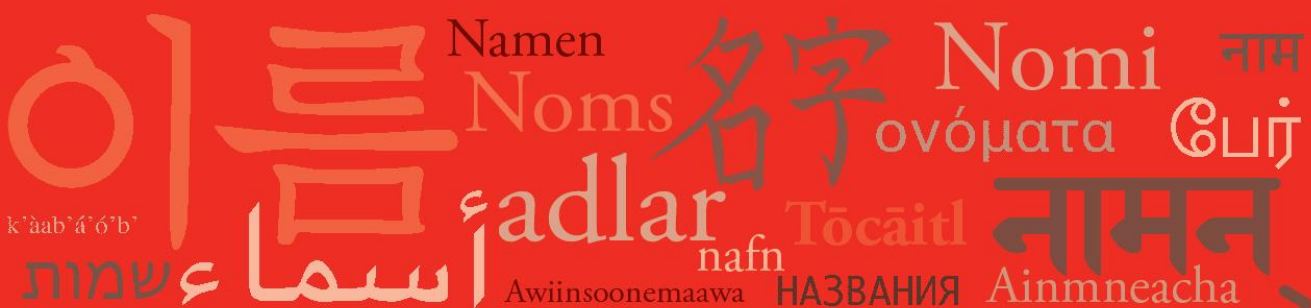




The Use of *Wanghong* as a Stigmatized By-name: A Note on Polarity Analysis and Topic Modeling Using NLP

Danyang Zheng

Tianjin University, CHINA

ans-names.pitt.edu

ISSN: 0027-7738 (print) 1756-2279 (web)

Vol. 74 No. 3, Summer 2026

DOI 10.5195/names.2026.2800



Articles in this journal are licensed under a [Creative Commons Attribution 4.0 International License](https://creativecommons.org/licenses/by/4.0/).



This journal is published by [Pitt Open Library Publishing](https://pitt-open-library-publishing.org/).

Abstract

This note applies NLP-based polarity analysis and topic modeling to explore the stigmatization of the Chinese occupational by-name *wanghong*. Using over 40,000 Weibo posts (2014-2023), the study reveals a shift toward negative sentiment and identifies key thematic patterns. It demonstrates how NLP methods can support socio-onomastic inquiry by tracing affective and ideological dynamics in digital naming practices, while also offering an efficient and objective methodological approach for investigating other socially salient by-names.

Keywords: anthroponymy, by-name, polarity, topic modeling, stigmatization, socio-onomastics, China

Introduction

The term *wanghong* ‘net red’, short for *wangluo hongren* ‘internet red person’, originated in early 21st-century China and refers to individuals who gain widespread attention through online or offline visibility. The referent of the term broadly corresponds to that of the English word *influencer*. In this paper, the term by-name is used in a socio-onomastic sense to refer to an auxiliary designation that supplements or replaces a formal name in order to convey social, relational, or evaluative meanings (McClure 2010; Bourdieu 1991). As an occupational by-name, *wanghong* encompasses emerging roles such as livestreamers, influencers, and bloggers. Beyond indicating profession, such by-names encode public perceptions and affective associations, functioning as socially saturated labels within contemporary discourse (Nick 2024; Tagg 2015; Barton & Lee 2013).

The term *wanghong*, formed by combining *wang* ‘net’ and *hong* ‘red’/fame, initially bore a positive connotation denoting digital popularity. However, as digital culture and online communication have evolved, the term has undergone significant semantic shifts, mirroring changing public perceptions and socio-discursive functions. In the era of social media, *wanghong* figures have garnered attention not only for their online influence but also for the increasing stigma² attached to the label—paralleling developments associated with the English term *influencer*. This stigma is frequently associated with sensationalism, commercial exploitation, and controversial conduct (Zhang 2019), prompting many content creators in both Chinese and Western contexts to disavow such designations³.

