

UC Merced

Proceedings of the Annual Meeting of the Cognitive Science Society

Title

A Quantitative Analysis of Cultural Universals and Variation in Moral Association

Permalink

<https://escholarship.org/uc/item/3km141jn>

Journal

Proceedings of the Annual Meeting of the Cognitive Science Society, 48(0)

Authors

Shah, Malaika

Xu, Yang

Ramezani, Aida

Publication Date

2026

Copyright Information

This work is made available under the terms of a Creative Commons Attribution License, available at <https://creativecommons.org/licenses/by/4.0/>

Peer reviewed

A Quantitative Analysis of Cultural Universals and Variation in Moral Association

Malaika Shah¹, Yang Xu^{2,3}, Aida Ramezani²

mshah213@utexas.edu, yang.xu@cs.utoronto.ca, armzn@cs.toronto.edu

¹ The University of Texas at Austin

² Department of Computer Science, University of Toronto

³ Cognitive Science Program, University of Toronto

Abstract

People's moral views may vary across cultures. Previous work has explored certain aspects of this cultural variation, but the extent to which moral variation and universals are reflected in the mental lexicon is underexplored. We present an analysis that quantifies cultural moral universals and variation based on large-scale word association constructed from English and Rioplatense Spanish speakers. We find that concepts related to 'wrongdoers' and 'medicine' are moralized more prominently among Spanish speakers, while 'conflict' and 'military action' are more moralized among English speakers. Furthermore, we show that the mental lexicon viewed through semantic categorization and structure explains 13% and 7% of the variances in the English and Spanish speakers in terms of moral association. Our study demonstrates that the mental lexicon, particularly the semantic network, offers a rich resource for characterizing the relations of morality and culture through the lens of their cognitive underpinnings.

Keywords: the lexicon; morality; cross-cultural analysis; word association; semantic network

Introduction

While morality is a fundamental domain of cognition, people from different cultures may differ in moral views. This cultural variation has important implications in various contexts ranging from moral decision-making regarding autonomous vehicles around the globe (Awad et al., 2018) to the moralization of everyday concepts in the lexicon, such as *smoking* or *cigarette* (Rozin, 1999; Ramezani et al., 2026). To what extent do cultures share or vary in their moral views, and what factors might predict such cross-cultural commonalities and differences? In this study, we explore these questions with a large-scale analysis that quantifies cultural moral universals and variation reflected in the mental lexicon.

Our starting point is the premise that word meaning, or how words relate to concepts in people's minds, can offer a window into understanding how we perceive things in a morally relevant context. This idea of linking moral relevance with the mental lexicon has received some empirical support. Recent work has shown that word associations or semantic networks can effectively capture English speakers' moral perception of certain concepts. For instance, knowing that a word like *cigarette* is frequently associated with terms including *health*, *bad*, *gross* offers a clue to the moral relevance of that word, and existing research has utilized such clues to characterize moral association for a wide range of concepts among English speakers in both contemporary (Ramezani & Xu, 2025) and

historical (Ramezani et al., 2026) settings. The success of this existing approach relies partly on the fact that many words in the lexicon signify moral foundations, or a set of moral modules like *care*, *fairness*, and *loyalty* and are demonstrated to be shared across cultures (Graham et al., 2009, 2013; Atari et al., 2020, 2023; Iyer et al., 2012). However, a critical limitation of this line of research is the restriction to the analysis of the English lexicon and therefore the moral views of English-speaking participants, which undercuts the diversity of languages and cultures around the world and therefore the moral cognition of non-English speakers (e.g., see Blasi et al. (2022) for a review).

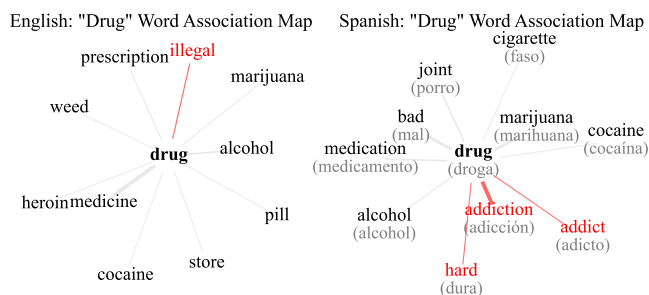


Figure 1: Simplified English and Spanish associations for the word *drug* (*droga* in Spanish) with red highlighting responses in the Moral Foundations Dictionary (MFD) (Graham et al., 2009; Frimer et al., 2019). The width of the line connecting the cue word and its respective responses indicates the relative frequency of each response word, where a thicker line means more frequent.

