

KIER DISCUSSION PAPER SERIES

KYOTO INSTITUTE OF ECONOMIC RESEARCH

Discussion Paper No. 904

On the Spatial Scale of Industrial Agglomerations

Tomoya Mori and Tony E. Smith

October 2014



KYOTO UNIVERSITY
KYOTO, JAPAN

On the Spatial Scale of Industrial Agglomerations

Tomoya Mori and Tony E. Smith^{*,†}

October 2014

Abstract

The standard approaches to studying industrial agglomeration have been in terms of summary measures of the “degree of agglomeration” within each industry. But such measures often fail to distinguish between industries that exhibit substantially different spatial scales of agglomeration. In a previous paper, Mori and Smith [45] proposed a new pair of quantitative measures for distinguishing both the scale and degree of industrial agglomeration based on an explicit method for detecting spatial clusters. The first, designated as the *global extent (GE)* of industrial clusters, measures the spatial spread of these clusters (within a given country) in terms of the areal size of their *essential containment*, defined to be the (convex-solid) region containing the most significant subset of these clusters. The second, designated as the *local density (LD)* of industrial clusters, measures the spatial extent of individual clusters *within their essential containment* in terms of the areal share of that containment occupied by clusters. The central purpose of the present paper is to apply these two measures to the manufacturing industries in Japan, and to demonstrate how they can be used in combination to distinguish both the relative scale and degree of agglomeration exhibited by cluster patterns for each industry. In addition, the information provided by this pair of measures (*GE, LD*) is systematically compared to that of the most prominent summary measures currently in use. Finally, it is shown that these measures also support certain predictions of new economic geography models in the sense that shipping distances for establishments in each industry tend to be negatively (positively) correlated with the *GE (LD)* measures of agglomeration in these industries.

JEL Classifications : C49, L60, R12, R14

Keywords : Industrial Agglomeration, Cluster analysis, Spatial patterns of agglomeration, Shipment distances, New economic geography

^{*}Mori: Institute of Economic Research, Kyoto University and Research Institute of Economy, Trade and Industry (RIETI) of Japan. Email: mori@kier.kyoto-u.ac.jp. Smith: Department of Electrical and Systems Engineering, University of Pennsylvania. Email: te-smith@seas.upenn.edu.

[†]This research is conducted as part of the project, *the formation of economic regions and its mechanism: theory and evidence*, undertaken at the Research Institute of Economy, Trade and Industry, and has been partially supported by the Grant in Aid for Research (Nos. 25285074, 25380294, 2624503) of MEXT of Japan.

1 Introduction

The standard approach to studying industrial agglomeration has focused on the overall degree of agglomeration in industries, and typically measures the discrepancy between industry-specific regional distributions of establishments (or employment) and a given hypothetical reference distribution representing “complete dispersion” in terms of some scalar index (e.g., Ellison and Glaeser [13], Duranton and Overman [10], Brühlhart and Traeger [6], Mori et al. [42]).¹ But even if industries are judged to be similar with respect to these indices, their spatial patterns of agglomeration may appear to be quite different. In particular, these aggregate measures often fail to distinguish between industries that exhibit substantially different spatial scales of agglomeration.

In a previous paper, Mori and Smith [45] proposed a new pair of quantitative measures for distinguishing both the scale and degree of industrial agglomeration based on an explicit method for detecting spatial clusters. The first, designated as the *global extent (GE)* of an industry’s clusters within a given country, measures the spatial spread of these clusters in terms of the areal share of their *essential containment* within that country, namely, the smallest “convex-solid” region containing all “significant” clusters (to be formally defined in Section 3.1). Smaller values of *GE* for industries imply that their major clusters are essentially confined to smaller regions of the country, while larger values indicate that these clusters are more dispersed. In contrast to this global measure of spread, the second measure, designated as *local density (LD)*, focuses solely on clusters *within the essential containment* for that industry, and measures their local density in terms of areal share within this containment. Larger (smaller) values of *LD* for an industry thus imply that its clusters tend to be more (less) spread out within this critical region.

These specific measures are largely inspired by theoretical results from the “new economic geography” (NEG)² where industrial location is modeled in continuous space.³ Here it has been shown that the spatial structure of agglomer-

¹Examples of such reference distributions are (1) the regional distribution of all-industry employment, used by Ellison and Glaeser [13], (2) the regional distribution all-industry establishments, used by Duranton and Overman [10], and (3) the regional distribution of economic area used by Mori et al. [42]. Brühlhart and Traeger [6] adopted (1) and (3).

