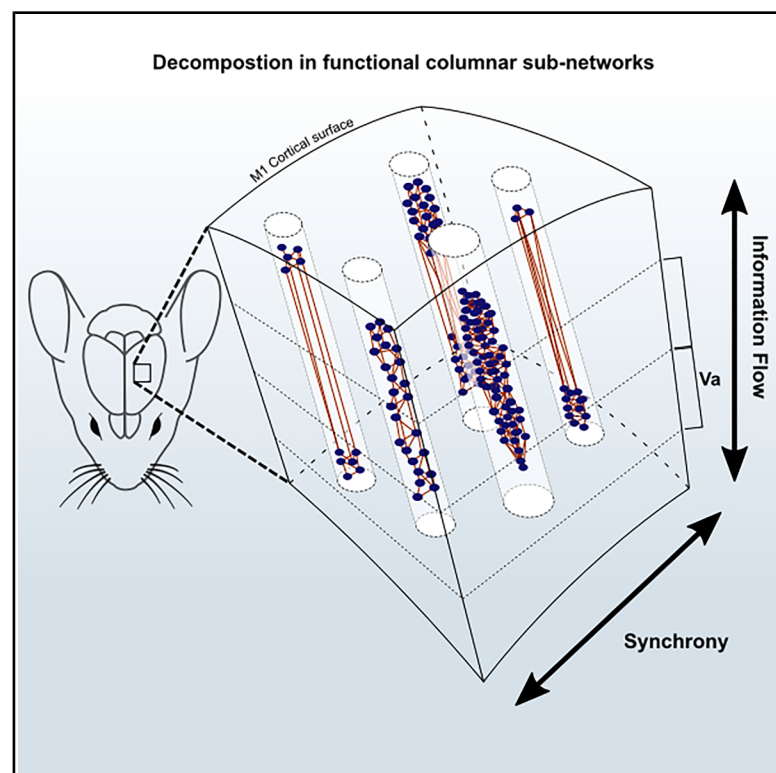


Graph-based analysis of volumetric image data reveals predominant layer Va-to-II/III feedback in mouse motor cortex

Graphical abstract



Authors

Philippe Aymard, Juan Carlos Boffi, Thibault Lagache, Hiroki Asari, Robert Prevedel, David Holcman

Correspondence

prevedel@embl.de (R.P.), david.holcman@ens.fr (D.H.)

In brief

Using fast volumetric two-photon Ca^{2+} imaging and a graph-inference pipeline, Aymard et al. reconstructed 3D functional networks in motor cortex in an awake mouse. Aymard et al. uncovered seven column microcircuit, directional feedback flow of information ($\text{Va} \rightarrow \text{II/III}$) and specific microcircuits correlated with spontaneous movement.

Highlights

- 2-photon volumetric imaging records $\sim 1,000$ motor cortex neurons at 4Hz
- Graph reconstruction method reveals 3D functional connectivity
- Extraction of ~ 30 column-like subnetworks of ~ 10 neurons
- Stronger feedback flow of information between layer II/III and Va



Article

Graph-based analysis of volumetric image data reveals predominant layer Va-to-II/III feedback in mouse motor cortex

Philippe Aymard,¹ Juan Carlos Boffi,² Thibault Lagache,³ Hiroki Asari,⁴ Robert Prevedel,^{5,6,*} and David Holcman^{1,7,8,*}¹Applied Mathematics and Computational Biology Group, IBENS, UMR8197, Ecole Normale Supérieure, PSL University, 75005 Paris, France²Crick Advanced Light Microscopy Science and Technology Platform, The Francis Crick Institute, London, UK³Institut Pasteur, Université de Paris Cité, CNRS UMR 3691, Biologie Analytique Unit, Paris, France⁴Scuola Internazionale Superiore di Studi Avanzati (SISSA), 34136 Trieste, Italy⁵Epigenetics and Neurobiology Unit, European Molecular Biology Laboratory, 00015 Monterotondo, Italy⁶Cell Biology and Biophysics Unit, European Molecular Biology Laboratory, 169117 Heidelberg, Germany⁷Churchill College, University of Cambridge, Cambridge, CB30DS UK⁸Lead contact*Correspondence: prevedel@embl.de (R.P.), david.holcman@ens.fr (D.H.)<https://doi.org/10.1016/j.celrep.2026.117618>

SUMMARY

How precise 3D interactions among cortical neurons underlie layer-specific computations remains elusive. We develop a graph-framework to infer functional connectivity from fast volumetric two-photon Ca²⁺ imaging of spontaneous activity in the awake mouse primary motor cortex. By converting deconvolved traces into binary spike trains, removing population bursts, and applying an adaptive, layer-specific threshold, we reconstruct a directed, weighted network of ~1,000 neurons. Decomposition into strongly connected components reveals ~30 sub-networks of ~10 neurons, predominantly in layer II/III and often bridging to layer Va. Across six 20-min recordings, we find that (1) layer II/III dominates connectivity, (2) feedback (Va → II/III) links exceed and outweigh feedforward (II/III → Va) ones, and (3) information flows in ≤6 synapses. We uncover seven geometrical and dynamical motifs with characteristic event sizes and durations, revealing diverse column-like microcircuits in M1 with a net ascending flow, suggesting that such sub-networks form elemental processing modules for motor control.

INTRODUCTION

Although cortical columns have long been proposed as the neo-cortex's elementary processing units,¹ it remains poorly understood how their three-dimensional (3D) functional pathways are spatially organized across layers. Unraveling this micro-scale architecture is essential for deciphering key processes such as perception, cognition, and refined motor control.^{2–4} In sensory areas, sequential two-photon imaging has provided indirect evidence of radial organization through receptive-field mapping,^{5–7} but such approaches cannot capture simultaneous activity across layers and thus fall short of revealing true functional connectivity.

Recent advances in volumetric *in vivo* imaging—most notably scanned temporal-focusing two-photon (sTeFo 2P) microscopy—now permit near-simultaneous recordings of thousands of neurons across hundreds of micrometers in depth.^{8–11} This dense, high-speed sampling opens the door to reconstructing 3D networks, yet a systematic, data-driven framework for inferring layer-to-layer connectivity from such recordings has not been established. Prior efforts to infer functional networks from calcium imaging¹² either relied on simulated data, two-dimensional analyses, or neglected cortical layer diversity and 3D topology.^{13–23}

Here, we introduce a graph-pipeline that transforms volumetric calcium time series into directed, weighted networks of cortical neurons that mirror neuronal networks. By converting deconvolved traces into binary spike trains, removing population bursts via a background-matching filter, and applying adaptive, layer-specific statistical thresholds, we infer high-confidence functional links between individual cells. Decomposition into strongly connected components (SCCs) reveals ~10 neurons per sub-network that predominantly form column-like motifs in layer II/III and often bridge to layer Va. Across six 20-min recordings of the spontaneous activity in the primary motor cortex (M1) of awake mice, we identified seven distinct SCC archetypes, demonstrate a net feedback (Va → II/III) information flow in ≤6 synaptic steps, and uncover highly connected hub neurons that reinforce local dynamics. Our approach confirms both long-standing hypotheses about cortical column organization^{1,24,25} and unveils previously unrecognized microcircuit motifs—providing a general toolkit for mapping functional architecture throughout the neocortex.

RESULTS

To unveil the fine-scale organization of functional subnetworks in M1, we start by developing a computational pipeline that



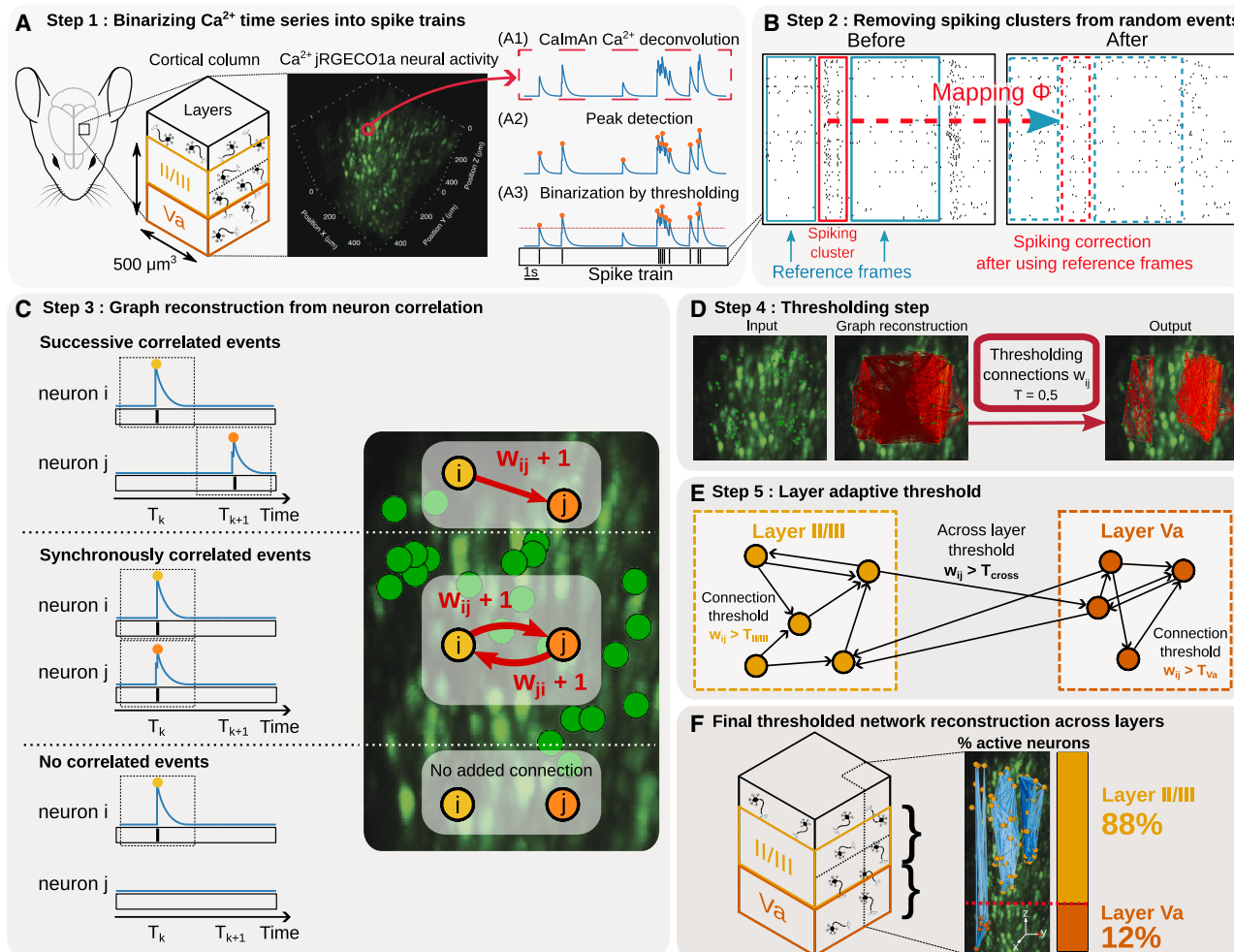


Figure 1. Network reconstruction

(A) Ca^{2+} signals recorded from a $470 \times 470 \times 450 \mu\text{m} = 0.1\text{mm}^3$ volume across layer II/III and Va from a mouse brain (left, Step 1) A1. jRGECO1a calcium time series deconvolved with CalmAn.²⁸ A2. Peak detection with prominence $p = 0.015$ (see Table 1). A3. Binarization of calcium events into spike trains using threshold $T_T = 0.1$ (see also Figure S1 for raw data).

(B) (Step 2) Removal of spiking cluster events using function ϕ (see STAR Methods and Figure S2) that matches bursting firing rates to the background activity (blue).

(C) (Step 3) Computing the matrix weight w_{ij} which represent temporal correlation between two firing neurons i, j : a successful correlation between neuron i and j is associated with a +1 increment of w_{ij} . For a simultaneous correlation event, both w_{ij} and w_{ji} are incremented by one. When no correlation occurs, the weight remains unchanged.

(D) (Step 4) Thresholding procedure to remove coincidental connections between neuron i and j when the weight $w_{ij} < T$.

(E) (Step 5) Thresholding procedure is adapted to each layer population (see STAR Methods and Figures S3 and S4).

(F) (Final output) Network reconstruction of cortical layers. Layer repartition: 88% for layer II/III vs. 12% for layer Va of active neurons with at least one connection.

converts volumetric calcium signals into graph representations of neuronal connectivity. Below, we summarize the key steps—from raw fluorescence traces to a directed, weighted network.

Volumetric two-photon imaging of large-scale mouse motor cortex activity in awake mice

We performed fast volumetric calcium imaging using sTeFo 2P microscopy⁸ in awake, head-fixed C57BL/6J mice. We recorded one session on 6 different animals resulting in 6 datasets of 1,000 neurons within a $470 \times 470 \times 450 \mu\text{m}^3$ volume at 4Hz. To target

the M1 region controlling face movements—and avoid intermingled motor areas—we centered our field of view on the snout-associated zone.²⁶ Because mouse M1 is agranular,²⁷ we classified somata deeper than $400 \mu\text{m}$ as layer Va and those above $400 \mu\text{m}$ as layer II/III (see STAR Methods for further details).

Raw $\Delta F/F_0$ time series were first motion- and neuropil-corrected, baseline-drift-adjusted, and denoised using the CalmAn toolbox.²⁸ We then deconvolved each trace and extracted spike times by identifying peaks with a prominence threshold of $p = 0.015$ (Figure 1A). Although mice remained

Table 1. Parameters used for network reconstruction and segmentation

Parameters	Symbol	Values
Peak detection prominence	ρ	0.015
Spike firing threshold	T_f	0.05
Activation regime threshold	T_c	2σ
Adaptive threshold confidence	α	0.05
SCC size threshold	T_{scc}	3
XY plan distance threshold	T_D	45
Event activation threshold of an SCC	T_{event}	2

passive during imaging, three of six datasets exhibited transient, high-amplitude population bursts. To prevent these episodes from biasing our connectivity estimates, we applied a background-matching filter that selectively down-samples burst-associated spikes while preserving overall event continuity (see [STAR Methods](#) and [Figure 1B](#)).