While sociological and media studies (Gerhards 2025; Yan & Zhou 2019) have explored the stigmatization of online figures, they often neglect the constitutive role of naming. A socio-onomastic perspective offers not merely a complementary angle but a foundational framework for analyzing how labels like *wanghong*, as occupational by-names, acquire stigmatized meanings through semantic and connotative shifts. Yet, this line of inquiry has received comparatively limited attention within broader discussions of social labeling.

Recent advances in Natural Language Processing (NLP), particularly in topic modeling, have significantly improved the ability to extract latent themes and semantic structures from large-scale text data. One notable technique is BERTopic (Bidirectional Encoder Representations from Transformers for Topic Modeling), which builds on BERT (Bidirectional Encoder Representations from Transformers), a deep learning model that generates context-sensitive word embeddings. Unlike earlier models with static representations, BERT captures nuanced semantic relationships by considering word order and surrounding context, making it especially effective in handling polysemy and subtle meaning shifts.

BERTopic leverages these contextual embeddings to enhance topic modeling in noisy or short texts, such as those found on social media. Rather than relying on traditional word co-occurrence patterns as in Latent Dirichlet Allocation (LDA), BERTopic combines deep semantic embeddings with dimensionality reduction techniques like Uniform Manifold Approximation and Projection (UMAP), clustering algorithms such as Hierarchical Density-Based Spatial Clustering of Application with Noise (HDBSCAN) or K-Means, and keyword extraction through class-based Term Frequency-Inverse Document Frequency (c-TF-IDF). This pipeline enables the identification of coherent and fine-grained topic clusters and has been shown to outperform conventional models like LDA in terms of interpretability and granularity (Egger & Yu 2022).

In this study, BERTopic is used to analyze online discourse about the by-name *wanghong*, identifying key themes and topic relationships. By applying topic modeling to a socially charged name, the study shows how NLP can advance socio-onomastic research by uncovering the discursive and cultural contexts that shape

stigma. This provides a scalable method for examining how names function as emotional and ideological symbols in digital discourse.

To complement the topic modeling analysis and provide empirical evidence for the development of stigma, this study employs polarity-based sentiment analysis—classifying texts as positive, neutral, or negative—to trace the evolution of public evaluations of *wanghong* over the past decade. By classifying over 40,000 social media posts into these three categories, the analysis reveals broad evaluative trends—most notably, the increasing stigmatization of the term. In socio-onomastics, tracing such affective trajectories is central to understanding how names shift in their social and symbolic functions. While multi-vector sentiment models (e.g., mood, intensity, aspect) offer greater emotional nuance (Enochson et al. 2020), they present higher computational demands and interpretive complexity, especially across large, diachronic datasets. In contrast, polarity analysis offers a practical and scalable approach that balances analytical clarity with methodological efficiency, making it well-suited for long-term, data-driven socio-onomastic inquiry. Building on NLP techniques, this study tests two hypotheses central to socio-onomastics: (1) that public discourse around *wanghong* has grown increasingly negative over the past decade, signaling its affective re-evaluation and stigmatization; and (2) that topic modeling can reveal the discursive and socio-cultural associations driving this shift, showing how naming reflects broader ideological trends in digital society.

Text Extraction and Polarity Analysis

This study collected posts containing the keyword “wanghong” from Sina Weibo⁴, covering a ten-year period from January 2014 to December 2023. Figure 1 outlines the key steps of data extraction and sentiment classification, which are detailed below.

