

Article

A Survey on Bias in Deep NLP

Ismael Garrido-Muñoz [†], Arturo Montejo-Ráez ^{*,†}, Fernando Martínez-Santiago [†] and L. Alfonso Ureña-López [†]

Centro de Estudios Avanzados en TIC (CEATIC), 230071 Jaén, Spain; igmunoz@ujaen.es (I.G.-M.); dofer@ujaen.es (F.M.-S.); laurena@ujaen.es (L.A.U.-L.)

* Correspondence: amontejo@ujaen.es; Tel.: +34-953-212-882

† These authors contributed equally to this work.

Abstract: Deep neural networks are hegemonic approaches to many machine learning areas, including natural language processing (NLP). Thanks to the availability of large corpora collections and the capability of deep architectures to shape internal language mechanisms in self-supervised learning processes (also known as “pre-training”), versatile and performing models are released continuously for every new network design. These networks, somehow, learn a probability distribution of words and relations across the training collection used, inheriting the potential flaws, inconsistencies and biases contained in such a collection. As pre-trained models have been found to be very useful approaches to transfer learning, dealing with bias has become a relevant issue in this new scenario. We introduce bias in a formal way and explore how it has been treated in several networks, in terms of detection and correction. In addition, available resources are identified and a strategy to deal with bias in deep NLP is proposed.

Keywords: natural language processing; deep learning; biased models



Citation: Garrido-Muñoz, I.; Montejo-Ráez, A.; Martínez-Santiago, F.; Ureña-López, L.A. A Survey on Bias in Deep NLP. *Appl. Sci.* **2021**, *11*, 3184. <https://doi.org/10.3390/app11073184>

Academic Editor: José Ignacio Abreu Salas

Received: 1 March 2021

Accepted: 22 March 2021

Published: 2 April 2021

Publisher's Note: MDPI stays neutral with regard to jurisdictional claims in published maps and institutional affiliations.



Copyright: © 2021 by the authors. Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<https://creativecommons.org/licenses/by/4.0/>).

1. Introduction

In sociology, bias is a prejudice in favor or against a person, group or thing that is considered to be unfair. Since, on one hand, it is a extremely pervasive phenomena, and on the other hand, deep neural networks are intended to discover patterns in existing data, it is known that human-like semantic biases are found when applying machine learning to ordinary human related results, such as computer vision [1], audio processing [2] and text corpora [3,4]. All these fields are relevant as constituents of automated decision systems. An “automated decision system” is any software, system or process that aims to automate, aid or replace human decision-making. Automated decision systems can include both tools that analyze datasets to generate scores, predictions, classifications or some recommended action(s) that are used by agencies to make decisions that impact human welfare, which includes but is not limited to decisions that affect sensitive aspects of life such as educational opportunities, health outcomes, work performance, job opportunities, mobility, interests, behavior and personal autonomy.

In this context, biased artificial intelligence models may make decisions that are skewed towards certain groups of people in these applications [5]. Obermeyer et al. [6] found that an algorithm widely used in US hospitals to allocate health care to patients has been systematically discriminating against black people, since it was less likely to refer black people than white people who were equally sick to programs that aim to improve care for patients with complex medical needs. In the field of computer vision, some face recognition algorithms fail to detect faces of black users [7] or labeling black people as “gorillas” [1]. In the field of audio processing, it is found that voice-dictation systems recognize a voice from a male more accurately than that from a female [2]. Moreover, regarding predicting criminal recidivism, risk assessment systems are likely to predict that people of some certain races are more presumably to commit a crime [8].

In the field of deep natural language processing (deep NLP), word embeddings and related language models are massively used nowadays. These models are often trained on

large databases from the Internet and may encode stereotyped biased knowledge and generate biased language. Such is the case of dialog assistants and chatbots when using biased language [9], or resume-review systems that ranks female candidates as less qualified for computer programming jobs because of biases present in training text, among other NLP applications. Caliskan et al. [10] propose the Word Embedding Association Test (WEAT) as a way to examine the associations in word embeddings between concepts captured in the Implicit Association Test (IAT) [11], in the field of social psychology, intended to assess implicit stereotypes held by test subjects, such as unconsciously associating stereotyped black names with words consistent with black stereotypes.

This problem is far from being solved, or at least attenuated. Currently, there are no standardized documentation procedures to communicate the performance characteristics of language models in spite of some efforts to provide transparent model reporting such model cards [12] or Data Statements [13]. In addition, the new models use document collections that are getting larger and larger during their training, and they are better able to capture the latent semantics in these documents, it is to be expected that biases will become part of the new model. This is the case of GPT-3 [14], a state-of-the-art contextual language model. GPT-3 uses 175 billion parameters, more than 100x more than GPT-2 [15], which used 1.5 billion parameters. Thus, Brown et al. [14] report findings in societal bias, more concisely regarding gender, race and religion. Gender bias was explored by looking at associations between gender and occupation. They found that 83% of 388 occupations tested were more likely to be associated with a male identifier by GPT-3. In addition, professions demonstrating higher levels of education (e.g., banker, professor emeritus) were heavily male leaning. On the other hand, professions such as midwife, nurse, receptionist, and housekeeper were heavily female leaning. Racial bias was explored by looking at how race impacted sentiment. The result: Asian race had a consistently high sentiment, while Black race had a consistently low sentiment. Finally, religious bias was explored by looking at which words occurred together with religious terms related to the following religions. For example, words such as “violent”, “terrorism”, and “terrorist” were associated with Islam at a higher rate than other religions. This finding is consistent with the work reported in [16]. When GPT-3 is given a phrase containing the word “Muslim” and asked to complete a sentence with the words that it thinks should come next, in more than 60% of cases documented by researchers, GPT-3 created sentences associating Muslims with shooting, bombs, murder or violence. As evidence of the growing interest in this problem, if we consider the bibliography prior to 2019 included in the present paper, English is the only language to be studied. However, from 2019 onward, it is possible to find papers dealing with the problem in more than 10 languages. Likewise, the number of articles has increased in recent years. Thus, we report just 1 paper in 2016 and 2017, 5 in 2018, 21 in 2019 and 18 papers in 2020. Not only it is a matter of the number of languages and papers, also new dimensions regarding bias are studied over the course of the last years. In this way, the early studies are focused on bias associated with gender. Thus, if we consider the studied papers in 2016 and 2017, there are only two articles out of seven that deal with other biases apart from the one associated with gender. By contrast, this number increases to 12 out of the 18 papers reviewed in 2020. This paper provides a formal definition of bias in NLP and a exhaustive overview on the most relevant works that have tackled the issue in the recent years. Main topics in bias research are identified and discussed. The rest of this paper is structured as follows: firstly, it is introduced a formal definition of bias, and its implication in machine learning in general, and language models in particular. Then, we present a review of the state-of-the-art in bias detection, evaluation and correction. Section 5 is our proposal for a general methodology for dealing with bias in deep NLP and more specifically in language model generation and application. We finalize with some conclusions and identify main research challenges.

2. Defining Bias

The study of different types of bias in cognitive sciences has been done for more than four decades. Since the very beginning, bias has been found as a innate human strategy for decision making [17]. When a cognitive bias is applied, we are presuming reality to behave according to some cognitive priors that may are not true at all, but with which we can form a judgment. A bias can be acquired by an incomplete induction process (a limited view over all possible samples or situations) or learned from others (educational or observed). In any case, a bias will provide a way of thinking far from logical reasoning [18]. There are more than 100 cognitive biases identified, which can be classified in several domains like social, behavioral, memory related and many more. Among them, there is one that we will focus on: stereotyping.

If a cognitive bias can be defined as a case in which human cognition reliably produces representations that are systematically distorted compared to some aspect of objective reality [19], *stereotyping* can be defined as the assumption of some characteristics applied to others on the basis of their national, ethnic or gender groups [20]. Therefore, stereotyping assigns certain characteristics to an individual because that individual pertains to a certain group. Somehow, it is like an ontology where certain classification rules are applied (so certain properties are presumed, like ignorance, weaknesses or criminal behavior) just because the individual possesses one specific value for a given property (she holds the “female” value for the property “gender”, or he holds the “African” value for the property “ethnicity”). As can be seen, stereotyping can be modeled at semantic level using a formal scheme like those provided by ontology languages in knowledge engineering.

We will first introduce fairness, as it is a well-know concept in machine learning (as it is, actually, equivalent to “zero-biases” systems), along with some of the measures used for its treatment. We will then discuss how fairness measures can help us to approach the bias problem in language models. To end this section, our proposal for a formal definition is provided.

2.1. The Bias Problem in Machine Learning

A concept that is intimately associated with bias is *fairness*. A system is considered to be “fair” when its outcomes are not discriminatory according to certain attributes, like gender or nationality. In machine learning evaluation, discrimination can be estimated looking at the confusion matrices for different protected groups. That is, we can compute confusion matrices and derived rates (positive rates, true positive rates, false positive rates and so on) for each subset of samples obtained as a segmentation of the full collection of samples on a certain feature (like “gender”). If these rates are far from being equal, that is a potential evidence of a prediction system with an “unfair” behavior, i.e., with a clear bias on how decisions are made depending on the values of that certain feature. Several measures have been proposed to study divergences among prediction rates over different population groups, and how to interpret them according to each system goal is now clearly identified [21]. From the large amount of biases derived from cognitive ones, about 24 are of interest in machine learning problems [5]. These latter two studies compile several measures that have been agreed in the analysis of the bias problem in machine learning systems, those measures are *demographic parity*, *equal opportunity*, *equalized odds* or *counterfactual fairness*, among others. Of course, these measures can be applied in many artificial intelligence subareas, like image recognition or natural language processing. Let us see the definition of one of them (demographic parity), as some elements can be transferred to our formal definition of bias in language modeling.

