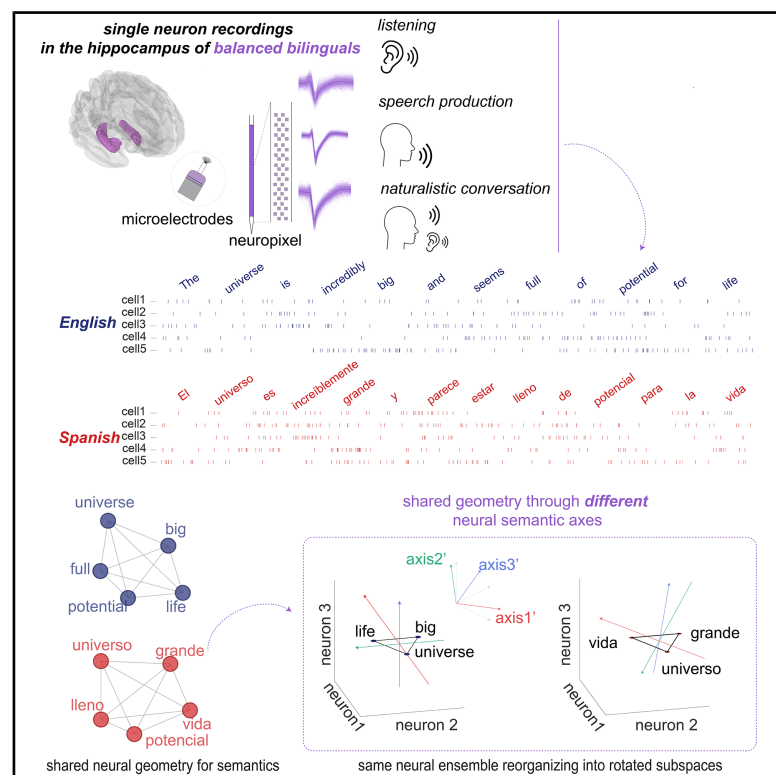


Shared neural geometries for bilingual semantic representations in human hippocampal neurons

Graphical abstract



Authors

Xinyuan Yan, Ana G. Chavez,
Melissa Franch, ...,
Seng Bum Michael Yoo,
Benjamin Y. Hayden, Sameer A. Sheth

Correspondence

xinyuan.yan@bcm.edu (X.Y.),
benjamin.hayden@bcm.edu (B.Y.H.),
sameer.sheth@bcm.edu (S.A.S.)

In brief

In bilinguals, hippocampal recordings during listening and speaking revealed overlapping neural geometry for English and Spanish. Each language accesses this shared meaning space through rotated neural readout axes, showing how the brain represents concepts across languages without interference.

Highlights

- Words with similar meanings in English and Spanish have shared geometry
- Hippocampal neurons have different semantic tunings for English and Spanish
- English and Spanish are differentiated using rotated readout axes
- The hippocampus supports translation of meaning for both speaking and listening

Yan et al., 2026, Cell 189, 1–16

August 6, 2026 © 2026 Elsevier Inc. All rights are reserved, including those for text and data mining, AI training, and similar technologies.

<https://doi.org/10.1016/j.cell.2026.05.020>

Article

Shared neural geometries for bilingual semantic representations in human hippocampal neurons

Xinyuan Yan,^{1,*} Ana G. Chavez,¹ Melissa Franch,¹ Kalman A. Katlowitz,¹ Ivy Gautam,¹ Brian Kim,¹ Aaditya Krishna,¹ Aadit Shrivastava,¹ Katie Van Arsdell,¹ James Belanger,¹ Assia Chericoni,¹ Taha Ismail,¹ Elizabeth A. Mickiewicz,¹ Danika Paulo,¹ Hanlin Zhu,¹ Alica M. Goldman,² Vaishnav Krishnan,^{2,3} Atul Maheshwari,² Eleonora Bartoli,^{1,3} Nicole R. Provenza,^{1,3,4} Seng Bum Michael Yoo,^{5,6,7} Benjamin Y. Hayden,^{1,3,9,10,*} and Sameer A. Sheth^{1,3,8,9,*}

¹Department of Neurosurgery, Baylor College of Medicine, Houston, TX 77030, USA

²Department of Neurology, Baylor College of Medicine, Houston, TX 77030, USA

³Neuroengineering Initiative and Electrical and Computer Engineering, Rice University, Houston, TX 77005, USA

⁴Department of Bioengineering, Rice University, Houston, TX 77005, USA

⁵Department of Biomedical Engineering, Sungkyunkwan University, Suwon 16419, South Korea

⁶Department of Intelligent Precision Health Convergence, Sungkyunkwan University, Suwon 16419, South Korea

⁷Center for Neuroscience Imaging Research, Sungkyunkwan University, Suwon 16419, South Korea

⁸Gordon and Mary Cain Pediatric Neurology Research Foundation Laboratories, Jan and Dan Duncan Neurological Research Institute, Texas Children's Hospital, Houston, TX, USA

⁹These authors contributed equally

¹⁰Lead contact

*Correspondence: xinyuan.yan@bcm.edu (X.Y.), benjamin.hayden@bcm.edu (B.Y.H.), sameer.sheth@bcm.edu (S.A.S.)

<https://doi.org/10.1016/j.cell.2026.05.020>

SUMMARY

The human brain has the remarkable ability to comprehend and express similar concepts in multiple languages. To understand how it does so, we examined responses of hippocampal neurons during passive listening, directed speaking, and spontaneous conversation in both English and Spanish in a small group of balanced bilinguals. We found a small number of putative “cross-language neurons,” whose responses to equivalent words (e.g., “tierra” and “earth”) are correlated. However, neurons’ semantic tunings differed substantially by language, suggesting language-specific neural implementations. Instead, the crucial driver of translation was a preserved geometric organization of neural responses between the two languages, one that did not depend on neuron-level functional overlap. Indeed, that geometry was implemented by a common set of neurons along distinct readout axes; this difference in readout may help prevent cross-language interference. Together, these results suggest that the hippocampus encodes a language-independent internal model for meaning.

INTRODUCTION

Humans have the remarkable capacity to understand and express the same thoughts in multiple languages without confusing them.¹ For example, in Arnhem Land (Northern Australia), speakers commonly manage six or seven languages.² A growing body of evidence suggests that bilingual speakers use at least partially shared neural substrates for each of their languages. For example, in bilingual brains, fMRI studies of listening and reading show broadly overlapping language-network responses across languages.^{3–12} Even constructed languages (e.g., Klingon and Na'vi) reliably recruit the same brain areas as natural languages.¹³ Regions that elicit overlapping responses include classic language regions, such as the inferior frontal gyrus (IFG) and posterior temporal cortex, as well as several others not as closely associated with language. However, these findings do not tell us *how* the brain matches corresponding concepts between languages while simultaneously maintaining a functional separation.

We hypothesized that the bilingual brain uses shared neural geometry for representing meaning in both languages. That is, the brain may have similar sets of coding distances between all pairs of words in high-dimensional neural space. Thus, if “cat” and “dog” elicit similar neural activation patterns and “door” is dissimilar, then “gato” and “perro” should evoke similar neural activation patterns and “puerta” should be dissimilar. Indeed, this is the strategy used by multilingual Large Language Models (LLMs), such as multilingual Bidirectional Encoder Representations from Transformers (mBERT).^{14,15}

A small body of neuroimaging work suggests that the brain may employ shared semantic geometry. In particular, a recent neuroimaging study shows that voxelwise encoding models trained to predict neural activity from language-model embeddings (mBERT) in one language can transfer across languages within the same bilingual participants.¹⁶ Other studies show cross-linguistic generalization across different monolingual groups listening to translations of the same story^{6,17,18} so that encoding models trained on neural responses from speakers

of one language predict neural responses in the other.^{6,17,18} One important study used electrocorticography (ECoG), which provides dense sampling of local field potentials (LFPs), with corresponding bilingual speech in one bilingual patient.¹⁹ This study found shared articulatory codes sufficient for cross-language decoding, although it did not investigate semantics. Moreover, none of these studies directly examined whether *shared semantic geometry* (i.e., representational geometry of semantics) exists.

If shared semantic geometry provides the scaffolding that links language to meaning, an important next question is how different languages access this common structure. The properties of representational geometry suggest that rotations of neural axes preserve pairwise relationships, leaving the overall geometric structure invariant.^{20,21} As a result, the bilingual brain could access the same semantic geometry through language-specific readout axes. In the motor cortex, rotation of neural activity planes enables the same population to produce diverse reaching movements.²² Motivated by that work, we reasoned that an analogous process may occur in language.

We are particularly interested in the hippocampus, which has a clear role in encoding word meanings.^{23–32} That role in turn is likely part of a broader role in flexibly representing and linking concepts.^{33–39} The hippocampus has often been ignored in language research, in part due to the difficulty of imaging it; however, our laboratory has demonstrated that single neurons in the hippocampus track meanings of words during both listening and monolingual speech production.^{25,27,29,40–42} Together, these findings suggest that hippocampal neurons form a high-dimensional “semantic geometry” in which the distances between neural activity patterns correspond to the relationships between word meanings.

We recorded single-unit activity from the hippocampus in a rare sample of four fully bilingual (English/Spanish) patients with treatment-resistant epilepsy. Using high-density microelectrodes ($n = 3$) and a Neuropixels probe ($n = 1$), we recorded single units across three studies: (1) passive listening to matched stories in English and Spanish (~120 min each), (2) reading and speaking matched phrases in English and Spanish (using the task developed by Silva et al.¹⁹), and (3) unconstrained, naturalistic conversations with English and Spanish conversation partners, respectively. Overall, we find evidence of a shared semantic geometry that is common to both listening and speaking and that is preserved across the three tasks. At the same time, individual neurons’ semantic tuning curves differed substantially between English and Spanish, and, moreover, each language constructs its semantic geometry through distinct readout configurations of the same neurons. Together, these results clarify how bilingual brains represent shared meanings and suggest that the brain uses a readout-based approach to avoiding interference.

RESULTS

We recorded responses of single neurons in the hippocampus of four bilingual individuals undergoing neurosurgical procedures (Figure 1A). All patients acquired both English and Spanish in early childhood (ages 4–5), used both regularly, and were fully balanced bilinguals. Bilingualism was also confirmed informally

through conversations in both English and Spanish with caregivers fluent in each language, as well as by family members. Patients (Pt.) 1–3 underwent stereo-electroencephalography (sEEG) for seizure detection in the epilepsy monitoring unit with an implant strategy that included microelectrodes (AdTech Behnke-Fried) in the hippocampus. Pt. 4 underwent anterior temporal lobectomy for medically refractory epilepsy under general anesthesia. Following resection of the neocortex, we recorded from the exposed hippocampus using a Neuropixels probe prior to the hippocampal resection.

All patients performed a passive listening task with equivalent speech in English and Spanish (Figure 1B) (603 neurons; stimulus set K: 4,530 English words, 4,375 Spanish words; stimulus set A: 4,855 English words, 4,745 Spanish words) (Figure 1C). Two patients (Pt. 2 and 3) also performed a speech production task in both languages (74 neurons total) in which they saw 99 matching phrases in English and Spanish (198 phrases total) on a computer screen and had to say those words out loud (same task and stimuli as Silva et al.¹⁹) (Figure 1D). These two patients also performed a conversation task²⁵ twice, once in each language, although no effort was made to ensure continuity of topic (Figure 1E). Specifically, these two participants took part in conversations in English (durations: 51.37 and 45.5 min) and in Spanish (32.42 and 99.98 min). Word-level, time-aligned transcripts of the listening and speaking content were produced by combining automated transcription with manual correction in Praat⁴³ by trained lab members (Figure 1F; see STAR Methods).

A small number of cross-language neurons

We first asked whether neurons show similar responses to co-translated words across languages. We initially tried cross-language word-by-word alignment using two open-source software packages (GIZA++ and fast_align) but found their alignments to be insufficiently accurate. We therefore performed a manual English-Spanish alignment (Figure 2A). We excluded from analysis words without corresponding meanings and, in many cases, made difficult judgment calls. For example, part of the English text in stimulus set A reads “When you’re travelling in India...” and is translated as “Al viajar por India...” In this context, “When” and “Al” have similar, but not identical, meanings, and were aligned (that is, treated as co-translated pairs). In contrast, “you’re” does not have an equivalent word in the Spanish translation and therefore was excluded from the analysis. Meanwhile, “travelling” and “viajar” have corresponding meanings but different conjugations; given our focus here on semantics over grammar, they were included.

Example responses (Figure 2B) illustrate neurons that show comparable evoked responses to co-translated words across languages. Example units showed significant positive relationships across thousands of matched items, indicating that words that elicit above-average (e.g., “friends”/“amigos”) or below-average (e.g., “and”/“y”) firing rates in English tend to do so in Spanish as well (Figure 2C). To systematically identify putative “cross-language neurons,” we computed a Pearson correlation across all matched words for every neuron (e.g., for stimulus set K, $n = 3,457$ matched words for Pt. 1–3; $n = 1,252$ matched words for Pt. 4) and used a binomial test to determine whether the proportion of significant correlations exceeded chance levels

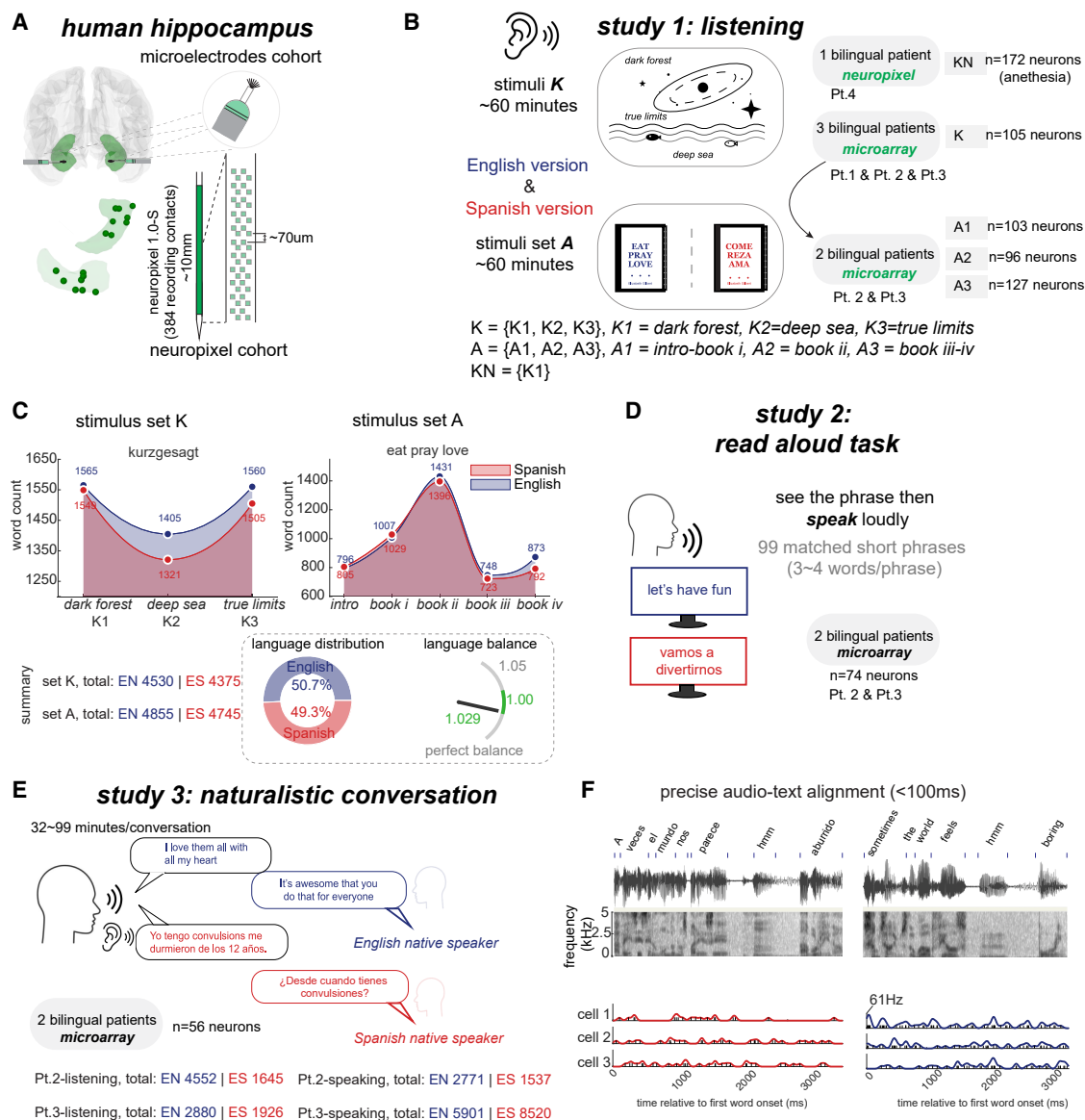


Figure 1. Experimental design for single-neuron recordings during bilingual semantic processing in the human hippocampus

(A) Schematic of recording setup showing bilateral human hippocampus with microelectrode and Neuropixels probe. Green: recording locations.
(B) Study 1: passive listening task involved audio from YouTube videos and an audiobook. Each stimulus set was presented in both English and Spanish versions.
(C) Two stimulus sets used for Spanish-English bilingual participants. Word-count distributions for stimulus sets K and A showing comparable totals across languages.
(D) Study 2: read-aloud task. Two patients viewed and spoke aloud 99 matched short phrases (3~4 words per phrase) in both English and Spanish (198 total phrases).
(E) Study 3: naturalistic conversation task (32~99 min). Two bilingual patients (Pt. 2 and Pt. 3) engaged in unconstrained conversations lasting 60~90 min per session with English and Spanish native speakers.
(F) Precise audio-text alignment. Top: audio waveform of sample English/Spanish texts. Middle: corresponding spectrogram (0~6 kHz) with word boundaries; timestamps are accurate to within 100 ms relative to consensus manual annotations. Bottom: example single-neuron responses from three simultaneously recorded cells showing firing rates. Time axis shows milliseconds relative to first word onset for both English (right, blue) and Spanish (left, red) text. Pt., patient.

(Figure 2D; STAR Methods). A small number of neurons with zero firing rates across all words in either language was excluded ($n = 1$ in dataset A3, listening task; $n = 1$ in read-aloud task).

A subset of neurons met the cross-language neuron criterion for both listening (Figure 2E) and speaking (Figure 2F; see also

Tables S1 and S2). Note that sessions of the listening task (study 1) and speech production (study 2) were collected on different days, so we cannot tell whether the same neurons contribute to speaking and listening. This difference between phrase level and word level likely reflects signal attenuation from averaging

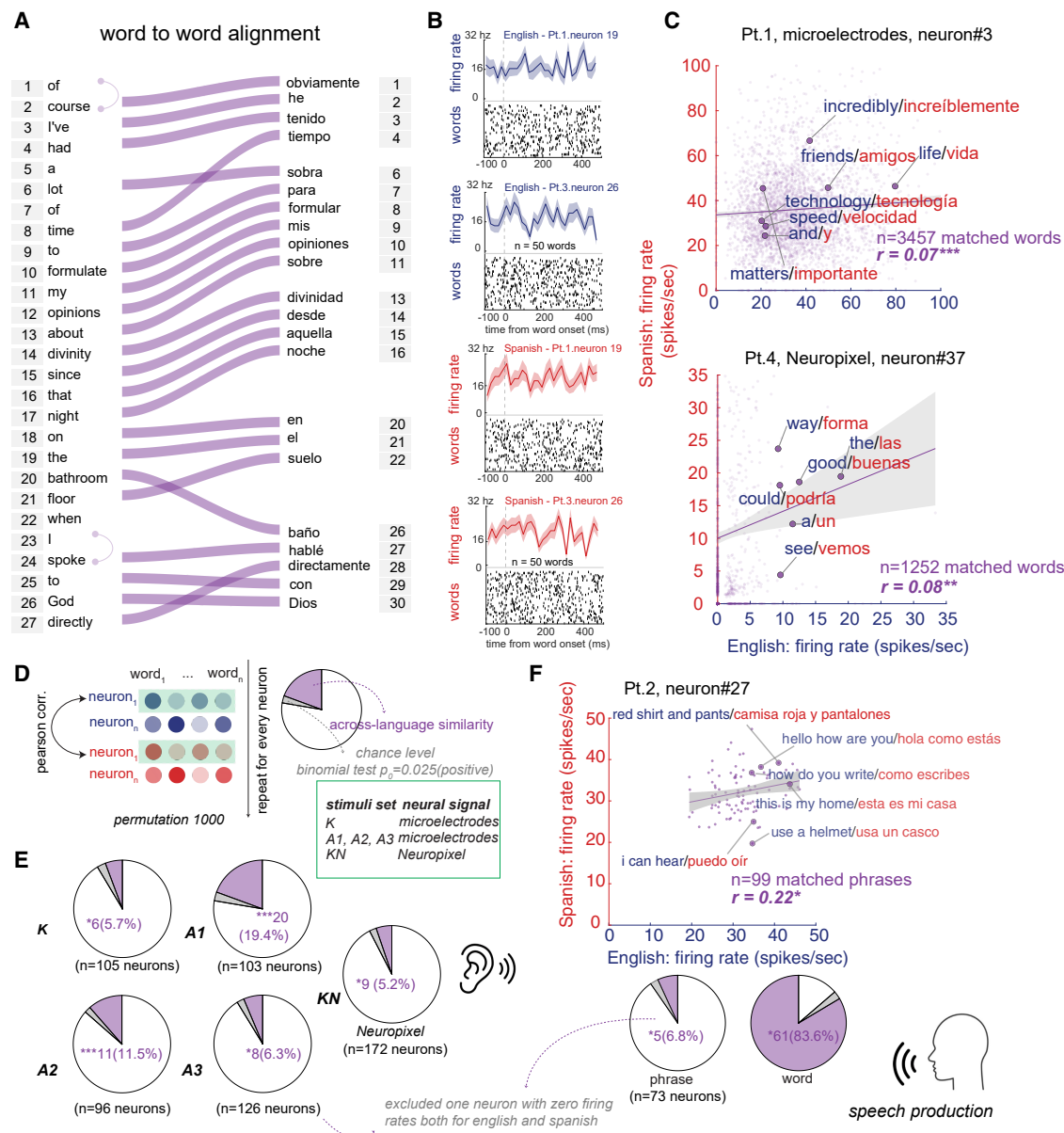


Figure 2. Cross-language single-neuron firing-rate similarity across matched items in passive listening and speech production

(A) Word-to-word alignment mapping between English (left) and Spanish (right) translations. Purple lines connect semantically equivalent words across languages. Numbers indicate word order in each language. Only matched items enter the scatter/correlation analyses below.

