

Recepción / Received: 4 de enero de 2023
Aprobación / Approved: 11 de abril de 2023



SOCIAL IMAGINARIES ON MENTAL ILLNESS: A COMPUTATIONAL APPROACH BASED ON TEXT MINING

Manuel Cebral-Loureda^a Manuel Torres-Cubeiro^b

^a Tecnológico de Monterrey. Monterrey, México. manuel.cebral@tec.mx

^b Universidade de Santiago de Compostela. Santiago de Compostela, España. manueltorres.cubeiro@usc.es

Abstract

This article presents research within an area of study between psychology and social communication. The study approaches how the term “mental illness” is used in academic communication within a dataset compared with other two online databases: one of newspaper articles and another of film abstracts. More than 5000 abstracts extracted from those databases have been analyzed using computational techniques with R programming: network relations, longitudinal analysis, correlations calculus and sentiment analysis. We have been able to describe the social imaginaries present in communication about mental illness in those datasets. Our findings revealed a significant gap between the scientific standards and common view on mental health, somehow related with the stigma linked to mental illness. Two social imaginaries of mental illness have been identified mining those three datasets: the academical and the popular social imaginary of mental illness. Only by understanding how complexity is simplified in social communication we would be able to manage better, not just the suffering of living with a mental illness, but also the stigma surrounding mental illness.

Key words: social imaginaries, mental illness, social communication, psychology, social behavior.

The present research treats a mixed area of study, between psychology and social communication, understanding that scientific psychology standards also need to be contrasted with the way in which they are present in social life. Particularly, the study approaches how the term “mental illness” is used in an academic dataset compared with other forms of social communication, in this case, two other datasets: one of newspaper articles and another of film abstracts.

On one hand, for the scientific standard, the psychiatric taxonomies are usually taken. APA’s DSM-5 and WHO’s ICD-10 define mental illness from a bio-psycho-social perspective (APA, 2013; WHO, 1993)¹. By acknowledging the bio-psycho-social complexity, both classifications recognize a

▼ Autor para la correspondencia

manueltorres.cubeiro@usc.es

¹ “A mental disorder is a syndrome characterized by clinically significant disturbance in an individual’s cognition, emotion regulation, or behavior that reflects a dysfunction in the psychological, biological, or developmental processes underlying mental functioning. Mental disorders are usually associated with significant distress in social, occupational, or other important activities” (DSM-V, APA 2013)

reality of mental illness with multiple faces, with ontological, psychological and social elements (Cockerham, 2017).

Nevertheless, on the other hand, communication on mental illness does not follow these scientific standards. Instead, social communication simplifies such complexity generating social imaginaries of mental illness (Torres-Cubeiro 2012; 2018). This significant gap between scientific standards and common view on mental health has been studied as *Mental Health Literacy* (Jorm et al., 1997; Torres-Cubeiro 2016); as well as in relation with the stigma linked to mental illness (Pescosolido, 2013). Consequently, mental illness doesn't work as a closed concept; on the contrary, it is communicated from different points of views which can be studied and compared, trying to detect what is more relevant for each one.

Theoretical framework

What we call society emerges in communication during human evolution as an intersubjective reality; nor as an ontological reality in the physical world, nor as a physical reality in psychology (Luhmann, 1992, 1995; Moeller, 2013; Spencer-Brown, 1979). Social imaginaries are part of the complexity simplification that society has the function of facilitating in communication (Pintos, 1995). Thus, a social imaginary makes possible to see reality, understand it, and act using the provided simplified scheme. That scheme produces, at the same time, an emotional and rational structure of what should be understood as real (Torres-Cubeiro, 2015). Social imaginaries emerge by communicative repetition until they become unquestioned truisms given by granted, thus providing meaning to reality.

Social Imaginaries are detected thanks to the binary communicative code (relevant / opaque) with which they emerge and work. A social imaginary makes visible and simple in other ways complex realities, because social imaginaries hide part of such complexity, concealing such part as *opaque* (Pintos, 2003). A social imaginary

emerges thanks to this relevant / opaque binary code. Social Imaginaries allow the possibility of perceiving, understanding and acting over such reality (Pintos, 1995). For example, by insisting on mental illness as a biochemical reality (the *relevant* side of the code), its biographical, psychological or social components (the *opaque* side of the code) are obviated. By insisting on mental illness as a biographical trauma (the *relevant*), its biochemical and social components (the *opaque*) are precluded, even eliminated from view (Torres-Cubeiro, 2012a).

There is not a unique single social imaginary but multiple competing social imaginaries. Because of this, society ensures not only its simplifying function but also the possibility of adapting to changes with a broad and available semantic pool from which to make sense of the real (Pintos et al., 2004; Torres-Cubeiro, 2012a). Although one social imaginary may appear as dominant at a given moment, it coexists with other multiple contradictory social imaginaries, ensures social adaptability to change.

Therefore, detection of social imaginaries needs to be done by discovering conflicting views in social communication (Pintos & Marticorena, 2012). Since each social imaginary hides the *opaque* in order to make perceptible what is selected as *relevant* and therefore real; only the *opaque* of a given social imaginary could be observed as *relevant* from another social imaginary that contradicts the former (Torres-Cubeiro, 2012a). Methodologically, a social imaginary is only detected from an observation of an observation, that is, from a “*second-order observation*” (Luhmann, 1992).

Hypothesis

In social communication, there is no one meaning of the term “mental illness”; on the contrary, there are several. Our starting hypothesis is that conflicting and contradictory imaginaries emerge in communication about mental illness. They can be designated as Social Imaginaries. By collecting a sample of three different sources of communication

on mental illness, we will be able to detect different and conflicting simplifications of mental illness. By contrasting what is seen as *relevant* and what is seen as *opaque* in one source, with what is *relevant* and what is *opaque* in a different source, the Social Imaginaries in use about mental illness could be observed.

Methods

Methodologically, the present research can be situated between different approaches. Firstly, the digital methods as Rogers (2013, p. 38) has defined them: in studying social media we are studying the whole society, not a part of it. Rogers proposes methods that combine different databases and uses techniques like web scraping and term frequencies (Rogers, 2019). In this sense, it is necessary to understand society as it is present in the places where it expresses itself. Similarly, Marres (2017) proposes a digital sociology. He affirms that new resources of social data available in digital format demand the use of new techniques which include computation as a sociological method.

Applying these computational techniques, the methods used are also related with text mining and text analytics. As Moreno and Redondo (2016) have stated, text analytics can be implemented through computational tools to many different kinds of corpus, like emails, blogs, tweets and forums, among others. In the present case, abstracts of different types of databases were put together by applying text analytics over them. Moreno and Redondo also quote techniques such as the sentiment analysis used in this study. Previous research with this approach is Cebral-Loureda et al. (2023).

In addition, it should be mentioned the digital humanities approach. As Moretti and many others following him had defended specifically (Kuhn, 2019; Martínez-Gamboa, 2017; Moretti, 2013), a distant reading is needed to handle the huge amount of data available nowadays. Through techniques like relationship network graphs,

it is possible to observe from a distant point of view what happens in a big corpus with several documents, specifically their trends. Finally, this distant reading approach could be complemented with a closer reading, analyzing how specific terms have been used in their more specific context. This double view has also been applied in the present case.