Demographic parity states that all the groups resulting from the different values of a protected class (e.g., gender) should receive the same rate of positive outcomes [22]. For example, if the system decides to concede a scholarship with the same rate to people in both male and females groups, then system shows demographic parity. Let $\hat{Y}$ be the predicted decision on whether a scholarship should be granted ($\hat{Y} = 1$) or denied ($\hat{Y} = 0$). Then, demographic parity can be defined as $P(\hat{Y} = 1|A = 0) = P(\hat{Y} = 1|A = 1)$, which is

equivalent to equal positive rates for both male and females $PR(A = 0) = PR(A = 1)$. Here, $\hat{Y}$ is the system prediction and A is the “protected” attribute/class. In our example, this is the gender and its possible values are 0 for male or 1 for female. Of course, this measure can be generalized to any protected class, like ethnicity or nationality. In that case, fairness is granted if positive rates for all possible population segments are equal. Where is bias here? It is right there, as the bias would be the deviation between groups resulting from different values of the protected attribute. Thus, the bias would be $bias = |P(\hat{Y} = 1|A = 0) - P(\hat{Y} = 1|A = 1)|$, which is equal to $bias = |PR(A = 0) - PR(A = 1)|$ using demographic parity as estimator. *Equal opportunity* is a good estimator of fairness. This one considers the equality between true positive rates. The rest of the measures are, as pointed out, variants on what we want to be equal from different scores.

In general, fairness is computed over the distribution $\langle X, A, Z, Y, \hat{Y} \rangle$, referring X to samples, A to protected attribute, Z to rest of attributes, Y the true labels for those samples and $\hat{Y}$ to the predicted labels by the model. This clear definition of fairness and how it is evaluated allows the introduction of correction mechanisms in the learning process, like those implemented in the FairTorch library (<https://fairtorch.github.io/FairTorch/> accessed on 13 January 2020). This way of approaching bias correction is close to what is known as *statistical bias*, as we have seen. We introduce the minimization of the bias as an additional constraint in the learning process.

Fairness is not a cognitive bias, this is something related to the estimation of parameters in statistical modeling, which is what neural networks do. Fairness is somehow the formalization of measures to reduce stereotyping in machine learning. According to Wikipedia ([https://en.wikipedia.org/wiki/Bias_\(statistics\)](https://en.wikipedia.org/wiki/Bias_(statistics)) accessed on 20 December 2020), a statistical bias is a feature of a statistical technique or of its results whereby the expected value of the results differs from the true underlying quantitative parameter being estimated. Fairness measures are, actually, measures of a statistical bias.

Therefore, whenever a *protected* feature is clearly identified or can be derived from sample features in the training set, it is possible to evaluate the model on its equity for generating similar distributions of predictions over groups resulting from different values of the protected feature. Even when in natural language processing many tasks can be defined in terms of machine learning, the challenge is when the protected attribute is not a clear feature in the dataset. How can we define bias/fairness when pre-training? How can we measure fairness over models like GPT-2 or BERT which have been trained following a language modeling approach? We propose an answer to these questions in the next section.

2.2. A Reflection on Bias in Language Models

A language model (LM) estimates the probability of a sequence of words $P(w_1, \dots, w_m)$. This allows for, given a sequence of words, estimating the next most probable word. The machinery behind the learning of model parameters can be used for solving many different tasks, like machine translation, text generation, text classification or token labeling (as for named entity recognition), among others [23]. Bias is present in language models as it is present in humans. Bias is intrinsic to human language, and it is not always source of unfairness. A car full of breakdowns is prone to accidents; fans of sci-fi movies are willing to watch similar movies; a patient with a chronic disease could have more risk of worsening, and so on. What we mark as “unfair” is established at a high semantic level. Remember that bias is not about prediction error, it is about skewed behavior regarding semantic expectations.

Definition 1. *The stereotyping bias in a language model is the undesired variation of the probability distribution of certain words in that language model according to certain prior words in a given domain.*

Those prior words are terms that can be linked to a protected attribute. Staying within the “gender” domain, those terms could be *actress*, *woman*, *girl*, etc. That is, in a language

model, we expect the distribution of probabilities after word *woman* to be equal (or very close) to that of the word *man* for certain words, like those related to professional skills. Both, *man* and *woman* are certain words in the *gender* domain (the protected attribute). It raises the problem of defining precisely the domain and those expected “certain” words. Following this example, the words within the gender domain would be split into two different classes where *stereotypes* are willing to occur, one class for men (actor, waiter...) and another class for women (actress, waitress...). Then, words regarding, in this case, attributes on professional skills (intelligence, efficiency, cleanness, creativity...) could be used to analyzed how they appear over the different probability distributions associated to each class, that is, when words in the domain are present, as priors of the distributions. Therefore, we could identify that the probability of the word “creativity” in the presence of a man is different from that of a woman.

This language modeling-based approach to the bias phenomena makes clear that bias is not a fault of the language model by itself, it is just the effect of the data from which this model was generated and of the desired behavior of the model at semantic level. Thus, is up to the language engineer to decide which domains and which expected distributions must be monitored or, eventually, corrected. To that end, stereotyped concepts must be identified within the domain and related attributes or concepts biased by those stereotypes must be selected. To overcome a clean definition of the bias problem, we propose an ontology-based approach, as the bias problem is firstly identified at a semantic level and, later on, treated at model-parameter level.

2.3. Definition of Bias at Semantic Level

Description logics [24] provides a complete set of elements for knowledge base structure, population and manipulation. It is, actually, the ontology formalization acquired by the Semantic Web and its high level ontological terminology OWL (Web Ontology Language) [25]. An OWL ontology has the following components: $\langle C, P, I, L \rangle$ classes C , properties P , individuals I and literal values L . For the sake of simplicity, we will summarize saying that individuals are instances of classes, instances are interrelated by properties and literals are associated to individuals by properties. For example, Christine is an individual which is of “type” Woman (belongs to class Woman). She works in a hospital (works-in would be a property). She has a job as a doctor (has-job would be another property). Woman is a class that can be defined through the expression has-gender “female” (this is called *class expression* in OWL), where has-gender is a property and “female” is a literal value. This simple knowledge can be graphically plotted as in Figure 1.

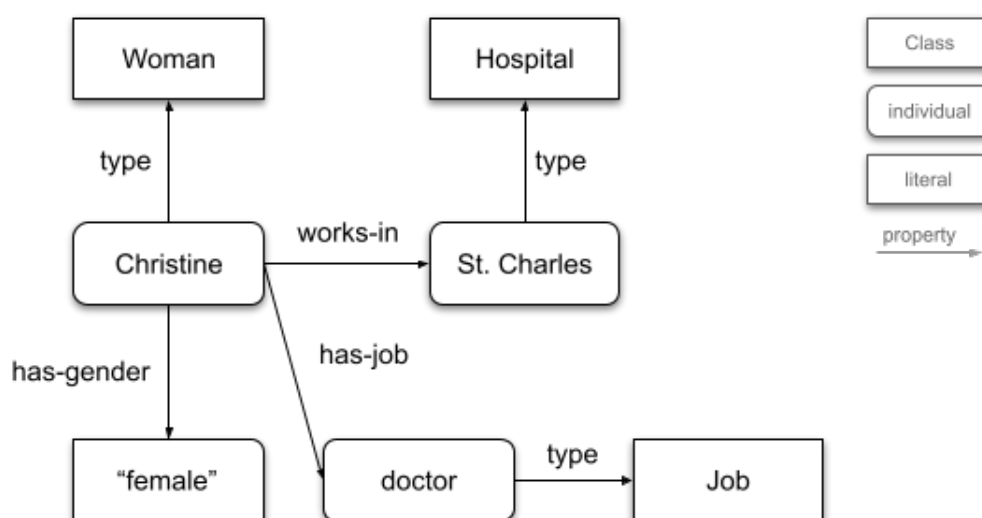


Figure 1. A simple example in Web Ontology Language (OWL).

Now it is time to borrow some terminology from fairness measures in machine learning and some elements from OWL.

Definition 2. A stereotyped knowledge is represented by the tuple $\langle C, P, I, L, p_p, P_s \rangle$ where C is the set of classes, P is the set of properties, I is the set of individuals, L the set of literals, $p_p \in P$ is the protected property and $P_s \subset P$ is the set of stereotyped properties. This expresses that groups of individuals resulting from different values of the protected property p_p could exhibit inequality in the distribution of values for stereotyped properties P_s .

In the example displayed in Figure 1, we could consider has-gender as the protected property p_p and $P_s = \{\text{has-job}\}$ as the set of stereotyped properties, so the tuple would be $\langle C, P, I, L, \text{has-gender}, \{\text{has-job}\} \rangle$. According to Definition 2, this means that values for property has-job could be not equally distributed over individuals of both classes defined by has-gender. For example, we may find that for individuals with has-gender ‘female’, it is more frequent to observe has-job nursery than has-job doctor, while the situation is the inverse for the class with individuals holding has-gender ‘male’. It is important to note that a *stereotyped knowledge* is only defining a potential bias, i.e., a bias we are sensible to.

2.4. Definition of Bias in Language Modeling

Once it is clear the semantic definition of the stereotyping bias, we can map that semantic identification down to word probabilities. This is straightforward, as a language model is nothing but a model able to compute a probability for a sequence of words $P(w_1, \dots, w_m)$.

Definition 3. A stereotyped language can be represented as the tuple $\langle C, P, I, L, p_p, P_s, T_p, T_s \rangle$, which contains a stereotyped knowledge and two terminology sets: protected terms T_p and stereotyped terms T_s . Protected terms T_p are those expressions (words or multi-words) in the vocabulary that can be unambiguously mapped to values of a protected property p_p . Stereotyped terms T_s are those expressions (words or multi-words) in the vocabulary that can be unambiguously mapped to values of stereotyped properties P_s .

T_s is the set of words or terms that represent possible values of stereotype properties P_s (for example, ‘high imagination’, ‘low sensibility’, ‘beauty’, ‘rational mind’, etc.). Examples of expressions in T_p would be any term defining gender, like ‘nurse’, ‘actress’, ‘woman’, ‘girl’ or alike. Once the stereotype is defined at the semantic level, we can consider that if the probability of a sequence of words containing expressions in T_s on stereotyped properties P_s is significantly different according to the value of p_p of the referenced individual, then the model is biased.

Now, we are ready for the final definition of stereotyping bias in language models.

Definition 4. Let $L_s = \langle C, P, I, L, p_p, P_s, T_p, T_s \rangle$ be the definition of a stereotyped language, stereotyping bias is defined as the distance d between probabilities $d(P(w_1, \dots, w_m | t_p^i), P(w_1, \dots, w_m | t_p^j))$, with $i \neq j$ where t_p^i and t_p^j are the expressions for two different values of the protected property p_p and $\exists w_k \in \{w_1, \dots, w_m\}$ so that $w_k \in T_s$.