(B) Peristimulus time histograms (PSTHs) from example neurons. For each panel, top: average response (firing rate, Hz) to 50 matched words in English (blue traces) and Spanish (red traces); bottom: raster plots below each trace show word-wise spikes.

(C) Pearson correlation analysis of firing rates for matched word pairs. Top: Pt. 1 neuron #3 (microelectrodes) showing strong correlation between English and Spanish firing rates for 3,457 matched words ($r = 0.07$, $p < 0.001$). Bottom: Pt. 4 neuron #37 (Neuropixels) displaying similar language-invariant responses. Each dot is one matched word pair; the x axis represents the unit's firing rate in English, and the y axis represents the unit's firing rate in Spanish. For visualization, several items are annotated (e.g., *friends/amigos* and *life/vida*). Shaded regions: $\pm$ SEM.

(D) For each neuron, we computed the Pearson correlation between firing rates across translation-matched English-Spanish word pairs and assessed significance with a one-tailed permutation test (positive direction; 1,000 shuffles of Spanish indices). We then counted neurons with $r > 0$ and two-tailed permutation $p < 0.05$ and tested whether this fraction exceeds chance using a right-tailed binomial test with ($p_0 = 0.025$).

(E) A set of "cross-language neurons" across all stimulus conditions. Pie charts show the proportion of neurons exhibiting significant cross-language similarity for each stimulus set (purple), with gray portions indicating the expected 2.5% false positive rate under the null hypothesis of no cross-language positive relationships.

(legend continued on next page)

across words within each phrase, as phrase-level effect sizes were substantially smaller than word-level effect sizes (mean $r = 0.03$ vs. 0.29 ; Wilcoxon signed-rank test, $p < 0.001$). All phrase-level cross-language neurons were also identified at the word level, and 67.2% of word-level-only cross-language neurons showed phrase-level correlations above their permutation null (binomial test, $p = 0.006$), confirming that the cross-language signal exists at the phrase level but may be too weak for per-neuron detection. We found essentially no evidence for “anti-cross-language” neurons (i.e., neurons with significant negative correlations between neural responses for matched words across languages). The lack of a *a priori* expectation of equally likely anti-cross-language neurons in all datasets supports the idea that cross-language neurons, while rare, are not simply a statistical artifact. However, the overlap in firing-rate magnitude does not necessarily imply that tuning is identical across languages nor that these neurons play a crucial role in translation.

English-Spanish tuning profiles rarely align at the single-neuron level

Individual neurons respond systematically to word meanings with “semantic tuning functions.”^{25,27,44} To test how semantic tuning functions relate across languages, we first extracted contextual word embeddings using mBERT (Figure 3A; STAR Methods), a transformer model that learns cross-linguistic representations through masked language modeling.⁴⁵

For stimulus set A, we combined excerpts from the audiobook presented across three recording sessions (A1, A2, and A3). Sentences were tokenized using mBERT preserving their original narrative sequence (STAR Methods). For every matched translation pair (e.g., “dog”/“perro”), we calculated the correlation between English and Spanish word embeddings (768-dimensional vectors from mBERT itself, not neural data) at each mBERT layer and computed the average correlation across words. Layer 11 (out of 13) yielded the highest cross-language similarity (Figure 3B). This result is consistent with reports that deeper layers carry semantic information and yield stronger cross-lingual alignment (e.g., Mousi et al.⁴⁶). We therefore used layer 11 embeddings for study 1 and study 3. For study 2 (read-aloud task), words were produced in short phrases without narrative context. We therefore extracted mBERT embeddings for each phrase independently (STAR Methods). Both layer 10 and layer 11 showed the highest cross-language similarity for matched translation pairs (Figure 3B, blue line); we used layer 11 for consistency with studies 1 and 3.

We projected both languages’ embeddings into a common space by concatenating English and Spanish embeddings and applying joint principal-component analysis (PCA), retaining 100 principal components (Figure 3C). This approach ensured that both languages were projected onto identical semantic axes. To validate this shared PCA space, we conducted two control analyses. First, for each shared component, we

computed the fractional contribution of each language to the total variance explained by that component, which yielded nearly balanced contributions (English = 0.509 ; Spanish = 0.491). Second, we confirmed that each principal component derived from the joint PCA showed a strong correlation between English and Spanish scores (average correlation, $r = 0.501$, all PCs, $p < 0.001$), indicating that both languages aligned along the same semantic axes. These results confirm that the shared PCA captured common axes for both languages.

For every neuron, we then fit ridge regressions to map these shared semantic features to firing rates separately for English and Spanish. To avoid overfitting, we selected a single ridge penalty (λ) per neuron via 5-fold cross-validation that minimized the average normalized error across both languages (see STAR Methods). A neuron’s cross-language tuning similarity was defined as the Pearson correlation between its English and Spanish tuning profiles (i.e., the tuning vectors across the shared PCA components from ridge regression). Significance was assessed by a 1,000-fold permutation test (Figure 3D).

Significant correlations between English and Spanish tuning vectors were rare, indicating that neurons do not generally have matching overall tuning in the two languages (Figures 3E and 3F; see also Figure S1). We next asked whether the divergence in tuning curves between English and Spanish was greater than within-language tuning similarity. For each study (study 1, listening; study 2: read aloud; study 3: conversation) separately, we split the English and Spanish words into two halves (independently, using 1,000 random splits), refit ridge weights on each half, and computed within-language split-half tuning curve correlations per neuron. The mean value cross-language tuning correlation fell at the extreme left of the within-language tuning curve correlation distribution (near the 0.3 percentile in study 1 and study 3). Cross-language tuning profile correlations were significantly lower than the within-language tuning profile correlation (Wilcoxon signed-rank test for all stimuli sets in study 1 and study 3; all z values < -2.20 , all $p < 0.02$) (Figures 3G and 3I). For study 2 (speech production) at the word level, however, the cross-language tuning correlation did not significantly differ from within-language tuning correlations (Wilcoxon signed-rank test, $p = 0.06$) (Figure 3H). This result is consistent with the high proportion of cross-language neurons observed at the word level in study 2 (83.6%) (Figure 2F). Cross-language neurons ($n = 115$) showed significantly higher cross-language tuning similarity (mean = 0.14 , SD = 0.11) than other neurons ($n = 560$, mean = -0.01 , SD = 0.15 ; Welch’s t test: $t(204) = 11.83$, $p < 0.001$), suggesting that when words are produced in isolation, hippocampal neurons may exhibit largely language-invariant semantic tuning. Taken together, semantic tuning profiles generally differ across languages, with the notable exception of isolated word production, where cross-language neurons maintain similar tuning across English and Spanish.

(F) A set of “cross-language neurons” during speech production (study 2, read-aloud task). Example neuron from Pt. 2 shows correlated firing patterns when patient was speaking matched English-Spanish phrase pairs. Shaded regions: $\pm$ SEM.

*** $p < 0.001$, ** $p < 0.01$, * $p < 0.05$.

See also Tables S1, S2, and S4.

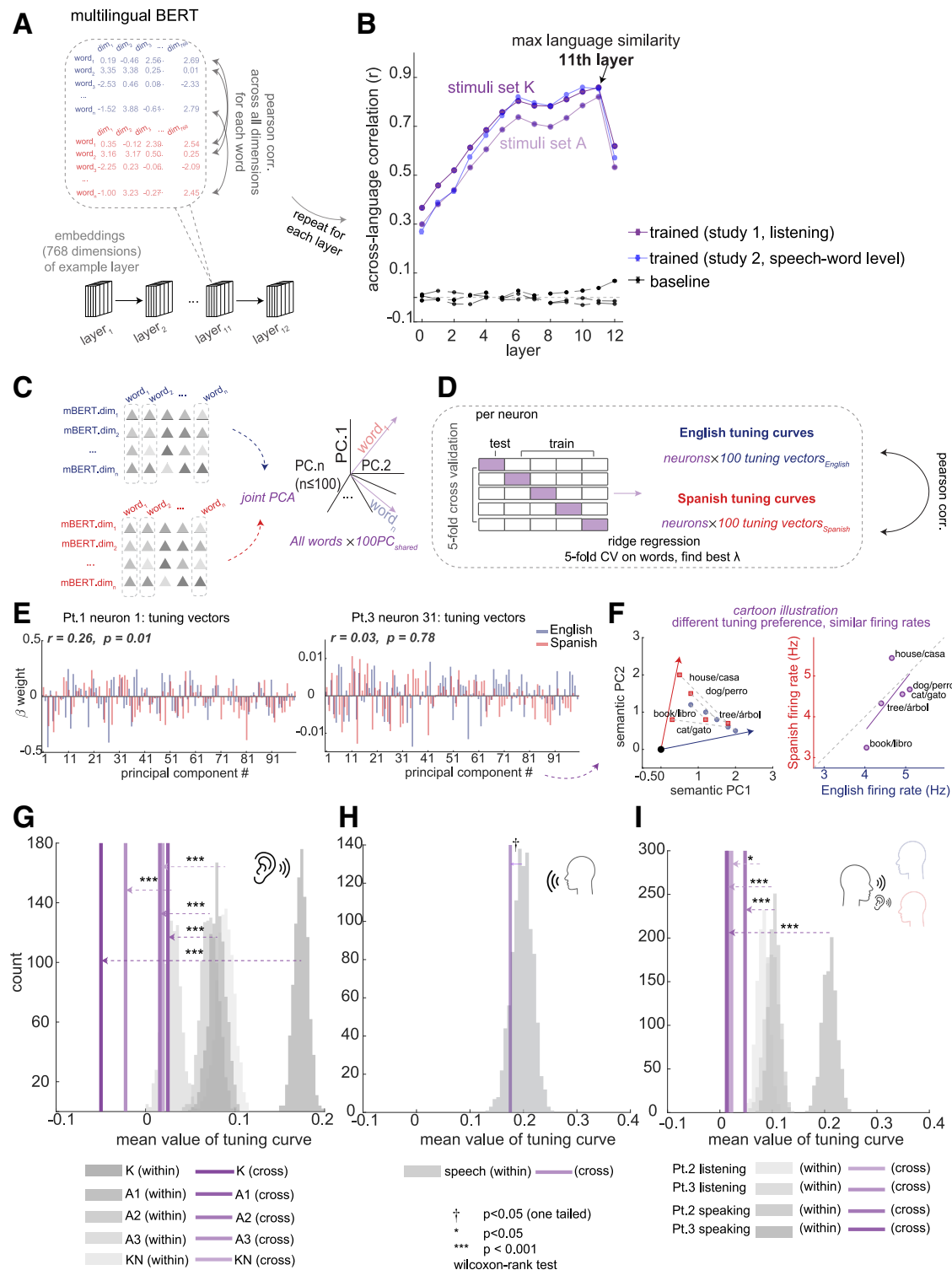


Figure 3. Semantic embeddings and tuning curve similarity at a single-neuron level

(A) Multilingual BERT (mBERT) architecture and contextual embedding extraction pipeline.

(B) Cross-language correlation (purple curves; two stimulus sets: K = *Kurzgesagt* podcasts; A = audiobook; blue curves, words from read-aloud task) rises with depth and peaks at layer 11, which we therefore used in all analyses in passive listening (study 1), speech production (study 2), and conversation (study 3).

(C) Fitting tuning curves in a shared semantic coordinate system. English and Spanish words' embeddings from each dataset are concatenated and projected onto a joint PCA, yielding a shared basis of 100 principal components (PCs).

(legend continued on next page)

Population geometry preserves cross-language semantic structure

We next asked whether “semantic geometry” is shared across English and Spanish. We utilized representational dissimilarity matrices (RDMs), which capture the pairwise distances between word representations and are invariant to axis swaps, rotations, and linear remixing.^{21,47,48} These transformations (e.g., rotation) can change single-neuron tuning functions without altering the population geometry, making RDMs ideal for testing whether English and Spanish maintain the same semantic geometry.

For both languages, we constructed RDMs from populational-level hippocampal neural activity by concatenating word-evoked firing-rate vectors across neurons and computing pairwise cosine distances (“neural semantic maps”) among words. In parallel, we computed mBERT RDMs from mBERT layer 11 embeddings of the same words (“LLM semantic maps”) (Figures 4A–4F; STAR Methods). We asked two questions: (1) how similar is the neural semantic map to the LLM semantic map (brain-model alignment)? (2) How similar is the neural semantic map across English and Spanish within the brain (cross-language neural RDM correlation)?

We used Spearman tests to analyze representational similarity, following best practices.⁴⁹ Brain-model alignment was significant within both languages. Specifically, across datasets in study 1, the neural semantic map significantly correlated with the LLM semantic map (Spearman’s $\rho = 0.132$ – 0.263 , all $p < 0.001$, permutation test) (Figure 4G). That is, word pairs that are distant in the LLM semantic map (e.g., “human” and “galaxy”) are also distant in the neural semantic map, while close word pairs (e.g., “cat” and “dog”) in the LLM semantic map are close in the neural semantic map. These alignments are robust in both English and Spanish. Study 2 (read-aloud task) also showed similar brain-model alignment (Figure 4H) (Spearman’s $\rho = 0.090$ – 0.103 , all $p < 0.01$, permutation test). Study 3 (conversation) showed similar brain-model alignment, despite using different sets of words: correlations were significant during listening (Spanish: Pt. 2, $\rho = 0.179$, Pt. 3, $\rho = 0.128$; English: Pt. 2, $\rho = 0.148$, Pt. 3, $\rho = 0.093$) and speaking (Spanish: Pt. 2, $\rho = 0.197$, Pt. 3, $\rho = 0.136$; English: Pt. 2, $\rho = 0.153$, Pt. 3, $\rho = 0.089$), all $p < 0.001$ (Figure 4I). Together, these results confirm that the hippocampus and mBERT share semantic geometry both for English and Spanish.

Crucially, the neural semantic map was preserved across languages (for example, see Figures 5A and 5B). We directly compared the neural RDMs of the two languages by correlating pairwise neural distances between words in English with the corresponding pairwise neural distances between their Spanish translations across all matched word pairs. We found significant positive correlations for all passive listening datasets ($\rho = 0.255$ – 0.477 , all $p < 0.001$) (Figures 5C and 5D; STAR Methods). Cross-language alignment was also observed in study 2 ($\rho = 0.247$, $p < 0.001$) (Figures 5E and 5F). Permutation controls showed that shuffled Spanish word labels produced null distributions centered near zero, with observed correlations falling far outside these distributions for both passive listening and speech production (all z values > 3.0 , all $p < 0.001$) (Figures 5D and 5F).

We also compared cross-language RDMs derived from mBERT embeddings. Given the strong brain-model alignment observed within each language (Figures 4G and 4H), it is unsurprising that mBERT exhibited similarly high correlations between English and Spanish RDMs ($\rho = 0.651$ for K; $\rho = 0.620$ for KN; $\rho = 0.504$ for A1; $\rho = 0.467$ for A2; $\rho = 0.529$ for A3; all $p < 0.001$).

Local geometric structure enables cross-language prediction

Previous studies have used learned encoding weights between word embeddings and neural responses to predict neural activity in another language.^{6,16,18} However, none has tested whether such prediction is possible at the single-word level. This question is crucial, because we can often infer the meaning of an unfamiliar word from its context, even without explicit translation knowledge, a form of zero-shot learning.^{50,51} For example, in English we know that “dog” and “wolf” are close in meaning, while both are far from “fork.”

To test whether preserved semantic structure could support single-word neural prediction, we used local Procrustes alignment (Figure 5G; STAR Methods). For each target word, we identified its k -nearest neighbors in English neural space (excluding the target itself for held-out prediction aim), learned a rotation matrix from these neighbors’ cross-language response pairs, and predicted the target’s Spanish response. Prediction accuracy was quantified as the angular error between predicted and actual vectors. We optimized anchor size ($k = 15$ – 19 across datasets) (Figure S3A) and found median

(D) For each neuron and each language separately, we fit ridge regression from the 100 PCs of all words to firing rates of all words, selecting λ by 5-fold cross-validation. The resulting beta is the neuron’s semantic tuning curve for that language (for each neuron, there are 100 β [i.e., 100 tuning vectors] for English and 100 for Spanish). Across-language tuning curve similarity of each neuron was defined by the Pearson correlation between English and Spanish tuning curves. (E) Different tuning preferences across languages in example neurons. Bars show β weights across the 100 shared PCs for two representative neurons (left panel from an example “cross-language” neuron; right panel from an example “non-cross-language” neuron). Neuron-wise English-Spanish tuning curve correlations are not significant.

(F) Cartoon illustration shows how different tuning preferences across languages may nevertheless result in similar firing rates. Left: English and Spanish use different readout directions in the semantic PC1-PC2 plane. Right: projections of the same words onto those directions are nevertheless similar, yielding similar firing rates. This visualizes “output-similar does not mean recipe-identical”: firing-rate similarity across languages can arise from different weights of semantic components.

(G–I) Mean value of across-language tuning curve similarity (purple vertical ticks) of all neurons compared with within-language tuning curve similarity in (G) passive listening (study 1), (H) speech production (study 2), and (I) naturalistic conversation (study 3). Histograms show the distribution of within-language tuning curve similarity. Across all tasks, the across-language tuning curve similarities are below the within-language similarity, demonstrating that tuning curves were not well aligned at the single-neuron level.

*** $p < 0.001$, * $p < 0.05$, † $p < 0.05$ (marginally significant if two-tailed, significant if one-tailed).

See also Figure S1.

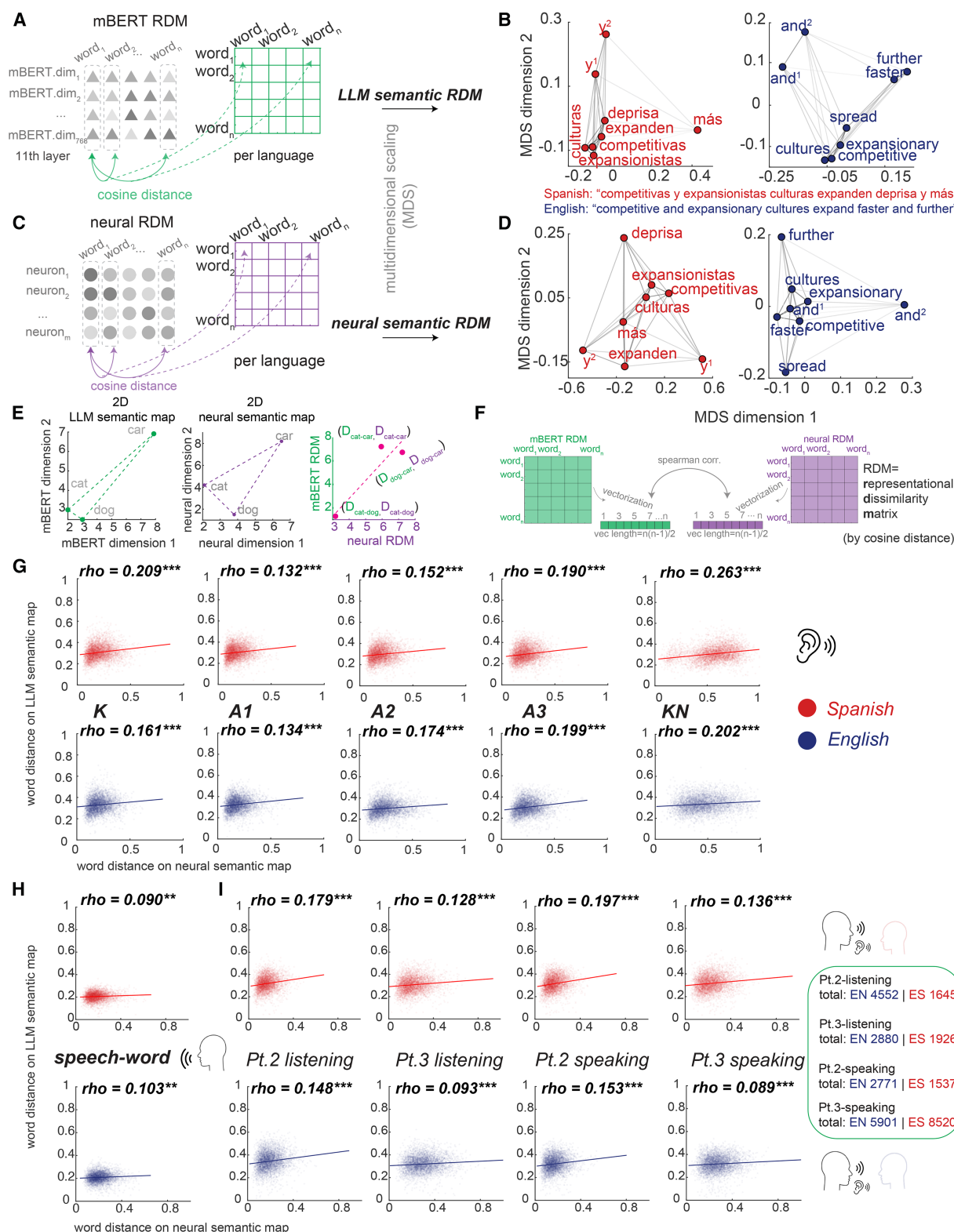


Figure 4. Neuronal population semantic geometry preserved across English and Spanish

(A) Construction of LLM semantic map. Word embeddings from mBERT are used to compute pairwise cosine distance matrices for each language. (B) Multidimensional scaling (MDS) visualization of the LLM semantic map for words from the example sentence. Gray lines represent distances between word pairs. Semantic geometry constructed from mBERT embeddings shows that words with close semantic distances in English maintain similar relationships in Spanish.