Data and techniques

For the present study, data has been gathered from three different datasets or resources: IMDb, Scopus and The New York Times. A total of 5650 abstracts were collected using the keyword “mental illness” during the first days of November 2020:

- *Scopus*: 2000 academic article abstracts including: title, abstract and year; corresponding with the years from 2000 to 2020. These data have been collected using the permissions which the Scopus platform itself allows for academic research. 2000 is the highest number of items that the platform allows with that configuration, but it is important to point out that this source is the only one among the three considered which has many more abstracts than recollected. So, it could be said that the research was made, just in this database, over a sample. The criteria applied to obtain was the *relevance* option that Scopus platform provides -not the date, nor the number of citations-.
- *The New York Times*: 1800 journal abstracts were obtained from one of the most popular newspapers in the United States. The data was captured using the R package *rtimes* (Chamberlain, 2019) with the function `as_search()`. The database generated includes the variables: title, abstract and publication date of published articles. The abstracts correspond to articles also published between the years 2000 and 2020. The number of abstracts obtained approximates very well the maximum allowed by Scopus.

- **IMDb:** it is a database of film and tv programs launched in 1990 and soon moved to the web (<https://www.imdb.com/>). For the present research there have been collected 1770 film abstracts including the variables: title, year, gender and rating. The data has been gathered using scraping techniques from their website through the R programming package *rvest* (Wickham, 2021). In this case, less amount of data was obtained, so we used all of them although they had dates earlier than 2000. We preferred a more accurate volume of data, than an exact period of publication. Anyway, the dataset has 1458 abstracts between 2000 and 2020 and just 312 abstracts before 2000, which means less than 18% of the data.

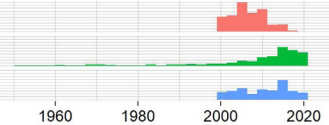
dataset	abstracts	words	distribution along time
academic	2000	210538	
films	1770	39363	
newspaper	1546	35578	

Table 1. Summary of the corpus with the structure of each database and the histogram of abstracts deployed along time. There is more data on the academic database but, as it has been explained, all the data has been normalized taking into account their proportions. Own elaboration.

These datasets were mined in five steps through R programming techniques. First, the package *tidytext* (Robinson & Silge, 2020), together with the *tidyverse* environment (Wickham, 2021), allowed easily tokenizing the texts (abstracts), extracting stopwords and counting most frequent words of each dataset. With this, a global approach was obtained to the whole corpus. Also, with the function *bind_tf_idf()*, it was possible to calculate the term frequency-inverse document frequency (abbreviated as *tf-idf*), that is, an useful measure to find the words that characterize the most each dataset being little present in the others (Sparck Jones, 1972).

A second group of techniques has to do with network relations, which have proven to be useful for the quantitative study of the Humanities (Ahnert et

al., 2020). Data was reorganized according to the *tidygraph* (Lin Pedersen, 2020) structure, which led to visualizing the corpus in terms of nodes and edges. In this case, each dataset became a main node from where edges depart to map the within words' frequencies. Words appeared working as secondary nodes for other words; therefore, connecting the domains of each dataset and allowing to visualize the network characteristic to the corpus.

In a third step, a longitudinal analysis was applied inquiring which words were trending from 2000 to 2020 for each dataset within the corpus. To visualize data's evolution management operations with the *tidyverse* package were applied.

In a fourth step, correlations between words were calculated, according to their appearance within each dataset. This was possible through the *widyr* package (Robinson, 2020), specifically using the *pairwise_cor()* function which finds correlations of pairs of items in a document. The calculation is based on the Pearson correlation giving a *phi* coefficient as result of the function (Silge & Robinson, 2017, pp. 63-66). This coefficient shows how often a pair of words appear together in the same section. In the present case, the words chosen to calculate their correlations were the APA terms that classify mental illness, that is: schizophrenia, psychopathy, depression and bipolar (APA, 2013). It is also worth pointing out that the unity where that correlations were found was the abstract, previously filtered or separated by dataset. So, the results show how the correlations are given inside each dataset according to the appearance of words together in the same abstract.

Finally, a sentiment analysis was applied: managed with *tidyverse* and *tidytext* packages, which in turn uses sentiment libraries such as NRC (Mohammad & Rubin, 2016). NRC library classifies words in eight sentiment categories: trust, fear, sadness, anger, surprise, disgust, joy and anticipation. Belonging to a sentiment category, each word adds one value point to that sentiment, given some words can be in different categories at the same time. NRC also classifies sentiments as positives or negatives:

trust, surprise, joy and anticipation as positives; and fear, sadness, anger and disgust as negatives. Then, the words with sentiment within the NRC dataset and the words of the corpus were matched. The sentiment value given in the NRC dataset was multiplied by the number of times that the word appeared within the studied corpus.

In order to compensate size differences among datasets, all calculations have been done using relative values; that is, counting total words amounts of each datasets and calculating their percentages -with the exception of the `bind_tf_idf()` function and the `pairwise_cor()` function, which do not allow this correction-. It is important to acknowledge that no lemmatization nor stemming have been used. Although these techniques could increase precision in some operations, they introduce errors when confusing some words. Also, many times the specific form of a word is valuable. However, for some forms of closer expressions that could contribute more together, some terms have been changed manually. For instance: the expressions “mental illness”, “mental health” and “new york” have been considered as a single token, given that they express a specific meaning, different from what they do independently. In a similar fashion, different tokens such as “disorders” and “disorder”, “study” and “studies” or “shooting” and “shootings” have been considered as the same word. There were some little adaptations to the keywords that APA uses for mental illness classification: “psychopathic” and “psychopath” were assimilated with “psychopathy”; “depressive” and “depressed” with “depression”; and “schizophrenic” with “schizophrenia”. Numbers and some abbreviations have been deleted as well as the stop words, in this case employing the snowball lexicon as it comes with the tidytext package (Robinson & Silge, 2020).

Results

Examining most frequent words results by dataset, in Figure 1, clearly, films and newspaper datasets share the use of closer semantic word groups, presenting related semantic word families that

contrast with those used in the academic dataset. Firstly, both datasets have in common the two most frequent words: “mentally” and “ill”, highlighting with a percentage of appearance over the 1% within the film’s dataset. Secondly, they share a vocabulary that describes common “people”, common places, like “new york”, and common actions, like “says”, in the newspaper dataset, similarly to the high frequency of words like “story”, “family”, “man”, “women”, “one”, “life”, “girl”, “husband” and “sister” in the film’s dataset. That contrasts with highly neutralized terms to refer to people in the academic dataset: terms like “study”, “results”, “objective”, “controlled”, “patients”, “sample” or “data”. Thirdly, regarding inverse document frequencies, films and newspaper dataset show a high use of the same term: “murder” -with a little higher inverse frequency within the newspaper dataset lexicon-. On the contrary, it is worth to observe that “murder” -or any related word- doesn’t appear in the high frequency terms of the academic dataset. In this sense, terms like “battles”, “shooting”, “gun”, “killed” or “fatally”, the most of them present with high frequencies in the newspaper dataset, appeal to a very tense reality, emphasizing not just a story to be told behind the mental illness, but cases of dangerous or violent conflicts. Just within the film’s dataset appear some words that make its stories softer, like “documentary” or “discovers” as its terms with the most inversely frequencies.