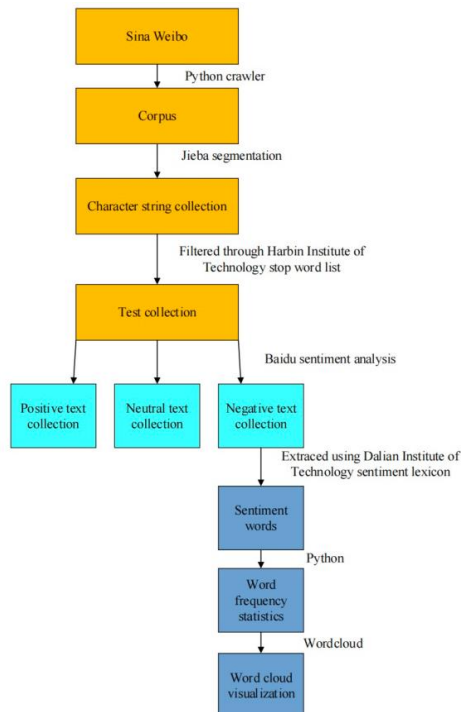


Figure 1: Workflow for Polarity Analysis and Word Cloud Generation of *Wanghong* Discussions

Using Python’s requests library to retrieve webpage content and the lxml library to apply XPath expressions, this study scraped and parsed 49,684 Weibo posts containing the keyword *wanghong*. To ensure consistency, a series of preprocessing steps were applied. These included converting full-width to half-width characters—an essential operation in Chinese text normalization, as mixed character widths can interfere with word segmentation and semantic parsing. Full-width characters are typically used for Chinese scripts, while half-width characters are common in Latin-based inputs; their coexistence can confuse tokenizers and reduce the accuracy of downstream tasks.

Further preprocessing steps involved lowercasing English characters, removing control symbols (e.g., \n, \t), and filtering out non-Chinese or semantically trivial posts (e.g., emojis, usernames, or punctuation-only strings). In addition, and following common practices in Chinese short-text processing (e.g., Zhang & Wallace 2015), this study excluded posts containing fewer than five characters, as ultra-short texts tend to be semantically sparse and contribute little to topic modeling or polarity analysis. After all preprocessing steps were completed, the dataset was reduced to 40,523 clean entries, ready for linguistic and computational analysis.

The dataset was tokenized using Python’s Jieba library⁵, and the tokenized results were filtered with the Harbin Institute of Technology (HIT) stop words list⁶. Word frequencies were calculated using Python, and word clouds were generated with the WordCloud library, a visualization tool that displays high-frequency words in varying font sizes based on their relative importance. This method provides an intuitive overview of dominant lexical items within the corpus, facilitating preliminary thematic exploration.

Sentiment polarity was assessed using Baidu’s Sentiment Analysis API⁷ via Python. Each comment received a polarity label (−1 for negative, 0 for neutral, 1 for positive) and associated probability scores for positive and negative sentiment. Labels were assigned based on the higher score exceeding an internal threshold.

To maintain consistency and minimize subjectivity, the API’s default output was used without manual adjustment. Additionally, sentiment-bearing keywords were extracted using the Dalian University of Technology’s Chinese Sentiment Lexicon⁸, enabling both category-level and lexical-level sentiment analysis. This hybrid approach allowed us to analyze both the predicted sentiment category and the linguistic elements associated with affective expression.

BERTopic Topic Modeling Analysis

Topic modeling analysis was conducted using the BERTopic model to identify and analyze latent topics. Table 1 illustrates the process of topic modeling.

Table 1: Workflow for Topic Modelling Using BERTopic

Steps	Description
Embed Documents	Convert text into vector representations using methods such as Word2Vec, TF-IDF, or BERT.
Reduce Dimensionality	Apply dimensionality reduction techniques (e.g., UMAP, PCA) to simplify embeddings.
Cluster Embeddings	Group the reduced embeddings into topic clusters using methods like HDBSCAN or KMeans.
Extract Keywords	Identify topic-specific keywords using the c-TF-IDF algorithm.
Fine-tune Topics	Refine and adjust the extracted keywords for improved topic clarity.