Reconstructing 3D functional connectivity across layers

To infer layer-to-layer functional links despite our relatively low 4 Hz sampling rate and 20 min recordings, we adopted a recurrence-based, statistical correlation framework. We computed pairwise weights by counting both synchronous and consecutive spike events between every neuron pair ([Figure 1C](#)), reasoning that true functional connections will reappear reliably over long recordings. By aggregating these co-activation counts across the entire session, we obtained a directed, weighted adjacency matrix in which stronger edges reflect more consistently recurring interactions. This approach ensures that only pathways with repeated temporal associations—rather than isolated coincidences—survive into the reconstructed 3D network.

To prune spurious co-activations, we derived an adaptive, layer-specific threshold from a data-driven null model. In this model (see [STAR Methods](#)), independent homogeneous Poisson processes are used to compute the distribution of coincidental co-activation count expected by chance using the most two active neurons as parameters ([Figures 1D–1E](#)). A threshold is derived by computing the 95 centile of each layer-specific distribution. Only edges whose weights exceed the threshold are retained, ensuring that our 3D network reflects repeatedly recurring interactions rather than isolated coincidences. Our threshold computation algorithm accounts for differences in baseline as we use the two most active neurons within the population during the entire recording duration, thereby also scaling with firing intensities. This adaptive threshold is therefore data-driven and accounts for differences in baseline activity across layers and datasets.

Applying this criterion across six datasets, we reconstructed functional graphs encompassing on average $33.9\% \pm 12.1\%$ of imaged neurons. This conservative fraction is likely due to both outside-volume connectivity and the stringent confidence threshold. Within this subnetwork, 88% of retained neurons reside in layer II/III versus 12% in layer Va—slightly enriched relative to their raw proportions (84% versus 16%), consistent with higher firing rates and stronger recurrent dynamics in the superficial layers ([Figure 1F](#)). Due to potential undersampling

of layer Va associated with the difficulties of deep tissue imaging. However, the change of these proportions must be taken cautiously.

Quantitative connectivity patterns in layers II/III and Va

We represent the reconstructed network as a directed, weighted graph with adjacency matrix $A \in \mathbb{R}^{N \times N}$, where N is the number of neurons ([Figure 2A](#)). Each entry $A_{ij} = w_{ij}$ counts the number of synchronous or successive co-activations between neuron i and neuron j ([Figure 1C](#)); absent connections have $w_{ij} = 0$. After applying our statistical threshold (see [STAR Methods](#)), we computed standard graph metrics.

The in- and out-degree (i.e., number of incoming and outgoing connections) distributions exhibit a monotonic decay but feature a small “bump” in the in-degree between 18 and 30 connections, shifting its mean (7.44 ± 9.30) slightly above the out-degree mean (7.49 ± 9.57) ([Figure 2B](#)). Layer-specific breakdown shows this excess in-degree is driven by neurons in layer II/III, consistent with locally clustered activity. We fitted both degree distributions with a double-exponential model ($y = Ae^{-Bx} + Ce^{-Dx}$), which outperformed single-exponential and power-law alternatives, suggesting two activation regimes: a background Poisson-like firing and a reinforced, hub-driven regime in layer II/III. This observation aligns with recent evidence that cortical networks rarely follow strict scale-free laws.²⁹

Edge weights—normalized by each dataset’s maximum—also follow an exponential decay ($y = Ae^{-Bx}$), with a mean weight of 29.0 ± 11.9 across layers ([Figure 2C](#)). Again, layer II/III dominates the high-weight tail. Together, these results underscore the prominent role of superficial layers in structuring both the density and strength of functional links in M1.

Another critical insight emerged when we segregated inter-layer edges by their directionality. We found that when we pooled all connections between layer Va and II/III 56% of them were ascending (Va $\rightarrow$ II/III), versus 44% descending (II/III $\rightarrow$ Va) ([Figure 2D](#), left). This result seems to be heterogeneous across datasets giving a less significant averaged result ($50.1\% \pm 13\%$). However, ascending links carried 10% more total weight than descending ones ([Figure 2D](#), right), an imbalance driven primarily by edges of intermediate strength (weights 20–40). These results reveal a robust feedback bias in functional information flow, with mid-range pathways disproportionately reinforcing the upward projection.

Identification of high-density functional modules via SCCs

To identify densely interconnected sub-networks, we decomposed the directed graph into SCCs, each defined by the property that every neuron is reachable from every other ([Figure 3A](#)). Across six recordings, we identified on average 29.2 SCCs per dataset (≈ 175 in total), with a mean size of 9.8 ± 9.2 neurons ([Figure 4A](#)). Most SCCs formed compact, column-like clusters, though a minority appeared as multi-column outliers. We used a clustering algorithm density-based spatial clustering of applications with noise (DBSCAN)³⁰ on SCC size and planar spread to isolate these outliers and then split them into elementary columns using a distance threshold $T_D = 45$ ([STAR Methods](#)). After discarding trivial SCCs (fewer or equal to a threshold $T_{scc} = 3$

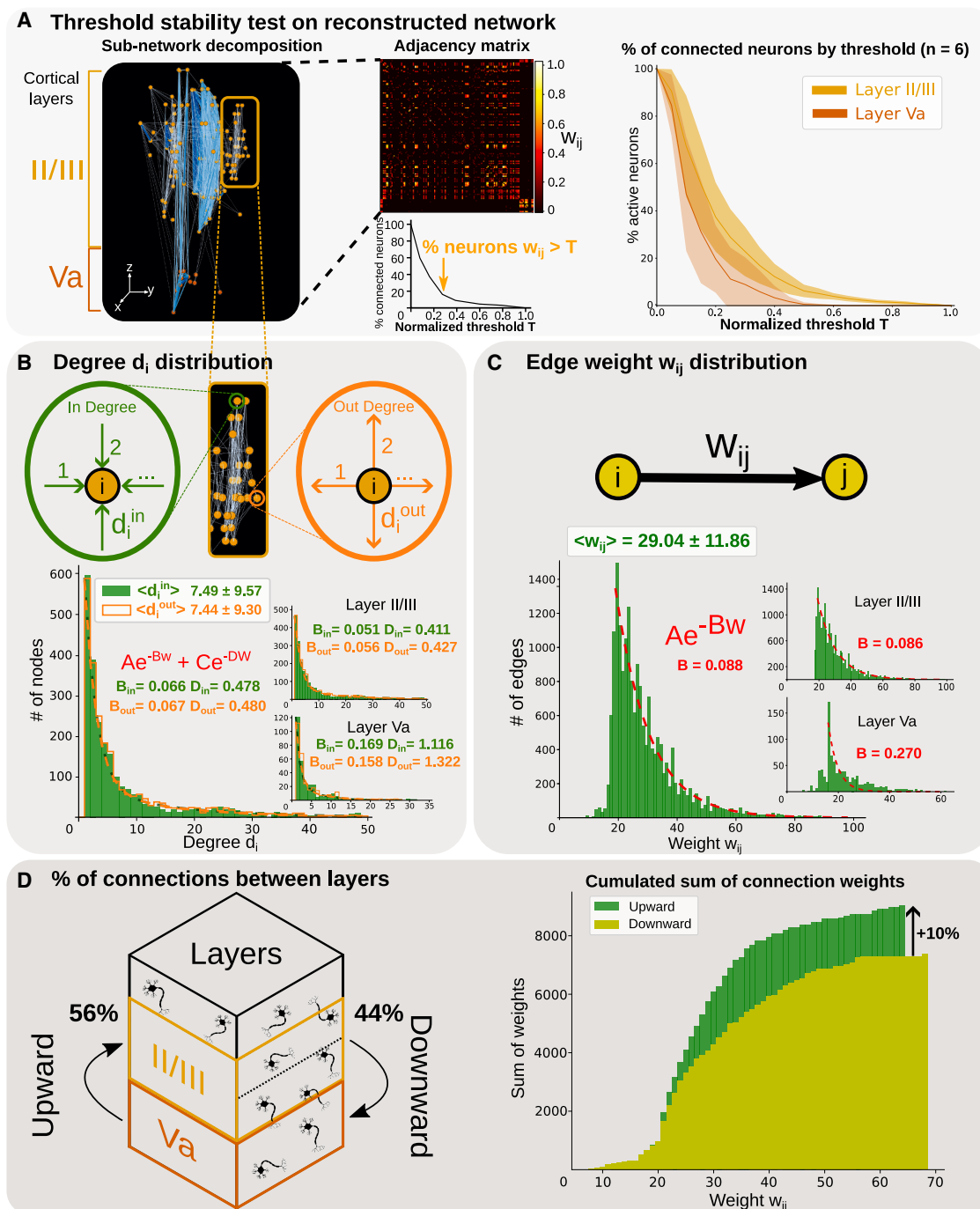


Figure 2. Structure and statistics of the neuronal network across layers II/III and Va

(A) Left: example of a network reconstruction with threshold $T = 0.25$. Middle: normalized adjacency matrix $w_{ij}/\max(w_{ij})$. Right: evolution of the proportion of connected neurons to the graph (i.e., with at least one connection) between layer II/III and layer Va when increasing threshold T over the 6 experiments (6 mice).

(B) Upper: in/out degree definition. Lower: in (green) and out (orange) degree distributions over the entire layers and within layer II/III and Va. A sum of two decaying exponential is fitted over each degree distribution.

(C) Edge weight w_{ij} distribution fitted by a single exponential (red).

(D) Left: proportion of upward (resp. downward) connections from layer Va to Layer II/III (resp. II/III to Va) in the network. Right: cumulated sum of upward connections (green) and downward connections (yellow) weights.

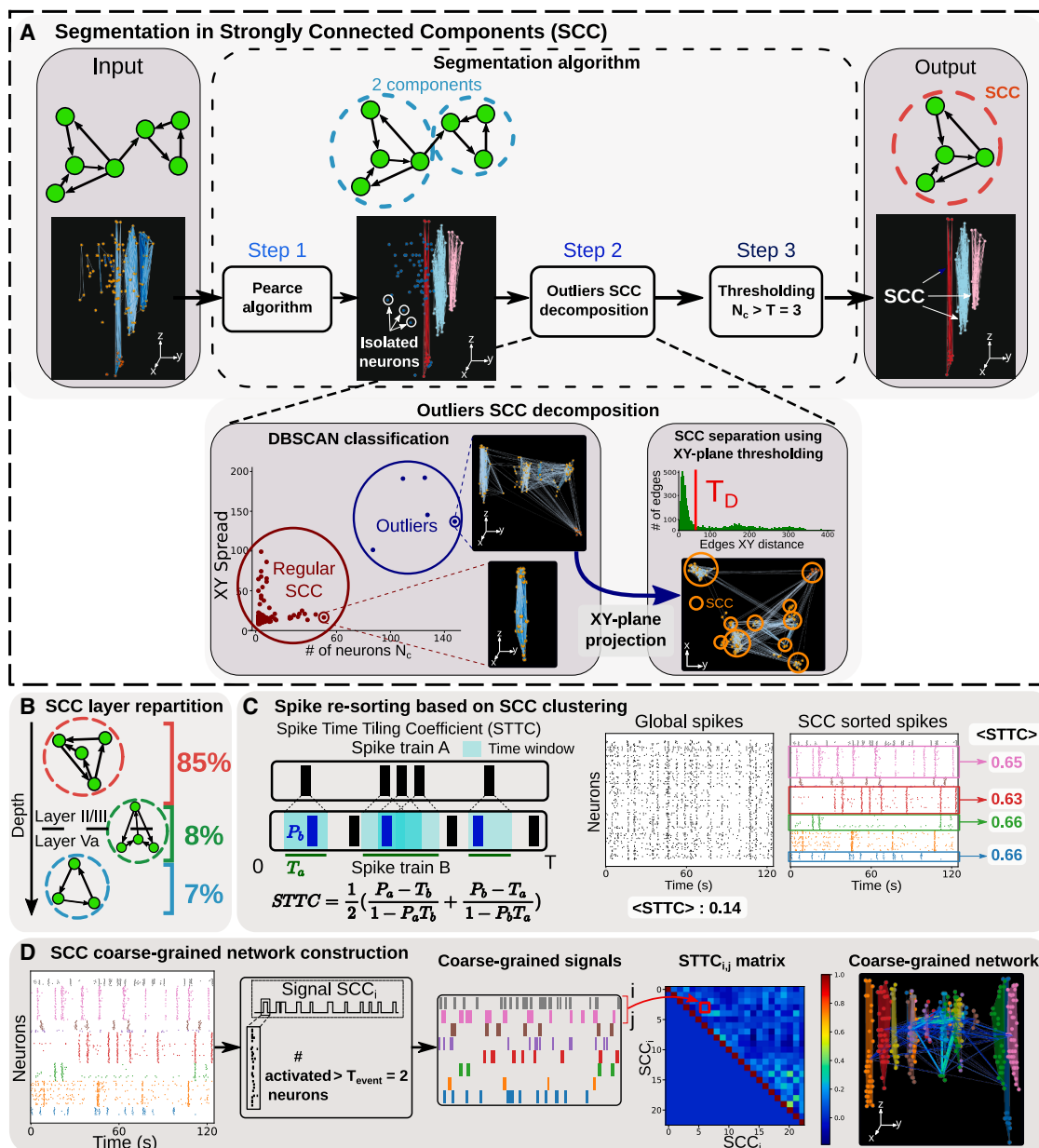


Figure 3. Decomposition of cortical layer network into sub-components

(A) Segmentation of the network in strongly connected components (SCC) in three steps. Step 1: pearce algorithm to decompose the network in SCC. Step 2: outlier detection using the DBSCAN algorithm based on a planar projection (see inset) (see STAR Methods (8)). Outliers are further decomposed into smaller SCC by removing edges with planar distance $d_{xy} < T_D = 45$ (see Table 1). Step 3: constructing the final set of sub-networks by keeping the components where the number of neurons $N_c > T = 3$.

(B) Distribution of sub-networks across layers: 85% in layer II/III, 8% across layers, and 5% in layer Va.

