

# Multilingual Contextual Affective Analysis of LGBT People Portrayals in Wikipedia

Chan Young Park,\* Xinru Yan,\* Anjalie Field,\* Yulia Tsvetkov

Language Technologies Institute, Carnegie Mellon University  
{chanyoun, anjalief, ytsvetko}@cs.cmu.edu, xinruyan@alumni.cmu.edu

## Abstract

Specific lexical choices in narrative text reflect both the writer’s attitudes towards people in the narrative and influence the audience’s reactions. Prior work has examined descriptions of people in English using *contextual affective analysis*, a natural language processing (NLP) technique that seeks to analyze how people are portrayed along dimensions of *power*, *agency*, and *sentiment*. Our work presents an extension of this methodology to *multilingual* settings, which is enabled by a new corpus that we collect and a new multilingual model. We additionally show how word connotations differ across languages and cultures, highlighting the difficulty of generalizing existing English datasets and methods. We then demonstrate the usefulness of our method by analyzing Wikipedia biography pages of members of the LGBT community across three languages: English, Russian, and Spanish. Our results show systematic differences in how the LGBT community is portrayed across languages, surfacing cultural differences in narratives and signs of social biases. Practically, this model can be used to identify Wikipedia articles for further manual analysis—articles that might contain content gaps or an imbalanced representation of particular social groups.

## Introduction

In 1952, Alan Turing was prosecuted for being gay and subsequently underwent a hormonal injection; two years later he committed suicide. Figure 1 shows parallel sentences drawn from his English, Spanish, and Russian Wikipedia pages. Although all three sentences describe the same situation, their connotations subtly differ. The English edition uses the verb *accepted*, which suggests that Turing had little control over the situation (low agency). In contrast, the verbs *chose* in Spanish and *preferred* in Russian imply that he actively made the decision (high agency). The verb *preferred* in Russian can even imply positive sentiment towards the injections, while the English connotation is more negative. Thus, Russian, Spanish, and English readers who search for Alan Turing on Wikipedia may form different impressions about this part in his life.

These subtle differences in phrasing can be indicative of social norms and perceptions about social roles (Eckert

### English Wikipedia:

He *accepted* the option of injections of what was then called stilboestrol.

### Spanish Wikipedia:

Finalmente escogió las inyecciones de estrógenos.  
Finally he *chose* estrogen injections.

### Russian Wikipedia:

Учёный предпочёл инъекции стилибэстрола  
The scientist *preferred* stilbestrol injections.

Figure 1: Example from Alan Turing’s biography page on Wikipedia in different languages. Verb choice in different languages can have subtly different connotations.

2000; Tannen 1994). In general, analyzing narratives about people sheds light on stereotypes and power structures (Hall and Braunwald 1981; Fournier, Moskowitz, and Zuroff 2002) and examining how these concepts differ across cultures is an important component of social-oriented analysis (Almeida et al. 2009; Balsam et al. 2011; Harris et al. 2013). In the example in Figure 1, these discrepancies in phrasing could indicate that stereotypes and bias about LGBT people manifest differently in Russian, English, and Spanish-speaking cultures. On Wikipedia, manifestations of stereotypes are violations of the platform’s “Neutral Point of View” policy.<sup>1</sup>

In this work, we develop computational methods that facilitate large-scale analyses of people described in multilingual narrative text. This technology can aid readers or writers, such as Wikipedia editors or journalists, in identifying biases in sets of articles that can be further analyzed to understand social stereotypes or edited in order to reduce bias.

Recent advances in NLP have analyzed stereotypes and biases in narratives in English (Bamman, O’Connor, and Smith 2013; Wagner et al. 2015; Sap et al. 2017). Field, Bhat, and Tsvetkov (2019) establish a framework called Contextual Affective Analysis (CAA) that focuses on affective dimensions of *power* (strength/weakness), *agency* (active-

\*Equal contribution

Copyright © 2021, Association for the Advancement of Artificial Intelligence (www.aaai.org). All rights reserved.

<sup>1</sup>[https://en.wikipedia.org/wiki/Wikipedia:Neutral\\_point\\_of\\_view](https://en.wikipedia.org/wiki/Wikipedia:Neutral_point_of_view)

ness/passiveness), and *sentiment* (goodness/badness). Analyzing portrayals of people along these dimensions (e.g., are men or women portrayed as more powerful?) has revealed stereotypes and bias in various domains, including movie scripts and newspaper articles (Rashkin, Singh, and Choi 2016; Sap et al. 2017; Field and Tsvetkov 2019; Field, Bhat, and Tsvetkov 2019; Antoniak, Mimno, and Levy 2019).

Measuring these affective dimensions relies on *connotation frames*—lexicons of verbs annotated to elicit implications (Rashkin, Singh, and Choi 2016; Sap et al. 2017). These connotations can be subtle, and the same verb often has different connotations in different contexts (Field, Bhat, and Tsvetkov 2019; Field and Tsvetkov 2019). Until now, these manually-annotated affective lexicons have existed only for English. While lexicons can be machine-translated (Rashkin et al. 2017; Mohammad 2018), no work has yet conducted in-language evaluations of connotation translatability nor attempted extensive analysis in other languages.

Our ultimate goal is to measure the power, agency, and sentiment of people described in multilingual text; we here focus on English, Spanish, and Russian. These measurements allow us to conduct both in-language analysis (in Russian text, are LGBT people portrayed as more powerful than non-LGBT people?) and cross-language analysis (is the power differential between LGBT people and non-LGBT people greater in English or in Russian?). To accomplish this, we first crowdsource annotations of connotation frames in English, Spanish, and Russian (§). We then analyze how connotations vary across contexts and languages in these data in order to demonstrate why existing English data sets are insufficient (§). With the new dataset, we develop multilingual CAA classifiers of power, agency, and sentiment (§).

Finally, we demonstrate the usefulness of our methodology in a semi-automated analysis §, by collecting a new corpus (LGBTBio) and analyzing how members of the LGBT community are portrayed on Wikipedia in different languages.<sup>2</sup> Our results show that the biography pages of LGBT people in our corpus typically contain more negative connotations than the pages of other people in Russian. In contrast, pages about the same LGBT people are more positive or neutral compared to other pages in English, and generally neutral or mixed in Spanish. These trends align with survey results about perceptions of LGBT people in English, Spanish, and Russian-speaking countries (Flores 2019).

The key contributions of our work are annotated datasets, machine learning models, and a general methodology that enables nuanced analyses of narratives about people across languages. We additionally present an analysis of LGBT people on Wikipedia, whereas extensive prior computational work on Wikipedia biography pages has focused primarily on male/female gender bias (Callahan and Herring 2011; Recasens, Danescu-Niculescu-Mizil, and Jurafsky 2013; Wagner et al. 2015; Chandrasekharan et al. 2017).<sup>3</sup>

<sup>2</sup>Our research focuses on the LGBT sub-community of the LGBTQIA+ community due to data scarcity of other groups.

<sup>3</sup>Code and data are publicly available at <https://github.com/chan0park/multilingual-affective-analysis>.

## Crowdsourcing Contextualized Connotation Frames

We first collected a corpus of multilingual connotation frames in English, Spanish, and Russian. Connotation frame annotations ask annotators to answer questions about the power, sentiment, and agency of the agent (approximated as the grammatical subject) and theme (approximated as the object) of verbs. At a high level, we seek to answer: (1) Does the subject have more/less/equal **power** as the object? (2) Does the subject have low/moderate/high **agency**? (3) Does the writer feel positive/negative/neutral about the subject (Sent<sub>subj</sub>)? (4) Does the writer feel positive/negative/neutral about the object (Sent<sub>obj</sub>)?