The text data were encoded using the sentence-transformers library with the paraphrase-multilingual-MiniLM-L12-v2 model, a pre-trained multilingual sentence embedding model capable of processing over 50 languages, including Chinese. This model generates 384-dimensional semantic vectors that capture the contextual meaning of each comment. Semantically similar texts are positioned closer in the resulting vector space, thereby facilitating effective clustering for subsequent topic modeling. UMAP (Uniform Manifold Approximation and Projection) was applied to reduce the dimensionality of text embeddings while preserving semantic relationships, enabling clearer pattern visualization and analysis. Thereafter, two procedural steps were followed. In Step 1, a similarity graph in high-dimensional space was built. UMAP began by evaluating pairwise similarities between data points in high-dimensional semantic space, where each point represents a text embedding. Local density parameters are applied to account for data variation, resulting in a weighted

The Use of *Wanghong* as a Stigmatized By-name

graph that reflects each point's proximity to its neighbors—laying the foundation for subsequent dimensionality reduction. In Step 2, mapping to a lower-dimensional space was undertaken. UMAP projected data into a 10-dimensional space by minimizing the divergence between original and reduced-space similarities. This optimization preserves both local and global structural relationships, facilitating coherent topic discovery and clustering.

After dimensionality reduction, K-means was applied to group similar comments into 10 clusters representing dominant themes. As an unsupervised algorithm, K-means iteratively assigns data points to the nearest centroid and updates cluster centers until convergence, minimizing intra-cluster variance. To identify the most representative keywords for each topic, we used the class-based TF-IDF (c-TF-IDF) algorithm, which is a core component of BERTopic. This method merged all comments in a topic into one document and scored each word based on its distinctiveness and relevance to that topic. The score for a word t in topic c is computed using the following formula:

$$c - \text{TF-IDF}(t, c) = \text{TF}(t, c) \times \log\left(\frac{N}{|\{c \in C: t \in c\}|}\right)$$

Here, $\text{TF}(t, c)$ refers to the total occurrence of word t within topic c . The log term, $\log\left(\frac{N}{|\{c \in C: t \in c\}|}\right)$, stands for the inverse category frequency of word t , which quantifies how uniquely t is associated with a specific topic as opposed to being common across many. A higher c-TF-IDF score indicates that the word is more distinctive to the given topic. In other words, a word is considered important to a topic if it occurs frequently in that topic but rarely in others. For example, if the term “scam” appears frequently in a cluster related to *wanghong* economy but only rarely in other clusters, it will receive a high c-TF-IDF score for that topic. In this way, c-TF-IDF helps extract the most distinctive lexical items that characterize each semantic region.

To enhance the interpretability and quality of topic representations, this study employed the Maximal Marginal Relevance (MMR) algorithm for keyword refinement. MMR is a re-ranking technique that selects keywords by optimizing a trade-off between relevance to the topic and novelty relative to already selected terms. Specifically, it penalizes the selection of keywords that are overly similar to one another, thus preventing redundancy. In doing so, MMR ensures that the final keyword list for each topic contains terms that are not only highly indicative of the topic's central theme but also diverse enough to reflect its semantic breadth. This enhances the descriptive richness of the topic and facilitates clearer interpretation, particularly in large-scale, unsupervised topic modeling tasks.

During the clustering process, some data points could not be meaningfully grouped with others. These were identified by HDBSCAN (Hierarchical Density-Based Spatial Clustering of Applications with Noise)—a hierarchical density-based clustering algorithm known for detecting clusters of varying densities—as outliers and were automatically labeled “-1”. These outlier topics were interpreted as noise and excluded from further analysis, as they do not form cohesive or interpretable thematic clusters.

Taken together, the steps described above form an integrated NLP pipeline for socio-onomastic inquiry. UMAP reduced the dimensionality of the semantic space, K-means identified dense thematic regions, c-TF-IDF extracted distinguishing lexical features for each cluster, and MMR refined these keywords to ensure diversity and clarity. Outlier removal via HDBSCAN further ensured that only thematically robust regions were retained. This layered approach enabled the construction of a coherent discourse map, setting the stage for the results section that follows, which examines the temporal trajectory and thematic structure of public sentiment surrounding the by-name *wanghong*.

