipate the scenario in Figure 6b.

This is worth stressing because R3 states that our discussion of generalization "is inaccurate because it mixes up two fundamentally different concepts: 1) collecting informative training data and 2) generalizing from task to task." We have heavily revised this section and improved it. However, it was never inaccurate. Under Figure 6b, these two concepts

absolutely are mixed up. If different tasks use different neural subspaces, then this requires collecting different “informative training data” for each. One cannot simply count on having explored the physical workspace.

Comment 2. MINT is shown to achieve competitive/high performance in highly stereotyped datasets with structured trials, but worse performance on MC_RTT, which is not based on repeated trials and is less stereotyped. This shows that MINT is valuable for decoding in repetitive stereotyped use-cases. However, it also highlights a limitation of MINT for BCIs, which is that MINT may not work well for real-world and/or less-constrained setups such as typing, moving a robotic arm in 3D space, etc. This is again due to MINT being a lookup table with a library of stereotyped trajectories rather than a model. Indeed, the authors acknowledge that the lower performance on MC_RTT (Figure 4) may be caused by the lack of repeated trials of the same type. However, real-world BCI decoding scenarios will also not have such stereotyped trial structure and will be less/unconstrained, in which MINT underperforms. Thus, the claim in the abstract or lines 480-481 that MINT is an “excellent” candidate for clinical BCI applications is not accurate and needs to be qualified. The authors should revise their statements according and discuss this issue. They should also make the use-case of MINT on BCI decoding clearer and more convincing.

We discussed, above, multiple changes and additions to the revision that were made to address these concerns. Here we briefly expand on the comment that MINT achieves “worse performance on MC_RTT, which is not based on repeated trials and is less stereotyped”. All decoders performed poorly on this task. MINT still outperformed the two traditional methods, but this was the only dataset where MINT did not also perform better (overall) than the expressive GRU and feedforward network. There are probably multiple reasons why. We agree with R3 that one likely reason is that this dataset is straining generalization, and MINT may have felt this strain more than the two machine-learning-based methods. Another potential reason is the structure of the training data, which made it more challenging to obtain library trajectories in the first place. Importantly, these observations do not support the view in Figure 1a. MINT still outperformed the Kalman and Wiener filters (whose assumptions align with Fig. 1a). To make these points we have added the following:

“Decoding was acceptable, but noticeably worse, for the MC_RTT dataset... As will be discussed below, every decode method achieved its worst estimates of velocity for the MC_RTT dataset. In addition to the impact of slower reaches, MINT was likely impacted by training data that made it challenging to accurately estimate library trajectories. Due to the lack of repeated trials, MINT used AutoLFADS to estimate the neural state during training. In principle this should work well. In practice AutoLFADS may have been limited by having only 10 minutes of training data. Because the random-target task involved more variable reaches, it may also have stressed the ability of all methods to generalize, perhaps for the reasons illustrated in Figure 1b.

The only dataset where MINT did not perform the best overall was the MC_RTT dataset, where it was outperformed by the feedforward network and GRU. As noted above, this may relate to the need for MINT to learn neural trajectories from training data that lacked repeated trials of the same movement (a design choice one might wish to avoid). Alternatively, the less-structured MC_RTT dataset may strain the capacity to generalize; all methods experienced a drop in velocity-decoding R2 for this dataset compared to the others. MINT generalizes somewhat differently than other methods, and may have been at a modest disadvantage for this dataset. A strong version of this possibility is that perhaps the perspective in Figure 1a is correct, in which case MINT might struggle because it cannot use forms of generalization that are available to other methods (e.g. generalization based on neuron-velocity correlations). This strong version seems unlikely; MINT continued to

significantly outperform the Wiener and Kalman filters, which make assumptions aligned with Figure 1a.”

Comment 3. Related to 2, it may also be that MINT achieves competitive performance in offline and trial-based stereotyped decoding by overfitting to the trial structure in a given task, and thus may not generalize well to online performance due to overfitting. For example, a recent work showed that offline decoding performance may be overfitted to the task structure and may not represent online performance (Deo et al. 2023). Please discuss.

We agree that a limitation of our study is that we do not test online performance. There are sensible reasons for this decision:

“By necessity and desire, all comparisons were made offline, enabling benchmarked performance across a variety of tasks and decoded variables, where each decoder had access to the exact same data and recording conditions.”