Consider the sentence: *The firefighter rescued the boy*. The verb *rescued* implies that the subject, *firefighter* has more power than the object, *boy*. The firefighter is also active and in control of his actions, which shows high agency. Because *rescuing* is a positive action (e.g. as opposed to *killing*), the writer likely feels positively about the subject. In the absence of information conveying positive or negative sentiment about *boy*, we can infer that the writer feels neutrally about the object.

Our connotation frames differ from existing lexicons in two primary ways: first, no prior work has collected connotation frames in languages other than English, and second, we collect all annotations in complete contexts drawn from newspaper articles, e.g., *the firefighter rescued the boy*, whereas prior work uses either simplified tuples or artificial placeholders, e.g., *X rescues Y* (Sap et al. 2017; Rashkin, Singh, and Choi 2016). Because these affective dimensions can be difficult to define, we took numerous steps to ensure annotations would be of high quality. We briefly summarize here and provide details in Appendix to facilitate reproducibility.

We first extracted frequent verbs and contexts that are representative of each verb’s most common usage, and then asked annotators to annotate verbs in these contexts. For each language, we extracted all (subject, object, verb) tuples from a *News Crawl* corpus.<sup>4</sup> We chose the 300 most frequent transitive verbs to annotate. For each verb, we took the three most common (subject, object, verb) tuples as the most representative context. We restricted tuples to have at least one human subject or object by using the list of words in *noun.person* category of WordNet (Fellbaum 2012).<sup>5</sup> We then pulled phrases containing the chosen tuples from the news corpus, which served as our data to be annotated.

We used the same interfaces as Rashkin, Singh, and Choi (2016) and Sap et al. (2017), with minor modifications based on feedback during pilot studies.<sup>6</sup> For non-English annotation tasks, a native speaker translated the task instructions into the target language. We restricted the pool of annotators to the United States for English, Russia for Russian, and to several South American countries for Spanish. For each of

<sup>4</sup>A large monolingual corpus of newspaper articles (Barrault et al. 2019). Throughout this work, parsing was done with SpaCy.

<sup>5</sup>Extensive list of English nouns denoting people (Miller 1995); we translated to other languages with Google Translate.

<sup>6</sup>We will provide full annotation interface in the released data.

|                   | English | Russian | Spanish |
|-------------------|---------|---------|---------|
| <b>Power</b>      | ↓29.2%  | ↓24.5%  | ↓26.7%  |
| <b>Agency</b>     | ↓31.1%  | ↓30.9%  | ↓35.2%  |
| <b>Sent(subj)</b> | ↓18.5%  | ↓18.4%  | ↓29.6%  |
| <b>Sent(obj)</b>  | ↓20.2%  | ↓23.6%  | ↓29.8%  |

Table 1: Assessment of how much information is lost when using verb-level (“rescues”) annotations instead of context-level (“the firefighter rescued the boy”). Accuracy decreases by nearly 20% for all languages.

the three target languages, we collected power, agency, and sentiment annotations for 300 verbs in three contexts each (900 instances). For each instance, we collected judgements from three annotators, leading to 32,400 total annotations.

Despite the steps taken to ensure annotation quality, we suspect that some annotators paid more attention to the task instructions and generated higher-quality judgements than others. We correct for this by discarding annotations from 14.4% of workers who frequently disagreed with other annotators (Appendix ). Krippendorff’s alpha, averaged across tasks, was 0.22 for English, 0.31 for Russian, and 0.26 for Spanish. These agreement scores are comparable to prior work (Rashkin, Singh, and Choi 2016; Sap et al. 2017) and reflect the subjective nature of connotation frames; very high agreement would suggest that we over-simplified the task, for example by choosing non-representative samples to annotate. Additionally, the most common case of annotator disagreement was when one annotator labeled an instance as neutral and another did not, meaning polar-opposite annotations were rare; if we only count polar-opposite annotations as disagreements, the average pairwise agreement is 92.5%. We describe additional restrictions used to ensure annotation quality and full agreement metrics in Appendix .

To aggregate annotations, we mapped each judgement to a  $(-1, 0, 1)$  value and averaged annotator scores. We then ternarized the aggregated scores by labeling connotations as positive  $[1, 0.35]$ , neutral  $(0.35, -0.35)$ , and negative  $[-0.35, -1]$ . With these boundaries, a connotation is only scored as positive or negative if at least two annotators labeled it with this polarity, and thus, samples where annotators disagreed were labeled as neutral—not clearly indicative of a positive or negative connotation.

## Crowdsourced Connotation Analysis

As described in §, our data differs from existing lexicons in two primary ways: (1) we collected annotations in context and (2) in various languages. We analyzed our data to assess how these differences impact lexicon quality.

**How important is contextualization?** We first address this question by examining how much accuracy would be lost if we use a single connotation score for each verb (e.g., one score for *deserves* in all contexts where it appears), rather than different scores when the verb appears in different contexts (e.g., different scores for *the boy deserves a reward* and *the boy deserves a punishment*) (Field and Tsvetkov 2019).

|                   | Russian | Spanish |
|-------------------|---------|---------|
| <b>Power</b>      | ↓37.6%  | ↓51.1%  |
| <b>Agency</b>     | ↓51.2%  | ↓47.8%  |
| <b>Sent(subj)</b> | ↓21.6%  | ↓34.8%  |
| <b>Sent(obj)</b>  | ↓30.4%  | ↓37.0%  |

Table 2: Assessment of information lost when using translated annotations instead of in-language annotations, evaluated over 125 Russian and 135 Spanish verbs that overlapped with English annotations.

To compute this, we “decontextualize” our lexicons by averaging annotations for each verb, regardless of the context it was annotated in, into a single verb-level score. We compare this decontextualized score with the contextualized score in our actual dataset, where we only average annotations for verbs annotated in the same context. Table 1 shows how often the verb-level score differs from the context-level score. If verbs had the same connotations in different contexts, all values in this table would be 0%. Instead, we see that ignoring contextualization can result in an  $> 30\%$  drop in accuracy. While Field and Tsvetkov (2019) have similar findings for sentiment connotations in English, Table 1 extends these findings to power and agency connotations and to Spanish and Russian.

**How important are in-language annotations?** We cannot directly compare contextualized annotations across languages because we annotate different contexts for different languages; however, there is some overlap between the most frequent verbs in each language. For each non-English language, we first aggregate annotations into the same decontextualized verb-level scores as in the previous paragraph. We then use Google Translate to translate verbs into English and intersect them with our annotated English verbs.<sup>7</sup> We discard any translation pairs that native Russian and Spanish speakers judged to be inaccurate or questionable (11 in Russian; 9 in Spanish). Finally, we measure how often the decontextualized English annotations differ from the original in-language annotations.

Table 2 reports the results. Because we assess the decontextualized scores, if word-level translation were effective for obtaining multilingual connotations, all scores would be similar to the ones in Table 1. Instead, they are substantially higher, showing that information is lost because of translation. Both Tables 1 and 2 suggest that our annotations can facilitate higher-quality analyses than ones collected in prior work. Because word-level translations are often inaccurate, in §, we also explore using a cross-lingual model and sentence-level translation to translate connotations.

