

Supplementary Materials

1. Stimulus characteristic-related temporal discriminability

To evaluate whether the observed language differences in neural informativeness could be attributed to differences in stimulus-driven temporal discriminability, we examined the intrinsic temporal separability of the narrative stimuli using LaBSE-based representations. Story-wise semantic similarity analysis showed that narratives derived from the same comic strips were highly aligned across languages (mean cross-language similarity = 0.80), with a strong second-order correlation between the semantic RSMs of L1 and L2 ($r = 0.78$, $p < .001$), confirming high semantic consistency between the two language versions. At the temporal level, we computed $TR \times TR$ dissimilarity matrices based on token-level embeddings extracted at multiple transformer layers (layers 3, 7, and 11) and compared the distributions of all pairwise TR dissimilarities between languages. Non-parametric comparisons revealed that temporal discriminability of the stimulus representations was slightly higher in L1 than in L2 across layers (all $ps < .009$), exhibiting a pattern opposite to that observed in the neural decoding results. These findings rule out the possibility that the higher neural informativeness observed during L2 processing was driven by greater temporal separability of the stimulus itself, and instead support the interpretation that the language differences reflect actual differences in neural representation rather than stimulus properties.

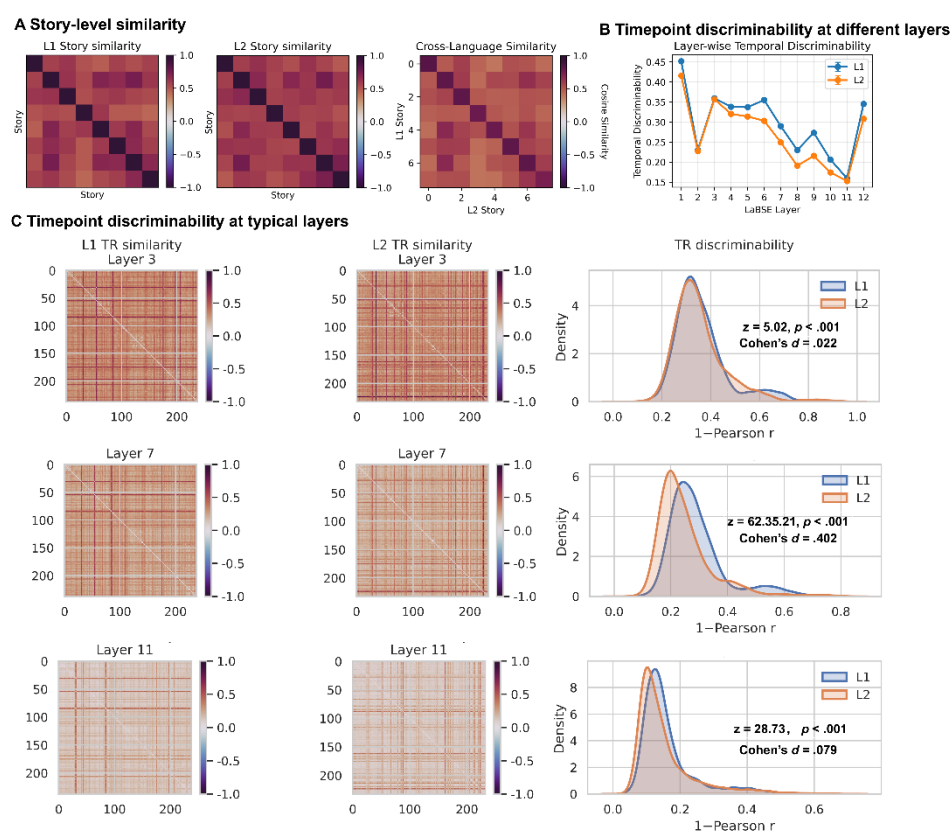


Fig. S1 Stimulus-induced temporal discriminability across languages. A. Within-language and cross-language story-level semantic similarity. B. Temporal discriminability of two languages at

different embedding levels of LABSE model. C. Temporal discriminability at 3, 7, 11 layers and the distribution of temporal embedding similarities.