²See, e.g., Fujita et al. [17] and Combes et al. [9] for an overview of the literature.

³See Fujita and Mori [19, 20] for a survey.

ation and dispersion can change at different scales of analysis, depending on a host of factors including plant-level increasing returns, product differentiation and transport costs. These effects are well illustrated by considering the spatial effects of transport costs in simple “core-periphery” models of industrial location (e.g., Tabuchi [53]; Murata and Thisse [46]). At very high levels of transport costs, the dispersion of consumers between the “core” and “periphery” regions leads to a corresponding dispersion of manufacturing. But as transport costs decrease and distance to consumers becomes less critical, manufacturing tends to concentrate (in the core region). Finally, at even lower levels of transport costs, the commuting costs dominate (together with congestion effects) in the core region and can induce a second phase of manufacturing dispersion (popularly referred to as “re-dispersion” or “revival”). Alternatively, the responses of the manufacturing industry to the different levels of transport costs may be interpreted as the responses of different (manufacturing) industries to a given transport cost level. In particular, the industries which are very sensitive (resp., insensitive) to transport costs tend to spatially disperse when transport costs are very high (resp., low), while those with intermediate sensitivity to transport costs tend to spatially agglomerate.

Indeed, these two dispersion patterns often appear to be exactly the same (i.e., a symmetric distribution of manufacturing between the two regions). But, in NEG models involving more general location spaces (e.g., Krugman [33]; Fujita and Mori [18]), these two dispersion phases have qualitatively different spatial patterns. In particular, while the dispersion of manufacturing at high levels of transport costs continues to be global (as in core-periphery models), the second phase of dispersion at low levels of transport costs is much more localized in nature, and perhaps better described as an expansion of existing core areas rather than re-dispersion to peripheral areas (see the formation of an “industrial belt” – a continuum of cities – in Mori [41] as an explicit example of this dispersion in continuous location spaces).⁴ Accordingly, the implications of these two types of dispersions are quite different.

Such theoretical findings raise important questions as to whether this diversity of patterns can in fact be identified empirically. Hence the specific measures proposed here are designed to quantify pattern differences both in terms

⁴See also Behrens [5] for a related discussion on the spatial extent of agglomeration in NEG models.

of their global and local properties. While the details of these measures require a more formal definition and construction of agglomeration patterns, the basic ideas can be illustrated by a preview of the types of patterns we have identified for Japanese manufacturing industries in 2001.

First, there are industries which clearly exhibit strong spatial concentration, such as the “compounding plastic materials, including reclaimed plastics” industry. The agglomeration pattern derived for this industry is shown in Figure 12(b), where the areas marked by the enclosed red regions denote industrial clusters.⁵ Notice that the main industrial concentration lies clearly in the Industrial Belt along the Pacific coast extending westward from Tokyo to Fukuoka. Moreover, the individual clusters of establishments within this belt are seen to be densely packed from end to end. We describe this type of agglomeration pattern as “globally confined” and “locally dense” (here with respect to the Industrial Belt). In particular, this pattern is reminiscent of the type of “second-phase” dispersion of manufacturing identified in the NEG models described above. But even globally dispersed industries often form small clusters at local scales. For example, the agglomeration pattern for the “manufactured ice” industry shown in Figure 9(b) is spread throughout the country, but exhibits a large number of local clusters. Such patterns, which we describe as “globally dispersed” and “locally sparse”, are closer in spirit to the “first-phase” dispersion of manufacturing in the NEG models above.

With respect to the aggregate measures of agglomeration above, it is not clear which of these two industries would be judged as “more agglomerated”, since the first industry exhibits agglomeration at the global scale but dispersion at the local scale, while the opposite is true for the second industry. In fact, this may not even be an appropriate comparison. Aside from these extremes, there are a variety of other patterns that can be identified, as discussed more fully in Sections 3 and 4 below. As will be shown in Section 5, industries generally exhibit wide variations in GE and LD . With respect to scalar measures of agglomeration, it will also be shown and that those of Ellison and Glaeser [13] and Mori et al. [42] can both be roughly equally represented by these two components, while that of Duranton and Overman [10] is primarily associated with GE .

Finally, by using micro data for the shipments of individual establishments in

⁵See Section 4.3 below for a more detailed discussion of these figures.

