

RESEARCH ARTICLE

Geospatial analysis of toponyms in geotagged social media posts

Takayuki Hiraoka^{1*}, Takashi Kirimura², Naoya Fujiwara^{3,4,5,6}

1 Department of Computer Science, Aalto University, Espoo, Finland, **2** Department of Kyoto Studies, Kyoto Sangyo University, Kyoto, Japan, **3** Graduate School of Information Sciences, Tohoku University, Sendai, Japan, **4** PRESTO, Japan Science and Technology Agency, Kawaguchi, Japan, **5** Institute of Industrial Science, The University of Tokyo, Tokyo, Japan, **6** Center for Spatial Information Science, The University of Tokyo, Kashiwa, Japan

* takayuki.hiraoka@aalto.fi



Abstract

Place names, or *toponyms*, play an integral role in human representation and communication of geographic space. In particular, how people relate each toponym with particular locations in geographic space should be indicative of their spatial perception. Here, we make use of an extensive dataset of georeferenced social media posts, retrieved from Twitter, to perform a statistical analysis of the geographic distribution of toponyms and uncover the relationship between toponyms and geographic space. We show that the occurrence of toponyms is characterized by spatial inhomogeneity, giving rise to patterns that are distinct from the distribution of common nouns. Using simple models, we quantify the spatial specificity of toponym distributions and identify their core-periphery structures. In particular, we find that toponyms are used with a probability that decays as a power law with distance from the geographic center of their occurrence. Our findings highlight the potential of social media data to explore linguistic patterns in geographic space, paving the way for comprehensive analyses of human spatial representations.

OPEN ACCESS

Citation: Hiraoka T, Kirimura T, Fujiwara N (2025) Geospatial analysis of toponyms in geotagged social media posts. PLoS One 20(6): e0325022.
<https://doi.org/10.1371/journal.pone.0325022>

Editor: Ciro Clemente De Falco, University of Naples Federico II: Università degli Studi di Napoli Federico II, ITALY

Received: October 4, 2024

Accepted: May 5, 2025

Published: June 5, 2025

Copyright: © 2025 Hiraoka et al. This is an open access article distributed under the terms of the [Creative Commons Attribution License](https://creativecommons.org/licenses/by/4.0/), which permits unrestricted use, distribution, and reproduction in any medium, provided the original author and source are credited.

Data availability statement: All data and code necessary to reproduce the findings of this paper are deposited and publicly available at 10.5281/zenodo.13860968 and <https://github.com/takayukihir/geotagged-tweets>. Due to Twitter's Terms of Service, we are unable to redistribute the raw text of the posts we have obtained from the Twitter API. Instead, the data on the number of geotagged posts aggregated at the level of basic grid square cells are available in the above repository. The resident

Introduction

When we speak or write about geographic spaces, we generally communicate them using the names of places, or toponyms. Although any point on Earth can be specified by geographic coordinates, one would rarely refer to a place in day-to-day conversation as, for example, “35.7°N, 139.7°E”; instead, we would usually represent it by a toponym, such as *Tokyo*. The use of toponyms reflects the way we perceive and mentally structure geographic space. Unlike geographic coordinates, which are objective and unambiguous, the area a toponym refers to is often vague and difficult to define; even for the names of administratively defined areas (such as municipalities), how people use them in colloquial settings may have only a loose correspondence with the officially demarcated boundaries [1,2]. This however does not mean that the extent that each toponym denotes can be arbitrarily defined by individuals; since the main function of toponyms is to effectively communicate geographic information, there must be a shared consensus within the population that determines the areas they represent.

and employed population data are based on 2015 Population Census and 2016 Economic Census for Business Activity, respectively, and are available from the Statistics Bureau of Japan (<https://www.stat.go.jp/english/data/kokusei/2015/summary.html>, <https://www.stat.go.jp/english/data/e-census/2016/outline.html>).

Funding: N.F. was supported by JSPS KAKENHI Grant Number 24K03007 (<https://www.jsps.go.jp/english/e-grants/>) and JST PRESTO Grant Number JPMJPR21RA (<https://www.jst.go.jp/kisoken/presto/en/index.html>), Japan. The funders did not play any role in the study design, data collection and analysis, decision to publish, or preparation of the manuscript of this work.

Competing interests: The authors have declared that no competing interests exist.

The objective of this study is to understand such a spontaneous and collective relationship between geographic space and the use of toponyms. In order to quantitatively examine this relationship, one needs to collect large-scale data on the locations people associate with each toponym. The abundance of user-generated content online, especially on social media, offers a promising opportunity for this purpose. In particular, Twitter (rebranded to X in 2023), one of the major social media platforms, had provided free access to posts on the platform through APIs until 2023, allowing researchers to perform statistical analysis and find patterns and trends in the data. On Twitter, users could opt to attach to each of their post geolocation metadata (a *geotag*) that represent the geographic coordinates of the GPS location of their device. These user-annotated geotags as well as the content of the posts allow us to explore the spatial dimension of user behavior and language use. For instance, by leveraging the fact that the set of vocabulary that appears in the text of posts varies according to the whereabouts of users, studies have found that toponyms in the text can be disambiguated [3] and the location of individual users can be identified [4,5], although there are limitations to this approach at high spatial resolution [6]. When aggregated at the population level, the data can be used to study dialectal variation and language evolution, making it of interest to sociolinguistics and linguistic geography [7–12].

We note that the geotag attached to a post does not necessarily correspond to the place referred to in the text [13,14]. This discrepancy can arise from inaccurate tagging [15], but it can also be a manifestation of the geographic awareness and identity of the user, i.e., how users perceive and choose to represent their location or the place under discussion [16–18]. Users may (mis)represent their location for various reasons, such as privacy concerns, a desire to be associated with a particular place, or simply because they are referring to a location that is not their own.

In this study, we focus on understanding the geographic distribution of geotagged posts and the toponyms they contain. Our goal is to understand the collective, population-level knowledge about the relationship between toponyms and geographic locations rather than to identify patterns of toponym usage at the individual user level. Similar questions have been addressed in the literature: Hollenstein and Purves [19] used user-annotated geotag data from the image hosting service Flickr to delineate city centers and neighborhoods; Hu and Janowicz [20] studied the names of points of interest (such as restaurants and shops) registered on Yelp, a user-generated local business review platform, in metropolitan areas in the United States. These studies were primarily focused on urban areas, and therefore the analysis was limited to a relatively small geographic scale. In addition, the results were rather descriptive and were not aimed at deriving general laws of toponymic occurrence. In contrast, we use a data-driven modeling approach to uncover the underlying principles governing the occurrence of toponyms of different granularities on a larger geographic scale, aiming to understand how these distributions reflect collective spatial cognition.

The rest of the paper is structured as follows. We start by introducing the geotagged Twitter dataset and how we collect, preprocess, and subsample it in the *Data* section. In the *Results* section, we first present the inhomogeneity of the geographic distribution of geotagged posts, and compare it to the population distributions. We then focus on the geographic distribution of individual toponyms. In the second subsection, we observe that the occurrence of each domestic toponym follows a characteristic pattern, hinting at its spatial specificity. To formalize this observation, we introduce a class of models called binomial models. In the third subsection, we use the simplest instance of this model to quantify the spatial specificity of each toponym. In the last subsection, we show that another variant of the binomial model,

which we call the core-periphery model, reproduces the essential elements of toponym occurrence patterns despite its simplicity. In the *Discussion and Conclusions* section, we discuss the implications of our work.

Data

We collected 395,268,777 geotagged Twitter posts from the Twitter API. These posts are annotated with coordinates within the bounding box of Japan (latitude between 20.43°N and 45.56°N and longitude between 122.93°E and 153.99°E) during the period from 1 February 2012 to 30 September 2018. Note that this bounding box also includes South and North Korea as well as parts of China and Russia. We note that the data collection and analysis method in this study complies with the terms and conditions of all data sources. Twitter's terms of service allow researchers to analyze and publish findings based on Twitter data, but prohibit the redistribution of raw data, such as the text or geotags of individual posts.

Data collection was followed by several preprocessing stages. First, we excluded posts made via location check-in services (e.g., Foursquare) as well as those generated by automated bots or manipulative users. Posts made through Foursquare are in specific text formats, such as "I'm at [the name of a place/point of interest]" or "[post text] (@ [the name of a place/point of interest])". When a user checks in to a place, they select the name of the place/point of interest, which may include toponyms, from a list of nearby places provided by Foursquare. However, users do not have independent control over the geographic coordinates tagged to the post; these coordinates are automatically determined by Foursquare's location database based on the name of the place chosen. Although these posts are created by personal users, they do not represent an organic association between toponyms and geographic coordinates. Including them in the analysis would bias the findings of this study.

In addition, the geotags attached to bot-generated posts may not reflect the geospatial behavior of individual users. Many non-personal accounts generate geotagged posts for various purposes, not necessarily with commercial or malicious intent; for example, they may be bots that send out weather or traffic alerts [21]. Most of them are identifiable by the 'source' metadata, which indicates the application or device used to make the post. Some accounts engage in more active geotag manipulation. Zhao and Sui [22] developed a technique to detect manipulated geotags and concluded that such manipulations accounted for 0.22% of their sample, with even lower percentages among posts from official Twitter clients for iPhone and Android.

Based on these considerations, we restricted our dataset to posts originating from official or general-use third-party mobile applications using the source metadata. This filtering process excluded 29.87% of the collected posts—20.06% from location check-in services and 9.81% from other sources. A full list of the applications considered for inclusion in the dataset is available in [S1 Appendix](#) in Supporting Information. We further excluded posts tagged outside the geographic range of our study but accidentally included in the collected data; this reduces the sample size by an additional 0.46%.

Additionally, we found evidence that some users likely manipulated the geotags of their posts and assigned random geographic coordinates in unnatural rectangular bounding boxes that partly extend over sea areas (see [S1 Appendix](#) in Supporting Information). These posts can be characterized by containing a large number of mentions to other accounts, typically seven or more. Although not all posts with many mentions are necessarily manipulated, we conservatively excluded all posts with seven or more mentions. We confirmed that removing these posts, which account for 0.06% of the sample, did not significantly alter any of our findings.

After these filtering steps, the dataset consists of 275,750,003 posts. The geolocation meta-data of each post is aggregated into grid cells based on the standard grid square system used in Japan's official spatial statistics. Each grid cell spans 30'' of latitude and 45'' of longitude, which is approximately a square with a side length of 1 km, although the east-west width of a grid cell varies slightly with latitude. In total, 293,444 grid cells contain a nonzero number of posts in the dataset.