2. Decoding accuracy of lesion of all possible network pairs

In addition to the lesion simulation of cortico-cerebellar network pairs, we performed the analysis after lesion of all possible network pairs and use the decoding accuracy from all the unrepeated edges as the original accuracy. The results revealed that lesion of any network pairs did not lead to significant change of decoding accuracy during L1 narrative comprehension ($p < 0.05$, FDR corrected). However, the lesions between the cortical DMN and other cortical networks including the DAN, VAN, FPN as well as within the DMN resulted in dramatically decrease of decoding accuracy during L2 narrative comprehension ($p < 0.05$, FDR corrected). Notably, no lesion effect of any network pair was significant after controlling feature number as compared to the original accuracy from the whole brain.

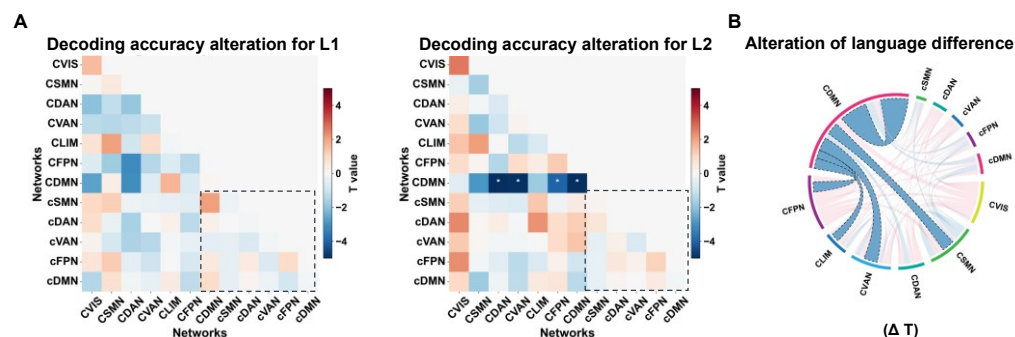


Fig. S2 Simulation lesions of all the network pairs. A. Decoding accuracy changes after lesion of specific network pair within languages. ***: $p < 0.001$, **: $p < 0.01$, *: $p < 0.05$ after FDR correction. B. Alterations between the language difference after the specific lesion and the original difference. No significant lesion effect or language differentiation effect was found after removing any specific network-pair using the FDR corrected p-values obtained from the permutation test.

3. Control analysis for the signal quality between the cortical and cerebellar nodes

To further examine whether the decoding accuracy differences among connectivity types, including intra-cerebellar, cortico-cerebellar, and intra-cortical connectivity, could be driven by signal-quality differences between cortical and cerebellar nodes, we performed an additional tSNR-matched control analysis. In this analysis, tSNR was computed at the voxel level using the full, non-standardized voxel-wise BOLD time series, defined as the temporal mean divided by the temporal standard deviation. ROI-level tSNR was then obtained by averaging voxel-wise tSNR values within each ROI and further averaging across runs and participants. This approach provides a more direct estimate of local BOLD signal quality than computing tSNR from ROI-averaged time series.

We then constructed a signal-quality-matched cortical node pool. Specifically, cortical ROIs were first restricted to those whose ROI-averaged voxel-wise tSNR fell within the range of cerebellar tSNR values. This initial pool was then iteratively refined using a greedy mean-matching procedure: at each step, we removed the cortical ROI

whose exclusion most reduced the absolute difference between the cortical-pool mean and the cerebellar mean tSNR. The refinement was stopped once the pool-level tSNR difference reached a negligible effect size, defined as Cohen's $d \leq 0.2$, thereby retaining as many cortical nodes as possible while achieving mean-level tSNR matching. The final matched cortical pool showed comparable mean voxel-wise tSNR to the cerebellar nodes, with a negligible mean difference: matched cortical nodes, $n = 177$, mean = 74.30, SD = 19.74, range = 29.04–101.23; cerebellar nodes, $n = 27$, mean = 72.28, SD = 17.85, range = 28.70–101.24. A Welch's t-test comparing the full matched cortical pool with the cerebellar nodes confirmed that the two node sets did not differ significantly in mean tSNR, $t(36.4) = 0.538$, $p = .594$, Cohen's $d = 0.107$. These results indicate that the cortical node pool used for subsequent sampling was well matched to the cerebellar nodes in average voxel-wise tSNR.

