

Single-Neuron Network Topology Governs Neural Computation and Learning in Primate Cortex

Corresponding Author: Dr Yang Zhou

This file contains all reviewer reports in order by version, followed by all author rebuttals in order by version.

Version 0:

Reviewer comments:

Reviewer #1

(Remarks to the Author)

This manuscript shows a solid empirical study, with substantial evidence and comprehensive analyses. It proposes a segregated role of interactions between units inside and outside the hubs. Within hubs, positive and stable high interactions between neurons correlate with the observed task variables. Simultaneously, the small-world topology maps performance. This dual role is new to my knowledge.

My main, general thought is that other authors have addressed similar questions leveraging neuronal correlations or information theory, and obtained results reminiscent of some of the ones proposed in this manuscript. For instance, Balaguer-Ballester et al. (2020), PLoS Computational Biology, suggest how positive noise correlations in ensembles map to task performance, and negative correlations hinder it, influence the stability of the ensemble, etc. Valente et al. (2021) in Nature Neuroscience obtained a similar result, and other authors. There is a review in Panzeri et al. (2022), Nature Reviews in Neuroscience. Less directly related, some recent studies, like Ibanez-Bergaza et al. (2025), provide mechanistic insights. I suggest investigating the connections.

Specific and minor comments:

Introduction:

Your hypotheses are accurate in my view, because they refer to network measures, which, to my knowledge (as a non-expert in network/graph-theory expert), are new. However, when you mention encoding and connectivity, there is other literature to consider.

Results and Discussion:

Page 3, Lines 26 and onwards (P2L26+, page numbers below are approximate). This result appears in other cognitively related regions; see previous papers for more details.

P4L34+, Fig 2. I think noise correlations within hub neurons may have, most of the time, positive correlations, which would relate to previous results I mentioned (Balaguer-Ballester et al., 2020; Panzeri et al., 2021, etc.). I think looking at the absolute value of pairwise Pearson, noise correlations (discounting the effect of the other neurons). Am I right? Perhaps you can be more explicit.

P5L20+. Your results show that the neurons are correlated, confirming previous results. Connecting this with the population activity dynamics seems to require more analysis.

P5L34+. Maybe this is because the hub neurons are more correlated and therefore more deterministic. This result seems natural, but it is still an interesting observation.

P6L4. I think you tested for normality in the supplementary. Anyway, with this p-value, the evidence is overwhelming.

P6L36. Here, and in this section of the results, there is a reminiscence with Balaguer-Ballester et al. (2020) regarding the

effect of negative correlations, also maybe the work on Information-Limiting correlations by Moreno-Bote et al. and other single-neuron studies utilising different approaches.

P8L17+, and then 25+. According to previous studies (see upper comments), it is relatively expected that within-module units, correlations relate more to performance. That said, I am unaware of this observation from the network topology standpoint.

P9:29+, P10, L2+, P11, L16+. If I am correct, these paragraphs pinpoint the main novelties of the work. I agree that the interest here is that you did segregate the types of correlations: with hubs and inside modules. This result is novel with respect to previous analyses.

I also understand and agree with your conclusions.

Methods:

They seem suitable, but I would make the correlations/partial correlations more explicit (I understand they are Pearson/person-partial correlations). Since these are known methods in the area, but not more broadly, I suggest referencing them.

By “principle” component name/section, you mean principal component, the main one from PCA? Later, you mean “variance explained by the top 10 components?” Perhaps I am misunderstanding it.

Cosmetic:

I suggest spelling out the acronyms early (the PPC area and FOV) and adding a space before the parentheses. Also, re-arranging some of your equations and numbering them.

(Remarks on code availability)

I skimmed over the MATLAB scripts and your results folder, and it is sensible, and I trust the code operates as expected.

However, I cannot run it in Code Ocean. It says I need to duplicate it for a reproducible run, which may not be possible. I blame my lack of familiarity with Code Ocean, though.

My comments so far are chiefly scientific, and I can run it at a later stage if the editor decides to take the manuscript forward.

Reviewer #2

(Remarks to the Author)

Overall Assessment

The manuscript aims to address an open and pressing question regarding the relationship between neural network topology at the single-neuron level, neural functional properties, and their modulation during learning.

On the one hand, the study demonstrates methodological and analytical rigor. The use of two-photon calcium imaging with multi-day neuronal tracking in behaving monkeys represents a technically challenging approach that allows the investigation of single-neuron network function in a realistic behavioral context. Functional connectivity combined with network measures is a standard approach for analyzing multivariate collective dynamics in neural networks, but its combination with two-photon imaging and multi-session recording during learning is novel. The analyses appear statistically robust, employing multiple graph-theoretic metrics and rigorous comparisons to null models.

On the other hand, I have several major concerns that prevent me from providing an overall positive review.

Major Concern 1

First, I believe that the manuscript presents results that do not address the hypotheses mentioned in the introduction. In fact, the introduction explicitly poses two major questions:

- Does network topology and structural connectivity (e.g., small-world organization) constrain neuronal function during behavior?

- Does network topology coevolve with learning?

As stated, “Understanding these structural and functional networks at multiple scales is essential for elucidating how neural circuits contribute to overall brain activity and behavior.” This seems to be the main question that the study aims to address. However, the manuscript presents analyses based only on functional connectivity and does not analyze data concerning neural network topology and structural connectivity.

Major Concern 2

Contrary to what is stated in the manuscript, the literature on the role of “hub” neurons in the brain is rich and continuously evolving. For example, one of the first papers demonstrating a relationship between a subpopulation of GABAergic hub neurons—displaying remarkably widespread axonal arborization (i.e., structural connectivity)—and the orchestration of network synchrony in developing hippocampal networks (functional properties) is the following:

<https://www.science.org/doi/10.1126/science.1175509>

I would suggest revising the introduction to take this literature into account.

Major Concern 3

At the methodological level, the use of the causality index does not mention a large field of research in information theory and signal processing related to well-established methods for assessing directionality between time series, such as Granger causality.

I would suggest revising the methods section to take into account recent advances in this field.

Overall, I feel that the results are solid and of interest, but the framing in the introduction is not well aligned with them. I would suggest to rewrite the framing and objective of the manuscript. I would be happy to review again the revised manuscript.

(Remarks on code availability)

Reviewer #3

(Remarks to the Author)

This paper examines functional connectivity (FC) in networks of the PPC area during a behavior task. Authors found a relationship between various measures of FC and task variables over the course of learning. Overall, this paper is well structured and presents interesting results linking patterns of FC to neural encoding. Here are some comments to help improve the work and its potential impact.

Major comments:

It is important to show how consistent is functional connectivity across recording sessions. Also, how consistent is small-world connectivity across sessions? And finally, does the membership of neurons to different modules remain consistent across sessions?

The use of partial correlations can be problematic. While partial correlations reduce the number of closed triangles caused by indirect or common-input effects, they do not eliminate them entirely. How well they break those spurious $A-B-C \rightarrow A$ triangles depends on several factors that should be discussed. Otherwise, you risk over-estimating clustering coefficient if too many triangles are closed. In many respects, a generalized linear coupling model (GLM) would be superior to partial correlations.

Authors mention that the “averaged short-path length of the network decreased”. Does this analysis control for the total number of connections within the network? As you know, networks with higher density will naturally have lower path length.

Assessing the causal role of hub neurons would involve functionally or anatomically ablating them, which is not possible here as it would be technically very challenging (which is something you mention at the end of your Discussion). However, you could run analyses where you remove the highest hubs from the functional connectivity networks and assess impact on small-world measures and task encoding. This would give you some indication of the causal role of these hubs.

Authors mention that “Future research should investigate whether different types of learning induce distinct network dynamics across spatial scales and brain regions.” There is a straightforward way to address this, which I recommend doing. You can do a very simple multi-scale analysis where you lump together nearby neurons into small groups (by averaging their activity) and compute the functional connectivity between groups. In this way, each “node” of your functional network would be a group instead of a single neuron. Manipulating the size of groups would provide an indication of multi-scale structure within functional networks. I’m curious what you think you would find if you did that, in terms of small-world measures, task encoding, etc.

In the Discussion, I encourage authors to speculate about the class subtype of hub neurons. I suspect that many such hubs may be inhibitory interneurons.

Also, in the Discussion please speculate about whether you believe that learning-driven changes in functional connectivity would revert with a prolonged period without training on the task.

Minor comments:

Authors mention that “Additionally, the connection weights did not significantly correlate with the spatial distance between neurons within most FOVs”. Clearly, this means that functional connectivity does not reflect monosynaptic connections. Please add a statement to that effect.

Fig.1B-C: figure legend should specify what is the meaning of “L-L”, “L-R”, etc. (x-axis)

Authors state that “sparsity of their connection matrix was 0.048”. Can you clarify: do you mean that 4.8% of all possible connections were present? If so, it seems a bit low compared to known anatomical connectivity within local cortical networks.

Author state that “with modularity coefficients substantially exceeding those of random networks (Fig. 1H).” I am confused here, I did not think that H was a random network. Please rephrase or clarify.

Please add a final paragraph at the end of your Discussion that summarizes your main findings and the contribution of your study in few punchy sentences.

Typos:

Please add line numbers to your submission to facilitate the review.

Figure 1G legend: “ramdom networks” should be “random”

“hub neurons exhibited significantly greater connection strength and degree compared to non-hub neurons” specify: in-degree or out-degree? Or both?

Figure 1: on the figure itself, the bold letters “I” and “L” appear exactly the same. Please rewrite “L” correctly.