We recently reported excellent online performance in the cycling task with a different algorithm

(Schroeder et al. 2022). In the course of that study, we consistently found that improvements in our offline decoding translated to improvements in our online decoding. We thus believe that MINT (which improves on the offline performance of our older algorithm) is a good candidate to work very well online. Yet we agree this still remains to be seen. We have added the following to the Discussion:

“With that goal in mind, there exist three important practical considerations. First, some decode algorithms experience a performance drop when used online. One presumed reason is that, when decoding is imperfect, the participant alters their strategy which in turn alters the neural responses upon which decoding is based. Because MINT produces particularly accurate decoding, this effect may be minimized, but this cannot be known in advance. If a performance drop does indeed occur, one could adapt the known solution of retraining using data collected during online decoding [13]. Another presumed reason (for a gap between offline and online decoding) is that offline decoders can overfit the temporal structure in training data [107]. This concern is somewhat mitigated by MINT’s use of a short spike-count history, but MINT may nevertheless benefit from data augmentation strategies such as including timelated versions of learned trajectories in the libraries”

Comment 4. Related to 2, since MINT requires firing rates to generate the library and simple averaging does not work for this purpose in the MC_RTT dataset (that does not have repeated trials), the authors needed to use AutoLFADS to infer the underlying firing rates. The fact that MINT requires the usage of another model to be constructed first and that this model can be computationally complex, will also be a limiting factor and should be clarified.

This concern relates to the computational complexity of computing firing-rate trajectories during training. Usually, rates are estimated via trial-averaging, which makes MINT very fast to train. This was quite noticeable during the Neural Latents Benchmark competition. As one example, for the “MC_Scaling 5 ms Phase”, MINT took 28 seconds to train while GPFA took 30 minutes, the transformer baseline (NDT) took 3.5 hours, and the switching nonlinear dynamical system took 4.5 hours.

However, the reviewer is quite correct that MINT’s efficiency depends on the method used to construct the library of trajectories. As we note, “MINT is a method for leveraging a trajectory library, not a method for constructing it”. One can use trial-averaging, which is very fast. One can also use fancier, slower methods to compute the trajectories. We don’t view this as a

negative – it simply provides options. Usually one would choose trial-averaging, but one does not have to. In the case of MC_RTT, one has a choice between LFADS and grouping into pseudo-conditions and averaging (which is fast). LFADS produces higher performance at the cost of being slower. The operator can choose which they prefer. This is discussed in the following section:

“For MINT, ‘training’ simply means computation of standard quantities (e.g. firing rates) rather than parameter optimization. MINT is thus typically very fast to train (Table 1), on the order of seconds using generic hardware (no GPUs). This speed reflects the simple operations involved in constructing the library of neural-state trajectories: filtering of spikes and averaging across trials. At the same time we stress that MINT is a method for leveraging a trajectory library, not a method for constructing it. One may sometimes wish to use alternatives to trial-averaging, either of necessity or because they improve trajectory estimates. For example, for the MC_RTT task we used AutoLFADS to infer the library. Training was consequently much slower (hours rather than seconds) because of the time taken to estimate rates. Training time could be reduced back to seconds using a different approach – grouping into pseudo-conditions and averaging – but performance was reduced. Thus, training will typically be very fast, but one may choose time-consuming methods when appropriate.”

Comment 5. I also find the statement in the abstract and paper that "computations are simple, scalable" to be inaccurate. The authors state that MINT's computational cost is $O(NC)$ only, but it seems this is achieved at a high memory cost as well as computational cost in training. The process is described in section "Lookup table of log-likelihoods" on line [978-990]. The idea is to precompute the log-likelihoods for any combination of all neurons with discretization $\times$ all delay/history segments $\times$ all conditions and to build a large lookup table for decoding. Basically, the computational cost of precomputing this table is $O(V^{\{N\tau\}} \times TC)$ and the table requires a memory of $O(V^{\{N\tau\}})$, where V is the number of discretization points for the neural firing rates, N is the number of neurons, τ is the history length, T is the trial length, and C is the number of conditions. This is a very large burden, especially the $V^{\{N\tau\}}$ term. This cost is currently not mentioned in the manuscript and should be clarified in the main text. Accordingly, computation claims should be modified including in the abstract.

As discussed above, the manuscript has been revised to clarify that our statement was accurate.