In other words, a language model is biased if distributions of probabilities of terms containing stereotyped expressions are different subject to existing protected expression priors. Following our simple example, the *stereotyped language* could be defined as $\langle C, P, I, L, \text{has-gender}, \{\text{has-job}\}, \{\text{girl}, \text{women}, \text{Christine}, \text{man}\}, \{\text{doctor}, \text{nurse}\} \rangle$.

As key elements in the definition of a stereotype are p_p, P_s, T_p and T_s , we can focus on these four components to characterize it. In Table 1, further examples in different bias domains are proposed.

Table 1. Simple examples on stereotype formalization.

Protected Property (p_p)	Stereotyped Properties (P_s)	Protected Terms (T_p)	Stereotyped Terms (T_s)
has-gender	{has-job}	{girl, women, Christine, man}	{doctor, nurse}
has-age	{responsibility, efficiency}	$\{t \in L: t < 25\}, \{t \in L: t > 25\}$	{responsible, efficient, irresponsible, inefficient}
has-religion	{confidence, crime-committed}	{Muslim, Christian, atheist}	{terrorist, dangerous, robbery, homicide}
has-profession	{social-skills, intelligence, aspect}	{IT specialist, physicist, cleaning lady, politician, athlete, lawyer, CEO, teacher}	{empathetic, friendly, kind, confident, handsome, strong, intelligent, attractive, powerful, influencer}

Now, consider this simple text:

Christine works as a nurse in the hospital. A man is the doctor.

The definition is open to any kind of distance. If we select absolute difference, the stereotyping bias of the language model trained on the text above could be:

$$|P(\text{works, as, a, nurse} | \text{Christine}) - P(\text{works, as, a, doctor} | \text{Christine})|$$

Another valid measure would be:

$$|P(\text{works, as, a, nurse} | \text{man}) - P(\text{works, as, a, doctor} | \text{man})|$$

As you can see, different distances can be computed depending on the sequence or the prior value of the protected property considered. An appropriate evaluation of a language model would imply, therefore, a battery of expressions like the ones above, with protected expressions as priors and stereotyped expressions in the sequence, from which an average distance could be calculated.

3. Overview on Bias Related Research

An exhaustive review of relevant papers on bias in natural language processing has been carried out. In order to provide a global view into the different studies and analysis found regarding bias detection and correction, a set of elements have been identified to characterize major issues over all the works compiled. This allows for a organization of up-to-date research work on the targeted matter. These elements are now introduced for better understanding of the following tables as dimensions over the different main aspects in bias related research.

The studies are grouped according to their contributions in order to show the progress in the main lines of work. Each line constitutes one of the tables in this section.

- **Year.** This column is the publication year in ascending order and will serve as timeline on research progress. It also serves to highlight the increase of interest in the research community over time. We can see that it was not until two years after [26] when the community began to actively work on the bias of word embeddings models.
- **Reference** points to the publication.
- **Domain(s)** show us in which category falls the studied bias. The most represented category is gender bias, usually showing difference treatment between male and female. The second most represented one is ethnicity bias, in this category we grouped bias against race, ethnicity, nationality or language. We also found work on bias related with age, religion, sexual orientation and disability. It is worth mentioning that there is some work done on political bias.
- **Model** will refer to the neural network model studied in the paper. When the bias is not a model but an application, we will refer to such an application. Bias is not only studied in open system but also in black box applications like Google Translate. It is

interesting how some studies are able to discover and measure bias in those system. Although they are not able to mitigate the bias directly, there are some samples that manage to reduce the bias without having access to the model by modifying strategically the input.

- **Data** will serve as a summary of what data were used. We will consider almost all the resources that have taken part in the study regarding: the information used to train the models to the corpus on which the models is applied, or the other dataset that helps to contextualize the technique used.
- **Language** column mainly shows that most of the work has been done on English datasets and models. Some approaches when working with bias in other languages usually have English as a reference point, involving the translation of the data or test sets from English to other languages with both automated tools and paid professionals. Another approach involves looking for analogies between different languages.
- **Evaluation** column shows the reader which was the technique for evaluation or for measuring the bias:
 - **Association tests.** The usage of association tests began with the appearance of WEAT tests by Caliskan et al. [10] based on a study outside of the computer science field by Greenwald et al. [11]. It aims to measure the strength of the connection between two words.
 - **Sentiment of association,** a common way to find biased terms is measuring the sentiment of sentences by changing just one word. The words that differ will belong to the two classes being compared. A term will be biased if one sentence has a strong negative sentiment regarding the complementary. This is also tested with text generation tasks where a given sentence start will produce a full sentence or text, just changing a word of each class.
 - **Analogies.** The use of analogies has been found useful to show the bias with simple examples. Word embeddings space is suited to this type of technique, as analogies can be studied from a geometric perspective.
 - **Representation.** The works that fall in this category compare the likelihood between two classes of the protected property. Some studies will consider the goal to achieve equal representation, but usually the likelihood of the classes is compared with real world data. For example, comparing the distribution of men and women in the United States for a occupation with the probability of a sentence to be completed with an attribute of each one of the genres. In this way, you can compare the model output representation with the demographic percentage.
 - **Accuracy.** It is common to find studies that measure accuracy in tasks like classification or prediction to find out how biased the model is. This is similar to the general approach in machine learning with fairness measures.
- **Mitigation** shows how the bias is removed or attenuated from the data or the model.
 - **Vector space manipulation** evolves from the work of Bolukbasi et al. [26] in which he proposes to find the vector representation of the gender to compensate for its deviation and equalize some terms with respect to the neutral gender. This technique is known as word embedding debiasing or hard debiasing. This proposal has been explored with substantial improvements to better capture the bias, trying to avoid causing a harm to the model.
 - **Data augmentation** by increasing the source corpus/data of the biases. For example, by adding examples that balance with respect to an attribute. Thus seeking to make the data represent that given attribute in a less biased way.
 - **Data manipulation** makes changes to the data to help the model capture a less biased reality. For example, removing named entities in such a way that the model cannot learn differences associated with named entities.
 - **Attribute protection** tries to prevent an attribute from containing bias. For this purpose, different techniques are used to manipulate the data, the model or

the training in order to avoid capturing information about that attribute. For example, if you remove proper names from phrases in a dataset and train a model, the model will not be able to associate proper names with other features such as jobs. If you train a model to analyze the sentiment of phrases and avoid proper nouns, the names will not have sentiment associated with them. You can find its application in the other techniques or as a combination of them. For example, eliminating proper names so that they do not capture gender information, duplicating all sentences that have gender (data modification) using the opposite gender (data augmentation) and finally training the model and manipulating it to eliminate the gender subspace (vector space manipulation).

- **Stage** column stands for mitigation stage, and indicates when the mitigation/bias correction work was done.
 - **Before.** Mostly altering or augmenting the source data to avoid bias or to balance the data that will be used in the model training.
 - **During/Train.** Changing the training process or fine tuning the model. For example by using a custom loss function.
 - **After.** Usually changing the model vector space after the learning stage.
- **Task.** This column outlines the field or scope in which the author is working. Since the appearance of [26], an important part of the studies will try to solve the novel problem of both “debiasing” and “bias evaluation”. Since both tasks are already reported in columns of the table itself, they will not appear in this column.

4. Discussion

Tables 1–6 introduce the detailing key aspects in bias research according to the dimensions identified, a deeper analysis of all this prolific research production is carried out now. We have divided the discussion into salient topics in the following.

4.1. Association Tests

There are several approaches to bias measurement and mitigation. Bolukbasi et al. [26] laid the foundations for much of the work that was to follow. The main contribution was on showing that embeddings captured the correlation between terms so that they could correctly resolve analogies such as *man:king* → *woman:queen*, but also some similar analogies were biased. For example, it associates *man:doctor* and *woman:nurse* while the association *woman:doctor* would be more adequate. Using this same mechanism, it was obtained a set of terms that were stereotyped to each gender, to prove that this was not an isolated case. To remove that bias, they proposed to find the gender vector subspace direction and adjust the vector to make the occupational terms gender-neutral.

Caliskan et al. [10] took the idea of measuring bias of using the Implicit Association Test and proposed the Word Embedding Association Test (WEAT). WEAT measures the similarity of words by using the cosine between the pair of vectors of those words. It was applied to GloVe (Global Vectors for word representation) [27] and also to Word2Vec (word vectors) [28] with very similar findings. Other extension to WEAT was proposed by Lauscher and Glavaš [29], Lauscher et al. [30] with the name XWEAT, a cross lingual extension for WEAT. XWEAT was later extend to Arab Lauscher et al. [31]. WEAT can also be applied to other models. Gonen and Goldberg [32] applied it to the gender neutral version of GloVe called GN-GloVe from Zhao et al. [33]. Jentzsch et al. [34] use it with a skip-gram network in the context of Question Answering and Decision Making and [35] creates an algorithm to discover offensive association related with gender, race and other attributes, generating WEAT tests for them. They called this technique Unsupervised Bias Enumeration (UBE). UBE is applied to Word2Vec, fastText and GloVe. Dev and Phillips [36] propose two complementary tests that measure the bias removal effect (ECT, EQT).

Table 2. Previous work on bias detection and treatment in natural language processing (NLP): vector space manipulation.

Year	Ref.	Stereotype(s)	Model	Data	Lang.	Evaluation	Mitigation	Stage	Task
2016	[26]	Gender	Word2Vec, GloVe	GoogleNews corpus (w2vNEWS), Common Crawl	English	Analogies/Cosine Similarity	Vector Space Manipulation	After	-
2018	[37]	Gender	GloVe [27]	OntoNotes 5.0, WinoBias , Occupation Data (BLS), B&L	English	Prediction Accuracy	Data Augmentation (Gender Swapping), Vector Space Manipulation	After	Coreference Resolution
2018	[33]	Gender	GloVe [27], GN-GloVe , Hard-GloVe	2017 English Wikipedia dump, SemBias (3)	English	Prediction Accuracy, Analogies (3)	Attribute Protection, Vector Space Manipulation(1), Hard-Debias (2)	Train (1), After (2)	Coreference resolution
2019	[38]	Ethnicity, Gender, Religion	Word2Vec	Reddit L2 corpus	English	PCA, WEAT, MAC, Clustering	Vector Space Manipulation	After	POS tagging, POS chunking, NER
2019	[39]	Gender	Spanish fastText	Spanish Wikipedia, bilingual embeddings (MUSE) [40]	English, Spanish	CLAT , WEAT	Vector Space Manipulation	After	-
2019	[41]	Gender	Transformer, GloVe, Hard-Debiased GloVe, GN-GloVe	United Nations [42], Europarl [43], newstest2012, newstest2013, Occupation data (BLS)	English, Spanish	BLEU [44]	Vector Space Manipulation (Hard-Debias)	Train, After	Translation
2020	[30]	*	CBOW, GloVe, FastText, DebiasNet	-	German, Spanish, Italian, Russian, Croatian, Turkish, English	WEAT, XWEAT, ECT, BAT, Clustering (KMeans) (BIAS ANALOGY TEST)	Vector Space Manipulation, DEBIE	After	-
2019	[45]	Gender	BERT (base, uncased), GPT-2 (small)	-	English	Visualization, Text Generation likelihood	-	-	-
2020	[46]	Gender	RoBERTa/GloVe (1)	Common Crawl (1)	English	WEAT* , SIRT	Vector Space Manipulation, OSCaR	Train	-
2020	[47]	Gender	BERT	Equity Evaluation Corpus, Gen-data	English	EEC, Gender Separability. Emotion/Sentiment Scoring	Vector Space Manipulation	Train	-

Table 3. Previous work on bias detection and treatment in NLP: measuring association (Part 1/2).