We present a novel framework that extends beyond existing work to characterize moral associations across cultures through word associations collected from speakers of different languages. We propose that examining cross-language patterns of word association may help both shared and varied moral views of those speakers. Figure 1 presents an illustrative example for our framework. Here, we compare the association network for the cue word *drug* in the English and Spanish word association networks, collected through crowd-sourced experiments (De Deyne et al., 2019; Cabana et al., 2024). As shown in the figure, the concept *drug* is strongly associated with the concept of *addiction* in Spanish, but this connection

is much weaker in the English network. Instead, the connection between *drug* and the concept *illegal* is stronger within the English network. How do cultures vary in their moral associations of various concepts, and what cognitive factors might explain such variation and commonalities? Our goal is to develop a methodological framework that supports automatic discoveries of (hidden) factors that may explain the rise of culturally shared and varied moral associations, clarifying why moral associations take different forms across speakers.

Our approach is related to a recent development in cognitive science, which utilizes large-scale semantic structures of the lexicon to inform cultural universals and variation. In particular, existing work has explored the use of word association data to reveal conceptual diversity across cultures and languages (Ke & De Deyne, 2024) and identify culturally distinctive terms (Lim et al., 2024). Other studies have also used word association or similar resources, such as bilingual dictionaries and word embeddings, to reveal universally shared properties of word meanings across languages (Xu et al., 2020; Brochhagen & Boleda, 2022; Brochhagen et al., 2023; Youn et al., 2016; Thompson et al., 2020; Khishigsuren et al., 2025). However, to our knowledge, the current study serves as an initial bridge to connect semantic networks with morality and culture.

We present our computational framework for investigating moral association across cultures. We compare moral associations of various semantic domains across English and Spanish speakers. To understand how the structure of word associations influences moral cognition, we extract numerical features from the word association networks that capture how words are positioned and connected in the mental lexicon. In particular, some features measure how easily a word can be accessed (its centrality in the network), while others capture how tightly it is clustered with neighboring concepts. These graph-based features allow us to systematically examine whether and how the structural organization of the mental lexicon shapes moral associations across cultures. Finally, we compare the effects of both semantic and graph-based features on moral association in English and Spanish speakers. For this initial investigation, we focus on a comparison of English-speaking and Spanish-speaking cultures. However, the framework we develop here can be applied in a more general setting to other languages and cultures in the future.

Data

We draw on two complementary word association networks in English and Spanish. In each, words are connected by directed links that represent the strength of association between concepts. Each link captures how strongly a cue word (source) triggers a response word (target) in participants' minds.

English word association

We used the word association dataset by De Deyne et al. (2019), known as the Small World of Words (SWOW-EN), for the English word association network. This dataset includes associations for more than 12,000 cue words and over 3 million

responses collected between 2011 and 2018. Each cue word has exactly 100 responses.

Spanish word association

Similar to the English data, we used the Small World of Words Rioplatense Spanish dataset (SWOW-ES) (Cabana et al., 2024). This dataset includes associations for more than 13,000 cue words and over 3.6 million responses, with the majority of participants from Argentina and Uruguay. Most cue words in SWOW-ES have exactly 70 associative responses, although a few have fewer.

While participants of these word association experiments could provide up to three responses for each cue word, we only considered the first response in both datasets to maintain consistency and prevent possible confounds from multiple responses. Previous work consistently reports a chaining phenomenon, where people generate their second and third responses based on their first one. This renders the second and third responses dependent on the first response, which is the primary one we analyze (De Deyne et al., 2019).

We processed association connections from both datasets through several preprocessing steps. Data were normalized to lowercase, stripped of whitespace, and lemmatized using spaCy (the `en_core_web_sm` package for English and `es_core_news_sm` for Spanish). We filtered out English words shorter than 4 letters and Spanish words shorter than 3 letters, and performed stopword removal. Non-word entries (e.g., booleans, NaN values) and multi-word phrases were removed, retaining only single words. Duplicate cue-response pairs after preprocessing were aggregated by summing their association frequencies.

Computational framework

We first review prior work on moral association graphs that we use in our analysis. We then explain our choices of the lexical and semantic features used to predict moral perception cross-culturally.