Comment 6. In addition to the above technical concerns, I also believe the authors should clarify the logic behind developing MINT better. From a scientific standpoint, we seek to gain insights into neural computations by making various assumptions and building models that parsimoniously describe the vast amount of neural data rather than simply tabulating the data. For instance, low-dimensional assumptions have led to the development of numerous dimensionality reduction algorithms and these models have led to important interpretations about the underlying dynamics (e.g., fixed points/limit cycles). While it is of course valid and even insightful to propose different assumptions from existing models as the authors do here, they do not actually translate these assumptions into a new model. Without a model and by just tabulating the data, I don't believe we can provide interpretation or advance the understanding of the fundamentals behind neural computations. As such, I am not clear as to how this library building approach can advance neuroscience or how these assumptions are useful. I think the authors should clarify and discuss this point.

As requested, a major goal of the revision has been to clarify the scientific motivations underlying MINT's design. In addition to many textual changes, we have added figures

(Figures 1a,b and 6a,b) to outline the two competing scientific perspectives that presently exist. This topic is also addressed by extensions of existing analyses and by new analyses (e.g. Figure 6c-g).

In our view these additions have dramatically improved the manuscript. This is especially true because we think R3's concerns, expressed above, are reasonable. If the perspective in Figure 1a is correct, then R3 is right and MINT is essentially a hack that fails to model the data. MINT would still be effective in many circumstances (as we show), but it would be unprincipled. This would create limitations, just as the reviewer argues. On the other hand, if the perspective in Figure 1b is correct, then MINT is quite principled relative to traditional approaches. Traditional approaches make assumptions (a fixed subspace, consistent neuron-kinematic correlations) that are not correct under Figure 1b.

We don't expect R3 to agree with our scientific perspective at this time (though we hope to eventually convince them). To us, the key is that we agree with R3 that the manuscript needs to lay out the different perspectives and their implications, so that readers have a good sense of the possibilities they should be considering. The revised manuscript is greatly improved in this regard.

Comment 7. Related to 6, there seems to be a logical inconsistency between the operations of MINT and one of its three assumptions, namely, sparsity. The authors state that neural states are sparsely distributed in some neural dimensions (Figure 1a, bottom). If this is the case, then why does MINT extend its decoding scope by interpolating known neural states (and behavior) in the training library? This interpolation suggests that the neural states are dense on the manifold rather than sparse, thus being contradictory to the assumption made. If interpolation-based dense meshes/manifolds underlie the data, then why not model the neural states through the subspace or manifold representations? I think the authors should address this logical inconsistency in MINT, especially since this sparsity assumption also questions the low-dimensional subspace/manifold assumption that is commonly made.

We agree this is an important issue, and have added an analysis on this topic (Figure 4d). The key question is simple and empirical: during decoding, does interpolation cause MINT to violate the assumption of sparsity? R3 is quite right that in principle it could. If spiking observations argue for it, MINT's interpolation could create a dense manifold during decoding rather than a sparse one. The short answer is that empirically this does not happen, in agreement with expectations under Figure 1b. Rather than interpolating between distant states and filling in large 'voids', interpolation is consistently local. This is a feature of the data, not of the decoder (MINT doesn't insist upon sparsity, even though it is designed to work best in situations where the manifold is sparse).

In addition to adding Figure 4d, we added the following (in an earlier section):

"The term mesh is apt because, if MINT's assumptions are correct, interpolation will almost always be local. If so, the set of decodable states will resemble a mesh, created by line segments connecting nearby training-set trajectories. However, this mesh-like structure is not enforced by MINT's operations. Interpolation could, in principle, create state-distributions that depart from the assumption of a sparse manifold. For example, interpolation could fill in the center of the green tube in Figure 1b, resulting in a solid manifold rather than a mesh around its outer surface. However, this would occur only if spiking observations argued for it. As will be documented below, we find that essentially all interpolation is local."

Recommendations for the authors:

Reviewer #1 (Recommendations For The Authors):

I appreciate the detailed methods section, however, more specifics should be integrated into the main text. For example on Line 238, it should additionally be stated how many minutes were used for training and metrics like the MAE which is used later should be reported here.

Thank you for this suggestion. We now report the duration of training data in the main text:

“Decoding R^2 was .968 over ~7.1 minutes of test trials based on ~4.4 minutes of training data.”

We have also added similar specifics throughout the manuscript, e.g. in the Fig. 5 legend:

“Results are based on the following numbers of training / test trials: MC_Cycle (174 train, 99 test), MC_Maze (1721 train, 574 test), Area2_Bump (272 train, 92 test), MC_RTT (810 train, 268 test).”