As of 2019, Japan is divided into 47 prefectures as the first level of administrative division, and 1741 municipalities (cities, towns, villages, and twenty-three special wards of Tokyo) as the second level of division. In addition, some large cities, referred to as "designated cities", have administrative, non-autonomous subdivisions known as wards. Prefectures are sometimes grouped into seven to thirteen regions, although regions are not official administrative units, and the name and extent of each region can be ambiguous.

From the full sample of the 276 million geotagged posts, we extract the subset of posts that contains each of 24 Japanese toponyms that refer to regions, prefectures, cities, and wards (special wards of Tokyo Metropolis and wards of designated cities) in Japan. The list of toponyms sampled in this work and the administrative areas they refer to can be found in Fig 1.

Even if the text of a post contains a string that matches a toponym, it does not necessarily mean that the user is referring to the place it denotes. For instance, the name of the city Hiroshima is a substring of the name of another city Kitahiroshima, which is located over 1200 km away. As a result, posts containing *Hiroshima* include not only posts that refer to Hiroshima but also posts that mention Kitahiroshima. To separate the references to Hiroshima from the references to Kitahiroshima, we need to exclude the posts containing *Kitahiroshima*. For each of the 24 toponyms we study in this paper, we discounted the posts that include the names of other regions, prefectures, cities with a population larger than 50,000, or wards that contain the toponym as a substring. We provide further details in S1 Appendix in Supporting Information.

In addition to these domestic toponyms, we extracted posts that contain twelve common Japanese nouns and six Japanese toponyms that refer to places outside Japan to construct reference datasets. We refer to the samples for individual keywords (toponyms and nouns) as *keyword subsamples*. In the following, we only use the information about the number of posts tagged inside each grid cell for each sample, and disregard the content of the text or user information. Refer to Table 1 for the notation and definition of variables used in this work.

Results

Spatial distribution of geotagged posts

Let us first study the geospatial distribution of the full sample of geotagged posts before looking at the distribution of each toponym subsample (Fig 2). The geolocations to which the sampled posts are tagged are not uniformly distributed within the observed area, as shown in Fig 2A. A large number of posts are concentrated in relatively few grid cells in large metropolitan areas, such as the city centers of Tokyo and Osaka, while few posts are found in most of the grid cells. The heterogeneity of the spatial distribution of geotagged posts is also evident from the heavy-tailed probability distribution of the number of posts sampled in each grid cell (Fig 2D). Namely, it is characterized by two power laws with different exponents: the bulk part is characterized by an exponent of approximately -1.2 while the tail part follows a steeper power law with an exponent of about -2.8 .

To investigate the origin of this heterogeneous distribution, we compare the geotagged post statistics with census data for the resident population in 2015 [25] and the employed

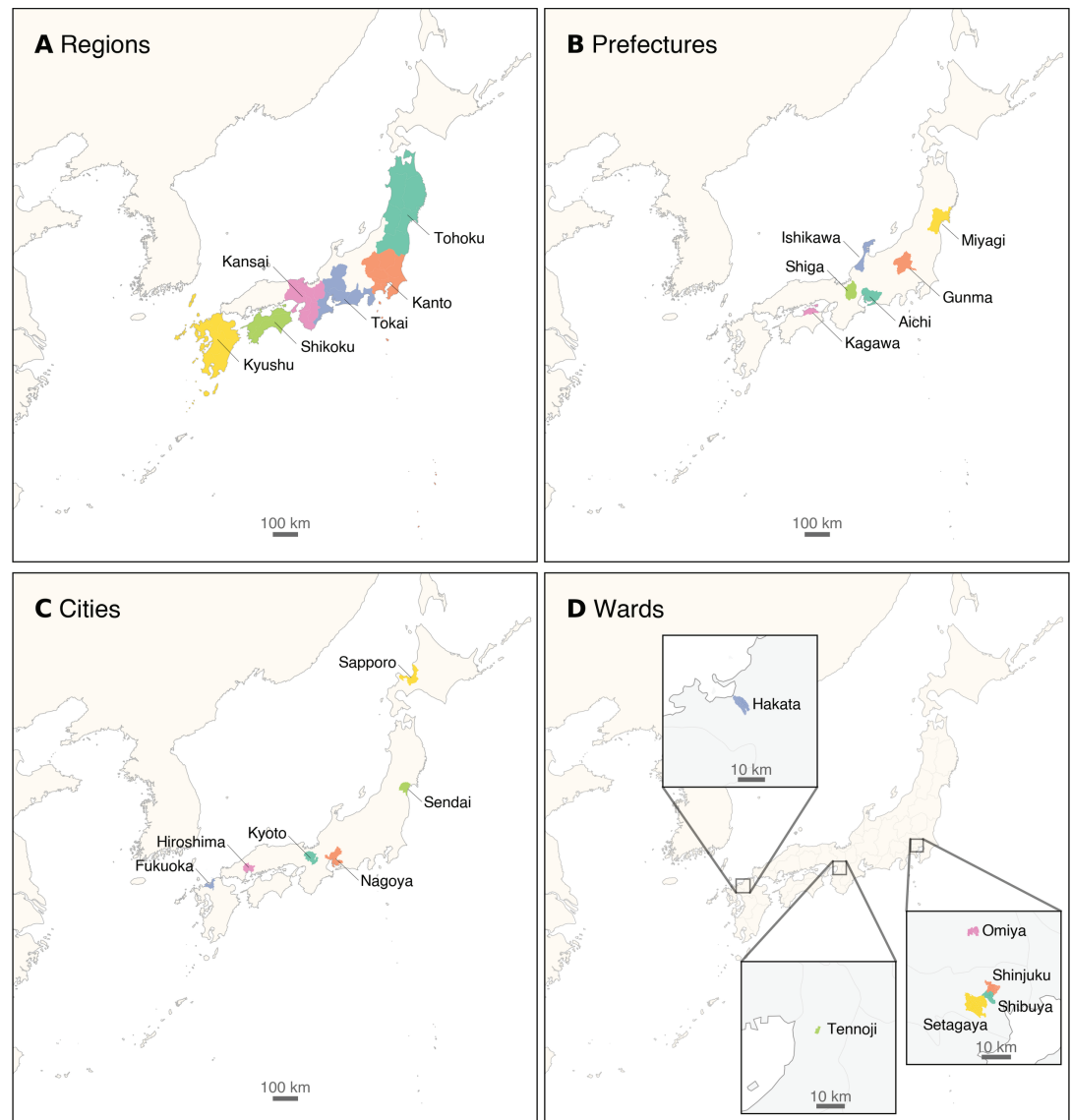


Fig 1. Places in Japan denoted by the toponyms studied in this paper. In this study, we sample 24 toponyms: the names of (A) six regions, (B) six prefectures, (C) six major cities, and (D) six wards (submetropolitan/submunicipal districts). The colored areas in each panel show the administratively defined geographic area (except for cities) denoted by each toponym. (A) The extent of each region is not uniquely defined. Here we show one of the commonly used classifications of regions. (C) The colored area shows the metropolitan employment area [23,24], which is considered to be more representative of urban activity than administratively defined city areas. Note that Kyoto, Hiroshima, and Fukuoka are used both as the names of the cities and as the names of the prefectures of which the cities are the capitals. (D) Prefectural boundaries are shown for visual guidance. Maps made with Natural Earth (<https://www.naturalearthdata.com/>).

<https://doi.org/10.1371/journal.pone.0325022.g001>

population in 2016 [26]. The employed population is defined as the number of permanent or temporary employees, self-employed individuals, contractors, and unpaid staff in family-owned businesses, whose main working site is located in each grid cell. From Fig 2B and 2C, one can see the similarity in the spatial distribution between the densities of geotagged posts, resident population, and employed population. In fact, the geotagged post density is strongly correlated with the population densities (Fig 2E, 2F). The correlation with the post

Table 1. Variables used in this work. Subscript c for the grid cell index may be omitted when it is clear.

Symbol	Definition
w	Keyword on which the dataset is subsampled
c	Index of grid cell
$n_{\text{all},c}$	Number of all posts tagged to grid cell c
$n_{w,c}$	Number of posts containing w and tagged to grid cell c
N_{all}	Total number of geotagged posts in the dataset
N_w	Total number of geotagged posts containing keyword w ; $N_w = \sum_c n_{w,c}$
A_c	Area of grid cell c
$\sigma_{\text{all},c}$	Density of all posts tagged to grid cell c ; $\sigma_{\text{all},c} := n_{\text{all},c}/A_c$
$\sigma_{\text{res},c}$	Resident population density in grid cell c
$\sigma_{\text{emp},c}$	Density of employed population (number of personnel working at business sites) in grid cell c
$\sigma_{w,c}$	Density of posts containing keyword w and tagged to grid cell c ; $\sigma_{w,c} := n_{w,c}/A_c$
$\phi_{w,c}$	Occurrence ratio of keyword w among all posts tagged to grid cell c ; $\phi_{w,c} := n_{w,c}/n_{\text{all},c}$
O_w	“Center” of the distribution of toponym w , defined by Eq (1)
$d_{w,c}$	Distance between the center O_w and grid cell c
$p_{w,c}$	Probability that keyword w is contained in a post tagged to grid cell c in the binomial model

<https://doi.org/10.1371/journal.pone.0325022.t001>

density is stronger for the employee population density than for the resident population density, both in terms of Pearson's correlation and Spearman's rank correlation. This may suggest that social media posts are made more in the daytime than at night. Fig 2G shows the probability density functions of the geotagged post density, the resident population density, and the employee population density collapse to one another when they are scaled by the average density, strongly suggesting that the uneven distribution of population gives rise to the heterogeneity in the spatial distribution of geotagged posts.

Spatial distributions of toponym subsamples

We now turn to our main question: How are geotagged posts containing each toponym, i.e., each *toponym subsample*, spatially distributed? To address this question, we show the results for *Fukuoka*, a toponym referring to the sixth largest city in Japan as of 2019, located in the southwestern part of the country, as an illustrative example.

Fig 3A shows the spatial distribution of the occurrence of *Fukuoka*, while Fig 3C shows the probability density functions of the gridwise occurrence density σ_w . Both figures suggest that the *Fukuoka* occurs heterogeneously across different grid cells. The same observation can be made for other toponyms, as shown in S1 Fig in Supporting Information. It is noteworthy that toponym subsamples of different sizes follow the same scaling, which implies that the spatial heterogeneity does not depend on the popularity of toponyms. On the other hand, each distribution is characterized by a single power law, which is a clear deviation from the scaling observed for the full sample.