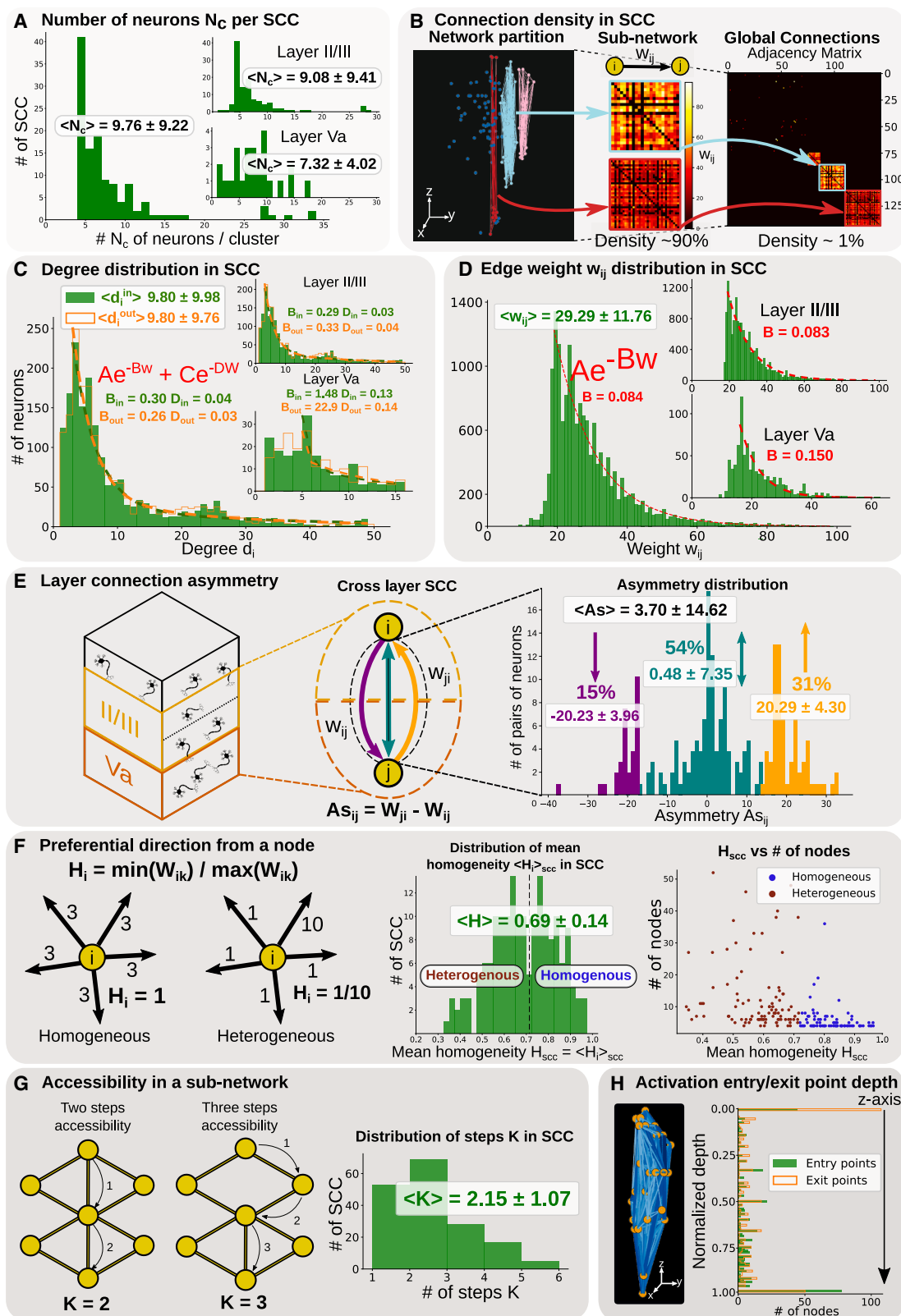
(C) Left: illustration of spike-time-tiling-coefficient (STTC).³¹ Middle: average STTC for all spike trains. Right: spike trains sorted by sub-figure A and their associated STTC (pink, red, green, and blue).

(D) Network coarse-graining pipeline. From left to right: spike trains sorted by SCC, binarization of activation event for a number of consecutive activated neurons $> T_{\text{event}} = 2$, binarization results leading to coarse-grained representation of time events, STTC matrix associated to coarse-grained network reduction and final 3D network decomposed into SCC (see also Figure S5 for subnetwork labelling).

neurons), 85% of modules localized entirely to layer II/III, 7% to layer Va, and 8% spanned both layers (Figure 3B).

This algorithm clearly extracts with high confidence temporally correlated subnetworks (see Figures 3C and S5). Its perfor-

mance was evaluated by comparing the spike-time-tiling-coefficient (STTC)³¹ computed between population-level spike trains and intra-subnetwork spike trains, revealing an increase in temporal coherence within the identified subnetworks.



(legend on next page)

Compared to the full network (edge density 0.010 ± 0.002), SCCs exhibited dramatically higher connectivity (mean density 0.725 ± 0.188), i.e., ~ 70 -fold increase in local connection density (Figure 4B). Within SCCs, both in- and out-degree (i.e., number of incoming and outgoing connections) distributions peaked at three connections—versus one in the global graph—reflecting the removal of dead-end neurons (Figure 4C). Although the average degree remained comparable, SCCs preserved the “bump” in in-degree at 18–30 links, confirming that the hub-driven regime arises from these tightly layer II/III modules. Edge-weight distributions within SCCs (Figure 4D) mirrored those of the overall network, indicating that our partitioning retains the full spectrum of functional link strengths.

Directional flow and accessibility within subnetworks

Focusing on SCCs that span both layers, we quantified inter-layer bias using an *asymmetry* index (STAR Methods), defined as the difference between the adjacent matrix coefficient $A_{ij}-A_{ji}$ for each cross-layer pair. Positive values denote net Va $\rightarrow$ II/III flow, negative values II/III $\rightarrow$ Va. Across all cross-layer SCCs, the asymmetry distribution is significantly shifted toward positive values (Figure 4E), confirming a predominant feedback information bias from layer Va to II/III within M1. Bidirectional edges—where $A_{ij} \sim A_{ji}$ —form a synchronously activating backbone that does not contribute to net flux.

To assess local wiring uniformity (whether or not each routing path are equally used), we introduced a *homogeneity* metric, the ratio of a node’s smallest to largest outgoing weight. SCCs are split into two groups: “homogeneous” modules with evenly balanced links, and “heterogeneous” ones with pronounced hub-to-periphery pathways (Figure 4F). Small SCCs (<8 neurons) are equally split between these types, whereas larger SCCs (>8 neurons) are predominantly heterogeneous, suggesting that extended modules develop stronger preferential routes over longer activation windows.

The two metrics introduced above should be interpreted with caution as they may reflect differences in excitability between neurons in response to a shared external input rather than true directionality. However, as we consider here a sufficiently long period of time to reach stationary excitability, biases in functional connectivity are more plausibly attributed to proportional differences in connectivity strength rather than exclusively higher excitability. We emphasize that, in the absence of electrophysiological measurements, it is not possible to unambiguously distinguish increased excitability from genuine functional interactions between neurons.

Finally, by normalizing the adjacent matrix A to a transition matrix P and computing the smallest iteration k such that all entries of the product P^k are positive,³² we estimate the maximal number of synaptic relays required for full spiking redistribution. The distribution of k peaks at ~ 2.2 (Figure 4G), with an upper bound at six, and thus much fewer intermediate neurons suffice to bridge any two neurons across network layers. Consistent across SCCs, event-entry neurons are disproportionately found in deeper (Va) positions, while exit neurons cluster superficially (II/III) (Figure 4H), mirroring the net upward flow.

Diverse column structural organization revealed by unsupervised clustering

To further catalog recurrent subnetwork archetypes, we applied an unsupervised classification to SCC feature vectors (Figure 5A; STAR Methods). Initially, SCCs were grouped by laminar extent—Va-only, II/III-only, or cross-layer—revealing that 85% reside in II/III (Figure 3B). Focusing on the overrepresented II/III components, we extracted 27 geometric and functional metrics (e.g., component size, spatial spread, link density, homogeneity, spectral gap) and reduced them via a PCA analysis. Then we applied a K-means clustering (elbow method, $K = 5$) and partitioned these into five II/III-specific classes, which together with the two laminar classes yielded seven distinct subnetwork types (Figure 5B; Table 2).

We next analyzed activation “events,” defined as sequences of ≥ 2 SCC neuron firings that end when no further spikes occur. Notably, type 4 SCCs alone produced events recruiting ~ 20 –30 neurons synchronously spiking—precisely accounting for the in-degree “bump” seen earlier. This behavior leads to highly correlated sub-networks where every neuron is strongly correlated. Other classes segregated into “short” motifs (compact, brief activations) and “elongated” motifs (longer, vertically extended sequences). Remarkably, when we normalized event size and duration by SCC neuron count, all seven types collapsed onto a common distribution, and both mean event size and duration scaled linearly with module size (Figure 5B). These results suggest that—under passive conditions—the temporal dynamics of cortical microcircuits depend primarily on subnetwork size rather than laminar identity.

Topological heterogeneity revealed by the statistics of transition between neurons

Beyond size and density, we probed subnetwork topology via the mathematical theory of spectrum focusing on the properties of the transition probability matrices P (Figure 6A; STAR Methods)

Figure 4. Structural and connectivity flow in cortical sub-networks

- (A) Distribution of SCC’s number of neurons N_c across layers.
- (B) High density adjacency matrices of two SCCs and adjacency matrix of the entire network sorted by sub-networks.
- (C) In (green) and out (orange) degree distributions for all SCCs ($n = 175$) over the entire layers and within layer II/III and Va. A sum of two exponential is fitted over each degree distribution.
- (D) Edge weight w_{ij} distribution for all SCCs fitted by a single exponential (red).
- (E) Distribution of the asymmetry weight As_{ij} between pairs of neurons of different layers within SCCs.
- (F) Left: Statistics of preferential direction in neurons characterized by the homogeneity $H_i = \min(w_{i,k})/\max(w_{i,k})$ coefficient. Middle: Distribution of the average homogeneity H_i within SCCs, separated into two populations at threshold $H_{SCC} = 0.7$ described either as more heterogeneous or homogeneous. Right: Homogeneity vs. heterogeneity with respect to the number of neurons in SCCs.
- (G) Left: minimal number of steps K to cross a graph: two examples. Right: distribution of K within SCCs.
- (H) Distribution of normalized entry and exit activation across SCCs in the z direction.

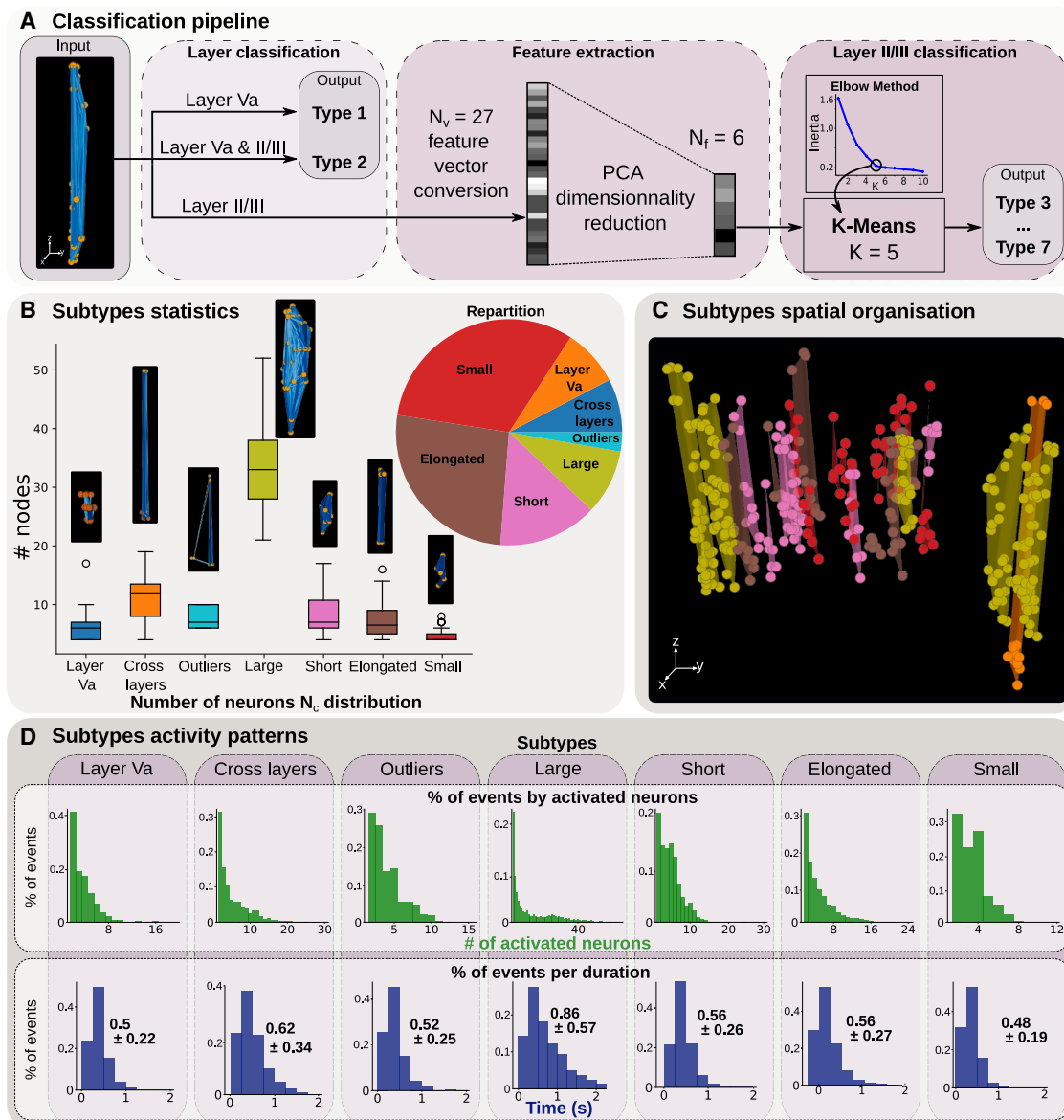


Figure 5. Subtypes network classification

(A) Classification pipeline of sub-networks in 7 types: Step 1: classifying sub-networks according to their layers. Step 2: extraction of 27 features (see Table 3 and STAR Methods) for sub-networks in layer II/III. Dimension reduction to dimension 6 using PCA. Step 3: clustering of layer II/III sub-networks using K-means algorithm parametrized with an elbow method ($K = 5$, see Figure S6).