## Classification of Connotations

Our goal is to develop methodology for analyzing how people are portrayed in different languages. The multilingual annotations alone are insufficient, as 900 contexts represent a

<sup>7</sup>Google Translate has previously been used to obtain multilingual lexicons (Mohammad and Turney 2013)

| Tgt | Src | Sent <sub>subj</sub> | Sent <sub>obj</sub> | Pow.         | Agen.        |
|-----|-----|----------------------|---------------------|--------------|--------------|
| EN  | EN  | <b>43.4*</b>         | 43.0                | <b>41.1</b>  | <b>48.2*</b> |
|     | ES  | 38.1                 | 43.4                | 29.5         | 43.4         |
|     | RU  | 41.1                 | <b>44.3</b>         | 40.1         | 41.4         |
| ES  | EN  | 38.9                 | 36.6                | 24.5         | 31.3         |
|     | ES  | <b>49.5*</b>         | <b>51.2*</b>        | <b>43.6*</b> | <b>43.6*</b> |
|     | RU  | 39.0                 | 42.2                | 34.0         | 38.9         |
| RU  | EN  | 43.6                 | 49.2                | 36.4         | 44.5         |
|     | ES  | 37.2                 | 49.3                | 38.2         | 42.7         |
|     | RU  | <b>46.4*</b>         | <b>54.9*</b>        | <b>45.3*</b> | <b>49.9*</b> |

Table 3: Macro F1 score of classifiers trained and evaluated with different target and source languages. Matching the language of the training and test data achieves better results than training on different languages. Asterisks are added to the best performing model in each (language, attribute) pair when it is significantly better than the second-best model (paired t-test, \*:  $p < 0.05$ ).

tiny subset of all verb usages in a corpus. Thus, we need a method to obtain connotation scores for unseen verbs and contexts. We describe our method here and provide reproducibility details in Appendix .

We follow prior work in developing a supervised classifier trained on our contextualized multilingual annotations that can predict a connotation frame label  $y_v$  for any unseen (in-context) verb  $v$  (Rashkin, Singh, and Choi 2016; Field, Bhat, and Tsvetkov 2019). Unlike prior models, ours is trained on contextualized annotations, and it leverages pre-trained cross-lingual language models (CLMs). CLMs produce language-agnostic feature representations, allowing us to combine different languages in the training and test data.

We obtain multilingual embedding representations ( $c_v$ ) of verbs in-context by extracting the last hidden layer of the pre-trained model called XLM, which achieves state-of-the-art performance in a variety of cross-lingual tasks (Conneau and Lample 2019). We then use  $c_v$  as features in a classifier.

Our primary classifier is a logistic regression model with sample weighting. The classifier is trained to predict a connotation frame label  $y_v$  from the input representation  $c_v$ . Sample weights are tuned over a dev set. While the classifier architecture is the same as in (Rashkin, Singh, and Choi 2016; Field, Bhat, and Tsvetkov 2019), inputs to our model differ. Prior connotation frame lexicons only contain verb-level annotations, and classifiers were trained using non-contextual embeddings (e.g., word2vec (Rashkin, Singh, and Choi 2016)) or decontextualized embeddings (e.g., ELMo (Field, Bhat, and Tsvetkov 2019)). With new annotations over verbs in-context, we directly train the classifier on contextual embeddings. Finally, we use the trained classifier to predict connotation frame labels for verbs provided with their context in a target language corpus.

## Connotation Classification Evaluation

We evaluate our model on the contextualized multilingual annotation data described in § using 5-fold cross-validation

| Tgt | Sent <sub>subj</sub> | Sent <sub>obj</sub> | Power | Agency |
|-----|----------------------|---------------------|-------|--------|
| ES  | 21.1                 | 20.1                | 31.7  | 27.3   |
| RU  | 18.9                 | 30.8                | 34.4  | 24.2   |

Table 4: Macro F1 for Machine Translation approach is strictly worse than for cross-lingual model (Table 3).

and splitting data into train, development, and test sets with the ratio of 6:2:2. Table 3 reports macro F1 scores<sup>8</sup> for single-language evaluations, where we train on the training set of one language (“Src”) and evaluate on the test set of a second language (“Tgt”). Unsurprisingly, training and testing on the same language achieves better performance than training on one language and testing on another. While the cross-lingual model works well in certain cases, in most cases transferring languages yields a substantial decrease in performance. For power in Spanish, F1 decreases by 19 points (44%) when the training language is English instead of Spanish. These results offer further evidence of the importance of in-language connotations.

**Machine Translated Classification** In Table 4, we consider an alternative approach to in-language annotations. We translate the Spanish and Russian test sentences into English through Google Translate, and then we use the model trained on English annotations to predict connotations in these translated sentences. This method simulates translating a corpus into English, and using a model trained on English for analysis. Machine translation performs strictly worse than a model training on in-language data, suggesting it cannot replace in-language data.

**Augmented Cross-Lingual Connotation Classification** While Table 3 suggests that in-language training data results in better performance than out-of-language training data, we also explore using in-language and out-of-language data in combination. This experimentation is based on the hypothesis that while connotations differ in different languages, there may be enough overlap for the cross-lingual model to learn useful signals from out-of-language data. Table 5 shows results. For all connotation dimensions, the best performing augmented models outperform the non-augmented models. In §, we use the best performing models from Table 5 to analyze biography pages of LGBT people.

## Case Study: Multilingual Affective Analysis of LGBT People

Finally, we demonstrate how the new corpus of annotations (§) and the cross-lingual model (§) facilitate multilingual analysis by examining portrayals of LGBT people on Wikipedia. We focus on LGBT people because discrimination against the LGBT community is an increasingly important global issue. Although pride marches are held  $\geq 158$

<sup>8</sup>We provide evidence that our model performs comparably with prior work in Appendix

| Tgt | Src    | S <sub>subj</sub> | S <sub>obj</sub> | Pow.         | Agen.        |
|-----|--------|-------------------|------------------|--------------|--------------|
| EN  | EN     | 43.4              | 43.0             | 41.1         | 48.2         |
|     | +ES    | 44.8              | <b>45.2*</b>     | 40.5         | 49.7         |
|     | +RU    | <b>46.5*</b>      | 43.2             | <b>41.8</b>  | 49.9         |
|     | +ES+RU | 45.0              | 44.3             | 41.7         | <b>50.0*</b> |
| ES  | ES     | 49.5              | 51.2             | <b>43.6</b>  | 43.6         |
|     | +EN    | 50.4              | 51.6             | 36.4         | 45.5         |
|     | +RU    | 51.0              | <b>55.0*</b>     | 42.1         | <b>45.6*</b> |
|     | +EN+RU | <b>51.8*</b>      | 54.8             | 40.8         | 44.9         |
| RU  | RU     | 46.4              | 54.9             | 45.3         | 49.9         |
|     | +EN    | 45.6              | 55.7             | 44.1         | 50.9         |
|     | +ES    | 46.0              | <b>59.2*</b>     | 42.1         | 49.8         |
|     | +EN+ES | <b>47.7</b>       | 53.7             | <b>46.9*</b> | <b>51.7*</b> |

Table 5: Macro F1 score of classifiers, where in-language training data is augmented with training data from other languages (e.g., Tgt=EN, Src=+ES indicates the model was trained on English and Spanish data). Augmentation improves performance in most cases. Asterisks are added to the best performing model in each (language, attribute) pair if there is a significant improvement made by data augmentation (paired t-test, \*:  $p < 0.05$ ).

cities world-wide (Lisitz 2017), social oppression is prevalent in many countries (Balsam et al. 2011; Doi and Stewart 2019). Nevertheless, little prior computational work has studied narratives about the community, likely due to data scarcity (Mendelsohn, Tsvetkov, and Jurafsky 2020). To facilitate analysis, we collect a new corpus, titled LGBTBio, which contains Wikipedia biography pages of LGBT people and pages of non-LGBT people (controls). We explain the motivation and methods behind our corpus collection in § and discuss results in §.

### LGBTBio Corpus

We collected a multilingual corpus of 1, 340 Wikipedia biography pages for people in the LGBT community using Wikipedia data properties and lists of LGBT people from Wikipedia (details in Appendix ). This corpus allows us to analyze how the same person is portrayed in different languages. However, we cannot draw conclusions from this corpus alone, because we need to control for overall language differences. For example, if we find that LGBT people have higher power in English pages than in Russian pages, we cannot determine if this difference occurs because LGBT people are actually described differently or because English verbs tend to have more positive power connotations than Russian verbs.