Similar additions were made to the legends for Fig. 6 and 8. Regarding the request to add MAE for the multitask network, we did not do so for the simple reason that the decoded variable (muscle activity) has arbitrary units. The raw MAE is thus not meaningful. We could of course have normalized, but at this point the MAE is largely redundant with the correlation. In contrast, the MAE is useful when comparing across the MC_Maze, Area2_Bump, and MC_RTT datasets, because they all involve the same scale (cm/s).

Regarding the MC_RTT task, AutoLFADS was used to obtain robust spike rates, as reported in the methods. However, the rationale for splitting the neural trajectories after AutoLFADS is unclear. If the trajectories were split based on random recording gaps, this might lead to suboptimal performance? It might be advantageous to split them based on a common behavioural state?

When learning neural trajectories via AutoLFADS, spiking data is broken into short (but overlapping) segments, rates are estimated for each segment via AutoLFADS, and these rates are then stitched together across segments into long neural trajectories. If there had been no recording gaps, these rates could have been stitched into a single neural trajectory for this dataset. However, the presence of recording gaps left us no choice but to stitch together these rates into more than one trajectory. Fortunately, recording gaps were rare: for the decoding analysis of MC_RTT there were only two recording gaps and therefore three neural trajectories, each ~2.7 minutes in duration.

We agree that in general it is desirable to learn neural trajectories that begin and end at behaviorally relevant moments (e.g. in between movements). However, having these trajectories potentially end midmovement is not an issue in and of itself. During decoding, MINT is never stuck on a trajectory. Thus, if MINT were decoding states near the end of a trajectory that was cut short due to a training gap, it would simply begin decoding states from other trajectories or elsewhere along the same trajectory in subsequent moments. We could have further trimmed the three neural trajectories to begin and end at behaviorally relevant moments, but chose not to as this would have only removed a handful of potentially useful states from the library.

We now describe this in the Methods:

“Although one might prefer trajectory boundaries to begin and end at behaviorally relevant moments (e.g. a stationary state), rather than at recording gaps, the exact boundary points are unlikely to be consequential for trajectories of this length that span multiple movements. If MINT estimates a state near the end of a long trajectory, its estimate will simply jump to

another likely state on a different trajectory (or earlier along the same trajectory) in subsequent moments. Clipping the end of each trajectory to an earlier behaviorally-relevant moment would only remove potentially useful states from the libraries.”

Are the training and execution times in Table 1 based on pure Matlab functions or Mex files? If it's Mex files as suggested by the code, it would be good to mention this in the Table caption.

They are based on a combination of MATLAB and MEX files. This is now clarified in the table caption:

“Timing measurements taken on a Macbook Pro (on CPU) with 32GB RAM and a 2.3 GHz 8-Core Intel Core i9 processor. Training and execution code used for measurements was written in MATLAB (with the core recursion implemented as a MEX file).”

As the method most closely resembles a Bayesian decoder it would be good to compare performance against a Naive Bayes decoder.

We agree and have now done so. The following has been added to the text:

“A natural question is thus whether a simpler Bayesian decoder would have yielded similar results. We explored this possibility by testing a Naïve Bayes regression decoder [85] using the MC_Maze dataset. This decoder performed poorly, especially when decoding velocity ($R^2 = .688$ and $.093$ for hand position and velocity, respectively), indicating that the specific modeling assumptions that differentiate MINT from a naive Bayesian decoder are important drivers of MINT’s performance.”

Line 199 Typo: The assumption of stereotypy trajectory also enables neural states (and decoded behaviors) to be updated in between time bins.

Fixed

Table 3: It's unclear why the Gaussian binning varies significantly across different datasets. Could the authors explain why this is the case and what its implications might be?

We have added the following description in the “Filtering, extracting, and warping data on each trial” subsection of the Methods to discuss how σ may vary due to the number of trials available for training and how noisy the neural data for those trials is:

“First, spiking activity for each neuron on each trial was temporally filtered with a Gaussian to yield single-trial rates. Table 3 reports the Gaussian standard deviations σ (in milliseconds) used for each dataset. Larger values of σ utilize broader windows of spiking activity when estimating rates and therefore reduce variability in those rate estimates. However, large σ values also yield neural trajectories with less fine-grained temporal structure. Thus, the optimal σ for a dataset depends on how variable the rate estimates otherwise are.”