(B) Left: distributions of number of neurons per subtypes with their associated shape examples. Right: proportion of clusters per subtype.

(C) Example of spatial organization of subnetworks colored by subtypes described in (B).

(D) Sub-network event distribution for the 7 subtypes: normalized distribution of activated neurons per events and duration (see Table 2 for more details).

constructed from the connectivity matrix introduced above. Interestingly the second largest eigenvalue modulus (SLEM) of such matrix serves as a proxy for the mixing time (time for spiking in a single neuron to reach any other neuron inside the cluster) and reveals features such as bottlenecks in the network, well mixed organization (expansion quality), and homogenize connections between the neighbors of each neurons in the network (degree regularity).^{33–36} Across all SCCs, SLEM values spanned a continuum (Figure 6B), delineating two principal regimes:

1. Low-SLEM, degree-regular modules: characterized by rapid spike mixing and uniform link patterns, indicative of tightly knit, homogeneous clusters.
2. High-SLEM, slow-mixing or expander-like modules: featuring bottlenecks or hierarchical hubs that prolong information retention and support more complex routing.

Importantly, SLEM showed no clear dependence on SCC size, suggesting that both small and large neuronal subnetworks can

Table 2. Table of sub-networks subtypes characteristics

	Type 1	Type 2	Type 3	Type 4	Type 5	Type 6	Type 7
Name	layer Va	cross layer	outliers	large	short	elongated	small
Layer	Va	both	II/III	II/III	II/III	II/III	II/III
# Neurons N_c	~ 7	~ 12	~ 8	$\gg 8$	~ 8	~ 8	$\ll 8$
Spatial organization	short	elongated	random	elongated	short	elongated	short
Frequency %	6.3	8.0	3.4	9.7	12.0	29.7	30.9

manifest either topology. This spectral diversity underscores the repertoire of functional motifs embedded within the cortical columnar scaffold (see also Figure S7).

Subnetworks as fundamental units of behavior correlation

After characterizing sub-networks as independent structures with distinct structure and activity patterns, a coarse-grained representation of the network was incorporated into the analysis (Figures 3C and 3D). In this representation, an undirected graph is constructed using the STTC³¹ as a measure of functional synchronicity. Each subnetwork was abstracted as a single node and associated with a binary activity signal that was set to 1 whenever two or more neurons fired consecutively. This coarse-grained view revealed that certain subnetworks tended to co-activate in parallel, whereas others activated independently.

Simultaneously with volumetric calcium imaging, facial movement videos were acquired and a continuous snout-movement signal was extracted as previously described.¹¹ Continuous cross-correlograms computed from population-averaged neural activity (Figure 7A) exhibited clear correlation peaks within a 5-s window around snout movements, as expected for the snout-associated region of M1. We then used the binary activity of sub-networks to compute cross-correlograms with the snout trace (see STAR Methods). Prior anatomical and slice-mapping studies report strong local recurrence in supragranular layers

and prominent feedforward L2/3→L5 pathways in sensory cortex, implying that supragranular neurons often precede deeper layer activation.^{37–39}

To test whether subnetwork activation sequences show a consistent layer-order (II/III → Va versus Va → II/III), we examined the first-active neuron within each SCC event (STAR Methods). We used a bootstrap procedure to identify significant peaks and troughs in the subnetwork-snout cross-correlograms and classified subnetworks into five motifs: pre-bump, pre-trough, post-bump, post-trough, and unclassified (Figure 7B). These motifs describe subnetworks that reliably increase or decrease activation before or after snout movement, and a large class that shows no significant association. Notably, post-bump sub-networks are roughly twice as frequent as the other significant motifs (Table 4). Because our recordings use 250 ms time bins, temporal ordering at fine (<250 ms) timescales is limited.

Finally, we quantified the asymmetry (STAR Methods) of edges between L2/3 and Va neurons for each snout-correlation class. The mean ($\pm$ SD) asymmetry values were: pre-bump, 0.00 ± 0.00 ; pre-trough, 11.67 ± 8.89 ; post-bump, -2.29 ± 16.15 ; post-trough, 0.00 ± 0.00 ; no correlation, 4.95 ± 13.78 . These results show that a minority of subnetworks are reliably associated with snout movements. In that subset, feedback-dominated edges (Va → II/III) are relatively suppressed prior to movement, whereas feed-forward-dominated edges (II/III → Va) tend to be enhanced after movement—a pattern

Table 3. Features extracted from strongly connected components

Feature type	Name	Symbol
Geometrical	number of neurons	N_c
	maximal distance and vector between neurons	d_{max}
	minimal distance and vector between neurons	d_{min}
	angle between the maximal and minimal vector between neurons	θ
	center of mass	m_1
	volume	V
	XY area	A_{xy}
Graph (method network statistics)	graph density	D
	mean in- and out- degree and standard deviation	d^{in} and d^{out}
	mean weight and standard deviation	$\bar{w}$
	mean homogeneity	$\bar{H}$
	number of traversal steps	$\hat{k}$

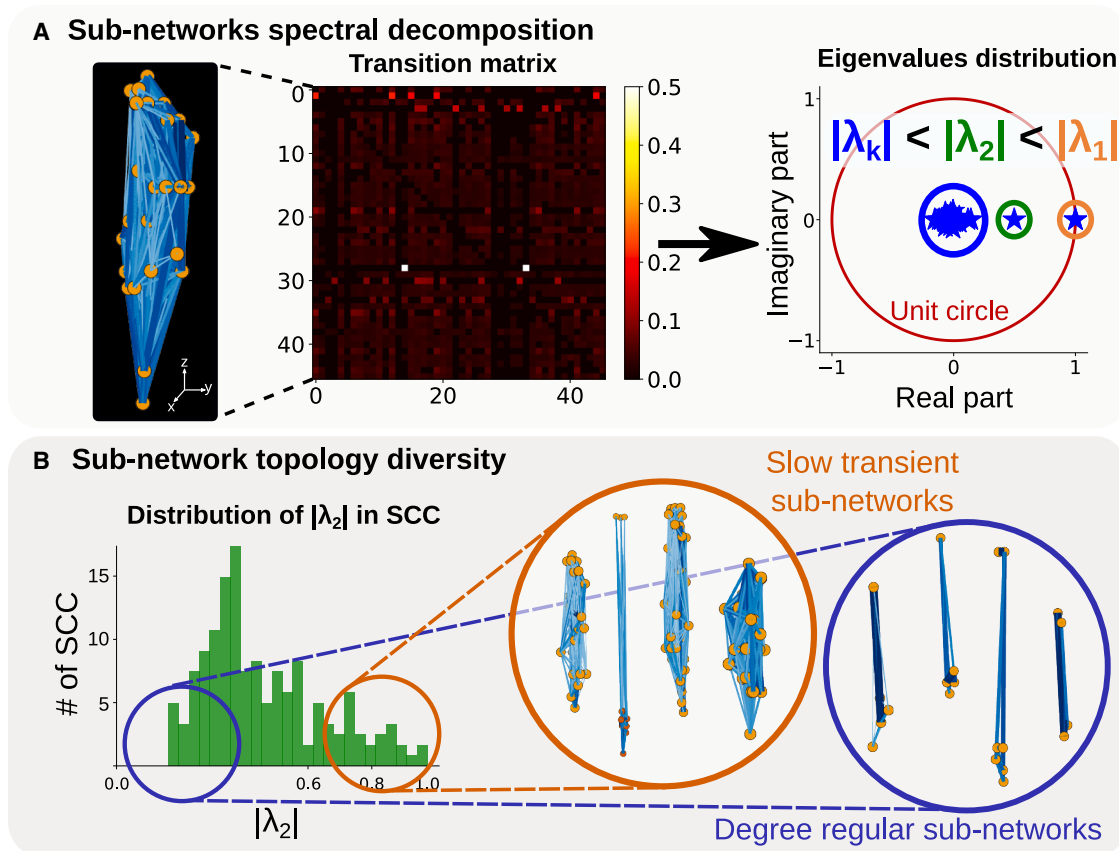


Figure 6. Spectral flow of neuronal activation across cortical sub-units

(A) Left: example of a sub-network representation with its normalized transition probability matrix P_{ij} accounting for the transition between each neuron of the network. Right: distribution of the eigenvalues λ_k of the matrix P_{ij} represented on the complex plane. The first eigenvalue is 1 and the second (green) represents the speed of activation flow.

(B) Left: distribution of second largest eigenvalue modulus $|\lambda_2|$ for all sub-networks. Right: sub-networks topology trends for small (degree regular network^{34,36}) and large (slow transient networks or expander graphs³⁵) $|\lambda_2|$.

consistent with reports of predominantly feedforward signaling during movement and of state-dependent L5→L2/3 influences *in vivo*.^{37,40–43} Non-associated subnetworks exhibit modest positive asymmetry on average. Because these analyses are correlational and constrained by the temporal resolution of calcium imaging, they do not establish causality; targeted perturbations (for example, two-photon holographic stimulation or paired electrophysiology) would be required to disentangle causal feedforward vs. feedback roles.

We additionally computed cross-correlograms for several neuronal subpopulations and subnetwork groups (Figure S8), including the unclassified subnetworks (57% of detected subnetworks). These analyses show that neurons outside the identified subnetworks can still correlate with snout movement, suggesting contributions either from single neurons to the behavior or from functional networks that extend beyond the imaged volume. Conversely, subnetworks uncorrelated with snout movement are likely engaged in other computations. Acquisition of evoked-potential or perturbation data would be required to further validate and causally dissect these observations.

DISCUSSION

In this study, we combined fast volumetric two-photon calcium imaging with a novel graph-based pipeline to reconstruct 3D functional connectivity in awake mouse primary motor cortex. From ~1,000 neurons imaged over six 20-min sessions each recorded on different animals, we obtain a novel decomposition of cortical column network in sub-population derived from reconstructing weighted graphs and decomposed them into SCCs. Such decomposition allowed to uncover seven reproducible subnetwork types—ranging from compact columns to diagonal and triangular motifs—predominantly in layer II/III but often bridging to layer Va. Quantitative analyses revealed a net feedback (Va → II/III) information bias, rapid access across neurons in ≤6 steps (with an average of 2), and both homogeneous and hub-dominated modules distinguished by their spectral profiles. The seven geometric organizational motifs are associated with specific correlation activity patterns (event sizes and durations) as well as specific response to spontaneous behavior. The present analysis reveals how microcircuits function within columns.

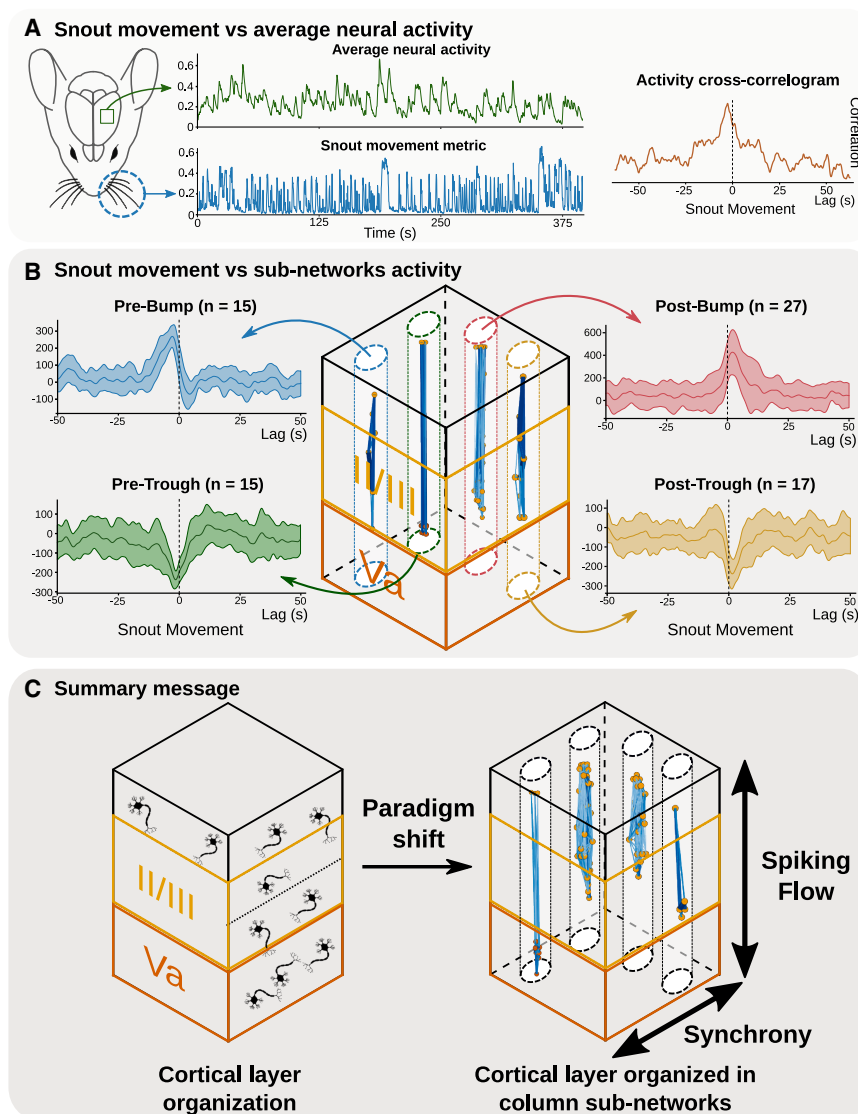


Figure 7. Correlation between snout activity and reconstructed subnetworks

(A) Left: example of population average neural activity and the labeled snout movement metric continuous signal. Right: cross-correlogram of the neural activity around the snout movement metric continuous signal.

(B) Decomposed cross-correlogram of sub-network activity around snout movement metric continuous signal in four classes: pre-bump ($n = 15$), pre-trough ($n = 15$), post-bump ($n = 27$) and post-trough ($n = 17$).

(C) **Summary message:** cortical layers can be decomposed into sub-networks organized in columns with a predominant feedback flow that can further fire or being suppressed in synchrony. These sub-networks appears as computational sub-units.