2000, we investigate the relationship between the shipment patterns and spatial scales of agglomeration. Here, we focus on the shipment distances for individual establishments. Since the establishments of industries which are sensitive to transport costs are expected to locate closer to their market, they are also expected to ship their products over shorter distances. We show that the shipment distances of individual industries are negatively correlated with GE and are positively correlated with LD of clusters, which provides some of the first direct evidence for the theoretical predictions of NEG models mentioned above. For the “globally dispersed and locally sparse” pattern (i.e., large GE and small LD) corresponding to the first phase of dispersion mentioned above, the data indicates small shipment distances, while for the “globally confined and locally dense” pattern (i.e., small GE and large LD) corresponding to the second phase of dispersion, the data indicates larger shipment distances, agreeing with these theoretical predictions.

To develop these ideas, the paper is organized as follows. In Section 2 below we develop the formal framework for analysis, and briefly sketch the cluster identification procedure developed in Mori and Smith [45]. This is followed in Section 3 with a development of our summary measures for analyzing and classifying the agglomeration patterns obtained. These methods are then applied in Section 4 to (i) identify establishment clusters for each manufacturing industry in Japan, and to (ii) identify the spatial scales of these agglomeration patterns. In Section 5, the relationship between existing scalar indices of agglomeration and our pair of measures, GE and LD , are discussed. Finally, the relationship between shipment distances and spatial scales of agglomeration is investigated in Section 6. The paper concludes with brief discussions of related research in Section 7.

2 Identification of Industrial Clusters

This section provides an overview of the cluster detection framework developed by Mori and Smith [45].⁶ We begin with a set, R , of *basic regions* (municipalities),

⁶All the relevant C++ and Python programs for the cluster detection introduced in this section can be downloaded from the web: http://www.mori.kier.kyoto-u.ac.jp/data/cluster_detection.html. Also, all the input and output data as well as map data for the application to Japanese manufacturing industries in Section 4 can be downloaded from the same site.

r , within which each industry can locate. An *industrial cluster* is then taken roughly to be a spatially coherent subset of regions within which the density of industrial establishments is unusually high. Since the explicit construction of such clusters will have consequences for the summary measures to be developed, it is appropriate to outline this construction more explicitly. The present notion of “spatial coherence” is taken to include the requirement that such regions be contiguous, and as close to one another as possible – where “closeness” is defined with respect to the relevant underlying regional network, where the nodes of this network are represented by the set R of basic regions, and the links are taken to represent pairs of regional “neighbors” in terms of the underlying regional network. By using travel distances between regional centers along this network, we define *shortest paths* between each pair of regions, r_i and r_j , to be sequences of intermediate regions, $(r_i, r_1, \dots, r_k, r_j)$ reflecting minimum travel distances with respect to the road network. Our key requirement for spatial coherence of a cluster is that it be *convex-solid* in the sense that it includes all shortest paths between its member regions (convexity), and allows no holes (solidity).⁷

2.1 Interregional Distance

Since the underlying interregional network will have direct impact on the industrial clusters to be identified, it is worth discussing our choice of the network structure. In our practical applications in Section 4 and thereafter, we adopt the actual road network, and hence the interregional distances are measured in terms of the travel distance along the road network. Note that it is possible to use more stylized interregional distances, based for example on Great-Circle distances (as is common in the literature⁸). But, the advantage of choosing road-network data is primarily to take into account the underlying topographical heterogeneity, which could hardly be reflected in the Great-Circle distances. In the case of Japan to be studied below, the Pearson’s correlation between the Great-Circle distances and road-network distances for all pairs of basic regions (municipalities) is as high as 0.976.⁹ But, as suggested by Duranton and Overman [10,

⁷The requirement of solidity is not essential. But, it provides a more cohesive view of clusters as areas of industrial agglomeration.

⁸In particular, Duranton and Overman [10] and all of their followers.

⁹The magnitude of the correlation is comparable to 0.97 for the case of the United Kingdom reported in Duranton and Overman [10, footnote 4].