While the spatial density (occurrence per unit area) of posts containing the word *Fukuoka* is high in *Fukuoka* and neighboring areas, it is also high in other large metropolitan areas such as Tokyo and Osaka, which are geographically distant from *Fukuoka*. However, this is presumably just by chance due to the large total number of posts tagged in these areas. To account for the difference in total sample size n_{all} in each grid cell, we normalize the toponym subsample size n_w by n_{all} and obtain the fraction of occurrence of the toponym. We clearly see that posts contain the word *Fukuoka* with higher probabilities in the region around the city of *Fukuoka* (Fig 3B).

To further understand the relationship between the full sample and the toponym subsamples, we create a scatter plot for each toponym as in Fig 3D, where the horizontal axis

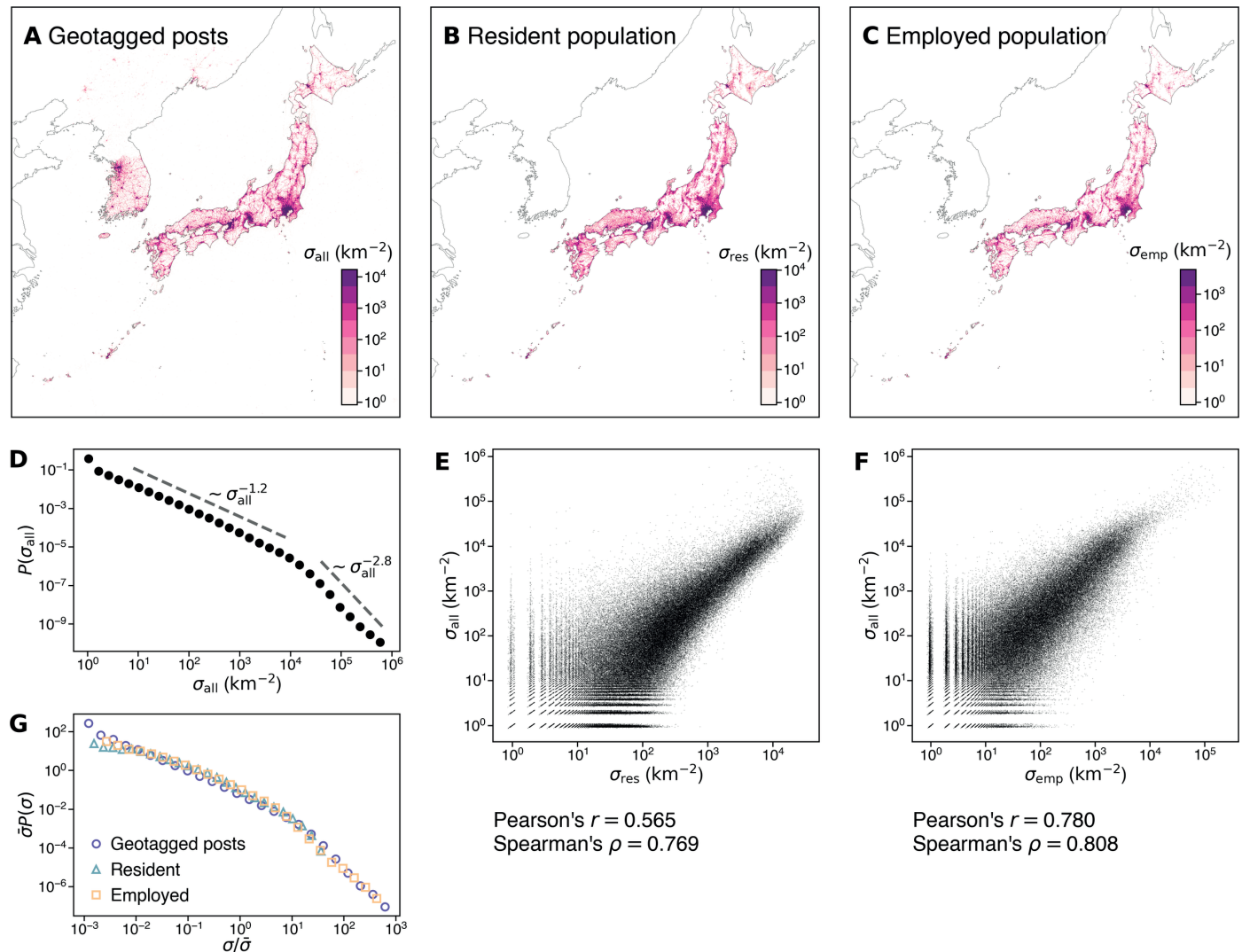


Fig 2. The spatial densities of geotagged posts, resident population, and employed population. (A–C) Geographic distribution of the three densities. Note that the population data are geographically limited inside Japan, while the geotagged posts are sampled in the bounding box of Japan, which also includes neighboring countries. Maps made with Natural Earth (<https://www.naturalearthdata.com/>). (D) Probability distribution of density (the number of geotagged posts per unit area). (E, F) Scatter plots showing the correlation between each of the population densities and the geotagged post density. Pearson and the Spearman correlation coefficients are shown below the plot. (G) Probability distributions of the three densities, each rescaled by its mean.

<https://doi.org/10.1371/journal.pone.0325022.g002>

represents n_{all} and the vertical axis corresponds to n_w . Apart from the general trend that n_w increases with n_{all} , we see that there are two distinct branches of scaling, with one increasing faster than the other. That is, n_w increases as a function of n_{all} in two different ways. This two-branch scaling behavior is not specific to *Fukuoka* but is widely seen for different toponyms across various degrees of popularity and granularity (S2 Fig in Supporting Information). In contrast, keyword subsamples for common nouns such as *wallet* and toponyms that refer to places outside the observed area such as *Hawaii* do not exhibit heterogeneity in spatial distribution, and the relationship between n_{all} and n_w only shows a single scaling (see Fig 4B and 4D; see also S2 Fig for the scatter plots for all keywords sampled in this study). This suggests that the two-branch scaling is unique to the toponyms that refer to places within the observed

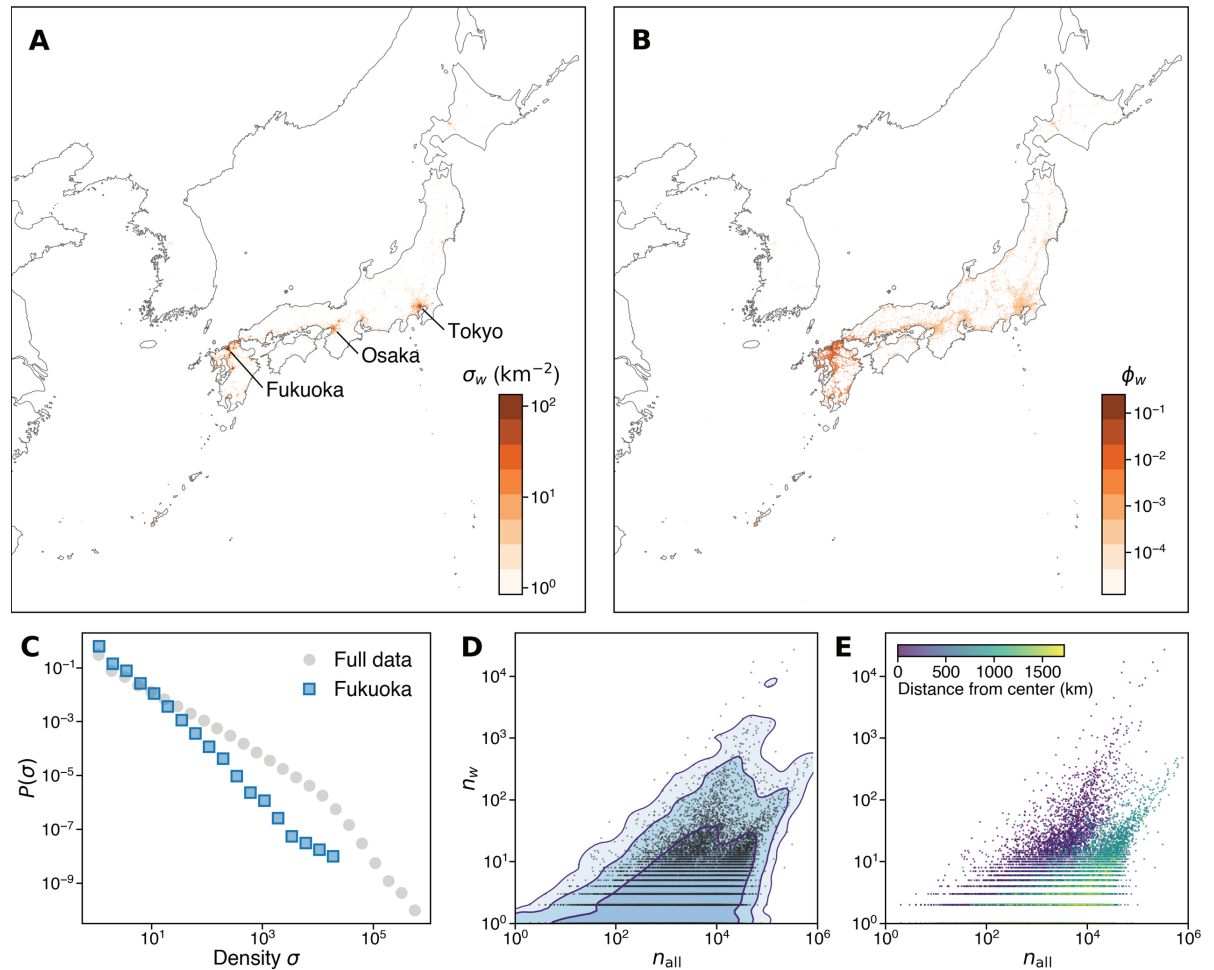


Fig 3. Occurrence pattern of Fukuoka. (A) Geographic distribution of density σ_w . (B) Geographic distribution of occurrence ratio ϕ_w . (C) Probability distributions of spatial density for all geotagged posts (in gray) and for Fukuoka (in blue). (D) Scatter plot showing the relationship between the total number of geotagged posts n_{all} and the number of posts containing Fukuoka n_w in each grid cell. The lines represent contours along which the density of points (kernel density estimate) on the double logarithmic scale is constant. The region inside each contour, from dark to light colors, contains 90.0%, 99.0%, and 99.9% of the data points, respectively. (E) The same scatter plot as panel D, but with points colored by the distance between the grid cell and the center O_w . Maps made with Natural Earth (<https://www.naturalearthdata.com/>).