(legend continued on next page)

angular errors of 30°–61° for listening and word-speech tasks (see example for stimulus set *K* in Figure 5H; see all results in Figure S3E), significantly exceeding null (all $p < 0.001$). Sliding window analysis showed that even distant neighbors (positions 50–100) predicted nearly as well as the nearest neighbors (<1° decay), indicating a near-global rather than strictly local transformation (Figures S3B and S3C).

We further found that the neighbor-cluster tightness strongly predicted accuracy (Spearman's $\rho = 0.37$ – 0.59) (Figure S3D). Most importantly, higher effective dimensionality of anchor sets, indicating rich structured rather than redundant neighborhoods (Figure 5I), was associated with better prediction (Spearman's $\rho = 0.09$ – 0.31 , all $p < 0.001$) (Figure 5J). This aligns with work showing that high-dimensional neural geometry supports cross-condition generalization.^{35,52} Structured neighborhoods spanning multiple neural dimensions provide richer geometric constraints for estimating the rotation, enabling better generalization to held-out words than redundant low-dimensional clusters. Note that study 3 was not included in cross-language neural comparisons (Figure 5) because conversational content differed between languages, precluding word-by-word matching.

Cross-language neurons are not essential for cross-language neural semantic geometry

We next confirmed that the previously identified cross-language neurons are not responsible for the cross-language geometry match. To test this, we repeated the analyses above excluding the cross-language neurons (that is, the ones that showed similar English-Spanish firing profiles, identified in Figure 2) ($n = 115$). Even after excluding these neurons, we still observed brain-model alignment within languages (all $p < 0.001$) and cross-language neural semantic map correlations (all $p < 0.001$). Critically, the drop in correlation upon removing cross-language neurons did not exceed what would be expected from removing a random subset of equal size (Table S3).

Shared semantic geometry emerges from language-specific semantic axes within a common neural ensemble

Similar pairwise distances can, in principle, be maintained even if the underlying coordinate axes are altered, leaving the geometry

unchanged²¹ (Figure 6A). Moreover, English and Spanish may engage distinct subsets of neurons to construct their respective semantic spaces; or they may recruit the same neuronal population but implement different rotations or re-weightings of the neuronal axes to form language-specific semantic axes. Our results support the latter account.

We applied eigendecomposition to the representational similarity matrices (RSMs, $1 - \text{RDM}$) (Figure 6B) to extract the principal axes (eigenvectors) and quantified cross-language subspace alignment using variance-capture alignment indices (Figure 6C; STAR Methods).⁵³ We term the eigenvectors of the semantic similarity matrix “semantic axes,” representing the principal directions of meaning in population space. These axes can be interpreted as “semantic readout axes” because they specify linearly decodable directions that downstream circuits could access.⁵⁴

Despite the similar semantic geometry, the manifolds constructed from semantic axes for English and Spanish were less aligned than the within-language subspaces both for mBERT (K, A1, A2, A3, KN, speech-word, all $p < 0.001$) and the hippocampus (K, A1, A2, A3, speech-word, all $p < 0.001$; KN-Neuropixel, $p > 0.05$) (Figure 6C). However, the absolute across-language subspace alignment value in mBERT is substantially larger than both within- and across-language subspace alignment values observed in the brain (Figure S4A). The reduction in alignment from within-language to cross-language conditions was comparable between the hippocampus and mBERT (hippocampus = 0.11 ± 0.06 vs. mBERT layer 11 = 0.13 ± 0.07 , $p = 0.44$, Wilcoxon signed-rank test) (Figure 6D). For both systems, the across-language alignments are significantly higher than randomly shuffled distributions (unmatched words) across all stimulus sets (all $p < 0.001$), indicating that English and Spanish implement similar semantic geometry through differently oriented, partially independent semantic axes.

We then examined the neural coordination structure through eigendecomposition of neuron $\times$ neuron covariance matrices, which quantify how neurons co-vary across words within each language (STAR Methods). Each neuron's participation strength (the sum of squared loadings across the top eigenvectors, reflecting how heavily it contributes to these coordination modes) was highly correlated between languages (Figure 6E)

(C) Construction of the neural semantic map. Pairwise neural-distance matrices are computed from population firing-rate vectors (neurons shown as circles) using cosine distance. The resulting matrices capture how the neural population differentiates between words: the more dissimilar two words are, the larger their distance in neural space.

(D) MDS visualization of the neural semantic map for words from the example sentence. Neural population geometry shows that neurons organize words using similar patterns in English and Spanish.

(E) Schematic illustrating the relationship between the LLM semantic map and the neural semantic map analyzed in (G)–(I). If “cat-dog” are close in semantic space, they also tend to be close in neural space, appearing near the diagonal line in the rightmost panel.

(F) Distance-matrix correlation computation. Semantic distance matrices (green, from word embeddings) and neural-distance matrices (purple, from population responses) are vectorized and correlated using Spearman correlation.

(G) Brain-model alignment in passive listening (study 1). Scatterplots show correlations between word distances in LLM-mBERT space (x axis) and neural space (y axis) for datasets K, A1, A2, A3, and KN-Neuropixel. All datasets show significant positive correlations ($r = 0.114$ – 0.237 , all $p < 0.001$).

(H) Brain-model alignment in the read-aloud task (study 2). Scatterplots show correlations between word distances in LLM-mBERT space (x axis) and neural space (y axis) for words in short phrases. The results show significant positive correlations ($\rho = 0.090$ for Spanish and $\rho = 0.103$ for English, all $p < 0.01$).

(I) Brain-model alignment in naturalistic conversation (study 3). Word distances in neural space aligned to word distances in LLM-mBERT during listening and speaking.

*** $p < 0.001$, ** $p < 0.01$.

See also Figure S2 and Table S3.

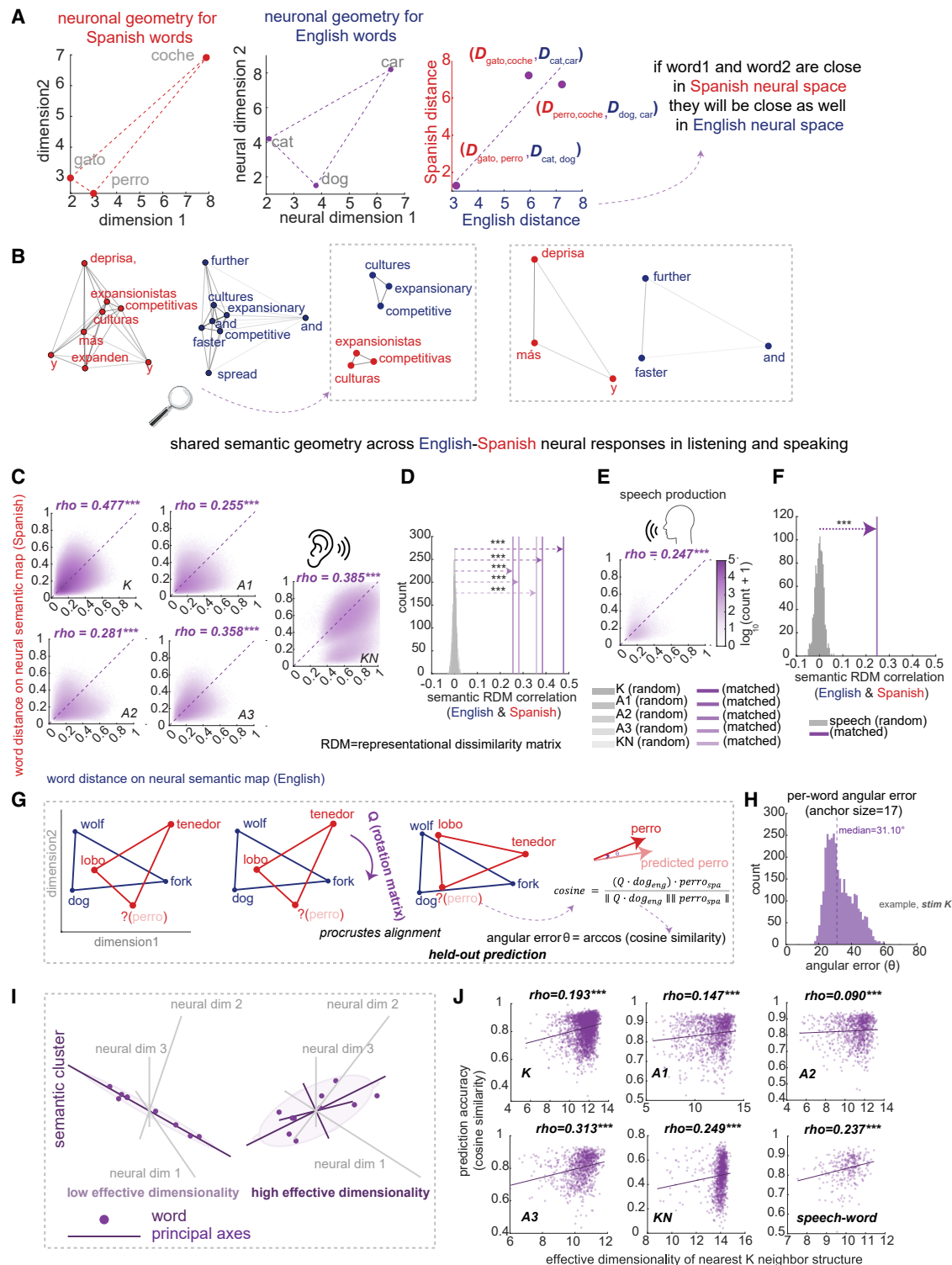


Figure 5. Cross-language alignment of neural population geometry reveals shared semantic geometry

(A) Schematic illustrating cross-language neural-distance preservation. Left: example words in Spanish neural space. Middle: same words in English neural space. Right: cross-language distance-correlation plot where each point represents a word pair.

(B) Zoomed views highlighting semantic clusters within each language from the example sentence in Figure 4. Words with close semantic relationships cluster together in both languages.

(legend continued on next page)

(K: $\rho = 0.87$; A1: $\rho = 0.93$; A2: $\rho = 0.93$; A3: $\rho = 0.88$; KN-Neuropixel: $\rho = 0.68$; speech-word, $\rho = 0.90$, all $p < 0.001$) (for mBERT results, see Figure S4B). These effects are not driven by a small subset of extreme neurons (see control analysis in Figure S5).

DISCUSSION

We examined responses of hippocampal neurons in four bilingual English-Spanish speakers with implanted microelectrodes. We find evidence for a shared cross-linguistic neural embedding space for both listening and speaking for the two languages. That is, words with similar meanings in English and Spanish evoke similar neural response patterns across the two languages, leading to shared semantic geometry. Each neuron had distinct semantic tuning for the two languages, and the shared semantic geometry arose from the different semantic readout axes of two languages. Thus, cross-language shared representation of concepts appears to be a complex emergent phenomenon that takes place not primarily through specialized neurons, or even a translation axis, but rather on the manifold through shared geometry.^{55–57} Together, these results suggest a specific solution to the problem of translation, one that gives a central role to shared coactivation patterns in driving the ability to move between languages.⁵⁸

Although important neuroimaging studies^{6,16,18,59,60} demonstrate that the brain encodes semantic features in a similar way across languages, our single-unit findings take a further step. We quantify the population neuron-level embeddings for each individual word and reveal how every word occupies a specific position within a high-dimensional semantic space. The across-language prediction based on geometric features revealed that the preserved semantic geometry in hippocampal population codes is sufficient to infer the embedding vectors of a hidden word, indicating that shared conceptual geometry is sufficient

for translation and that direct word-to-word mapping is not necessary.

We also consider the relationship between mBERT and the brain. At the population level, the representational geometry of mBERT (layers 9–11) (Figures S2A–S2C) resembles that of the hippocampus. Both represent words as high-dimensional vectors with linearly decodable structure and exhibit cross-language geometric alignment in which translation equivalents occupy similar relative positions.^{15,45,61} Importantly, this alignment is layer selective: brain-mBERT correlations were negative in early layers, became positive around layers 4–6, and peaked in layers 9–11 (Figure S2), indicating that the hippocampus aligns specifically with higher semantic layers rather than mBERT's full hierarchy. Both systems achieve shared geometry through language-specific subspaces constructed from overlapping populations (neurons in the brain, embedding dimensions in mBERT) (Figure S4), though mBERT's cross-language subspace alignment far exceeds that observed in the brain at every layer. At the single-unit level, a fundamental difference emerges: the vast majority of “units” are cross-language correlated, and 100% of units in higher layers of mBERT ($\geq$ layer 6) are cross-language correlated (Table S5). Meanwhile, neurons with similar tuning functions (“cross-language neurons”) are rare in our dataset (5.2%–19.4%, but 83.6% in speech production only). What might account for this difference? One potential reason is a scale mismatch: mBERT similarities are computed over thousands of dimensions that average across many representational factors, whereas a single neuron may reflect a narrow, idiosyncratic subset of features. A second possibility is that mBERT is partly driven by surface-form regularities—shared subword structure, orthography, and distributional cues (for example, “universe” and “universo”), whereas neurons in semantic regions may be more strictly organized around conceptual content. Finally, there may be a deeper representational difference: mBERT encodings are relatively linear and factorized, so related words align in clean

(C) Shared neural semantic geometry between English and Spanish. The two-dimensional histograms show the joint distribution of pairwise neural distances measured in English (x axis) and in Spanish (y axis) for corresponding word pairs. Color intensity represents $\log_{10}(\text{count} + 1)$ of word pairs falling within each bin (60×60 bins, color scale 0–5; STAR Methods). Purple gradient from white (sparse) to dark purple (dense) indicates concentration of data points. The diagonal dashed line represents perfect cross-language correspondence. Passive listening datasets (K, A1, A2, A3, and KN-Neuropixel) show significant correlations.

(D) Permutation-test validation for passive listening. Histogram shows the null distribution of cross-language neural-distance correlations generated by randomly shuffling word-pair assignments (1,000 permutations, gray bars).

(E) Speech production (read-aloud task) shows significant alignment ($\rho = 0.247$, $p < 0.001$). Dense clustering along the diagonal indicates that word pairs maintaining specific distances in English neural space preserve similar distances in Spanish neural space.

(F) Permutation test for speech production (read-aloud task). Similar analysis showing that the observed cross-language correlation (purple vertical line) significantly exceeds the null distribution from random word-pair assignments (gray histogram), validating semantic specificity of neural geometric alignment during speech.

(G) Schematic of local Procrustes prediction method. The aim is to predict the target word (e.g., “perro” in light red) in Spanish given known structure from matched English words. For the translation-equivalent target word in English (e.g., “dog”), we identified its k -nearest semantic neighbors in English neural space (left, for visualization, only 3 words shown; STAR Methods), learned a rotation matrix Q from the cross-language response pairs of these neighbors (middle), and applied Q to predict the target word's Spanish neural response (right). The predicted Spanish response (light red) is compared with the true Spanish response vector (dark red) via cosine similarity and angular error θ ($\arccos(\text{cosine similarity})$).

(H) Per-word angular-error distribution (example from stimulus set K). Histogram of angular error ($\theta = \arccos(\text{cosine similarity})$) between predicted and actual Spanish response vectors at optimal anchor size ($k = 17$) for stimulus set K. Median angular error = 31.10° . See Figure S3E for all datasets.

(I) Effective-dimensionality schematic. Low effective dimensionality (left): anchors cluster (neighborhoods cluster) along a single neural dimension, providing limited geometric constraints. High effective dimensionality (right): anchors span multiple neural dimensions.

(J) Structured neighborhoods predict better. Relationship between effective dimensionality of anchor sets and prediction accuracy across datasets. Higher dimensionality is associated with better prediction (Spearman $\rho = 0.09$ – 0.31 , all $p < 0.001$), indicating that structured neighborhoods outperform redundant ones.

*** $p < 0.001$. RDM, representational dissimilarity matrix (by cosine distance).

See also Figure S3 and Table S3.

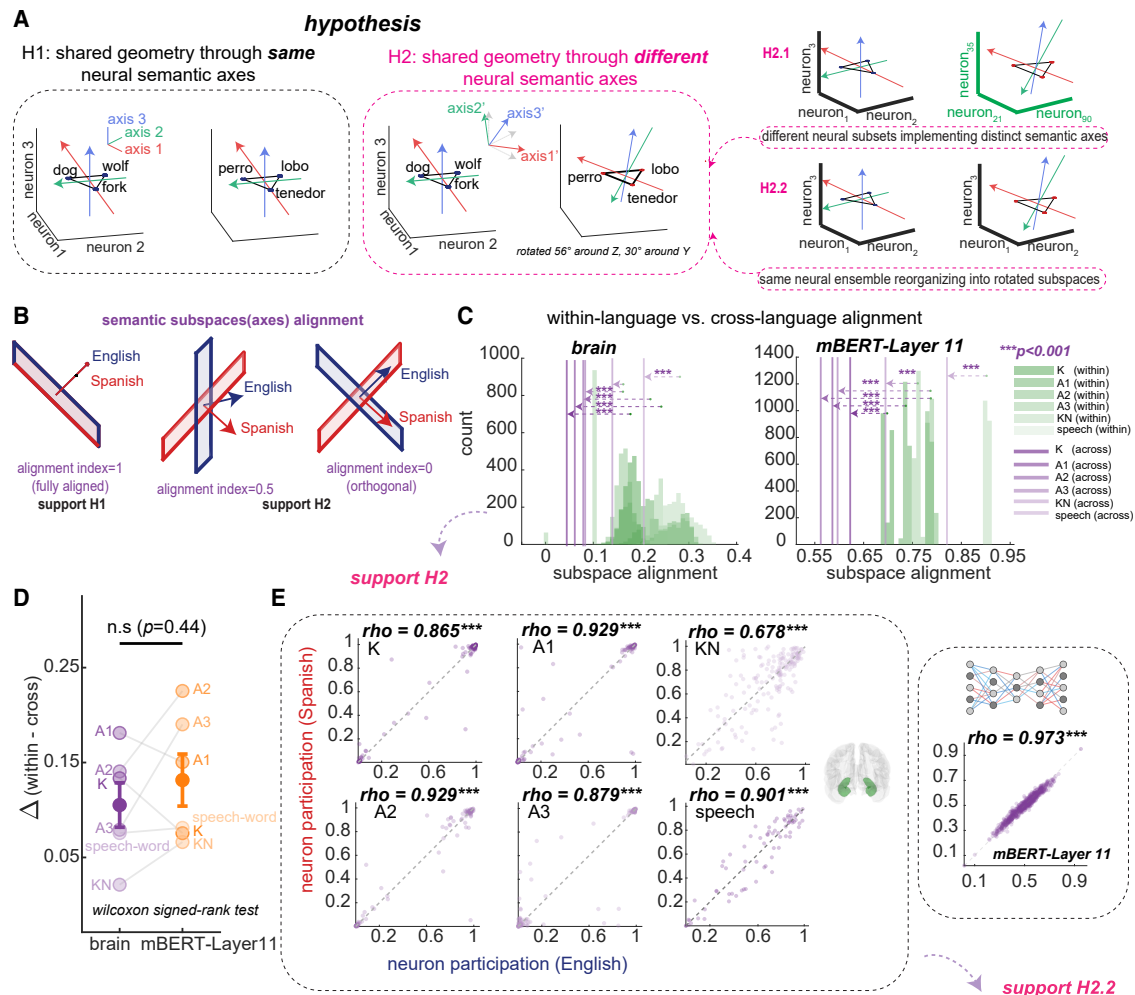


Figure 6. English and Spanish recruit the same group of neurons but different semantic readout axes to support a shared semantic geometry

(A) Conceptual illustration showing that the shared semantic geometry across languages can be preserved in two ways: either by relying on the same semantic axes (left, hypothesis H1) or by altering the underlying semantic axes (right, red and green axes rotated 56° in X-Y plane, blue axes tilted 30°, supporting hypothesis H2, angles are for illustration purposes). If H2 is correct, then there are two possible mechanisms that could give rise to different neural semantic axes across languages: H2.1, English and Spanish may engage different subsets of neurons to construct their semantic spaces, or H2.2, they may recruit the same subset of neurons but use rotations or reweighting of the neuronal axes to form language-specific semantic axes.