§4] the size of most industrial clusters are within the 40km range, and almost all are within 100km range. So this broad correlation over all scales is not sufficiently informative to gauge the relevance of the network distance. Namely, as is clear from the frequency distributions of interregional distances shown in Figure 1, the majority in the entire set of municipality pairs are simply too distant from one another to constitute to meaningful clusters. More specifically, the municipality pairs within 100km range account for less than 10% of all the municipality pairs in both Great-Circle and road-network distances. In fact, the correlations reduce to 0.711, 0.633, 0.485 and 0.480, for the municipality pairs within 100km, 50km, 20km and 10km ranges (in terms of the Great-Circle distance), respectively. The corresponding correlations reduce even to 0.539, 0.461, 0.338 and 0.249, respectively, for the municipality pairs along the sea coast, where most of the major clusters are to be identified.

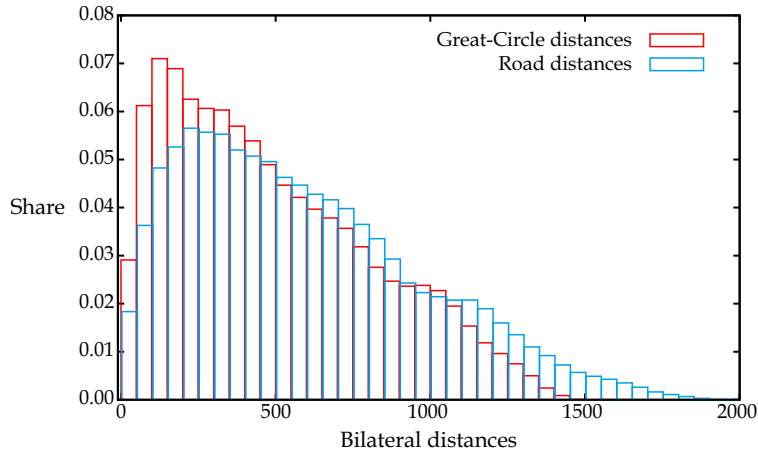


Figure 1: Distribution of interregional distances

These results are not specific to Japan. In the case 4626 unions of the continental Germany (which has comparable areal size with Japan) in 2008, while the Pearson's correlation for all the 10,697,625 union pairs is 0.911, it reduces to 0.803, 0.777, 0.630 and 0.382 for all the union pairs within 100km, 50km, 20km and 10km ranges, respectively.¹⁰ In the case of 3106 counties in the continental US in 2007, while the same correlation for all the 4,822,065 county pairs is as high as 0.928, it reduces to as low as 0.216, 0.136, 0.073 and 0.019

¹⁰The distances between these unions are computed in terms of those between union offices. We thank Wolfgang Dauth and Jens Südekum for sharing Germany data with us.

for all the county pairs within 100km, 50km, 20km and 10km ranges, respectively.¹¹ Although the values of correlations are not directly comparable between Japan/German and the US, since the sizes of these regions greatly differ,¹² the discrepancy in the US case appears to be far more serious than the cases of Japan and Germany.

These evidences suggest that it is important to adopt realistic distance data to obtain reliable results on agglomeration patterns. In addition, it is a simple matter to compute bilateral distances along a given network in R by using a GIS software. In ArcGIS (ver.10.2) of ESRI, for instance, all the interregional distances can be automatically computed by utilizing the “network analyst” extension. Thus, today, there is no strong reason to choose simplistic distance data.¹³

2.2 Cluster Schemes

Most industries consist of multiple clusters in R that together define the agglomeration pattern for that industry. In fact, the spacing between such clusters is a topic of considerable economic interest (as discussed further in Section 7.1 below). Hence it is essential to model such patterns as explicit spatial arrangements of multiple clusters. The model proposed in Mori and Smith [45] is a *cluster scheme*, $\mathbf{C} = (R_0, C_1, \dots, C_{k_C})$, that partitions R into one or more disjoint clusters (convex solids), $C_1, \dots, C_{k_C}$, together with the residual set, R_0 , of all non-cluster regions in R . The individual clusters are implicitly taken to be areas in R where industry density is unusually high. But within each cluster, C_j , all that is assumed for modeling purposes is that location probabilities for randomly sampled industrial establishments are uniform across the feasible locations in C_j . More precisely, if the *feasible area* as defined in Section 4.1.2 below for locations in each region, $r \in R$, is denoted by a_r , so that the total area of C_j is given by $a_{C_j} = \sum_{r \in C_j} a_r$, then location probabilities in C_j are taken to be uniform over a_{C_j} . In particular, this implies that the conditional probability of an

¹¹The distances between the US counties are computed in terms of those between the county courthouses, or some other public facilities if there are no courthouses.

¹²The US counties on average more than twenty times larger in areal size than Japanese municipalities.