Year	Ref.	Stereotype(s)	Model	Data	Lang.	Evaluation	Mitigation	Stage	Task
2017	[10]	Gender, Ethnicity	GloVe, Word2Vec	Common Crawl, Google News Corpus, Occupation Data (BLS)	English	Association Tests (WEAT , WEFAT)	-	-	-
2019	[32]	Gender	HARD-DEBIASED [26], GN-GloVe [33]	Google News, English Wikipedia	English	WEAT, Clustering	-	-	-
2019	[34]	Gender, Crime, Moral	Skip-Gram	Google's News	English	WEAT	-	-	Question answering, Decision making
2019	[35]	-	Word2Vec (1), fastText (2), GloVe (3)	Google News (1), Web data (2,3), First Names (SSA)	English	WEAT	-	-	Unsupervised Bias Enumeration
2019	[36]	Gender, Age, Ethnicity	GloVe	Wikipedia Dump, WSim-353, SimLex-999, Google Analogy Dataset	English	WEAT, EQT , ECT	Vector Space Manipulation	-	- -
2019	[38]	Ethnicity, Gender, Religion	Word2Vec	Reddit L2 corpus	English	PCA, WEAT, MAC, Clustering	Vector Space Manipulation	After	POS tagging, POS chunking, NER
2019	[29]	Gender	CBOW (1), GloVe (1,2), FastText (1), Dict2Vec(1)	English Wikipedia (1), Common Crawl (2), Wikipedia (2), Tweets (2)	English, German, Spanish, Italian, Russian, Croatian, Turkish	WEAT, XWEAT	-	-	-
2019	[39]	Gender	Spanish fastText	Spanish Wikipedia, bilingual embeddings (MUSE)[40]	English, Spanish	CLAT , WEAT	Vector Space Manipulation	After	-
2019	[48]	Gender	Skip-Gram (1,2,3), FastText (4)	Google News (1), PubMed (2), Twitter (3), GAP-Wikipedia (4) [49]	English	WEAT, Clustering (K-Means++)	-	-	-
2019	[50]	Gender, Ethnicity	BERT(large, cased), CBoW-GloVe (Web corpus version), InferSent, GenSen, USE, ELMo, GPT	-	English	SEAT	-	-	-
2019	[51]	Gender, Race	BERT(base cased, large cased), GPT-2 (117M, 345M), ELMo, GPT	-	English	Contextual SEAT	-	-	-

Table 3. Cont.

Year	Ref.	Stereotype(s)	Model	Data	Lang.	Evaluation	Mitigation	Stage	Task
2019	[52]	Gender	CBOW	English Gigaword, Wikipedia, Google Analogy, SimLex-999	English	Analogies, WEAT, Sentiment Classification, Clustering	Hard-Debiasing, CDA, CDS	Train	-
2020	[30]	*	CBOW, GloVe, FastText, DebiasNet	-	German, Spanish, Italian, Russian, Croatian, Turkish, English	WEAT, XWEAT, ECT, BAT, Clustering(KMeans) (BIAS ANALOGY TEST)	Vector Space Manipulation, DEBIE	After	-
2020	[31]	Gender, Ethnicity	AraVec CBOW (1), CBOW (2), AraVec Skip-Gram (3) and FASTTEXT (4), FastText (5)	translated WEAT test set, Leipzig news (2), Wikipedia (1,3,5), Twitter (1,3,4), CommonCrawl (5)	Modern Arabic. Egyptian Arabic	WEAT, XWEAT, AraWEAT , ECT, BAT	-	-	-

Table 4. Previous work on bias detection and treatment in NLP: measuring association (Part 2/2).

Year	Ref.	Stereotype(s)	Model	Data	Lang.	Evaluation	Mitigation	Stage	Task
2020	[46]	Gender	RoBERTa/GloVe (1)	Common Crawl (1)	English	WEAT* , SIRT	Vector Space Manipulation, OSCaR	Train	-
2020	[53]	Gender, Profession, Race, Religion	BERT, GPT-2, RoBERTa, XLNet	StereoSet	English	CAT Context Association Test	-	-	Language Modeling
2020	[54]	Intersectional Bias (Gender, Ethnicity)	GloVe, ELMo, GPT, GPT-2, BERT	CommonCrawl, Billion Word Benchmark, BookCorpus, English Wikipedia dumps, BookCorpus, WebText, Bert-small-cased?	English	WEAT, CEAT	-	-	-
2020	[55]	Gender	Word2Vec	Wikipedia-es 2006	Spanish	Analogies	-	-	-
2020	[56]	Gender	CBOW	British Library Digital corpus, The Guardian articles	English	Association, Prediction likelihood, Sentiment Analysis	-	-	-
2020	[57]	Gender	BERT	GAP, BEC-Pro , Occupation Data (BLS)	English, German	Association Test (like WEAT)	Fine-tuning, CDS	Train	-
2018	[58]	Gender	Deep Coref. [59]	WinoGender , Occupation Data (BLS), B&L	English	Prediction Accuracy	-	-	Coreference Resolution
2018	[33]	Gender	GloVe [27], GN-GloVe , Hard-GloVe	2017 English Wikipedia dump, SemBias (3)	English	Prediction Accuracy, Analogies (3)	Attribute Protection, Vector Space Manipulation (1), Hard-Debias (2)	Train (1), After (2)	Coreference resolution
2018	[60]	Gender	e2e-coref [61], deep-coref [62]	CoNLL-2012, Wikitext-2	English	Coreference score (1), likelihood (2)	Data Augmentation (CDA), WED [26],	Before, Train, After	Coreference Resolution (1), Language Modeling (2)
2019	[63]	Gender	ELMo, GloVe	One Billion Word Benchmark, WinoBias, OntoNotes 5.0	English	PCA, Prediction Accuracy	Data Augmentation (1), Attribute Protection (gender swapping averaging) (2)	Train (1), After (2)	Coreference Resolution

Table 4. Cont.

Year	Ref.	Stereotype(s)	Model	Data	Lang.	Evaluation	Mitigation	Stage	Task
2020	[14]	Gender, Race, Religion	GPT-3	Common Crow, WebText2, Books1, Books2, Wikipedia	English	Text generation	-	-	-
2020	[64]	Ideological, Political, Race	GPT-3	Common Crow, WebText2, Books1, Books2, Wikipedia	English	QA, Text Generation	-	-	-
2020	[65]	Race	GPT-3	Common Crow, WebText2, Books1, Books2, Wikipedia	English	Text Generation	-	-	Question Answering
2020	[66]	Gender	Google Translate	United Nations [42], Europarl [43], Google Translate Community	English, Hungarian	Prediction accuracy -	-	Translation	
2021	[16]	Ethnicity	GPT-3	Common Crow, WebText2, Books1, Books2, Wikipedia, Humans of New York images	English	Analogies, associations, Text Generation	Positive Contextualization	After	-

Table 5. Previous work on bias detection and treatment in NLP: data manipulation.

Year	Ref.	Stereotype(s)	Model	Data	Lang.	Evaluation	Mitigation	Stage	Task
2018	[37]	Gender	GloVe [27]	OntoNotes 5.0, WinoBias , Occupation Data (BLS), B&L	English	Prediction Accuracy	Data Augmentation (Gender Swapping), Vector Space Manipulation	After	Coreference Resolution
2018	[33]	Gender	GloVe [27], GN-GloVe , Hard-GloVe	2017 English Wikipedia dump, SemBias (3)	English	Prediction Accuracy, Analogies (3)	Attribute Protection, Vector Space Manipulation (1), Hard-Debias (2)	Train (1), After (2)	Coreference resolution
2019	[67]	Gender, Ethnicity, Disability, Sexual Orientation	Google Perspective API	WikiDetox, Wiki Madlibs, Twitter, WordNet	English	Classification Accuracy, likelihood	Data correction, Data Augmentation, Attribute Protection	Before	Hate Speech Detection
2019	[68]	Gender	fastText, BoW, DRNN with Custom Dataset	Common Crawl, Occupation Data (BLS)	English	Prediction Accuracy	Attribute protection (Removing Gender and NE)	Before	Hiring
2019	[69]	Account Age, user features	Graph Embeddings [70]	WikiData	English	Accuracy	Attribute Protection (Remove user information)	Train	Vandalism Detection
2019	[38]	Ethnicity, Gender, Religion	Word2Vec	Reddit L2 corpus	English	PCA, WEAT, MAC, Clustering	Vector Space Manipulation	After	POS tagging, POS chunking, NER
2019	[63]	Gender	ELMo, GloVe	One Billion Word Benchmark, WinoBias, OntoNotes 5.0	English	PCA, Prediction Accuracy	Data Augmentation (1), Attribute Protection (gender swapping averaging) (2)	Train (1), After (2)	Coreference Resolution
2019	[52]	Gender	CBOW	English Gigaword, Wikipedia, Google Analogy, SimLex-999	English	Analogies, WEAT, Sentiment Classification, Clustering	Hard-Debiasing, CDA, CDS	Train	-
2020	[71]	Ethnicity	GPT-2	Science fiction story corpus, Plotto, ROCstories, toxic and Sentiment datasets	English	Classification Accuracy	Loss function modification	Fine tuning	Normative text Classification

Table 6. Previous work on bias detection and treatment in NLP: task-specific accuracy/scoring.