<https://doi.org/10.1371/journal.pone.0325022.g003>

area and is not present in other types of words. Naturally, we expect that such a pattern stems from the geospatial specificity of toponym use.

In particular, we hypothesize that one of the two scaling branches observed for toponym subsamples corresponds to the distribution of the toponym that is spatially specific to the area it refers to, while the other branch represents the use of the toponym that is not spatially specific, similar to that of common nouns. To test this hypothesis, we first need to identify the area that the toponym refers to. Instead of relying on external sources such as gazetteers, we do this in a data-driven way: we define the *center* of the geographic distribution of toponym w as the grid cell with the highest frequency of w :

$$O_w = \operatorname{argmax}_{c \in S_w} n_{w,c} \quad (1)$$

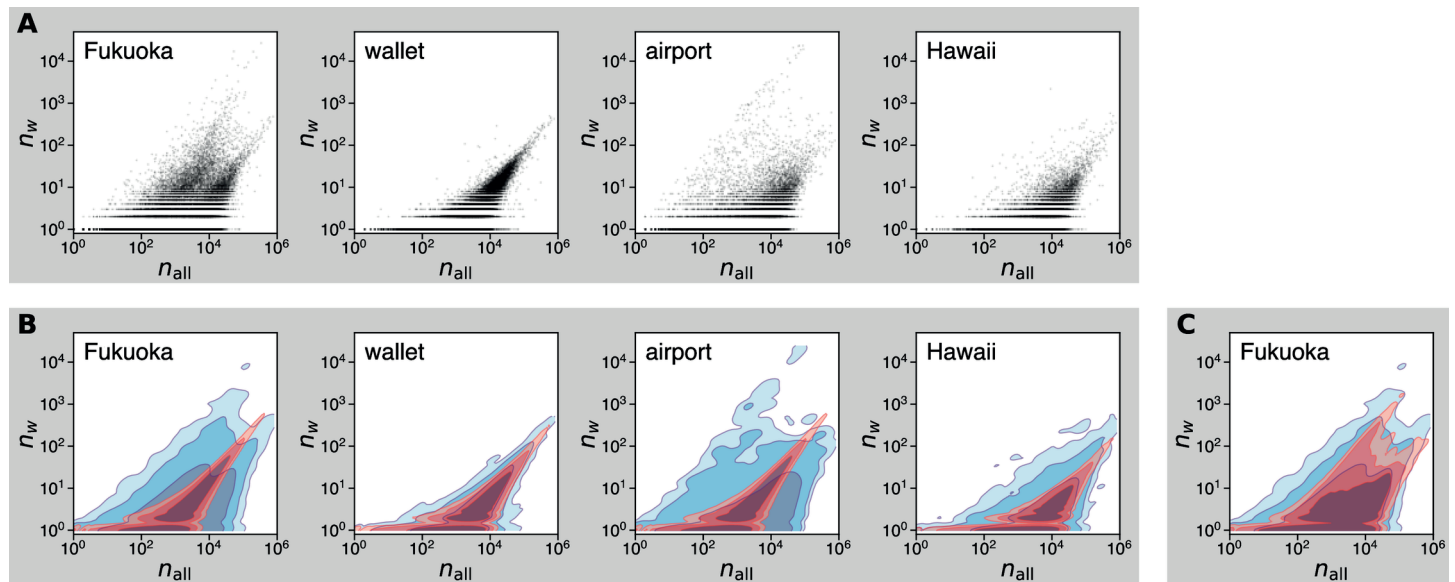


Fig 4. Relationship between n_w and n_{all} in empirical and model distributions. (A) Scatter plots for different keywords. (B, C) Kernel density profiles of empirical and model distributions. In each panel, the red contours show the density profile obtained from the location-independent model (B) or the core-periphery model (C), fitted to the empirical occurrence pattern of each word, represented by the blue contours. As in Fig 3D, the region inside each contour, from dark to light colors, contains 90.0%, 99.0%, and 99.9% of the data points, respectively. Note that the scatter plot of empirical data and its density profile for *Fukuoka* are identical to those in Fig 3D.

<https://doi.org/10.1371/journal.pone.0325022.g004>

where $S_w = \{c \mid \phi_{w,c} \geq \epsilon\}$ is the set of cells with an occurrence ratio equal to or greater than ϵ . This condition prevents a cell from being selected as the center merely due to the large total number of posts. Here we set $\epsilon = 0.01$, i.e., toponym w must appear in at least 1% of all posts tagged to a cell for the cell to be included in S_w . As we can see in Fig 3E, the grid cells that constitute the two branches are clearly distinct in terms of distance from the center of the toponym subsample. The branch with faster scaling is geographically closer to the center while the branch with slower scaling is relatively far from the center. This corroborates our hypothesis that the two-branch scaling arises from the spatial specificity of the toponym.

Location-independent model

The results in the previous section give us an intuitive understanding of the geographic distribution of toponyms. In the next two subsections, we present a model-based analysis to validate our intuition and to characterize the empirical observation in a quantitative way. Specifically, we introduce a model where each post tagged to location c contains word w with probability $p_{w,c}$ that may depend on c . We assume that the occurrence of word w in each post is an independent event, i.e., a Bernoulli trial. In this model, the number $n_{w,c}$ of posts in grid cell c that contains word w out of $n_{all,c}$ total posts follows a binomial distribution:

$$P(n_{w,c} \mid n_{all,c}, p_{w,c}) = \binom{n_{all,c}}{n_{w,c}} p_{w,c}^{n_{w,c}} (1 - p_{w,c})^{n_{all,c} - n_{w,c}}. \quad (2)$$

This model essentially posits that the underlying mechanism of the occurrence of a word can be summarized by probability $p_{w,c}$ specific to each grid cell c . Importantly, it relies on the simplifying assumption that the occurrence of a word in one post is independent of its occurrence in other posts, and that the dependence between the occurrence of different words

within the same post can also be neglected. Variations of this binomial model are obtained by adopting different formulations of occurrence probability $p_{w,c}$.

We first examine if the empirical pattern can be explained by the simplest version of the binomial model which assumes that $p_{w,c} = p_w$ for all grid cell c ; that is, the toponym occurs with a constant probability independent of location. The unbiased and maximum likelihood estimator for p_w can be obtained simply by dividing the size of the toponym subsample by the number of all geotagged posts: $\hat{p}_w = N_w/N_{\text{all}}$.

Fig 4A shows the distribution of occurrence of *Fukuoka* against n_{all} in empirical data and the expectation from the location-independent binomial model. We observe that the empirical distribution is not consistent with the model. In particular, the model does not reproduce the two-branch scaling behavior seen in the empirical data and exhibits a single scaling instead. We confirm the same kind of discrepancy for all the domestic toponyms we studied (see S3 Fig in Supporting Information for the full results).

The implication of this observation becomes apparent through juxtapositions against patterns for common nouns and foreign nouns. The location-independent model shows a close agreement with the empirical data for *wallet* (Fig 4B), implying that the word indeed occurs at a constant rate in every grid cell. In general, we find that the empirical data and the location-independent model are in fairly good agreement for many common nouns that refer to objects (e.g., *telephone*) or abstract concepts (e.g., *society*). For a comparison between the model and the empirical distributions for these words, see S3 Fig in Supporting Information. On the other hand, there is a large deviation between the empirical and model distributions for the word *airport* (Fig 4C), suggesting that the occurrence at a constant rate is not a shared characteristic among all common nouns. This can be explained by the fact that common words such as *airport* are often used in combination with toponyms (e.g. *Narita Airport*), and are therefore more likely to be used in specific places. Finally, for foreign toponyms, such as *Hawaii* shown in Fig 4D, we observe the single scaling pattern in both the empirical and model distributions, although the empirical distributions generally show slightly greater variance than the model. While these toponyms are not semantically associated with a specific place in the observed area, they may occur with higher probability around international airports from which people travel to the places these toponyms refer.

Beyond visual inspection, the discrepancy between the location-independent model P and the data can be quantified using relative entropy, also known as the Kullback–Leibler (KL) divergence. Intuitively, relative entropy measures the dissimilarity from one probability distribution to another. Here, we wish to quantify the dissimilarity of the empirical data from the model, each of which is a probability distribution on the total number of posts in each cell and the number of posts with word w in each cell. Specifically, we employ a modified version of relative entropy, denoted by $D_{\text{KL}}(\tilde{Q}_w \parallel P_w)$, that takes into account the cells with at least one occurrence of w . We elaborate on this particular definition and provide the rationale for our choice in S1 Appendix in Supporting Information.

As shown in Fig 5, relative entropy $D_{\text{KL}}(\tilde{Q}_w \parallel P_w)$ clearly differentiates domestic toponyms and place nouns from foreign toponyms and common nouns without place connotations. All domestic toponyms except *Kanto* and all place nouns (*park*, *university*, *hotel*, *airport*, and *shrine*) are characterized by relatively large values of relative entropy, namely $D_{\text{KL}}(\tilde{Q}_w \parallel P_w) > 4$, implying that the empirical distribution is highly dissimilar from the location-independent binomial model for these words. Common nouns such as *trip* and *vegetable* are on the borderline ($D_{\text{KL}}(\tilde{Q}_w \parallel P_w) \simeq 4$); relative entropy for other common nouns and foreign toponyms are significantly smaller.