¹³See Combes and Lafourcade [8] for more sophisticated definition of interregional distances which takes into account more general costs for travel such as time and fuel costs. See also Özak [47] for the derivation of the least-cost routes based various topographical and climatic characteristics.

establishment locating in $r \in C_j$ given that it is located in C_j is simply a_r/a_{C_j} . With this assumption, the only unknown probabilities are the marginal location probabilities, $p_{\mathbf{C}}(j)$, for clusters C_j in $\mathbf{C}$. Hence each cluster scheme, $\mathbf{C}$, generates a possible *cluster probability model*, $p_{\mathbf{C}} = [p_{\mathbf{C}}(j) : j = 1, \dots, k_{\mathbf{C}}]$, of establishment locations for the industry.¹⁴ If there are n establishments in the given industry, then each cluster probability model, $p_{\mathbf{C}}$, amounts formally to multinomial sampling model with sample size, n , and outcomes given by the $k_{\mathbf{C}} + 1$ sets in cluster scheme, $\mathbf{C}$, with respect to samples of size n . Finally, since the observed relative frequencies, $\hat{p}_{\mathbf{C}} = [\hat{p}_{\mathbf{C}}(j) = n_j/n : j = 1, \dots, k_{\mathbf{C}}]$, of establishments in each cluster are well known to be the maximum-likelihood estimates of these (multinomial) probabilities, such estimates yield a family of well-defined candidate probability models for describing the agglomeration patterns of each industry.

2.3 Cluster-Detection Procedure

The key question remaining is how to find a “best” cluster-scheme for capturing the observed distribution of industry establishments. It is argued in Mori and Smith [45] that the *Bayes Information Criterion (BIC)* offers an appropriate measure of model fit in the present setting. In particular, for any given cluster scheme, $\mathbf{C}$, the (multinomial) log-likelihood of $\hat{p}_{\mathbf{C}}$ is given by

$$L_{\mathbf{C}}(\hat{p}_{\mathbf{C}}) = \sum_{j=0}^{k_{\mathbf{C}}} n_j(x) \ln \left(\frac{n_j(x)}{n} \right) + \sum_{j=0}^{k_{\mathbf{C}}} \sum_{r \in C_j} n_r \ln \left(\frac{a_r}{a_{C_j}} \right) \quad (1)$$

and that in terms of $L_{\mathbf{C}}(\hat{p}_{\mathbf{C}})$, the appropriate value of *BIC* is given for each candidate cluster scheme, $\mathbf{C}$, by

$$BIC_{\mathbf{C}} = L_{\mathbf{C}}(\hat{p}_{\mathbf{C}}) - \frac{k_{\mathbf{C}}}{2} \ln(n) . \quad (2)$$

Hence *BIC* is a “penalized likelihood” measure, where the second term in (2) essentially penalizes cluster schemes with a large number of clusters, $k_{\mathbf{C}}$, to avoid “over fitting” the data.

Given this criterion function, the present *cluster-detection procedure* amounts to a systematic way of searching the space of possible cluster probability models to

¹⁴This probability model is completed by the condition that $p_{\mathbf{C}}(R_0) = 1 - \sum_j p_{\mathbf{C}}(j)$.

find a cluster scheme, $\mathbf{C}^*$, with a maximum value of $BIC_{\mathbf{C}^*}$.¹⁵ While the details of this search procedure will play no role in the present analysis, the results of this procedure for Japanese industries will play a crucial role. Hence it is appropriate to illustrate these results in terms of the “livestock products” industry in Japan, shown in Figure 8 in Section 4.3.1 below.

Here Figure 8(a) shows the relative density of “livestock products” establishments in each municipality of Japan, where darker patches correspond to higher densities.¹⁶ The red patches surrounded by a solid curve in Figure 8(b) show the cluster scheme, $\mathbf{C}^*$, that was produced for the “livestock products” industry by this cluster-detection procedure.¹⁷ Here it is seen that not all isolated patches of density are clusters. But the highest density areas do indeed yield significant clusters. Notice also that the convex solidification procedure above has produced easily recognizable clusters that do seem to reflect the shapes of these high density areas.¹⁸

2.4 A Test of Spurious Clusters

Finally, it should be emphasized that even random locational patterns will tend to exhibit some degree of clustering. So there remains the statistical question of whether the “locally best” cluster scheme, $\mathbf{C}^*$, found for an industry by the above procedure is significantly better (in terms of BIC values) than would be expected in a random location pattern. Basically, a “random” location pattern is taken to be one in which location probabilities in all regions, $r \in R$, are proportional to their feasible areas, a_r . Hence a Monte Carlo test can be constructed by (i) generating N random location patterns for the establishments of a given industry, (ii) determining the locally optimal values, say BIC_s^* , for each simulated pattern,

¹⁵However, it should be emphasized that this space of probability models is very large, and hence that one can only expect to find *local* maxima (with respect to the particular perturbations defined by the search procedure itself).