Bridging simultaneous volumetric imaging and functional inference

Prior sequential two-photon studies in sensory areas provided key evidence for cortical columns via receptive-field mapping,^{6,7} but their non-simultaneous sampling limits inference of temporal and causal interactions. *In vitro* calcium-imaging first demonstrated how optical signals could reveal functional microcircuits⁴⁴ and later how cortical layer II/III from primary visual cortex consists of an on- and off-group of neurons leading to higher orientation selectivity and narrower bandwidth than individual neurons.⁴⁵ Finally, by applying our pipeline to volumetric data

recorded by sTeFo 2P *in vivo*, we now capture real-time, layer-spanning functional columnar motifs in M1 associated with specific activity patterns and presenting different forms of correlation to spontaneous behavior.

In the barrel cortex, layer IV barrels show whisker-specific clustering, while those in layer II/III more diffuse tuning,^{46,47} highlighting departures from strict radial modules. Our findings in M1 similarly reveal not only compact column-like subnetworks but also elongated diagonal and triangular motifs. These diverse geometries echo anatomical observations of oblique intracortical projections⁴⁸ and point to a balance

Table 4. Sub-networks correlation type per class

	Type 1	Type 2	Type 3	Type 4	Type 5	Type 6	Type 7	Total
Pre-bump	1	0	2	0	2	8	2	15
Pre-trough	0	1	0	0	1	3	10	15
Post-bump	4	4	0	2	2	8	7	27
Post-trough	0	1	1	2	1	4	8	17
Not correlated	7	8	3	16	11	29	27	101

between vertical integration and lateral diversification in neocortical organization.

Intrinsic versus task-evoked dynamics

Unlike task-driven mapping studies that combine encoding of kinematics or forces,⁴⁹ the present functional networks are computed from ongoing, spontaneous activity, probably similar to the UP-state attractor dynamics originally characterized in slice preparations.⁴⁴ Manifestations of spontaneous activity thus provide an echo of responses⁵⁰ in ongoing cortical ensembles that reflect similar activity to sensory evoked one in groups of neurons with strong synaptic connectivity,⁵¹ as also shown by our reconstruction procedure (see Figure 7) from spontaneous activity that allows anticipating motion during snout activation. Overlaying task-evoked patterns onto network graph ensembles in future work should clarify how behavior recruits and reshapes specific subnetwork channels.

Functional implications of subnetwork diversity and comparison with anatomical connectivity

The net information flow (not the connectivity strength) depends on the presence of layer II/III–V loop in M1 that has a stronger feedforward connectivity from layer II/III to layer V than the feedback from layer V to layer II/III.³⁸ Here, we revealed a stronger feedback than feedforward signaling for spontaneous M1 dynamics. Such vertical dynamics were not found in the inferior colliculus (see Figure S4), supporting that our finding in M1 was not an experimental or analysis artifact (e.g., due to low frame rates).

Notably, while many slice and optogenetic mapping studies emphasize feedforward layer II/III to layer V routing,^{37–39} recent *in vivo* work shows that layer V can exert strong influence over superficial layers in specific states and areas.^{40–43} Our observation of a net $V_a \rightarrow II/III$ bias therefore likely reflects the operating regime of spontaneously active, awake M1 and highlights the necessity of causal stimulations across states to disambiguate anatomical potential from state-dependent information flow. Our method leads to a quantification 56 versus 44% suggesting a 10% difference toward feedback signaling pathway. Such results should be further confirmed in future studies to reveal the underlying specific network connectivity.

Interestingly, the flow of information is structured by the seven SCC archetypes that exhibit distinct size, density, temporal profiles, and spectral signatures. Homogeneous modules may serve rapid feedforward relay, whereas heterogeneous, hub-dominated clusters could implement flexible routing or gain control, similar to dense inhibitory subnetworks.⁵² Low-SLEM (fast-mixing) motifs may support rapid sensorimotor

loops, while high-SLEM (slow-mixing) modules facilitate temporal integration.

Anatomical connectivity has long been a gold-standard for neuronal circuit reconstruction, inferred from fixed-tissue reconstructions or transsynaptic tracing.^{53–57} These approaches provide static wiring diagrams but little insight into dynamic ensemble activity. While connectomic maps offer structural constraints,⁵⁸ they themselves do not necessarily reveal which connections are functional *in vivo*, nor how neuronal ensembles encode stimuli over time.⁵⁹ Alternative approaches include functional circuit inference based on electrophysiology or functional magnetic resonance imaging (fMRI) data.^{60,61} Descriptive statistics, such as cross-correlogram, have been traditionally used to reveal underlying interactions between neurons, or information flow and hierarchy between brain areas.^{62–64} As the number of recorded neurons increases,⁶⁵ interpretation of such statistical structures becomes non-trivial,^{66,67} even with the aid of anatomical model-based constraints.⁶⁸ More recent studies build on advances in calcium imaging and holographic optogenetics to infer and functionally validate cortical network structure from the dynamics of population activity. By selectively inactivating neurons and measuring the resulting effects on sensory encoding, it turns out that functional ensembles are not just subsets of anatomically connected cells, but are defined by their causal contribution to circuit computations,⁴⁵ and related work.^{51,69,70}

These efforts, complemented by algorithmic advances in spike inference and signal deconvolution,⁷¹ provide a comprehensive framework for mapping and interrogating dynamic network structure in behaving animals. Finally, our approach provides a dynamic, behaviorally relevant reconstruction of functional connectivity,⁷² enabling the identification of ensemble-specific roles that anatomical methods alone cannot resolve. However, the functional connectivity inferred from spontaneous activity cannot uniquely identify the underlying anatomical substrate. Here the recurrence-based thresholding and SCC decomposition is associated with multiple repeated events that could be due to recurring pathways, or common inputs. A causal validation could be possible using two-photon targeted holographic stimulation/optogenetics, paired whole-cell recordings as future experiments.

Limitations of the study

The present adaptive, layer-specific reconstruction provides high-confidence links but may omit weaker connections, transient interactions, and under-represented neurons projecting outside the imaged volume. The 4 Hz sampling rate constrains temporal precision for very fast synaptic events, and deconvolution accuracy depends on calcium indicator's kinetics. Combining this framework with higher speed indicators or simultaneous electrophysiology could further refine connectivity estimates.

This statistical approach also neglects activity that could vary with brain state or some task-related networks, where we could expect specific networks to be activated. It would be interesting to apply the present reconstruction algorithm when the animal is engaged in a task, where we expect different types but similar networks for similar tasks. This will allow us to quantify the size of the network, and compare it across different tasks and further relate it with spontaneous networks.

Future directions

The present network reconstruction pipeline could be applied to other cortical and subcortical regions, across developmental stages, or in disease models (e.g., Parkinson's, stroke) to track microcircuit motif evolution, or organization driven by a given frequency stimulation. Integrating cell-type-specific perturbations will disentangle the roles of interneurons and projection neurons in each motif. Ultimately, correlating these subnetwork archetypes with behavior could elucidate elementary processing units of the cortex and their disruption in pathology.

RESOURCE AVAILABILITY

Lead contact

Requests for further information and resources should be directed to and will be fulfilled by the lead contact, David Holcman (david.holcman@ens.fr).

Materials availability

This study did not generate new materials.

Data and code availability

- All the raw imaging data and CalmAn intermediate files (motion corrected data, rois, and signal extracted), together with the face video data supporting the paper are available at the link: <https://www.ebi.ac.uk/biostudies/bioimages/studies/S-BIAD2355>.
- All original code has been deposited at Zenodo under the associated doi <https://zenodo.org/records/15691025>
- Any additional information required to reanalyze the data reported in this paper is available from the [lead contact](#) upon request.

ACKNOWLEDGMENTS

D.H. is supported by ANR AstroXcite and the European Research Council (ERC) under the European Union's Horizon 2020 research and innovation program (no. 882673). J.C.B. acknowledges supporting fellowships from the EMBL Interdisciplinary Postdoc (EIPOD) Program under Marie Skłodowska Curie Cofund Actions MSCA-COFUND-FP (664726). J.C.B. and R.P. were funded by the European Molecular Biology Laboratory (EMBL) and the Deutsche Forschungsgemeinschaft (DFG, German Research Foundation)—projects 458898724 and 425902099 awarded to R.P.

AUTHOR CONTRIBUTIONS

Conceptualization, P.A. and D.H.; methodology, P.A. and D.H.; investigation, P.A.; writing – original draft, P.A., D.H., and J.C.B.; writing – review and editing, P.A., D.H., T.L., J.C.B., and H.A.; funding acquisition, D.H.; resources, J.C.B., H.A., and R.P.; supervision, D.H., R.P., and H.A.

DECLARATION OF INTERESTS

The authors declare no competing interests.

DECLARATION OF GENERATIVE AI AND AI-ASSISTED TECHNOLOGIES IN THE WRITING PROCESS

During the preparation of this work, the authors used ChatGPT5.1 in order to improve some aspect of writing. After using this tool or service, the authors reviewed and edited the content as needed and take full responsibility for the content of the publication.

STAR★METHODS

Detailed methods are provided in the online version of this paper and include the following:

- **KEY RESOURCES TABLE**
- **EXPERIMENTAL MODEL AND STUDY PARTICIPANT DETAILS**
 - Animals and ethics statement
- **METHOD DETAILS**
 - Surgery
 - *In vivo* Ca²⁺-imaging of M1 populations
- **QUANTIFICATION AND STATISTICAL ANALYSIS**
 - Ca²⁺-imaging signal extraction
 - Time series pre-processing based on background statistics
 - Identifying time correlation between neurons
 - Reconstructing the neural network
 - Threshold computation criteria
 - Adapting the threshold to layer dynamics
 - Random networks as a statistical comparison
 - High-connectivity sub-networks obtained from graph decomposition
 - Coarse-grained network using sub-networks as elementary components
 - Subtypes of cortical layer sub-networks revealed by K-means classification
 - Events activation within sub-networks
 - Sub-network activity correlation with snout movements
 - Network statistics

SUPPLEMENTAL INFORMATION

Supplemental information can be found online at <https://doi.org/10.1016/j.celrep.2026.117618>.

Received: August 1, 2025

Revised: May 6, 2026

Accepted: June 9, 2026

REFERENCES

1. Mountcastle, V.B. (1957). Modality and topographic properties of single neurons of cat's somatic sensory cortex. *J. Neurophysiol.* 20, 408–434.
2. Miller, E.K. (2000). The prefrontal cortex and cognitive control. *Nat. Rev. Neurosci.* 1, 59–65.
3. Peters, A.J., Liu, H., and Komiyama, T. (2017). Learning in the rodent motor cortex. *Annu. Rev. Neurosci.* 40, 77–97. <https://doi.org/10.1146/annurev-neuro-072116-031407>.
4. Gilbert, C.D., and Sigman, M. (2007). Brain states: top-down influences in sensory processing. *Neuron* 54, 677–696.
5. Hubel, D.H., Wiesel, T.N., and Stryker, M.P. (1977). Orientation columns in macaque monkey visual cortex demonstrated by the 2-deoxyglucose autoradiographic technique. *Nature* 269, 328–330.
6. Ohki, K., Chung, S., Kara, P., Hübener, M., Bonhoeffer, T., and Reid, R.C. (2006). Highly ordered arrangement of single neurons in orientation pinwheels. *Nature* 442, 925–928.
7. Tischbirek, C.H., Noda, T., Tohmi, M., Birkner, A., Nelken, I., and Konnerth, A. (2019). *In vivo* functional mapping of a cortical column at Single-Neuron resolution. *Cell Rep.* 27, 1319–1326.e5.
8. Prevedel, R., Verhoef, A.J., Pernia-Andrade, A.J., Weisenburger, S., Huang, B.S., Nöbauer, T., Fernández, A., Delcour, J.E., Golshani, P., Baltuska, A., and Vaziri, A. (2016). Fast volumetric calcium imaging across multiple cortical layers using sculpted light. *Nat. Methods* 13, 1021–1028.
9. Weisenburger, S., Prevedel, R., and Vaziri, A. (2017). Quantitative evaluation of two-photon calcium imaging modalities for high-speed volumetric calcium imaging in scattering brain tissue. Preprint at bioRxiv. <https://doi.org/10.1101/115659>.
10. Boffi, J.C., Wiessalla, T., and Prevedel, R. (2021). Primary motor cortex activity traces distinct trajectories of population dynamics during