Bottom of p.6: “after the transitioned” replace with “transition”

p.8: “within-module pair” replace with “pairs”

p.14: “(suite2p,)” remove the extra comma

“sessions\days” replace with a forward slash “/”

« hub\non-hub” use a forward slash

(Remarks on code availability)

Reviewer #4

(Remarks to the Author)

Dear authors,

Your manuscript entitled “Single-neuron network topology governs neural computation and learning in primate cortex” presents an in-depth network analysis with refined methodology of two-photon calcium imaging of parietal cortex neurons imaged from superficial layers of 7a in awake behaving macaques performing a sensorimotor association task. Some of the results presented on hub vs. non-hub neurons are somehow expected by definition e.g. modular computation within PPC network. The most interesting aspect of your study is the observed dynamics on the hub status of neurons. In particular, the fact that status transition (hub to non-hub and vice-versa) results from the network’s influence and inter-modular connectivity during learning. This novel aspect could be given more emphasis. The paper is well written and clearly describe the methods used to obtain the results which support the conclusions. However, some aspects of the manuscript warrant correction and/or clarifications and are described below:

1. The statement in the last sentence of the abstract shall be revised. Indeed, the imaging of superficial cortical layers has been made presumably on a 5mm diameter window over PPC and therefore the reference to orchestrating global computations shall be removed or changed with local.
2. In the paragraph p11 lines 16-30 of the Discussion. The authors might also want to consider that discrepancies could also arise from the limitation of the presented results to superficial cortex.
3. On Figures 2C and S10F, I; the mention of ‘Time to target off’ is really confusing. It seems, from the Methods section that the time 0 mention time of target offset. Therefore it would be less confusing to change the X axis label to ‘Time after target off’. Especially in comparison to e.g. Figure 2 B axis.
4. The repeated references to ‘(...) our previous study (...)’ p3, lines 11-14; p4 lines 41-42; p11 line 41; p12 lines 17 and line 32; p13 line 15. The terms previous study can be misleading for the reader, indirectly implying published work. The reference #52 is a bioRxiv paper submitted in August 2025, that has not been peer-reviewed. Hence, I strongly suggest you refer to it in the present tense rather than the past tense e.g. ‘our work also shows’ or ‘we are also showing’.
5. Related to the above point, the reference to ‘our previous study’ in the Material and Methods section refers to a wrong publication #1 on p4 lines 41-42; p11 line 41; p12 lines 17 and line 32; p13 line 15. Also on p13 line 3, it seems you actually refers to other publications than reference #1 and #2. Please correct.
6. This study resonates with previous work from Dann et al., eLife 2016 <https://doi.org/10.7554/eLife.15719> and would gain including other references and comparing their results in the discussion to similar network topology analyses on parietal cortex neurons

Minor points:

Please include the macaque species used in your study.

p13 line 7, you refer the ‘monkey brain atlas’. Would be more accurate to specify that 7a was located based on the gyral and sulcal landmarks in comparison to an atlas with a reference to it.

Could you please describe the quality criteria used in Neuronal data analysis? p14 line 24

(1) In the Discussion, please add a sentence or two to speculate about what synaptic plasticity mechanisms may be responsible for (a) small-world functional connectivity networks in PPC; and (b) gain or loss of hub status. Around line 561, you state "Future research should investigate whether different types of learning..." but unless you specify what potential mechanisms (Hebbian plasticity? STDP? Homeostatic plasticity?) it is difficult to link your results to synaptic function. Even if this is speculative, at least provide some ideas of how your results on small-world and hub status may arise in a biologically plausible fashion.

(2) A small typo: line 527: *A* previous study

(Remarks on code availability)

Reviewer #4

(Remarks to the Author)

Dear authors,

you have satisfactorily addressed all my concerns and I have no further comments.

Sincerely

(Remarks on code availability)

I reviewed briefly the code provided but could not run it in Code Ocean.

Due to default in the activation of my Code Ocean account, I could not verify if the code runs well on the Day1 preprocessed dataset provided.

However, codes seems legit and basic README files are provided with enough instructions and details to be able to run the code (including toolboxes and environment required, brain connectivity toolbox and other functions or m files used for the analysis) and reproduces the figures listed. The authors shall engage to ensure response to any request in case of missing data/information.

Open Access This Peer Review File is licensed under a Creative Commons Attribution 4.0 International License, which permits use, sharing, adaptation, distribution and reproduction in any medium or format, as long as you give appropriate credit to the original author(s) and the source, provide a link to the Creative Commons license, and indicate if changes were made.

In cases where reviewers are anonymous, credit should be given to 'Anonymous Referee' and the source.

The images or other third party material in this Peer Review File are included in the article's Creative Commons license, unless indicated otherwise in a credit line to the material. If material is not included in the article's Creative Commons license and your intended use is not permitted by statutory regulation or exceeds the permitted use, you will need to obtain permission directly from the copyright holder.

To view a copy of this license, visit <https://creativecommons.org/licenses/by/4.0/>

Authors' Response to Reviewers

We sincerely thank all the reviewers for their thoughtful and detailed feedback on the earlier version of our manuscript. Their comments and questions have been invaluable, and we have addressed each one in the responses below while incorporating the corresponding revisions into the manuscript. As a result, the manuscript has been significantly improved. The revised version includes updates to all sections of the text, revised Figures 2-3, and 5 revised supplementary figures. Notably, we have substantially revised the Introduction and Discussion sections to better narrow and refine our research questions, ensuring they align more closely with the scope of functional neuronal networks. Additionally, we have offered a more thorough contextualization of our findings within the existing literature on noise correlations to enhance the interpretation of our results and highlight the novelty of our contributions. In the response below, reviewers' comments are shown in black, with our point-by-point replies in blue.

REVIEWER COMMENTS

Reviewer #1 (Remarks to the Author):

This manuscript shows a solid empirical study, with substantial evidence and comprehensive analyses. It proposes a segregated role of interactions between units inside and outside the hubs. Within hubs, positive and stable high interactions between neurons correlate with the observed task variables. Simultaneously, the small-world topology maps performance. This dual role is new to my knowledge.

My main, general thought is that other authors have addressed similar questions leveraging neuronal correlations or information theory, and obtained results reminiscent of some of the ones proposed in this manuscript. For instance, Balaguer-Ballester et al. (2020), PLoS Computational Biology, suggest how positive noise correlations in ensembles map to task performance, and negative correlations hinder it, influence the stability of the ensemble, etc. Valente et al. (2021) in Nature Neuroscience obtained a similar result, and other authors. There is a review in Panzeri et al. (2022), Nature Reviews in Neuroscience. Less directly related, some recent studies, like Ibanez-Bergaza et al. (2025), provide mechanistic insights. I suggest investigating the connections.

Response: We thank the reviewer for the thoughtful comments and the valuable suggestion. We agree that our original manuscript did not sufficiently contextualize our findings relative to the existing body of literature concerning neuronal activity correlation and its influence on behavior.

The established research on noise correlation, which analyzes the shared

trial-to-trial variability in neural activity between neuron pairs, has provided profound insight into the functional connectivity within neural circuits and their linkage to neural information encoding/readout and animal’s behavior performance. We agree that explicitly connecting our current results to these previous noise correlation studies will substantially enhance the mechanistic interpretation of our results and clarify the novelty of our contributions

Our approach differs conceptually in that we leverage neuronal activity correlations—specifically, partial noise correlations—not merely as a statistical measure, but as the basis for defining and analyzing the mesoscale functional neuronal network based on graph theoretical approaches,. We systematically examine how the network’s topological features relate to neuronal encoding, inter-neuron communication, the evolution of encoding across days, and the animal’s behavioral performance during learning. Whereas prior studies have largely treated correlations as covariance statistics linking population activity to encoding, readout, or behavioral outcomes, we instead use correlation structure to interrogate an explicit network architecture that constrains neuronal function, communication, plasticity, and behavior. This conceptual shift reveals several findings not reported in previous work on noise correlations. For example, Balaguer-Ballester et al. (2020) showed that positive correlations within neuronal ensembles facilitate task performance, while negative correlations impede it by destabilizing ensemble activity, and Valente et al. (2021) demonstrated that noise correlations contribute to information readout in association cortex. In contrast, our study demonstrates that functional network topology not only constrains the role of individual neurons—where hub neurons exert disproportionate influence on task-related encoding and inter-neuron communication—but also shapes long-term plasticity, reflected in the evolution of hub versus non-hub status and neuronal encoding across learning days, and co-evolves with behavioral performance throughout learning.

A detailed discussion of the relationship between network topology and these classical noise correlation studies is indeed essential for fully highlighting the conceptual novelty of our mesoscale analysis. Therefore, we have revised both the Introduction and Discussion sections to provide a robust contextualization of our findings relative to the previous literature on noise correlation, shown as follows.

Introduction section:

‘Although reconstructing the structural neuronal network in the behaving brain remains challenging, several approaches have been developed to infer functional connectivity. Among these, noise correlation—capturing shared trial-to-trial variability in neural activity—has become a widely used metric for estimating functional coupling²⁵⁻²⁷. This measure has generated important insights into circuit-level interactions and their relationship to neural encoding,

information readout, and ultimately behavioral performance²⁶⁻³¹. However, prior studies have reported conflicting findings regarding whether higher noise correlations facilitate or impair behavioral performance^{27, 28, 31-35}. One potential reason is that earlier studies typically treated noise correlation as a covariance statistic between pairs or small subsets of neurons, overlooking the broader network topology that arises at the population level. This has left critical gaps in our understanding of the mesoscale functional architecture that supports neural computation and learning. More recently, researchers have proposed that the fine-grained structure of correlations—such as those within versus across neuronal ensembles—may account for these divergent results^{26, 27, 32, 33, 35}, highlighting the need to study network topology as an integrated whole. ’

Discussion section:

‘Functional connectivity within neuronal ensembles has been widely investigated using diverse statistical approaches^{3, 10, 30, 33, 42, 47, 64, 66, 67}. Prior work commonly conceptualizes activity correlations as covariance-based measures that impact neural computation and behavioral performance^{25, 26, 28}. Particularly, distinct patterns of noise correlations have been shown to shape population-level information encoding and readout^{26, 27, 31, 34}; positive correlations within neuron ensembles can facilitate task performance, whereas negative correlations may impair it by destabilizing ensemble dynamics³². Building on these foundational insights, our study introduces a conceptual shift by leveraging sophisticated correlation structure to define and interrogate an explicit mesoscale neural network architecture based on graph theoretical approaches. We systematically investigated how the network’s topological features such as hub status and modular organization, as well as their dynamics, relate to neuronal encoding, inter-neuron communication, long-term plasticity, and behavioral performance throughout learning. Recent studies have also begun to interpret functional connectivity through the lens of network topology. For example, an electrophysiological study in behaving monkeys demonstrated that functional networks in the frontoparietal cortex exhibit small-world organization and highlighted the pivotal role of rich-club hub neurons in driving oscillatory dynamics³⁰. Together, these studies and our findings highlight the necessity of understanding sophisticated network architecture to fully characterize neural computation and its behavioral consequences. ’

In the following, we have addressed all the comments described below, which were very helpful in improving the manuscript.

Specific and minor comments:

Introduction:

Your hypotheses are accurate in my view, because they refer to network measures, which, to my knowledge (as a non-expert in network/graph-theory expert), are new. However, when you mention encoding and connectivity, there is other literature to consider.

Response: We appreciate the reviewer's comment. We realize that the initial version of the Introduction did not provide adequate referencing or connection to previous foundational work in neural encoding and connectivity.