(B) Conceptual illustration of possible subspace alignments between languages. Fully aligned (1.0, support H1), partially aligned (0.5, support H2), and orthogonal (alignment = 0, support H2).

(C) Eigendecomposition of the centered RSM (representational similarity matrix, 1 – RDM) yields the principal semantic axes (eigenvectors) for each language. Then we build the subspaces based on the top eigenvectors (which captured 95% of the variance) and test the across-language alignment (STAR Methods). The principal semantic axes for English and Spanish were less aligned compared with within-language subspaces, both for mBERT layer 11 and the hippocampus.

(D) The reduction in alignment from within-language to cross-language conditions was comparable between the hippocampus and mBERT (layer 11).

(E) Per-neuron participation strength in the top subspace (explaining >95% of the variance) was highly correlated between languages both for the hippocampus and mBERT (layer 11), showing that both English and Spanish rely on nearly the same group of neurons to form their semantic readout axes, even though their semantic subspaces differ in direction.

*** $p < 0.001$.

See also Figures S4 and S5 and Table S5.

vector subspaces, whereas neural codes may be more curved, multi-lobed, and context dependent.⁵⁶

Cross-language neurons

We identify a class of what we call “cross-language neurons,” whose responses to co-translated words are correlated. Such neurons would have similar responses to, for example, “maybe”

and “quizás.” These neurons may contribute to translation; however, they do not play an essential or even a major role: even when we exclude them from our analysis, cross-language semantic geometry is robustly maintained (Table S3). Moreover, although their tuning functions are more aligned than other neurons, they are not perfectly aligned, indicating that they are not simply amodal conceptual representations. Consistent with

this, their response fields span multiple semantic categories (e.g., Figure 2C), which differentiates them from the canonical amodal concept-related neurons (concept cells).^{38,62} Word-level decomposition confirms that their cross-language correlations are driven by a relatively sparse subset of words. These words are constrained along fewer latent semantic axes in mBERT space than those for other neurons that did not show cross-language correlation (Tables S1 and S2). We speculate that they may be the extreme end of a larger distribution of neurons whose shared population-level geometry contributes to translation.

How geometry provides simultaneous shared meaning and language-specific readout

Bilinguals need to resolve a control problem that monolinguals do not face: they must not only select a word that matches their intended meaning but also select between available equivalent words in different languages. That is, a bilingual speaker who wants to order a glass of water may choose a different equivalent term in El Paso and Ciudad Juárez. That may feel effortless, but it requires a specialized representational scheme that allows for simultaneous, readout-dependent semantic coding, while also allowing for translatability. In other words, bilingual representations are a manifestation of a broader category of generalization/differentiation problems. The need for simultaneous generalization and differentiation is one that is difficult to solve with single neurons but that is readily solved with population codes.^{35,63–68} We speculate that language context (provided by surrounding words, interlocutor cues, and even general knowledge) could gate which readout axes' rotations are active.⁶⁹

Limitations of the study

One limitation of our study is the low number of patients that we recorded. The main reason for this low number is the rarity of balanced bilingual patients who acquired both languages early in life with access to single units in hippocampus. However, our core finding of shared semantic geometry was replicated across all listening datasets, extended to speech production (i.e., read-aloud task), and was significant in all four patients individually. Second, we did not test additional languages. Although English and Spanish are different, they are genetically related and many words are cognate; it would be valuable to test these ideas on unrelated languages. Third, our data cannot address how the shared semantic geometry is formed. Future studies that track participants as they learn a new language could reveal how semantic geometry emerges and becomes aligned across languages. Likewise, recordings from patients who learned their languages at different life stages could shed light on how existing languages serve as a scaffold for learning new ones.⁷⁰ Finally, some of our data come from an anesthetized patient whose anterior temporal lobe had been resected prior to recording (but the Spanish data from the same patient have never been published or analyzed). Both the anesthesia and the resection may have altered patterns of responding (but see Katlowitz et al.²⁹), although we find largely similar results between this patient and our other three.

RESOURCE AVAILABILITY

Lead contact

Further information and requests for resources and reagents should be directed to and will be fulfilled by the lead contact, Benjamin Y. Hayden (benjamin.hayden@bcm.edu).

Materials availability

This study did not generate new unique reagents.

Data and code availability

The analysis code is available on https://github.com/psywalkeryanxy/bilingual_hippocampal_neurons.

ACKNOWLEDGMENTS

We thank Joshua Adkinson, Justin Fine, Victoria Gates, Raissa Mathura, Suzanne Kemmer, George Kokalas, and Steven Piantadosi for their assistance. This research was supported by the McNair Foundation, the NIH (R01 MH129439, U01 NS121472, and UE5NS070694-15), the SNS Allan Friedman RUNN Research Grant, the NLM Training Program in Biomedical Informatics & Data Science for Predoctoral & Postdoctoral Fellows (T15LM007093-33), the Gordon and Mary Cain Pediatric Neurology Research Foundation, and the National Research Foundation of South Korea (RS-2023-00211018).

AUTHOR CONTRIBUTIONS

Conceptualization, X.Y., B.Y.H., and S.A.S.; investigation, X.Y., A.G.C., M.F., K.A.K., I.G., B.K., A.K., A.S., K.V.A., J.B., A.C., T.I., E.A.M., D.P., H.Z., A.M.G., V.K., A.M., E.B., N.R.P., S.B.M.Y., B.Y.H., and S.A.S.; writing – original draft, X.Y. and B.Y.H.; writing – review and editing, X.Y., B.Y.H., and S.A.S.

DECLARATION OF INTERESTS

S.A.S. has consulting agreements with Boston Scientific, Zimmer Biomet, Koh Young, Abbott, Neuropace, Varian Medical, and Sensoria Therapeutics and is a co-founder of Motif Neurotech.

STAR★METHODS

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

- **KEY RESOURCES TABLE**
- **EXPERIMENTAL MODEL AND STUDY PARTICIPANT DETAILS**
 - Ethics approval
 - Patient recruitment
 - Unit of analysis and sample size
 - Study design
- **METHOD DETAILS**
 - Experimental stimuli in study 1
 - Experimental stimuli in study 2
 - Neuropixels data acquisition setup and intraoperative recordings in the hippocampus
 - Microelectrodes recording in the hippocampus
 - Firing rate responses to words
 - Audio transcription
 - Temporal dynamics of semantic encoding (Study 2)
 - Electrode visualization
 - Methods for cross-language neurons identification
 - Word-level contribution analysis for cross-language neurons
 - Whether there are language-specific neurons
 - Contextual embedding extraction from multilingual BERT
 - mBERT embeddings extraction for short phrase in read aloud task
 - Computing similarity between English-Spanish word embeddings in all layers
 - Single neuron's tuning vector analyses

- Constructing embedding dissimilarity matrix and neural dissimilarity matrix
- Neural-semantic correspondence quantification
- Cross-language semantic map similarity analysis
- Multidimensional Scaling (MDS) network visualization
- Two-dimensional histogram visualization
- Local geometric cross-language prediction
- Cross-language subspace alignment analysis
- Neuron-space coordination structure analysis
- QUANTIFICATION AND STATISTICAL ANALYSIS
- ADDITIONAL RESOURCES

SUPPLEMENTAL INFORMATION

Supplemental information can be found online at <https://doi.org/10.1016/j.cell.2026.05.020>.

Received: November 26, 2025

Revised: April 30, 2026

Accepted: May 14, 2026

REFERENCES

1. Bialystok, E. (2017). The bilingual adaptation: How minds accommodate experience. *Psychol. Bull.* 143, 233–262. <https://doi.org/10.1037/bul0000099>.
2. Evans, N. (2011). *Dying Words, First Edition* (Wiley-Blackwell).
3. Abutalebi, J., and Green, D. (2007). Bilingual language production: The neurocognition of language representation and control. *J. Neurolinguist.* 20, 242–275. <https://doi.org/10.1016/j.neuroling.2006.10.003>.
4. Brignoni-Pérez, E., Jamal, N.I., and Eden, G.F. (2022). Functional neuroanatomy of English word reading in early bilingual and monolingual adults. *Hum. Brain Mapp.* 43, 4310–4325. <https://doi.org/10.1002/hbm.25955>.
5. Cao, F., Tao, R., Liu, L., Perfetti, C.A., and Booth, J.R. (2013). High proficiency in a second language is characterized by greater involvement of the first language network: evidence from Chinese learners of English. *J. Cogn. Neurosci.* 25, 1649–1663. https://doi.org/10.1162/jocn_a_00414.
6. de Varda, A.G., Malik-Moraleda, S., Tuckute, G., and Fedorenko, E. (2025). Multilingual computational models capture a shared meaning component in brain responses across 21 languages. Preprint at bioRxiv. <https://doi.org/10.1101/2025.02.01.636044>.
7. Klein, D., Milner, B., Zatorre, R.J., Meyer, E., and Evans, A.C. (1995). The neural substrates underlying word generation: a bilingual functional-imaging study. *Proc. Natl. Acad. Sci. USA* 92, 2899–2903. <https://doi.org/10.1073/pnas.92.7.2899>.
8. Liu, H., and Cao, F. (2016). L1 and L2 processing in the bilingual brain: A meta-analysis of neuroimaging studies. *Brain Lang.* 159, 60–73. <https://doi.org/10.1016/j.bandl.2016.05.013>.
9. Liu, H., Hu, Z., Guo, T., and Peng, D. (2010). Speaking words in two languages with one brain: neural overlap and dissociation. *Brain Res.* 1316, 75–82. <https://doi.org/10.1016/j.brainres.2009.12.030>.
10. Malik-Moraleda, S., Ayyash, D., Gallée, J., Affourtit, J., Hoffmann, M., Minero, Z., Jouravlev, O., and Fedorenko, E. (2022). An investigation across 45 languages and 12 language families reveals a universal language network. *Nat. Neurosci.* 25, 1014–1019. <https://doi.org/10.1038/s41593-022-01114-5>.
11. Sulpizio, S., Del Maschio, N., Fedeli, D., and Abutalebi, J. (2020). Bilingual language processing: A meta-analysis of functional neuroimaging studies. *Neurosci. Biobehav. Rev.* 108, 834–853. <https://doi.org/10.1016/j.neubiorev.2019.12.014>.
12. Tan, L.H., Spinks, J.A., Feng, C.-M., Siok, W.T., Perfetti, C.A., Xiong, J., Fox, P.T., and Gao, J.-H. (2003). Neural systems of second language reading are shaped by native language. *Hum. Brain Mapp.* 18, 158–166. <https://doi.org/10.1002/hbm.10089>.
13. Malik-Moraleda, S., Taliaferro, M., Shannon, S., Jhingan, N., Swords, S., Peterson, D.J., Frommer, P., Okrand, M., Peterson, J., Cardwell, R., et al. (2025). Constructed languages are processed by the same brain mechanisms as natural languages. *Proc. Natl. Acad. Sci. USA* 122, e2313473122. <https://doi.org/10.1073/pnas.2313473122>.
14. Artetxe, M., and Schwenk, H. (2019). Massively multilingual sentence embeddings for zero-shot cross-lingual transfer and beyond. *Trans. Assoc. Comput. Linguist.* 7, 597–610. https://doi.org/10.1162/tacl_a_00288.
15. Pires, T., Schlinger, E., and Garrette, D. (2019). How multilingual is Multilingual BERT?. *Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics (Association for Computational Linguistics)*, pp. 4996–5001.
16. Chen, C., Gong, X.L., Tseng, C., Klein, D.L., Gallant, J.L., and Deniz, F. (2026). Bilingual language processing relies on shared semantic representations that are modulated by each language. *Proc. Natl. Acad. Sci. USA* 123, e2503721123. <https://doi.org/10.1073/pnas.2503721123>.
17. Dehghani, M., Boghrati, R., Man, K., Hoover, J., Gimbel, S.I., Vaswani, A., Zevin, J.D., Immordino-Yang, M.H., Gordon, A.S., Damasio, A., et al. (2017). Decoding the neural representation of story meanings across languages. *Hum. Brain Mapp.* 38, 6096–6106. <https://doi.org/10.1002/hbm.23814>.
18. Zada, Z., Nastase, S.A., Li, J., and Hasson, U. (2025). Brains and language models converge on a shared conceptual space across different languages. Preprint at arXiv.
19. Silva, A.B., Liu, J.R., Metzger, S.L., Bhaya-Grossman, I., Dougherty, M.E., Seaton, M.P., Littlejohn, K.T., Tu-Chan, A., Ganguly, K., Moses, D.A., et al. (2024). A bilingual speech neuroprosthesis driven by cortical articulatory representations shared between languages. *Nat. Biomed. Eng.* 8, 977–991. <https://doi.org/10.1038/s41551-024-01207-5>.
20. Kriegeskorte, N., and Kievit, R.A. (2013). Representational geometry: integrating cognition, computation, and the brain. *Trends Cogn. Sci.* 17, 401–412. <https://doi.org/10.1016/j.tics.2013.06.007>.
21. Kriegeskorte, N., and Wei, X.-X. (2021). Neural tuning and representational geometry. *Nat. Rev. Neurosci.* 22, 703–718. <https://doi.org/10.1038/s41583-021-00502-3>.
22. Sabatini, D.A., and Kaufman, M.T. (2024). Reach-dependent reorientation of rotational dynamics in motor cortex. *Nat. Commun.* 15, 7007. <https://doi.org/10.1038/s41467-024-51308-7>.
23. Binder, J.R., Desai, R.H., Graves, W.W., and Conant, L.L. (2009). Where is the semantic system? A critical review and meta-analysis of 120 functional neuroimaging studies. *Cereb. Cortex* 19, 2767–2796. <https://doi.org/10.1093/cercor/bhp055>.
24. Kurczek, J., Brown-Schmidt, S., and Duff, M. (2013). Hippocampal contributions to language: evidence of referential processing deficits in amnesia. *J. Exp. Psychol. Gen.* 142, 1346. <https://doi.org/10.1037/a0034026>.
25. Chavez, A.G., Franch, M., Mickiewicz, E.A., Baltazar, W., Belanger, J.L., Devara, D., Etta, M., Hamre, T., Ismail, T., Joiner, B., et al. (2025). Mirror manifolds: partially overlapping neural subspaces for speaking and listening. Preprint at bioRxiv. <https://doi.org/10.1101/2025.09.20.677504>.
26. Duff, M.C., and Brown-Schmidt, S. (2012). The hippocampus and the flexible use and processing of language. *Front. Hum. Neurosci.* 6, 69. <https://doi.org/10.3389/fnhum.2012.00069>.
27. Franch, M., Mickiewicz, E.A., Belanger, J.L., Chericoni, A., Chavez, A.G., Katlowitz, K.A., Mathura, R., Paulo, D., Bartoli, E., Kemmer, S., et al. (2025). A vectorial code for semantics in human hippocampus. Preprint at bioRxiv. <https://doi.org/10.1101/2025.02.21.639601>.
28. Karkowski, K., Kehl, M.S., Qin, Y., Darcher, A., Karkowski, P., Borger, V., Surges, R., Hebart, M.N., Reber, T.P., Dehaene, S., et al. (2025). Semantic Tuning of Single Neurons in the Human Medial Temporal Lobe. Preprint at bioRxiv.

29. Katlowitz, K.A., Cole, E.R., Mickiewicz, E.A., Shah, S., Franch, M.C., Adkinson, J., Belanger, J.L., Mathura, R.K., Meszéna, D., McGinley, M., et al. (2025). Plasticity and Language in the Anesthetized Human Hippocampus. Published online May 6, 2026. *Nature*.
30. Piai, V., Anderson, K.L., Lin, J.J., Dewar, C., Parvizi, J., Dronkers, N.F., and Knight, R.T. (2016). Direct brain recordings reveal hippocampal rhythm underpinnings of language processing. *Proc. Natl. Acad. Sci. USA* **113**, 11366–11371. <https://doi.org/10.1073/pnas.1603312113>.
31. van de Ven, V., Esposito, F., and Christoffels, I.K. (2009). Neural network of speech monitoring overlaps with overt speech production and comprehension networks: A sequential spatial and temporal ICA study. *NeuroImage* **47**, 1982–1991. <https://doi.org/10.1016/j.neuroimage.2009.05.057>.
32. Dijksterhuis, D.E., Self, M.W., Posselt, J.K., Peters, J.C., van Straaten, E.C.W., Idema, S., Baaijen, J.C., van der Salm, S.M.A., Aarnoutse, E.J., van Klink, N.C.E., et al. (2024). Pronouns reactivate conceptual representations in human hippocampal neurons. *Science* **385**, 1478–1484. <https://doi.org/10.1126/science.adr2813>.
33. Bakermans, J.J.W., Warren, J., Whittington, J.C.R., and Behrens, T.E.J. (2025). Constructing future behavior in the hippocampal formation through composition and replay. *Nat. Neurosci.* **28**, 1061–1072. <https://doi.org/10.1038/s41593-025-01908-3>.
34. Bakker, A., Kirwan, C.B., Miller, M., and Stark, C.E.L. (2008). Pattern separation in the human hippocampal CA3 and dentate gyrus. *Science* **319**, 1640–1642. <https://doi.org/10.1126/science.1152882>.
35. Bernardi, S., Benna, M.K., Rigotti, M., Munuera, J., Fusi, S., and Salzman, C.D. (2020). The geometry of abstraction in the hippocampus and prefrontal cortex. *Cell* **183**, 954–967.e21. <https://doi.org/10.1016/j.cell.2020.09.031>.
36. Courellis, H.S., Minxha, J., Cardenas, A.R., Kimmel, D.L., Reed, C.M., Valiante, T.A., Salzman, C.D., Mamelak, A.N., Fusi, S., and Rutishauser, U. (2024). Abstract representations emerge in human hippocampal neurons during inference. *Nature* **632**, 841–849. <https://doi.org/10.1038/s41586-024-07799-x>.
37. Leutgeb, S., Leutgeb, J.K., Barnes, C.A., Moser, E.I., McNaughton, B.L., and Moser, M.-B. (2005). Independent codes for spatial and episodic memory in hippocampal neuronal ensembles. *Science* **309**, 619–623. <https://doi.org/10.1126/science.1114037>.
38. Quiroga, R.Q. (2012). Concept cells: the building blocks of declarative memory functions. *Nat. Rev. Neurosci.* **13**, 587–597. <https://doi.org/10.1038/nrn3251>.
39. Yassa, M.A., and Stark, C.E.L. (2011). Pattern separation in the hippocampus. *Trends Neurosci.* **34**, 515–525. <https://doi.org/10.1016/j.tins.2011.06.006>.
40. Katlowitz, K.A., Belanger, J.L., Ismail, T., Chavez, A.G., Chericoni, A., Franch, M., Mickiewicz, E.A., Mathura, R.K., Paulo, D.L., Bartoli, E., et al. (2025). Attention is all you need (in the brain): semantic contextualization in human hippocampus. Preprint at bioRxiv. <https://doi.org/10.1101/2025.06.23.661103>.
41. Mickiewicz, E.A., Franch, M., Katlowitz, K.A., Chavez, A.G., Zhu, H., Chericoni, A., Yan, X., Belanger, J.L., Ismail, T., Paulo, D.L., et al. (2025). A semantotopic map in human hippocampus. Preprint at bioRxiv. <https://doi.org/10.1101/2025.06.31.685959>.
42. Zhu, H., Franch, M., Mickiewicz, E.A., Belanger, J.L., Cowan, R.L., Katlowitz, K.A., Chavez, A.G., Chericoni, A., Paulo, D., Yan, X., et al. (2026). A geometric foundation for word meaning in the brain. Preprint at bioRxiv. <https://doi.org/10.64898/2026.01.28.702241>.
43. Boersma, P. (2001). Praat, a system for doing phonetics by computer. *Glott Int.* **5**, 341–345.
44. Jamali, M., Grannan, B., Cai, J., Khanna, A.R., Muñoz, W., Caprara, I., Paulk, A.C., Cash, S.S., Fedorenko, E., and Williams, Z.M. (2024). Semantic encoding during language comprehension at single-cell resolution. *Nature* **631**, 610–616. <https://doi.org/10.1038/s41586-024-07643-2>.
45. Devlin, J., Chang, M.-W., Lee, K., and Toutanova, K. (2019). BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding. In *Proceedings of the 2019 Conference of the North, J. Burstein, C. Doran, and T. Solorio, eds. (Association for Computational Linguistics)*, pp. 4171–4186.
46. Mousi, B., Durrani, N., Dalvi, F., Hawasly, M., and Abdelali, A. (2024). Exploring alignment in shared cross-lingual spaces. In *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, L.-W. Ku, A. Martins, and V. Srikumar, eds. (Association for Computational Linguistics), pp. 6326–6348. <https://doi.org/10.18653/v1/2024.acl-long.344>.
47. Diedrichsen, J., and Kriegeskorte, N. (2017). Representational models: A common framework for understanding encoding, pattern-component, and representational-similarity analysis. *PLoS Comput. Biol.* **13**, e1005508. <https://doi.org/10.1371/journal.pcbi.1005508>.
48. Kriegeskorte, N., Goebel, R., and Bandettini, P. (2006). Information-based functional brain mapping. *Proc. Natl. Acad. Sci. USA* **103**, 3863–3868. <https://doi.org/10.1073/pnas.0600244103>.
49. Nili, H., Wingfield, C., Walther, A., Su, L., Marslen-Wilson, W., and Kriegeskorte, N. (2014). A toolbox for representational similarity analysis. *PLoS Comput. Biol.* **10**, e1003553. <https://doi.org/10.1371/journal.pcbi.1003553>.
50. Pourpanah, F., Abdar, M., Luo, Y., Zhou, X., Wang, R., Lim, C.P., Wang, X.-Z., and Wu, Q.M.J. (2023). A review of generalized zero-shot learning methods. *IEEE Trans. Pattern Anal. Mach. Intell.* **45**, 4051–4070.
51. Socher, R., Ganjoo, M., Manning, C.D., and Ng, A. (2013). Zero-shot learning through cross-modal transfer. *Advances in Neural Information Processing Systems* **26**, 935–943.
52. Fusi, S., Miller, E.K., and Rigotti, M. (2016). Why neurons mix: high dimensionality for higher cognition. *Curr. Opin. Neurobiol.* **37**, 66–74. <https://doi.org/10.1016/j.conb.2016.01.010>.
53. Elsayed, G.F., Lara, A.H., Kaufman, M.T., Churchland, M.M., and Cunningham, J.P. (2016). Reorganization between preparatory and movement population responses in motor cortex. *Nat. Commun.* **7**, 13239. <https://doi.org/10.1038/ncomms13239>.
54. Okazawa, G., Hatch, C.E., Mancoo, A., Machens, C.K., and Kiani, R. (2021). Representational geometry of perceptual decisions in the monkey parietal cortex. *Cell* **184**, 3748–3761.e18. <https://doi.org/10.1016/j.cell.2021.05.022>.
55. Ebitz, R.B., and Hayden, B.Y. (2021). The population doctrine in cognitive neuroscience. *Neuron* **109**, 3055–3068. <https://doi.org/10.1016/j.neuron.2021.07.011>.
56. Perich, M.G., Narain, D., and Gallego, J.A. (2025). A neural manifold view of the brain. *Nat. Neurosci.* **28**, 1582–1597. <https://doi.org/10.1038/s41593-025-02031-z>.
57. Vyas, S., Golub, M.D., Sussillo, D., and Shenoy, K.V. (2020). Computation through neural population dynamics. *Annu. Rev. Neurosci.* **43**, 249–275. <https://doi.org/10.1146/annurev-neuro-092619-094115>.
58. Conneau, A., Lample, G., Ranzato, M., Denoyer, L., and Jégou, H. (2017). Word translation without parallel data. Preprint at arXiv. <https://doi.org/10.48550/arXiv.1710.04087>.
59. Correia, J., Formisano, E., Valente, G., Hausfeld, L., Jansma, B., and Bonte, M. (2014). Brain-based translation: fMRI decoding of spoken words in bilinguals reveals language-independent semantic representations in anterior temporal lobe. *J. Neurosci.* **34**, 332–338. <https://doi.org/10.1523/JNEUROSCI.1302-13.2014>.
60. Van de Putte, E., De Baene, W., Brass, M., and Duyck, W. (2017). Neural overlap of L1 and L2 semantic representations in speech: A decoding approach. *Neuroimage* **162**, 106–116. <https://doi.org/10.1016/j.neuroimage.2017.08.082>.
61. Wu, S., Van Durme, B., and Dredze, M. (2022). Zero-shot Cross-lingual Transfer is Under-specified Optimization. In *Proceedings of the 7th*