- p>spontaneous facial motor sequences in mice. Preprint at bioRxiv.
- <https://www.biorxiv.org/content/early/2021/02/16/2021.02.15.431209>
- .
11. Boffi, J.C., Bathellier, B., Asari, H., and Prevedel, R. (2024). Noisy neuronal populations effectively encode sound localization in the dorsal inferior colliculus of awake mice. *eLife* 13, RP97598. <https://doi.org/10.7554/eLife.97598.3>.
 12. Rusakov, D.A. (2012). Depletion of extracellular Ca^{2+} prompts astroglia to moderate synaptic network activity. *Sci. Signal.* 5, pe4.
 13. Stetter, O., Battaglia, D., Soriano, J., and Geisel, T. (2012). Model-free reconstruction of excitatory neuronal connectivity from calcium imaging signals. *PLoS Comput. Biol.* 8, e1002653. <https://doi.org/10.1371/journal.pcbi.1002653>.
 14. Chen, X., Mu, Y., Hu, Y., Kuan, A.T., Nikitchenko, M., Randlett, O., Chen, A.B., Gavornik, J.P., Sompolinsky, H., Engert, F., and Ahrens, M.B. (2018). Brain-wide organization of neuronal activity and convergent sensorimotor transformations in larval zebrafish. *Neuron* 100, 876–890.e5. <https://doi.org/10.1016/j.neuron.2018.09.042>.
 15. Nelson, C.J., and Bonner, S. (2021). Neuronal graphs: A graph theory primer for microscopic, functional networks of neurons recorded by calcium imaging. *Front. Neural Circ.* 15, 662882. <https://doi.org/10.3389/fncir.2021.662882>.
 16. Blevins, A. S., Bassett, D. S., Scott, E. K., and Vanwalleghem, G. C. (2022). From calcium imaging to graph topology. URL: <https://doi.org/10.1162/netn>. doi:10.1162/netn.
 17. Chen, X., Ginoux, F., Carbo-Tano, M., Mora, T., Walczak, A.M., and Wyart, C. (2023). Granger causality analysis for calcium transients in neuronal networks, challenges and improvements. *eLife* 12, e81279. <https://doi.org/10.7554/eLife.81279>.
 18. Zonca, L., Bellier, F.C., Milior, G., Aymard, P., Visser, J., Rancillac, A., Rouach, N., and Holcman, D. (2025). Unveiling the functional connectivity of astrocytic networks with astronot, a graph reconstruction algorithm coupled to image processing. *Commun. Biol.* 8, 114. <https://doi.org/10.1038/s42003-024-07390-0>.
 19. Kim, S., Putrino, D., Ghosh, S., and Brown, E.N. (2011). A granger causality measure for point process models of ensemble neural spiking activity. *PLoS Comput. Biol.* 7, e1001110. <https://doi.org/10.1371/journal.pcbi.1001110>.
 20. Bickel, P.J., and Sarkar, P. (2015). Hypothesis testing for automated community detection in networks. *J. R. Stat. Soc. Ser. B Stat. Methodol.* 78, 253–273. <https://doi.org/10.1111/rssb.12117>.
 21. Mölter, J., Avitan, L., and Goodhill, G.J. (2018). Detecting neural assemblies in calcium imaging data. *BMC Biol.* 16, 143. <https://doi.org/10.1186/s12915-018-0606-4>.
 22. Savtchenko, L.P., and Rusakov, D.A. (2014). Regulation of rhythm genesis by volume-limited, astroglia-like signals in neural networks. *Philos. Trans. R. Soc. Lond. B Biol. Sci.* 369, 20130614.
 23. Magloire, V., Savtchenko, L.P., Jensen, T.P., Sylantsev, S., Kopach, O., Cole, N., Tyurikova, O., Kullmann, D.M., Walker, M.C., Marvin, J.S., et al. (2023). Volume-transmitted gaba waves pace epileptiform rhythms in the hippocampal network. *Curr. Biol.* 33, 1249–1264.e7.
 24. Jones, E.G. (2000). Microcolumns in the cerebral cortex. *Proc. Natl. Acad. Sci. USA* 97, 5019–5021.
 25. Buxhoeveden, D.P., and Casanova, M.F. (2002). The minicolumn hypothesis in neuroscience. *Brain* 125, 935–951. <https://doi.org/10.1093/brain/awf110>.
 26. Tennant, K.A., Adkins, D.L., Donlan, N.A., Asay, A.L., Thomas, N., Kleim, J.A., and Jones, T.A. (2011). The organization of the forelimb representation of the C57BL/6 mouse motor cortex as defined by intracortical microstimulation and cytoarchitecture. *Cereb. Cortex* 21, 865–876.
 27. Muñoz-Castañeda, R., Zingg, B., Matho, K.S., Chen, X., Wang, Q., Foster, N.N., Li, A., Narasimhan, A., Hirokawa, K.E., Huo, B., et al. (2021). Cellular anatomy of the mouse primary motor cortex. *Nature* 598, 159–166.
 28. Giovannucci, A., Friedrich, J., Gunn, P., Kalfon, J., Brown, B.L., Koay, S.A., Taxis, J., Najafi, F., Gauthier, J.L., Zhou, P., et al. (2019). Caiman an open source tool for scalable calcium imaging data analysis. *eLife* 8, e38173. <https://doi.org/10.7554/eLife.38173>.
 29. Broido, A.D., and Clauset, A. (2019). Scale-free networks are rare. *Nat. Commun.* 10, 1017. <https://doi.org/10.1038/s41467-019-08746-5>.
 30. Ester, M., Kriegel, H.-P., Sander, J., and Xu, X. (1996). A density-based algorithm for discovering clusters in large spatial databases with noise. In *Proceedings of the Second International Conference on Knowledge Discovery and Data Mining (KDD'96 AAAI Press)*, pp. 226–231.
 31. Cutts, C.S., and Egle, S.J. (2014). Detecting pairwise correlations in spike trains: An objective comparison of methods and application to the study of retinal waves. *J. Neurosci.* 34, 14288–14303. <https://www.jneurosci.org/content/34/43/14288.full.pdf>.
 32. Karlin, S., and Taylor, H.A. (1981). *Second Course in Stochastic Processes*. Elsevier Science 2. <https://books.google.fr/books?id=nGy0nAxwO10C>.
 33. Boyd, S., Diaconis, P., and Xiao, L. (2004). Fastest mixing markov chain on a graph. *SIAM Rev.* 46, 667–689. <https://doi.org/10.1137/S0036144503423264>.
 34. Cvetković, D., and Simić, S. (1995). The second largest eigenvalue of a graph (a survey). *Filomat* 9, 449–472. <http://www.jstor.org/stable/43999231>.
 35. Hoory, S., Linial, N., and Wigderson, A. (2006). Expander graphs and their applications. *Bull. Am. Math. Soc.* 43, 439–561. <https://doi.org/10.1090/s0273-0979-06-01126-8>.
 36. Koledin, T., and Stanić, Z. (2013). Regular graphs with small second largest eigenvalue. *Appl. Anal. Discrete Math.* 7, 235–249. <http://www.jstor.org/stable/43660713>.
 37. Lefort, S., Tómm, C., Floyd Sarria, J.-C., and Petersen, C.C.H. (2009). The excitatory neuronal network of the C2 barrel column in mouse primary somatosensory cortex. *Neuron* 61, 301–316.
 38. Weiler, N., Wood, L., Yu, J., Solla, S.A., and Shepherd, G.M.G. (2008). Top-down laminar organization of the excitatory network in motor cortex. *Nat. Neurosci.* 11, 360–366.
 39. Hooks, B.M., Mao, T., Gutnisky, D.A., Yamawaki, N., Svoboda, K., and Shepherd, G.M.G. (2013). Organization of cortical and thalamic input to pyramidal neurons in mouse motor cortex. *J. Neurosci.* 33, 748–760. <https://www.jneurosci.org/content/33/2/748.full.pdf>.
 40. Beltramo, R., D'Urso, G., Dal Maschio, M., Farisello, P., Bovetti, S., Clovis, Y., Lassi, G., Tucci, V., De Pietri Tonelli, D., and Fellin, T. (2013). Layer-specific excitatory circuits differentially control recurrent network dynamics in the neocortex. *Nat. Neurosci.* 16, 227–234.
 41. Marshel, J.H., Kim, Y.S., Machado, T.A., Quirin, S., Benson, B., Kadmon, J., Raja, C., Chibukhchyan, A., Ramakrishnan, C., Inoue, M., et al. (2019). Cortical layer-specific critical dynamics triggering perception. *Science* 365, eaaw5202.
 42. Hage, T.A., Bosma-Moody, A., Baker, C.A., Kratz, M.B., Campagnola, L., Jarsky, T., Zeng, H., and Murphy, G.J. (2022). Synaptic connectivity to I2/3 of primary visual cortex measured by two-photon optogenetic stimulation. *eLife* 11, e71103.
 43. Onodera, K., and Kato, H.K. (2022). Translaminar recurrence from layer 5 suppresses superficial cortical layers. *Nat. Commun.* 13, 2585.
 44. Cossart, R., Aronov, D., and Yuste, R. (2003). Attractor dynamics of network up states in the neocortex. *Nature* 423, 283–288. <https://doi.org/10.1038/nature01614>. <https://pubmed.ncbi.nlm.nih.gov/12748641>.
 45. Pérez-Ortega, J., Akrouh, A., and Yuste, R. (2024). Stimulus encoding by specific inactivation of cortical neurons. *Nat. Commun.* 15, 3192. <https://doi.org/10.1038/s41467-024-47515-x>.
 46. Jia, H., Varga, Z., Sakmann, B., and Konnerth, A. (2014). Linear integration of spine Ca^{2+} signals in layer 4 cortical neurons in vivo. *Proc. Natl. Acad. Sci.* 111, 9277–9282. <https://doi.org/10.1073/pnas.1408525111>.
 47. Varga, Z., Jia, H., Sakmann, B., and Konnerth, A. (2011). Dendritic coding of multiple sensory inputs in single cortical neurons in vivo. *Proc.*

- Natl. Acad. Sci. USA 108, 15420–15425. <https://doi.org/10.1073/pnas.1112355108>.
48. Amirikian, B., and Georgopoulos, A.P. (2003). Modular organization of directionally tuned cells in the motor cortex: Is there a short-range order? Proc. Natl. Acad. Sci. USA 100, 12474–12479. <https://doi.org/10.1073/pnas.2037719100>.
49. Hatsopoulos, N.G. (2010). Columnar organization in the motor cortex. Cortex 46, 270–271.
50. Kenet, T., Bibitchkov, D., Tsodyks, M., Grinvald, A., and Arieli, A. (2003). Spontaneously emerging cortical representations of visual attributes. Nature 425, 954–956.
51. Carrillo-Reid, L., Yang, W., Bando, Y., Peterka, D.S., and Yuste, R. (2016). Imprinting and recalling cortical ensembles. Science 353, 691–694. <https://doi.org/10.1126/science.aaf7560>.
52. Fino, E., and Yuste, R. (2011). Dense inhibitory connectivity in neocortex. Neuron 69, 1188–1203. <https://doi.org/10.1016/j.neuron.2011.02.025>.
53. Smith, S.J. (2007). Circuit reconstruction tools today. Curr. Opin. Neurobiol. 17, 601–608. <https://doi.org/10.1016/j.conb.2007.11.004>.
54. Helmstaedter, M., Briggman, K.L., and Denk, W. (2008). 3d structural imaging of the brain with photons and electrons. Curr. Opin. Neurobiol. 18, 633–641. <https://doi.org/10.1016/j.conb.2009.03.005>.
55. Kleinfeld, D., Bharioke, A., Blinder, P., Bock, D.D., Briggman, K.L., Chklovskii, D.B., Denk, W., Helmstaedter, M., Kaufhold, J.P., Lee, W.-C.A., et al. (2011). Large-scale automated histology in the pursuit of connectomes. J. Neurosci. 31, 16125–16138.
56. Morgan, J.L., and Lichtman, J.W. (2013). Why not connectomics? Nat. Methods 10, 494–500.
57. Luo, L., Callaway, E.M., and Svoboda, K. (2018). Genetic dissection of neural circuits: A decade of progress. Neuron 98, 265–281.
58. Real, E., Asari, H., Gollisch, T., and Meister, M. (2017). Neural circuit inference from function to structure. Curr. Biol. 27, 189–198.
59. Zamora-López, G., Zhou, C., and Kurths, J. (2011). Exploring brain function from anatomical connectivity. Front. Neurosci. 5, 83.
60. Stevenson, I.H., Rebesch, J.M., Miller, L.E., and Körding, K.P. (2008). Inferring functional connections between neurons. Curr. Opin. Neurobiol. 18, 582–588.
61. Engel, T.A., Schölvinck, M.L., and Lewis, C.M. (2021). The diversity and specificity of functional connectivity across spatial and temporal scales. Neuroimage 245, 118692.
62. Perkel, D.H., Gerstein, G.L., and Moore, G.P. (1967). Neuronal spike trains and stochastic point processes. II. simultaneous spike trains. Biophys. J. 7, 419–440.
63. Levick, W.R., Cleland, B.G., and Dubin, M.W. (1972). Lateral geniculate neurons of cat: retinal inputs and physiology. Invest. Ophthalmol. 11, 302–311.
64. Bair, W., Zohary, E., and Newsome, W.T. (2001). Correlated firing in macaque visual area MT: time scales and relationship to behavior. J. Neurosci. 21, 1676–1697.
65. Jun, J.J., Steinmetz, N.A., Siegle, J.H., Denman, D.J., Bauza, M., Barbarits, B., Lee, A.K., Anastassiou, C.A., Andrei, A., Aydin, Ç., et al. (2017). Fully integrated silicon probes for high-density recording of neural activity. Nature 551, 232–236.
66. Brody, C.D. (1999). Correlations without synchrony. Neural Comput. 11, 1537–1551.
67. Ostojic, S., Brunel, N., and Hakim, V. (2009). How connectivity, background activity, and synaptic properties shape the cross-correlation between spike trains. J. Neurosci. 29, 10234–10253.
68. Siegle, J.H., Jia, X., Durand, S., Gale, S., Bennett, C., Graddis, N., Heller, G., Ramirez, T.K., Choi, H., Luviano, J.A., et al. (2021). Survey of spiking in the mouse visual system reveals functional hierarchy. Nature 592, 86–92.
69. Carrillo-Reid, L., Han, S., Yang, W., Akrouh, A., and Yuste, R. (2019). Controlling visually guided behavior by holographic recalling of cortical ensembles. Cell 178, 447–457.e5.
70. Yang, W., and Yuste, R. (2017). In vivo imaging of neural activity. Nat. Methods 14, 349–359. <https://doi.org/10.1038/nmeth.4230>. <https://pubmed.ncbi.nlm.nih.gov/28362436>.
71. Pnevmatikakis, E.A., Soudry, D., Gao, Y., Machado, T.A., Merel, J., Pfau, D., Reardon, T., Mu, Y., Lacefield, C., Yang, W., et al. (2016). Simultaneous denoising, deconvolution, and demixing of calcium imaging data. Neuron 89, 285–299.
72. Hill, D.N., Curtis, J.C., Moore, J.D., and Kleinfeld, D. (2011). Primary motor cortex reports efferent control of vibrissa motion on multiple timescales. Neuron 72, 344–356.
73. Boffi, J.C., Knabbe, J., Kaiser, M., and Kuner, T. (2018). Kcc2-dependent steady-state intracellular chloride concentration and pH in cortical layer 2/3 neurons of anesthetized and awake mice. Front. Cell. Neurosci. 12, 7. <https://doi.org/10.3389/fncel.2018.00007>.
74. Holtmaat, A., Bonhoeffer, T., Chow, D.K., Chuckowree, J., De Paola, V., Hofer, S.B., Hübener, M., Keck, T., Knott, G., Lee, W.-C.A., et al. (2009). Long-term, high-resolution imaging in the mouse neocortex through a chronic cranial window. Nat. Protoc. 4, 1128–1144.
75. Wang, Q., Ding, S.-L., Li, Y., Royall, J., Feng, D., Lesnar, P., Graddis, N., Naeemi, M., Facer, B., Ho, A., et al. (2020). The allen mouse brain common coordinate framework: A 3d reference atlas. Cell 181, 936–953.e20. <https://doi.org/10.1016/j.cell.2020.04.007>.
76. Paxinos, G., and Franklin, K.B.J. (2019). The Mouse Brain in Stereotaxic Coordinates, 5th ed. (Academic Press).
77. Pnevmatikakis, E.A., and Giovannucci, A. (2017). Normcor: An online algorithm for piecewise rigid motion correction of calcium imaging data. J. Neurosci. Methods 291, 83–94. <https://doi.org/10.1016/j.jneumeth.2017.07.031>.
78. Schindelin, J., Arganda-Carreras, I., Frise, E., Kaynig, V., Longair, M., Pietzsch, T., Preibisch, S., Rueden, C., Saalfeld, S., Schmid, B., et al. (2012). Fiji: an open-source platform for biological-image analysis. Nat. Methods 9, 676–682.
79. Virtanen, P., Gommers, R., Oliphant, T.E., Haberland, M., Reddy, T., Cournapeau, D., Burovski, E., Peterson, P., Weckesser, W., Bright, J., et al. (2020). SciPy 1.0: fundamental algorithms for scientific computing in python. Nat. Methods 17, 261–272.
80. Dana, H., Mohar, B., Sun, Y., Narayan, S., Gordus, A., Hasseman, J.P., Tsegaye, G., Holt, G.T., Hu, A., Walpita, D., et al. (2016). Sensitive red protein calcium indicators for imaging neural activity. eLife 5, e12727. <https://doi.org/10.7554/eLife.12727>.
81. Rubinov, M., and Sporns, O. (2010). Complex network measures of brain connectivity: Uses and interpretations. Neuroimage 52, 1059–1069. <https://doi.org/10.1016/j.neuroimage.2009.10.003>.
82. Pearce, D. J. An improved algorithm for finding the strongly connected components of a directed graph (2005):URL: <https://api.semanticscholar.org/CorpusID:7781255>.
83. Rousseeuw, P.J. (1987). Silhouettes: A graphical aid to the interpretation and validation of cluster analysis. J. Comput. Appl. Math. 20, 53–65. [https://doi.org/10.1016/0377-0427\(87\)90125-7](https://doi.org/10.1016/0377-0427(87)90125-7).
84. Steinley, D. (2004). Properties of the Hubert-Arabie adjusted rand index. Psychol. Methods 9, 386–396.
85. Newman, M. (2018). Networks, 2nd ed. (Oxford University Press, Inc.).