To isolate the effects of our variable of interest (i.e., sexual orientation), we built a control corpus by matching each LGBT person with a non-LGBT person who has similar characteristics using a matching method from Field, Park, and Tsvetkov (2020). Our goal is not to construct perfectly-

<sup>9</sup>All of our analysis is conducted over publicly available data. We do not expect our results to have any negative effects on the individuals analyzed or the broader LGBT community.

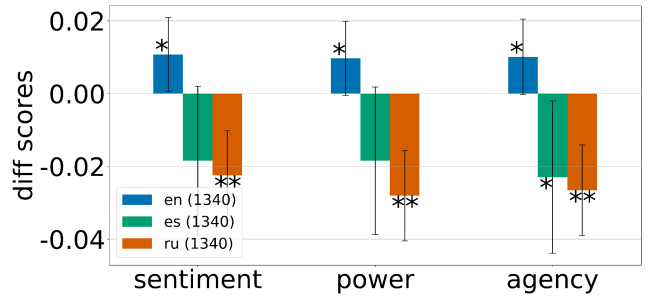


Figure 2: Average differences in affective scores in narratives about LGBT people vs. matched control people across languages. In Russian and Spanish, LGBT people are consistently portrayed more negatively, with lower power, and with lower agency, whereas in English, they are portrayed more positively. For all figures, asterisks indicate scores are statistically different from zero (paired t-test, \*:  $p < 0.05$  and \*\*:  $p < 0.01$ ) and brackets denote 0.95 confidence intervals. Numbers in the legend or in the title indicate the number of biographies in each group.

matched pairs, and we do not analyze any single pair. Rather, we seek to build a control corpus that has a similar distribution of characteristics as the main corpus *on aggregate*, excepting of sexual orientation, which is a common practice in causal inference studies (Stuart 2010).

To identify matches, for each person  $Y$  in the LGBT corpus, we first scraped their associated Wikipedia categories (e.g., *American fashion businesspeople*). We then identified all other people listed in these categories and constructed weighted TF-IDF vectors from each person’s associated categories. We selected the control person as the one who has the most similar category vector as  $Y$ . We refer to Appendix and Field, Park, and Tsvetkov (2020) for details about the corpus and matching algorithm. The 1, 340 pages about LGBT people and their matched controls together constitute the corpus.

### Contextual Affective Analysis of Narratives Describing LGBT People

We first use our model to compute high-level trends that identify directions for further exploration and then leverage our model to extract samples that exemplify these trends. Given the difficulty of the task and the subjective nature of connotations, we do not recommend using our model to draw black-box inferences, but rather suggest it is most useful in facilitating manual analyses that would otherwise be prohibitively expensive given the volume of Web-scale data. We also note that Wikipedia is constantly changing, and we observed several articles in our data set change over the course of this work. Thus, our results may not generalize to data collected at other times or to other articles beyond the specific ones included in our corpus.

For each language, we used the best performing model from Table 5 to predict contextualized verb annotations for all sentences that contain a target person’s name or pronoun

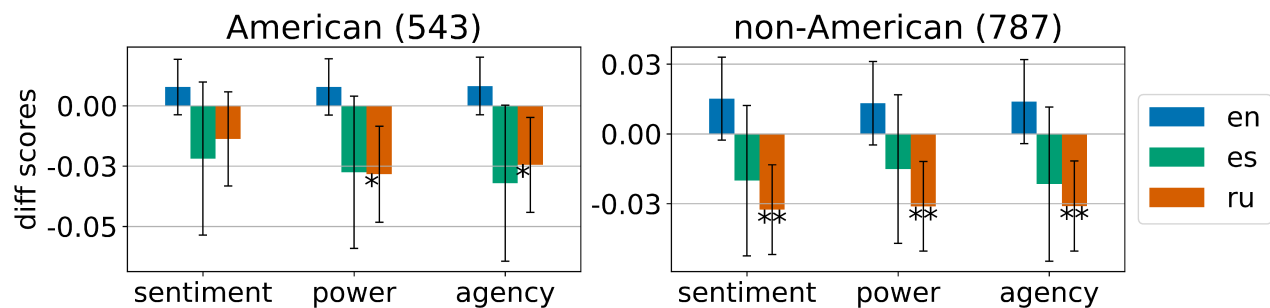


Figure 3: Average differences in affective scores for narratives about nationality subgroups in LGBTBio.

as the subject of a verb.<sup>10</sup> We then mapped these verb scores to entities (LGBT/matched control people) following Field, Bhat, and Tsvetkov (2019). We report *diff score* as the difference between sentiment/power/agency scores in each article about an LGBT person and its matched control, averaged across all pairs (e.g. “average treatment effect”); a positive score means the articles about LGBT people had a higher affect connotation in aggregate.

Figure 2 shows diff scores across the entire corpus. In English, all connotations are significantly positive, whereas in Russian all connotations are significantly negative. All connotations are also negative in Spanish, though only agency is significant. The differences in connotations across languages is surprising, considering that we examine articles about the exact same set of people in all languages.

These trends reflect global perceptions about LGBT people identified by other studies. Data from 2006 and 2011 suggest that many English Wikipedia editors are from the United States, many Russian Wikipedia editors are from Russia, and many Spanish Wikipedia editors are from Spain and (to a lesser extent) other countries, including Argentina, Chile, Netherlands, Mexico, and Venezuela.<sup>11,12</sup> While Wikipedia editor demographics may have since changed, this data shows historical contributions that have been made to Wikipedia and also reflects the countries where these languages are commonly spoken.

While homosexuality was de-criminalized in Russia in 1993, homophobia is still prevalent in the country and can manifest, for instance, in laws against “gay propaganda” (Wilkinson 2014; Buyantueva 2018). A report by the Wilkins Institute analyzed survey data from 174 countries to measure their social acceptance of LGBT people (Flores 2019). Based on data from 2014-2017, the United States ranked 21<sup>st</sup> (acceptance score of 7.2), while Russia ranked 120<sup>th</sup> (score of 3.4). Spain ranked highly (rank: 5, score: 8.1), while other Spanish-speaking countries ranked lower: Argentina (23; 6.9), Chile (27; 6.7), Mexico (32; 6.3), Venezuela (39; 5.7), Peru (53; 5.3). The results in Figure 2 suggest that these perceptions are reflected in the LGBTBio corpus: LGBT people are portrayed with more negative connotations in Russian, but not in English. Results are more mixed in Spanish, which is commonly spoken (and edited by on Wikipedia) people from a diverse range of countries.

In order to further examine possible cultural differences and in recognition that sexual orientation does not reflect an individual’s entire identity, we divide our corpus according to nationality, birth year, and occupation and test if additional social theories manifest in our data.

Figure 3 displays diff scores for each language over American people and non-American people in the corpus. While further subdivisions in nationality could be more informative, we do not report them, out of concern that results over small data sizes could be misleading (e.g. 34 people in our corpus have Russian nationality and 48 have South American nationalities.). In this figure, we investigate the “local heros” hypothesis, which suggests biography pages tend to be longer and more positive for people whose nationality match the language the page is written in (Callahan and Herring 2011; Eom et al. 2015). Our data does not show evidence of a local bias, in that trends are nearly identical across articles about American and non-American people, even in English, where we might expect to see different portrayals of Americans and non-Americans.

In Figure 4, we divide our corpus according to the year of birth of each article’s subject. Our primary research question in this figure is if the general increase in global acceptance of LGBT people over time is reflected in our data set (Flores 2019). Trends in Russian and Spanish do not significantly change across biography pages for people with different birth years. However, in English, portrayals are neutral/insignificantly negative for people born before 1900, but significantly positive for people born 1900-1960 and (insignificantly) positive for people born after 1960. Thus, while not all results are significant, our data does offer some evidence that changing global perceptions of LGBT people are reflected in their English Wikipedia biography pages.