Results

Temporal Evolution of the Stigmatization of Wanghong

A polarity analysis of 40,523 posts referencing *wanghong* from January 2014 to December 2023 reveals a marked shift in public sentiment. The proportion of positive polarity declined substantially from 79.80% in 2014 to 52.60% in 2023, while negative polarity increased significantly from 18.60% to 44.10% over the same period (see Figure 2). Overall, the negative-to-positive polarity ratio increased by 3.6 times, highlighting rising public skepticism and growing stigmatization of *wanghong* as a socially charged by-name.

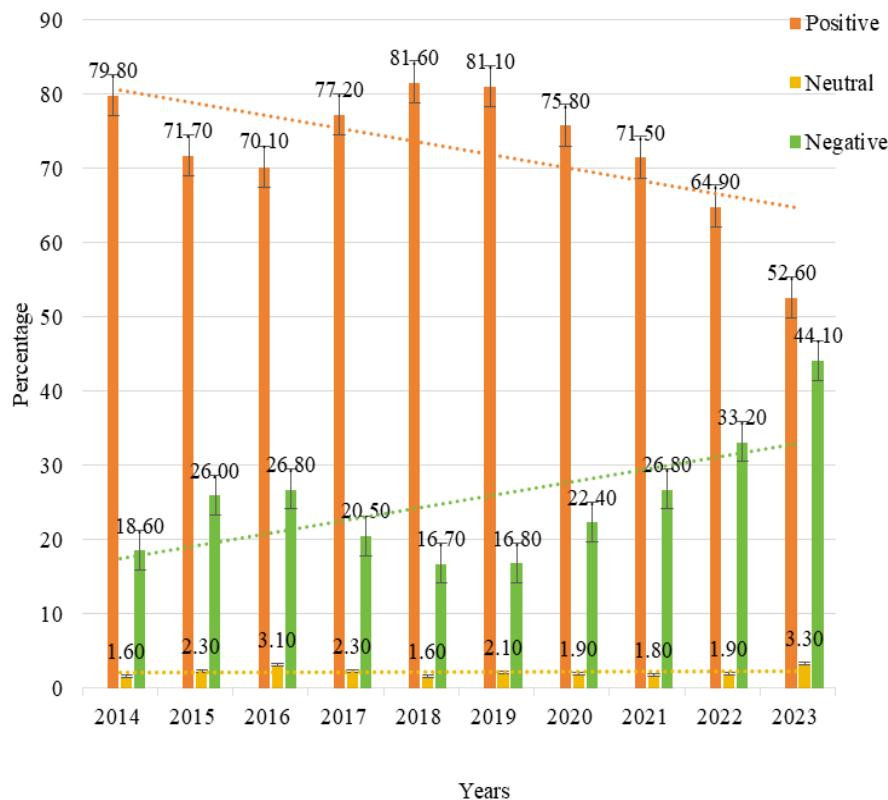


Figure 2: Temporal Changes in Sentiment Weights for *Wanghong* Posts (2014–2023)

From 2014 to 2023, the proportion of neutral polarity exhibited slight year-to-year fluctuations. However, the results of the regression analysis indicate that this variation is not statistically significant ($F = 0.613$, $p = 0.459$; see Table 2). This suggests that the level of neutral polarity remained relatively stable over the observed period, without a clear upward or downward trend over time.

Table 2: ANOVA Table for Regression Analysis of Neutral Polarity Proportion

Model	Sum of Squares	df	Mean Square	F	Sig.
Regression	0.140	1	0.140	0.613	0.459
Residual	1.600	7	0.229	—	—
Total	1.740	8	—	—	—

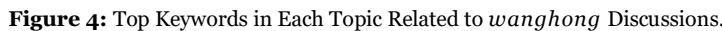
These trends offer empirical insight into the evolving sociolinguistic perception of *wanghong*. Using NLP-based polarity tracking, this study demonstrates how occupational by-names are shaped by shifting evaluative discourse and public identity over time. The approach contributes a scalable, data-driven tool for analyzing affective dimensions of naming in digital spaces.