On the contrary, observing the academic dataset, the semantic prevalence used is very technical, with high percentages for words like “results”, “study” or “studies”, “treatment”, “conclusions” or “clinical”. This trend can be also confirmed by attending to the document inverse frequencies, where the terms “reserved”, “sample”, again “conclusions”, “data”, “objective” or “controlled” have high frequencies, expressing indexes but also isolated environments disconnected from the real and vivid world. The high values of the document inverse frequencies in the academic dataset also points out that it is the one that differs the most from the rest, that is, the one that has a more differentiated terminological use.

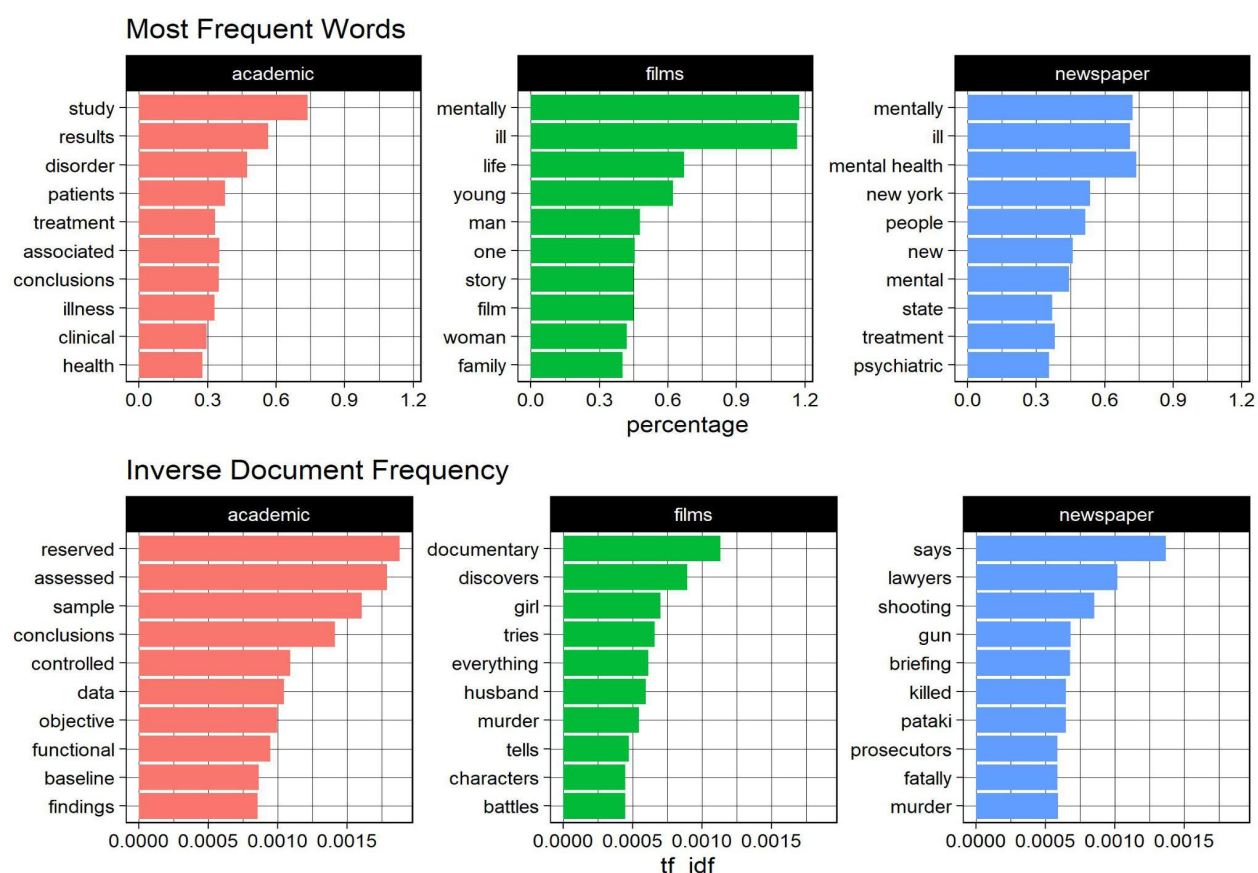


Figure 1. Ten most frequent words according to a normal count (over) and according to the inverse document frequency (below). Values were normalized using datasets size percentages. The token “mental illness” has been removed. Own elaboration using tidytext and tidyverse R programming packages.

Figure 2 adds new nuances that confirm precedent findings. Firstly, films and newspaper datasets have stronger connections, which imply more constant occurrences. With this network, it is easier to observe that they share only common terms describing cases and situations: “man”, a “family”, “one”, “help” or “year”, and the modal verb “can”. However, the terms that characterize most these situations are separated. On the newspaper’s dataset side, terms like “death”, “shooting” or “police” add fatal connotations to the common words shared; on the film’s dataset side, terms like “suffering”, “struggles” or even “love” refer also to difficult situations, but with a more processual aspects and even with positive feelings.

Between the newspaper dataset and the academic dataset there are also many words in common,

although with weaker connections. It is worth highlighting that they share a technical semantic field, with words like “study”, “treatment”, “disorder” or “patients”. All the words they have in common appeal to the name of generic processes. Meanwhile the newspaper dataset, as it was said before, concrete these cases with tense and fatal adjectives, the academic dataset uses a very neutralized vocabulary: the words which concrete the cases are “age”, “background”, “evidence”, “symptoms” and similar words. Again, these kinds of words isolate the processes of their contexts, treating life and the people who suffer mental ill as data without associated feelings.

There are no words in common -over the 0.2 minimum percentage of frequency established- between films dataset and academic dataset. This

suggests the idea that they are the furthest. Finally, it can be observed the words shared by the three datasets. It can be seen that the connections of these terms with the academic dataset are much thinner. However, in addition to the expected words they would share -"mental illness" or "mental"-, they

also have in common the temporary approach, a kind of duration, with the words “life” and “years”. Of course, this duration or processual aspect has, as it has been indicated before, very different connotations within each dataset.