Finally, in Figure 5, we subdivide the LGBT corpus by occupation. As in Figure 3, we focus only on the two most common occupations identified in our corpus (Entertainer and Artist) in order to ensure sufficient sample size. Survey and behavioral studies have suggested that LGBT people are

tences where the individual was an object are rare. All subgroups analyzed contained at least 280 verbs.

<sup>11</sup>[https://meta.wikimedia.org/wiki/Edits\\_by\\_project\\_and\\_country\\_of\\_origin](https://meta.wikimedia.org/wiki/Edits_by_project_and_country_of_origin)

<sup>12</sup>[https://commons.wikimedia.org/w/index.php?title=File:Editor\\_Survey\\_Report\\_-\\_April\\_2011.pdf&page=2](https://commons.wikimedia.org/w/index.php?title=File:Editor_Survey_Report_-_April_2011.pdf&page=2)

<sup>10</sup>We omitted Sent(Obj) and focused on Sent(Subj) since sen-

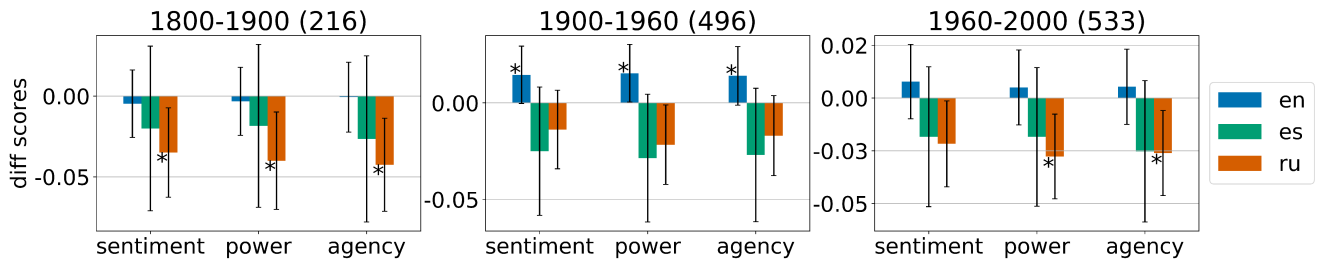


Figure 4: Average sentiment/power/agency diff scores for narratives about age subgroups in LGBTBio.

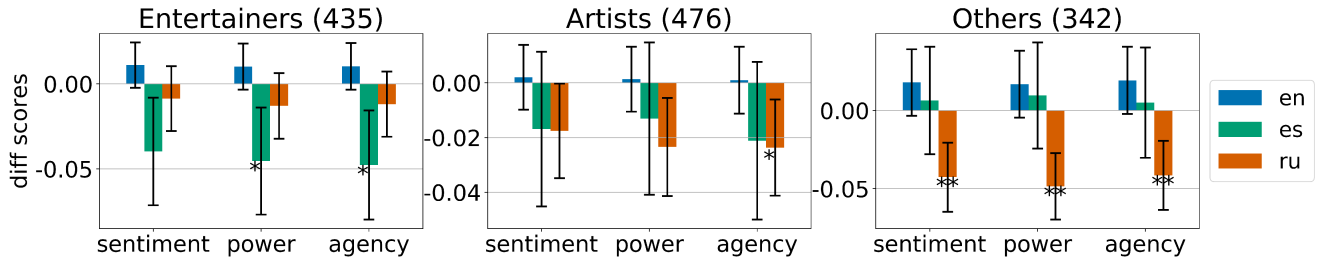


Figure 5: Average sentiment/power/agency diff scores for narratives about occupation subgroups in LGBTBio.

perceived as better suited to some occupations than others (Tilcsik, Anteby, and Knight 2015; Clarke and Arnold 2018). In Figure 5, we see little difference in how LGBT people of different occupations are portrayed on Wikipedia in English. However, in Spanish, articles about entertainers have significantly negative connotations, but articles about people of other occupations do not. Conversely, in Russian, while entertainers have near-neutral connotations, people of other occupations, such as politicians, scientists, and activists, are portrayed with significantly negative connotations. While more investigation is needed, our data offers evidence suggesting that perceptions about occupational stereotypes of LGBT people may differ across cultures and languages.

In general, the trends in Figure 2 remain consistent across different divisions of the data, with only slight variations. Articles about LGBT people are more negative than controls in Russian, but not in English, and results in Spanish are mixed.

**Identification of imbalanced content** One of the intended use-cases of our model is identifying sets of articles that could benefit from manual investigation and possibly revisions. In order to test this use-case, we examined the 10 article pairs that had the greatest difference in power scores between English and Spanish and similarly the 10 pairs with the greatest English-Russian power differentials, where the English article is more positive for both sets. We identify differences in English-Spanish and English-Russian versions of articles about LGBT people, both in what information is included and how it is phrased. We provide three examples drawn from this analysis:

*Cleve Jones* The Russian article about LGBT activist Cleve Jones spends several sentences describing the AIDS epidemic in the United States including his reaction to it:

Эта информация поразила Джонса, потому как многие из умерших были его друзьями или жили с ним по соседству в районе Кастро, он понял, что находится в центре ужасного невидимого бедствия.

(This information startled Jones, because many of the dead were his friends or lived with him in the Castro area, he realized that he was at the center of a terrible invisible disaster.)

Our model scores this sentence with negative power and agency and with neutral sentiment, reflecting Jones' role as a passive observer in this particular statement. In contrast, the English version of the article does not include this sentence, nor any detailed descriptions of the U.S. AIDS epidemic. Instead it focuses on projects that Jones initiated or worked on such as the AIDS Memorial Quilt. The primary difference between language editions in this example is the editors' decision whether to include information that frames Jones more positively or more negatively. Our model also scores this sentence with negative power and agency and with neutral sentiment:

Клиф Джонс публично заявил по телевидению о том, что является ВИЧ-инфицированным, и тотчас стал получать угрозы в свой адрес, а однажды даже подвергся нападению двух бандитов, попытавшихся его убить.

(Cleve Jones publicly announced on television that he was HIV-positive, and immediately began to receive threats against him, and once even was attacked by two bandits who tried to kill him.)

The English article does say "Jones described his status as HIV+" but makes no mention of threats or attacks.

*Hugh Walpole*<sup>13</sup> The Russian article about Hugh Walpole, an English novelist, describes:

В книге «Пироги и пиво» Мозем изобразил Уол-  
пола в качестве бездарного писателя-карьериста,  
которому честолюбие заменяет талант.  
In “Pies and Beers”, Maugham portrayed Walpole as a  
talentless careerist writer whose ambition replaces talent.

Our model scores this sentence with negative power and agency, and neutral sentiment. In contrast, the English article states: “His reputation in literary circles took a blow from a malicious caricature in Somerset Maugham’s 1930 novel *Cakes and Ale*: the character Alroy Kear, a superficial novelist of more pushy ambition than literary talent, was widely taken to be based on Walpole.” Our model does not score this sentence, as it does not refer to Walpole directly, but rather his reputation and a character based on him. The English article also includes a footnote explaining that Maugham originally denied the character was based on Walpole and only later recanted this denial. The Russian article does not make the distinction between a depiction of Walpole and a character based on Walpole, instead implying that Maugham’s critique directly applies to Walpole. The Russian article is also much shorter than the English article, and omits many of his career accomplishments and references to two biographies that portray him positively. For reference, the Spanish article, which is of comparable detail as the Russian article, does not mention Maugham’s book.