¹⁶These municipalities are mapped in Figure 3 in Section 4.1.1.

¹⁷The red area within each cluster contains establishments of the “livestock products” industry, while there is no establishments in the pink area which has been incorporated into the cluster through convex-solidification. See also for the refinement of cluster scheme proposed by Mori and Smith [45, §5.3] which constructs a set of *agglomerations*, each of which consists of a set of contiguous clusters with a single peak of establishment density.

¹⁸A complementary clustering approach has recently been proposed by Kerr and Kominers [31] which identifies establishment clusters based on maximal interaction distances. This distance approach is particularly useful when relevant interactions can be documented, as in the case of patent citations within research-intensive industries.

$s = 1 \dots, N$, and (iii) comparing the value, $BIC_{\mathbf{C}^*}$, with this sampling distribution of BIC values. If $BIC_{\mathbf{C}^*}$ is sufficiently large (say in the top 1% of these values), then one may conclude that the clustering captured by $\mathbf{C}^*$ is significantly higher than what would be expected under randomness. Otherwise, $\mathbf{C}^*$ is said to involve *spurious clustering*. Results of this testing procedure for the application to Japanese manufacturing industries will be discussed in Section 4.2 below.

3 Spatial Scales of Agglomeration

As emphasized in the Introduction, the main strength of our cluster detection approach is to identify cluster schemes in a manner that preserves their two-dimensional spatial properties. By so doing, it is possible to analyze the spatial patterns of industrial agglomeration in more detail. As we will see for the case of Japanese manufacturing industries in Section 4, agglomerations of given industries often tend to concentrate within specific subregions of the country, i.e., are themselves “spatially contained”. Hence our first task below is to construct an operational definition of such containments, designated as the *essential containment* (*e-containment*) for each industry. Our next task is to construct a measure of the relative size of these *e*-containments, designated as *global extent*. Industries with small global extent can be regarded as relatively “confined”, and those with large global extent can be regarded as relatively “dispersed”. Finally, industries can also differ with respect to their patterns of agglomeration *within* these *e*-containments. Some patterns may be “dense” and others “sparse”. To compare such patterns, we construct a measure of the *local density* of clusters within each *e*-containment. This will yield a useful classification of agglomeration patterns in terms of their spatial scales to be discussed in Section 3.2.

3.1 Essential Containment

To formalize the notion of an industry’s essential containment, we start by assuming that an optimal cluster scheme, $\mathbf{C} = \mathbf{C}^*$, has been identified for the industry.¹⁹ The main idea is to identify an appropriate subset of “most significant” clusters in $\mathbf{C}$, and then take *essential containment* to be the convex solidification of this

¹⁹For notational simplicity we drop the asterisk in $\mathbf{C}^*$.

set of clusters in R . To identify “most significant” clusters, we proceed recursively by successively adding those clusters in $\mathbf{C}$ with maximum incremental contributions to BIC .²⁰ This recursion starts with the “empty” cluster scheme represented by $\mathbf{C}_0 \equiv \{R_{0,0}\}$ where $R_{0,0}$ denotes the full set of regions, R . If the set of (non-residual) clusters in $\mathbf{C}$ is denoted by $\mathbf{C}^+ \equiv \mathbf{C} \setminus \{R_0\}$, then we next consider each possible “one-cluster” scheme created by choosing a cluster, $C \in \mathbf{C}^+$, and forming $\mathbf{C}_0(C) = \{R_{0,0}(C), C\}$, with $R_{0,0}(C) = R_{0,0} \setminus C$. The “most significant” of these, denoted by $\mathbf{C}_1 = \{R_{1,0}(C), C_{1,1}\}$, is then taken to be the cluster scheme with the *maximum BIC value* (defined below). If this is called *stage* $t = 1$, and if the *most significant cluster scheme* found at each stage $t \geq 1$ is denoted by

$$\mathbf{C}_t \equiv \{R_{t,0}, C_{t,1}, \dots, C_{t,t}\} , \quad (3)$$

then the recursive construction of these schemes can be defined more precisely as follows.