In response, we have added a dedicated paragraph to the Introduction section of the revised manuscript. This new section acknowledges the contributions of prior research on neural encoding and connectivity, particularly highlighting the established roles of noise correlation in shaping population neural encoding and behavior. This ensures that the context for our mesoscale network findings is properly established within the existing literature, emphasizing the transition from established single-unit and correlation studies to our novel approach focused on network topology. The added text was shown in bellow:

‘Although reconstructing the structural neuronal network in the behaving brain remains challenging, several approaches have been developed to infer functional connectivity. Among these, noise correlation—capturing shared trial-to-trial variability in neural activity—has become a widely used metric for estimating functional coupling²⁵⁻²⁷. This measure has generated important insights into circuit-level interactions and their relationship to neural encoding, information readout, and ultimately behavioral performance²⁶⁻³¹. However, prior studies have reported conflicting findings regarding whether higher noise correlations facilitate or impair behavioral performance^{27, 28, 31-35}. One potential reason is that earlier studies typically treated noise correlation as a covariance statistic between pairs or small subsets of neurons, overlooking the broader network topology that arises at the population level. This has left critical gaps in our understanding of the mesoscale functional architecture that supports neural computation and learning. More recently, researchers have proposed that the fine-grained structure of correlations—such as those within versus across neuronal ensembles—may account for these divergent results^{26, 27, 32, 33, 35}, highlighting the need to study network topology as an integrated whole.’

Results and Discussion:

Page 3, Lines 26 and onwards (P2L26+, page numbers below are approximate). This result appears in other cognitively related regions; see previous papers for more details.

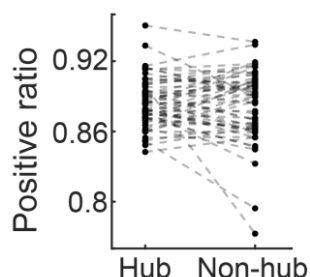
We thank the reviewer for highlighting this important point. It has indeed been

reported that neurons exhibiting similar functional properties, such as encoding preferences, tend to display higher correlated activity (noise correlation). Our finding directly reinforces this fundamental principle, even though our quantification utilizes a different metric (partial noise correlation) to delineate the functional network architecture. In the revised manuscript, we have added a sentence to the revised manuscript to explicitly contextualize our finding within this existing literature:

‘Notably, connection weights within the PPC neuron networks were stronger between neurons that exhibited similar encoding selectivity during monkeys’ task performance, compared to neurons with opposite encoding preferences (Figure. 1, b, c), which is consistent with the previous finding that neurons exhibited similar neural encoding were more correlated^{28, 63}.’

P4L34+, Fig 2. I think noise correlations within hub neurons may have, most of the time, positive correlations, which would relate to previous results I mentioned (Balaguer-Ballester et al., 2020; Panzeri et al., 2021, etc.). I think looking at the absolute value of pairwise Pearson, noise correlations (discounting the effect of the other neurons). Am I right? Perhaps you can be more explicit.

Response: We thank the reviewer for the insightful comments. The reviewer is correct that hub neurons showed positive connection weights (partial noise correlations) with most of their connected neurons. However, this pattern was not unique to hub neurons: non-hub neurons also exhibited predominantly positive connections, as negative correlations were relatively rare across the entire population. In our study, hub and non-hub neurons were defined based on the number and strength of their positive connections. Thus, hub neurons did not necessarily form exclusively positive connections, nor did non-hub neurons necessarily form negative ones. In the revised manuscript, we have now quantified the ratio of positive and connections separately for both hub and non-hub neurons. This analysis revealed that the proportion of positive connections (connection strength) did not differ significantly between hub and non-hub neurons across all constructed networks ($P = 0.476$). We have added this result as Supplementary Fig. 3c, shown as follows:



Regarding the sentence mentioned by the reviewer, we quantified the

relationship between neuronal encoding strength and the centrality of each neuron (our measure of hubness). This analysis showed that hub neurons were more likely to exhibit stronger encoding of task variables—an observation not reported in the studies cited by the reviewer.

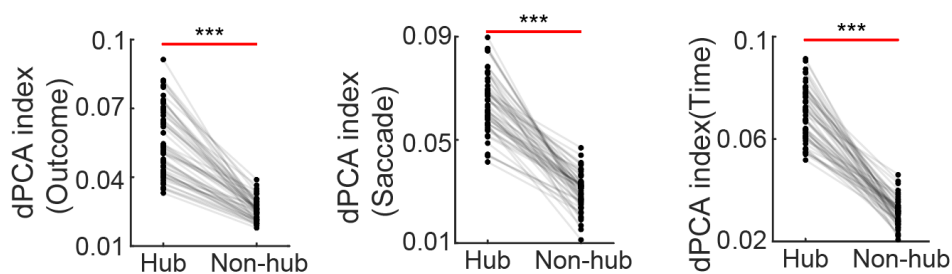
We also clarify that centrality was computed using the absolute values of pairwise Pearson noise correlations. In the revised manuscript, we have updated the sentence for clarity as follows:

“We found that the encoding strength for both task variables was positively correlated with each neuron’s centrality (Supplementary Fig. S8), a fundamental measure of hubness, defined as the sum of the absolute values of its connections.”

P5L20+. Your results show that the neurons are correlated, confirming previous results. Connecting this with the population activity dynamics seems to require more analysis.

Response: We thank the reviewer for the thoughtful comments. Indeed, our finding—that hub neurons explain the largest proportion of variance in the population activity—suggests that hub neurons are more strongly correlated with the overall neural dynamics during task performance. We also clarify that the network construction was based on partial noise correlations measured during the fixation (baseline) period, rather than the condition-averaged activity used to quantify population dynamics during task performance. Our results indicate that neurons with stronger baseline noise correlations tend to exhibit more similar activity patterns throughout a task trial, consistent with—and extending—previous results.

We agree that further analyses are needed to more directly quantify the relationship between hub/non-hub status and population dynamics. The original PCA did not dissociate task-variable-related components from time-dependent dynamics. To address this, we have now performed a de-mixed PAC analysis in the revised Results section, separating saccade-direction encoding, outcome encoding, and time-dependent activity. This new analysis shows that hub neurons contribute substantially more to both task-variable encoding and temporal dynamics. These results have been added to revised Fig. 3, as illustrated below.



We have added a concise paragraph summarizing these results in the revised Results section, as shown below.

“To further decompose task-related components of activity, we conducted a de-mixed PCA analysis, which isolated saccade-direction encoding, outcome encoding, and time-dependent activity for each recording session. Consistently, hub neurons demonstrated markedly stronger contributions to both task-variable encodings and temporal dynamics than non-hub neurons (Figure. 3b-d). Together, these findings indicate that hub neurons exert a dominant influence over population-level activity dynamics during the monkeys’ task performance. ”

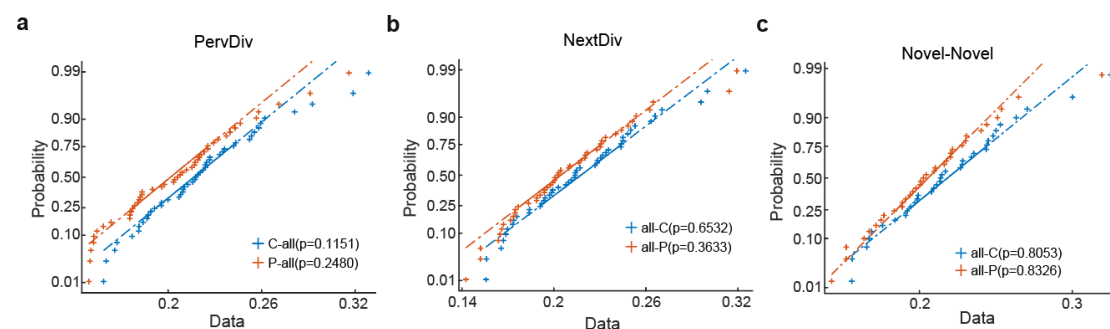
P5L34+. Maybe this is because the hub neurons are more correlated and therefore more deterministic. This result seems natural, but it is still an interesting observation.

Response: We thank the reviewer for this insightful comment. We think this might be a potential mechanical explanation for this result. We have included a sentence to discuss this point in the revised Results section:

“This may be because hub neurons exhibit stronger and broader correlations within the network, making them more influential in driving population activity dynamic”

P6L4. I think you tested for normality in the supplementary. Anyway, with this p-value, the evidence is overwhelming.

Response: We have verified that the data distributions shown in this figure do not significantly deviate from normality. Normality tests performed on the data underlying the referenced panels (original Fig. 2n - p), which is shown in below, indicate that these distributions are not significantly different from a normal distribution.



The figures above show the results of the normality test for the data presented in the original Fig. 2N-P. In each panel, the quantiles of the

theoretical normal distribution (y-axis) are plotted against the actual data values (x-axis). The fitted normal distributions are shown, with a solid reference line connecting the first and third quartiles of the data, and a dashed reference line extending the solid line to the ends of the data. Each condition is represented by a different color. The data points closely align with the reference lines, suggesting that the data do not significantly deviate from a normal distribution for each condition. P-values are derived from Lilliefors tests.

In addition, we performed non-parametric analyses (Friedman tests) for the mentioned test, which yielded qualitatively consistent results. ($P = 0.0001$)

P6L36. Here, and in this section of the results, there is a reminiscence with Balaguer-Ballester et al. (2020) regarding the effect of negative correlations, also maybe the work on Information-Limiting correlations by Moreno-Bote et al. and other single-neuron studies utilising different approaches.

Response: We appreciate the reviewer's comment and recognize that our initial description of these findings may have lacked sufficient clarity, potentially leading to confusion regarding the novelty of our approach relative to prior noise correlation studies. In this specific section, we analyzed the mesoscale network dynamics by studying the **transition of hub status** across two consecutive learning days. This process allowed us to identify two key populations: "*hub-losing*" neurons (transitioning from a network hub on Day 1 to a non-hub on Day 2) and "*hub-gaining*" neurons (the reverse). We established two major findings concerning these transitions:

1. **Causal Influence:** We quantified the mutual influence between these transitional neurons and the rest of the network on the evolution of neuronal encoding across days. We found that hub-losing neurons exerted a **greater causal influence** on the encoding changes of other neurons than they received. Conversely, hub-gaining neurons were significantly **more influenced** by the network's encoding evolution than they influenced it.
2. **Topological Predictors:** We identified a structural basis for these transitions: hub-losing neurons exhibited the strongest dominance of positive over negative inter-module connections among all neuron subsets, while hub-gaining neurons also showed a greater dominance compared to stable non-hub neurons.

Novelty Relative to Previous Correlation Studies:

Our results represent a significant conceptual shift from the studies the reviewer cited, which primarily focused on relating simple noise correlation statistics to information encoding and behavior. The novelty of our approach lies in three key aspects:

1. **Dynamic Network Topology and Neuronal Encoding:** This study is the first

to examine how dynamic shifts in a neuron’s topological position—reflected by changes in its sophisticated correlation structure—relate to the long-term evolution of its encoding properties across learning.