*Audrey Tang* The English article about Audrey Tang, a Taiwanese software programmer, describes “Tang was a child prodigy reading works of classical literature before the age of five...and they began to learn Perl at age 12. Two years later, they dropped out of junior high school, unable to adapt to student life.” which our model scores with neutral sentiment, power, and agency. The Spanish article states:

Con un cociente intelectual de 180, Tang tuvo dificultades  
para adaptarse a la educación formal desde su infancia, por  
lo que es autodidacta.  
(With an IQ of 180, Tang had a difficult time adjusting to  
formal education from childhood, making him self-taught.)

Our model scores this sentence with negative power, agency, and sentiment, as it focuses on Tang’s difficult time and suggests these challenges forced them to drop out, rather than suggesting they did so voluntarily. In general, we found that differences between Spanish and English articles were more subtle than differences between English and Russian articles.

## Related Work

Our work follows on a series of prior work: Rashkin, Singh, and Choi (2016) introduced sentiment connotation frames, Sap et al. (2017) extended them to power and agency, and Field, Bhat, and Tsvetkov (2019) introduced the CAA framework. Connotation frames have been used to analyze films,

newspaper articles, and online stories (Rashkin, Singh, and Choi 2016; Sap et al. 2017; Field, Bhat, and Tsvetkov 2019; Antoniak, Mimno, and Levy 2019). Rashkin et al. (2017) extend connotation frames to other languages through mapped embeddings, but they do not conduct evaluations against in-language annotations nor provide multilingual annotations.

Our work is generally consistent with existing literature on cross-cultural biases and online biographies. Dong et al. (2019) show that perceptions of social roles differ across cultures, while De-Arteaga et al. (2019) reveal gender bias in online biographies. Other work has examined biases in Wikipedia. Wagner et al. (2015) show that portrayals of men and women differ across languages, and Callahan and Herring (2011) reveal systematic cultural biases, particularly in biography pages.

Several studies in social science literature have analyzed biases and their effects on the LGBTQIA+ community, for example, examining mental health (Almeida et al. 2009) microaggressions (Balsam et al. 2011), and sociopolitical involvement (Harris et al. 2013). With a few exceptions (Schmidt and Wiegand 2017; Fast and Horvitz 2016; Dinakar et al. 2012; Mendelsohn, Tsvetkov, and Jurafsky 2020), biased language about or against the LGBTQIA+ community has not been examined and analyzed extensively in automated analyses. The closest study to ours is an examination of gender, race, and LGBT portrayals in 700 popular films (Smith et al. 2015).

## Conclusion

Our work provides methodology and datasets that extend the capabilities of affective analysis to multilingual settings. While we focus on Wikipedia, our methodology could be used to conduct analyses in any English, Russian, and Spanish narrative text, which can aid writers in obtaining a neutral point of view and provide insight into social stereotypes, especially when used in combination with other methods. This framework supports the investigation of a wide range of research questions, and offers multiple avenues for future work such as improving the multilingual model, expansion to additional languages, investigation of Wikipedia edit histories, and the incorporation of additional connotations and existing linguistic databases.

## Acknowledgements

We gratefully acknowledge Alan Black, Artidoro Pagnoni, Maria Ryskina, Santiago Castro, Luis Figueroa, Shuly Wintner and our anonymous reviewers for providing feedback and aiding with translations. AF gratefully acknowledges support from the Google PhD Fellowship. This material is based upon work supported by the National Science Foundation (NSF) Graduate Research Fellowship Program under Grant No. DGE1745016 and by NSF Grants No. IIS1812327 and IIS2040926. Any opinions, findings, and conclusions or recommendations expressed in this material are those of the authors and do not necessarily reflect the views of the NSF.

<sup>13</sup>[https://en.wikipedia.org/wiki/Hugh\\_Walpole](https://en.wikipedia.org/wiki/Hugh_Walpole)

## Multilingual Annotation Collection

We provide additional details of data collection to facilitate reproducibility and fully describe data quality.

**Language Choice** All annotations were crowdsourced through the Figure 8 platform.<sup>14</sup> Our original target languages for this task were English, Russian, Mandarin Chinese, and French. We constructed three rounds of in-house pilot studies for English and one round for Chinese before we released the annotation tasks. After releasing these tasks, we received almost no annotations in French or Chinese, despite increasing payment, expanding the number of target countries, and relaunching tasks. We ultimately dropped Chinese and French in favor of Spanish, for which we were able to obtain annotations.

**Instructions** Task instructions were originally written in English and then translated to other languages by native speakers. For each language, a second native speaker checked the translation. The same native speakers also examined the data samples to be annotated, in order to ensure that contexts were grammatical and representative. Heuristics for generating samples were revised according to their feedback.

While power and sentiment are generally well-known terms, agency is unfamiliar to most people and can be difficult to define. Additionally, our Russian, Chinese, and French translators determined that there is no single-word translation for “agency” in these languages, and they instead used combinations of other words to define the concept. Thus, in the annotation task, following Sap et al. (2017) we provided three agency “priming questions” to the annotators.

**Task Settings** We placed several restrictions on the pool of annotators in order to ensure annotation quality.

Each annotation task included five to eight examples in the task instructions, which were then turned into “quiz questions”. Annotators needed to answer an initial eight quiz questions with 70% accuracy to begin the task. As the task proceeded, an additional 10 quiz questions were interspersed with examples to be annotated, and annotators needed to maintain a 70% score on these questions in order to continue the task. We made our questions extremely similar to the examples given in the task instructions, because affective connotations can be subjective, and it is difficult to construct quiz questions that we can expect all high-quality annotators to consistently answer correctly. Instead, our quiz questions ensured that annotators read and understood the instructions.

We released data in batches of 100-200 annotation examples, as we found that larger batch sizes did not complete. Small batches also allowed us to block annotators who failed quiz questions in earlier batches from attempting to complete later batches. Additionally, we disabled the chrome translation plugin for all non-English tasks.

For Spanish, we restricted annotators to people from nine South America countries: Argentina, Bolivia, Chile, Colombia, Ecuador, Paraguay, Peru, Uruguay and Venezuela. We restricted English annotators to the United States and Russian annotators to Russia. Payment for annotation tasks was set based on the time taken to complete the task during pilot studies and the minimum wage of target countries. We also

adjusted pay based on survey feedback from early batches. The final rates are in Table 6. In general we paid agency substantially higher than the other two tasks because we had three additional priming questions that annotators needed to annotate for each instance.

| Task          | English | Russian | Spanish |
|---------------|---------|---------|---------|
| <b>Power</b>  | 20      | 4       | 5       |
| <b>Agency</b> | 40      | 8       | 8       |
| <b>Sent</b>   | 20      | 6       | 5       |

Table 6: Task pay rates in cents per five instances

Finally, as mentioned in §, we screened out annotations from lower-quality annotators after data collection. For each annotator, we computed how often that annotator judged an instance differently than the other two annotators who judged that instance. We then removed annotations from any annotators whose disagreement rate was greater than one standard deviation away from the mean disagreement rate. In our final dataset, we keep only instances that have at least two judgements after removing these annotators. Table 7 shows the full agreement scores for each annotation task after post-processing and Table 8 shows the number of annotated instances for each language.

|                   | English | Russian | Spanish |
|-------------------|---------|---------|---------|
| <b>Power</b>      | 0.27    | 0.33    | 0.25    |
| <b>Agency</b>     | 0.20    | 0.23    | 0.24    |
| <b>Sent(subj)</b> | 0.20    | 0.27    | 0.22    |
| <b>Sent(obj)</b>  | 0.22    | 0.39    | 0.31    |

Table 7: Krip.’s Alpha per task, after post-processing.

|                   | English | Russian | Spanish |
|-------------------|---------|---------|---------|
| <b>Power</b>      | 837     | 880     | 877     |
| <b>Agency</b>     | 888     | 879     | 888     |
| <b>Sent(subj)</b> | 860     | 868     | 808     |
| <b>Sent(obj)</b>  | 860     | 868     | 808     |

Table 8: Num. of annotated instances, after post-processing.