STAR★METHODS

KEY RESOURCES TABLE

REAGENT or RESOURCE	SOURCE	IDENTIFIER
Bacterial and virus strains		
AAV1-Syn-jRGECO1a	Addgene	Cat#: #100854-AAV1
Deposited data		
Volumetric sTeFo 2p dataset	This paper	Data available at : https://www.ebi.ac.uk/biostudies/bioimages/studies/S-BIAD2355
Experimental models: Cell lines		
C57Bl6/j mice of either sex	EMBL core colony	https://www.jax.org/strain/000664
Software and algorithms		
CorticalSubNetworksReconstruction	This paper	Github : https://github.com/Umanoides/CorticalSubNetworksReconstruction Zenodo : https://zenodo.org/records/15691025
Python	Python.org	RRID: SCR_008394; https://python.org
SciPy	Virtanen et al., 2020	RRID: SCR_008058; https://scipy.org/
CalmAn	Giovannucci et al., 2019	RRID: SCR_021152; https://github.com/flatironinstitute/CalmAn
Fiji	Schindelin et al., 2012	RRID: SCR_002285; https://doi.org/10.1038/nmeth.2019

EXPERIMENTAL MODEL AND STUDY PARTICIPANT DETAILS

Animals and ethics statement

This work complies with the European Communities Council Directive (2010/63/EU) to minimize animal pain and discomfort. Experimental procedures were approved by EMBL's committee for animal welfare and institutional animal care and use (IACUC), under protocol number RP170001. 8–16 week old C57Bl6/j mice from the EMBL Heidelberg core colonies were used for experiments, housed in groups of 1–5 in makrolon type 2L cages on ventilated racks at room temperature and 50% humidity with a 12 h light cycle. Food and water were available *ad libitum*.

METHOD DETAILS

Surgery

Surgeries were performed as detailed elsewhere^(10,11,73). Briefly, 7–8 week old mice of either sex were used for cranial window surgeries. Animals were anesthetized with a mixture of 40 μ L fentanyl (1 mg/mL; Janssen), 160 μ L midazolam (5 mg/mL; Hameln), and 60 μ L medetomidin (1 mg/mL; Pfizer), dosed at 3.1 μ L/g body weight and injected intraperitoneally. After loss of pain reflexes, the fur over the scalp was removed with hair removal cream, eye ointment was applied (Bepanthen, Bayer) and 1% xylocaine (AstraZeneca) was injected under the scalp as pre-incisional local anesthesia. The mouse was then placed in a stereotaxic apparatus (David Kopf Instruments, model 963) equipped with a heating pad (37°C) to preserve body temperature. The dorsal cranium was exposed by removing the scalp and periosteum with fine forceps and scissors. Next, a 4 mm diameter circular craniectomy was made over M1 using a dental drill (Microtorque, Harvard Apparatus), centered at 1.75 mm anterior and 0.5 mm lateral to bregma, avoiding damage to the dura and bleeding. For Ca²⁺ indicator expression at M1, recombinant AAV vectors (rAAVs, serotype 1) encoding jRGECO1a under the control of the synapsin promoter (Addgene #100854-AAV1) were stereotactically delivered at the center of the craniectomy using glass pipettes lowered to depths of 300, 400, and 500 μ m, at a rate of $\sim$ 4 μ m L/h using a syringe. Approximately 300 nL were injected per spot. After injection, the craniectomy was covered by a round 4 mm coverslip ($\sim$ 170 μ m thick, disinfected with 70% ethanol) with a drop of saline between the glass and the dura. The cranial window was sealed with dental acrylic (Hager Werken Cyano Fast and Paladur acrylic powder) and a custom head fixation bar was also cemented. The surgical wound was closed with dental acrylic. Mice were single-housed after surgery to prevent compromising the implanted windows and had a recovery period of at least 4 weeks before imaging, allowing Ca²⁺ indicator expression and resolution of inflammation associated with surgery.⁷⁴ Post-surgical care included pain relief (Metacam, Boehringer Ingelheim) and antibiotic (Baytril, Bayer) subcutaneous injections (0.1 and 0.5 mg/mL respectively, dosed at 10 μ L/g body weight).

***In vivo* Ca²⁺-imaging of M1 populations**

We used a self-built two-photon microscope based on scanned temporal focusing (sTeFo) enabling fast volumetric *in vivo* Ca²⁺ imaging of large M1 neuronal populations. The technical details and working principles are described in detail elsewhere⁸; however, for our study, we utilized a higher repetition laser (10 MHz, FemtoTrain, Spectra Physics) and operated the microscope with the following parameters: The scanning laser power was kept below 191 mW after the objective to avoid excessive heating of brain tissue.⁸ Fields of view of 470 μm^2 were imaged at 128 px² resolution and spaced in 15 μm steps in the axial direction to cover 585 μm in depth (39 z steps) at a volume rate of 3.99 Hz. Frames acquired during objective flyback were discarded. Typically, Ca²⁺ signals with good signal-to-noise were detected up to a depth of 500 μm from the pial surface. Mice were briefly (<1 min) anesthetized with 5% isoflurane in O₂ for quick head fixation on a custom-built stage, with their bodies inside a 5 cm acrylic tube, under our custom two-photon microscope. After head fixation, mice fully recovered from the isoflurane anesthesia in under a minute, displaying clear blinking reflexes, whisking, sniffing behaviors, and normal body posture and movements. We waited at least 5 min before starting experiments to ensure full recovery from the brief anesthesia. Before Ca²⁺-imaging experiments in the awake condition, a pilot group of mice was habituated to the head-fixed condition in daily 20 min sessions for 3 days. We noted no marked behavioral difference between habituated and unhabituated mice during these relatively short 20 min imaging sessions. Imaging sessions typically lasted 20 min, after which the mouse was returned to its home cage. Typically, mice were imaged a total of four times (sessions), once per week. Among the four datasets, the one with the largest number of neurons detected (peak of AAV expression and cranial window quality) was considered for further analysis. We use a 400 μm depth boundary to account for the large, prominent to layer V, that characterizes the mammalian M1, (300 μm thick in the mouse) and the complex curvature of the cortex at this lateral-medial location. We defined 400 μm deep as a boundary where most cells in the imaging plane would belong to layer Va and cannot be mixed up with cells from layer II/III at that depth. Although this cutoff may bias the classification modestly toward deeper assignments in curved cortex, it ensures that most somata classified as “Va” are not superficial neurons. Refined depth localization can be extracted from standardized atlases.^{75,76}

QUANTIFICATION AND STATISTICAL ANALYSIS

Functional neuronal network connectivity reconstruction based on calcium time series.

We present here a step-by-step description of the computational approach we developed to transform 3D volumetric calcium neuronal data into 3D neuronal functional network (Figure 1).

- 1. Binarizing calcium time series:** we filter calcium responses after motion and neuropil correction using NoRMCorre and CalmAn.^{28,77} The deconvolved Ca²⁺ signal is binarized by extracting large peaks while excluding low-amplitude peaks through thresholding (Figure 1A).
- 2. Removing clustered population spiking based on background statistics:** to avoid counting population responses, we introduce a mapping function Φ that randomly removes spikes present in population spiking intervals using the non-bursting time intervals statistics (Figure 1B).
- 3. Identifying time correlation between neurons:** we develop a sequential analysis on all calcium time series to define the successive activation and co-activation events between neurons. The result is a temporal correlation measure for each pair of neurons i, j based on their sequential and synchronous relationship (Figure 1C).
- 4. Graph reconstruction of the network:** the neural network is constructed by creating a graph with an adjacency matrix A_{ij} . Each element A_{ij} corresponds to the functional temporal correlation measure and represents the strength of the connection between neurons i and j (Figure 1D).
- 5. Removing sparse correlation by thresholding the correlation matrix:** to eliminate spurious connections, we created a statistical null-hypothesis informed by dataset-specific parameters. This model defines a maximum coincidental weight for random connections used as a threshold to distinguish genuine connections from spurious ones (Figure 1D).
- 6. Neuronal network reconstruction using an adaptive threshold for each layer:** to account for variations in firing rate distributions and inter-layer connectivity dynamics, we threshold the graph depending on each neuronal layer statistics (Figure 1E).
- 7. Graph partition into high-connectivity sub-networks:** to identify highly correlated sub-networks, we decompose the graph of the network into strongly connected components (SCCs), see Figure 3A. Outlier sub-networks are identified using a machine learning approach and re-decomposed based on the connections statistics (Figure 3B). Finally, smaller SCCs are excluded to retain only functionally relevant clusters of activity.
- 8. Construction of coarse-grained network:** we extract binary activation signals from SCCs and construct a coarse-grained network using Spike Time Tiling Coefficient³¹ of pairs of signals as connection weight between sub-networks (Figure 3D).
- 9. Classification of sub-networks in subtypes:** after extraction of relevant biological, geometrical and functional features using PCA, SCCs are classified in seven types based on their characteristics using machine learning (Figure 5A).
- 10. Determine correlation of sub-networks with spontaneous movement:** SCC activity is then correlated with continuous snout movement signal quantified from face movement videos (Figure 7).

We now describe the details of each step.

Ca²⁺-imaging signal extraction

Volumetric imaging data were visualized and rendered using FIJI.⁷⁸ Motion correction of *in vivo* awake recordings was performed using NoRMCorre⁷⁷ on each imaging plane in the volumetric database. jRGECO1a signal was extracted, filtered, and neuropil-corrected from each individual imaging plane using CalmAn.²⁸ For each neuron, its corresponding deconvoluted time-series is processed by detecting its peaks using the `find_peaks` function of the *scipy* library⁷⁹ with prominence $p = 0.015$. A spike is detected when the peak reaches the threshold $T_f = 0.05$ (see Table 1). The peaks are used to binarize the signal and create a spike train where the neurons is either activated or not (Figure 1A).

Time series pre-processing based on background statistics

Despite the head-fixed condition of the mice, three out of six datasets contained recordings of clustered spike events of neurons corresponding to different regimes of activation. These regimes are characterized by different mean activity rates, which, with our network reconstruction approach, are more likely to introduce noise rather than convey informative data. However, we believe relevant clues about the underlying network structure can be retrieved from those regimes. Rather than discarding these events,¹³ we propose the following denoising method that preserves the continuity of neuron activation across different regimes. Random clustered spike events are detected from the mean spike train by using the threshold $T_c = 2\sigma$, with σ as the standard deviation of the signal (see Table 1). This approach segregates the time series into two distinct set of intervals: the background regimes and the clustered spike events, each characterized by a different mean firing rate $\bar{\lambda}$. To align the clustered spike intervals with the background regimes, we introduce a mapping function ϕ . For each neuron i , ϕ adjusts the clustered spike intervals by randomly removing spikes until the interval matches the firing rate of the surrounding background regime (Figure 1B; Figure S2). This non-parametric approach is resilient to variation of the background activation and preserves continuity within the signal of a neuron independently of others. This continuity is a necessary hypothesis to define the thresholding method explained later on (STAR Methods Threshold).

Identifying time correlation between neurons

Due to the fast firing rate of neurons compared the limited image acquisition frequency of the calcium signaling, genuine time ordering of spikes may not be feasible when they occur in the same time bin. To address this limitation, the following approach is suggested (Figure 1C). We define the strength of the connection between two neurons i and j as the weight $w_{i,j}$. When $w_{i,j}$ is null, no connections exists between i and j .