2. **Sophisticated Network-Based Metrics:** We characterize population-level correlation structure from an explicit network perspective (employing graph theoretical approaches), employing graph-based metrics such as the diversity coefficient to capture inter-modular connectivity patterns (not global connectivity or within-modular connectivity). This approach captures topological properties—for instance, the predominance of positive over negative inter-module connections—that fundamentally differ from the pairwise or higher-order correlation measures commonly used in earlier work.
3. **A Distinct Phenomenon Beyond Pairwise Correlations:** While previous research has shown that distinct noise correlation patterns can influence population encoding and behavior in different manners, our study demonstrates that sophisticated network connectivity constrains the dynamics of network topology, inter-neuronal communication, and neuronal encoding within the network itself.

To better clarify the novelty of our findings, we have added a paragraph to the revised discussion section that directly compares our mesoscale network perspective to these previous correlation studies. As demonstrated in the response above and below:

“Functional connectivity within neuronal ensembles has been widely investigated using diverse statistical approaches^{3, 10, 30, 33, 42, 47, 64, 66, 67}. Prior work commonly conceptualizes activity correlations as covariance-based measures that impact neural computation and behavioral performance^{25, 26, 28}. Particularly, distinct patterns of noise correlations have been shown to shape population-level information encoding and readout^{26, 27, 31, 34}; positive correlations within neuron ensembles can facilitate task performance, whereas negative correlations may impair it by destabilizing ensemble dynamics³². Building on these foundational insights, our study introduces a conceptual shift by leveraging sophisticated correlation structure to define and interrogate an explicit mesoscale neural network architecture based on graph theoretical approaches. We systematically investigated how the network’s topological features such as hub status and modular organization, as well as their dynamics, relate to neuronal encoding, inter-neuron communication, long-term plasticity, and behavioral performance throughout learning. Recent studies have also begun to interpret functional connectivity through the lens of network topology. For example, an electrophysiological study in behaving monkeys demonstrated that functional networks in the frontoparietal cortex exhibit small-world organization and highlighted the pivotal role of rich-club hub neurons in driving oscillatory dynamics³⁰. Together, these studies and our findings highlight the necessity of

understanding sophisticated network architecture to fully characterize neural computation and its behavioral consequences. ’ ”

P8L17+, and then 25+. According to previous studies (see upper comments), it is relatively expected that within-module units, correlations relate more to performance. That said, I am unaware of this observation from the network topology standpoint.

Response: We thank the reviewer for their careful attention to the clarity of our results regarding network topology and modularity. We are not entirely certain whether the reviewer is suggesting that within-module correlations are more expected to relate to behavioral performance than inter-module correlations. In our study, each neuron is assigned to a specific module; thus, any pair of neurons can be classified as either a within-module pair or an inter-module pair. In the section in question, we reported that within-module neuron pairs exerted stronger mutual influences on their activity dynamics and showed greater mutual impacts on the evolution of neuronal encoding during task performance and learning, compared with inter-module pairs. This result is conceptually distinct from the claim that within-module correlations are more strongly related to behavioral performance than inter-module correlations.

That said, we acknowledge that the predictive influence of modularity—a core property of small-world networks—on instantaneous functional similarity is, in a general sense, an expected finding. Importantly, however, the significance and novelty of our work go well beyond this basic observation. Our contribution lies in the specific network-based analytical framework we employed and, critically, in revealing how modular network structure constrains the long-term, dynamic evolution of neuronal representations across learning.

In this section, we rigorously explored how network modularity, which we identified through complex noise correlation patterns, influences fundamental functional aspects: neuronal encoding, inter-neuron communication, and, critically, the **evolution of neuronal encoding across multiple learning days**.

Our novelty is threefold:

1. **Network-Centric Quantification:** We quantify modularity using a **mesoscale network approach** based on partial noise correlation, which defines modules as explicit structural components of the functional architecture. This is distinct from previous correlation studies that typically relied on simple pairwise correlation or functional similarity.
2. **Dynamic Consequence of Modularity:** While we confirm that neurons within the same module are functionally more similar and exert stronger mutual influence on activity dynamics, the key novel finding is that this topological structure **constrains the day-to-day evolution of neuronal**

encoding during long-term learning. This mechanistic insight has not been reported in the cited or other previous studies.

3. **Functional Dissociation:** We constructed the network using partial noise correlation measured *only during the fixation period* (a non-choice epoch) but then demonstrated that this derived structure governed neuronal encoding and activity dynamics *during active task performance and learning*. This functional dissociation implies that modular architecture defines an intrinsic scaffolding whose rules govern cognitive processing across disparate task phases.

As described in the response above, to better clarify the conceptual distinction and novelty of our mesoscale network perspective compared to previous correlation studies, we have incorporated a dedicated paragraph into the revised Discussion section.

P9:29+, P10, L2+, P11, L16+. If I am correct, these paragraphs pinpoint the main novelties of the work. I agree that the interest here is that you did segregate the types of correlations: with hubs and inside modules. This result is novel with respect to previous analyses.

I also understand and agree with your conclusions.

Response: We thank the reviewer for the positive comments.

Methods:

They seem suitable, but I would make the correlations/partial correlations more explicit (I understand they are Pearson/person-partial correlations). Since these are known methods in the area, but not more broadly, I suggest referencing them.

Response: We thank the reviewer for the comment. In the revised Methods section, we have explicitly specified the use of Pearson correlation and added several relevant references.

By “principle” component name/section, you mean principal component, the main one from PCA? Later, you mean “variance explained by the top 10 components?” Perhaps I am misunderstanding it.

Response: Thanks for pointing out this typo. It should be ‘principal component’ we have corrected it in the revised manuscript.

Cosmetic:

I suggest spelling out the acronyms early (the PPC area and FOV) and adding a space before the parentheses. Also, re-arranging some of your equations and

numbering them.

Response: Thank you for the suggestion. We have revised the acronyms in the early paragraphs and added spaces before all parentheses. In addition, we have numbered all equations in the revised Methods section.

Reviewer #1 (Remarks on code availability):

I skimmed over the MATLAB scripts and your results folder, and it is sensible, and I trust the code operates as expected.

However, I cannot run it in Code Ocean. It says I need to duplicate it for a reproducible run, which may not be possible. I blame my lack of familiarity with Code Ocean, though.

My comments so far are chiefly scientific, and I can run it at a later stage if the editor decides to take the manuscript forward.

Response: We apologize for the inconvenience. We have reviewed our capsule on Code Ocean, and it appears to be the same version of the code that worked well on our own system. We suspect that the issue might be related to the Peer Review Copy to some extent. In the meantime, we have published the capsule, and it can be accessed via this DOI: <https://doi.org/10.24433/CO.8493005.v1>. Please feel free to try the code from this link at your convenience. We hope this resolves the issue.

Reviewer #2 (Remarks to the Author):

Overall Assessment

The manuscript aims to address an open and pressing question regarding the relationship between neural network topology at the single-neuron level, neural functional properties, and their modulation during learning.

On the one hand, the study demonstrates methodological and analytical rigor. The use of two-photon calcium imaging with multi-day neuronal tracking in behaving monkeys represents a technically challenging approach that allows the investigation of single-neuron network function in a realistic behavioral context. Functional connectivity combined with network measures is a standard approach for analyzing multivariate collective dynamics in neural networks, but its combination with two-photon imaging and multi-session recording during learning

is novel. The analyses appear statistically robust, employing multiple graph-theoretic metrics and rigorous comparisons to null models.

On the other hand, I have several major concerns that prevent me from providing an overall positive review.

Response: We thank the reviewer for the comments! Please see our point-to-point response below.

Major Concern 1

First, I believe that the manuscript presents results that do not address the hypotheses mentioned in the introduction. In fact, the introduction explicitly poses two major questions:

- Does network topology and structural connectivity (e.g., small-world organization) constrain neuronal function during behavior?
- Does network topology coevolve with learning?

As stated, “Understanding these structural and functional networks at multiple scales is essential for elucidating how neural circuits contribute to overall brain activity and behavior.” This seems to be the main question that the study aims to address. However, the manuscript presents analyses based only on functional connectivity and does not analyze data concerning neural network topology and structural connectivity.

Response: We thank the reviewer for this valuable comment. We agree that the original introduction did not clearly articulate the central research questions, and that some of the questions presented were seemly over broad or not directly addressed in the current manuscript. As the reviewer noted, our study specifically centers on functional connectivity. In this work, we systematically analyzed the topology of the functional neuronal network and investigated its role in shaping neuronal encoding, inter-neuron communication, the evolution of encoding, and behavior performance across days of associative learning. In particular, we examined how two small-world topology features of the primate PPC functional network (hub/non-hub status and modularity organization) constrain neuronal function (i.e., task-variable encoding and contribution to network dynamics) and the learning process (i.e., the evolution of encoding across learning days). Notably, we assessed how different network-topology features relate to behavioral performance during learning.

In response, we have revised the Introduction section to narrow and refine our research questions, making them more directly aligned with the scope of functional neuronal networks. The revised Introduction part is shown below:

‘Although reconstructing the structural neuronal network in the behaving brain remains challenging, several approaches have been developed to infer

functional connectivity. Among these, noise correlation—capturing shared trial-to-trial variability in neural activity—has become a widely used metric for estimating functional coupling²⁵⁻²⁷. This measure has generated important insights into circuit-level interactions and their relationship to neural encoding, information readout, and ultimately behavioral performance²⁶⁻³¹. However, prior studies have reported conflicting findings regarding whether higher noise correlations facilitate or impair behavioral performance^{27, 28, 31-35}. One potential reason is that earlier studies typically treated noise correlation as a covariance statistic between pairs or small subsets of neurons, overlooking the broader network topology that arises at the population level. This has left critical gaps in our understanding of the mesoscale functional architecture that supports neural computation and learning. More recently, researchers have proposed that the fine-grained structure of correlations—such as those within versus across neuronal ensembles—may account for these divergent results^{26, 27, 32, 33, 35}, highlighting the need to study network topology as an integrated whole.

Network topology—defined by connectivity patterns and node positioning—systematically shapes the functional roles of its elements^{9, 36}. For instance, structural properties like degree distribution and modularity directly influence node dynamics³⁶, centrality measures (e.g., betweenness) predict functional importance in information flow or resource allocation³⁷, and highly connected ‘hub nodes’ are often functionally critical³⁸. For instance, in vitro studies have shown that hippocampal GABAergic interneurons, characterized by their extensive axonal arborization, serve as functional hubs that drive large-scale network dynamics and synchrony during the development of hippocampal circuits³⁹. However, how the functioning of individual neurons depends on their topology position in the neuron network remains largely unexplored. Moreover, learning is a dynamic process that fundamentally alters the structure and function of neuron networks⁴⁰⁻⁴⁵. During learning, specific patterns of neuronal activity would lead to the remodeling of neural connections, enhancing the efficiency of information processing within these networks^{43, 46}. Conversely, whether and how the functional network topology shapes the evolution of neuronal activity during learning remains unclear.’