The comparison with the location-independent model reveals how sensitive and specific the occurrences of toponyms and common nouns are to geographic locations. For toponyms,

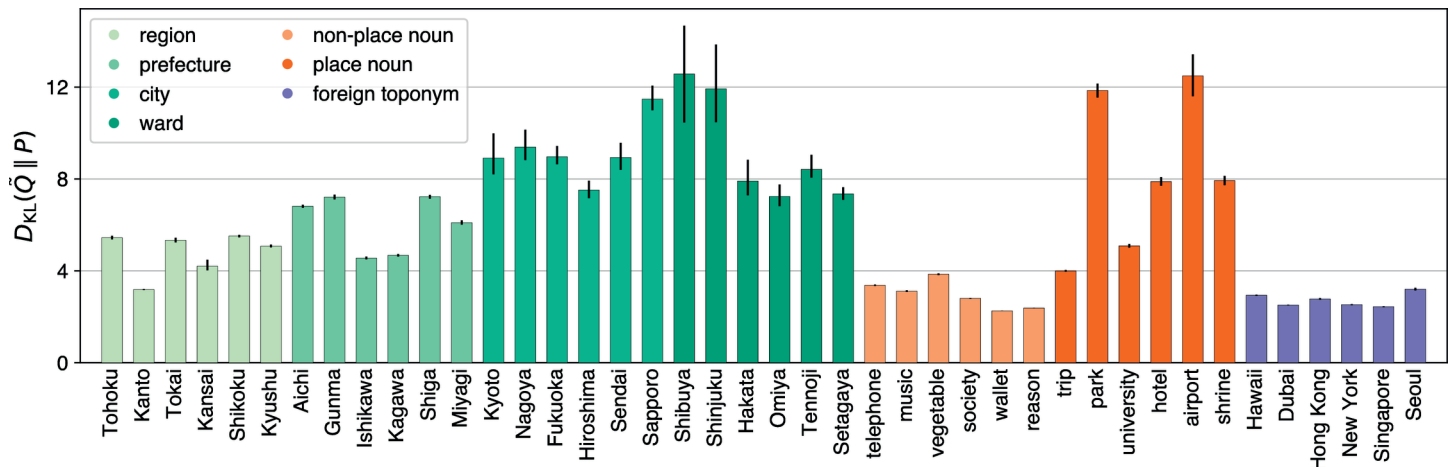


Fig 5. Dissimilarity of empirical data from the location-independent model. The dissimilarity is evaluated by relative entropy $D_{KL}(\tilde{Q}_w \parallel P_w)$. For each word, 3000 grid cells are randomly sampled 50 times. The error bar shows the 95% confidence interval.

<https://doi.org/10.1371/journal.pone.0325022.g005>

the spatial specificity may be readily observable by visualizing the geographic distribution of occurrence ratio ϕ_w , as in Fig 3B. However, identifying location-specific common nouns may be more subtle, as these nouns are usually associated not with a single place but with multiple places across the observed area. As a result, their geographic distributions may appear visually indistinguishable from those of non-place nouns. In such cases, quantifying the dissimilarity from the binomial model through $D_{KL}(\tilde{Q}_w \parallel P_w)$ can serve as a good indicator of the geospatial specificity of word occurrence.

Core-periphery model

So far, we have shown that the location-independent binomial model cannot reproduce the large variance and two-branch scaling behavior seen in the empirical distribution of domestic toponyms. Let us now take a step toward realism and discuss a more flexible modeling framework to account for the empirically observed geographic distribution of toponyms. We consider the *location-dependent* binomial model, that is, we assume that the occurrence probability $p_{w,c}$ in Eq (2) can vary for different grid cells.

Let us recall that the results in Fig 3E suggest that a post in grid cell c is more likely to contain a toponym w if c is within the area that w refers to. Indeed, the occurrence ratio $\phi_{w,c}$ plotted against the geodesic distance $d_{w,c}$ between the center O_w and cell c shows an overall decreasing trend (Fig 6A). The binned averages of $\phi_{w,c}$ imply that the occurrence probability decays as a power-law function of distance from the center, especially at long distances. Moreover, by grouping toponyms according to their administrative level as shown in Fig 6B, one can see that the onsets of the power law differ according to the granularity of each toponym. For toponyms referring to higher-level units, such as regions and prefectures, which are typically larger in area, the power-law decay starts at larger distances while the average occurrence ratios are relatively stable at smaller distances. Conversely, toponyms denoting small administrative units, such as wards, exhibit power-law behavior starting at short distances with no noticeable plateau regime.

Motivated by these observations, we propose a simple location-dependent variant of the binomial model that assumes the presence of a *core*, the area in which the occurrence probability $p_{w,c}$ is high, and a *periphery*, grid cells that are geographically distant from the

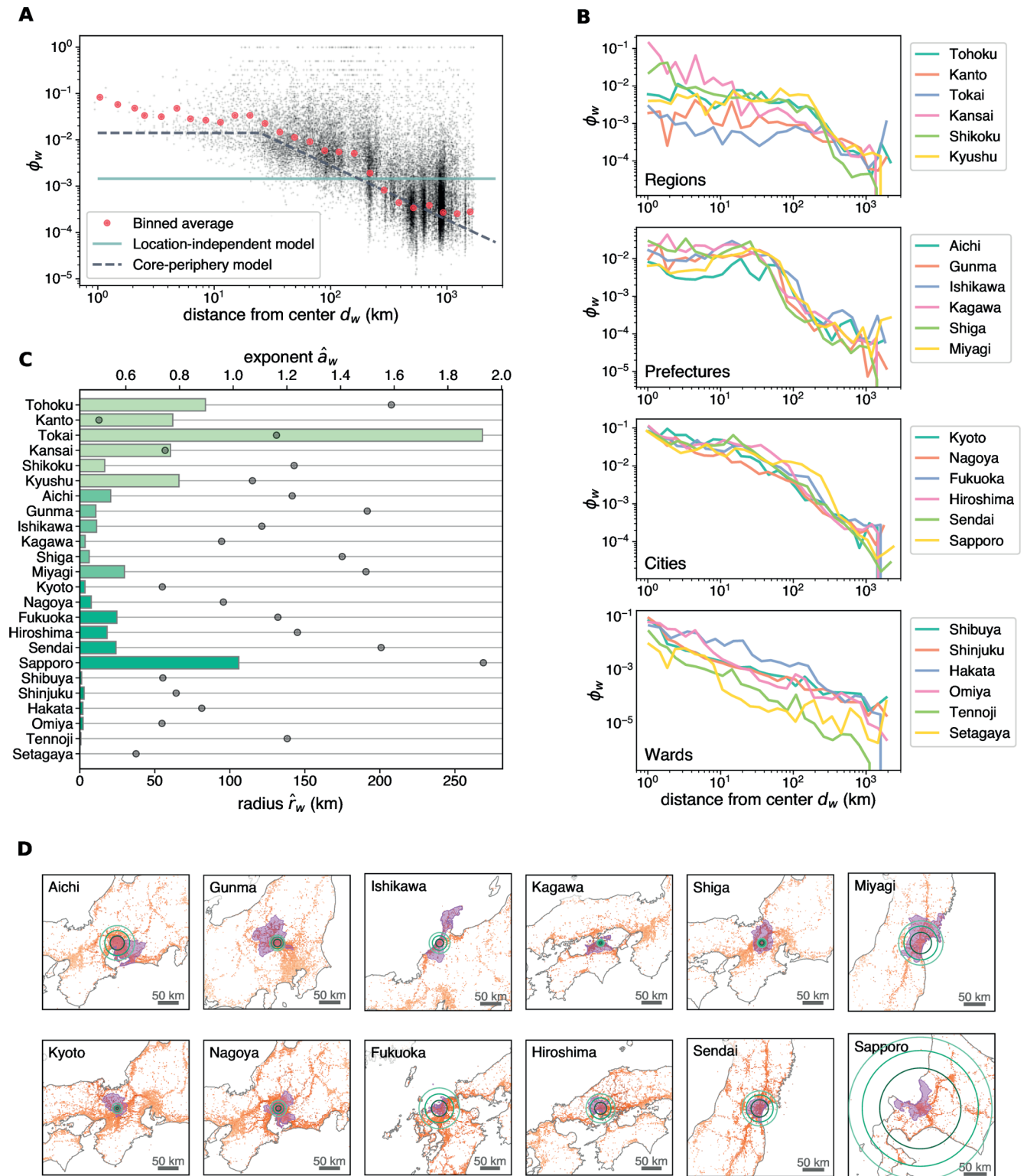


Fig 6. Core-periphery patterns of toponym occurrence. (A) Occurrence ratio ϕ_w of Fukuoka against distance d_w from center O_w (small black dots), overlaid with the average for each logarithmic bin (red circles). The solid and dashed lines represent the maximum likelihood fits of the location-independent and core-periphery binomial models. (B) Average occurrence ratio as a function of d_w for all the domestic toponyms studied in this work. (C) Maximum likelihood estimator of the core-periphery model parameters for each toponym. We represent estimated radius $\hat{r}_w$ by bars colored according to the category of the toponym (lower axis) and estimated exponent $\hat{a}_w$ by gray circles (upper axis). The standard errors are omitted as they are too small to be

meaningfully visualized. (D) The fitted core-periphery model compared to the administrative/metropolitan area. The innermost circles in dark green represent the core boundary (distance r_w from the center O_w) and the two outer circles in lighter green denote the distance at which the occurrence probability $p_{w,c}$ is equal to one half and one third of the probability in the core q_w , respectively. The areas shaded in purple indicate the administrative area of each prefecture (top row) and the metropolitan employment area of each city (bottom row). Maps made with Natural Earth (<https://www.naturalearthdata.com/>).

<https://doi.org/10.1371/journal.pone.0325022.g006>

center and characterized by lower $p_{w,c}$. Namely, the occurrence probability $p_{w,c}$ in Eq (2) is assumed to be a function of $d_{w,c}$ that is constant within a certain distance from the center and to decrease as a power law as a function of distance from the center outside of this range:

$$p_{w,c} = \begin{cases} q_w & \text{for } d_{w,c} \leq r_w, \\ q_w \left[\frac{d_{w,c}}{r_w} \right]^{-a_w} & \text{for } d_{w,c} > r_w. \end{cases} \quad (3)$$

This model has three free parameters: r_w is the radius of the core, q_w is the occurrence probability in cells within distance r_w from the center O_w , and a_w denotes the exponent of the power-law decay outside the range. Fig 7 shows a schematic diagram of this model.

For each toponym w , the values of these parameters can be estimated by maximizing the log-likelihood function

$$\mathcal{L}(q_w, r_w, a_w) = \sum_c \log P(n_{w,c} | n_{\text{all},c}, p_{w,c}). \quad (4)$$

To perform maximum likelihood estimation, numerical optimization is carried out using SciPy's `scipy.optimize.minimize` function with the Nelder-Mead algorithm as the solver. To reduce the risk of the solution being trapped in a local minimum, the parameter r_w is initialized with four distinct values: 10 km, 20 km, 40 km, and 80 km. An example of the fitted model is visualized in Fig 6A, along with the values of ϕ_w and the fitted location-independent model.