1. **Successive correlated events.** When two spikes are occurring one after the other, a single directed edge is created from i to j (Figure 1C - upper), leading in an increment of $w_{i,j}$ by one.
2. **Synchronous correlated events.** In instances where two spikes are occurring simultaneously, no temporal order can be determined within this *same bin interaction*.¹³ Consequently, two directional edges are added between neurons i and j : one from i to j and one from j to i (Figure 1C - center). This results in an increase of $w_{i,j}$ and $w_{j,i}$ by one.

Considering the spatial scale of the data ($\sim 500\mu m^3$), no spatial distance can be determined to influence the causal relationship between subsequent firing neurons. Thus no choice can be made to preferably connect one neuron to another. Moreover, self-connections are excluded, ensuring no neuron forms a connection with itself. Considering that each bin in our dataset is 250 ms long-reflecting the trade-off required to scan a $500\mu m^3$ volume using sTeFo technology-a neuron could exhibit two or more consecutive events within the same bin or could be active only at the beginning or end of the bin, given that calcium transient onset typically last ~ 100 ms.⁸⁰ Successive correlated events approach was therefore employed so that potential true neuronal interactions spanning two bins could be captured, despite the associated increase in error rate (i.e., the introduction of connections between neurons whose activity is mediated by intermediate neurons). This strategy makes sure that the true functional path is included in the set of all created paths. Given the long duration of the recordings (~ 20 min), we assumed that true activation pathways would reoccur consistently over time and compensate the increased error due to successive correlated events. By a thresholding step, only the repeatedly occurring (i.e., strongly correlated) connections are kept.

Reconstructing the neural network

A directed weighted graph G is generated to represent the neural network. The graph is associated with an adjacency matrix A of dimension N^2 with N , the number of neurons. First, the weight between every pair of neurons i and j is computed by applying an iterative algorithm using the previously mentioned principles over the total recording duration T :

$$w_{i,j}^{t+1} = w_{i,j}^t + \begin{cases} 1 & \text{if a correlated event occurs} \\ 0 & \text{otherwise} \end{cases} \quad (\text{Equation 1})$$

Then, the final weight w_{ij}^T is stored in each corresponding A_{ij} . In order to simplify the calculation, the graph can be computed using the spike trains representation as a set of activation times $s = \{s_1, s_2, \dots, s_n\}$. Then, our expression of A_{ij} for a directed network is determined by the spike trains s_i, s_j of neurons i and j and a time shift $\delta T = 1$:

$$A_{ij} = \sum_{t=0}^{\delta T} |(s_i + t) \cap s_j|. \quad (\text{Equation 2})$$

Threshold computation criteria

The resulting adjacent matrix A_{ij} from the previous step contains numerous irrelevant coincidental connections characterized by small weights. We decided to create a deterministic adaptive threshold to remove these spurious connections and thus manage to keep only the ones with the highest probability to be genuine functional pathways (Figure 1D). However, defining such threshold is not trivial.⁸¹ We propose a threshold estimation approach taking into account the dataset characteristics and the difference among neuronal layer populations. This approach computes the maximal coincidental connection that appears from independent Poisson processes, used as a statistical null-hypothesis. Each neuron i in a population L is characterized by a firing rate λ_i . If all neurons in L follows an independent Poisson process, then for any two neurons $i, j \in L^2$, the probability that both neurons fire synchronously or consecutively within ΔT , is given by:

$$Pr\{\text{successive event}\} = Pr\{\text{synchronous event}\} = \lambda_i \lambda_j (\Delta t)^2. \quad (\text{Equation 3})$$

By summing these probabilities, the probability of increasing the weight connection between i and j in our graph construction is:

$$Pr\{\text{creating connection}\} = 2\lambda_i \lambda_j (\Delta t)^2. \quad (\text{Equation 4})$$

Given the small magnitude of this probability, the number of connections can be approximated by a Poisson distribution with parameter $\hat{\lambda}$, which depends on the maximal connection probability and the number of samples. To determine the maximal coincidental connection probability, we consider the two most active (i.e., highest firing rates) neurons inside the population L , characterized by firing rates λ_m, λ_{m-1} . The number of samples $n = \frac{T}{\Delta t}$ is calculated using the total recording duration T and $\Delta t = 1$:

$$\hat{\lambda} = 2\lambda_m \lambda_{m-1} (\Delta t)^2 n = 2\lambda_m \lambda_{m-1} T. \quad (\text{Equation 5})$$

Thus the probability of observing k or less connections between the most active neurons is computed from the cumulative distribution function (CDF) of $Poisson(\hat{\lambda})$ distribution:

$$P(X \leq k) = \sum_{p=0}^k \frac{e^{-\hat{\lambda}} \hat{\lambda}^p}{p!}. \quad (\text{Equation 6})$$

Using the CDF, we calculate the smallest k for which $P(X \geq k) < \alpha$, with the confidence parameter $\alpha = 0.05$ (see Table 1). This approach ensures the choice of threshold k to filter out random connections and can dynamically adapt the threshold to data-specific properties by varying the chosen population and the number of samples. However, this method assumes that the number of samples is sufficient, the firing rates small and stationary.

Adapting the threshold to layer dynamics

The thresholding procedure removes invalid connections appearing fewer than k times in a population L . The threshold acts as a maximum coincidence boundary, filtering out spurious connections and potentially some real connections, while ensuring that the retained connections are statistically robust and unlikely to occur by chance. In our datasets, this threshold is not uniform but varies across layers and between them, reflecting differences in firing rate distributions and connectivity among neuronal layer populations (Figures 1E and 2A). For inter-layer connections, the threshold is derived from the maximal firing rates of the respective populations. By employing a layer adaptive null-hypothesis with a fixed data-independent parameter α , this approach avoids the need for manually defining a hard threshold. However, this strict filtering process also results in some neurons being excluded from the active network, as their activity lacks sufficient correlation with other neurons within the region to meet the thresholding criteria (Figure 2A).

Random networks as a statistical comparison

To compare network quantities when the threshold is increased (see Figures S4, S9, and S10), we simulated random networks where each neurons follows an homogeneous Poisson point process with their average firing rate as parameter. For each dataset, 50 simulations were computed of equal duration as the recordings. A measure of convergence r_{sc} is established to compare the emergence of sub-networks within the networks (see Figure S5) and is computed with the number of sub-network edges $card(E_{sc})$ (see STAR Methods Sub-network Decomposition) and the total number of edges $card(E)$:

$$r_{sc} = \frac{card(E_{sc})}{card(E)}. \quad (\text{Equation 7})$$

High-connectivity sub-networks obtained from graph decomposition

In order to find relevant sub-networks within the created graph, a strongly connected component (SCC) partitioning is performed using Pearce algorithm⁸² (Figure 3A). A directed graph $G=(V,E)$ is called strongly connected if for every vertices $i,j \in V^2$ there is a directed path from i to j . The strongly connected components of a directed graph are a partition into sub-graphs each of which is strongly connected (Figure 3A). From a Markovian perspective, we are dividing the graph into irreducible Markov chains. Each SCC represents a set of neurons that seems to have similar pattern of activity in space and time. After applying the algorithm, we found out that most of the SCC had a column shape except for a few outliers. Those outliers were defined by several columnar shaped clusters strongly connected together. For statistical comparison purposes, we decided to decompose these outliers into elementary columns that remains strongly connected. The approach is the following:

1. **Outlier detection.** A DBSCAN algorithm³⁰ was employed to automatically find these outliers among SCCs using the number of neurons N_c and the spread S_{SCC} on the XY plane projection (see Equation 8) as discriminative features (Figure 3B).
2. **Separation in columns.** Lengthy connections between neurons were removed based on their pairwise XY-plane distances to separate the cluster into columns. The distance threshold, denoted as $T_D = 45$, was empirically set based on the distribution of distances between neurons in the outliers (see Figure 3B; Table 1).

This two-step process allowed us to dissect the large synchronous components into elementary substructures comparable to other SCCs, facilitating a more detailed characterization of the network's functional organization. Only sub-networks with a number of neurons $N_c > T_{SCC} = 3$ (see Table 1) were kept after applying our method (Figure 3A).

$$S_{SCC} = \sqrt{\frac{\sum_{i=1}^{N_c} \|n_i - c\|_{(XY)}^2}{N_c}} \quad (\text{Equation 8})$$

Coarse-grained network using sub-networks as elementary components

An event within a strongly connected component (SCC) is defined as a continuous sequence of activations involving a minimum of $T_{event} = 2$ (see Table 1) within the sub-network. An event concludes when no further neuronal activation occurs following the last activation. With this definition, we extract a signal of activation of each SCC. Using this signal, we create a pairwise functional matrix representing a coarse-grained network between SCCs (Figure 3D). This matrix is created using the Spike Time Tiling Coefficient (STTC),³¹ a measure of correlation usually between pairs of neurons that is independent of firing rate. In order to fit in the STTC computation, we only use the start of the events to describe the SCC spike train. The time window is empirically set to four time bins (~ 1 s) because most the events in the SCC are a second long. A coarse-grained graph that represents the synchrony between SCC is thus obtained using the pairwise functional matrix.

Subtypes of cortical layer sub-networks revealed by K-means classification

After the partitioning approach was applied, strongly connected components (SCCs) were classified into several subtypes based on both geometrical and functional properties. First, we use the biological layer information to classify SCCs into three types: SCCs with neurons exclusively in Layer Va, exclusively in layer II/III, and SCCs in both layers. However, as illustrated in Figure 3C, the distribution of these classes was found to be highly imbalanced, with the majority of SCCs originating from layer II/III. To refine the classification further for layer II/III clusters, we define a vector of 27 features that are listed below in Table 3. We use a principal component analysis (PCA) to reduce the dimensionality from 27 to 6 dimensions. The number of output features was selected to retain between 70% and 80% of the explained variance ratio. The K-Means clustering algorithm was applied on the layer II/III SCCs feature vector to identify more specific subtypes. The K-Means optimal number of clusters was determined using the elbow method, mean silhouette score,⁸³ bootstrap stability⁸⁴ and on the interpretability of the clusters (see Figure S6). The Elbow method resulted in $K = 4-6$ on average (see Figure 5A; Figure S6). The silhouette score highlighted similar results $K = 5-6$ on average (see Figure S7). We added a bootstrap stability method that calculated the rand index adjusted for chance (ARI^{84}) for every K giving the highest stability for $K = 3-4$. In the end, we chose $K = 5$ because it added an explainable outlier cluster that accounts for the loss in stability, resulting in the identification of five distinct subtypes of columnar SCCs within Layer II/III.

Events activation within sub-networks

SCCs activation events (see STAR Methods Coarse-grained network) also enables the characterization of activation regimes associated with SCC subtypes by computing event statistics for each class, as presented in Figure 5B. The number of activations can be quantified for neurons serving as the first or last in an event sequence. This approach facilitates the comparison of the normalized relative depths of the starting and ending neurons of events, as illustrated in Figure 4H.

Sub-network activity correlation with snout movements

Frame-by-frame classification of snout movement was performed using videos recorded simultaneously with M1 neuronal activity. Because no visually determined threshold yielded a reliable binary classification of movement, a continuous metric of snout displacement was extracted, following the method described in previous work.¹¹ Using this continuous signal, cross-correlation analyses

were performed between population-averaged neural activity (Figure 7A) and snout movement, as well as between the activity of each sub-network (see STAR Methods Coarse-grained network using sub-networks as elementary components) and the movement trace.

To assess the significance of peaks and troughs in the cross-correlograms, a bootstrap procedure was applied. After Z score normalization of the signals, 1000 surrogate cross-correlograms were generated by introducing a random temporal circular shift in the range -50 s and 50 s for the sub-network activity and snout movement. The first and last centile of the bootstrap distribution were used as lower and upper confidence bounds. Peaks and troughs in the empirical cross-correlograms were then identified as deviations exceeding these significance bounds. Detected features were subsequently categorized according to their temporal alignment relative to snout movement within a time-window of 6s, creating five classes: pre-bump, pre-trough, post-bump, post-trough (Figure 7B) and unclassified. Table 4 was made to present the proportion of those classes and unclassified subnetworks with regards to their subtypes.

Network statistics

In this section, we define the statistical quantities we used here. The neuronal network is represented by directed weighted graph G associated with an adjacency matrix $A_{ij} \in \mathbb{R}^{N \times N}$ ^{15,85}, made of N neurons.

1. **The degree** is the number of in and out neurons given respectively by

$$d_i^{\text{in}} = \sum_{j=1}^N 1_{A_{ji} > 0} \quad \text{and} \quad (Equation 9)$$

$$d_i^{\text{out}} = \sum_{j=1}^N 1_{A_{ij} > 0}. \quad (Equation 10)$$

2. **Connection density** is the number of connections over the total number. It can be computed using the adjacency matrix as

$$\text{Density} = \frac{1}{N(N-1)} \sum_{i=1}^N \sum_{j=1}^N 1_{A_{ij} > 0}. \quad (Equation 11)$$

3. **Asymmetry** As is the difference in connection strength between two layers. For a neuron i in layer L_1 and j in layer L_2 , we have

$$As_{ij} = A_{ij} - A_{ji}. \quad (Equation 12)$$

4. **Homogeneity** represents the preferential connections to some particular nodes compared to other. It is defined as the fraction between the maximal and the minimal outward connection strengths for a given node i :

$$H_i = \frac{\inf_j \{A_{ij}\}}{\sup_j \{A_{ij}\}}. \quad (Equation 13)$$

5. **Graph traversal number of steps** $\hat{k}$ is the minimal number of iteration to transform the transition probability matrix P into a strictly positive matrix. We recall the adjacency matrix can be normalized into a transition matrix of probability P to represent the graph as a Markov Chain.³² This number describe the accessibility of each neurons in the graph defined as

$$\hat{k} = \inf_{k > 0} \{P^k \gg 0\}. \quad (Equation 14)$$

The symbol $\gg$ means that all coefficient of the matrix are strictly positive.

6. **Second largest eigenvalue modulus (SLEM)**: $\lambda = \{\lambda_1, \lambda_2, \dots, \lambda_k\}$ is the set of eigenvalues of the previously defined transition probability matrix P such that $1 = |\lambda_1| > |\lambda_2| \geq |\lambda_3| \geq \dots \geq |\lambda_k|$. The modulus of λ_2 is proportional to the convergence rate of the Markov Chain defined by P^{33} and is a proxy of the topology of the underlying graph (see Figure S7).