Workshop on Representation Learning for NLP (Association for Computational Linguistics), pp. 236–248.

62. Quiroga, R.Q., Reddy, L., Kreiman, G., Koch, C., and Fried, I. (2005). Invariant visual representation by single neurons in the human brain. *Nature* 435, 1102–1107. <https://doi.org/10.1038/nature03687>.
63. Barak, O., Rigotti, M., and Fusi, S. (2013). The sparseness of mixed selectivity neurons controls the generalization-discrimination trade-off. *J. Neurosci.* 33, 3844–3856. <https://doi.org/10.1523/JNEUROSCI.2753-12.2013>.
64. Johnston, W.J., Fine, J.M., Yoo, S.B.M., Ebitz, R.B., and Hayden, B.Y. (2024). Semi-orthogonal subspaces for value mediate a binding and generalization trade-off. *Nat. Neurosci.* 27, 2218–2230. <https://doi.org/10.1038/s41593-024-01758-5>.
65. Libby, A., and Buschman, T.J. (2021). Rotational dynamics reduce interference between sensory and memory representations. *Nat. Neurosci.* 24, 715–726. <https://doi.org/10.1038/s41593-021-00821-9>.
66. Tang, C., Herikstad, R., Parthasarathy, A., Libedinsky, C., and Yen, S.-C. (2020). Minimally dependent activity subspaces for working memory and motor preparation in the lateral prefrontal cortex. *eLife* 9, e58154. <https://doi.org/10.7554/eLife.58154>.
67. Tian, Z., Chen, J., Zhang, C., Min, B., Xu, B., and Wang, L. (2024). Mental programming of spatial sequences in working memory in the macaque frontal cortex. *Science* 385, eadp6091. <https://doi.org/10.1126/science.adp6091>.
68. Xie, Y., Hu, P., Li, J., Chen, J., Song, W., Wang, X.-J., Yang, T., Dehaene, S., Tang, S., Min, B., et al. (2022). Geometry of sequence working memory in macaque prefrontal cortex. *Science* 375, 632–639. <https://doi.org/10.1126/science.abm0204>.
69. Rigotti, M., Barak, O., Warden, M.R., Wang, X.-J., Daw, N.D., Miller, E.K., and Fusi, S. (2013). The importance of mixed selectivity in complex cognitive tasks. *Nature* 497, 585–590. <https://doi.org/10.1038/nature12160>.
70. Odlin, T. (2003). Cross-linguistic influence. In *The Handbook of Second Language Acquisition*, C.J. Doughty and M.H. Long, eds. (Wiley Online Library), pp. 436–486. <https://doi.org/10.1002/9780470756492.ch15>.
71. Chaure, F.J., Rey, H.G., and Quiñ Quiroga, R. (2018). A novel and fully automatic spike-sorting implementation with variable number of features. *J. Neurophysiol.* 120, 1859–1871. <https://doi.org/10.1152/jn.00339.2018>.
72. Dale, A.M., Fischl, B., and Sereno, M.I. (1999). Cortical surface-based analysis. *Neuroimage* 9, 179–194. <https://doi.org/10.1006/nimg.1998.0395>.
73. Jenkinson, M., and Smith, S.M. (2001). A global optimisation method for robust affine registration of brain images. *Med. Image Anal.* 5, 143–156. [https://doi.org/10.1016/S1361-8415\(01\)00036-6](https://doi.org/10.1016/S1361-8415(01)00036-6).
74. Jenkinson, M., Bannister, P., Brady, M., and Smith, S. (2002). Improved optimization for the robust and accurate linear registration and motion correction of brain images. *Neuroimage* 17, 825–841. <https://doi.org/10.1006/nimg.2002.1132>.
75. Joshi, A., Scheinost, D., Okuda, H., Belhachemi, D., Murphy, I., Staib, L.H., and Papademetris, X. (2011). Unified framework for development, deployment and robust testing of neuroimaging algorithms. *Neuroinformatics* 9, 69–84. <https://doi.org/10.1007/s12021-010-9092-8>.
76. Groppe, D.M., Bickel, S., Dykstra, A.R., Wang, X., Mégevand, P., Mercier, M.R., Lado, F.A., Mehta, A.D., and Honey, C.J. (2017). iELVis: An open source MATLAB toolbox for localizing and visualizing human intracranial electrode data. *J. Neurosci. Methods* 281, 40–48. <https://doi.org/10.1016/j.jneumeth.2017.01.022>.
77. Magnotti, J.F., Wang, Z., and Beauchamp, M.S. (2020). RAVE: Comprehensive open-source software for reproducible analysis and visualization of intracranial EEG data. *Neuroimage* 223, 117341. <https://doi.org/10.1016/j.neuroimage.2020.117341>.
78. Coughlin, B., Muñoz, W., Kfir, Y., Young, M.J., Meszéna, D., Jamali, M., Caprara, I., Hardstone, R., Khanna, A., Mustroph, M.L., et al. (2023). Modified Neuropixels probes for recording human neurophysiology in the operating room. *Nat. Protoc.* 18, 2927–2953. <https://doi.org/10.1038/s41596-023-00871-2>.
79. Charest, I., Pernet, C.R., Rousselet, G.A., Quiñones, I., Latinus, M., Fillon-Bilodeau, S., Chartrand, J.-P., and Belin, P. (2009). Electrophysiological evidence for an early processing of human voices. *BMC Neurosci.* 10, 127. <https://doi.org/10.1186/1471-2202-10-127>.
80. Li, X., Morgan, P.S., Ashburner, J., Smith, J., and Rorden, C. (2016). The first step for neuroimaging data analysis: DICOM to NIfTI conversion. *J. Neurosci. Methods* 264, 47–56. <https://doi.org/10.1016/j.jneumeth.2016.03.001>.
81. Mead, A. (1992). Review of the development of multidimensional scaling methods. *Statistician* 41, 27. <https://doi.org/10.2307/2348634>.
82. Borg, I., and Groenen, P. (2005). *Modern Multidimensional Scaling: Theory and Applications*. In *Springer Series in Statistics, Second Edition*, 2 (Springer), p. 614.
83. Hofmann, T., Schölkopf, B., and Smola, A.J. (2008). Kernel methods in machine learning. *Ann. Statist.* 36, 1171–1220. <https://doi.org/10.1214/009053607000000677>.

STAR★METHODS

KEY RESOURCES TABLE

REAGENT or RESOURCE	SOURCE	IDENTIFIER
Deposited data		
Analysis code	This paper	Github: https://github.com/psywalkeryanxy/bilingual_hippocampal_neurons
Software and algorithms		
MATLAB	MathWorks	https://www.mathworks.com ; RRID:SCR_001622
Python version 3.11	Python Software Foundation	https://www.python.org ; RRID:SCR_008394
Multilingual BERT (mBERT; bert-base-multilingual-cased)	Devlin et al. ⁴⁵	https://huggingface.co/bert-base-multilingual-cased
HuggingFace Transformers	HuggingFace	https://huggingface.co/docs/transformers
Wave_clus spike sorting	Chaure et al. ⁷¹	https://github.com/csn-le/wave_clus
SpikeGLX 3.0	Bill Karsh, Janelia Research Campus	https://billkarsh.github.io/SpikeGLX/
OpenEphys version 0.6x	Open Ephys	https://open-ephys.org ; RRID:SCR_021624
AssemblyAI (speech transcription)	AssemblyAI	https://www.assemblyai.com
Praat (speech analysis)	Boersma ⁴³	https://www.fon.hum.uva.nl/praat/ ; RRID:SCR_016564
MNE-Python version 0.24.0	MNE Developers	https://mne.tools ; RRID:SCR_005972
FreeSurfer version 7.4.1	Dale et al. ⁷²	https://surfer.nmr.mgh.harvard.edu ; RRID:SCR_001847
FSL (FMRIB Software Library)	Jenkinson and Smith ⁷³ ; Jenkinson et al. ⁷⁴	https://fsl.fmrib.ox.ac.uk ; RRID:SCR_002823
Biolmage Suite version 3.5β1	Joshi et al. ⁷⁵	https://bioimagesuiteweb.github.io/webapp/
iELVis (electrode visualization toolbox)	Groppe et al. ⁷⁶	https://github.com/iELVis/iELVis
RAVE (R Analysis and Visualization of iEEG)	Magnotti et al. ⁷⁷	https://rave.wiki ; RRID:SCR_019040
Other		
Neuropixels 1.0-S probe (384 channels, 70μm width, 10mm length)	IMEC	https://www.neuropixels.org
PXIe Acquisition Module (PXIe-1071)	IMEC / National Instruments	N/A
National Instruments DAQ (PXI 6133)	National Instruments	N/A
Ad-Tech Behnke-Fried sEEG depth electrodes (8 microwires per probe)	Ad-Tech Medical	https://www.adtechmedical.com
Blackrock Microsystems Neuroport system (512-channel, 30 kHz)	Blackrock Microsystems	https://blackrockneurotech.com/products/neuroport/
AlphaOmega microdrive	AlphaOmega	https://www.alphaomega-eng.com
ROSA ONE Brain robotic arm	Zimmer Biomet	https://www.zimmerbiomet.com
Kinevo 900 surgical microscope	Carl Zeiss Meditec	N/A
LabJack U6	LabJack	https://labjack.com
Ethylene oxide sterilization system	Bioseal	N/A

EXPERIMENTAL MODEL AND STUDY PARTICIPANT DETAILS

Ethics approval

The experiment's procedures were conducted by the standards set by the Declaration of Helsinki and approved by the Institutional Review Board for Baylor College of Medicine, Houston, TX (H-50885 and H-18112). Participants provided written informed consent after the experimental procedure had been fully explained and were reminded of their right to withdraw at any time during the study.

Patient recruitment

Four fully balanced Spanish-English bilingual epilepsy patients (daily use of both languages since age 4) participate in the study and were aware that participation was voluntary and would not affect their clinical course. Patients' age ranged from 26–54 years old (mean $\pm$ s.e.m = 38.50 ± 7.08 years). All four participants were female (sex assigned at birth) and identified as women. All four participants self-identified as Hispanic (ethnicity) and as White (race). Ancestry was not separately recorded. None of the patients reported explicit memory of intraoperative events after the case when discussed in the post-operative care unit or while recovering in the hospital the next day.

Unit of analysis and sample size

Because human single-unit electrophysiological recordings during clinically indicated neurosurgical procedures provide an exceptionally rare opportunity to access individual hippocampal neurons in awake bilingual patients, the unit of analysis in this study is the individual neuron rather than the individual participant. Across the four participants we recorded many hundreds of single units, and each unit contributed many hundreds to thousands of word-aligned firing rate observations per language, enabling well-powered population-level inference even with a small participant pool. A formal a priori sample size estimation at the participant level was therefore neither performed nor applicable: the participant count was constrained by the availability of patients meeting our strict inclusion criteria (balanced early-acquired Spanish-English bilingualism, clinically indicated hippocampal recording, and consent for research participation), while statistical power for all primary analyses derives from neuron-level and word-level resampling against permutation-derived null distributions (1000 iterations unless otherwise specified; see [QUANTIFICATION AND STATISTICAL ANALYSIS](#)). We were unable to evaluate the influence of sex or gender on the reported findings. Whether the cross-language semantic geometry reported here generalizes across sexes will require replication in a larger and more diverse cohort, and we note this as a limitation of the present study's generalizability.

Study design

This study used a fully within-subject design and did not divide participants into experimental groups. Each participant served as their own control across languages: the same participants and the same neurons were recorded during both English and Spanish conditions, allowing direct cross-language comparison at the level of single neurons and neural populations.

METHOD DETAILS

Experimental stimuli in study 1

Neural activity was recorded using either Neuropixels probe (P04) or microelectrodes enabling single neuron isolation (P01–P03) ([Figures 1B and 1C](#)). During passive listening tasks (Study1), participants were presented with naturalistic audio stimuli in both Spanish and English. To assess the robustness of our findings across different content domains, we employed two distinct stimulus sets (Set A (audiobook: “eat, pray, love”) and Set K (*kurzgesagt*)) that varied in their linguistic properties and semantic content. The stimuli consisted of four distinct recordings: three educational podcast episodes from *Kurzgesagt* (K: K1–“dark forest”; K2–“true limits”; K3–“deep sea”) and one complete audiobook (A: audiobook “Eat, Pray, Love” by Elizabeth Gilbert). The stimulus presentation varied by participant based on recording availability: P04 (with Neuropixels probe) listened to K1 only (we termed this dataset *KN-Neuropixel*), P01 listened to K1–K3, while P02 and P03 listened to both stimulus sets (K1–K3 and A). Due to clinical care requirements, the audiobook stimulus (A) was recorded across three non-consecutive days for P02 and P03 (Day 1: A1 (books 0–1); Day 2: A2 (book 2); Day 3: A3 (books 3–4)), whereas all podcast stimuli (K1–K3) were collected in single sessions.

Experimental stimuli in study 2

We adapted the bilingual speech neuroprosthesis paradigm from Silva et al.¹⁹ to investigate cross-linguistic neural encoding during active speech production (in our study, we called it as “read aloud task”). The experimental design employed a carefully curated vocabulary of bilingual short phrases presented in an isolated-target task paradigm, where participants attempted to produce individual words in either English or Spanish upon visual cue presentation. The core vocabulary comprised 99 matched English and Spanish short phrases (3–4 words; see Supplementary Table 7 of Silva et al.¹⁹). During this task, participants viewed a single target word displayed on a screen, then the participants attempted to produce the word overtly. The task design incorporated several key features to optimize neural signal quality. First, phrases were presented in randomized order across both languages to prevent anticipatory effects and ensure that neural responses reflected language-specific processing rather than sequential patterns. Second, the visual presentation format remained consistent across all trials, with text displayed in the same font, size, and screen position to minimize visual processing confounds.

Neuropixels data acquisition setup and intraoperative recordings in the hippocampus

Neuropixels 1.0-S probes (IMEC) with 384 recording channels (total recording contacts = 960, usable recording contacts = 384) were used for recordings (dimensions: 70 μ m width, 100 μ m thickness, 10mm length). The Neuropixels probe, consisting of both the recording shank and the headstage, were individually sterilized with ethylene oxide (Bioseal, CA).⁷⁸ Our intraoperative data acquisition system included a custom-built rig including a PXI chassis affixed with an IMEC/Neuropixels PXIe Acquisition module

(PXle-1071) and National Instruments DAQ (PXI466 6133) for acquiring neuronal signals and any other task-relevant analog/digital signals respectively. Our recording rig was certified by the Biomedical Engineering at Baylor St. Luke's Medical Center, where the intraoperative recording experiments were conducted. A high-performance computer (10-core processor) was used for neural data acquisition using open-source software such as SpikeGLX 3.0 and OpenEphys version 0.6x for data acquisition (AP band (spiking data)), band-pass filtered from 0.3kHz to 10kHz was acquired at 30kHz sampling rate; LFP band, band-pass filtered from 0.5Hz to 500Hz, was acquired at 2500Hz sampling rate). We used a "short-map" probe channel configuration for recording, selecting the 384 contacts located along the bottom 1/3 of the recording shank.

Audio was played via a separate computer using pre-generated wav files and captured at 30kHz or 1000kHz on the NIDAQ via a coaxial cable splitter that sent the same signal to speakers adjacent to the patient. MATLAB (MathWorks, Inc.; Natick, MA) in conjunction with a LabJack (LabJack U6; Lakewood, CO) was used to generate a continuous TTL pulse whose width was modulated by the current timestamp and recorded on both the neural and audio datafiles. Online synchronization of the AP and LFP files was performed by the OpenEphys recording software. Offline synchronization of the neural and audio data was performed by calculating a scale and offset factor via a linear regression between the time stamps of the reconstructed TTL pulses and confirmed with visual inspection of the aligned traces.

Acute intraoperative recordings were conducted in brain tissue designated for resection based on purely clinical considerations. The probe was positioned using a ROSA ONE Brain (Zimmer Biomet) robotic arm and lowered into the brain 5–6mm from the ependymal surface using an AlphaOmega microdrive. The penetration was monitored via online visualization of the neuronal data and through direct visualization with the operating microscope (Kinevo 900). Reference and ground signals on the Neuropixels probe were acquired separately by connecting to a sterile microneedle placed in the scalp (separate needles inserted at distinct scalp locations for ground and reference respectively).

Motion analysis

The motion-corrected location estimates were obtained at a 250Hz sampling frequency using the DREDge algorithm. This signal was down sampled to 10Hz. The power spectrum of the calculated motion was then estimated using Welch's overlapped segment averaging estimator for frequencies between 0.1 and 3Hz. The amount of motion was defined as the root mean square error of the location trace of the probes center relative to its average location.

Thresholding method

Spike detection was performed using an adaptive thresholding algorithm designed to achieve a target firing rate of approximately 20 Hz per channel. A threshold was set dynamically for each channel based on signal characteristics and noise level. For each channel noise level was estimated using the mean absolute deviation (MAD), calculated as $MAD = \text{median}(|x|)/0.6745$, where x is the band-passed filtered signal. Channels with MAD values below 1×10^{-6} were classified as dead channels and excluded. An iterative binary search algorithm was used to determine the optimal threshold multiplier for each channel. Different values were tested within a range. Negative-going threshold crossings were considered a neuronal spike. If the iterative firing rate was too high, the threshold was made more negative and vice versa for a low firing rate. A refractory period of 1ms was also used to eliminate some level of noise.

Microelectrodes recording in the hippocampus

The hippocampus was not a seizure focus area in any of the patients included in the study. Single neuron data were recorded from stereo-EEG (sEEG) probes, specifically Ad-Tech Medical probes in a Behnke-Fried configuration. Each patient had an average of three probes terminating in the left and right hippocampus. Electrode locations were verified by co-registered pre-operative MRI and post-operative CT scans. Each probe includes 8 microwires, each with 1 contact, designed explicitly for recording single-neuron activity. Single neuron data were recorded using a 512-channel *Blackrock Microsystems Neuroport* system sampled at 30 kHz. To identify single neuron action potentials, the raw traces were spike sorted using the *Wave_clus* sorting algorithm⁷¹ and then manually evaluated and modified to improve sorting. Noise was removed, and each signal was classified as multi or single unit using several criteria: consistent spike waveforms, waveform shape (slope, amplitude, trough-to-peak), and an exponentially decaying ISI histogram with no ISI shorter than the refractory period (1 ms). The analyses here used all single and multiunit activity.