Year	Ref.	Stereotype(s)	Model	Data	Lang.	Evaluation	Mitigation	Stage	Task
2018	[72]	Gender, Ethnicity	-	EEC, Tweets (SemEval-2018)	English	Sentiment, Emotion of Association	-	-	Sentiment Scoring
2019	[63]	Gender	ELMo, GloVe	One Billion Word Benchmark, WinoBias, OntoNotes 5.0	English	PCA, Prediction Accuracy	Data Augmentation (1), Attribute Protection(gender swapping averaging) (2)	Train (1), After (2)	Coreference Resolution
2019	[67]	Gender, Ethnicity, Disability, Sexual Orientation	Google Perspective API	WikiDetox, Wiki Madlibs, Twitter, WordNet	English	Classification Accuracy, likelihood	Data correction, Data Augmentation, Attribute Protection	Before	Hate Speech Detection
2019	[73]	Gender	Google Translate API (1)	United Nations and European Parliament transcripts (1), Translate Community (1), Occupation Data (BLS), COCA	Malay, Estonian, Finish, Hungarian, Armenian, Bengali, English, Persian, Nepali, Japanese, Korean, Turkish, Yoruba, Swahili, Basque, Chinese	Prediction accuracy -	-	Translation	
2019	[74]	Race, Gender, Sexual Orientation	LSTM, BERT, GPT-2 (small), GoogleLM1b (4)	One Billion Word Benchmark(4)	English	Sentiment Score (VADER [75]), Classification accuracy	Train LSTM/BERT	Train	Text Generation
2019	[76]	Gender	Google Translate, Microsoft Translator, Amazon Translate, SYSTRAN, Model of [77]	-	English, French, Italian, Russian, Ukrainian, Hebrew, Arabic, German	WinoMT (WinoBias + WinoGender), Prediction Accuracy	Positive Contextualization	After	Translation

Table 6. Cont.

Year	Ref.	Stereotype(s)	Model	Data	Lang.	Evaluation	Mitigation	Stage	Task
2019	[78]	Gender	ELMo	English-German news WTM18	English	cosine similarity, clustering, KNN	-	-	-
2019	[52]	Gender	CBOW	English Gigaword, Wikipedia, Google Analogy, SimLex-999	English	Analogies, WEAT, Sentiment Classification, Clustering	Hard-Debiasing, CDA, CDS	Train	-
2020	[47]	Gender	BERT	Equity Evaluation Corpus, Gen-data	English	EEC, Gender Separability, Emotion/Sentiment Scoring	Vector Space Manipulation	Train	-
2020	[79]	Ethnicity	GPT-2 (small), DISTILBERT	TwitterAAE [80], Amazon Mechanical Turk annotators (SAE)	English (AAVE/SAE)	Text generation, BLEU, ROUGE, Sentiment Classification, VADER [75]	-	-	-
2020	[56]	Gender	CBOW	British Library Digital corpus, The Guardian articles	English	Association, Prediction likelihood, Sentiment Analysis	-	-	-
2020	[81]	Gender, Race, Religion, Disability	BERT(1)	Wikipedia(1), Book corpus(1), Jigsaw identity toxic dataset, RtGender, GLUE	English	Cosine Similarity, Accuracy, GLUE	Fine tuning	Fine tuning	Decision Making
2020	[82]	Gender, Race	SqueezeBERT	Wikipedia, BooksCorpus	English	GLUE	-	-	-
2020	[83]	Disability	BERT, Google Cloud sentiment model	Jigsaw Unintended Bias	English	Sentiment Score	-	-	Toxicity prediction, Sentiment analysis.

Manzini et al. [38] extend WEAT to measure the bias in a multi-class setting and use it over a Word2Vec model trained with Reddit L2 corpus. Dev et al. [46] adapted it to work with two sets of words at a time instead of just two words, naming its variant WEAT*.

The appearance of models such as BERT [84] led to the adaptation of the technique to work at phrase level (SEAT May et al. [50]) and to work with contextualized embeddings in Guo and Çalışkan [54] named CEAT. CEAT was tested on BERT, GPT, GPT-2 and EIMo.

As part of the association-based bias study, Nadeem et al. [53] presents StereoSet and evaluates BERT, BPT-2, RoBERTa, XLNET models. For the evaluation, it confronts three terms in the same context, one stereotyped, one anti-stereotyped and one unrelated term. It measures the probability that a sentence is completed with each of them. In the sentences there is a token which is the one against which we measure the bias. This technique is called Contextual Association Test.

From this test, the sentiment associated with stereotyped and non-stereotyped sentences can be analyzed. Measuring the sentiment of an association to quantify bias is not new, it can be found in the work of Kiritchenko and Mohammad [72] where he evaluates race and gender bias. Sheng et al. [74] further measure bias associated with sexual orientation by comparing the associated sentiment. We also have the extensive study by Leavy et al. [56] on CBOW trained on articles from The Guardian journal and the British Digital Library. In addition, Hutchinson et al. [83] studied the perception of models towards disabled people and Bhardwaj et al. [47] combined the study of gender bias on BERT by sentiment analysis with gender separability.

4.2. Translation

Previously, we have seen XWEAT for the detection of bias in languages other than English. Although there are also alternatives that work with multiple languages, one of the areas of study is what occurs when translating a text, such as the work of Escudé Font and Costa-jussà [41], which seeks and mitigates the bias in English-Spanish translations with the three versions of GloVe previously discussed (base, gender neutral, hard-debiased).

Not only has translation bias been studied in open models, but also the bias in final products such as Google translator, Microsoft translator, Amazon translator, among others, has been evaluated in the study of Stanovsky et al. [76].

Farkas and N'émeth [66] pose a mismatch when using Google Translator for translating from languages such as Hungarian with neutral gender into English. The inferred gender does not proportionally represent the actual distribution of workers when making inferences about professions, using Google Translator. This same mismatch appears in Google Translator in more languages, such as Hungarian, Chinese, Yoruba and others when translated into English. In this case, [73] shows a very strong correlation between the science, technology, engineering and mathematics (STEM) family of academic disciplines and men.

According to Davis [85], Google has fixed the problem it had with Google Translator inferring gender when translating from non-gendered languages into English. In the Google AI blog, Johnson [86] develops the first approach to the problem. This solution was put into production in 2018. They trained a CNN with human categorized examples and further divided the training set into three chunks, one for masculine another feminine and another for neutral. To the sentences of each chunk they added in front a token of the type "<2MALE>". Therefore, "<2MALE> O birt doktor" would translate to "HE is a doctor". Allowing this to use all 3 prefixes with the user input to give an unbiased response. This resulted in a recall of 60%.

The next approach would come in 2020. Johnson [87] would firstly translate the phrase obviating the gender and secondly, would look for occurrences of the translated phrase from the same query but with the complementary gender. If only the gender changes when compared to the original translation, the phrase is returned to the user on both genders.

4.3. Coreference Resolution

Two of the first studies for gender bias were published as part of the Proceedings of the 2018 Conference on Empirical Methods in Natural Language Processing. The first of Zhao et al. [37] proposes WinoBias, a balanced Male/Female dataset for the evaluation of gender bias in Coreference Resolution tasks. The second, with a similar title, was that of Rudinger et al. [58] who introduced the WinoGender schemes for the study of bias also in Coreference Resolution tasks. Later, Stanovsky et al. [76] combined both resources to study the bias in Machine Translation, thus creating WinoMT.

The study of bias in coreference resolution does not stop here, Zhao et al. [33] studied gender bias in GloVe and develops 2 derived models, GN-GloVe (gender neutral) and HD-GloVe (hard-debiased). Lu et al. [60] tried to reduce the detected bias by using the data augmentation technique Contextual Data Augmentation CDA which consists of adding a complementary gender phrase to the sentences of the initial dataset. Based on CDA, in 2019 Hall Maudslay et al. [52] will develop Contextual Data Substitution (CDS). It proposes to eliminate the bias associated with proper names by adding a phrase with a complementary gender name in a balanced way. CDS will later be used by Bartl et al. [57] together with fine-tuning for BERT.

4.4. GPT-3 and Black-Box Models

The recent GPT-3 also seems to suffer from bias. Actually, in the very study that presents the model Brown et al. [14] already address the issue for gender, race and religion. The authors themselves discover an important tendency between terms such as violent, terrorism and terrorist with Islam. It will also be studied in Decision Making and Question Answering tasks by McGuffie and Newhouse [64], which will show how GPT-3 is better than GPT-2 at generating extremist narratives and suggest that it could be used for the radicalization of individuals. For the study, questions are asked to GPT-3 on specific topics and its responses are studied.

The alternative to studying bias in this model by asking questions is through the model's ability to generate text from the beginning of a given sentence that will serve as a context to the model. Floridi and Chiriatti [65] studied the model in this way and found that although the model is able to complete sentences and text, it lacks perspective or intelligence when dealing with topics.

Abid et al. [16] make a valuable contribution, finding that it is possible to alleviate the bias in the responses and text generated by GPT-3 by trying to guide the response with a positive context. If instead of asking the model to complete a sentence referring to Muslims, you should add a positive adjective such as hard-working or meticulous. This way the model's responses will move away from topics such as violence.

All studies on GPT-3 are conducted as a black-box model since it has not been released. This is why its web interface or API is used for its study.

4.5. Vector Space

The main debiasing techniques try to eliminate model bias. Different approaches are used to find the direction of the gender as proposed by Bolukbasi et al. [26] and try to correct the deviation between classes. The simplest papers define techniques to find the gender direction and adjust it between male/female pairs. Some, such as Zhou et al. [39], propose that there is not one but two gender directions given the characteristics of Spanish, considering one direction as semantic and the other as grammatical. In such a way that words like perfume in Spanish are masculine but strongly associated with the feminine gender, so it will try to eliminate the bias by considering both components. This is why, in our formal definition of bias, "stereotyped expressions" is preferred, rather than just mentioning words or isolated terms.

Gonen and Goldberg [32] suggest that the debiasing techniques that work with gender direction are not sufficient and that the bias is only superficially eliminated. There are

multiple approaches that try to improve by trying to identify gender as a space rather than as a direction, such as the work of Basta et al. [78] on ElMo.