Moral association

Building on prior work characterizing moral perception with word association connections, we quantify the moral association of concepts through the Moral Association Graph (MAG) framework (Ramezani & Xu, 2025). Formally, the MAG score captures the proportion of morally-relevant associations for a given concept in the word association network as follows:

$$\text{MAG}(c) = \frac{\sum_{t \in T(c) \cap M} \mathbf{A}_{c,t}}{\sum_{t \in T(c)} \mathbf{A}_{c,t}}. \quad (1)$$

Here, $\mathbf{A}$ represents the adjacency matrix of the word association network, c is the query concept represented as a node in the graph (e.g., *drug*), $T(c)$ is the set of all possible responses for the query, and M is a set of morally-relevant words (a moral lexicon) such as *wrong*, *immoral*, *addiction*, and *gross*. Each entry $\mathbf{A}_{c,t}$ captures the number of participants in the word association network who associated the query word c

with the response (target) word t , serving as a measure of the normalized associative strength. For example, if among the 100 responses provided for the term *drug* in SWOW-EN, 20 are from this moral lexicon, then the MAG of *drug* would be 0.2. To maintain comparative power across English and Spanish word association networks, we applied z -score standardization of MAG scores to each network individually.

We use the Moral Foundations Dictionary (MFD), an expert-curated moral lexicon developed to facilitate data-driven investigations of Moral Foundations Theory (Graham et al., 2009, 2013; Frimer et al., 2019). To extend our results to Spanish, we used the Google Cloud Translation API to translate the MFD words to Spanish, then applied reverse-translation techniques to correct automated translation errors. Words in both moral lexicons were further filtered to include only single words, excluding stopwords and tokens shorter than three characters (English) or two characters (Spanish). Each term in the word association networks was considered to be from the moral lexicon if either its original or lemmatized form appeared in the moral dictionary. We acknowledge that using a Spanish translation of the English MFD is not optimal. While versions of the Spanish MFD have been created and released for public use (Carvalho & Guedes, 2022), we developed our own translated version to gain clarity on the process of creating the dictionary. Future work can include generating a native lexicon for non-English languages and their respective MFDs.

Semantic features

To study how semantic categories are uniquely moralized across English and Spanish speakers, we draw on the WordNet database (Princeton University, 2010) to extract the semantic taxonomies of the words in the English and Spanish word association networks. These features allow us to capture the word’s semantic category (or semantic domains) to explore what types of topics and concepts are more prone to becoming morally charged among English and Spanish speakers.

We made use of WordNet, which organizes concepts based on a systematic taxonomy, including hypernym (“is-a”) relations that ascend from more specific synsets (i.e., cognitive synonyms that represent distinct concepts) to increasingly abstract categories, with the most abstract level rooted in the *entity* synset. For example, WordNet categorizes the *drug* (noun) synset as: *drug* $\rightarrow$ *agent* $\rightarrow$ *causal agent* $\rightarrow$ *physical entity* $\rightarrow$ *entity*. This hierarchical structure naturally allows us to group concepts in the word association networks into various semantic domains.

To maintain consistency between synsets in English and Spanish, we used the Open Multilingual WordNet (OMW) (Bond & Foster, 2013) through the *omw-1.4* database available in the *nlTK* package in Python. The OMW project extends the structure of WordNet to more than 200 languages through aligned synsets. Crucially, OMW ensures identical structure across languages: the Spanish word for *drug* (*droga*) maps directly to the same English WordNet synset *drug.n.01* and follows the identical hypernym path. This

synset alignment enables consistent cross-linguistic categorization at equivalent abstraction levels. For each cue, we traced its hypernym path upward to hierarchical depth $d = 5$, which was empirically found to provide sufficient granularity for domain-specific moralization patterns without excessive fragmentation. After pruning synset groups with fewer than 10 words intersecting with our word association networks, we identified a total of 142 semantic categories for English and 73 categories for Spanish. However, this approach does not consider polysemous cases.

Graph-based features

To explore how the structure of the mental lexicon influences moral association, we extracted a variety of features for each node in the graph (i.e., each concept) that capture its local and global structure within the mental lexicon. We used the *networkX* library to extract these features from the adjacency matrix of the word association networks in English and Spanish separately. Table 1 provides the name and description of the features we extracted from the word association networks. Due to space limitations, we refer readers to the original publications for more details on each feature. In addition to these features, we also included word corpus frequency using the *wordfreq* software (Speer, 2022). Specifically, this package provides word frequencies in different languages based on a mixture of corpus data sources, including Wikipedia, Google Ngrams books, movie subtitles, and social media.