Major Concern 2

Contrary to what is stated in the manuscript, the literature on the role of “hub” neurons in the brain is rich and continuously evolving. For example, one of the first papers demonstrating a relationship between a subpopulation of GABAergic hub neurons—displaying remarkably widespread axonal arborization (i.e., structural connectivity)—and the orchestration of network synchrony in developing hippocampal networks (functional properties) is the following:

<https://www.science.org/doi/10.1126/science.1175509>

I would suggest revising the introduction to take this literature into account.

Response: We thank the reviewer for pointing out this important reference. It provides valuable experimental evidence regarding the structural and cell-type basis of hub neurons, as well as their functional roles. We have now incorporated this reference into the Introduction section as follows:

“For instance, in vitro studies have shown that hippocampal GABAergic interneurons, characterized by their extensive axonal arborization, serve as functional hubs that drive large-scale network dynamics and synchrony during the development of hippocampal circuits³⁹ However, how the functioning of individual neurons depends on their topology position in the neuron network remains largely unexplored. Moreover, learning is a dynamic process that fundamentally alters the structure and function of neuron networks^{40–45}. During learning, specific patterns of neuronal activity would lead to the remodeling of neural connections, enhancing the efficiency of information processing within these networks^{43, 46}. Conversely, whether and how the functional network topology shapes the evolution of neuronal activity during learning remains unclear.”

Major Concern 3

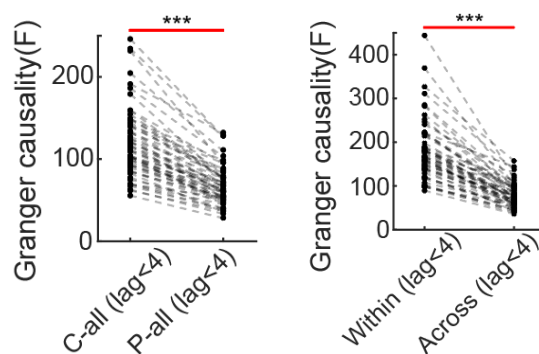
At the methodological level, the use of the causality index does not mention a large field of research in information theory and signal processing related to well-established methods for assessing directionality between time series, such as Granger causality.

I would suggest revising the methods section to take into account recent advances in this field.

Response: We thank the reviewer for raising this important point. As noted, we used a causality index to quantify the direction and strength of influence between two time series. We agree that employing more widely established methods can strengthen the analysis. As described in the Methods section, the causality index used here is a variant of Granger causality. The key difference is that, in our study, it was applied to condition-averaged activity rather than single-trial activity. We chose this approach instead of classical Granger causality for a specific reason: classical Granger causality requires simultaneously recorded time series and is therefore well suited for assessing mutual influences on instantaneous activity dynamics between neurons, but it cannot be readily applied to quantify influences on neuronal encoding (between different neurons) across learning days.

To validate our approach, we additionally applied the classical Granger causality test to single-trial activity data. This analysis yielded results that were highly consistent with those obtained using the causality index. Below shows the results of the classical Granger causality analysis on single-trial data

corresponding to the original Fig. 2L and Fig. 4H.



Left panel: Granger causality analysis was used to validate the results shown in the original Fig. 2L, which quantified the influence of each neuron's activity at a given time point on the activity of other neurons at later time points within the same task trial. Right panel: Granger causality analysis was used to validate the results shown in the original Fig. 4H, which quantified the mutual influences on activity dynamics between neurons under within-module and across-module conditions. In both cases, the results were highly consistent with those reported in the original manuscript.

In the revised Methods section, we have added a paragraph clarifying the differences between the causality index and Granger causality:

“The causality index used in this study is conceptually derived from Granger causality. However, unlike the classical Granger approach—which operates on single-trial time series—we applied the causality index to condition-averaged activity. This render to us to use the same metric to quantify both the mutual influence on the instantaneous activity dynamics between neurons and the influence on the neuronal encodings between different neurons across different learning days. To validate this approach, we also applied the Granger causality test to quantify the mutual influence on the instantons activity dynamics among different neurons, which yielded qualitatively similar results.”

Overall, I feel that the results are solid and of interest, but the framing in the introduction is not well aligned with them. I would suggest to rewrite the framing and objective of the manuscript. I would be happy to review again the revised manuscript.

Response: We thank the reviewer for these helpful comments, which have greatly contributed to improving the manuscript. As described above, we have extensively revised the Introduction to clarify and refine our research questions, ensuring that they are more directly aligned with the results presented in the manuscript.

Reviewer #3 (Remarks to the Author):

This paper examines functional connectivity (FC) in networks of the PPC area during a behavior task. Authors found a relationship between various measures of FC and task variables over the course of learning. Overall, this paper is well structured and presents interesting results linking patterns of FC to neural encoding. Here are some comments to help improve the work and its potential impact.

Response: Thanks for your thoughtful comments! Please see our point-to-point responses below.

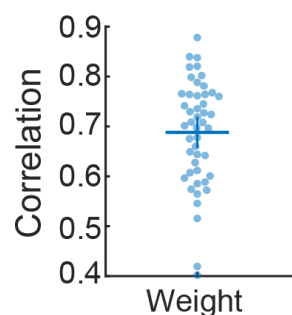
Major comments:

It is important to show how consistent is functional connectivity across recording sessions. Also, how consistent is small-world connectivity across sessions? And finally, does the membership of neurons to different modules remain consistent across sessions?

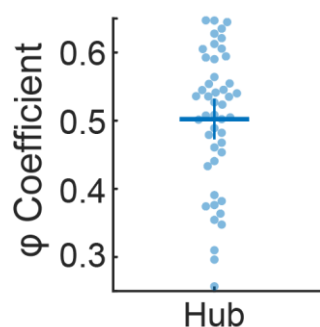
Response: We thank the reviewer for this important suggestion. We realized that the original manuscript did not fully address the consistency of network topology (functional connectivity) across learning sessions. In summary, we observed that: 1) The constructed functional neuronal networks consistently exhibited small-world topology across learning sessions. 2) The hub/non-hub status of some neurons transitioned across learning days. 3) The small-world topology features evolved alongside the monkeys' behavioral performance across learning days.

In the revised manuscript, we have included three new analyses to further examine the consistency of functional connectivity across learning days:

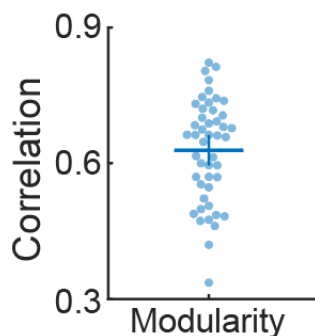
Network similarity across learning sessions: For each FOV, we assessed the stability of the functional network across days by computing the correlation between connectivity matrices derived from different learning sessions. This analysis revealed that connectivity patterns were highly consistent across consecutive learning days ($r = 0.688 \pm 0.030$). This result is presented in Supplementary Fig. 2a and also shown below.



Stability of hub/non-hub status: We quantified the similarity of hub/non-hub assignments for all neurons within each FOV across learning sessions. We found that the hub/non-hub status of the neurons within each network are partially consistent across consecutive learning days ($\phi = 0.502 \pm 0.029$), as some neurons exhibited transitions in the hub/non-hub status. This result is presented in Supplementary Fig. 11a and also shown below.



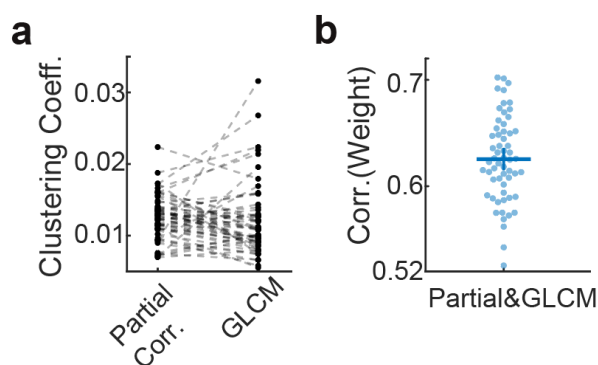
Consistency of modularity: We assessed the stability of network modularity by calculating the correlation of the probability that neuron pairs belong to the same module across learning sessions. This analysis showed that the module organization of neurons within each network was consistent across consecutive learning days ($r = 0.629 \pm 0.032$). This result is presented in Supplementary Fig. 7a and also shown below.



The use of partial correlations can be problematic. While partial correlations reduce the number of closed triangles caused by indirect or common-input effects, they do not eliminate them entirely. How well they break those spurious $A-B-C \rightarrow A$ triangles depends on several factors that should be discussed. Otherwise, you risk over-estimating clustering coefficient if too many triangles are closed. In many respects, a generalized linear coupling model (GLM) would be superior to partial correlations.

Response: We thank the reviewer for this important technical comment. In the original manuscript, we did not sufficiently consider the potential contamination

of network analyses—particularly clustering coefficient estimates—by local triangular correlation structures. Following the reviewer’s suggestion, we applied a generalized linear coupling model (GLCM) to estimate connection weights between neurons across all recording sessions. Across sessions, the resulting connection weights were highly correlated with those obtained from partial correlation analysis ($r = 0.626 \pm 0.010$). In addition, we computed the clustering coefficients of functional neuronal networks constructed using the GLCM. These clustering coefficients were comparable to those derived from partial correlation-based networks ($p = 0.444$), suggesting that our original approach did not substantially overestimate clustering in the present dataset. These results are shown below.



In the above figure, panel **a** shows a comparison of clustering coefficients between networks derived from the generalized linear coupling model (GLCM) and those constructed from partial noise correlation analysis. Each dashed line connects data from the same recording session. Panel **b** shows the Pearson correlation coefficients between connection weights estimated from noise correlation analysis and those derived from the GLCM.

Authors mention that the “averaged short-path length of the network decreased”. Does this analysis control for the total number of connections within the network? As you know, networks with higher density will naturally have lower path length.

Response: This is an important point, and we have addressed it in the revised manuscript. In the original version, we reported that the average shortest-path length of the network decreased as the monkeys’ performance accuracy improved across learning days. As the reviewer noted, such a decrease could potentially be driven by an increase in connection density.