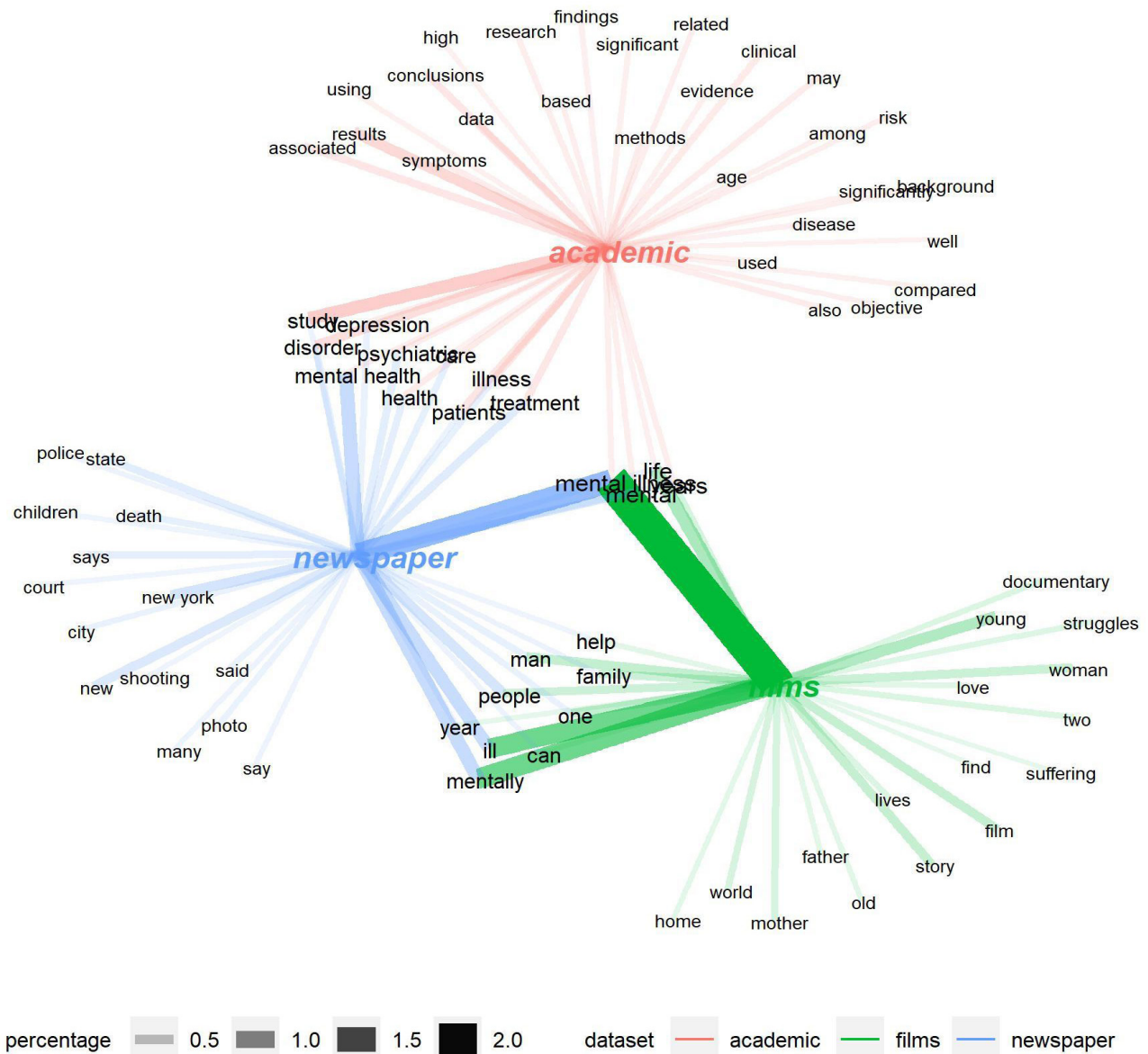


Figure 2. Corpus network visualizing words shared between datasets with a minimum over 0.2% of frequency. Values were normalized using percentages respecting the total size of each dataset. The width of the edges indicates a more or less percentage of connection. Also, label sizes point out the alpha centrality grade of each node within the network. Own elaboration using *igraph* and *tidygraph* packages in R programming.

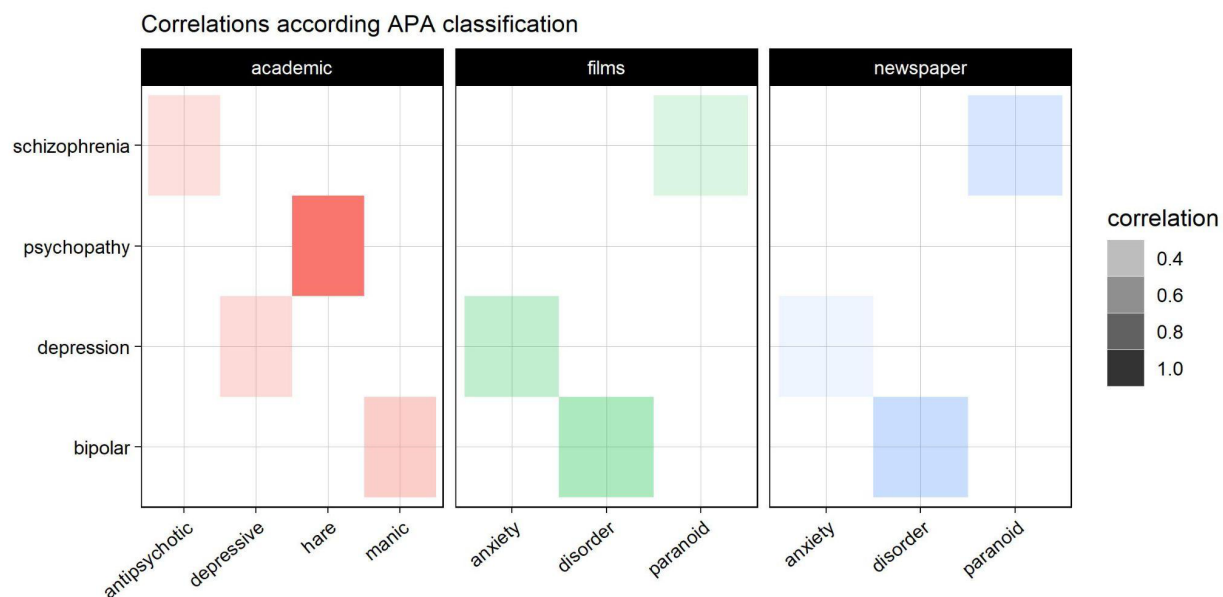
Figure 3 analyzes the presence of the four terms that, according to APA classification, define the different types of mental illness: schizophrenia, bipolar, psychopathy and depression. Figure 3a plots the correlations of these terms by dataset, filtering the most correlated word for each dataset and APA keyword. There were also removed words that appeared just one time. As it can be seen, the academic dataset has the highest correlations, being the only one that has correlations with the term “psychopathy” -as well as its derivatives “psychopathic” and “psychopathic”-, which appears very little within the whole corpus. The term highly correlated with psychopathy is “hare” which, checking in detail where it comes from, belongs to the name of the most common checklist used to assess the presence of psychopathy in individuals. Complementary, the APA classification keyword “bipolar” has the most distributed coefficients of correlation of the corpus, followed by “depression” -with its derivatives “depressive” and “depressed”-.

However, another finding highlights even more: the results on films and newspaper dataset are exactly the same, just the level of correlation changes very

little. This is a very important finding which implies that both datasets treat APA categories in a very similar way, with the same main words correlated and in the same order.

Regarding not the correlations, but the most frequent words appearing together with the APA classification terms, figure 3b can be seen. So, both plots work similarly as it happens in figure 1: the most specific terms are complemented with the most frequent. Here, it draws attention to how there are more frequent words in academic dataset related with APA terminology, especially terms which appear together with depression: “study” and “results” with the highest frequencies. Followed by the academic dataset, the films dataset has also many frequent words related with the APA mental illness classification. In this case, the frequencies are a little lower, highlighting words associated with depression, in the first place, and schizophrenia.

Again, psychopathy has very few associations, although regarding the most frequent words in figure 3b, it has a little presence in films dataset. All these findings will be discussed in the next section.