Results

In total, we identify 8772 terms in English and 4588 in Spanish with a complete profile of both features. The Spanish mean MAG score is 0.03 ($SD = 1.04$), and the English mean MAG score is 0.01 ($SD = 1.01$). Table 2 provides a list of the top and bottom ten words by moral association graph score. Words with the highest MAG scores are associated with the highest number of moral terms.¹

Semantic Analysis of Moral Association

Using semantic features from OMW, we examined the most morally associated WordNet categories in English and Spanish datasets. Figure 2 presents the top five categories with the highest average moral association in each language. For example, semantic categories of ‘christian,’ ‘conflict,’ and ‘military action’ have the strongest moral associations in English, while semantic categories ‘wrongdoer,’ ‘good,’ and ‘medicine’ show the highest association in Spanish. Despite these differences, cross-linguistic comparison revealed substantial agreement across the two datasets, where the moral association scores for equivalent semantic categories are statistically correlated between the two languages (Pearson’s $r = 0.78$, $p < 0.0001$, $n = 62$). Figure 4 further shows similar but not identical patterns of moral association in English and Spanish. For this test (unlike Figure 2), we only included categories with at least 10 members in both languages.

¹Code and data available at https://github.com/AidaRamezani/mag_culture/.

Feature	Description
Eigen-centrality (Wei, 1952; Kendall, 1955)	Measures a concept’s influence within the associative network. Concepts with high eigen-centrality are strongly associated with other highly influential concepts. Importantly, concepts with a few influential associations are typically more central than those with many peripheral associations. In the English word association network, <i>woman</i> and <i>water</i> are among the most eigen-central concepts. In Spanish, <i>amor</i> and <i>felicidad</i> are the most eigen-central concepts.
PageRank score (Page et al., 1999)	A normalized variant of eigen-centrality that accounts for the overall network structure. For example, concepts such as <i>money</i> and <i>water</i> have the highest PageRank scores in English.
Average neighbor PageRank	Provides the mean PageRank score of a concept’s direct associates, indicating whether a concept is embedded within an influential or peripheral region of the semantic network.
Clustering coefficient (Saramäki et al., 2007)	Measures the extent to which a concept’s associates are also associated with each other. A high clustering coefficient indicates that when people think of this concept, they activate a tightly interconnected semantic neighborhood. Words like <i>sahara</i> and <i>rodents</i> have the highest clustering coefficients among English speakers.
Average neighbor clustering coefficient	Measures the mean clustering coefficient of a concept’s associates, reflecting whether the concept is situated within a tightly or loosely organized semantic neighborhood.
Edge-cut centrality (Esfahanian, 2013)	Quantifies how well a concept bridges semantically distinct conceptual domains. For example, <i>mouse</i> bridges the domains of animals and technology. High edge-cut centrality suggests a concept serves as a cognitive connector between otherwise separate knowledge domains. Words like <i>romans</i> and <i>ducks</i> have the highest edge-cut centrality in English.
Coreness (Batagelj & Zaversnik, 2003)	Represents the largest k-core that a node belongs to, where a k-core is a maximal subgraph in which every node has degree k or more. Nodes with high coreness scores are embedded in densely interconnected subgraphs where all members maintain high connectivity. For example, words like <i>abandon</i> and <i>zoom</i> in English have the highest coreness scores.
Response frequency variance	Captures the distribution of association strengths for a given cue. Low variance indicates that responses are evenly distributed across associates, while high variance indicates that some associates are much more strongly activated than others. For instance, if all responses to a cue are directed to a single target word, the variance is 0.
Unique response ratio	The proportion of unique responses to total responses for a given cue. High ratios indicate that a concept activates a diverse set of associates rather than repeatedly eliciting the same few responses.
Triangle count	Quantifies the number of three-node cycles containing a given concept. Higher triangle counts indicate greater semantic cohesion, where a concept and its associates form densely interconnected clusters. Words like <i>food</i> and <i>bad</i> have the highest triangle counts.
Laplacian eigenvectors	Capture the spectral properties of the network, reflecting each concept’s structural position within the broader associative network. We use the values from the second (Fiedler vector) and third eigenvectors at every node (Fiedler, 1989). Nodes with similar eigenvector values tend to occupy similar structural positions and belong to similar communities in the graph.
In-degree vs. out-degree ratio	Compares backward associations (how often a concept is given as a response to other cues) with forward associations (how many responses the concept elicits as a cue). High ratios indicate that a concept is more frequently retrieved as an associate than it generates associations itself. In English, words like <i>money</i> and <i>food</i> have the highest ratio, indicating that participants often thought of these words as responses to the random cues they were given (e.g., <i>monopoly</i> → <i>money</i> , <i>mortgage</i> → <i>money</i> , <i>motivated</i> → <i>money</i>).

Table 1: Numerical graph-based features extracted from word association networks.