A word cloud generated from 9,624 negative posts reveals strong dissatisfaction (shown in Figure 3), with frequent words like “foodie”, “not ok”, “disgusting”, and “boycott”. Topics include food quality issues (e.g., “greasy”, “waste”), moral criticism (e.g., “suspected”, “fraud”, “deceive”), and emotional conflicts (e.g., “arrogant”, “annoying”). Overall, these negative sentiments reflect public dissatisfaction with food safety, emotionally biased attitudes, and unethical behavior, reflecting strong criticism and discontent from consumers or certain social groups regarding certain incidents or practices.

A word cloud visualization of negative words and phrases. The words are arranged in a dense, overlapping manner, with colors ranging from dark blue to light blue. The words include: temptation, negative, overturn, the heaven, slander, lack, fuss, complain, disturbing, greedy, angry, secretive, translation, not found, abuse, outburst, conservative, late, aggressive, messy, expose, shadows, excuse, pretend, threat, trick, luxury, deceive, lie, arrogant, overbearing, fault, inferior, annoy, ugly, subvert, show off, fake, delay, overdone, trouble, restraint, funny, power, violence, willful, black gold, phlopie, fool, goss, exclusive, violation, liar, vulgar, accident, waste, bold, baffling, temper, jealousy, useless, allergy, frivolous, template.

By mapping the evaluative and ideological weight attached to these labels, the word cloud reveals how *wanghong* naming has become imbued with shifting moral and affective connotations over time, thereby offering a valuable socio-onomastic perspective on the intersection of naming practices, public sentiment, social criticism, and cultural evaluation.

The BERTopic analysis of negative texts identifies ten thematic categories, each characterized by representative keywords (see Figure 4). Keyword importance is measured using the c-TF-IDF score produced by the BERTopic model.



The topic distribution reveals how negative discourse surrounding the by-name *wanghong* is thematically fragmented yet socially saturated. Prominent clusters such as Topic 2 (Influencer Economy) and Topic 4 (Food Trends) highlight public skepticism toward commercialization and consumer deception, while Topics 3 and 10 expose deeper trust and reputational crises. Themes like fan-media tension (Topic 6), personal image (Topic 8), and live entertainment (Topic 7) reflect cultural critiques of appearance-driven fame. Meanwhile, more isolated topics such as Topic 5 (Tourism hotspots) and Topic 9 (Pet-related Controversies) indicate niche but emotionally charged discussions. Overall, the keywords portray *wanghong* not merely as a digital occupational label but as a contested social signifier, indexing public anxieties over authenticity, commodification, and moral decline in algorithmic culture. The topic similarity matrix is presented in Figure 5.