The previously cited techniques are also extrapolated to try to tackle the problem in other languages. Zhou et al. [39] suggest that for Spanish there is not one gender direction but two: grammar and semantics. A term like “perfume” is semantically more masculine but is grammatically more strongly associated with the feminine gender. Therefore, gender measurement and mitigation will have to seek to balance between these two dimensions. He also proposes cross lingual analogy tasks (CLAT) to assess bias in Spanish. Given a pair $a:b$ in English and a word c in Spanish, the Spanish term associated with d must be predicted for $a:b = c:d$.

Alternatively, Díaz Martínez et al. [55] launched a proposal similar to Bolukbasi et al. [26] but for Spanish. It detects that there are indeed terms strongly associated to one of the genres in a Word2Vec model trained with Wikipedia articles.

4.6. Deep Learning Versus Traditional Machine Learning Algorithms

The bias problem is present throughout machine learning approaches. As long as model is trained on biased information, the model will adopt those deviations in its learned parameters. This is the expected behavior of a machine learning algorithm: to mimic reality by identifying patterns in data. Why differentiate bias in deep neural networks from previous algorithms? We consider that there are strong reasons that emphasize the problem when dealing with deep language models.

1. Deep models are very greedy of data. Large datasets must be fed into the learning process. Gigabytes (and even terabytes) of texts are consumed during the training process. Most of the models are pre-trained with a language model approach (masking, sentence sequence) over corpora generated from available sources in the Internet, like movie subtitles, Wikipedia articles, news, tweets and so on. As these texts come from open communities but with specific cultural profiles from western world, bias expressions are naturally present in such a collection of texts.
2. As the demand of larger corpora for larger models grows, the bias digs deeper its footprint. Some of the most powerful architectures, like GPT-3, are examples of such a situation. Bigger models, bigger bias (and more stereotype patterns found).
3. Intensive research is being carried out in order to compress models, so the fit into environments with limited memory, latency and energy capabilities. Quantization, pruning or teacher–student approaches are enabling deep learning models to operate in more restricted infrastructures at a negligible cost in terms of performance [88]. It has been found that bias is, far from expected, emphasized by this reduction methods [89].

The rise of deep learning algorithms has promoted the study of stereotypes in language models. Therefore, specific research on bias for this type of architectures is populating scientific production, becoming a major issue nowadays.

4.7. Complementary Works

There are similar approaches that show that it is possible to detect gender bias in models such as GPT-2 [15]. Vig [45] reviews the interior of these networks and evidences the strong connections between “she, nurse” and “he, doctor” and suggests that it would be possible to detect and control it. For all this, he relies on a tool that allows to visualize the interior of transformer networks such as BERT [84] or GPT-2.

5. A General Methodology for Dealing with Bias in Deep NLP

Up to this point, it is possible to conclude that the bias problem is of relevance for industrial deployments of artificial intelligence solutions. When putting a language model into production for a defined task like classification, dialogue or whatsoever, the engineer has to ensure that no bias could affect the expected behavior of the system so future troubles due to stereotyped decisions are prevented. This paper has made an effort in showing

the state-of-the-art in bias related research, specially on deep learning models for natural language processing. In addition, a clear definition of this phenomena has been provided, detailing all the elements involved in spotting the bias with precision and completeness.

In this section, we propose the use of all those elements to help the engineer to identify them in the subject of her study and to follow a structured method to tackle it in a software engineering process. With that purpose in mind, the following steps are proposed as a general methodology for dealing with stereotyping bias in deep language models generation and application:

1. Define the **stereotyped knowledge**. This implies to identify one or more protected properties and all the related stereotyped properties. For each protected property, you have to develop its own ontology.
2. With the previous model at hand, we can overcome the task of identifying **protected expressions** and **stereotyped expressions**, so your **stereotyped language** is defined. It is equivalent to the process of *populating your ontology* (i.e., your stereotyped knowledge). There are some corpora available, like the ones mentioned in this work, but you may need to define your own expressions in order to capture all the potential biases that may harm your system. Anyhow, it is here when different resources could be explored to obtain a set of expressions as rich as possible.
3. The next step is to **evaluate your model**. Choose a distance metric and compute overall differences in sequence probabilities containing stereotyped expressions with protected expressions as priors. Detail the benchmarked evaluation framework used.
4. **Analyze** the results of the evaluation to identify which expressions or categories of expressions result in higher bias.
5. Design a corrective mechanism. You have to decide which strategy fits better with your problem and with your available resources: data augmentation, a constraint in the learning process, model parameters correction, etc.
6. **Re-evaluate** your model and loop over these last three steps until an acceptable response is reached, or though out your model if behavior is not what is desired. Rethink the whole process (network architecture, pre-training approach, fine-tuning, etc.).
7. Report the result of this procedure by attaching model cards or similar document formalism in order to achieve **transparent model reporting**.

Following these steps may help in getting a final system you understand better and with predictions not affected or marginally affected by stereotyping bias. For sure, this method can be adapted or extended according to the requirements of each specific AI project.

6. Conclusions and Challenges

In this work, we focused on deep NLP techniques and how these techniques are affected by bias as a consequence of the advent of more challenge data sets and methods. We found that gender bias for English language when using word embedding related technologies is the most frequent scenario that is faced in those methods developed to mitigate bias in different tasks. This can be achieved in three different ways: by modifying the training corpora, the training algorithm or the results obtained according to the given task. We proposed to systematize the evaluation of the impact of bias as part of the design of systems relying on deep NLP techniques and resources. The focus of the proposed procedure is the identification and management of stereotyped expressions apart from protected expressions, both concepts introduced in Section 2.3. As future challenges, apart from digging deeper in the detection and softening of bias, it is our view that there are some aspects that deserve more attention than given nowadays. The first one is related with the effect of bias mitigation in both the global system performance and the management of other terms and features different from stereotyped expressions. Is it possible that the main task to be solved by deep NLP systems could be damaged by the intervention to mitigate stereotyped expressions? In the same way, we propose to study the impact of a

preventive strategy rather than a corrective one. That is, in the case of having transparent language models (i.e., accompanied by model reports), we consider measuring how the choice of different language models that are free of bias compared to those that do present some degree of bias, affects the final performance of the system. In any case, although it is clear the fact that there is no biased algorithms but biased corpora and language models, there is little effort in describing characteristics of corpora and making transparent language models by means of the inclusion of model reporting, related with demographic or phenotypic groups, environmental conditions, instrumentation or environment, *inter alia*. As a consequence, it is needed further effort to characterize, to make transparent the language model or corpora to be chosen regarding a given task.

Another interesting approach would be to apply the techniques studied to systems in production and perform different measurements that allow us to know the impact of the changes made on the model. Applying this work to real applications will allow us to see if the changes are really effective, to see how they affect other aspects on the application's performance and, above all, to discover which aspects have not been taken into account.

As a matter of engineering processes, resources should be put on the focus of the problem. Additional benchmarks and tests for different stereotypes over different languages are, in our opinion, in the way to a consistent management of biases for final applications.

Author Contributions: Conceptualization, I.G.-M., A.M.-R., F.M.-S. and L.A.U.-L., contextualization, F.M.-S.; formalism and methodology, A.M.-R., research overview, I.G.-M.; analysis, I.G.-M., A.M.-R. and F.M.-S. All authors have read and agreed to the published version of the manuscript.

Funding: This study is partially funded by the Spanish Government under the LIVING-LANG project (RTI2018-094653-B-C21).

Institutional Review Board Statement: Not applicable.

Informed Consent Statement: Not applicable.

Conflicts of Interest: The authors declare no conflict of interest.

Abbreviations

The following abbreviations are used in this manuscript:

AAVE	African American Vernacular English
AI	Artificial Intelligence
BAT	Bias Analogy Test
BERT	Bidirectional Encoder Representations from Transformers
BLS	Bureau of Labor Statistics
CAT	Context Association Test
CDA	Counterfactual Data Augmentation
CDS	Counterfactual Data Substitution
CEAT	Contextualized Embedding Association Test
CLAT	Cross-lingual Analogy Task
ECT	Embedding Coherence Test
EQT	Embedding Quality Test
GPT	Generative Pre-Training Transformer
IAT	Implicit Association Tests
LM	Language Model
LSTM	Long-Short Term Memory
NER	Named Entities Recognition
NLP	Natural Language Processing
PCA	Principal Component Analysis
POS	Part of Speech
SAE	Standard American English

SEAT	Sentence Encoder Association Test
SIRT	Sentence Inference Retention Test
UBE	Universal Bias Encoder
USE	Universal Sentence Encoder
WEAT	Word Embedding Association Test
WEFAT	Word Embedding Factual Association Test
XWEAT	Multilingual and Cross-Lingual WEAT