Language	Top ten words
English	convent, cuss, deity, enforcement, regulations, disposal, untruth, disciples, ebola, fib
Spanish	(convent) convento, (sanitation) sanidad, (parish) parroquia, (sanctity) santidad, (Christianity) cristianismo, (waste) residuos, (prayer) plegaria, (believer) creyente, (theft) hurto, (dirty) mugriento
Language	Bottom ten words
English	april, context, dwarfs, duvet, dusk, duration, duplicate, duo, dunk, a
Spanish	(hole) hueco, (hole) hoyo, (today) hoy, (accommodation) hospedaje, (to accelerate) acelerar, (volcano) volcán, (glass) vidrio, (thumb) pulgar, (position) puesto, (fan) abanico

Table 2: Top and bottom ten words by moral association score.

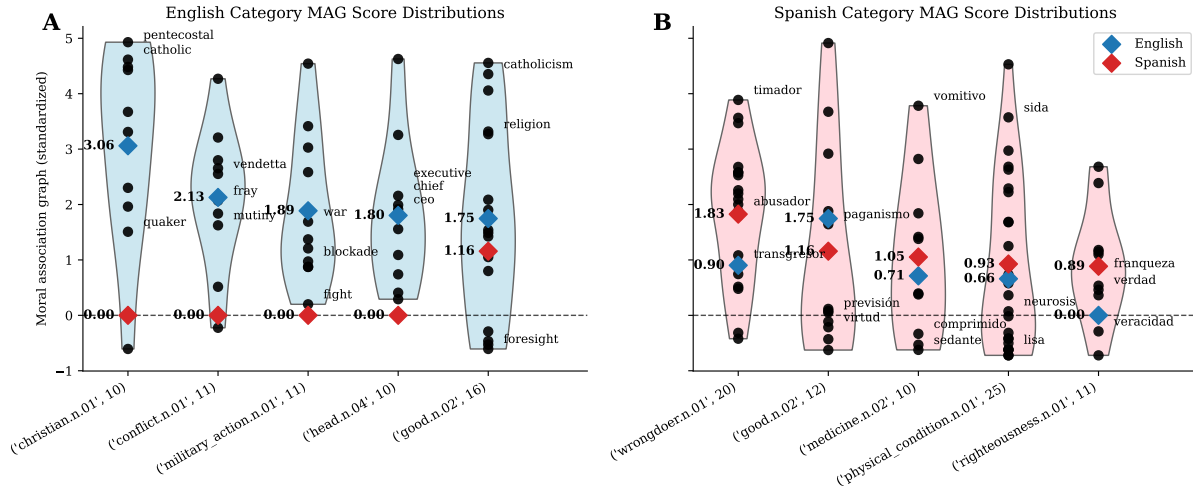


Figure 2: Distribution of moral association graph scores across semantic categories. Violin plots show moral association distribution within the top five WordNet categories with the highest mean scores in (A) English and (B) Spanish datasets. Diamond markers indicate mean scores for each category in English (blue) and Spanish (red). Word members are randomly selected as display representatives for each category. Numbers in parentheses show the total count of words for each category in the respective language. MAG value of zero indicates less than 10 members existed in the category.

Prediction and Interpretation of Moral Association

We used ordinary least squares regression with LASSO cross-validation to systematically examine which graph-based and semantic features predict moral association (i.e., MAG score) in English and Spanish. LASSO cross-validation allowed us to identify features with non-zero predictive influence while addressing multicollinearity. On the holdout data (20% of the total sample), we achieved $R^2 = 0.13$ ($n = 1,755$) for English and $R^2 = 0.07$ ($n = 918$) for Spanish. The predictive accuracy is relatively low, which suggests that other factors may play a role in determining moral associations.

Graph-based feature	English β	Spanish β
Laplacian e3	0.1040***	-0.3171***
PageRank	0.2914**	-0.1510
Eigen-centrality	-0.2449***	—
Triangle count	0.0290	0.2275***
Laplacian e2 (Fidler)	-0.1195***	-0.2194***
Clustering coefficient	0.1162***	0.1622***
Coreness	0.0709***	0.1109***
In-degree vs. out-degree ratio	-0.1055*	—
Unique response ratio	0.0343	0.0831
Average clustering coefficient	-0.0348**	-0.0833***
Average PageRank	-0.0286*	—
Corpus frequency	-0.0236	-0.0246
Edge cut	0.0174	-0.0144
Response variance	0.0005	0.0204
Response uniqueness	0.0003	-0.0530

Table 3: OLS regression β coefficients for graph-based features in English and Spanish datasets. Asterisks indicate levels of statistical significance for * $p < 0.05$, ** $p < 0.01$, *** $p < 0.001$. Features removed by LASSO regularization are indicated by ‘—’.