We then repeated the decoding comparisons using this restricted tSNR-matched cortical node pool while keeping the number of connectivity features matched across conditions using the same sampling method reported in the main text. The same pattern of results remained significant in both languages: tSNR-matched cortico-cerebellar connectivity still yielded higher decoding accuracy than intra-cerebellar connectivity, and intra-cortical connectivity still yielded higher decoding accuracy than cortico-cerebellar connectivity, with all p values $< .001$. These control analyses suggest that the observed differences between connectivity types cannot be explained solely by lower average voxel-wise tSNR in cerebellar nodes.

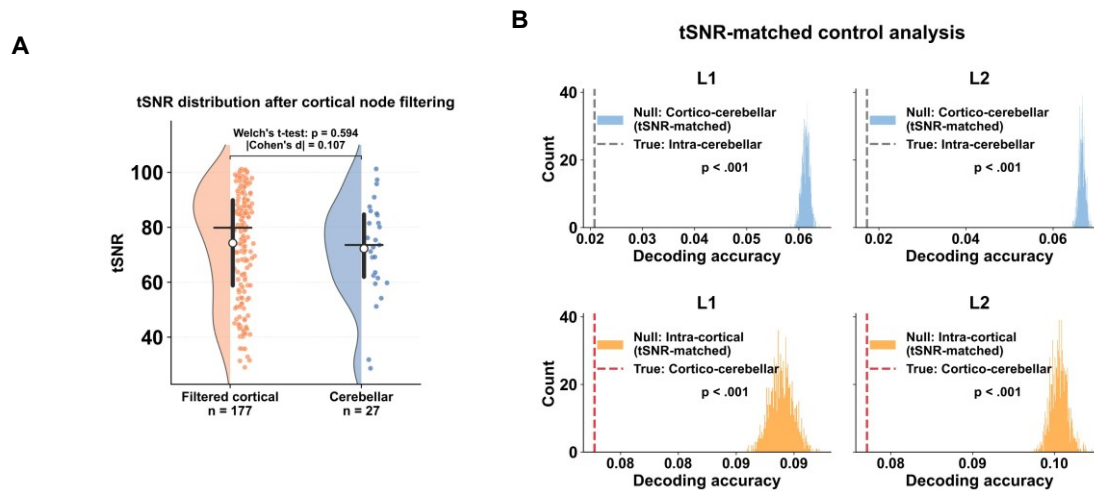


Fig. S3 Control analysis for the signal quality. A. Distribution of temporal signal-to-noise ratio (tSNR) values for the tSNR-matched cortical node pool and all cerebellar nodes. tSNR was computed at the voxel level using the full, non-standardized voxel-wise BOLD time series, and ROI-level tSNR was obtained by averaging voxel-wise tSNR values within each ROI and across runs and participants. The matched cortical pool was constructed by first restricting cortical ROIs to those whose ROI-averaged voxel-wise tSNR fell within the range of cerebellar tSNR values, followed by iterative refinement to match the mean tSNR of the cerebellar nodes while retaining as many cortical nodes as possible. Points indicate individual ROIs, half-violins show the estimated distributions, vertical bars indicate the interquartile range, and white circles indicate group means. The final matched cortical pool showed comparable mean tSNR to the cerebellar nodes, with a negligible mean difference confirmed by Welch's t-test and Cohen's d . B. Decoding control analyses performed using the tSNR-matched cortical node pool. The upper panels compare the null

distributions of tSNR-matched cortico-cerebellar connectivity with the observed decoding accuracy of intra-cerebellar connectivity in L1 and L2. The lower panels compare the null distributions of tSNR-matched intra-cortical connectivity with the observed decoding accuracy of cortico-cerebellar connectivity in L1 and L2. Histograms indicate null distributions generated by repeated random sampling with matched feature numbers, and dashed vertical lines indicate the observed decoding accuracy of the corresponding reference condition. The results remained significant in both languages, indicating that the observed decoding differences among connectivity types cannot be explained solely by lower average voxel-wise tSNR in cerebellar nodes.