An implementation of the method in an open-source programming language could further enhance the widespread use of the tool.

We agree this would be useful, but have yet not implemented the method in any other programming languages. Implementation in Python is still a future goal.

Reviewer #2 (Recommendations For The Authors):

- Figures 4 and 5 should show the error bars on the horizontal axis rather than portraying them vertically.

[Note that these are now Figures 5 and 6]

The figure legend of Figure 5 now clarifies that the vertical ticks are simply to aid visibility when symbols have very similar means and thus overlap visually. We don't include error bars (for this analysis) because they are very small and would mostly be smaller than the symbol sizes. Instead, to indicate certainty regarding MINT's performance measurements, the revised text now gives error ranges for the correlations and MAE values in the context of Figure 4c. These error ranges were computed as the standard deviation of the sampling distribution (computed via resampling of trials) and are thus equivalent to SEMs. The error ranges are all very small; e.g. for the MC_Maze dataset the MAE for x-velocity is 4.5 +/- 0.1 cm/s. (error bars on the correlations are smaller still).

Thus, for a given dataset, we can be quite certain of how well MINT performs (within ~2% in the above case). This is reassuring, but we also don't want to overemphasize this accuracy. The main sources of variability one should be concerned about are: 1) different methods can perform differentially well for different brain areas and tasks, 2) methods can decode some behavioral variables better than others, and 3) performance depends on factors like neuron-count and the number of training trials, in ways that can differ across decode methods. For this reason, the study examines multiple datasets, across tasks and brain areas, and measures performance for a range of decoded variables. We also examine the impact of training-set-size (Figure 8a) and population size (solid traces in Fig. 8b, see R2's next comment below).

There is one other source of variance one might be concerned about, but it is specific to the neuralnetwork approaches: different weight initializations might result in different performance. For this reason, each neural-network approach was trained ten times, with the average performance computed. The variability around this average was very small, and this is now stated in the Methods.

"For the neural networks, the training/testing procedure was repeated 10 times with different random seeds. For most behavioral variables, there was very little variability in performance across repetitions. However, there were a few outliers for which variability was larger. Reported performance for each behavioral group is the average performance across the 10 repetitions to ensure results were not sensitive to any specific random initialization of each network."

- For Figure 6, it is unclear whether the neuron-dropping process was repeated multiple times. If not, it should be since the results will be sensitive to which particular subsets of neurons were "dropped". In this case, the results presented in Figure 6 should include error bars to describe the variability in the model performance for each decoder considered.

A good point. The results in Figure 8 (previously Figure 6) were computed by averaging over the removal of different random subsets of neurons (50 subsets per neuron count), just as the reviewer requests. The figure has been modified to include the standard deviation of performance across these 50 subsets. The legend clarifies how this was done.

Reviewer #3 (Recommendations For The Authors):

Other comments:

(1) [Line 185-188] The authors argue that in a 100-dimensional space with 10 possible discretized values, 10^{100} potential neural states need to be computed. But I am not clear on this. This argument seems to hold only in the absence of a model (as in MINT). For a model, e.g., Kalman filter or AutoLFADS, information is encoded in the latent state. For example, a simple Kalman filter for a linear model can be used for efficient inference. This 10^{100} computation isn't a general problem but seems MINT-specific, please clarify.

We agree this section was potentially confusing. It has been rewritten. We were simply attempting to illustrate why maximum likelihood computations are challenging without constraints. MINT simplifies this problem by adding constraints, which is why it can readily provide data likelihoods (and can do so using a Poisson model). The rewritten section is below:

“Even with 1000 samples for each of the neural trajectories in Figure 3, there are only 4000 possible neural states for which log-likelihoods must be computed (in practice it is fewer still, see Methods). This is far fewer than if one were to naively consider all possible neural states in a typical rate- or factor-based subspace. It thus becomes tractable to compute log-likelihoods using a Poisson observation model. A Poisson observation model is usually considered desirable, yet can pose tractability challenges for methods that utilize a continuous model of neural states. For example, when using a Kalman filter, one is often restricted to assuming a Gaussian observation model to maintain computational tractability “

(2) [Figure 6b] Why do the authors set the dropped neurons to zero in the "zeroed" results of the robustness analysis? Why not disregard the dropped neurons during the decoding process?