Zooming in on semantic categorization from WordNet, Fig-

ure 3 shows categories with statistically significant influence on the moral association of concepts. Panel A shows that words falling within categories such as ‘belief,’ ‘wrongdoer,’ ‘physical condition,’ and ‘evidence’ are more likely to be morally associated in both English and Spanish datasets. Panels B and C display semantic features that are statistically significant in only one of the two languages. For example, in the Spanish data, the semantic categories of ‘group action’ (e.g., *riña*: quarrel, *contienda*: contest), ‘district’ (e.g., *Paraguay*, *Inglaterra*: England, *patria*: homeland), and ‘righteousness’ (e.g., *respetuoso*: respectful, *sinceridad*: sincerity) independently affect moral association strength. What these results suggest is that while certain semantic domains universally predict moral association across English and Spanish speakers, other domains exhibit language-specific effects that may reflect cultural and lexical variation.

Moving to graph-based features and their influence on moral association, Table 3 shows the effects of statistically significant features in English and Spanish models independently. Similar to semantic features, we observe both universal and language-specific patterns. For example, both languages show that local network cohesion strongly predicts moral association. Clustering coefficient (measuring how interconnected a word’s neighbors are) and coreness (measuring embeddedness in densely connected network cores) both have positive effects. This suggests that morally-relevant concepts form tight semantic neighborhoods, where morally significant words tend to be surrounded by other interconnected concepts.

However, the role of node centrality and influence varies between languages. In English, PageRank (global influence) positively predicts moral association ($\beta = 0.2914$), suggesting that concepts connected to influential network hubs are

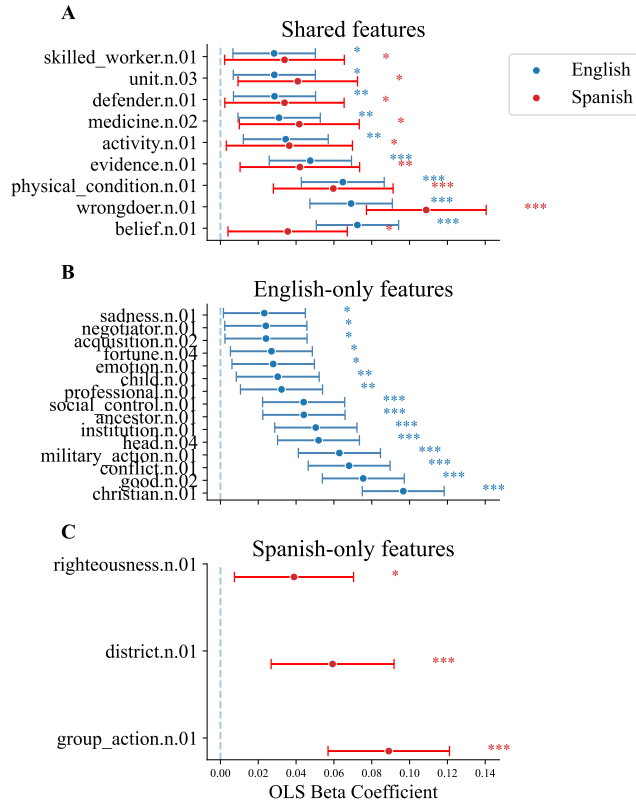


Figure 3: WordNet categories predicting MAG scores (all $p < 0.05$, positive β coefficients). Error margins show the 95% confidence intervals and asterisks indicate levels of statistical significance for * $p < 0.05$, ** $p < 0.01$, *** $p < 0.001$. (A) Shared categories between English and Spanish, ranked by β coefficient magnitude. (B, C) Language-specific categories. Asymmetry between English (B) and Spanish (C) categories might reflect differences in data coverage or lexical diversity.

perceived as more morally significant, but this feature is not statistically significant in Spanish. Conversely, eigenvector centrality in English shows a strong negative association ($\beta = -0.2449$), suggesting that morally-charged words occupy peripheral rather than bridging positions in English networks. These features do not appear as significant predictors in Spanish. We also note that in English, the in-degree to out-degree ratio negatively predicts moral association, suggesting that words which primarily point to other concepts (high out-degree) rather than being pointed to (high in-degree) have stronger moral associations. This asymmetry is absent in Spanish data. Laplacian eigenvectors further capture different aspects of network structure at a global level. For example, the Fiedler vector (e_2), which identifies the primary connectivity division in the network, negatively predicts moral association in both languages. This suggests that morally-charged words tend to fall on one side of the network’s main structural division rather than spanning it.