However, this explanation is unlikely. Our linear regression analysis did not reveal a significant correlation between connection density and behavioral performance. To further rule out this confound, we normalized the average shortest-path length of each network by its corresponding connection density and repeated the regression analysis using this normalized measure. This analysis produced results that were consistent with those reported in the original

manuscript, which is shown below.

Behavior	Variable	Estimate	SE	tStat	pValue
Accuracy	Qmod	-0.8988	0.5408	-1.6618	0.1022
	Lambda	-13.5238	3.0974	-4.3661	5.6321e-05
	Qmod:LMD	18.7266	4.2691	4.3865	5.2543e-05
W_{stimulus}	Qmod	-5.7048	1.8650	-3.0588	0.0034
	Lambda	-61.7814	16.9183	-3.6517	0.0006
	Qmod:LMD	80.1984	19.0357	4.2131	9.4317e-05

In the revised manuscript, we have added a clarifying sentence to make this point more explicit, as follows:

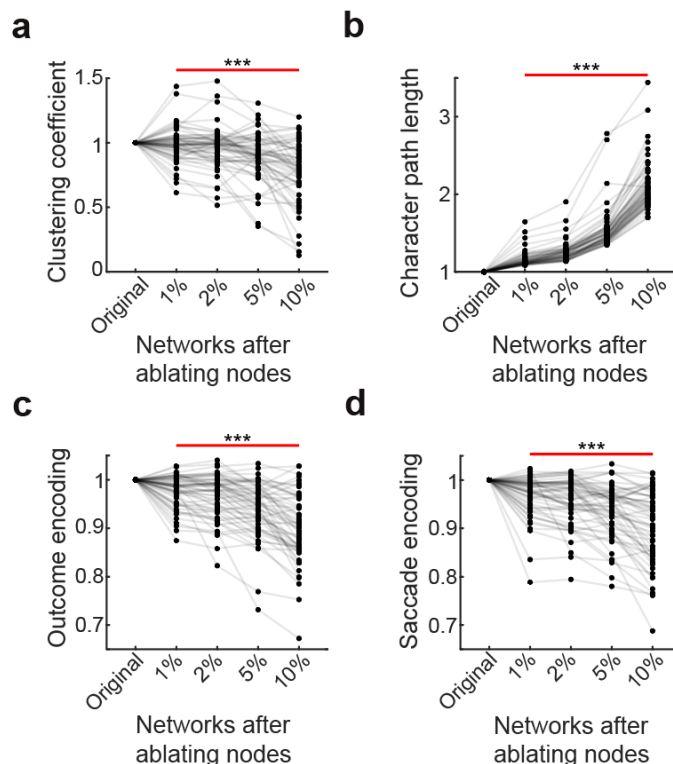
“The observed decrease in network short path length with increasing behavioral accuracy was not driven by changes in connection density over the course of learning, as similar results were obtained after normalizing short path length by the average connection density of the entire network.”

Assessing the causal role of hub neurons would involving functionally or anatomically ablating them, which is not possible here as it would be technically very challenging (which is something you mention at the end of your Discussion). However, you could run analyses where you remove the highest hubs from the functional connectivity networks and assess impact on small-world measures and task encoding. This would give you some indication of the causal role of these hubs.

Response: We thank the reviewer for this valuable suggestion. Indeed, selectively removing the most highly connected hub neurons and then reanalyzing the resulting network topology and population-level encoding provides an intuitive way to evaluate the potential causal role of hub neurons in supporting small-world organization and task-variable encoding. Following this suggestion, we conducted additional analyses in which we ablated the top 1%-10% highest-centrality hub neurons and recalculated the network metrics and population encoding performance. We found that both small-world network measures (clustering coefficient and shortest-path length) and the strength of neuronal encoding (saccade direction and trial outcome) were significantly changed relative to the intact network, indicating a potential causal contribution of hub neurons to maintaining small-world topology and encoding capacity. The results are shown below:

However, we acknowledge that this conclusion should be interpreted with caution. In this analysis, the functional network structure was treated as static after node deletion, whereas in vivo the network may reorganize dynamically when a subset of neurons is perturbed or removed. In the revised manuscript, we have incorporated these results into the revised Figure. 2 and supplementary figures

(Supplementary Fig. 5a-b) and referenced them in the Results section as follows:



The above figures illustrate changes in both small-world network metrics (a – b) and neuronal encoding strength (c – d) following the ablation of 1% to 10% of the top hub neurons in each network. Each dashed line connects results from an individual network. All metric values were normalized to those of the corresponding intact network. Ablation of hub neurons led to a reduction in the clustering coefficient, accompanied by a progressive increase in short path length as a larger fraction of hub neurons was removed. In parallel, both reward–outcome encoding and saccade–direction encoding showed significant declines following hub neuron ablation. Panels a and b have been included in the revised manuscript as supplementary Fig. 5a–b, whereas panels c and d are now presented as Fig. 2g–h in the revised manuscript.

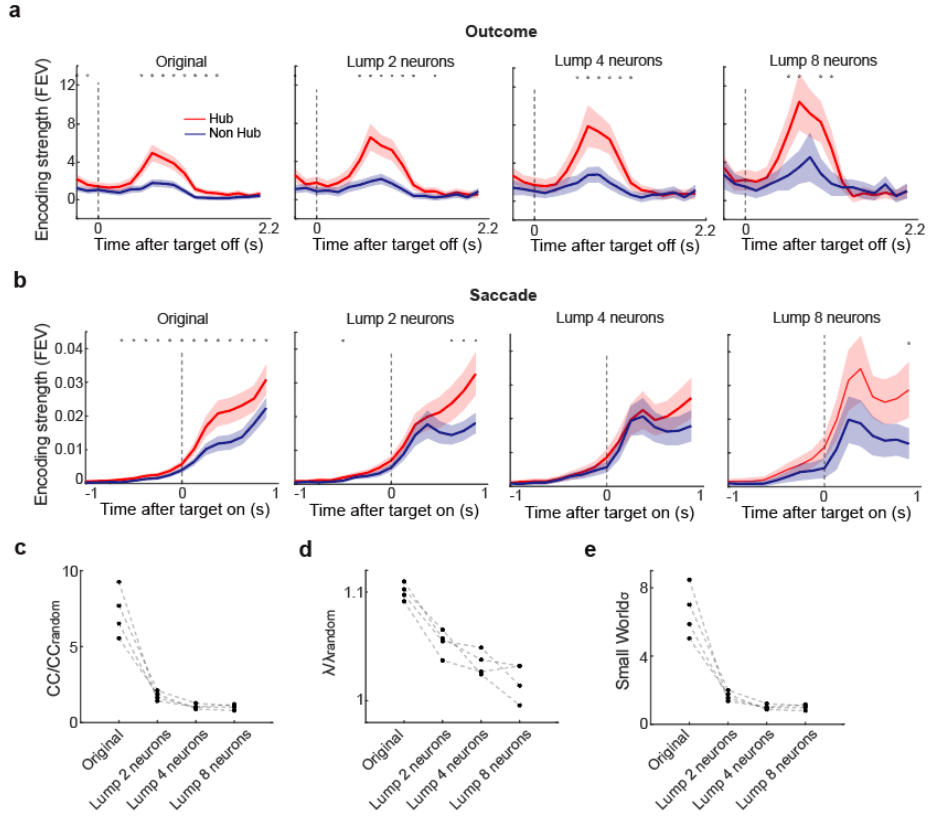
Additionally, we have revised the Discussion to reflect these new analyses and clarify the limitations of this approach, shown as:

“While virtual node ablation in network models can, in principle, be used to probe the theoretical contributions of specific topological elements to neural computation, in vivo neural networks may dynamically reorganize when subsets of neurons are perturbed or removed. Consequently, causal manipulations through targeted single-neuron interventions, combined with structural network mapping, are essential to uncover how network topology enables interactions among neural elements, providing mechanistic insights into the emergence of diverse forms of neural computation”

Authors mention that “Future research should investigate whether different types of learning induce distinct network dynamics across spatial scales and brain regions.” There is a straightforward way to address this, which I recommend doing. You can do a very simple multi-scale analysis where you lump together nearby neurons into small groups (by averaging their activity) and compute the functional connectivity between groups. In this way, each “node” of your functional network would be a group instead of a single neuron. Manipulating the size of groups would provide an indicate of multi-scale structure within functional networks. I’ m curious what you think you would find if you did that, in terms of small-world measures, task encoding, etc.

Response: We thank the reviewer for this suggestion. While averaging the activity of neuron groups of different sizes may provide some perspective on network dynamics across spatial scales, it offers limited insight given the constraints of our dataset. Our discussion of spatial scale focuses on two key considerations: (1) functional connectivity may systematically vary with inter-neuronal distance, and (2) network dynamics would likely differ substantially in a much larger population. In our recordings, each field of view contained only a few hundred neurons within an approximately $800 \times 800 \mu\text{m}$ region; therefore, averaging over small groups is unlikely to capture the broader spatial organization or dynamics of larger-scale networks, such as those involving millions of neurons across more distributed cortical areas. Furthermore, neurons were not evenly distributed across most FOVs, which prevented us from applying a consistent set of criteria to separate neurons into spatial groups across all recordings.

Nevertheless, following the reviewer’ s suggestion, we averaged the activity of neuron groups of different sizes (2, 4, and 8 neurons) and reconstructed the corresponding networks. Below, we present an example result from one session:



The above figures illustrate neuronal encoding (a – b) and small-world network metrics (c) for an example network after progressively aggregating nearby neurons. From left to right, the panels show results from the original network and from networks in which 2, 4, or 8 neighboring neurons were lumped together, respectively. As shown, the difference in neuronal encoding between hub and non-hub nodes showed a trend of gradual decrease as larger numbers of neurons were aggregated. Notably, small-world properties were largely abolished after aggregating nearby neurons. This suggests that small-world features are primarily expressed at the single-neuron level in local functional neuronal networks and are not preserved when neuronal activity is averaged across small local groups. This effect likely arises because different neuronal types do not exhibit strong spatial clustering at the spatial scale of our recordings. We did not include these results in the revised manuscript, as we do not consider averaging nearby neuronal activity to be an appropriate approach for probing network topology across spatial scales in our study.

In the Discussion, I encourage authors to speculate about the class subtype of hub neurons. I suspect that many such hubs may be inhibitory interneurons.

Response: We thank the reviewer for this insightful suggestion. Previous studies have identified GABAergic interneurons, characterized by extensive axonal arborization, as functional hubs in the hippocampus – entorhinal regions. We agree that many of the hubs in our study could similarly be inhibitory interneurons.