The initial number of neurons was simply the total number of neurons that we collected. All preprocessing was done by a lab member blind to the goals of the study. In preprocessing (spike sorting), we only kept the isolated neurons. We analyzed all collected neurons, and did not restrict analyses to task-responsive neurons.

Firing rate responses to words

In Study 1 (listening) and Study 3 (conversation-listening), firing rates were computed from spikes occurring 200 ms after word onset to 200 ms after word offset, accounting for transmission delays to hippocampus.^{25,79} In Study 2 (read aloud task) and Study 3 (conversation-speaking), the analysis window was shifted 200 ms earlier (from 200 ms before word onset to 200 ms before word offset) to capture neural activity during speech preparation.²⁵ Spike counts were divided by window duration and multiplied by 1000 to yield spikes per second.

Audio transcription

After experiments, the audio.wav file was automatically transcribed using Assembly AI, a state-of-the-art AI model trained to transcribe speech. Speaker diarization was achieved with Assembly AI's internal speaker diarization model. The transcribed words,

corresponding timestamps, and speaker labels output from these models were converted to a TextGrid and then loaded into Praat, a well-established software for speech analysis (<https://www.fon.hum.uva.nl/praat/>). The original .wav file was also loaded into Praat. Trained lab members manually inspected the spectrograms, timestamps, and speaker labels, correcting each word to ensure that onset and offset times were precise and assigned to the correct speaker (this process took about 4–5 hours per minute of speech). The TextGrid output of corrected words and timestamps from Praat was converted to an Excel file and loaded into Python (version 3.11) for further analysis.

Temporal dynamics of semantic encoding (Study 2)

To characterize the timing of semantic encoding during reading aloud, we performed a sliding window analysis aligned to two events: visual onset (when the phrase appeared on screen) and speech onset (when the participant began speaking). We used 50-ms non-overlapping windows spanning -500 to 1000 ms relative to each phrase and each patient. Following previous research,²⁵ for each window, we extracted spike counts and fit Poisson generalized linear models (GLM) with ridge regularization to predict neural responses from semantic features. Predictors included the first 100 principal components of mBERT embeddings (z-scored), phrase duration (z-scored), and their interactions, yielding 201 features total. Model performance was quantified as Δ LLH, the difference between the actual model log-likelihood and the median log-likelihood from 100 null models in which phrase-neural correspondence was shuffled. The analysis revealed that semantic encoding peaked $\sim$ 200–250 ms after visual onset and $\sim$ 300–400 ms before speech onset, indicating that hippocampal neurons encode semantic content during both retrieval and pre-articulatory preparation. Although our timing analysis (based on Patient 2 and Patient 3) suggested peaks at 300–400 ms before speech onset, we adopted the more conservative 200 ms pre-shift established by Chavez et al.²⁵ using a larger cohort ($n = 10$ patients) to ensure robustness across patients.

Electrode visualization

Electrode visualization for Figure 1A was performed using MNE-Python (version 0.24.0) with FreeSurfer's fsaverage template brain. The visualization pipeline created an ultra-transparent glass brain (3% opacity) with highlighted hippocampal structures to clearly display electrode positions within deep brain regions. For each patient, DICOM images of the preoperative T1 anatomical MRI and the postoperative Stealth CT scans were acquired and converted to NIfTI format.⁸⁰ The CT was aligned to MRI space using FSL.^{73,74} The resulting co-registered CT was loaded into BiImage Suite (version 3.5 β 1),⁷⁵ and the electrode contacts were manually localized. Electrodes' coordinates were converted to the patient's native space using iELVis MATLAB functions⁷⁶ and plotted on the FreeSurfer (version 7.4.1) reconstructed brain surface.⁷² Microelectrode coordinates are taken from the first (deepest) macro contact on the Ad-Tech Behnke Fried depth electrodes. RAVE⁷⁷ was used to transform each patient's brain and electrode coordinates into MNI152 average space. The coordinates were plotted together on a glass brain with the hippocampus segmentation.

Methods for cross-language neurons identification

For each dataset group, we computed Pearson correlations between English and Spanish firing-rate responses across all matched words for each neuron (Figure 2). Neurons with zero variance in either language were excluded from inference and carried as NaN. For every remaining neuron, we assessed significance with a permutation test. Specifically, we generated a null distribution by randomly shuffling the Spanish word indices 1000 times and recomputing the correlation for each shuffle. The per-neuron p -value was the proportion of null correlations whose absolute value was greater than or equal to the absolute observed correlation. Neurons with $p < 0.05$ were classified as significant and then split by sign into significant positive (cross-language similar) and significant negative (cross-language anticorrelated). We set the random seed at the start of the analysis to ensure reproducibility.

At the population level, we summarized each group with two complementary tests. First, to ask whether the number of significant positive (and, separately, significant negative) neurons exceeded what would be expected from the per-neuron test's false-positive rate, we used right-tailed binomial tests with chance probability $p_0 = 0.025$. This value corresponds to the expected false-positive rate from the two-tailed per-neuron permutation test ($\alpha = 0.05$), in which only one tail represents positive correlations. The proportion of significant neurons in each dataset was visualized using pie charts (Figure 2), which displayed the percentage of neurons with significant positive cross-language correlations (purple), the expected proportion by chance (grey), and the remaining non-significant neurons (white).

We performed the same analysis at each layer of mBERT (see Table S5).

Word-level contribution analysis for cross-language neurons

For each cross-language neuron in Study 1, across all stimulus sets, we decomposed the cross-language Pearson correlation into per-word contributions and defined “driver words” as the minimal set accounting for 80% of the positive contribution mass (Tables S1 and S2). We compared across-language neurons to size-matched bootstrap samples of non-across-language neurons (100 iterations per group) on three properties of these driver words: (1) selectivity (the percentage of the driver words), (2) semantic clustering (mean pairwise cosine similarity of driver word mBERT embeddings, tested against permutation null, permutation=1000, we named it as cluster index), and (3) effective dimensionality, the participation ratio (PR) of the driver words' embedding matrix. For each neuron, we extracted the mBERT embeddings (Layer 11) of its driver words, performed PCA via singular value decomposition, and computed $PR = (\sum \lambda_i)^2 / \sum \lambda_i^2$, where λ_i are the eigenvalues. PR ranges from 1, when all variance is concentrated along a single

principal axis, to the number of non-zero eigenvalues, when variance is uniformly distributed across all axes. It thus provides a continuous measure of how many independent semantic dimensions the driver words effectively span. A lower PR indicates that the driver words cluster along fewer embedding dimensions (i.e., occupy a more semantically constrained subspace), whereas a higher PR indicates that they spread across a broader range of semantic directions. To ensure a fair comparison, PR values of cross-language neurons were compared against size-matched bootstrap samples of non-cross-language neurons (100 iterations per group), so that any difference in participation ratio reflects genuine differences in the semantic geometry of driver words rather than differences in sample size.

In the read aloud task, we performed the same driver word analysis as in Study 1 (listening). Because the majority of results were reported at the word level, we did not extend this analysis to the phrase level. Because cross-language neurons vastly outnumbered non-cross-language neurons (61 vs. 12), the bootstrap procedure was inverted: in each of 100 iterations, 12 cross-language neurons were randomly subsampled (without replacement) and compared to the fixed set of 12 non-cross-language neurons.

Whether there are language-specific neurons

To test whether individual hippocampal neurons were driven by the same words across languages or by different words in each language (language-specific or “separate” firing), we performed a word-level overlap analysis (Table S4). For each neuron, firing rates were z-scored within each language across all words. A neuron was considered responsive to a word if its firing rate exceeded a fixed threshold ($z > 1$, that is, $\sim$ top 16% responsive words). This yielded two binary response sets per neuron: English-responsive words and Spanish-responsive words. Responsive words were partitioned into three mutually exclusive categories: words eliciting responses in both languages (shared), words eliciting responses only in English, and words eliciting responses only in Spanish. To assess whether the observed overlap between English- and Spanish-responsive word sets exceeded chance expectations given each neuron’s marginal response rates, we used an exact hypergeometric test (without replacement). For each neuron, the null distribution assumed random overlap between English-responsive and Spanish-responsive word sets, conditioned on the total number of words, the number of English-responsive words, and the number of Spanish-responsive words. We computed the probability of observing an overlap greater than or equal to the empirical overlap under this null. Neurons with $p < 0.05$ were classified as showing overlap enrichment beyond chance. Based on this analysis, neurons were categorized as *pure-shared* (only shared responses), *pure-separate* (only language-specific responses), or *mixed* (both shared and language-specific responses).

Contextual embedding extraction from multilingual BERT

We extracted contextual word embeddings from multilingual BERT (mBERT) for complete English and Spanish versions of the transcripts for both stimulus sets, processing each language’s full stimulus independently before identifying matched word pairs for cross-linguistic analysis (Figure 3). The embedding extraction process was designed to maintain the contextual nature of transformer representations while ensuring that each word received naturalistic linguistic context from its own language stream. For each sentence in the transcripts (282 sentences in set K, 296 sentences in set A), we tokenized the text using mBERT’s tokenizer (bert-base-multilingual-cased) without adding special tokens initially, creating a flat sequential representation of all tokens across the entire stimuli. This approach allowed us to maintain accurate tracking of token positions and their relationships to original words in the transcript.

To preserve contextual information is crucial for mBERT’s contextualized representations, we implemented a sliding window strategy that provided each sentence with up to 509 tokens of preceding context.¹⁸ For each sentence, we calculated the available context window by subtracting the sentence length from the maximum allowable tokens (512 tokens minus 3 special tokens: [CLS], [SEP], [SEP]). The input sequence was then constructed following the format [CLS] + context_tokens + [SEP] + sentence_tokens + [SEP], where context tokens were drawn from immediately preceding text up to the calculated limit. This ensured that each word received rich contextual information from surrounding discourse while staying within mBERT’s architectural constraints.

The model processed each constructed input sequence to generate hidden state representations across all 13 layers (layers 0–12), with each layer capturing increasingly abstract linguistic features. We extracted embeddings specifically for the sentence tokens, excluding the context tokens and special tokens, by indexing into the appropriate positions of the output tensors. For words that were split into multiple subword tokens by the tokenizer, we averaged the embeddings across all constituent tokens to obtain a single word-level representation, ensuring that our analysis operated at the word level despite the subword tokenization scheme. This averaging preserved the dimensional structure (768 dimensions) while consolidating information from word pieces that collectively represent single lexical units.

The extraction process was performed identically for both English and Spanish transcripts, yielding parallel sets of contextualized embeddings where each word in the English version had a corresponding translation-equivalent word in the Spanish version. For each of the three podcast episodes and both languages, this resulted in a three-dimensional tensor of shape ($n_{\text{layers}} \times n_{\text{words}} \times 768$), where $n_{\text{layers}} = 13$ and n_{words} varied by stimuli set length. These layer-wise embeddings captured the hierarchical processing of linguistic information, from surface-level features in early layers to more abstract semantic representations in deeper layers, enabling our subsequent analysis of cross-linguistic alignment patterns. The identical processing pipeline for both languages ensured that any observed differences in representation patterns could be attributed to linguistic rather than methodological factors.

mBERT embeddings extraction for short phrase in read aloud task

We extracted mBERT embeddings using the pretrained *bert-base-multilingual-cased* model from HuggingFace Transformers. Each phrase was tokenized and passed through the model independently (i.e., without cross-phrase context), with input formatted as [CLS] phrase [SEP]. We extracted hidden states from Layer 11 (of 13 transformer layers), as middle-to-late layers have been shown to capture more semantic rather than syntactic information. For phrases containing multiple words, we obtained word-level embeddings by averaging the hidden states of all subword tokens belonging to each word. The resulting embeddings were 768-dimensional vectors for each word.

Computing similarity between English-Spanish word embeddings in all layers

To investigate how multilingual language models represent semantic equivalences across languages, we analyzed the similarity between English and Spanish word embeddings extracted from parallel translated content. We obtained contextual embeddings for matched word pairs across three podcast episodes in both English and Spanish versions, extracting representations from multiple transformer-based models: multilingual BERT (mBERT)^{15,45} with 13 layers (layers 0-12). For each word pair, embeddings were extracted from every layer of these models to capture the hierarchical development of cross-linguistic representations.

For each stimulus set (Set A (audiobook: “eat, pray, love”) and Set K (*kurzgesagt*)), we computed Pearson correlations between the embedding vectors of each English-Spanish word pair (e.g., “human-humano”), yielding a correlation coefficient for every matched word across all layers. This approach quantified how similarly the models represent semantically equivalent words in different languages at various stages of processing depth.

To control for spurious similarities that might arise from random initialization or architectural properties rather than learned linguistic knowledge, we established a baseline using untrained versions of the same models with identical architectures but randomized weights. These baseline models underwent the same embedding extraction and correlation computation procedures as their trained counterparts, maintaining near-zero correlations across all layers as expected from random representations.

We aggregated correlations across all word pairs by computing means and standard errors for each layer, with error propagation applied when calculating the relative similarities to account for uncertainty in both trained and baseline measurements properly. The resulting layer-wise profiles revealed how cross-linguistic semantic alignment emerges and evolves through the depth of each model, with both stimulus sets showing progressive increases in correlation through the middle layers before reaching maximum alignment at Layer 11 (Figure 3B).

Single neuron’s tuning vector analyses

For each dataset (K, A1, A2, A3, KN-Neuropixel, read aloud task, conversation-listening, conversation-speaking), we concatenated single-unit firing-rate matrices across all words to form word $\times$ neuron matrices separately for English and Spanish (Figure 3). We extracted mBERT token embeddings for the matched words in each language (Layer 11, see above), then concatenated English and Spanish embeddings and ran a single PCA on the combined matrix to obtain a shared embedding basis. The resulting PCA scores were split back into English and Spanish, yielding a common set of semantic predictors with identical axes across languages (100 components).

Then, for every neuron, we fit two ridge regressions that map shared PCA embeddings to firing rates, one in English and one in Spanish. The ridge penalty was selected per neuron by cross-validation, using a grid of penalties and 5-fold CV jointly across languages (the chosen penalty minimized the average normalized prediction error over English and Spanish folds). With the selected penalty, we refit the English and Spanish ridge models and extracted, for each neuron, we got a *tuning vector*, that is, the vector of regression coefficients on the shared PCA semantic features. Each tuning vector (length = number of PCA components) specifies how strongly that neuron weights each semantic dimension to predict its firing rate.

We quantified neuron-wise similarity by correlating, for each neuron, its English tuning vector with its Spanish tuning vector (over PCA dimensions). To test whether English-Spanish similarity exceeded chance, we generated a null distribution by randomly permuting the Spanish word indices (1000 permutations, this procedure is same as we did in Figure 2), refitting the Spanish ridge model for every neuron with the same cross-validated penalty, and recomputing neuron-wise correlations. Per-metric p-values were computed as the proportion of null correlations whose absolute value was greater than or equal to the observed value (two-sided $\alpha = 0.05$ test). Same as the method we used for Figure 2 (identify the cross-language neurons), we used right-tailed binomial tests with chance probability $p_0 = 0.025$ (corresponding to one tail of the per-neuron two-sided $\alpha = 0.05$ test).

To evaluate whether the observed cross-language tuning similarity could be explained by within-language variability, we quantified the split-half reliability of tuning vectors for each neuron within English and within Spanish. For every iteration, word indices were randomly split into two equal halves, and independent ridge models were fit on each half using the neuron’s cross-validated penalty. The two resulting tuning vectors were correlated to obtain a within-language reliability value per neuron. This procedure was repeated 1000 times, generating a distribution of split-half correlations for each neuron and, when pooled across all neurons and both languages, a global distribution of within-language reliability. The vertical red line in Figure 3 (last panel) marks the mean cross-language tuning similarity across all neurons, shown relative to this within-language reliability distribution. We ran a paired Wilcoxon signed-rank test across neurons, testing whether per-neuron mean within-language reliability exceeded cross-language similarity.

Constructing embedding dissimilarity matrix and neural dissimilarity matrix

For both neural and semantic spaces, we computed *Representational Dissimilarity Matrices* (RDMs)^{20,47,49} using cosine distance as the dissimilarity metric (Figure 4). The cosine distance between two vectors $\mathbf{x}$ and $\mathbf{y}$ is defined as:

$$d_{\text{cosine}}(\mathbf{x}, \mathbf{y}) = 1 - \frac{\mathbf{x} \cdot \mathbf{y}}{\|\mathbf{x}\| \|\mathbf{y}\|}$$

For a dataset with n words, this yields a condensed distance vector of length $n(n-1)/2$ containing all unique pairwise distances. The neural RDM for language l was computed as:

$$\mathbf{D}_{\text{neural}}^{(l)} = \left\{ d_{\text{cosine}}(\mathbf{n}_i^{(l)}, \mathbf{n}_j^{(l)}) : 1 \leq i < j \leq n \right\}$$

Where $\mathbf{n}_i^{(l)}$ represents the neural population vector for word i in language l . Similarly, the semantic RDM was computed as:

$$\mathbf{D}_{\text{semantic}}^{(l)} = \left\{ d_{\text{cosine}}(\mathbf{s}_i^{(l)}, \mathbf{s}_j^{(l)}) : 1 \leq i < j \leq n \right\}$$

Where $\mathbf{s}_i^{(l)}$ represents the mBERT embedding vector for word i in language l .

Neural-semantic correspondence quantification

The correspondence between neural and semantic representational geometries was quantified by computing Spearman correlation between the vectorized RDMs (i.e., $\mathbf{D}_{\text{semantic}}^{(l)}$ and $\mathbf{D}_{\text{neural}}^{(l)}$). This correlation coefficient quantifies the degree to which neural dissimilarity patterns mirror semantic dissimilarity patterns. A high positive correlation indicates that words with similar meanings (small semantic distance) tend to evoke similar neural responses (small neural distance), supporting the hypothesis of semantic encoding in the neural population.

Cross-language semantic map similarity analysis

We performed a cross-language similarity analysis comparing neural response patterns for matched English-Spanish word pairs to investigate the correspondence between neural representations of semantic content across languages (Figure 5). This analysis quantifies whether words that elicit similar neural responses in one language also elicit similar responses in another, providing evidence for shared or distinct neural geometry for semantic representation across languages.

Neural similarity between word representations was quantified using cosine distance, which measures the angular separation between neural population vectors independent of their magnitude. For each language, we computed pairwise cosine distances between all word pairs, this yielded similarity matrices $\mathbf{S}^{(E)}$ and $\mathbf{S}^{(S)}$ for English and Spanish, respectively, with values ranging from 0 (identical patterns) to 1 (maximally dissimilar patterns).

To assess the preservation of similarity structure across languages, we extracted the upper triangular portions of both similarity matrices (excluding the diagonal), yielding vectors of unique pairwise similarities. Then we use the Spearman correlation⁴⁹ to quantify the cross-language correspondence. A high positive correlation indicates that word pairs with similar neural representations in English tend to have similar representations in Spanish, suggesting a shared neural geometry for organizing semantic content across languages.

To evaluate whether the observed cross-language correlation exceeded chance, we generated a null distribution by randomly shuffling the Spanish word labels 1000 times and recomputing the correlation after each shuffle. This procedure destroys the correspondence between matched English-Spanish word pairs while preserving each language's internal structure. The resulting distribution of shuffled correlations represents the range of values expected if there were no true relationship across languages. The p -value was then calculated as the proportion of shuffled correlations whose absolute value was equal to or greater than the observed correlation (two-sided test, $\alpha = 0.05$).

Multidimensional Scaling (MDS) network visualization

In Figures 4B and 4D, we employed classical Multidimensional Scaling (MDS)⁸¹ combined with network visualization techniques to visualize the geometric organization of neural representations for matched word pairs across languages. This approach reveals the topological structure of semantic representations in neural space by projecting high-dimensional neural patterns onto interpretable low-dimensional manifolds while preserving pairwise dissimilarities. The resulting 2D coordinates were visualized as network graphs where nodes represent words and weighted edges encode distance between words.

Two-dimensional histogram visualization

We constructed a two-dimensional histogram with 60×60 bins spanning the range $[0, 1]$ for both dimensions to visualize the joint distribution of similarity values across languages (Figures 5C and 5E). The bin counts were log-transformed using:

$$N_{\text{display}} = \log_{10}(N + 1)$$

Where N represents the raw count in each bin, this transformation enhances visibility of the full density range while preventing singularities from empty bins.

Local geometric cross-language prediction

To test whether preserved semantic geometry could support single-word neural decoding across languages, we implemented a local Procrustes alignment approach (Figures 5G–5J). All neural firing rate vectors were L2-normalized to unit length, which ensures that prediction accuracy reflects directional alignment in neural state space rather than differences in firing rate magnitude.

Anchor size optimization

We first determined the optimal number of anchor words by systematically varying K from 3 to 100. For each target word, we identified its k nearest semantic neighbors in English neural space (excluding the target itself), learned a local Procrustes rotation from these anchor pairs, and evaluated prediction accuracy on the held-out target.