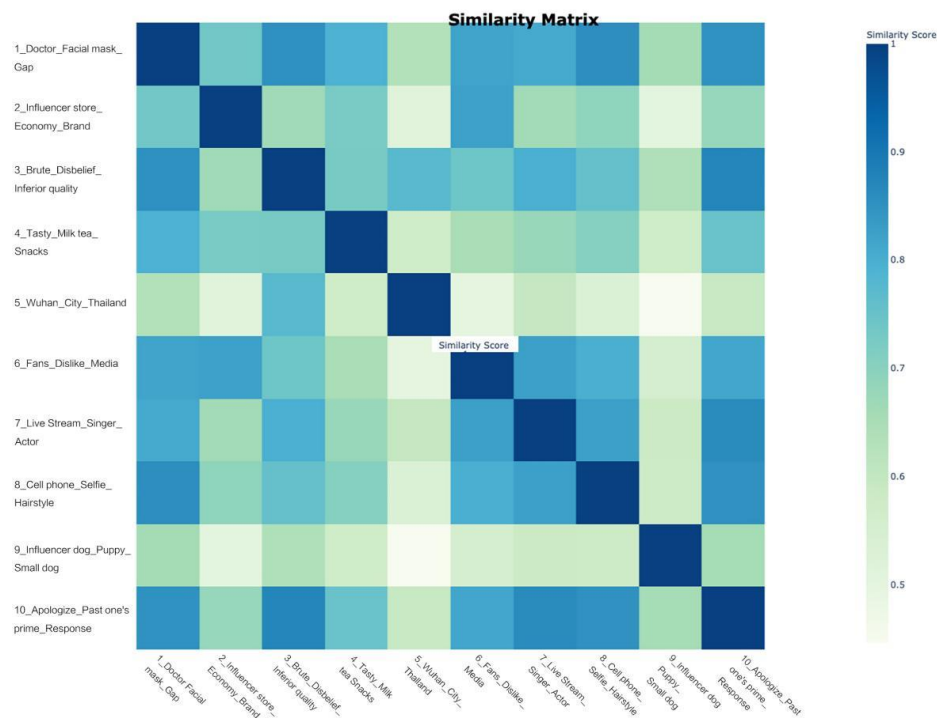


Figure 5: Topic Similarity Matrix for *Wanghong* Discussions

The heatmap in Figure 5 visualizes the structural relationships and similarity patterns between topics, offering insight into how different themes related to *wanghong* cluster or diverge in discourse. Topics 6 and 7 form a cohesive group centered on celebrity culture and fan dynamics, while Topics 3 and 10 share concerns related to product quality and reputational decline. In contrast, Topics 5 and 9 appear more peripheral, indicating isolated discussions on specific events or niche subdomains, such as pet influencers or regional controversies.

These patterns suggest that *wanghong*, as an occupational by-name, functions not merely as a referential term for internet figures but as a socially saturated sign embedded in evaluative discourse. Clusters around commerce, entertainment, and moral judgment reflect layered public attitudes. From a socio-onomastic perspective, *wanghong* reveals how contemporary by-names mediate affective positioning, shifting identities, and collective anxieties about authenticity and social value in digital culture. The matrix thus illustrates the role of by-names in organizing social meaning and negotiating communal norms in algorithmically shaped public discourse.

Summary

This note applied NLP techniques—polarity analysis and BERTopic topic modeling—to examine the evolving stigma surrounding the occupational by-name *wanghong* in China. The analysis revealed a clear rise in negative sentiment over the past decade, driven by factors such as low-quality content, aggressive marketing, and perceived social disruption. By mapping emotional shifts associated with a high-frequency by-name, the study demonstrated how NLP can contribute to socio-onomastic research, offering scalable methods for analyzing how names operate as socially responsive symbols in digital discourse.

Notes

1. In Chinese culture, the color “red” symbolizes happiness, good fortune, prosperity, and other positive values. Its rich semantic range imbues it with auspicious connotations and enduring social significance.
2. Goffman (1986) defines stigma as a “spoiled identity” rooted in discrediting attributes tied to one’s body, character, or group. These attributes diminish social status, cause psychological harm, and lead to societal exclusion and unfair treatment.
3. In both Chinese and Western contexts, influencers have increasingly distanced themselves from the label due to its associations with superficiality, commercialism, and attention-seeking behavior. In China, some content creators have argued that the term *wanghong* undermines their professional identity by reducing their work to gimmickry and marketing tactics. For example, several Chinese bloggers have publicly rejected the label, asserting that it trivializes their sustained creative efforts and reinforces stereotypes of superficiality and ostentation. Similarly, in Western media discourse, many influencers have expressed unease with the reductive nature of the term. Coverage in BBC and The Washington Post has shown that “influencer” is increasingly perceived as a pejorative label, often linked to inauthentic self-promotion and para-social marketing practices.
4. Link to Sina Weibo:
<https://weibo.com/newlogin?tabtype=weibo&gid=102803&openLoginLayer=0&url=https://weibo.com/>
5. Jieba is a Python library for Chinese word segmentation—an essential task in Chinese NLP due to the lack of spaces between words.
6. The Harbin Institute of Technology (HIT) stop words list contains high-frequency Chinese function words (e.g., conjunctions, prepositions) with limited semantic value. In NLP tasks, these are commonly removed to enhance processing efficiency and accuracy.
7. Baidu’s Sentiment Analysis API, part of its NLP services, uses AI to classify text as positive, negative, or neutral. It supports accurate sentiment analysis across Chinese-language data. The API is available at the following link: https://cloud.baidu.com/product/nlp_apply/sentiment_classify.
8. The Chinese Sentiment Lexicon developed by Dalian University of Technology is a widely used resource that contains a large number of sentiment-bearing Chinese words. Each entry is labeled with its sentiment polarity (e.g., positive or negative) and emotion category (e.g., anger, joy), enabling more detailed analysis of emotional expressions in text.