The fitted model shows, in general, a better agreement to the data in terms of the relationship between n_{all} and n_w , as shown in Fig 4B. For the results for all domestic toponyms studied in this work, see S4 Fig in Supporting Information. In particular, this model reproduces the two-branch scaling behavior of the empirical data. This significant improvement from the location-independent model is also evidenced quantitatively by the decrease in the Akaike Information Criterion (AIC) for all the toponyms studied in this work (Fig 8).

The advantage of the core-periphery model is that it conforms to an intuitive interpretation: the radius r_w can be seen as the extent of the area to which the toponym refers, i.e. the core. Outside this core area, users refer to the place less often, but the decrease in probability is gradual as a function of distance and slow enough to be modeled as a power law (compared to, e.g., an exponential decay). To verify this interpretation, we compare the estimated core (the area within the estimated radius $\hat{r}_w$ from the center O_w) of each toponym with the extent of the administrative unit or metropolitan area to which it refers (Fig 6D). For all toponyms, the model identifies the center within the area to which each toponym refers, although the sizes of the cores vary and do not necessarily coincide with the administrative boundaries. For many toponyms, the core area detected by the model is smaller than the administrative area. This may indicate that users are more likely to make geotagged posts with these toponyms when they are in the central city of a prefecture or in the central area of a city. For *Fukuoka* and *Sendai*, the core has a geographic scale similar to the metropolitan area, suggesting that

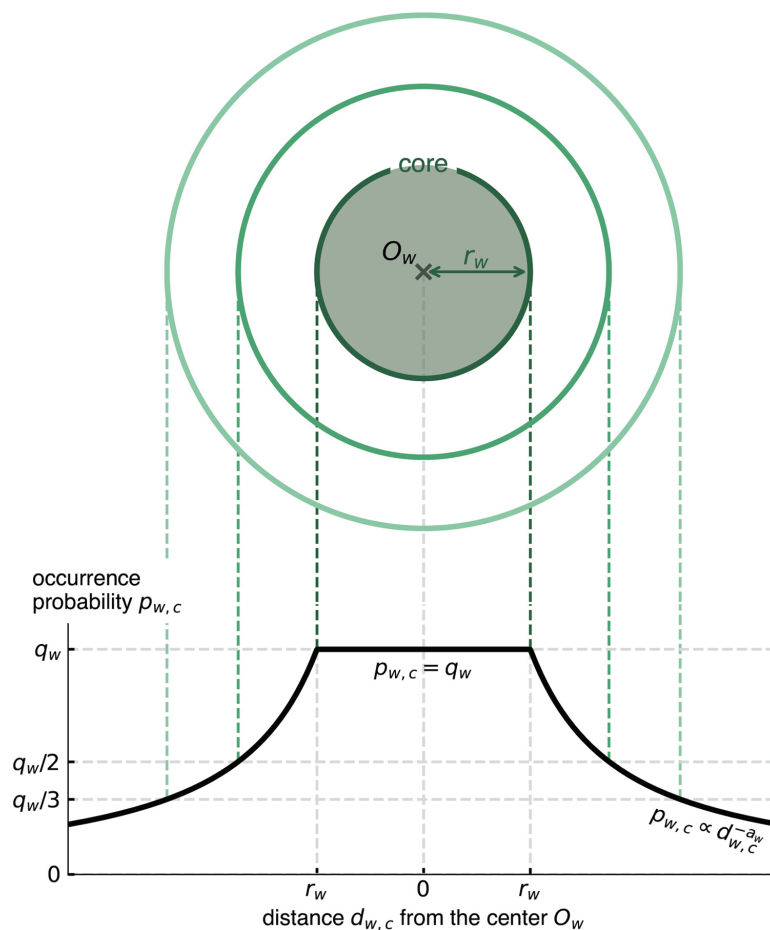


Fig 7. A schematic diagram of the core-periphery model. Top: The occurrence probability $p_{w,c}$ of toponym w at grid cell c is equal along each contour line. The model is isotropic in geographical space, meaning the equiprobability lines form concentric circles centered at O_w . The inside of inner most circle of radius r_w is the core. Bottom: The occurrence probability profile. Inside the core, $p_{w,c}$ is constant at q_w , while it decays outside the core as a power law with exponent a_w as a function of distance $d_{w,c}$ from O_w .

<https://doi.org/10.1371/journal.pone.0325022.g007>

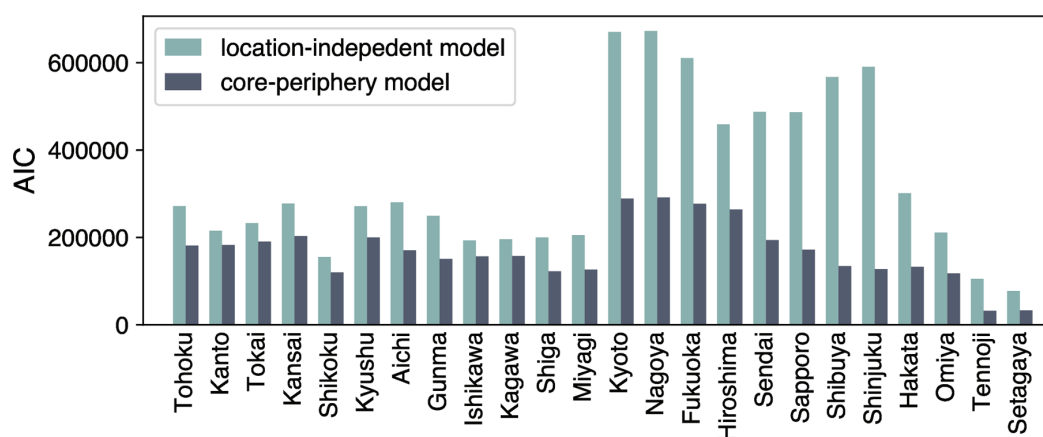


Fig 8. Goodness of fit evaluated by Akaike information criterion (AIC).

<https://doi.org/10.1371/journal.pone.0325022.g008>

the use of these two toponyms is aligned with the extent of the corresponding metropolitan area. In the case of *Sapporo*, the core is larger than the metropolitan area. This could indicate that users tend to associate the toponym with a larger area; however, it is also possible that this is because the estimated parameter represents one of the many local maxima in the likelihood landscape.

In Fig 6C, we show the estimated values of the radius and exponent for each of the 24 domestic toponyms we study. The radii of the toponyms vary according to the size of the area they denote. Region names are associated with larger cores, typically ranging from 50 km to 100 km in radius, which is consistent with the spatial scale of regions. In contrast, ward names are characterized by much smaller cores, with radii less than 5 km. Prefectures and cities fall in between, with core radii typically between 10 km and 30 km. The value of the exponent of the decay outside the core also varies from one toponym to another; however, it does not seem to correlate clearly with other quantities, such as the frequency of toponym occurrence.

Discussion and conclusions

In this article, we investigated the geographic patterns of toponym occurrence in social media using a dataset of geotagged Twitter posts. We found that the heterogeneous geographic distribution of geotagged posts is highly correlated with the population, especially with the employed population. This implies that the geotagged posts are, on the whole, representative of the language use in the population. The occurrence of each toponym in these geotagged posts is also characterized by geographic heterogeneity. Moreover, we found that the relationship between the total number of posts and the number of posts containing toponyms shows a distinctive scaling pattern. Comparison with patterns for common nouns and foreign toponyms suggests that this scaling pattern originates from the spatial specificity of toponym occurrence, which is successfully quantified by the dissimilarity from the location-independent model. Finally, we presented the core-periphery model, which assumes a location-dependent occurrence probability of toponyms. Despite its simplicity with only three fitting parameters, this model can reasonably reproduce the empirically observed geographic distributions of toponym occurrence. This implies the following: First, each toponym has a core, i.e., a geographic area in which the toponym occurs with the highest probability, which can be regarded as the area that users collectively identify with the toponym. This interpretation is supported by the fact that the core is larger for the names of regions than for the names of cities and wards. Second, outside this core, the occurrence probability decreases slowly with distance following a power law.

Our findings may indicate that human attention, cognition, and representation of geographic space respond nonlinearly to distance [27]. It is reminiscent of Tobler's first law of geography: "everything is related to everything else, but near things are more related than distant things" [28]. In this context, our findings can be seen as another example of the distance decay phenomenon, which has been observed in various aspects of human behavior such as commuting [29–31], tourism [32–34], and crime [35–37]. The concept of distance decay, and its more sophisticated formulation, the gravity model, have also been used in archaeology and history [38,39], attesting to its universal applicability in describing human activities. Inspired by developments in geography, distance decay and gravity models have been adopted in linguistics to explain variations in pronunciation between different dialects [40,41] and language evolution via lexical replacement [42]. However, unlike other language elements that are generally geography-neutral, each toponym is intrinsically associated with a specific geographic point or area. As such, our results represent a unique variant of Tobler's first law in language, one that cannot be characterized by autocorrelation or other similarity measures [43,44].

Our findings point to the need for research that addresses the micro-foundations (such as individual behavior) of the emergence of empirically observed patterns of toponym occurrence. Namely, a relevant question is: what prompts individuals to use a particular toponym (i.e., to make a post that contains a toponym)? One obvious factor is the location of the individual, or more precisely, the name of the place as perceived by the individual. Indeed, the high occurrence probability inside the core, which coincides with the geographic area to which the toponym refers, implies that the mere fact that an individual is in the place makes them more likely to use the toponym. However, this alone cannot explain the power-law decrease of the occurrence probability outside the core.

We can envisage several hypothetical mechanisms, which are not necessarily mutually exclusive, to explain the distance decay behavior. The social network is one of the plausible explanations; when users interact with each other, they may refer to each other's location in their conversations. It has been repeatedly shown that social ties are strongly influenced by geographic proximity in social networks, and that the probability of forming a tie decreases as a power-law function of distance [45–48], suggesting a possible link with toponym occurrence patterns. Place identity and place attachment may also play a major role [18,49–51]. Identity and attachment may extend to places far from the person's current residence, such as the hometown where they grew up. In this case, the use of toponyms outside the core is driven by individuals who have emigrated from the place, i.e., their long-term mobility. Another possible mechanism is related to the activity space of individuals. Activity space is defined as “a spatiotemporal construct that captures the set of places individuals encounter as a result of their routine activities in everyday life”, which “include—but are typically not limited to—individuals' residential areas” [52]. The more interaction an individual has with a place in their day-to-day activities, the more detailed information about the place they have in their cognitive map, which in turn can increase the likelihood of mentioning the toponym that refers to that place. This mechanism is driven by the short-term mobility of individuals. In this sense, mentioning the destination of a trip that an individual takes before arriving there could also be seen as a variation of this mechanism. The role each mechanism plays in shaping individual and population-wide patterns of toponym use would require text analysis of geotagged posts; we leave this for future work.