## Model Details

Each logistic regression classification model has 3,075 parameters. We did a grid search over the weights given for each class and chose the final weights based on the validation set F1 score. We release all trained models and their hyperparameter configuration as a part of our codebase.

Table 9 validates that the performance of our XLM-based model is comparable to prior work by evaluating our model on the same data used in Rashkin, Singh, and Choi (2016) and Field, Bhat, and Tsvetkov (2019). In this setting where we evaluate our model on uncontextualized annotations, we decontextualize the XLM embeddings in the same way as Field, Bhat, and Tsvetkov (2019). Thus, the primary difference between our model and theirs is the use of cross-lingual

<sup>14</sup>Figure 8 is now called Appen (<https://appen.com/>)

XLM embeddings instead of ELMo embeddings. Our model performs similarly with Rashkin, Singh, and Choi (2016). Although our model is slightly worse than Field, Bhat, and Tsvetkov (2019), this drop is not surprising and considered a cost of making our model language-agnostic.

|                   | Sent <sub>subj</sub> | Sent <sub>obj</sub> | Power | Agency |
|-------------------|----------------------|---------------------|-------|--------|
| Majority baseline | 26.2                 | 28.7                | 27.4  | 29.5   |
| Rashkin (2016)    | 66.6                 | 37.4                | 51.8  | 46.5   |
| Field (2019)      | 61.1                 | 40.4                | 56.0  | 48.8   |
| Our model         | 54.6                 | 45.0                | 47.4  | 45.0   |

Table 9: Classification evaluation results in macro F1 score. Majority baseline always outputs the most frequent label in the data. Numbers in the Field (2019) row are directly borrowed from the original paper.

## LGBTBio Corpus Construction

As described in §, we collected Wikipedia biography pages about LGBT people using lists of LGBT people from English<sup>15</sup> and Spanish<sup>16</sup> Wikipedia. We additionally include people with Wikidata property P91 values of Q6636 (homosexuality), Q6649 (lesbian), Q43200 (bisexuality), or Q592 (gay), and people with Wikidata property P21 of Q1052281 (transgender female) or Q2449503 (transgender male). We removed people who do not have pages in all target languages (English, Spanish, Russian) and who have < 3 sentences to be analyzed in any language. We define sentences to be analyzed as ones containing the person’s name or pronoun as a subject or object. We automatically inferred pronouns based on whether “he” or “she” is more frequent in the article text, which we expect to be effective even for transgender people because of Wikipedia’s Manual of Style/Gender identity.<sup>17</sup>

After filtering, the LGBT-half of the LGBTBio corpus contains 1,340 biography pages.

We identify matched control biography pages using the algorithm from Field, Park, and Tsvetkov (2020). We apply the same filters to candidate control pages as to LGBT pages (discarding articles with < 3 sentences to be analyzed in English, Russian, or Spanish). The matching algorithm using TF-IDF vectors constructed from biography page categories as matching features. The vectors contain a pivot-slope correction term, which is intended to prevent the method from favoring pages with fewer categories (Singhal, Buckley, and Mitra 1996). Following the recommendations in Field, Park, and Tsvetkov (2020), we set the pivot to the average number of categories per article in our data set. We then tune the slope until the LGBT-half and the matched controls have the same number of average categories (excluding LGBT-specific categories). We try pivot values in 0.1 increments

<sup>15</sup>[https://en.wikipedia.org/wiki/List\\_of\\_gay,\\_lesbian\\_or\\_bisexual\\_people](https://en.wikipedia.org/wiki/List_of_gay,_lesbian_or_bisexual_people)

<sup>16</sup>[https://es.wikipedia.org/wiki/Anexo:Famosos\\_que\\_han\\_salido\\_del\\_armario](https://es.wikipedia.org/wiki/Anexo:Famosos_que_han_salido_del_armario)

<sup>17</sup>[https://en.wikipedia.org/wiki/Wikipedia:Manual\\_of\\_Style/Gender\\_identity](https://en.wikipedia.org/wiki/Wikipedia:Manual_of_Style/Gender_identity)

between [0, 0.5], and fix the pivot as 0.1. In Table 10 we provide examples of constructed pairs.

| LGBT<br>(Non-LGBT pair)        | Common Categories (sampled three)                                                                                      |
|--------------------------------|------------------------------------------------------------------------------------------------------------------------|
| Tim Cook<br>(Steve Jobs)       | <i>Apple_Inc._executives</i><br><i>American_computer_businesspeople</i><br><i>21st-century_American_businesspeople</i> |
| Plato<br>(Aristotle)           | <i>4th-century_BC_philosophers</i><br><i>Ancient_Greek_political_philosophers</i><br><i>4th-century_BC_writers</i>     |
| Lily Allen<br>(Dua Lipa)       | <i>English_female_singer-songwriters</i><br><i>Brit_Award_winners</i><br><i>Electropop_musicians</i>                   |
| Tom Ford<br>(Anna Sui)         | <i>American_fashion_businesspeople</i><br><i>Luxury_brands</i><br><i>Parsons_School_of_Design_alumni</i>               |
| Francis Bacon<br>(Roger Bacon) | <i>Philosophers_of_ethics_and_morality</i><br><i>English_philosophers</i><br><i>Empiricists</i>                        |

Table 10: A sampled list of paired-people from LGBTBio.

## References

- Almeida, J.; Johnson, R. M.; Corliss, H. L.; Molnar, B. E.; and Azrael, D. 2009. Emotional Distress Among LGBT Youth: The Influence of Perceived Discrimination Based on Sexual Orientation. *Journal of Youth and Adolescence* 38(7): 1001–1014. ISSN 1573-6601. doi:10.1007/s10964-009-9397-9.
- Antoniak, M.; Mimno, D.; and Levy, K. 2019. Narrative Paths and Negotiation of Power in Birth Stories. In *Proc. ACM Hum.-Comput. Interact.*, volume 3, 88. ACM.
- Balsam, K. F.; Molina, Y.; Beadnell, B.; Simoni, J.; and Walters, K. 2011. Measuring Multiple Minority Stress: The LGBT People of Color Microaggressions Scale. *Cultur Divers Ethnic Minor Psychol.* 17(2): 163–174.
- Bamman, D.; O’Connor, B.; and Smith, N. A. 2013. Learning Latent Personas of Film Characters. In *Proc. of ACL*, 352–361. Sofia, Bulgaria: Association for Computational Linguistics.
- Barraut, L.; Bojar, O.; Costa-jussà, M. R.; Federmann, C.; Fishel, M.; Graham, Y.; Haddow, B.; Huck, M.; Koehn, P.; Malmasi, S.; Monz, C.; Müller, M.; Pal, S.; Post, M.; and Zampieri, M. 2019. Findings of the 2019 Conference on Machine Translation (WMT19). In *Proc. of WMT*, 1–61. Florence, Italy: Association for Computational Linguistics. doi:10.18653/v1/W19-5301.
- Buyantueva, R. 2018. LGBT Rights Activism and Homophobia in Russia. *Journal of Homosexuality* 65(4): 456–483. doi:10.1080/00918369.2017.1320167. URL <https://doi.org/10.1080/00918369.2017.1320167>. PMID: 28409697.
- Callahan, E. S.; and Herring, S. C. 2011. Cultural bias in Wikipedia content on famous persons. *Journal of the American society for information science and technology* 62(10): 1899–1915.