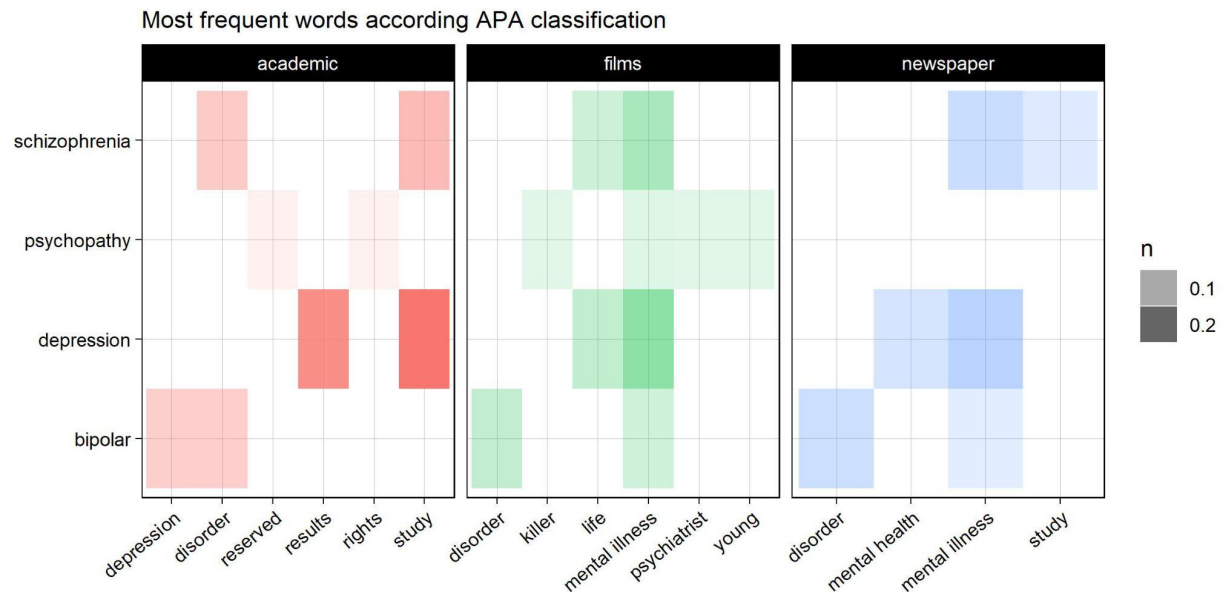


Figure 3a and 3b. More correlated words for the APA classification terms separated by each dataset, based on the Pearson correlation (above). The plot is complemented with the most frequent words appearing together with the APA classification terms and filtered by dataset (below). In 3b picture, the measures have been calculated taking into account the total amount of words, which means that are percentage values. Own elaboration using the *widyr* package in R programming.

Figure 4 shows frequency distribution over years. A clear change can be observed from the first decade (2000-2010) to the second (2010-2020): meanwhile before 2010 academic and newspaper datasets have higher frequencies, after 2010 the films dataset highlights over the others. Forcefully, mental illness issues are treated in films during the last decade, reaching the highest frequencies of the corpus in 2017 and 2019. Of course, it should be observed that the capture on Scopus was made attending the relevance order that the platform provides. This fact affects, without doubt, to the higher articles published some years ago, with more time to acquire relevance. However, it shouldn't hide at all the general trends of the data along the period observed, especially the increase of frequency within the films dataset.

Regarding the kind of words detected, again the technical terms appear within the academic dataset

in the first decade, also refer to a mediated and distant approach. Within the newspaper dataset, in addition to the term "mental illness", which has the highest frequencies, there can be observed situations that have an instant impact. Words like "photo", with a very high presence in the year 2000, as well as the terms "gun", "shooting", "death", "said" or "killed" appeal to very concrete, tense and fatal moments. They are quite opposed to the terms that more characterize films dataset: although they share a vivid and popular semantic field, the situations which appear in films dataset, overall after 2010, are more about gradual and experienced processes. Terms like "struggles", "story", "battles", "abuse", "addiction", or even verbs like "finds", "become" or "can", together with words previously detected, like "documentary" and "discover" with high inverse frequencies, refer to experiences that impact deeply in persons.

Most Frequent Words by Year

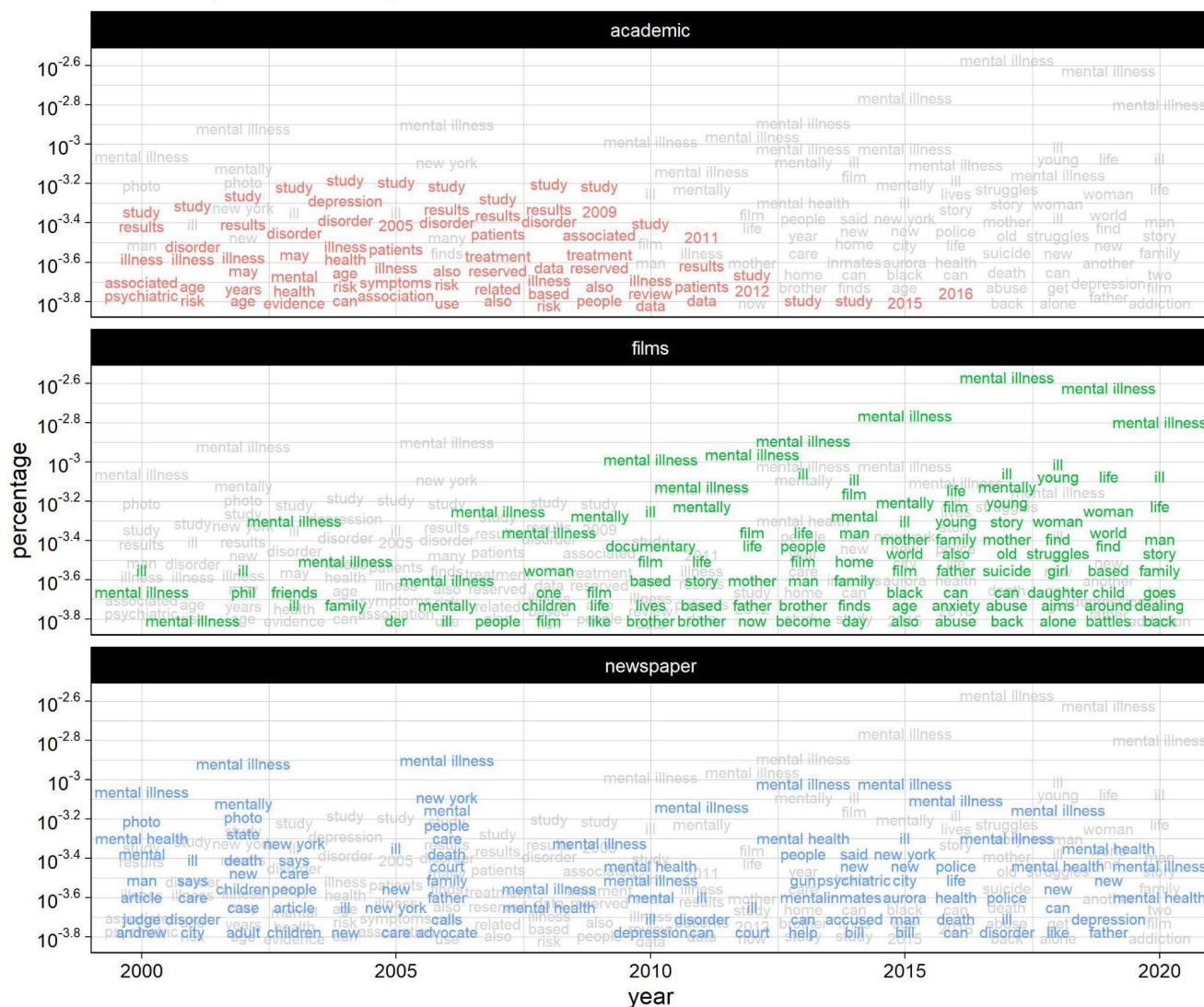


Figure 4. Most frequent words by year. Results were normalized using percentages respecting the size of each dataset. Own elaboration using *tidytext* and *tidyverse* packages in R programming.