However, in our two-photon calcium imaging recordings, we captured both excitatory and inhibitory neurons, and classified them as hub or non-hub neurons based on their functional connectivity patterns within each FOV, with roughly equal numbers in each group. Both hub and non-hub neurons exhibited positive and negative functional connectivity with other neurons. We therefore also speculate that some identified hubs may be excitatory neurons with extensive network connections. In response to this comment, we have added a sentence in the Discussion section to speculate on the potential subtypes of hub neurons, as follows:

“we identified hub neurons based on the architecture of functional connectivity, which may include both excitatory and inhibitory subtypes, each exhibiting extensive connectivity within the network.”

Also, in the Discussion please speculate about whether you believe that learning-driven changes in functional connectivity would revert with a prolonged period without training on the task.

Response: This is an important and valuable suggestion. Previous studies have shown that neuronal encoding in the PPC of rodents can gradually drift across days, even when animals repeatedly perform a well-trained task (Laura N Driscoll et. Al., 2017; Rule, E. D., et al., 2020). In line with this observation, we speculate that the functional neuron network in primate PPC may also undergo slow, inherent evolution across days, independent of active learning. However, we expect this intrinsic drift to occur at a much slower rate than the network changes driven by learning. Therefore, we speculate that learning-induced changes in functional connectivity may partially revert after an extended period without task training, while still exhibiting slow, spontaneous drift over time. In the revised discussion section, we have added a new sentence to discuss it as follows:

“We speculate that learning-induced changes in network topology may partially revert after a long period without task training, while still exhibiting slow, spontaneous drift over time. “

Minor comments:

Authors mention that “Additionally, the connection weights did not significantly correlate with the spatial distance between neurons within most FOVs”. Clearly, this means that functional connectivity does not reflect monosynaptic connections. Please add a statement to that effect.

Response: We thank the reviewer for this helpful suggestion. We agree that this result indicates that functional connectivity does not necessarily reflect monosynaptic connections. In the Results section, when introducing partial noise

correlation, we have already noted this point (page 3, lines 22–24 in the original manuscript; page 3, lines 126–127 in the revised manuscript):

“While it does not necessarily represent monosynaptic connections, it reflects the influence of one neuron on another, significantly reducing the impact of common input and task-related modulation.”

Furthermore, we also speculate that the spatial extent of our recording area may be insufficient for physical distance to play a major role. In response to the reviewer’s comment, we have added a sentence in the revised manuscript to clarify this interpretation.

“This suggests that the inferred connectivity between neurons may not correspond to monosynaptic connections, or that the spatial extent of each FOV is too limited for physical distance to have a significant effect.”

Fig.1B-C: figure legend should specify what is the meaning of “L-L”, “L-R”, etc. (x-axis)

Response: We have added one sentence in the revised figure legend to specify them as follows:

“Each column represents a different pair of neuron groups (L: left-preferring neurons; R: right-preferring neurons; C: correct-preferring neurons; W: wrong-preferring neurons; N: neurons with no significant selectivity, L-L: left-preferring neurons to left-preferring neurons, etc.”

Authors state that “sparsity of their connection matrix was 0.048”. Can you clarify: do you mean that 4.8% of all possible connections were present? If so, it seems a bit low compared to known anatomical connectivity within local cortical networks.

Response: The reviewer is correct that, on average, 4.8% of all possible connections were present within the network. This sparsity is influenced by our thresholding criteria. Using our current threshold ($p < 0.001$), approximately 5% of all connections were deemed significant. Because the network is non-directional, however, the true underlying connectivity could range between 4.8% and 9.6%.

Author state that “with modularity coefficients substantially exceeding those of random networks (Fig. 1H).” I am confused here, I did not think that H was a random network. Please rephrase or clarify.

Response: Thank you for catching this typographical error. We have corrected it

in the revised manuscript — it should refer to Fig. 1G.

Please add a final paragraph at the end of your Discussion that summarizes your main findings and the contribution of your study in few punchy sentences.

Response: We thank the reviewer for this valuable suggestion. In the original manuscript, we used the first paragraph of the Discussion section to summarize our main findings and contributions; however, we now recognize that this may not have been sufficient. In response to the reviewer's advice, we have added a concluding paragraph that explicitly summarizes our key findings and contributions to the field, as shown below:

‘Together, our study demonstrates that small-world functional network topology in behaving animals actively shapes functional roles of single neuron, inter-neuron communication, and long-term plasticity that support task performance and learning. These findings establish a fundamental mechanistic link between mesoscale functional connectivity, neural computation, and behavior from an integrated network perspective, and introduce a new framework for tracking single-neuron plasticity at the functional graph level to uncover the mechanisms underlying learning.’

Typos:

Please add line numbers to your submission to facilitate the review.

Response: We've overhauled the line numbers.

Figure 1G legend: “random networks” should be “random”

Response: Thanks for pointing out this typo. We have corrected it.

“hub neurons exhibited significantly greater connection strength and degree compared to non-hub neurons” specify: in-degree or out-degree? Or both?

Response: This is an important question. However, it cannot be addressed with the current method used to construct the functional neuron network. In this study, we generated non-directional graphs based on partial noise correlations. As a result, the degree measure reflects a combined in-degree and out-degree without distinguishing between them. We have revised the Methods section to explicitly clarify this point, as follows:

“This undirected matrix was used as the weighted adjacency matrix in all subsequent analyses of the non-directional graph”

Figure 1: on the figure itself, the bold letters “I” and “L” appear exactly the same. Please rewrite “L” correctly.

Bottom of p.6: “after the transitioned” replace with “transition”

p.8: “within-module pair” replace with “pairs”

p.14: “(suite2p,)” remove the extra comma

“sessions\days” replace with a forward slash “/”

« hub\non-hub” use a forward slash

Response: Thanks for correcting these typos. We have corrected them in the revised manuscript.

Reviewer #4 (Remarks to the Author):

Dear authors,

Your manuscript entitled “Single-neuron network topology governs neural computation and learning in primate cortex” presents an in-depth network analysis with refined methodology of two-photon calcium imaging of parietal cortex neurons imaged from superficial layers of 7a in awake behaving macaques performing a sensorimotor association task. Some of the results presented on hub vs. non-hub neurons are somehow expected by definition e.g. modular computation within PPC network. The most interesting aspect of your study is the observed dynamics on the hub status of neurons. In particular, the fact that status transition (hub to non-hub and vice-versa) results from the network’s influence and inter-modular connectivity during learning. This novel aspect could be given more emphasis. The paper is well written and clearly describes the methods used to obtain the results which support the conclusions. However, some aspects of the manuscript warrant correction and/or clarifications and are described below:

1. The statement in the last sentence of the abstract shall be revised. Indeed, the imaging of superficial cortical layers has been made presumably on a 5mm diameter window over PPC and therefore the reference to orchestrating global computations shall be removed or changed with local.

Response: Thank you for this comment. We recognize that the original abstract may not have been sufficiently clear and could potentially confuse readers. In particular, the phrase “orchestrating global computations” should be used more cautiously, as our study focused on neuronal networks in the superficial layers of a relatively small region within PPC. Here, “local” and “global” refer to the observed local network within PPC. We have revised the sentence in the abstract as follows:

“These findings reveal how single-neuron-resolution brain networks, through

small-world organization, orchestrate both global and modular neural computations within local network to mediate behavior and shape learning. ”

2. In the paragraph p11 lines 16–30 of the Discussion. The authors might also want to consider that discrepancies could also arise from the limitation of the presented results to superficial cortex.

Response: We thank the reviewer for this valuable suggestion. In the revised manuscript, we have added a sentence in this paragraph discussing a potential reason for the discrepancy between our findings and previous studies. Specifically, this difference may arise from the limitation that two-photon calcium imaging primarily tracks neuronal activity in the superficial cortical layers (layers 2–3), as shown below:

“While our findings appear to diverge from previous reports, these discrepancies may reflect differences in spatial scale—our micro-scale measurements versus macroscopic network observations—as well as differences in brain regions and learning paradigms. Additionally, our recordings were restricted to superficial cortical layers, which may exhibit distinct plastic properties compared to deeper layers or whole-region averages. ”

3. On Figures 2C and S10F, I; the mention of ‘Time to target off’ is really confusing. It seems, from the Methods section that the time 0 mention time of target offset. Therefore it would be less confusing to change the X axis label to ‘Time after target off’ . Especially in comparison to e.g. Figure 2 B axis.

Response: Following the reviewer’ s suggestion, we have revised the corresponding figures and updated the X-axis label to “Time after target off.”

4. The repeated references to ‘(…) our previous study (…)’ p3, lines11–14; p4 lines 41–42; p11 line 41; p12 lines 17 and line 32; p13 line 15. The terms previous study can be misleading for the reader, indirectly implying published work. The reference #52 is a bioRxiv paper submitted in August 2025, that has not been peer-reviewed. Hence, I strongly suggest you refer to it in the present tense rather than the past tense e.g. ‘our work also shows’ or ‘we are also showing’ .

Response: We thank the reviewer for this helpful comment. We acknowledge that the way we referred to this unpublished work in the original manuscript may have caused confusion. The referenced paper has recently been published (in November, 27, 2025). In accordance with the reviewer’ s suggestion, we have revised the text to use the present tense rather than the past tense when referring to this study.

5. Related to the above point, the reference to ‘our previous study’ in the Material and Methods section refers to a wrong publication #1 on p4 lines 41–42; p11 line 41; p12 lines 17 and line 32; p13 line 15. Also on p13 line3, it seems your actually refers to other publications than reference #1 and #2. Please correct.

Response: We thank the reviewer for pointing this out. In the original manuscript, the reference numbering was reset in the Methods section while the references were merged in the Reference list, which caused inconsistency. In the revised manuscript, we have corrected the reference numbering in the Methods section.

6. This study resonates with previous work from Dann et al., eLife 2016 <https://doi.org/10.7554/eLife.15719> and would gain including other references and comparing their results in the discussion to similar network topology analyses on parietal cortex neurons

Response: We thank the reviewer for this valuable suggestion. We agree that the work referenced is highly relevant to our study. Although it was cited in the original manuscript, we recognize that it was not given sufficient attention considering its significance. In the revised manuscript, we have added a short paragraph in the Discussion to compare our findings with this and other mesoscale neural network studies, as shown below.

“We systematically investigated how the network’s topological features such as hub status and modular organization, as well as their dynamics, relate to neuronal encoding, inter-neuron communication, long-term plasticity, and behavioral performance throughout learning. Recent studies have also begun to interpret functional connectivity through the lens of network topology. For example, an electrophysiological study in behaving monkeys demonstrated that functional networks in the frontoparietal cortex exhibit small-world organization and highlighted the pivotal role of rich-club hub neurons in driving oscillatory dynamics³⁰. Together, these studies and our findings highlight the necessity of understanding sophisticated network architecture to fully characterize neural computation and its behavioral consequences.”