We note that our modeling approach does not aim to precisely replicate empirical observations of toponym occurrence. Rather, our models serve as a reference against which empirical data can be compared. As such, they simplify some aspects of the real-world toponym occurrence patterns. For example, in the core-periphery model, the model is isotropic, that is, the core has a circular boundary and the occurrence probability decreases uniformly in all directions. In reality, however, the area denoted by the toponym may have elongated or irregular shapes and the decay outside the core may not be isotropic because of geographic entities (such as the sea, mountains, and rivers), administrative boundaries, and transportation networks (such as roads and railways). We also assumed no correlations between the occurrence of different toponyms, although there may be competitive or synergistic interactions between them that influence their occurrence patterns. It is because of these simplifications that our approach is able to provide an interpretable framework that allows us to identify the essential elements underlying the geospatial distributions of toponyms. Nevertheless, extending the model to capture the impact of these additional ingredients remains an open challenge that could be addressed in future work. Prior research has explored how natural landscapes shape the evolution of population distributions and transport networks [53], suggesting that similar approaches could be applied to the study of toponym distributions.

We also remark on the possibility that geotagged posts on Twitter may not be an unbiased, representative sample of the language use of the general population. This issue can be divided

into two questions: whether the geotagged Twitter users can be considered a good proxy for the population at large [13,54–56], and whether the language use in geotagged Twitter posts is consistent with language use in other contexts [57,58]. The effect of the first problem on our results is presumably relatively small compared to other work using geotagged Twitter posts to study language use, for two reasons: (i) we focus only on the text of the posts without correlating it with user demographics, and (ii) the age and socioeconomic status of users are unlikely to significantly affect their use of toponyms. It is however possible that urban toponyms are overrepresented due to population biases, which could affect the geographic distribution of toponyms. In this work, we focused on relatively large cities and wards within them, but whether our findings generalize to the names of smaller towns and villages needs to be carefully examined in future research.

The second question concerns the generalizability of our findings to the use of toponyms in other contexts of language use. While capturing the totality of language use is challenging, this question can be partly addressed, for example, by comparing language use on different social media platforms. Our analysis focuses solely on Twitter data, which presents a limitation of our study. To the best of our knowledge, only a few previous works have addressed if there are significant differences in the language used in different contexts or platforms. One of them documents that descriptions of natural landscapes can vary between different data sources, finding that Flickr (a photo-sharing social media platform) tends to contain a higher proportion of toponyms than free lists obtained through field interviews and hiking blogs [58]. Another study suggests that the same user may exhibit different styles on different social media platforms for certain linguistic features, such as emoji and hashtags [59]. However, neither study addresses whether individual toponyms are used differently across contexts and platforms. We speculate that, while contextual differences may affect the overall volume of toponyms, they do not strongly influence their geographic patterns of occurrence. However, this is a hypothesis that needs to be tested in future research.

Lastly, we note that gaining a comprehensive understanding of the use of toponyms—how they reflect the interaction between people and the environment, how they shape and reinforce people's identity, and how they are affected by urban planning and place branding—would require historical, ecological, cultural, and economic analyses [49,60–66]. These aspects are abstracted away in the present study, where our focus is to establish general empirical laws that govern the spatial distribution of toponyms. The simplified models we propose are designed to serve the purpose of quantitative, so-called *extensive* analysis [67]. Future work should aim to integrate these qualitative factors with our quantitative framework to investigate the dynamics underlying toponym distribution.

Supporting information

S1 Appendix. Definition of relative entropy, data preprocessing.
(PDF)

S1 Fig. Probability density functions of occurrence density σ_w .
(PDF)

S2 Fig. Scatter plot of n_w versus n_{all} for all toponyms and nouns studied in this work.
(PDF)

S3 Fig. Comparison between the empirical data and the location-independent model.
Each set of contours represents the kernel density plot of n_w versus n_{all} of empirical data (blue) and the location-independent model (red).
(PDF)

S4 Fig. Comparison between the empirical data and the core-periphery model. Each set of contours represents the kernel density plot of n_w versus n_{all} of empirical data (blue) and the core-periphery model (red).
(PDF)

Acknowledgments

T.H. acknowledges the computational resources provided by the Aalto Science-IT project.

Author contributions

Conceptualization: Takayuki Hiraoka, Takashi Kirimura, Naoya Fujiwara.

Data curation: Takayuki Hiraoka, Takashi Kirimura.

Formal analysis: Takayuki Hiraoka.

Funding acquisition: Naoya Fujiwara.

Investigation: Takayuki Hiraoka, Takashi Kirimura, Naoya Fujiwara.

Methodology: Takayuki Hiraoka, Takashi Kirimura, Naoya Fujiwara.

Resources: Takashi Kirimura.

Software: Takayuki Hiraoka.

Supervision: Naoya Fujiwara.

Visualization: Takayuki Hiraoka.

Writing – original draft: Takayuki Hiraoka.

Writing – review & editing: Takayuki Hiraoka, Takashi Kirimura, Naoya Fujiwara.

References

1. Montello DR, Goodchild MF, Gottsegen J, Fohl P. Where's downtown?: Behavioral methods for determining referents of vague spatial queries. *Spatial Cognit Comput*. 2003;3(2–3):185–204. <https://doi.org/10.1080/13875868.2003.9683761>
2. Jones CB, Purves RS, Clough PD, Joho H. Modelling vague places with knowledge from the Web. *Int J Geograph Inf Sci*. 2008;22(10):1045–65. <https://doi.org/10.1080/13658810701850547>
3. DeLozier G, Baldridge J, London L. Gazetteer-independent toponym resolution using geographic word profiles. *AAAI*. 2015;29(1):2382–8. <https://doi.org/10.1609/aaai.v29i1.9531>
4. Cheng Z, Caverlee J, Lee K. You are where you tweet. In: *Proceedings of the 19th ACM International Conference on Information and Knowledge Management*. ACM. 2010. <https://doi.org/10.1145/1871437.1871535>
5. Li W, Serdyukov P, de Vries AP, Eickhoff C, Larson M. The where in the tweet. In: *Proceedings of the 20th ACM International Conference on Information and Knowledge Management*. ACM. 2011. <https://doi.org/10.1145/2063576.2063995>
6. Hahmann S, Purves R, Burghardt D. Twitter location (sometimes) matters: exploring the relationship between georeferenced tweet content and nearby feature classes. *JOSIS*. 2014;(9):1–36. <https://doi.org/10.5311/josis.2014.9.185>
7. Eisenstein J, O'Connor B, Smith N, Xing E. A latent variable model for geographic lexical variation. In: *Proceedings of the 2010 Conference on Empirical Methods in Natural Language Processing*. 2010. p. 1277–87.
8. Huang Y, Guo D, Kasakoff A, Grieve J. Understanding U.S. regional linguistic variation with Twitter data analysis. *Comput Environ Urban Syst*. 2016;59:244–55. <https://doi.org/10.1016/j.compenvurbsys.2015.12.003>

9. Abitbol JL, Karsai M, Magué J-P, Chevrot J-P, Fleury E. Socioeconomic dependencies of linguistic patterns in Twitter. In: Proceedings of the 2018 World Wide Web Conference on World Wide Web - WWW 2018. ACM Press. 2018. p. 1125–34. <https://doi.org/10.1145/3178876.3186011>
10. Louf T, Gonçalves B, Ramasco JJ, Sánchez D, Grieve J. American cultural regions mapped through the lexical analysis of social media. *Humanit Soc Sci Commun*. 2023;10(1):133. <https://doi.org/10.1057/s41599-023-01611-3>
11. Louf T, Ramasco J, Sánchez D, Karsai M. When dialects collide: how socioeconomic mixing affects language use. *arXiv preprint* 2023. <https://arxiv.org/abs/2307.10016>
12. Morin C, Grieve J. The semantics, sociolinguistics, and origins of double modals in American English: New insights from social media. *PLoS One*. 2024;19(1):e0295799. <https://doi.org/10.1371/journal.pone.0295799> PMID: 38265988
13. Pavalanathan U, Eisenstein J. Confounds and consequences in geotagged Twitter Data. In: Proceedings of the 2015 Conference on Empirical Methods in Natural Language Processing. Association for Computational Linguistics. 2015. <https://doi.org/10.18653/v1/d15-1256>
14. Johnson IL, Sengupta S, Schöning J, Hecht B. The geography and importance of localness in geotagged social media. In: Proceedings of the 2016 CHI Conference on Human Factors in Computing Systems. ACM. 2016. p. 515–26. <https://doi.org/10.1145/2858036.2858122>
15. Oliveira MG de, Campelo CEC, Baptista C de S, Bertolotto M. Gazetteer enrichment for addressing urban areas: a case study. *J Locat Based Serv*. 2016;10(2):142–59. <https://doi.org/10.1080/17489725.2016.1196755>
16. Xu C, Wong DW, Yang C. Evaluating the “geographical awareness” of individuals: an exploratory analysis of twitter data. *Cartograph Geograph Inf Sci*. 2013;40(2):103–15. <https://doi.org/10.1080/15230406.2013.776212>
17. Han SY, Tsou M-H, Clarke KC. Do global cities enable global views? Using Twitter to quantify the level of geographical awareness of U.S. Cities. *PLoS One*. 2015;10(7):e0132464. <https://doi.org/10.1371/journal.pone.0132464> PMID: 26167942
18. Arthur R, Williams HTP. The human geography of Twitter: quantifying regional identity and inter-region communication in England and Wales. *PLoS One*. 2019;14(4):e0214466. <https://doi.org/10.1371/journal.pone.0214466> PMID: 30986213
19. Purves R, Hollenstein L. Exploring place through user-generated content: using Flickr to describe city cores. *JOSIS*. 2010; 1:21–48. <https://doi.org/10.5311/josis.2010.1.3>
20. Hu Y, Janowicz K. An empirical study on the names of points of interest and their changes with geographic distance. In: *LIPICs*. 2018. p. 5:1-5:15. <https://doi.org/10.4230/LIPICs.GISCIENCE.2018.5>
21. Guo D, Chen C. Detecting non-personal and spam users on geo-tagged Twitter network. *Trans GIS*. 2014;18(3):370–84. <https://doi.org/10.1111/tgis.12101>
22. Zhao B, Sui DZ. True lies in geospatial big data: detecting location spoofing in social media. *Annal GIS*. 2017;23(1):1–14. <https://doi.org/10.1080/19475683.2017.1280536>
23. Kanemoto Y, Tokuoka K. Proposal for the standards of metropolitan areas of Japan. *J Appl Reg Sci*. 2002;7:1–15.
24. 2015 Metropolitan Employment Area. [cited 2023 Jan 25]. https://www.csis.u-tokyo.ac.jp/UEA/index_e.htm
25. Statistics Bureau, Ministry of Internal Affairs and Communications, Japan. 2015 Population Census; 2017 [cited 2023 Aug 02; Japanese]. <https://www.e-stat.go.jp/gis/statmap-search?page=1&type=1&toukeiCode=00200521&toukeiYear=2015&aggregateUnit=S&surveyId=S002005112015&statsId=T000846>
26. Statistics Bureau M of IA and CJ. 2016 economic census for business activity. 2019. <https://www.e-stat.go.jp/gis/statmap-search?page=1&type=1&toukeiCode=00200553&toukeiYear=2016&aggregateUnit=S&surveyId=S002005112016&statsId=T000917>
27. Montello DR. Cognitive geography. *International encyclopedia of human geography*. Elsevier. 2009. p. 160–6. <https://doi.org/10.1016/b978-008044910-4.00668-4>
28. Tobler WR. A computer movie simulating urban growth in the detroit region. *Econ Geography*. 1970;46:234. <https://doi.org/10.2307/143141>
29. Iacono M, Krizek K, El-Geneidy A. Access to destinations: how close is close enough? estimating accurate distance decay functions for multiple modes and different purposes. 2008–11. Minnesota Department of Transportation. 2008. <https://hdl.handle.net/11299/151329>
30. Helminen V, Rita H, Ristimäki M, Kontio P. Commuting to the centre in different urban structures. *Environ Plann B Plann Des*. 2012;39(2):247–61. <https://doi.org/10.1068/b36004>
31. Halás M, Klapka P, Kladivo P. Distance-decay functions for daily travel-to-work flows. *J Transp Geography*. 2014;35:107–19. <https://doi.org/10.1016/j.jtrangeo.2014.02.001>