Finally, Figure 5 presents graphically the sentiment analysis of the datasets. As it can be seen, fear is the dominant sentiment within the corpus, followed by sadness. Then trust and anger come, with a very notable presence of the first positive sentiment within the academic dataset. In general, the higher values are presented in the newspaper and film databases, predominating, by far, the negative sentiment scores. Just the academic dataset has a

positive sentiment average, with high scores of trust and anticipation. However, within the film dataset also stands out the sentiment of anticipation and, even more, joy; comparing with the others datasets. Inclusively, the film's dataset's positive sentiment average is not so far from the negative one. Finally, newspaper abstracts have a very low sentiment of joy as well as the higher contrasts between positive and negative emotions.

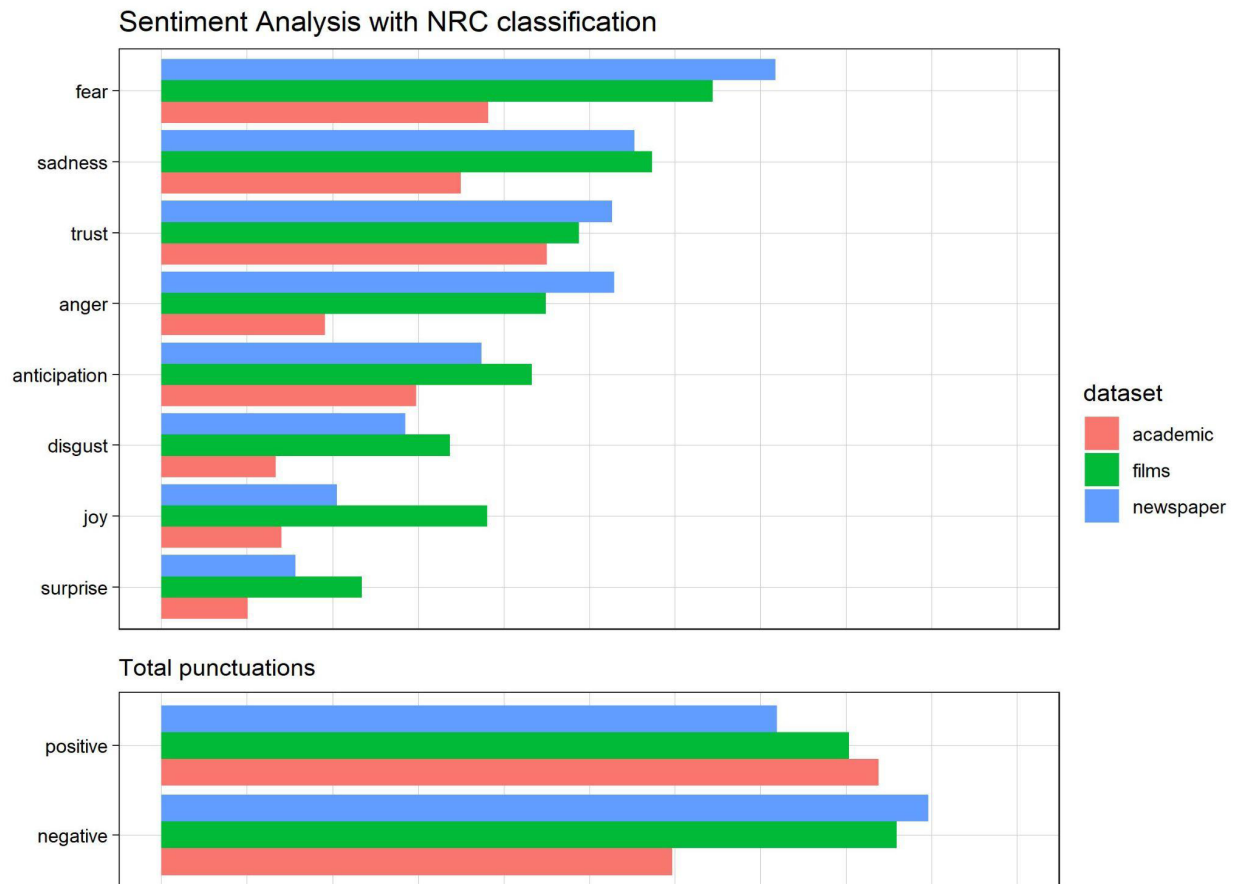


Figure 5: Sentiment scores obtained applying NRC classifier over the corpus. The plot above shows scores by sentiment and dataset; below, the sentiments aggregated as positive or negative. Results were normalized using percentages respecting the size of each dataset. Own elaboration using tidytext and tidyverse packages in R programming.

Discussion

The collected data show how communication about mental illness happens using mainly two Social Imaginaries: one that can be denominated *academic social imaginary* of mental illness, closer to medical scientific standards; and a second one, a *popular social imaginary*, which is away from those standards, and it is mainly represented, in the case of study, by the film's dataset. The newspaper dataset flows between these two social imaginaries, although often having more in common with films. For example, coincidences were found between newspapers and films in the most frequent words and inverse document frequency analysis, also in the APA terms correlation analysis, as well as in

the sentiment analysis. However, also found were common terms shared between academic and newspaper in the network relationship analysis, and more generally, considering the semantics used in both datasets.

More specifically, while academic and newspaper datasets share a type of communication about mental illness with a more technical vocabulary, newspapers share with films a more secular or popular semantic description. Clearly, mental illness in academic media focuses on medical and clinical descriptions, obviating biographical or social context and creating a more isolated, impersonal and distant discourse; on the other hand, mental disorders in films abstracts and

also in newspapers conceal medical clinical data to focus on biographical and social information, even giving very personal details within the film's dataset, which appeal to personal and deep processes affecting people's lives. This difference is also strongly supported in the correlations analysis shown in picture 3, where the technical APA terms to refer to mental illness -that is: "schizophrenia", "bipolar", "psychopathy" and "depression"- are more present and with higher correlations in the academic dataset. Again, the popular social imaginary would be conformed with the view given in the newspaper and the film's datasets, talking about mental disorders and their relations with suffering, crime and violence present in common life.

Academic social imaginary of mental illness

The academic imaginary offers a simplified description of those living a mental diagnosis as patients with a diagnostic label, following (or not) a treatment, and with a clinical prognosis in a controlled environment. From this perception, few sentiments or emotions are involved, as it was shown in figure 5. In addition, a so-called objective description is offered, with high values for terms like "sample", "data", "factors", "depression", "mortality" or "prevalence" referring to the implicated subjects; or "study", "results", "symptoms", "risks", "clinical" or "cognitive" referring to the processes. All these terms always define a distant approach, mediated by quantitative and technical vocabulary.