Furthermore, we have revised the text of the third and fourth paragraphs in the Discussion to more explicitly reference the cited work and to better relate our findings to these previous studies.

Minor points:

Please include the macaque species used in your study.

Response: We have added the information about the macaque species in the revised Methods section as follows:

“Two male Rhesus Macaques, aged 4–5 years and weighing 5–6 kg, were trained on the behavioral task.”

p13 line 7, you refer the ‘monkey brain atlas’. Would be more accurate to specify that 7a was located based on the gyral and sulcal landmarks in comparison to an atlas with a reference to it.

Response: Thank you for this valuable suggestion. In the revised manuscript, we have added a sentence specifying the anatomical landmarks used to identify area 7a, and included an additional reference to support this description, as shown below:

“We located area 7a according to the macaque brain atlas, where it is defined as the posterior portion of the inferior parietal lobule⁸⁶. Anatomically, area 7a lies posterior to and inferior to the intraparietal sulcus, anterior to the parieto-occipital sulcus, and dorsal to the superior temporal sulcus. Our recordings targeted the exposed cortical surface of area 7a, which is readily accessible for two-photon calcium imaging.”

Could you please describe the quality criteria used in Neuronal data analysis?
p14 line 24

Response: We thank the reviewer for pointing this out. In the revised Results section, we have added a sentence clarifying the quality criteria used to identify ROIs, as follows:

“We manually filtered the results by excluding ROIs that were either smaller than 5 pixels in diameter, smaller than 15 pixels² in area, or displayed clearly abnormal, non-cellular shapes.”

Authors' Response to Reviewers

We sincerely thank the reviewers for their positive feedback on our previous revision and for their continued constructive comments. We have addressed the remaining points in the responses below and have updated the manuscript accordingly. Furthermore, we have substantially restructured the MATLAB code and documentation on Code Ocean. In this document, the reviewers' comments are shown in black, and our point-by-point responses are highlighted in blue.

Reviewer #1 (Remarks to the Author):

Thank you very much for addressing my concerns. I understand your explanations, and I have no further input.

Cosmetic suggestions:

- On Eqn. 23, the index of the summation (j) is omitted.
- On Eqn. 25, I think the "narrative" may still be in math mode, not within `\text{}` or `\textrm{}` brackets (or a similar command).
- On supplementary fig. 4 (Fig. S4), I think there is a missing space after AUC in the y-axis label.

Response: We thank the reviewer for pointing out these issues regarding the formulas and the figures. We have revised the above content according to the suggestions.

Reviewer #1 (Remarks on code availability):

I didn't read all the files, but I skimmed over the structure and ran the "Code Ocean" capsule demo. I obtained the intended results, and the walkthrough was clear from the demo ".txt" file. I suggest that authors check some of the text in the comments and ReadMe files for typos, although it is overall clear, and I had no difficulties running it.

Response: We thank the reviewer for their comments regarding the Code Ocean capsule and the associated README files. We have restructured the code and verified the technical details mentioned in the feedback. Accordingly, the README file and the code capsule have been updated to ensure clarity and reproducibility.

Reviewer #2 (Remarks to the Author):

Dear Authors,

The revised version of the manuscript addressed most of the concerns I raised. However, I feel that Major Concern 3 was not address in enough detail.

I suggest to add the relevant references associated with the field of Granger causality analysis and add the results shown in the rebuttal in the Supplementary material.

For what concerns the references I would suggest the following:

Granger, C.W.J. (1980) 'Testing for causality: A personal viewpoint', *Journal of economic dynamics & control*, 2, pp. 329–352.

Brovelli, A. et al. (2004) 'Beta oscillations in a large-scale sensorimotor cortical network: directional influences revealed by Granger causality', *Proceedings of the National Academy of Sciences of the United States of America*, 101(26), pp. 9849–9854.

Bressler, S.L. and Seth, A.K. (2011) 'Wiener–Granger Causality: A well established methodology', *NeuroImage*, 58(2), pp. 323–329

Best regards

Response: We sincerely appreciate the reviewer's valuable suggestion. We agree that a more thorough comparison between our Causality Index and well-established methods, such as Granger Causality, was needed to better contextualize our findings.

As noted in the Methods section, the Causality Index we employed is a variant of Granger Causality; we provided a detailed justification for this specific choice in our previous response. To further validate our approach as suggested, we have now included Supplementary figures (Supplementary Fig. 10d and Supplementary Fig 12), which demonstrates that our method yields qualitatively similar results to classic Granger Causality analysis. Additionally, we have revised the corresponding Results section (page [Xxx]) to incorporate the suggested references and explicitly relate our index to established directional measures. The added text is as follows:

Line 243 to line 255: *'We then measured how each neuron's current activity influenced the future activity of other neurons within the local network using a causality index. This index—a variant of Granger causality⁶⁵⁻⁶⁷—is defined as the partial correlation between a neuron's activity in the current time period and the activity of other neurons in future time periods within the same task trial, conditioned on the current activity of these neurons (see Methods). Like Granger causality, this index quantifies directional influence strength between time series, with higher absolute value indicating greater influence of a neuron's activity on the activity dynamics of other neurons over time. We computed the causality index across successive time frames for both hub and non-hub neurons. Across networks, hub neurons exhibited significantly higher outgoing causality (influencing other neurons) but lower incoming causality (being influenced) compared to non-hub neurons. (Figure 3e-f, Supplementary Fig. 10a-c). The observed causal lead from hub neuron to non-hub neurons was independently corroborated by classic Granger causality analysis (Supplementary Fig. 10d).'*

Line 372 to line 379: *"To test this, we calculated the causality index separately for within-module and across-module pairs using condition-averaged activity across successive time frames within a task trial. We found a higher causality index for the within-module pairs compared to the across-module pairs in most networks (Figure 5h, $P(+1f) = 2.3e-25$, $t(58) = 18.0$, $P(+2f) = 3.4e-23$, $t(58) = 16.2$, $P(+3f) = 3.7e-20$, $t(58) = 13.9$, paired*

t-test). This heightened influence within modules was independently corroborated by Granger causality analysis (Supplementary Fig. 12). These results suggest that within-module pairs had a stronger mutual influence on their activity dynamics during monkeys' task performance. ”.

Reviewer #2 (Remarks on code availability):

The code is an unstructured set of directories with tens of Matlab codes.

The Matlab codes are not clearly commented and the structure of the code is opaque.

The code does not seem to make reference to the results shown in the main text.

I think the code is not user friendly, and it should require some rewriting and re-structuring.

Response: We sincerely appreciate the reviewer's critical feedback regarding our code. We recognize that the previous version lacked the necessary structure and clarity for external use. To address these concerns, we have performed a comprehensive overhaul of the codebase:

1. **Structural Reorganization:** We have refactored the previously unstructured directories into a standardized, hierarchical format. By modularizing the core scripts, utility functions, and plotting routines, the code can now seamlessly replicate the primary analyses presented in the main figures.
2. **Enhanced Documentation:** We have added more detailed comments to all MATLAB scripts and provided demo scripts to guide users through the entire analysis pipeline.
3. **Data & Reproducibility:** To accommodate the computational and storage constraints of the Code Ocean platform, we have provided a representative validation dataset consisting of recordings from two successive learning days. While the full raw dataset (exceeding several terabytes) requires significant high-performance computing (HPC) resources, this demo allows users to verify the code's functionality and generate results that are qualitatively and procedurally consistent with the manuscript's findings within a reasonable timeframe.

Finally, we have updated the README file to include a direct mapping between specific code modules and their corresponding figures and analytical results in the manuscript. The format of analytical results (.txt files) has also been updated for clear readout.

Reviewer #3 (Remarks to the Author):

The reviewers have done a great job addressing my comments from round 1. I only have two remaining comments:

- (1) In the Discussion, please add a sentence or two to speculate about what synaptic plasticity mechanisms may be responsible for (a) small-world functional connectivity

networks in PPC; and (b) gain or loss of hub status. Around line 561, you state "Future research should investigate whether different types of learning..." but unless you specify what potential mechanisms (Hebbian plasticity? STDP? Homeostatic plasticity?) it is difficult to link your results to synaptic function. Even if this is speculative, at least provide some ideas of how your results on small-world and hub status may arise in a biologically plausible fashion.

Response:: We thank the reviewer for this insightful suggestion. We agree that bridging the gap between our network-level observations and potential synaptic mechanisms enhances the biological relevance of our findings.

Following the reviewer's suggestion, we have expanded the Discussion to speculate on the underlying synaptic plasticity rules. Specifically, we propose that the emergence of small-world topology may be driven by a combination of local Hebbian reinforcement and homeostatic scaling. Furthermore, we suggest that Spike-Timing-Dependent Plasticity (STDP) could provide a "rich-get-richer" mechanism for the gain of hub status, while synaptic pruning or homeostatic down-regulation might account for its loss. We have added one short paragraph to discuss these points as follows:

Line 558 to line 567: *'Future research should explore the specific synaptic rules governing these topological changes during learning. We propose that the small-world topology may arise from a synergy between local Hebbian reinforcement of functional clusters and homeostatic pruning of redundant long-range paths. Additionally, Spike-Timing-Dependent Plasticity could facilitate "rich-get-richer" dynamics — a form of preferential attachment where synchronized firing drives the emergence of hubs^{88,89}. Conversely, hub status may be lost via homeostatic scaling or metabolic constraints to prevent runaway excitation and maintain network economy^{5,90}. Bridging these micro-scale synaptic mechanisms with meso-scale and macro-scale network dynamics is essential for a comprehensive understanding of learning'*

(2) A small typo: line 527: *A* previous study

Response: Thank you for catching this typographical error. We have corrected it in the revised manuscript.

Reviewer #4 (Remarks to the Author):

Dear authors,

you have satisfactorily addressed all my concerns and I have no further comments.

Sincerely

Response: We sincerely thank the reviewers for the positive feedback on our previous revision.

Reviewer #4 (Remarks on code availability):

I reviewed briefly the code provided but could not run it in Code Ocean.

Due to default in the activation of my Code Ocean account, I could not verify if the code runs well on the Day1 preprocessed dataset provided.

However, codes seems legit and basic README files are provided with enough instructions and details to be able to run the code (including toolboxes and environment required, brain connectivity toolbox and other functions or m files used for the analysis) and reproduces the figures listed. The authors shall engage to ensure response to any request in case of missing data/information.

Response: We thank the reviewer for the suggestions regarding the codes. We have once again reorganized and corrected the code, and it is now able to run successfully in Code Ocean. We also updated the ReadMe file and the comments in the code simultaneously to make them more understandable and easier to run.