

Smart brands, trustworthy names, and evil pseudowords: form-meaning associations for existing, possible, and impossible Dutch strings

Joris van Eijkelenburg¹, Aron Joosse^{2,3}, Erik-Jan Van Kesteren^{3,4}, Giovanni Cassani¹

¹ Research Center for Cognitive Science and Artificial Intelligence, Tilburg University, Netherlands

² Atlas Research, Netherlands

³ ODISSEI Social Data Team (SoDa), Netherlands

⁴ Department of Methodology and Statistics, Utrecht University, Netherlands

Background: Systematic form-meaning mappings are well attested for a variety of semantic dimensions (shape [1], size [2], gender, age [3], valence [4], emotions [5]). They have been investigated relying on (i) pseudo-words, i.e., strings that follow the orthotactic patterns of a language but lack conventional meaning; (ii) given names [6], and (iii) brand names in marketing [7, 8]. These cases allow us to study whether stable form-meaning associations are affected by the degree of experience with a stimulus (given names are experienced, pseudowords and novel brand names are not) or by its plausibility (names and pseudowords follow the patterns of a language, brand names need not [9]). We studied associations for Dutch names, pseudowords, and novel brand names with four semantic dimensions with different degrees of conventionalization (gender, evilness, trustworthiness, smartness), that are also relevant in studies on hiring bias [10] and marketing [8].

Method: We collected ratings using best-worst scaling. In each trial, participants saw four items (out of 200 the same type) and judged which fitted a target semantic dimension (e.g., femininity) best and worst. Given names were sampled from the Dutch Census using stratified sampling by popularity bins. Pseudowords were taken from the Dutch Lexicon Project [11]. Novel company names were generated using a character-based 3-gram Markov model trained on a large store of real Dutch company names. 66 students enrolled at a Dutch university participated: each provided 192 ratings. After excluding inattentive participants, fast (<3s) and slow (>4SD above the mean) responses, we were left with ~10,000 ratings. We used a Bayesian rank-ordered logit model [12] to convert ratings into a numerical value, separately for each semantic dimension. This procedure allowed us to quantify the uncertainty due to rater disagreement, sample size, and the randomness in the item pairings. We then obtained 10,000 draws of the log-worth parameter of each item type-by-dimension combination. Finally, we embedded each item using a custom-trained FastText model [13] validated on semantic relatedness ratings [14] and computed the model-based associations of each item [15]. For each semantic dimension, we subtracted the average cosine similarity between an item and several words chosen to capture a certain semantic dimension (e.g., *man*, *boy*, *son* for MASCULINE) from the average cosine similarity to words capturing the opposite end of a semantic differential (e.g., *woman*, *girl*, *daughter* for FEMININE). We correlated this model-based score with the 10,000 draws from the posterior to ascertain to what extent distributional patterns pattern with human ratings. We conducted several robustness checks by varying the algorithm and the smallest and largest n-gram size to ensure results were consistent.

Results: In Figure 1, we see a sizable number of items with robust associations (i.e., whose 90% Credible Interval derived from the 10,000 draws from the posterior does not include 0). On the contrary, no item exhibits a robust association under 20 simulations with random responses, confirming that participants did agree on the connotations of many items because of the presence of a true signal. The FastText-based measure was significantly positively correlated with human ratings (Figure 2), typically outperforming a symbolic baseline leveraging the overlap in letter bigrams between items and reference words, i.e., how many bigrams are shared between, e.g., *Julie* and *boy*, *man*, *son* vs. *girl*, *woman*, *daughter*.

Discussion: Experience with a stimulus led to more robust associations on femininity and evilness (many given names show robust associations), but not on trustworthiness and smartness. The degree of plausibility of a stimulus, on the contrary, did not affect ratings: pseudowords and made-up company names resulted in a similar number of items with robust associations on evilness and femininity, *contra*

[9]. The prominence of a semantic dimension also mattered: gender and polarity (related with evilness), both cardinal dimensions of meaning and often lexicalised in inflectional and derivational morphology, exhibit the strongest associations across word types.

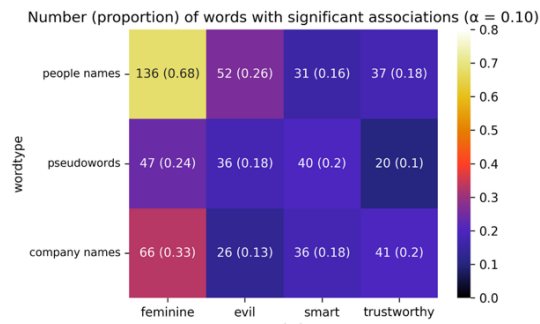


Figure 1: Number (and proportion) of items with robust associations per item-type (rows) by association.