The academic imaginary almost doesn't mention social or biographic data about those with a psychiatric diagnosis. In another publication, it has been pointed out how the biomedical imaginary on mental disorders is dominant in social communication (Torres 2012b). Within the present corpus, it has been shown how the academic dataset was more correlated with the APA technical terms, treating more depression and schizophrenia.

The academic imaginary describes a patient without much reference to his/her suffering or those issues involving his social or family environment (Figure 2). No terminology related to violence or crime is in use in this social imaginary. Jorm (1997) points out the distance between academic scientific standards describing mental health and the more common popular view that associates mental disorders with crime and violence. That distance is strongly linked with the negative stigma associated with these disorders (Pescosolido, 2013) and it deserves further considerations and research.

Popular social imaginary of mental illness

The popular imaginary describes not patients but people: brothers, sisters, husbands, sons, or siblings suffering a mental illness. Although diagnostic labels are not forgotten, this imaginary focuses more on personal concrete data to describe mental illness. The judicial elements are relevant, as well as those related with criminal and violent behaviors. Mental illness is related in these datasets with words as "story", "family", "man", "women", "one", "life", "girl", "husband" or "sister" referring to the implicated subjects; or terms like "death", "gun" or "shooting" referring to fatal and tense instants. Particularly, film abstracts provide biographies with personal and social background, obviating more medical terminology (figure 2). Also, APA psychiatric labels are scarcely in use among film abstracts, being more frequent than in the newspaper dataset.

It is significant that in the final years of data collection, there is a substantial increase of the presence of "mental illness" in film abstracts, when compared with the other two datasets; with the terminology associated with the popular social imaginary more present. That increase is less evident within newspaper abstracts, and not present at all among academic data.

If it is observed in sentiment analysis (figure 5), films and newspapers show a considerable presence of

vocabulary associated with fairness and sadness, something coherent with the frequency of crime, suffering and violent terminology, which is very frequent in both datasets. The academic imaginary offers a perspective emotion clean or prophylactic, it is neutralized by the technical terms and with the lack of personal life references. While, the popular imaginary offers an emotional sad and fearsome atmosphere, coherent with film narratives and crime newspaper outlines.

It is important to highlight, in this context, that the APA technical term “psychopathy” doesn’t appear at all in the newspaper dataset and just a very little in the film’s dataset. However, it doesn’t necessarily mean that there is no presence of this kind of mental illness. There is evidence to interpret that the technical formulation has been replaced by the criminal and violent behavior found throughout all the analysis carried out in the present study.

Conclusions

Two main social imaginaries of mental illness have been identified in social communication analyzing data mined from three datasets widely used on the global internet during November 2020. Firstly, we detected what we could call an *academical* social imaginary of mental illness. This social imaginary offers a perception of mental illness as a medical entity scientifically defined: a patient with a differential diagnosis, a clinical development and a medical treatment. That is what is relevant for this social imaginary. But the academic imaginary offers no much access to the biographical or social aspects involved, neither emotions or sentiments. It is here where its opacity lies. This social imaginary, closer to scientific standards, is dominant in the academic dataset, with some presence in the newspaper abstracts and almost non-existence among the film’s dataset.

A second social imaginary on mental illness has been identified: the *popular* social imaginary of mental illness. This social imaginary offers a

perception of mental illness as a biographical or social happening: someone who suffers from a sad and difficult situation, related to different forms of violence and linked to crime or punishable behaviors. That is what is relevant for this social imaginary. The popular imaginary does not offer much access to medical, clinical or psychiatric data. It is here where its opacity lies. This social imaginary is dominant in the film dataset, with a clear presence in the newspaper abstract.

These two imaginaries have shown few commonalities, opposing the technical to the biographic, the neutral to the sentimental, the cognitive processes to the suffering and pain. So, the relevant and the opaque are mutually exclusive. Just the newspaper dataset works, sometimes, as a bridge between them, at least in terms of shared words. The issue is that precisely this dataset, the newspaper, is the one which has the highest values of negative sentiment, with very strong values of fear, sadness and anger. In this sense, it shouldn’t be considered as a good candidate to mediate between the polarization we have described, at least in the form of the newspaper communication found between 2000 and 2020.

This issue worsens when the timeline analysis is taken into account. As it has been shown in newspapers and, even more, in films, the presence of mental illness has increased during the last decade. Especially when we consider films, where there have been more discussions about mental illness since 2010, *being the most used media among the three treated in this study*. However, this could be considered positive as it supposes a way to popularize mental illness problems, taking it out of the controlled and isolated environment of the psy-practices and labs. Also, there are quite positive sentiment scores in the film dataset, even with words that appeal to living and vivid processes which invite one to discover and transform oneself. The documentary genre, which appears in the inverse document term frequency, offers an approach that combines objectivity with a biographical and personal view. Although almost no technical term

accompanying it was found. That lack is often replaced with stories of addiction, suffering and violence.

Further analysis

Social imaginaries construct common sense in social communication. In our topic common sense refers to how in social communication mental illness is understood, perceived and activated; or, in a wider sense, mental health. The analysis shows how it is possible to distinguish two social imaginaries in use when social communication focuses on mental illness: the *academic* and the *popular*. It could be argued that both simplified so much the complex reality involved in mental health, and that both contribute to feeding our illiteracy on the topic, contributing to explain the social stigma almost inevitably linked to these experiences.

The present investigation can be expanded to research deeply these possible limitations, it would be suitable to analyze with more attention each dataset. Also, there are variables that have not been touched, like audience votes within film abstracts or the number of citations within the academic dataset. It hasn't been carried out because the difficulty to find a way to normalize them. However, these variables can be taken into account in a specific approach to each dataset, studying separately what tendencies or worries are more relevant for each of them and even modelling their topics.

Further research also should be conducted to comprehend not only the medical and clinical processes, or the biographical and social elements, but the communication processes involved, more and more charged with polarization on current social networks. Only by understanding how such complexity is simplified in our social communication, we would be able to manage better, not just the suffering of living with a mental illness, but the stigma that seems to be inevitably associated with it.

Disclosure statement

No potential conflict of interest is reported by authors.

Funding

This paper has not received any funding.

References

- Ahnert, R., Ahnert, S. E., Coleman, C. N., & Weingart, S. B. (2020). *The Network Turn: Changing Perspectives in the Humanities* (1.^a ed.). Cambridge University Press. <https://doi.org/10.1017/9781108866804>
- American Psychiatric Association, & American Psychiatric Association (Eds.). (2013). *Diagnostic and statistical manual of mental disorders: DSM-5* (5th ed). American Psychiatric Association.
- Cebral-Loureda, M. Hernández-Baqueiro, A., & Tamés-Muñoz, E. (2023). A text mining analysis of human flourishing in Twitter. *Scientific Reports*, 13(1), 3403. <https://doi.org/10.1038/s41598-023-30209-7>
- Chamberlain, S. (2019). *rtimes: Client for New York Times «APIs»* (0.5.0) [Computer software]. <https://github.com/ropengov/rtimes>
- Cockerham, W. C. (2017). *Sociology of mental disorder* (Tenth edition). Routledge/Taylor & Francis Group.
- Jorm, A. F., Korten, A. E., Jacomb, P. A., Christensen, H., Rodgers, B., & Pollitt, P. (1997). "Mental health literacy": A survey of the public's ability to recognise mental disorders and their beliefs about the effectiveness of treatment. *Medical Journal of Australia*, 166(4), 182-186. <https://doi.org/10.5694/j.1326-5377.1997.tb140071.x>
- Kuhn, J. (2019). Computational text analysis within the Humanities: How to combine working