Together, these results show that while local neighborhood

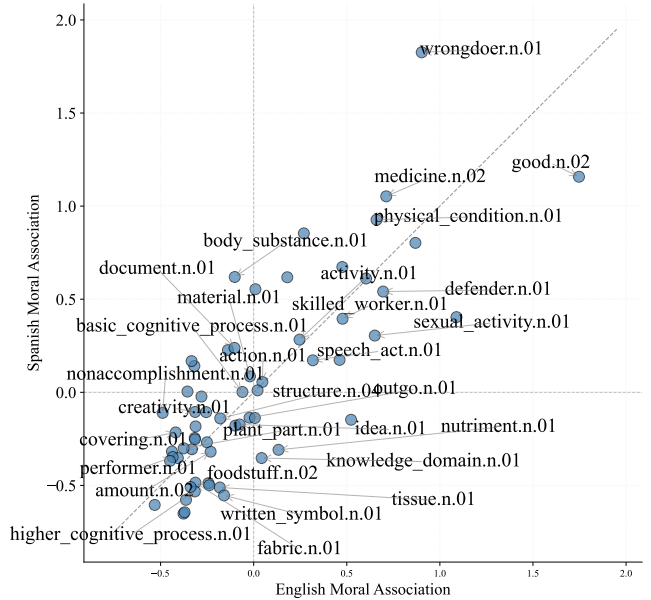


Figure 4: Moral association scores for WordNet categories. Each point represents a category containing at least 10 members in each language. The diagonal reference line shows perfect cross-linguistic agreement. Points above the line show stronger moral associations in Spanish, and points below show stronger moral associations in English.

cohesion predicts moral relevance in both English and Spanish, the relationship between global network position and moral association is language-specific. Even when controlling for word semantics (i.e., WordNet category), these structural properties of concepts within the mental lexicon (e.g., how accessible they are, how central or peripheral, how they divide semantic communities) significantly predict moral association, but the nature of this relationship depends on the linguistic and cultural context.

Conclusion

We present the first quantitative analysis of cross-cultural moral association based on word association, demonstrating that both what concepts mean and how they are positioned in semantic memory contribute to moral associations in culturally specific ways. Our analysis of English and Spanish word association data reveals both shared and language-specific tendencies in moral association. We also find that certain semantic domains, such as ‘wrongdoer’ and ‘belief’, universally predict moral association, while other domains exhibit language-specific associations. At the structural level, we observe that local network cohesion predicts moral relevance across both languages, while global centrality and node influence measures have differential effects. Future work should expand this framework to additional languages with native moral lexicons and, therefore, increase the diversity and quality of the cross-cultural analysis.

Acknowledgments

M.S. was supported by the Fields Undergraduate Summer Research Program. Y.X. is supported by The U of T Arts & Science Tri-Council Bridge Fund. A.R. was supported by the Data Science Institute at the University of Toronto. We thank Maike Anthony for his contributions during the early stages of this project as part of the Fields Undergraduate Summer Research Program. We thank Professor Jennifer Stellar for her feedback on the project and the manuscript.