References

- Howard, A.; Borenstein, J. Trust and Bias in Robots: These elements of artificial intelligence present ethical challenges, which scientists are trying to solve. *Am. Sci.* **2019**, *107*, 86–90. [\[CrossRef\]](#)
- Rodger, J.A.; Pendharkar, P.C. A field study of the impact of gender and user's technical experience on the performance of voice-activated medical tracking application. *Int. J. Hum. Comput. Stud.* **2004**, *60*, 529–544. [\[CrossRef\]](#)
- Bullinaria, J.A.; Levy, J.P. Extracting semantic representations from word co-occurrence statistics: A computational study. *Behav. Res. Methods* **2007**, *39*, 510–526. [\[CrossRef\]](#) [\[PubMed\]](#)
- Stubbs, M. *Text and Corpus Analysis: Computer-Assisted Studies of Language and Culture*; Blackwell: Oxford, UK, 1996.
- Mehrabi, N.; Morstatter, F.; Saxena, N.; Lerman, K.; Galstyan, A. A Survey on Bias and Fairness in Machine Learning. *arXiv* **2019**, arXiv:1908.09635.
- Obermeyer, Z.; Powers, B.; Vogeli, C.; Mullainathan, S. Dissecting racial bias in an algorithm used to manage the health of populations. *Science* **2019**, *366*, 447–453. [\[CrossRef\]](#)
- Kendall, A.; Gal, Y. What Uncertainties Do We Need in Bayesian Deep Learning for Computer Vision? *arXiv* **2017**, arXiv:1703.04977.
- Tolan, S.; Miron, M.; Gómez, E.; Castillo, C. Why machine learning may lead to unfairness: Evidence from risk assessment for juvenile justice in catalonia. In Proceedings of the Seventeenth International Conference on Artificial Intelligence and Law, Montreal, QC, Canada, 17–21 June 2019; pp. 83–92.
- Xu, J.; Ju, D.; Li, M.; Boureau, Y.L.; Weston, J.; Dinan, E. Recipes for Safety in Open-domain Chatbots. *arXiv* **2020**, arXiv:2010.07079.
- Caliskan, A.; Bryson, J.; Narayanan, A. Semantics derived automatically from language corpora contain human-like biases. *Science* **2017**, *356*, 183–186. [\[CrossRef\]](#)
- Greenwald, A.G.; McGhee, D.E.; Schwartz, J.L. Measuring individual differences in implicit cognition: The implicit association test. *J. Personal. Soc. Psychol.* **1998**, *74*, 1464. [\[CrossRef\]](#)
- Mitchell, M.; Wu, S.; Zaldivar, A.; Barnes, P.; Vasserman, L.; Hutchinson, B.; Spitzer, E.; Raji, I.D.; Gebru, T. Model Cards for Model Reporting. In Proceedings of the Conference on Fairness, Accountability, and Transparency, New York, NY, USA, 23–24 February 2018.
- Bender, E.M.; Friedman, B. Data Statements for Natural Language Processing: Toward Mitigating System Bias and Enabling Better Science. *Trans. Assoc. Comput. Linguist.* **2018**, *6*, 587–604. [\[CrossRef\]](#)
- Brown, T.B.; Mann, B.; Ryder, N.; Subbiah, M.; Kaplan, J.; Dhariwal, P.; Neelakantan, A.; Shyam, P.; Sastry, G.; Askell, A.; et al. Language Models are Few-Shot Learners. *arXiv* **2020**, arXiv:2005.14165.
- Radford, A.; Wu, J.; Child, R.; Luan, D.; Amodei, D.; Sutskever, I. Language models are unsupervised multitask learners. *OpenAI Blog* **2019**, *1*, 9.
- Abid, A.; Farooqi, M.; Zou, J. Persistent Anti-Muslim Bias in Large Language Models. *arXiv* **2021**, arXiv:2101.05783.
- Kahneman, D.; Tversky, A. On the psychology of prediction. *Psychol. Rev.* **1973**, *80*, 237. [\[CrossRef\]](#)
- Gigerenzer, G. Bounded and rational. In *Philosophie: Grundlagen und Anwendungen/Philosophy: Foundations and Applications*; Mentis: Paderborn, Germany, 2008; pp. 233–257.
- Haselton, M.G.; Nettle, D.; Murray, D.R. The evolution of cognitive bias. *The Handbook of Evolutionary Psychology*; John Wiley & Sons: Hoboken, NJ, USA, 2015; pp. 1–20.
- Schneider, D.J. *The Psychology of Stereotyping*; Guilford Press: New York, NY, USA, 2005.
- Gajane, P.; Pechenizkiy, M. On Formalizing Fairness in Prediction with Machine Learning. *arXiv* **2017**, arXiv:1710.03184.
- Verma, S.; Rubin, J. Fairness definitions explained. In Proceedings of the 2018 IEEE/ACM International Workshop on Software Fairness (Fairware), Gothenburg, Sweden, 29 May 2018; pp. 1–7.
- Qiu, X.; Sun, T.; Xu, Y.; Shao, Y.; Dai, N.; Huang, X. Pre-trained models for natural language processing: A survey. *Sci. China Technol. Sci.* **2020**, *63*, 1872–1897. [\[CrossRef\]](#)
- Baader, F.; Horrocks, I.; Sattler, U. Description logics. In *Handbook on Ontologies*; Springer: Berlin/Heidelberg, Germany, 2004; pp. 3–28.
- Antoniou, G.; Van Harmelen, F. Web ontology language: Owl. In *Handbook on Ontologies*; Springer: Berlin/Heidelberg, Germany, 2004; pp. 67–92.
- Bolukbasi, T.; Chang, K.W.; Zou, J.Y.; Saligrama, V.; Kalai, A. Man is to Computer Programmer as Woman is to Homemaker? Debiasing Word Embeddings. *arXiv* **2016**, arXiv:1607.06520.
- Pennington, J.; Socher, R.; Manning, C.D. GloVe: Global Vectors for Word Representation. In Proceedings of the 2014 Conference on Empirical Methods in Natural Language Processing (EMNLP), Doha, Qatar, 25–29 October 2014; pp. 1532–1543.

28. Mikolov, T.; Chen, K.; Corrado, G.; Dean, J. Efficient Estimation of Word Representations in Vector Space. *arXiv* **2013**, arXiv:1301.3781.
29. Lauscher, A.; Glavaš, G. Are We Consistently Biased? Multidimensional Analysis of Biases in Distributional Word Vectors. In Proceedings of the Eighth Joint Conference on Lexical and Computational Semantics (*SEM 2019), Minneapolis, MN, USA, 6–7 June 2019; pp. 85–91. [\[CrossRef\]](#)
30. Lauscher, A.; Glavas, G.; Ponzetto, S.P.; Vulic, I. A General Framework for Implicit and Explicit Debiasing of Distributional Word Vector Spaces. *arXiv* **2020**, arXiv:1909.06092.
31. Lauscher, A.; Takieddin, R.; Ponzetto, S.P.; Glavaš, G. AraWEAT: Multidimensional Analysis of Biases in Arabic Word Embeddings. In Proceedings of the Fifth Arabic Natural Language Processing Workshop, Barcelona, Spain, 12 December 2020; pp. 192–199.
32. Gonen, H.; Goldberg, Y. Lipstick on a Pig: Debiasing Methods Cover up Systematic Gender Biases in Word Embeddings But do not Remove Them. *arXiv* **2019**, arXiv:1903.03862.
33. Zhao, J.; Zhou, Y.; Li, Z.; Wang, W.; Chang, K.W. Learning Gender-Neutral Word Embeddings. In Proceedings of the 2018 Conference on Empirical Methods in Natural Language Processing, Brussels, Belgium, 31 October–4 November 2018; pp. 4847–4853. [\[CrossRef\]](#)
34. Jentzsch, S.; Schramowski, P.; Rothkopf, C.; Kersting, K. Semantics Derived Automatically from Language Corpora Contain Human-like Moral Choices. In Proceedings of the 2019 AAAI/ACM Conference on AI, Ethics, and Society (AIES'19), New York, NY, USA, 27–28 January 2019; pp. 37–44. [\[CrossRef\]](#)
35. Swinger, N.; De-Arteaga, M.; Heffernan, N.T., IV; Leiserson, M.D.; Kalai, A.T. What Are the Biases in My Word Embedding? In Proceedings of the 2019 AAAI/ACM Conference on AI, Ethics, and Society, New York, NY, USA, 27–28 January 2019; pp. 305–311. [\[CrossRef\]](#)
36. Dev, S.; Phillips, J. Attenuating Bias in Word vectors. In Proceedings of the The 22nd International Conference on Artificial Intelligence and Statistics, Okinawa, Japan, 16–18 April 2019; Volume 89, pp. 879–887.
37. Zhao, J.; Wang, T.; Yatskar, M.; Ordonez, V.; Chang, K. Gender Bias in Coreference Resolution: Evaluation and Debiasing Methods. *arXiv* **2018**, arXiv:1804.06876.
38. Manzini, T.; Yao Chong, L.; Black, A.W.; Tsvetkov, Y. Black is to Criminal as Caucasian is to Police: Detecting and Removing Multiclass Bias in Word Embeddings. In Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Minneapolis, MN, USA, 2–7 June 2019; Volume 1, pp. 615–621. [\[CrossRef\]](#)
39. Zhou, P.; Shi, W.; Zhao, J.; Huang, K.H.; Chen, M.; Chang, K.W. *Analyzing and Mitigating Gender Bias in Languages with Grammatical Gender and Bilingual Word Embeddings*; ACL: Montréal, QC, Canada, 2019.
40. Conneau, A.; Lample, G.; Ranzato, M.; Denoyer, L.; Jégou, H. Word Translation Without Parallel Data. *arXiv* **2018**, arXiv:1710.04087.
41. Escudé Font, J.; Costa-jussà, M.R. Equalizing Gender Bias in Neural Machine Translation with Word Embeddings Techniques. In Proceedings of the First Workshop on Gender Bias in Natural Language Processing, Florence, Italy, 1–2 August 2019; pp. 147–154. [\[CrossRef\]](#)
42. Ziemski, M.; Junczys-Dowmunt, M.; Pouliquen, B. The United Nations Parallel Corpus v1.0. In Proceedings of the Tenth International Conference on Language Resources and Evaluation (LREC'16), Portorož, Slovenia, 23–28 May 2016; pp. 3530–3534.
43. Koehn, P. Europarl: A Parallel Corpus for Statistical Machine Translation. 2005. Available online: <https://www.statmt.org/europarl/> (accessed on 1 October 2019).
44. Papineni, K.; Roukos, S.; Ward, T.; Zhu, W.J. Bleu: A Method for Automatic Evaluation of Machine Translation. In Proceedings of the 40th Annual Meeting of the Association for Computational Linguistics, Philadelphia, PA, USA, 7–12 July 2002; pp. 311–318. [\[CrossRef\]](#)
45. Vig, J. A Multiscale Visualization of Attention in the Transformer Model. *arXiv* **2019**, arXiv:1906.05714.
46. Dev, S.; Li, T.; Phillips, J.M.; Srikumar, V. OSCaR: Orthogonal Subspace Correction and Rectification of Biases in Word Embeddings. *arXiv* **2020**, arXiv:2007.00049.
47. Bhardwaj, R.; Majumder, N.; Poria, S. Investigating Gender Bias in Bert. *arXiv* **2020**, arXiv:2009.05021.
48. Chaloner, K.; Maldonado, A. Measuring Gender Bias in Word Embeddings across Domains and Discovering New Gender Bias Word Categories. In Proceedings of the First Workshop on Gender Bias in Natural Language Processing, Florence, Italy, 1–2 August 2019; pp. 25–32. [\[CrossRef\]](#)
49. Webster, K.; Recasens, M.; Axelrod, V.; Baldridge, J. Mind the GAP: A Balanced Corpus of Gendered Ambiguous Pronouns. *Trans. Assoc. Comput. Linguist.* **2018**, *6*, 605–617. [\[CrossRef\]](#)
50. May, C.; Wang, A.; Bordia, S.; Bowman, S.R.; Rudinger, R. On Measuring Social Biases in Sentence Encoders. In Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Minneapolis, MN, USA, 2–7 June 2019; Volume 1, pp. 622–628. [\[CrossRef\]](#)
51. Tan, Y.; Celis, L. Assessing Social and Intersectional Biases in Contextualized Word Representations. *arXiv* **2019**, arXiv:1911.01485.
52. Hall Maudslay, R.; Gonen, H.; Cotterell, R.; Teufel, S. It's All in the Name: Mitigating Gender Bias with Name-Based Counterfactual Data Substitution. In Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing (EMNLP-IJCNLP), Hong Kong, China, 3–7 November 2019; pp. 5267–5275. [\[CrossRef\]](#)