For each $t \geq 1$ let $\mathbf{C}_{t-1}^+$ denote the (non-residual) clusters in $\mathbf{C}_{t-1}$ (so that for $t = 1$ we have $\mathbf{C}_{t-1}^+ = \mathbf{C}_0^+ = \emptyset$), and for each cluster not yet included in $\mathbf{C}_{t-1}$, i.e., each $C \in \mathbf{C}^+ \setminus \mathbf{C}_{t-1}^+$, let $\mathbf{C}_{t-1}(C)$ be defined by,

$$\mathbf{C}_{t-1}(C) = (R_{t-1,0}(C), C_{t-1,1}, \dots, C_{t-1,t-1}, C) , \quad (4)$$

where

$$R_{t-1,0}(C) = R_{t-1,0} \setminus C . \quad (5)$$

Then the *most significant additional cluster*, $C_t (\equiv C_{t,t}) (\in \mathbf{C}^+ \setminus \mathbf{C}_{t-1}^+)$, at stage $t \geq 1$ is defined by

$$C_t \equiv \arg \max_{C \in \mathbf{C}^+ \setminus \mathbf{C}_{t-1}^+} L(\widehat{p}_{\mathbf{C}_{t-1}(C)} | \mathbf{C}_{t-1}) , \quad (6)$$

where $L(\widehat{p}_{\mathbf{C}_{t-1}(C)} | \mathbf{C}_{t-1})$ is the *estimated maximum log-likelihood value* for model $p_{\mathbf{C}_{t-1}(C)}$ given [in a manner paralleling expression (1) above] by

$$L(\widehat{p}_{\mathbf{C}_{t-1}(C)} | \mathbf{C}_{t-1}) = \sum_{C' \in \mathbf{C}_{t-1}(C)} n_{C'} \ln \left(\frac{n_{C'}}{n} \right) + \sum_{C' \in \mathbf{C}_{t-1}(C)} \sum_{r \in C'} n_r \ln \left(\frac{a_r}{a_{C'}} \right) , \quad (7)$$

²⁰At this point it should be emphasized that the following procedure for identifying “significant clusters” in $\mathbf{C}$ is different from the one used to identify $\mathbf{C}$ in Section 2.3. In particular, the only candidate clusters now being considered are those in $\mathbf{C}$ itself.

where $n_{C'} \equiv \sum_{r \in C'} n_r$ and $n \equiv \sum_{r \in R} n_r$. Thus, at each stage $t \geq 1$ the likelihood-maximizing cluster, C_t , is removed from the residual region, $R_{t-1,0}$, and added to the set of significant clusters in $\mathbf{C}_{t-1}$. The resulting BIC value at each stage t is then given by

$$BIC_{\mathbf{C}_t} = L_{\mathbf{C}_t} - \frac{t}{2} \ln(n) \quad (8)$$

with

$$L_{\mathbf{C}_t} = \sum_{C \in \mathbf{C}_t} n_C \ln\left(\frac{n_C}{n}\right) + \sum_{C \in \mathbf{C}_t} \sum_{r \in C} n_r \ln\left(\frac{a_r}{a_C}\right). \quad (9)$$

Finally, the *incremental contribution* of each new cluster, C_t , to BIC is given by the increment for its associated cluster scheme, $\mathbf{C}_t$, as follows:

$$\Delta BIC_t \equiv BIC_{\mathbf{C}_t} - BIC_{\mathbf{C}_{t-1}}. \quad (10)$$

To identify the relevant set of “significant clusters” in $\mathbf{C}$, relevant requirements would depend on the objectives. For our present purpose of distinguishing the spatial scale of agglomeration, it suffices to impose a simple requirement that the sum of BIC contributions by the first $t^e \geq 1$ essential clusters accounts for at least a given *share*, $\lambda \in (0, 1]$, in that of $\mathbf{C}$:²¹

$$\sum_{t=1}^{t^e} \Delta BIC_t \geq \lambda BIC_{\mathbf{C}}. \quad (11)$$

If the set of *essential clusters* in $\mathbf{C}$ is now defined to be $\mathbf{C}^e = \mathbf{C}_{t^e}^+$, then the desired *essential containment* (*e-containment*), $ec(\mathbf{C})$, for an industry with cluster scheme $\mathbf{C}$ is taken to be the smallest convex-solid set in R containing $\mathbf{C}^e$, i.e., the convex solidification of $\mathbf{C}^e$.²²