- Chandrasekharan, E.; Pavalanathan, U.; Srinivasan, A.; Glynn, A.; Eisenstein, J.; and Gilbert, E. 2017. You can't stay here: The efficacy of reddit's 2015 ban examined through hate speech. *Proc. ACM Hum.-Comput. Interact.* 1(CSCW): 31.
- Clarke, H. M.; and Arnold, K. A. 2018. The Influence of Sexual Orientation on the Perceived Fit of Male Applicants for Both Male- and Female-Typed Jobs. *Frontiers in Psychology* 9: 656. doi:10.3389/fpsyg.2018.00656.
- Conneau, A.; and Lample, G. 2019. Cross-lingual Language Model Pretraining. In *Advances in Neural Information Processing Systems*, 7057–7067.
- De-Arteaga, M.; Romanov, A.; Wallach, H.; Chayes, J.; Borgs, C.; Chouldechova, A.; Geyik, S.; Kenthapadi, K.; and Kalai, A. T. 2019. Bias in bios: A case study of semantic representation bias in a high-stakes setting. In *Proc. of FAT*, 120–128. ACM.
- Dinakar, K.; Jones, B.; Havasi, C.; Lieberman, H.; and Picard, R. 2012. Common sense reasoning for detection, prevention, and mitigation of cyberbullying. *ACM Transactions on Interactive Intelligent Systems (TiiS)* 2(3): 18.
- Doi, K.; and Stewart, P. 2019. Interview: The Invisible Struggle of Japan's Transgender Population. *Human Rights Watch*.
- Dong, M.; Jurgens, D.; Banea, C.; and Mihalcea, R. 2019. Perceptions of Social Roles Across Cultures. In *International Conference on Social Informatics*, 157–172. Springer.
- Eckert, P. 2000. *Language variation as social practice: The linguistic construction of identity in Belten High*. Wiley-Blackwell.
- Eom, Y.-H.; Aragón, P.; Laniado, D.; Kaltenbrunner, A.; Vigna, S.; and Shepelyansky, D. L. 2015. Interactions of cultures and top people of Wikipedia from ranking of 24 language editions. *PloS one* 10(3).
- Fast, E.; and Horvitz, E. 2016. Identifying Dogmatism in Social Media: Signals and Models. In *Proc. of EMNLP*, 690–699. Austin, Texas: Association for Computational Linguistics. doi:10.18653/v1/D16-1066.
- Fellbaum, C. 2012. WordNet. *The encyclopedia of applied linguistics*.
- Field, A.; Bhat, G.; and Tsvetkov, Y. 2019. Contextual Affective Analysis: A Case Study of People Portrayals in Online #MeToo Stories. In *Proc. of ICSWM 2019*.
- Field, A.; Park, C. Y.; and Tsvetkov, Y. 2020. Controlled Analyses of Social Biases in Wikipedia Bios. ArXiv preprint arXiv:2101.00078.
- Field, A.; and Tsvetkov, Y. 2019. Entity-Centric Contextual Affective Analysis. In *Proc. of ACL*, 2550–2560. Florence, Italy: Association for Computational Linguistics. doi:10.18653/v1/P19-1243.
- Flores, A. R. 2019. Social Acceptance of LGBT People in 174 Countries. <https://williamsinstitute.law.ucla.edu/publications/global-acceptance-index-lgbt/>. Accessed: 2021-04-17.
- Fournier, M. A.; Moskowitz, D.; and Zuroff, D. C. 2002. Social rank strategies in hierarchical relationships. *Journal of Personality and Social Psychology* 83(2): 425.
- Hall, J. A.; and Braunwald, K. G. 1981. Gender Cues in Conversations. *Journal of Personality and Social Psychology* 40(1): 99.
- Harris, A.; Battle, J.; Pastrana, Antonio (., J.; and Daniels, J. 2013. The Sociopolitical Involvement of Black, Latino, and Asian/Pacific Islander Gay and Bisexual Men. *Journal of Men's Studies* 21(3): 236–254.
- Lisitz, A. 2017. History of Pride Parades in the U.S. *Teen-Vogue*.
- Mendelsohn, J.; Tsvetkov, Y.; and Jurafsky, D. 2020. A Framework for the Computational Linguistic Analysis of Dehumanization. *Front. Artif. Intell.* doi:10.3389/frai.2020.00055.
- Miller, G. A. 1995. WordNet: A Lexical Database for English. *Commun. ACM* 38(11): 39–41. ISSN 0001-0782. doi:10.1145/219717.219748.
- Mohammad, S. 2018. Obtaining Reliable Human Ratings of Valence, Arousal, and Dominance for 20,000 English Words. In *Proc. of ACL*, 174–184. Melbourne, Australia: Association for Computational Linguistics. doi:10.18653/v1/P18-1017.
- Mohammad, S. M.; and Turney, P. D. 2013. Crowdsourcing a Word-Emotion Association Lexicon. *Computational Intelligence* 29(3): 436–465.
- Rashkin, H.; Bell, E.; Choi, Y.; and Volkova, S. 2017. Multilingual Connotation Frames: A Case Study on Social Media for Targeted Sentiment Analysis and Forecast. In *Proc. of ACL*, 459–464. Vancouver, Canada: Association for Computational Linguistics. doi:10.18653/v1/P17-2073.
- Rashkin, H.; Singh, S.; and Choi, Y. 2016. Connotation Frames: A Data-Driven Investigation. In *Proc. of ACL*, 311–321. Berlin, Germany: Association for Computational Linguistics. doi:10.18653/v1/P16-1030.
- Recasens, M.; Danescu-Niculescu-Mizil, C.; and Jurafsky, D. 2013. Linguistic Models for Analyzing and Detecting Biased Language. In *Proc. of ACL*, 1650–1659. Sofia, Bulgaria: Association for Computational Linguistics.
- Sap, M.; Prasetio, M. C.; Holtzman, A.; Rashkin, H.; and Choi, Y. 2017. Connotation Frames of Power and Agency in Modern Films. In *Proc. of EMNLP*, 2329–2334. Copenhagen, Denmark: Association for Computational Linguistics. doi:10.18653/v1/D17-1247.
- Schmidt, A.; and Wiegand, M. 2017. A Survey on Hate Speech Detection using Natural Language Processing. In *Proc. of SocialNLP*, 1–10. Valencia, Spain: Association for Computational Linguistics. doi:10.18653/v1/W17-1101.
- Singhal, A.; Buckley, C.; and Mitra, M. 1996. Pivoted document length normalization. In *Proc. of SIGIR*, volume 51, 176–184. ACM New York, NY, USA.
- Smith, S.; Choueiti, M.; Pieper, K.; Gillig, T.; Lee, C.; and DeLuca, D. 2015. Inequality in 700 Popular Films: Examining Portrayals of Gender, Race, & LGBT Status from 2007 to

2014. *Institute for Diversity and Empowerment at Annenberg*

Stuart, E. 2010. Matching methods for causal inference: A review and a look forward. *Stat Sci* 25(1): 1–21. doi:doi:10.1214/09-STS313.

Tannen, D. 1994. *Gender and discourse*. Oxford University Press.

Tilcsik, A.; Anteby, M.; and Knight, C. R. 2015. Concealable Stigma and Occupational Segregation: Toward a Theory of Gay and Lesbian Occupations. *Administrative Science Quarterly* 60(3): 446–481. doi:10.1177/0001839215576401. URL <https://doi.org/10.1177/0001839215576401>.

Wagner, C.; Garcia, D.; Jadidi, M.; and Strohmaier, M. 2015. It's a man's Wikipedia? Assessing gender inequality in an online encyclopedia. In *Proc. of ICWSM*.

Wilkinson, C. 2014. Putting “Traditional Values” Into Practice: The Rise and Contestation of Anti-Homopropaganda Laws in Russia. *Journal of Human Rights* 13(3): 363–379. doi:10.1080/14754835.2014.919218. URL <https://doi.org/10.1080/14754835.2014.919218>.