We agree the terminology we had used in this section was confusing. We have altered the figure and rewritten the text. The following, now at the beginning of that section, addresses the reviewer's query:

“It is desirable for a decoder to be robust to the unexpected loss of the ability to detect spikes from some neurons. Such loss might occur while decoding, without being immediately detected. Additionally, one desires robustness to a known loss of neurons / recording channels. For example, there may have been channels that were active one morning but are no longer active that afternoon. At least in principle, MINT makes it very easy to handle this second situation: there is no need to retrain the decoder, one simply ignores the lost neurons when computing likelihoods. This is in contrast to nearly all other methods, which require retraining because the loss of one neuron alters the optimal parameters associated with every other neuron.”

The figure has been relabeled accordingly; instead of the label ‘zeroed’, we use the label ‘undetected neuron loss’.

(3) Authors should provide statistical significance on their results, which they already did for Fig. S3a,b,c but missing on some other figures/places.

We have added error bars in some key places, including in the text when quantifying MINT's performance in the context of Figure 4. Importantly, error bars are only as meaningful as the source of error they assess, and there are reasons to be careful given this. The standard method for putting error bars on performance is to resample trials, which is indeed what we now report. These error bars are very small. For example, when decoding horizontal velocity for the MC_Maze dataset, the correlation between MINT's decode and the true velocity had a mean and SD of the sampling distribution of 0.963 ± 0.001 . This means that, for a given dataset and target variable, we have enough trials/data that we can be quite certain of how

well MINT performs. However, we want to be careful not to overstate this certainty. What one really wants to know is how well MINT performs across a variety of datasets, brain areas, target variables, neuron counts, etc. It is for this reason that we make multiple such comparisons, which provides a more valuable view of performance variability.

For Figure 7, error bars are unavailable. Because this was a benchmark, there was exactly one test-set that was never seen before. This is thus not something that could be resampled many times (that would have revealed the test data and thus invalidated the benchmark, not to mention that some of these methods take days to train). We could, in principle, have added resampling to Figure 5. In our view it would not be helpful and could be misleading for the reasons noted above. If we computed standard errors using different train/test partitions, they would be very tight (mostly smaller than the symbol sizes), which would give the impression that one can be quite certain of a given R^2 value. Yet variability in the train/test partition is not the variability one is concerned about in practice. In practice, one is concerned about whether one would get a similar R^2 for a different dataset, or brain area, or task, or choice of decoded variable. Our analysis thus concentrated on showing results across a broad range of situations. In our view this is a far more relevant way of illustrating the degree of meaningful variability (which is quite large) than resampling, which produces reassuringly small but (mostly) irrelevant standard errors.

Error bars are supplied in Figure 8b. These error bars give a sense of variability across resamplings of the neural population. While this is not typically the source of variability one is most concerned about, for this analysis it becomes appropriate to show resampling-based standard errors because a natural concern is that results may depend on which neurons were dropped. So here it is both straightforward, and desirable, to compute standard errors. (The fact that MINT and the Wiener filter can be retrained many times swiftly was also key – this isn't true of the more expressive methods). Figure S1 also uses resampling-based confidence intervals for similar reasons.

(4) [Line 431-437] Authors state that MINT outperforms other methods with the PSTH R^2 metric (trial-averaged smoothed spikes for each condition). However, I think this measure may not provide a fair comparison and is confounded because MINT's library is built using PSTH (i.e., averaged firing rate) but other methods do not use the PSTH. The author should clarify this.

The PSTH R^2 metric was not created by us; it was part of the Neural Latents Benchmark. They chose it because it ensures that a method cannot 'cheat' (on the Bits/Spike measure) by reproducing fine features of spiking while estimating rates badly. We agree with the reviewer's point: MINT's design does give it a potential advantage in this particular performance metric. This isn't a confound though, just a feature. Importantly, MINT will score well on this metric only if MINT's neural state estimate is accurate (including accuracy in time). Without accurate estimation of the neural state at each time, it wouldn't matter that the library trajectory is based on PSTHs. This is now explicitly stated:

"This is in some ways unsurprising: MINT estimates neural states that tend to resemble (at least locally) trajectories 'built' from training-set-derived rates, which presumably resemble test-set rates. Yet strong performance is not a trivial consequence of MINT's design. MINT does not 'select' whole library trajectories; PSTH R^2 will be high only if condition (c), index (k), and the interpolation parameter (a) are accurately estimated for most moments."

<https://doi.org/10.7554/eLife.89421.2.sa0>