Held-out prediction

For each target word with English response vector $\mathbf{e}_{\text{target}}$, we identified its nearest K semantic neighbors (K was the optimal anchor size from above analysis step) in English neural space using cosine distance (explicitly excluding the target word itself), then extracted these neighbors' responses in both languages to form matrices $\mathbf{X}_E$ (English, $n_{\text{neurons}} \times k$) and $\mathbf{X}_S$ (Spanish, $n_{\text{neurons}} \times k$). After centering both matrices by subtracting column means (μ_E and μ_S), we computed the cross-covariance matrix $\mathbf{M} = (\mathbf{X}_S - \mu_S)(\mathbf{X}_E - \mu_E)^T$ and performed singular value decomposition $\mathbf{M} = \mathbf{U}\Sigma\mathbf{V}^T$. The optimal orthogonal rotation matrix was obtained as $\mathbf{Q} = \mathbf{U}\mathbf{V}^T$, and the predicted Spanish response was computed as $\hat{\mathbf{y}} = \mu_S + \mathbf{Q}(\mathbf{e}_{\text{target}} - \mu_E)$, then L2-normalized to unit length. Prediction accuracy was quantified as cosine similarity and angular error $\theta = \arccos(\hat{\mathbf{y}} \cdot \mathbf{y}_{\text{true}})$ in degrees ($\arccos$ of cosine similarity).

Statistical significance

We tested predictions for all words in each dataset (datasets with matched English-Spanish words, including all stimulus sets in Study 1 and words from short phrases in Study 2). Statistical significance was assessed against 1000 null permutations where we (1) randomly shuffled English-Spanish word pairings using derangement (ensuring no correct translations), (2) randomly selected Spanish prediction targets, and (3) used randomly selected words (excluding the target) rather than true nearest neighbors as anchors for learning rotation matrix $\mathbf{Q}$, thereby destroying both semantic correspondence and neighborhood structure while preserving neural response distributions. P-values were computed as the proportion of null means achieving equal or greater accuracy than the observed mean.

Sliding window analysis

To test whether the cross-linguistic rotation is local or global, we applied a sliding window approach. Instead of using the fixed nearest neighbors (positions 1 to K , where K represents the optimal anchor size), we shifted the window to use progressively more distant neighbors (for example, if $K=17$, then we have positions 2–18, 3–19, ..., up to positions 84–100). If the rotation were strictly local, prediction accuracy should decay sharply beyond the immediate neighborhood.

Per-word trajectory analysis

To determine whether accuracy decay was gradual or sharp at the individual word level, we computed sliding window trajectories for each word separately. For each word, we calculated: (1) the overall slope of accuracy versus window position, and (2) the maximum single-step drop across adjacent windows. Sharp drops were defined as single-step decreases exceeding 2 standard deviations below the mean derivative across all words and windows.

Local geometry features

To examine what properties of local neighborhoods predict accuracy, we computed two metrics for each target word at optimal K (neighbors at positions 1 to k , excluding the target itself): (1) neighbor-cluster tightness, defined as the mean pairwise neural cosine similarity among the K neighbors. (2) The effective dimensionality of each anchor set using the participation ratio from PCA on the centered anchor matrix. The effective dimensionality reflects the neural geometric structure of the local neighborhood: low values indicate redundant anchors that span a low-dimensional neural subspace, while high values indicate structured anchors that span multiple dimensions. This metric allowed us to test the hypothesis that highly structured neighborhoods might predict better than redundant ones.

Cross-language subspace alignment analysis

To examine how English and Spanish neural population representations relate to each other in geometric structure, we performed a cross-language subspace alignment analysis based on representational dissimilarity matrix (RDM, by cosine distance) and eigendecomposition (Figures 6C and 6D).

To enable eigendecomposition, we first converted cosine distances to a similarity matrix via 1 minus the RDM (where the subtraction is elementwise) to get the RSM (representation similarity matrix). We then applied double-centering (kernel centering) to remove the mean structure and obtain a centered Gram matrix^{82,83} directly in MATLAB (Statistics and Machine Learning Toolbox).

We then performed the eigendecomposition. Eigenvalues were sorted in descending order, and the corresponding eigenvectors were arranged accordingly. The eigenvectors define the principal semantic axes (or we called the “semantic readout axes” in results sections) in the high-dimensional word space along which the neural population responses vary most strongly across words. That is, they are the directions that maximize variance in the centered word-word similarity structure.

To quantify representational dimensionality, we used the cumulative variance threshold, defined as the smallest number of components capturing at least 95% of the total variance. The top- k eigenvectors from each language formed orthonormal bases $\mathbf{U}_E$ and $\mathbf{U}_S$ representing the English and Spanish *semantic subspaces*, respectively. To quantify their geometric relationship, we adopted the alignment index from Elsayed et al.⁵³ We computed the variance-capture alignment index, which measures how much variance in one dataset's representational structure is captured by another dataset's subspace, normalized by the variance that the dataset could optimally capture on its own. The directional alignment is defined as:

$$\text{Alignment}(A \leftarrow B) = \frac{\text{Tr}(\mathbf{U}_B^T \mathbf{G}_A \mathbf{U}_B)}{\sum_{i=1}^k \lambda_A^{(i)}}$$

where $\mathbf{G}_A$ denotes the centered cosine Gram matrix for language A, which is mathematically equivalent to the second-moment (covariance-like) structure of neural activity in the cosine similarity space. $\text{Tr}(\cdot)$ is the matrix trace. Conceptually, $\mathbf{G}_A$ plays the same role as a covariance matrix, it captures how population responses co-vary across stimuli, which is computed from cosine similarities. $\mathbf{U}_B$ is the matrix of source language B's top- k eigenvectors, and denominator is the sum of language A's top- k eigenvalues. This metric ranges from 0 (no overlap) to 1 (perfect alignment), measuring what fraction of A's top- k variance can be explained by projecting onto B's top- k subspace. We computed bidirectional alignments $\text{Alignment}(\text{Eng} \leftarrow \text{Spa})$ and $\text{Alignment}(\text{Spa} \leftarrow \text{Eng})$, and reported their mean as a symmetric measure of overall subspace overlap between English and Spanish.

$$\text{Alignment}_{\text{sym}} = \frac{1}{2} [\text{Alignment}(\text{Eng} \leftarrow \text{Spa}) + \text{Alignment}(\text{Spa} \leftarrow \text{Eng})]$$

Permutation test

To test whether cross-language alignment exceeded chance, we performed 1000 permutation. In each permutation, the word order of the Spanish Gram matrix was randomly shuffled, disrupting the one-to-one correspondence between matched English-Spanish word pairs while preserving each language's internal structure. For each shuffled instance, we recomputed the symmetric alignment index, producing a null distribution of alignment scores expected by chance. The empirical alignment was compared to this null distribution, and a permutation p -value was then calculated as the proportion of shuffled alignment whose absolute value was equal to or greater than the observed alignment (two-sided test, $\alpha = 0.05$)

Compare to within-language subspace alignment

To assess the reliability of the across-language alignment, we computed within-language subspace alignment as a control. Because different subsets of words can produce distinct cosine-distance structures, splitting the words would alter the geometry of the representational dissimilarity matrix and thus would not reflect the reproducibility of the underlying subspace. Instead, we randomly divided the neurons into two equal halves (50% each) and repeated the procedure 1000 times. For each half, we recomputed the centered cosine Gram matrix, performed eigendecomposition, and re-estimated subspace alignment using the same alignment index calculation. To statistically evaluate whether the across-language alignment was significantly lower than within-language alignment, we compared the observed cross-language alignment value against the distribution of within-language alignments pooled across English and Spanish halves. We verified that the gap between cross-language and within-language alignment is stable when both are evaluated at the same fixed subspace dimensionality, swept from 5 to the dimensionality needed to explain 95% of the variance of full-neuron data (whichever language required fewer). At every matched dimensionality, cross-language alignment remained significantly below the within-language split-half ceiling (except for the Neuropixels dataset, which also showed no significant cross-language alignment in the main text analyses). This confirms that the comparison is not confounded by differences in effective dimensionality between split-half within-language and full-neuron cross-language analyses. We computed a z -score as the deviation of the observed value from the within-language mean, normalized by the within-language standard deviation. Significance thresholds followed standard criteria: $|z| > 1.64$ ($p < 0.10$), $|z| > 1.96$ ($p < 0.05$), and $|z| > 2.58$ ($p < 0.01$).

Neuron-space coordination structure analysis

To determine whether English and Spanish engage the same neural ensemble with similar coordination patterns, we analyzed the covariance structure across neurons for each language (Figure 6E). This analysis addresses whether cross-language semantic similarity arises from different neural subsets forming language-specific representations, or the same neurons reorganizing through different linear readouts.

Constructing neuron by neuron covariance matrix

For each language, we computed how pairs of neurons covary in their firing rates across all words. We first centered each neuron's responses by subtracting its mean firing rate across words, then computed the covariance between all neuron pairs. This neuron-by-neuron covariance matrix quantifies patterns of coordinated neural activity: neurons with high positive covariance tend to increase their firing rates together across words, while neurons with negative covariance show opposite response patterns. We then performed eigendecomposition of each neuron covariance matrix to identify the dominant modes of coordinated neural activity. Each eigenvector represents a pattern of how neurons covary together, that is, neurons with high loadings on the same eigenvector tend to increase or decrease their firing rates together across words. The eigenvalues reflect the variance explained by each

coordination mode, sorted from largest to smallest. We retained eigenvectors capturing greater than 95% of total variance. The number of eigenvectors required to reach this threshold was determined separately for each language. For cross-language comparisons, we used the minimum number of eigenvectors needed across both languages to ensure comparable dimensionality. We also quantified effective dimensionality using the participation ratio, which measures how many coordination modes contribute meaningfully to neural activity patterns. Then we quantified whether English and Spanish recruit the same patterns of neural coordination using the same alignment index for semantic subspaces alignment analyses.

Per-neuron participation strength

Also, to quantify each neuron's contribution to the main coordination subspace, we computed participation strength as the sum of squared loadings across all retained eigenvectors. This metric reflects how strongly each neuron participates in the dominant coordination patterns. We computed the Spearman correlation between English and Spanish participation strengths to test whether the same neurons serve as the similar role of coordination in both languages.

Statistical validation via neuron-identity permutation

To test whether cross-language neural coordination alignment exceeds chance expectations, we performed 1000 permutations by randomly shuffling neuron identities between languages. This permutation test addresses the null hypothesis that neuron correspondences between English and Spanish are arbitrary. For each permutation, we randomly reordered the neurons in the Spanish data, breaking the neuron-to-neuron correspondence between languages while preserving each language's word-evoked activity patterns. Then we recomputed the Spanish neuron covariance matrix from the permuted data, performed eigendecomposition and extracted the top coordination modes. Finally, we computed the alignment index and participation strength correlation under permutation. The permutation p -value was then calculated as the proportion of shuffled alignment whose absolute value was equal to or greater than the observed alignment (two-sided test, $\alpha = 0.05$). If the observed values significantly exceeding the permutation distribution indicate that English and Spanish specifically recruit the same neurons with matched coordination patterns.

QUANTIFICATION AND STATISTICAL ANALYSIS

All statistical analyses were performed in MATLAB R2024b (MathWorks; Statistics and Machine Learning Toolbox) and Python 3.11 (NumPy, SciPy, scikit-learn). The statistical tests, sample sizes, and significance thresholds for each individual analysis are described in the corresponding [method details](#) subsection and reported in the relevant figure legends and main text. This section provides only the conventions that apply across the manuscript.

Definition of n . The unit of analysis is the individual neuron (single units) for all neuron-level analyses; the matched English-Spanish word pair for all RDM, multidimensional scaling, subspace alignment, and Procrustes prediction analyses; the phrase or trial for read-aloud temporal dynamics analyses; and the participant ($n = 4$) for participant-level demographics only. The exact value of n for each analysis is reported in the corresponding figure legend.

Center, dispersion, and precision. Group-level statistics are reported as mean $\pm$ SEM unless otherwise noted; bootstrap confidence intervals (95%, 100-1000 iterations as specified) are reported for effect-size estimates in [Tables S1](#) and [S2](#). Error bars and shaded bands in figures denote SEM unless otherwise noted in the figure legend.

Inference framework and assumption checking. Inference relies on permutation or bootstrap procedures (typically 1000 iterations) rather than parametric tests, so normality and equal-variance assumptions are not required and were not formally tested. Where Pearson correlations were used (e.g., single-neuron cross-language similarity), results were verified for consistency with their non-parametric Spearman counterparts. Permutations were applied at the appropriate level for each test (word identities, neuron identities, or word-pair derangement, as specified per analysis), preserving each language's internal structure while disrupting only the cross-language correspondence. Random seeds were fixed at the start of each permutation analysis to ensure reproducibility. Significance threshold was $\alpha = 0.05$ (two-tailed) unless otherwise noted; p -value notation in tables and figures: $*p < 0.05$, $**p < 0.01$, $***p < 0.001$.

Software and code availability. Specific software packages and versions are listed in the [key resources table](#). Other quantitative and statistical methods are described above within the context of individual analyses in the [method details](#) section.

ADDITIONAL RESOURCES

This study did not generate additional online resources beyond the analysis code repository listed in the [key resources table](#). This study is not a registered clinical trial.

Supplemental information

**Shared neural geometries for bilingual semantic
representations in human hippocampal neurons**

Xinyuan Yan, Ana G. Chavez, Melissa Franch, Kalman A. Katlowitz, Ivy Gautam, Brian Kim, Aaditya Krishna, Aadit Shrivastava, Katie Van Arsdell, James Belanger, Assia Chericoni, Taha Ismail, Elizabeth A. Mickiewicz, Danika Paulo, Hanlin Zhu, Alica M. Goldman, Vaishnav Krishnan, Atul Maheshwari, Eleonora Bartoli, Nicole R. Provenza, Seng Bum Michael Yoo, Benjamin Y. Hayden, and Sameer A. Sheth

Supplemental figures

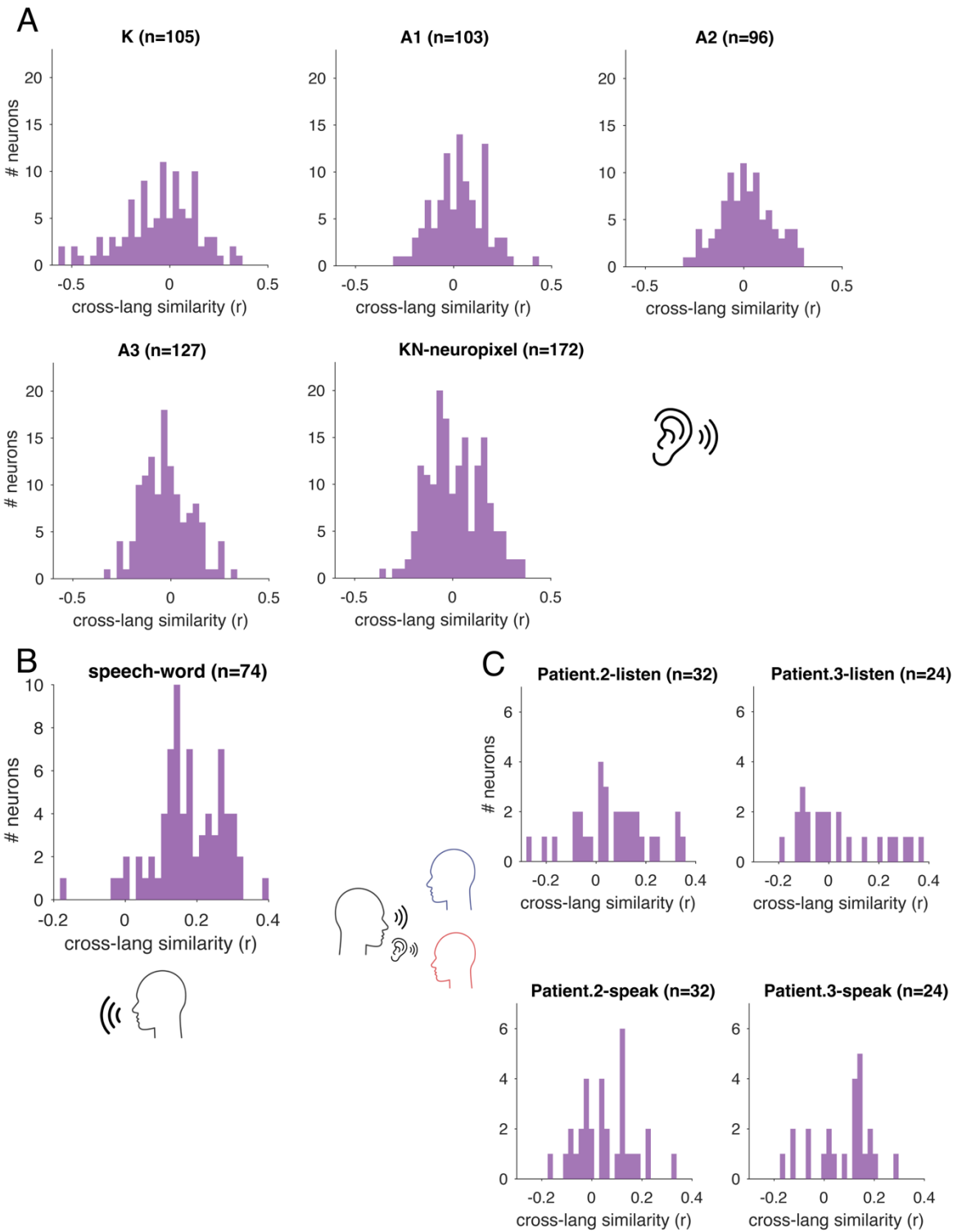


Figure S1. The distribution of cross-language tuning similarity (Pearson correlation between English and Spanish tuning vectors) across all neurons for each study, related to Figure 3.

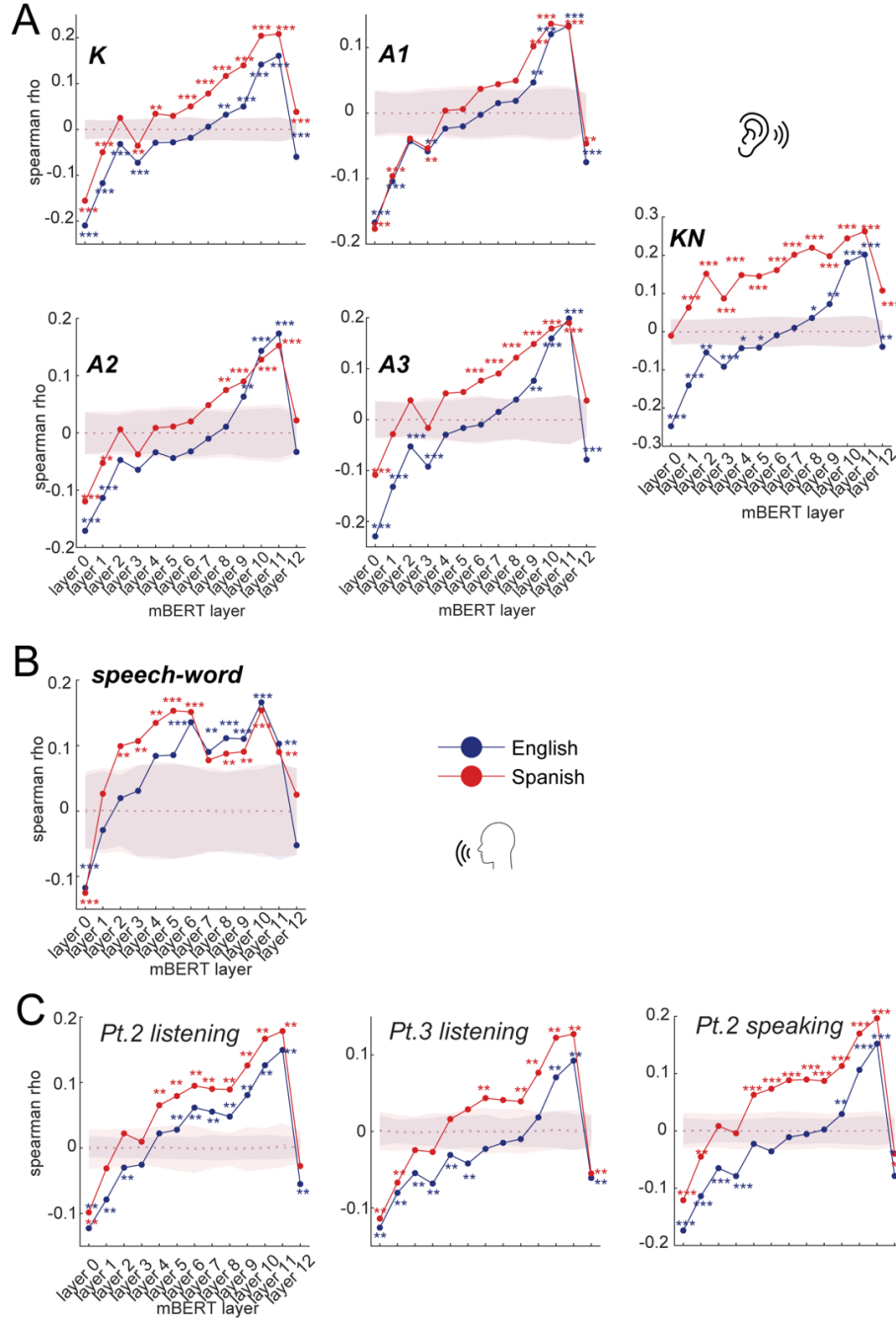


Figure S2. Brain-mBERT alignment at each mBERT layer, related to Figure 4.