References

- Atari, M., Graham, J., & Dehghani, M. (2020). Foundations of morality in iran. *Evolution and Human Behavior*, 41(5), 367–384.
- Atari, M., Haidt, J., Graham, J., Koleva, S., Stevens, S. T., & Dehghani, M. (2023). Morality beyond the weird: How the nomological network of morality varies across cultures. *Journal of Personality and Social Psychology*, 125(5), 1157.
- Awad, E., Dsouza, S., Kim, R., Schulz, J., Henrich, J., Shariff, A., . . . Rahwan, I. (2018). The moral machine experiment. *Nature*, 563(7729), 59–64.
- Batagelj, V., & Zaversnik, M. (2003). An o(m) algorithm for cores decomposition of networks. *arXiv preprint cs/0310049*.
- Blasi, D. E., Henrich, J., Adamou, E., Kemmerer, D., & Majid, A. (2022). Over-reliance on english hinders cognitive science. *Trends in Cognitive Sciences*, 26(12), 1153–1170.
- Bond, F., & Foster, R. (2013). Linking and extending an open multilingual wordnet. In *Proceedings of the 51st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)* (pp. 1352–1362).
- Brochhagen, T., & Boleda, G. (2022). When do languages use the same word for different meanings? the goldilocks principle in colexification. *Cognition*, 226, 105179.
- Brochhagen, T., Boleda, G., Gualdoni, E., & Xu, Y. (2023). From language development to language evolution: A unified view of human lexical creativity. *Science*, 381(6656), 431–436.
- Cabana, Á., Zugarramurdi, C., Valle-Lisboa, J. C., & De Deyne, S. (2024). The "small world of words" free association norms for rioplatense spanish. *Behavior Research Methods*, 56(2), 968–985.
- Carvalho, F., & Guedes, G. (2022). Dicionario de Fundamentos Morais em Espanhol. In M. Nussbaum, C. Infante, & J. Sanchez (Eds.), *Nuevas ideas en informática educativa, volumen 16* (pp. 287–291). Universidad de Chile.
- De Deyne, S., Navarro, D. J., Perfors, A., Brysbaert, M., & Storms, G. (2019). The "small world of words" word association norms for over 12,000 cue words. *Behavior Research Methods*, 51(3), 987–1006.
- Esfahanian, A.-H. (2013). Connectivity algorithms. *Topics in Structural Graph Theory*, 268–281.
- Fiedler, M. (1989). Laplacian of graphs and algebraic connectivity. *Banach Center Publications*, 25(1), 57–70.
- Frimer, J. A., Boghrati, R., Haidt, J., Graham, J., & Dehghani, M. (2019). Moral foundations dictionary for linguistic analyses 2.0. *Unpublished manuscript*.
- Graham, J., Haidt, J., Koleva, S., Motyl, M., Iyer, R., Wojcik, S. P., & Ditto, P. H. (2013). Moral Foundations Theory: The Pragmatic Validity of Moral Pluralism. In *Advances in Experimental Social Psychology* (Vol. 47, pp. 55–130).
- Graham, J., Haidt, J., & Nosek, B. A. (2009). Liberals and conservatives rely on different sets of moral foundations. *Journal of Personality and Social Psychology*, 96(5), 1029.
- Iyer, R., Koleva, S., Graham, J., Ditto, P., & Haidt, J. (2012, 08). Understanding libertarian morality: The psychological dispositions of self-identified libertarians. *PLOS ONE*, 7(8), 1–23.
- Ke, J., & De Deyne, S. (2024). Conceptual diversity across languages and cultures: A study on common word meanings among native english and chinese speakers. In *Proceedings of the Annual Meeting of the Cognitive Science Society* (Vol. 46).
- Kendall, M. G. (1955). Further contributions to the theory of paired comparisons. *Biometrics*, 11(1), 43–62.
- Khishigsuren, T., Regier, T., Vylomova, E., & Kemp, C. (2025). A computational analysis of lexical elaboration across languages. *Proceedings of the National Academy of Sciences*, 122(15), e2417304122.
- Lim, Z. W., Stuart, H., De Deyne, S., Regier, T., Vylomova, E., Cohn, T., & Kemp, C. (2024). A computational approach to identifying cultural keywords across languages. *Cognitive Science*, 48(1), e13402.
- Page, L., Brin, S., Motwani, R., & Winograd, T. (1999). The pagerank citation ranking : Bringing order to the web. In *The web conference*. Retrieved from <https://api.semanticscholar.org/CorpusID:1508503>
- Princeton University. (2010). *About WordNet*. Retrieved 2026-01-21, from <https://wordnet.princeton.edu/>
- Ramezani, A., Stellar, J., Feinberg, M., et al. (2026). Historical reconstruction of human moralization with word association and text corpora. *Nature Communications*, 17, 3412. Retrieved from <https://doi.org/10.1038/s41467-025-67891-2> doi: 10.1038/s41467-025-67891-2
- Ramezani, A., & Xu, Y. (2025). Moral association graph: A cognitive model for automated moral inference. *Topics in Cognitive Science*, 17, 120–138.
- Rozin, P. (1999). The process of moralization. *Psychological Science*, 10(3), 218–221.
- Saramäki, J., Kivelä, M., Onnela, J.-P., Kaski, K., & Kertesz, J. (2007). Generalizations of the clustering coefficient to weighted complex networks. *Physical Review E—Statistical, Nonlinear, and Soft Matter Physics*, 75(2), 027105.

- Speer, R. (2022, September). *rspeer/wordfreq: v3.0*. Zenodo. Retrieved from <https://doi.org/10.5281/zenodo.7199437> doi: 10.5281/zenodo.7199437
- Thompson, B., Roberts, S. G., & Lupyan, G. (2020). Cultural influences on word meanings revealed through large-scale semantic alignment. *Nature Human Behaviour*, 4(10), 1029–1038.
- Wei, T.-H. (1952). *Algebraic foundations of ranking theory*. Unpublished doctoral dissertation, University of Cambridge. Retrieved from <https://www.repository.cam.ac.uk/handle/1810/250988>
- Xu, Y., Duong, K., Malt, B. C., Jiang, S., & Srinivasan, M. (2020). Conceptual relations predict colexification across languages. *Cognition*, 201, 104280.
- Youn, H., Sutton, L., Smith, E., Moore, C., Wilkins, J. F., Maddieson, I., . . . Bhattacharya, T. (2016). On the universal structure of human lexical semantics. *Proceedings of the National Academy of Sciences*, 113(7), 1766–1771.