AI Disclosure Statement

No AI tools or technologies were used in conducting this research or writing this article.

References

- Bourdieu, Pierre. 1991. *Language and Symbolic Power*. Cambridge: Polity Press.
- Egger, Roman, and Joanne Yu. 2022. "A Topic Modeling Comparison Between LDA, NMF, Top2Vec, and BERTopic to Demystify Twitter Posts". *Frontiers in Sociology*, no.4: 1-16.
- Enochson, Kristine, Roberts Gregory, Sorah Michael, and Thompson Jamie. 2020. *Analyzing Newswire and Social Media Data Using Multi-Vector Sentiment Analysis*. Herndon, VA: Rosoka Software, Inc
Retrieved from
https://f.hubspotusercontent40.net/hubfs/8586094/Rosoka_October2020/pdf%20files/AnalyzingNewswireAndSocialMediaData.pdf
- Gerhards, Claudia. 2025. "Social Media Influencers as 'Dirty Workers': An Explorative Study on How They Use Strategies to Reduce the Moral Taint." *Social Media+ Society* 11.2: 20563051251348917.
- Goffman, Erving. 1986. *Stigma: Note on the Management of Spoiled Identity*. Sutton Valence: Touchatone Books: 1-24.
- Lee Carmen, and David Barton. 2013. *Language Online: Investigating Digital Texts and Practices*. London: Routledge.
- McClure, Peter. 2010. "Middle English Occupational Bynames as Lexical Evidence: A Study of Names in the Nottingham Borough Court Rolls 1303-1455. Part 1, Methodology." *Transactions of the Philological Society* 108.2: 164-177.
- Nick, I. M. 2024. "Names, Naming, Identity and the Law: A Basic Introduction." In *Names, Naming, and the Law: Onomastics, Identity, Power, and Policy*, edited by I. M. Nick, 1-18. New York: Routledge.
- Tagg, Caroline. 2015. *Exploring Digital Communication: Language in Action*. London: Routledge.
- Yan, Fangjie, and Yingjia Zhou. 2019. "Cong 'Wanghong' yu 'Wanghei' de Bianzouqu Kan Qingnian Gexing Fazhan Taishi "[A perspective on youth personality development through the variations of 'Wanghong' and 'Wanghei']". *Wangluo Sizheng*, (5), 77-81.
- Zhang, Tuo. 2019. "Wuminghua: Xinwen Baodao Zhong de Wangluo NvzhuboXingxiang Yanjiu" [Stigmatization: A Study on the Image of Online Female Anchors in News Reports]. *Shiting*, no.4: 142-143.
- Zhang, Ye, and Byron Wallace. 2015. "A Sensitivity Analysis of (and Practitioners' Guide to) Convolutional Neural Networks for Sentence Classification." *arXiv preprint arXiv: 1510.03820*.

Notes on Contributor:

Danyang Zheng, PhD, is a Lecturer in the School of Foreign Languages at Tianjin University, China. She received her PhD in Linguistics from the University of Auckland, New Zealand. Her research interests include pragmatics, semantics, and socio-onomastics, with a particular focus on naming practices and evaluative discourse in digital environments.

Correspondence to: Danyang Zheng, School of Foreign Languages, Tianjin University, Tianjin, China.
Email: danyang.zheng@tju.edu.cn