We repeated the brain-mBERT alignment analyses (Figure 4) across all 13 mBERT layers (layers 0-12) for each stimulus set. Across all five stimulus sets (K, A1, A2, A3, and KN), speech production (word-level) as well as the conversation task, the correlation between neural and semantic RDMs was negative in layers 0-1, progressively shifted toward positive values around layers 4-6, and peaked with statistical significance in layers 9-11.

* $p < .05$, ** $p < .01$, *** $p < .001$, permutation test (1000 iterations)

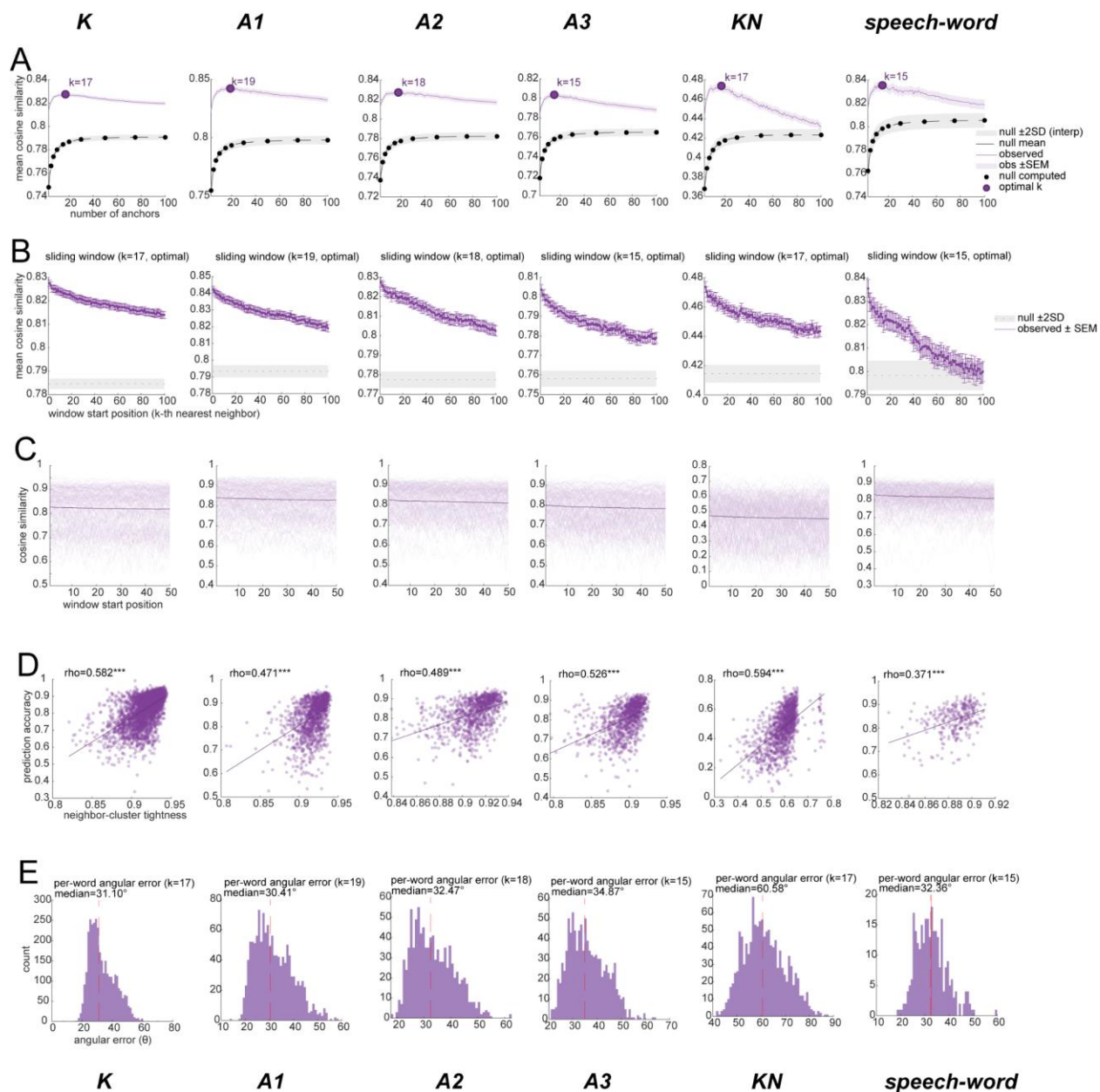


Figure S3. Extended results for local geometric structure enables cross-language prediction, related to Figure 5.

(A) Optimal anchor size. Mean cosine similarity between predicted and actual Spanish neural response vectors as a function of anchor set size ($k = 3-100$). Purple line shows observed prediction accuracy ($\pm SEM$); black dashed line shows null distribution mean with ± 2 SD band (gray shading). Black dots indicate anchor sizes where null was explicitly computed (1000 permutations). Large purple dot marks optimal k for each dataset. Prediction accuracy peaks at $k = 15-19$ depending on dataset and remains stable across anchor sizes. The high null baseline reflects shared global covariance between the two neural spaces; the consistent observed advantage above null reflects the additional precision gained from true semantic correspondence.

(B) Sliding window analysis. Mean cosine similarity as a function of window start position at optimal k . Position 1 uses neighbors ranked 1- k ; position n uses neighbors ranked n to $n+k-1$.

30 Accuracy decays by less than 1° across 100 window positions and remains far above null
 31 throughout, indicating the cross-linguistic transformation approximates a near-global rotation
 32 rather than a strictly local one.
 33 (C) Per-word sliding window trajectories. Individual word trajectories across window positions.
 34 Near-horizontal profiles indicate gradual decay; absence of clustered sharp drops suggests no
 35 discrete neighborhood boundaries.
 36 (D) Neighborhood tightness predicts accuracy. Relationship between neighbor-cluster tightness
 37 (mean pairwise neural cosine similarity among k anchors) and prediction accuracy for each
 38 word. Strong positive correlations (Spearman $\rho = 0.37\text{-}0.59$, all $p < 0.001$) indicate that words
 39 in tighter, more coherent neighborhoods are predicted more accurately.
 40 (E) Per-word angular error distribution. Histograms of angular error ($\theta = \arccos(\text{cosine}$
 41 $\text{similarity})$) at optimal k (i.e., anchor size in panel A).
 42 Columns: K = Kurzgesagt (3457 words); A1, A2, A3 = audio books sets 1-3; KN = Kurzgesagt-
 43 Neuropixel; speech-word = read-aloud task.
 44 $*p < .05$, $**p < .01$, $***p < .001$, permutation (1000 iterations)
 45

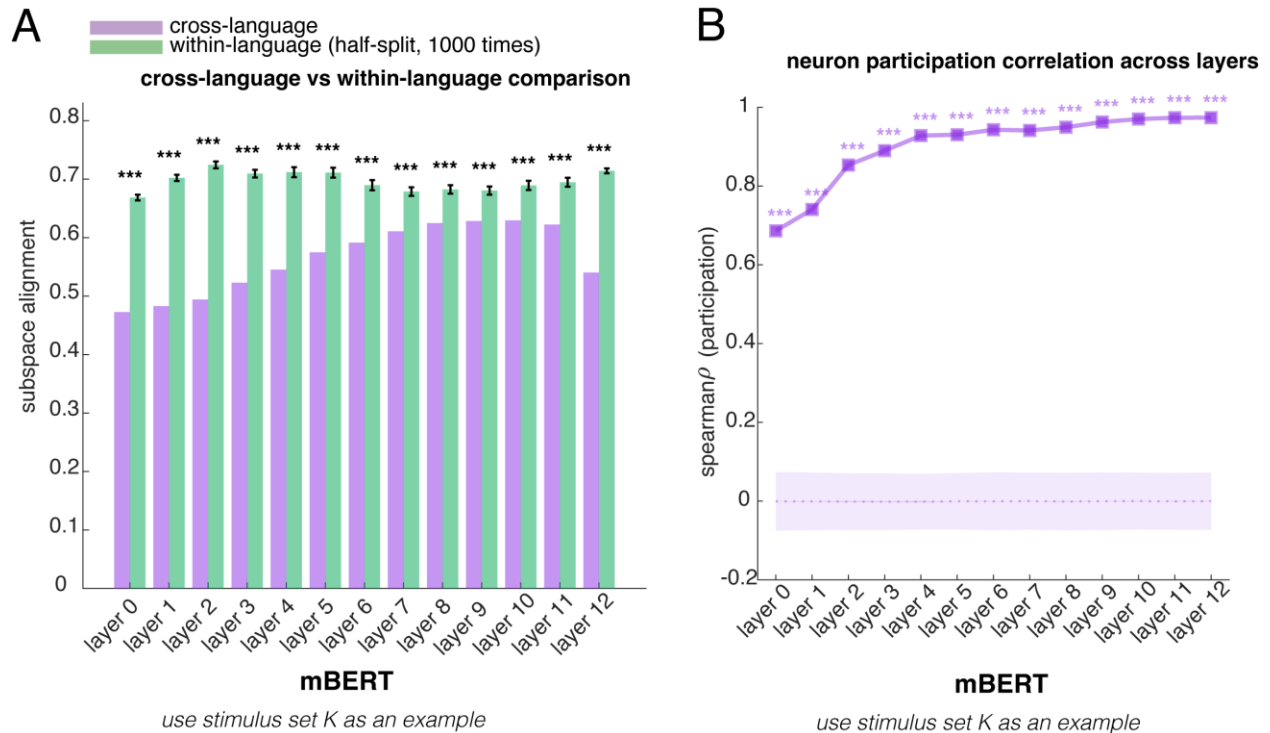


Figure S4. Cross-language subspace alignment and “neuron” participation in subspace construction at every layer of mBERT, related to Figure 6.

(A) Cross-language subspace alignment (using stimulus set K as an example) is significantly smaller than within-language half-split baselines at every layer of mBERT.

(B) If we treated the embedding dimension of mBERT as a neuron, the per-artificial neuron participation strength in the top subspace (explaining > 95 % variance) was highly correlated between languages (all $\rho > 0.650$, all $p < 0.001$) at every layer of mBERT, showing that mBERT utilized the same “hardware” to construct language-specific semantic readout axes to achieve the shared meaning.

* $p < .05$, ** $p < .01$, *** $p < .001$, permutation test (1000 iterations)

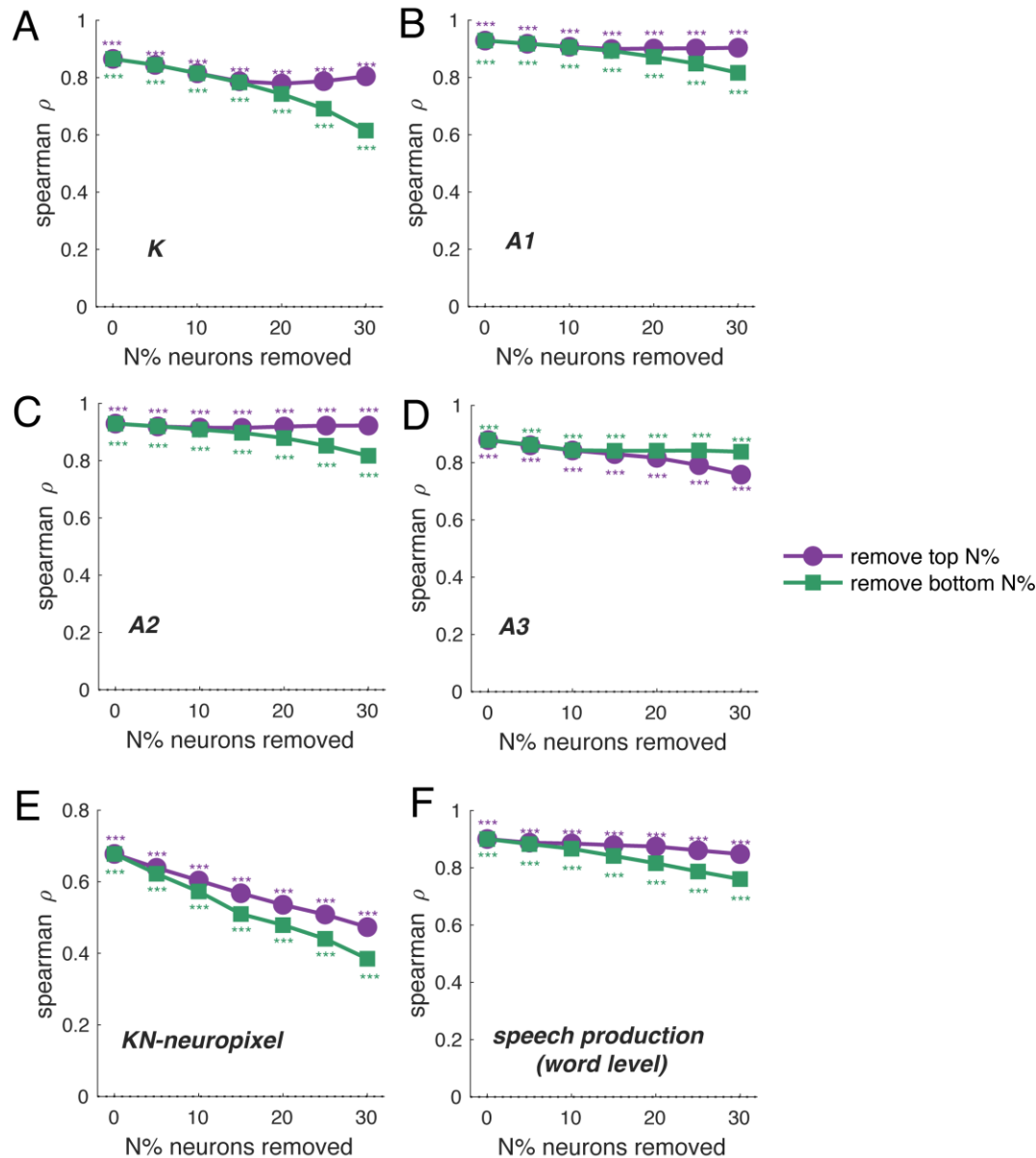


Figure S5. The robustness of per-neuron participation strength, related to Figure 6.

(A-E) We progressively removed the top 5-30% highest-participation neurons and recomputed the Spearman correlation at each step. The Spearman rank correlation remained very significant even after removing the top 30% of neurons (all stimulus sets, K, A1, A2, A3, KN-neuropixel in Study 1 and word speech in Study 2, $p < 0.001$; purple line). We then progressively removed the bottom 5-30% lowest-participation neurons and recomputed the correlation at each step. The Spearman rank correlation remained very significant (all stimulus sets, K, A1, A2, A3, KN-neuropixel in Study 1, and word speech in Study 2, $p < 0.001$; green line). Together, these results demonstrate that the effect is not driven by a small subset of extreme neurons.

* $p < .05$, ** $p < .01$, *** $p < .001$, permutation (1000 iterations)

Supplemental Tables

Table S1. Driver word properties of cross-language versus non-cross-language neurons (Study1), related to Figure 2.

Stimulus	<i>n</i> (cross-language)	Selectivity (K80%)			Cluster Index (<i>similarity</i>)			Participation Ratio (<i>dimensionality</i>)		
		Cross-Language neurons	Non-cross-language neurons	<i>d</i>	Cross-Language neurons	Non-cross-language neurons	<i>d</i>	Cross-Language neurons	Non-cross-language neurons	<i>d</i>
K	6	21.6 ± 6.4	18.3 ± 5.9	0.56	0.663 ± 0.002	0.664 ± 0.003	-0.092	66.5 ± 3.4	65.6 ± 3.4	0.27
A1	20	16.3 ± 5.8	25.6 ± 14.1	-0.66**	0.666 ± 0.011	0.667 ± 0.004	-0.295	55.3 ± 12.6	62.7 ± 6.5	-1.13***
A2	11	19.2 ± 9.1	24.3 ± 14.6	-0.35	0.696 ± 0.004	0.691 ± 0.005	0.927**	47.4 ± 4.4	47.9 ± 5.3	-0.10
A3	8	17.2 ± 11.1	30.0 ± 17.1	-0.75*	0.693 ± 0.019	0.690 ± 0.006	0.536	41.0 ± 17.4	49.7 ± 6.7	-1.28***
KN-neuropixel	9	26.6 ± 3.7	28.4 ± 5.2	-0.35	0.651 ± 0.004	0.652 ± 0.004	-0.294	56.9 ± 2.6	57.6 ± 3.4	-0.19
Pooled	54	19.3 ± 8.0	25.7 ± 13.4	-0.48***	0.673 ± 0.019	0.673 ± 0.015	0.017	53.4 ± 12.4	57.2 ± 8.7	-0.44**

Table Note. Values in each cell are mean ± SD.

d = Cohen's *d* (negative indicates cross-language neurons < non-cross-language neurons).

Selectivity (K80%): percentage of words needed to explain 80% of positive cross-language correlation mass.

Cluster Index: mean pairwise cosine similarity among driver words in mBERT embedding space.

Participation Ratio: effective dimensionality of driver word embeddings.

p* < .05, *p* < .01, ****p* < .001

Table S2. Driver word properties of cross-language versus non-cross-language neurons (read aloud task, Study 2, word level), related to Figure 2.

Stimulus	<i>n</i> (cross-language)	Selectivity (K80%)			Cluster Index (<i>similarity</i>)			Participation Ratio (<i>dimensionality</i>)		
		Cross-Language neurons	Non-cross-language neurons	<i>d</i>	Cross-Language neurons	Non-cross-language neurons	<i>d</i>	Cross-Language neurons	Non-cross-language neurons	<i>d</i>
Word-speech	61	17.4 ± 3.0	33.8 ± 16.4	-2.35***	0.807 ± 0.006	0.797 ± 0.004	1.722** *	28.4 ± 4.0	34.7 ± 5.6	-1.47***

Table Note. Values in each cell are mean ± SD.

d = Cohen's *d* (negative indicates cross-language neurons < non-cross-language neurons).

Selectivity (K80%): percentage of words needed to explain 80% of positive cross-language correlation mass.

Cluster Index: mean pairwise cosine similarity among driver words in mBERT embedding space.

Participation Ratio: effective dimensionality of driver word embeddings.

* $p < .05$, ** $p < .01$, *** $p < .001$.

Table S3. Cross-language neurons did not affect shared representation geometry, related to Figure 4 and Figure 5.

neural-LLM		K	A1	A2	A3	KN
r_with	spanish	r=0.209***	r=0.132***	r=0.152***	r=0.190***	r=0.263***
	english	r=0.161***	r=0.134***	r=0.173***	r=0.199***	r=0.202***
r_without	spanish	r=0.208***	r=0.164***	r=0.153***	r=0.188***	r=0.259***
	english	r=0.160***	r=0.174***	r=0.171***	r=0.199***	r=0.199***
Steiger's Z (r_without vs. r_with)	spanish	z = 0.329, p = 0.742	z=-21.751, p<0.001	z=-0.563, p=0.573	z = 0.846, p = 0.397	z=2.105, p=0.035
	english	z = 0.675, p = 0.500	z=-28.278, p<0.001	z=1.029, p=0.303	z = 0.723, p = 0.469	z=2.136, p=0.033
Permutation p- value for (Δr vs. Δr_{random})	spanish	p=0.751	p=1.00	p=0.764	p=0.511	p=0.206
	english	p=0.747	p=1.00	p=0.636	p=0.506	p=0.041

English-Spanish neural geometry similarities		K	A1	A2	A3	KN
r_with		r=0.477***	r=0.255***	r=0.281***	r=0.358***	r=0.385***
r_without		r=0.474***	r=0.341***	r=0.223***	r=0.316***	r=0.379***
Steiger's Z (r_without vs. r_with)		Z = 3.758, p = 0.0001	Z = -62.502, p<0.0001	Z = 2.447, p = 0.01	Z = 3.963, p < 0.0001	Z = 3.656, p = 0.0002
Permutation p-value for (Δr vs. Δr_{random})		p=0.773	p=1.000	p=0.736	p=0.486	p=0.606

Table Note. r_with = Cross-language RDM similarity (with cross-language neurons);
r_without = Cross-language RDM similarity (without cross-language neurons);
 Δr = r_with - r_without;
 Δr_{random} = mean Δr from random neuron exclusion (of equal size) (mean $\pm$ SD)
* $p < .05$, ** $p < .01$, *** $p < .001$

Table S4. No language-specific neurons, related to Figure 2.

	stimulus-K	stimulus-A1	stimulus-A2	stimulus-A3	stimulus-KN	Speech-word
mixed	105	101	92	119	172	72
Pure shared	0	0	0	0	0	0
Pure selective	0	2	4	8	0	2
selective-neurons above chance?	p=1.000	p=0.967	p=0.713	p=0.303	p=1.000	p=0.890
cross-language overlap exceeded chance	p<0.001	p<0.001	p<0.001	p<0.001	p<0.001	p<0.001

Table S5. “Cross-language neurons” in mBERT (all layers), related to Figure 2 and Figure 6.

Layer	K	A1	A2	A3	KN	Speech-word
layer0	99.7 %	99.5 %	97.3 %	97.0 %	99.2 %	83.1%
layer1	99.9 %	100%	99.2 %	100%	99.9 %	90.4%
layer2	100%	100%	100%	100%	100%	94.1%
layer3	100%	100%	100%	100%	100%	97.8%
layer4	100%	100%	100%	100%	100%	99.5%
layer5	100%	100%	100%	100%	100%	99.7%
>=layer6	100%	100%	100%	100%	100%	100%