- practices from the contributing fields? *Language Resources and Evaluation*, 53(4), 565-602. <https://doi.org/10.1007/s10579-019-09459-3>
- Lin Pedersen, T. (2020). *tidygraph: A Tidy API for Graph Manipulation (1.2.0)* [Computer software]. <https://cran.r-project.org/package=tidygraph>
- Luhmann, N. (1992). The Concept of Society. *Thesis Eleven*, 31(1), 67-80. <https://doi.org/10.1177/072551369203100106>
- Luhmann, N. (1995). *Social systems*. Stanford University Press.
- Marres, N. (2017). *Digital sociology: The reinvention of social research*. Polity.
- Martínez-Gamboa, R. (2017). Big Data en humanidades digitales: De la escritura digital a la “lectura distante”. [Big data and humanities: from digital writing to a “distant reading”] *Revista Chilena de Literatura*, 94. <https://revistaliteratura.uchile.cl/index.php/RCL/article/view/44969>
- Moeller, H.-G. (2012). *The radical Luhmann*. Columbia University Press.
- Mohammad, S., & Rubin, P. (2016). NRC Word-Emotion Association Lexicon. National Research Council Canada. <http://saifmohammad.com/WebPages/NRC-Emotion-Lexicon.htm>
- Moreno, A., & Redondo, T. (2016). Text Analytics: The convergence of Big Data and Artificial Intelligence. *International Journal of Interactive Multimedia and Artificial Intelligence*, 3(6), 57. <https://doi.org/10.9781/ijimai.2016.369>
- Moretti, F. (2013). *Distant reading*. Verso.
- Nielsen, F. A. (2011). AFINN. Informatics and Mathematical Modelling, Technical University of Denmark. <http://www2.imm.dtu.dk/pubdb/pubs/6010-full.html>
- Pescosolido, B. A. (2013). The Public Stigma of Mental Illness: What Do We Think; What Do We Know; What Can We Prove? *Journal of Health and Social Behavior*, 54(1), 1-21. <https://doi.org/10.1177/0022146512471197>
- Pintos, J. L. (2003). El metacódigo “relevancia/opacidad” en la construcción sistémica de las realidades. [The “Relevance / Opacity” Metacode in the systemic construction of realities] *Revista de Investigaciones Políticas y Sociológicas (RIPS)*, 2(1-2), 21-34. <https://www.redalyc.org/pdf/380/38020202.pdf>
- Pintos, J. L., & Marticorena, J. R. (2012). Análisis sociocibernético del discurso. La explotación de datos y los procedimientos informatizados en las investigaciones sobre Imaginarios Sociales. Un caso. [Sociocybernetic discourse analysis. The exploitation of data and computerized procedures in research on Social Imaginaries. A case]. *Revista de Investigaciones Políticas y Sociológicas (RIPS)*, 11(2), 47-79. <https://revistas.usc.gal/index.php/rips/article/view/376>
- Pintos-Juan-Luis. (1995). Los imaginarios sociales la nueva construcción de la realidad social. [Social Imaginaries: the New Construction of Social Reality] *Sal Terrae*. https://www.academia.edu/20690963/Los_imaginarios_sociales_la_nueva_construcci%C3%B3n_de_la_realidad_social
- Pintos, Juan Luis, Marticorena, J.R., Rey, R. (2004). El rol del paciente en salud mental. [The role of the patient in mental health] Copia impresa, Memoria de investigación USC [Unpublished University memory, University of Santiago de Compostela].
- Robinson, D. (2020). *widyr: Widen, Process, then Re-Tidy Data (0.1.3)* [Computer software]. <https://cran.r-project.org/package=widyr>
- Robinson, D., & Silge, J. (2020). *tidytext: Text Mining using «dplyr», «ggplot2», and Other Tidy Tools (0.2.4)* [Computer software]. <https://cran.r-project.org/package=tidytext>

- Rogers, R. (2013). *Digital methods* (First paperback edition). The MIT Press.
- Rogers, R. (2019). *Doing digital methods* (1st edition). SAGE Publications.
- Silge, J., & Robinson, D. (2017). *Text mining with R: A tidy approach* (First edition). O'Reilly. <https://www.tidytextmining.com/>
- Sparck Jones, K. (1972). A Statistical Interpretation of Term Specificity and Its Application in Retrieval. *Journal of Documentation*, 28(1), 11-21. <https://doi.org/10.1108/eb026526>
- Spencer Brown, G. (1979). *Laws of form*. Dutton.
- Torres-Cubeiro, M. (2012a). Complejidad social y locura en Galicia: Medicina, psiquiatría, familiares, personas diagnosticadas de enfermedad mental grave: el sentido de los sinsentidos. [Social complexity and insanity in Galicia. Medicine, psychiatry, relatives, people diagnosed with serious mental illness: the sense of nonsense.] Editorial Académica Española. <https://www.abebooks.com/9783848478736/Torres-Cubeiro-Complejidad-Social-Locura-3848478730/plp>
- Torres-Cubeiro, M. (2012b). Imaginarios sociales de la enfermedad mental. [Social imaginaries of mental disorder] *Revista de Investigaciones Políticas y Sociológicas (RIPS)*, 11(2), 101-113. <https://revistas.usc.gal/index.php/rips/article/view/378>
- Torres-Cubeiro, M. (2015). La evolución del concepto de imaginarios sociales en la obra publicada de Juan Luis Pintos de Cea Naharro. [The evolution of the concept of social imaginaires in published works of Juan Luis Pintos de Cea Naharro] *Imagonautas*, 6, 1-14. <https://imagonautas.webs.uvigo.gal/index.php/imagonautas/article/view/5/6>
- Torres-Cubeiro, M. (2016). Alfabetización en salud mental, estigma e imaginarios sociales. [Mental Health Literacy, stigma and social imaginaires] *Imagonautas*, 8, 50-63. <https://imagonautas.webs.uvigo.gal/index.php/imagonautas/article/view/62/38>
- Torres-Cubeiro, M. (2018). The complexity of non-sense making: A proposal of complex medical sociology of mental disorders. *Sociología y Tecnociencia*, 8(2), 67. <https://doi.org/10.24197/st.2.2018.67-91>
- Wickham, H. (2021a). *rvest: Easily Harvest (Scrape) Web Pages* (1.0.0) [Computer software]. <https://cran.r-project.org/package=rvest>
- Wickham, H. (2021b). *tidyverse: Easily Install and Load the «Tidyverse»* (1.3.1) [Computer software]. <https://cran.r-project.org/package=tidyverse>
- World Health Organization (Ed.). (1993). *The ICD-10 classification of mental and behavioural disorders: Diagnostic criteria for research*. World Health Organization.

Cita recomendada

Cebral-Loureda, M. y Torres-Cubeiro, M. (2023). Social Imaginaries on Mental Illness: A Computational Approach Based on Text Mining. En: *Imagonautas*, Nº 17 (12), pp. 11-26.