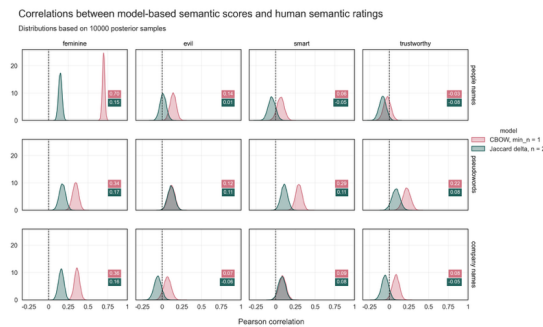


Figure 2: Correlations between FastText-based semantic associations (pink) and human ratings, per item-type (rows) by association (columns) combination. In green, the correlations for a purely symbolic baseline considering the overlap in letter bigrams between target items and reference words for each semantic dimension.

References

- [1] Ćwiek, A., Fuchs, S., Draxler, C., Asu, E., Dediu, D., Hiovain, K., & Winter, . (2022). The bouba/kiki effect is robust across cultures and writing systems. *Philosophical Transactions of the Royal Society B: Biological Sciences*, 377.
- [2] Sapir, E. (1929). A study in phonetic symbolism. *Journal of Experimental Psychology*, 12(3), 225–239. <https://doi.org/10.1037/h0070931>
- [3] Joosse, A., Kuscu, G., & Cassani, G. (2024). You sound like an evil young man: A distributional semantic analysis of systematic form-meaning associations for polarity, gender, and age in fictional characters. *Journal of Experimental Psychology: Learning, Memory, and Cognition*. <https://doi.org/10.31234/osf.io/wdsur>
- [4] Gatti, D., Raveling, L., Petrenco, A., & Günther, F. (2024). Valence without meaning: Investigating form and semantic components in pseudowords valence. *Psychonomic Bulletin & Review*, 31(5), 2357–2369. <https://doi.org/10.31234/osf.io/sfzgr>
- [5] Sabbatino, V., Troiano, E., Schweitzer, A., & Klinger, R. (2022). “splink” is happy and “phruth” is scary: Emotion Intensity Analysis for Nonsense Words. Association for Computational Linguistics. <https://doi.org/10.18653/v1/2022.wassa-1.4>
- [6] Sidhu, D., & Pexman, P. (2019). The Sound Symbolism of Names. *Current Directions in Psychological Science*, 28(4), 398–402. <https://doi.org/10.1177/0963721419850134>
- [7] Motoki, K., Park, J., Pathak, A., & Spence, C. (2022). The connotative meanings of sound symbolism in brand names: A conceptual framework. *Journal of Business Research*, 150, 365–373. <https://doi.org/10.1016/j.jbusres.2022.06.013>
- [8] Spence, C. (2012). Managing sensory expectations concerning products and brands: Capitalizing on the potential of sound and shape symbolism. *Journal of consumer psychology*, 22(1), 37–54. <https://doi.org/10.1016/j.jcps.2011.09.004>
- [9] Delgado, J., Pereira, R., Ferreira, M., Farinha-Fernandes, A., Guerreiro, J., Faustino, B., & Ventura, P. (2020). O simbolismo sonoro é modulado pela experiência linguística. *Revista Da Associação Portuguesa De Linguística*, (7), 137–150. <https://doi.org/10.26334/2183-9077/rapln7ano2020a9>
- [10] Newman, L., Tan, M., Caldwell, T., Duff, K., & Winer, E. (2018). Name norms: A guide to casting your next experiment. *Personality and Social Psychology Bulletin*, 44(10), 1435–1448. <https://doi.org/10.1177/0146167218769858>
- [11] Brysbaert, M., Stevens, M., Mandera, P., & Keuleers, E. (2016). The impact of word prevalence on lexical decision times: Evidence from the Dutch Lexicon Project 2. *Journal of Experimental Psychology: Human Perception and Performance*, 42(3), 441. <https://doi.org/10.1037/xhp0000159>
- [12] Glickman, M., & Hennessy, J. (2015). A stochastic rank ordered logit model for rating multi-competitor games and sports. *Journal of Quantitative Analysis in Sports*, 11(3), 131–144. <https://doi.org/10.1515/jqas-2015-0012>
- [13] Bojanowski, P., Grave, E., Joulin, A., & Mikolov, T. (2017). Enriching word vectors with subword information. *Transactions of the association for computational linguistics*, 5, 135–146. https://doi.org/10.1162/tacl_a_00051
- [14] Brans, L., & Bloem, J. (2024). SimLex-999 for Dutch. *Proceedings of the 2024 Joint International Conference on Computational Linguistics, Language Resources and Evaluation (LREC-COLING 2024)*, 14832–14845. <https://aclanthology.org/2024.lrec-main.1292/>
- [15] Caliskan, A., Bryson, J., & Narayanan, A. (2017). Semantics derived automatically from language corpora contain human-like biases. *Science*, 356(6334), 183–186. <https://doi.org/10.1126/science.aal4230>