32. McKercher B, Chan A, Lam C. The impact of distance on international tourist movements. *J Travel Res.* 2008;47(2):208–24. <https://doi.org/10.1177/0047287508321191>
33. Hooper J. A destination too far? Modelling destination accessibility and distance decay in tourism. *GeoJournal.* 2014;80(1):33–46. <https://doi.org/10.1007/s10708-014-9536-z>
34. McKercher B. The impact of distance on tourism: a tourism geography law. *Tourism Geograph.* 2018;20(5):905–9. <https://doi.org/10.1080/14616688.2018.1434813>
35. Rengert GF, Piquero AR, Jones PR. Distance decay reexamined. *Criminology.* 1999;37(2):427–46. <https://doi.org/10.1111/j.1745-9125.1999.tb00492.x>
36. Kent J, Leitner M, Curtis A. Evaluating the usefulness of functional distance measures when calibrating journey-to-crime distance decay functions. *Comput Environ Urban Syst.* 2006;30(2):181–200. <https://doi.org/10.1016/j.compenvurbsys.2004.10.002>
37. Townsley M, Sidebottom A. All offenders are equal, but some are more equal than others: variation in journeys to crime between offenders*. *Criminology.* 2010;48(3):897–917. <https://doi.org/10.1111/j.1745-9125.2010.00205.x>
38. Tobler W, Wineburg S. A cappadocian speculation. *Nature.* 1971;231(5297):39–41. <https://doi.org/10.1038/231039a0> PMID: 16062545
39. Renfrew C. Alternative models for exchange and spatial distribution. *Exchange systems in prehistory.* Elsevier. 1977. p. 71–90. <https://doi.org/10.1016/b978-0-12-227650-7.50010-9>
40. Trudgill P. Linguistic change and diffusion: description and explanation in sociolinguistic dialect geography. *Lang Soc.* 1974;3(2):215–46. <https://doi.org/10.1017/s0047404500004358>
41. Nerbonne J, Heeringa W. Geographic distributions of linguistic variation reflect dynamics of differentiation. *Roots.* Mouton de Gruyter. 2007. p. 267–98. <https://doi.org/10.1515/9783110198621.267>
42. Cavalli-Sforza LL, Wang WS-Y. Spatial distance and lexical replacement. *Lanaguage.* 1986;62(1):38–55. <https://doi.org/10.1353/lan.1986.0115>
43. Miller HJ. Tobler's first law and spatial analysis. *Annals Assoc Am Geograph.* 2004;94(2):284–9. <https://doi.org/10.1111/j.1467-8306.2004.09402005.x>
44. Waters N. Tobler's first law of geography. *International encyclopedia of geography.* Wiley. 2018. p. 1–15. <https://doi.org/10.1002/9781118786352.wbieg1011.pub2>
45. Lambiotte R, Blondel VD, de Kerchove C, Huens E, Prieur C, Smoreda Z, et al. Geographical dispersal of mobile communication networks. *Phys A: Statist Mech Appl.* 2008;387(21):5317–25. <https://doi.org/10.1016/j.physa.2008.05.014>
46. Onnela J-P, Arbesman S, González MC, Barabási A-L, Christakis NA. Geographic constraints on social network groups. *PLoS One.* 2011;6(4):e16939. <https://doi.org/10.1371/journal.pone.0016939> PMID: 21483665
47. McGee J, Caverlee JA, Cheng Z. A geographic study of tie strength in social media. In: *Proceedings of the 20th ACM International Conference on Information and Knowledge Management.* ACM. 2011. 2333–6. <https://doi.org/10.1145/2063576.2063959>
48. Lengyel B, Varga A, Ságvári B, Jakobi Á, Kertész J. Geographies of an online social network. *PLoS One.* 2015;10(9):e0137248. <https://doi.org/10.1371/journal.pone.0137248> PMID: 26359668
49. Hakala U, Sjöblom P, Kantola S-P. Toponyms as carriers of heritage: implications for place branding. *J Prod Brand Manag.* 2015;24(3):263–75. <https://doi.org/10.1108/jpbm-05-2014-0612>
50. Cardoso RV, Meijers EJ. The metropolitan name game: the pathways to place naming shaping metropolitan regions. *Environ Plan A.* 2016;49(3):703–21. <https://doi.org/10.1177/0308518x16678851>
51. Di Masso A, Williams DR, Raymond CM, Buchecker M, Degenhardt B, Devine-Wright P, et al. Between fixities and flows: navigating place attachments in an increasingly mobile world. *J Environ Psychol.* 2019;61:125–33. <https://doi.org/10.1016/j.jenvp.2019.01.006>
52. Cagney KA, York Cornwell E, Goldman AW, Cai L. Urban mobility and activity space. *Annu Rev Sociol.* 2020;46(1):623–48. <https://doi.org/10.1146/annurev-soc-121919-054848>
53. Aoki T, Fujiwara N, Fricker M, Nakagaki T. A model for simulating emergent patterns of cities and roads on real-world landscapes. *Sci Rep.* 2022;12(1):10093. <https://doi.org/10.1038/s41598-022-13758-1> PMID: 35710781
54. Hecht B, Stephens M. A tale of cities: urban biases in volunteered geographic information. *ICWSM.* 2014;8(1):197–205. <https://doi.org/10.1609/icwsm.v8i1.14554>
55. Malik M, Lamba H, Nakos C, Pfeffer J. Population bias in geotagged tweets. *ICWSM.* 2021;9(4):18–27. <https://doi.org/10.1609/icwsm.v9i4.14688>
56. Anselin L, Williams S. Digital neighborhoods. *J Urbanism: Int Res Placemak Urban Sustain.* 2015;9(4):305–28. <https://doi.org/10.1080/17549175.2015.1080752>

57. Li L, Goodchild MF, Xu B. Spatial, temporal, and socioeconomic patterns in the use of Twitter and Flickr. *Cartograph Geograph Inf Sci*. 2013;40(2):61–77. <https://doi.org/10.1080/15230406.2013.777139>
58. Wartmann FM, Acheson E, Purves RS. Describing and comparing landscapes using tags, texts, and free lists: an interdisciplinary approach. *Int J Geograph Inf Sci*. 2018;32(8):1572–92. <https://doi.org/10.1080/13658816.2018.1445257>
59. Marko K, Reitbauer M, Pickl G. Same person, different platform. *RS*. 2022;4(2):202–31. <https://doi.org/10.1075/rs.22006.mar>
60. Conedera M, Vassere S, Neff C, Meurer M, Krebs P. Using toponymy to reconstruct past land use: a case study of 'brüsáda' (burn) in southern Switzerland. *J Historic Geography*. 2007;33(4):729–48. <https://doi.org/10.1016/j.jhg.2006.11.002>
61. Radding L, Western J. What's in a name? Linguistics, geography, and toponyms*. *Geograph Rev*. 2010;100(3):394–412. <https://doi.org/10.1111/j.1931-0846.2010.00043.x>
62. Rose-Redwood R, Alderman D, Azaryahu M. Geographies of toponymic inscription: new directions in critical place-name studies. *Prog Hum Geography*. 2009;34(4):453–70. <https://doi.org/10.1177/0309132509351042>
63. Light D, Young C. Toponymy as commodity: exploring the economic dimensions of urban place names. *Int J Urban Regional Res*. 2014;39(3):435–50. <https://doi.org/10.1111/1468-2427.12153>
64. Capra GF, Ganga A, Filzmoser P, Gaviano C, Vacca S. Combining place names and scientific knowledge on soil resources through an integrated ethnopedological approach. *CATENA*. 2016;142:89–101. <https://doi.org/10.1016/j.catena.2016.03.003>
65. Atik M, Swaffield S. Place names and landscape character: a case study from Otago Region, New Zealand. *Landsc Res*. 2017;42(5):455–70. <https://doi.org/10.1080/01426397.2017.1283395>
66. Rose-Redwood R, Vuolteenaho J, Young C, Light D. Naming rights, place branding, and the tumultuous cultural landscapes of neoliberal urbanism. *Urban Geography*. 2019;40(6):747–61. <https://doi.org/10.1080/02723638.2019.1621125>
67. Tent J. Approaches to research in toponymy. *Names*. 2015;63(2):65–74. <https://doi.org/10.1179/0027773814z.000000000103>