53. Nadeem, M.; Bethke, A.; Reddy, S. StereoSet: Measuring stereotypical bias in pretrained language models. *arXiv* **2020**, arXiv:2004.09456.
54. Guo, W.; Çalışkan, A. Detecting Emergent Intersectional Biases: Contextualized Word Embeddings Contain a Distribution of Human-like Biases. *arXiv* **2020**, arXiv:2006.03955.
55. Díaz Martínez, C.; Díaz García, P.; Navarro Sustaeta, P. Hidden Gender Bias in Big Data as Revealed by Neural Networks: Man is to Woman as Work is to Mother? *Rev. Esp. Investig. Sociol.* **2020**, *172*, 41–60.
56. Leavy, S.; Meaney, G.; Wade, K.; Greene, D. Mitigating Gender Bias in Machine Learning Data Sets. In *Bias and Social Aspects in Search and Recommendation*; Boratto, L., Faralli, S., Marras, M., Stilo, G., Eds.; Springer International Publishing: Cham, Switzerland, 2020; pp. 12–26.
57. Bartl, M.; Nissim, M.; Gatt, A. Unmasking Contextual Stereotypes: Measuring and Mitigating BERT's Gender Bias. In Proceedings of the Second Workshop on Gender Bias in Natural Language Processing, Florence, Italy, 12–13 December 2020.
58. Rudinger, R.; Naradowsky, J.; Leonard, B.; Van Durme, B. Gender Bias in Coreference Resolution. In Proceedings of the 2018 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, New Orleans, LA, USA, 1–6 June 2018; Volume 2, pp. 8–14. [\[CrossRef\]](#)
59. Clark, K.; Manning, C. Deep Reinforcement Learning for Mention-Ranking Coreference Models. *arXiv* **2016**, arXiv:1609.08667.
60. Lu, K.; Mardziel, P.; Wu, F.; Amancharla, P.; Datta, A. Gender Bias in Neural Natural Language Processing. *arXiv* **2018**, arXiv:1807.11714.
61. Lee, K.; He, L.; Lewis, M.; Zettlemoyer, L. End-to-end Neural Coreference Resolution. In Proceedings of the 2017 Conference on Empirical Methods in Natural Language Processing, Copenhagen, Denmark, 7–11 September 2017; pp. 188–197. [\[CrossRef\]](#)
62. Clark, K.; Manning, C.D. Deep Reinforcement Learning for Mention-Ranking Coreference Models. In Proceedings of the 2016 Conference on Empirical Methods in Natural Language Processing, Austin, TX, USA, 1–5 November 2016; pp. 2256–2262. [\[CrossRef\]](#)
63. Zhao, J.; Wang, T.; Yatskar, M.; Cotterell, R.; Ordonez, V.; Chang, K.W. Gender Bias in Contextualized Word Embeddings. In Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Minneapolis, MN, USA, 2–7 June 2019; Volume 1, pp. 629–634. [\[CrossRef\]](#)
64. McGuffie, K.; Newhouse, A. The Radicalization Risks of GPT-3 and Advanced Neural Language Models. *arXiv* **2020**, arXiv:2009.06807.
65. Floridi, L.; Chiriatti, M. GPT-3: Its Nature, Scope, Limits, and Consequences. *Minds Mach.* **2020**, *30*, 681–694. [\[CrossRef\]](#)
66. Farkas, A.; N'emeth, R. How to Measure Gender Bias in Machine Translation: Optimal Translators, Multiple Reference Points. *arXiv* **2020**, arXiv:2011.06445.
67. Badjatiya, P.; Gupta, M.; Varma, V. Stereotypical Bias Removal for Hate Speech Detection Task using Knowledge-based Generalizations. In Proceedings of the World Wide Web Conference, San Francisco, CA, USA, 13–17 May 2019; pp. 49–59. [\[CrossRef\]](#)
68. De-Arteaga, M.; Romanov, A.; Wallach, H.; Chayes, J.; Borgs, C.; Chouldechova, A.; Geyik, S.; Kenthapadi, K.; Kalai, A.T. Bias in Bios: A Case Study of Semantic Representation Bias in a High-Stakes Setting. In Proceedings of the Conference on Fairness, Accountability, and Transparency, Atlanta, GA, USA, 29–31 January 2019; pp. 120–128. [\[CrossRef\]](#)
69. Heindorf, S.; Scholten, Y.; Engels, G.; Potthast, M. Debiasing Vandalism Detection Models at Wikidata. In Proceedings of the World Wide Web Conference, San Francisco, CA, USA, 13–17 May 2019; pp. 670–680. [\[CrossRef\]](#)
70. Zuckerman, M.; Last, M. Using Graphs for Word Embedding with Enhanced Semantic Relations. In Proceedings of the Thirteenth Workshop on Graph-Based Methods for Natural Language Processing (TextGraphs-13), Hong Kong, China, 3–7 November 2019; pp. 32–41. [\[CrossRef\]](#)
71. Peng, X.; Li, S.; Frazier, S.; Riedl, M. Reducing Non-Normative Text Generation from Language Models. In Proceedings of the 13th International Conference on Natural Language Generation, Dublin, Ireland, 7–10 September 2020; pp. 374–383.
72. Kiritchenko, S.; Mohammad, S. Examining Gender and Race Bias in Two Hundred Sentiment Analysis Systems. In Proceedings of the Seventh Joint Conference on Lexical and Computational Semantics, New Orleans, LA, USA, 5–6 June 2018; pp. 43–53. [\[CrossRef\]](#)
73. Prates, M.O.; Avelar, P.H.; Lamb, L.C. Assessing gender bias in machine translation: A case study with google translate. *Neural Comput. Appl.* **2019**, *32*, 1–19. [\[CrossRef\]](#)
74. Sheng, E.; Chang, K.W.; Natarajan, P.; Peng, N. The Woman Worked as a Babysitter: On Biases in Language Generation. In Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing (EMNLP-IJCNLP), Hong Kong, China, 3–7 November 2019; pp. 3407–3412. [\[CrossRef\]](#)
75. Hutto, C.; Gilbert, E. VADER: A Parsimonious Rule-based Model for Sentiment Analysis of Social Media Text. In Proceedings of the International AAAI Conference on Web and Social Media, Oxford, UK, 26–29 May 2015.
76. Stanovsky, G.; Smith, N.A.; Zettlemoyer, L. Evaluating Gender Bias in Machine Translation. In Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics, Florence, Italy, 28 July–2 August 2019; pp. 1679–1684. [\[CrossRef\]](#)
77. Ott, M.; Edunov, S.; Grangier, D.; Auli, M. Scaling Neural Machine Translation. In Proceedings of the Third Conference on Machine Translation: Research Papers, Brussels, Belgium, 31 October–1 November 2018; pp. 1–9. [\[CrossRef\]](#)

78. Basta, C.; Costa-jussà, M.R.; Casas, N. Evaluating the Underlying Gender Bias in Contextualized Word Embeddings. In Proceedings of the First Workshop on Gender Bias in Natural Language Processing, Florence, Italy, 1–2 August 2019; pp. 33–39. [CrossRef]
79. Groenwold, S.; Ou, L.; Parekh, A.; Honnavalli, S.; Levy, S.; Mirza, D.; Wang, W.Y. Investigating African-American Vernacular English in Transformer-Based Text Generation. In Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing (EMNLP), online, 16–20 November 2020; pp. 5877–5883. [CrossRef]
80. Blodgett, S.L.; Green, L.; O'Connor, B. Demographic Dialectal Variation in Social Media: A Case Study of African-American English. In Proceedings of the 2016 Conference on Empirical Methods in Natural Language Processing, Austin, TX, USA, 1–5 November 2016; pp. 1119–1130. [CrossRef]
81. Babaeianjelodar, M.; Lorenz, S.; Gordon, J.; Matthews, J.N.; Freitag, E. Quantifying Gender Bias in Different Corpora. In *Companion Proceedings of the Web Conference 2020*; ACM: New York, NY, USA, 2020.
82. Iandola, F.; Shaw, A.; Krishna, R.; Keutzer, K. SqueezeBERT: What can computer vision teach NLP about efficient neural networks? In Proceedings of the SustaiNLP: Workshop on Simple and Efficient Natural Language Processing, online, 11 November 2020; pp. 124–135. [CrossRef]
83. Hutchinson, B.; Prabhakaran, V.; Denton, E.; Webster, K.; Zhong, Y.; Denuyl, S. Social Biases in NLP Models as Barriers for Persons with Disabilities. In Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics, online, 5–10 July 2020; pp. 5491–5501. [CrossRef]
84. Devlin, J.; Chang, M.W.; Lee, K.; Toutanova, K. BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding. In Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Minneapolis, MN, USA, 2–7 June 2019; Volume 1, pp. 4171–4186. [CrossRef]
85. Davis, J. Gender Bias In Machine Translation. 2020. Available online: <https://towardsdatascience.com/gender-bias-in-machine-translation-819ddce2c452> (accessed on 1 January 2021).
86. Johnson, M. Providing Gender-Specific Translations in Google Translate. Google AI Blog. 2018. Available online: <https://ai.googleblog.com/2018/12/providing-gender-specific-translations.html> (accessed on 1 August 2020).
87. Johnson, M. A Scalable Approach to Reducing Gender Bias in Google Translate. Google AI Blog. 2018. Available online: <https://ai.googleblog.com/2020/04/a-scalable-approach-to-reducing-gender.html> (accessed on 1 September 2020).
88. Gupta, M.; Agrawal, P. Compression of Deep Learning Models for Text: A Survey. *arXiv* **2020**, arXiv:2008.05221.
89. Hooker, S.; Moorosi, N.; Clark, G.; Bengio, S.; Denton, E. Characterising Bias in Compressed Models. *arXiv* **2020**, arXiv:2010.03058.