These concepts can be illustrated by the stylized location patterns in Figure 2 below. For example, if the relevant cluster scheme, $\mathbf{C}$, for a given industry corresponds to the five clusters (shown in black) in Figure 2(a), and if the subset of essential clusters, $\mathbf{C}^e$, consists of the three largest clusters on the left, then the *e-containment*, $ec(\mathbf{C})$, for this industry is given by the filled square containing

²¹It would seem the most natural to simply add clusters as long as the increments are positive. But from the original construction of $\mathbf{C}$, it should be clear that these increments may often be positive for *all* $t = 1, \dots, k_{\mathbf{C}}$. See Mori and Smith [45, §4.4] for alternative requirements.

²²In terms of the d -convex solidification operator, σ_{cd} , defined in Mori and Smith [45, eq. (26)] (with respect to shortest-path travel distance, d), the formal definition of *e-containment* is given by $ec(\mathbf{C}) = \sigma_{cd}(\mathbf{C}^e)$.

these three clusters. Similar interpretations can be given to the filled rectangles of Figures 2(b,c,d).

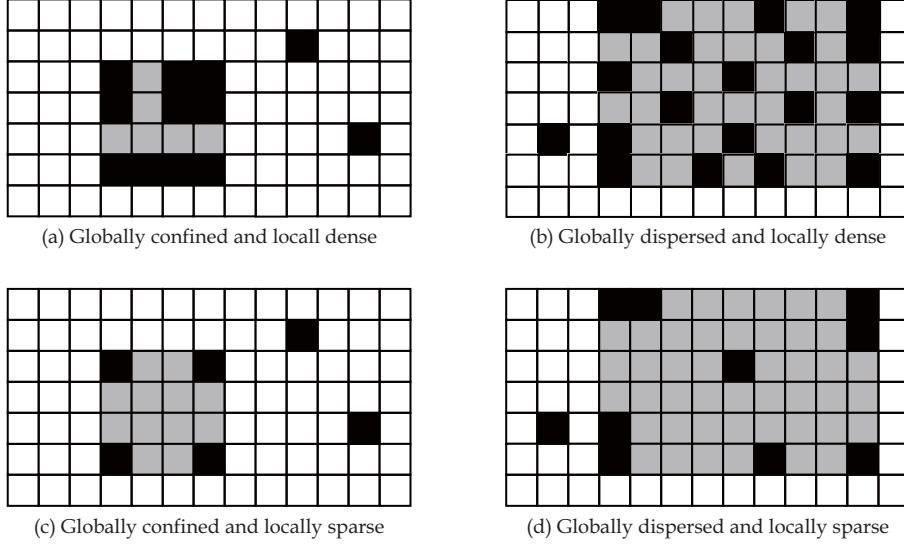


Figure 2: Classifications of agglomeration patterns

3.2 Global Extent and Local Density

With these definitions we next seek to compare e -containments for different industries in terms of their relative sizes. In order to reflect the actual spatial extent of such containments, it is now more appropriate to measure “size” in terms of total *geographic area* rather than the more limited notion of *feasible area* (employed for modeling the potential locations of individual establishments, as in Sections 2.2 above). Hence if we now let A to denote *geographic area*, then the economic areas for *basic regions* (a_r), *clusters* (a_C), and the *entire country* (a), are here replaced by A_r , A_C , and A , respectively. With these conventions, the *global extent* (GE) of an industry is now taken to be simply the total area of its e -containment, $ec(\mathbf{C})$, relative to that of the entire country:

$$GE(\mathbf{C}) = \frac{\sum_{r \in ec(\mathbf{C})} A_r}{A} \in (0, 1] . \quad (12)$$

Industries with relatively small global extents might be classified as “globally confined” industries [illustrated by the industries in Figures 2(a,c)]. Similarly, industries with substantially larger global extents might be classified as “globally dispersed” industries [illustrated by those in Figures 2(b,d)].²³

Finally, we consider the relative denseness of essential clusters within the e -containment for each industry. As a parallel to global extent, we now define the *local density* (LD) of a given industry to be simply the total area of its essential clusters, $\mathbf{C}^e$, relative to that of its e -containment, $ec(\mathbf{C})$, i.e.,

$$LD(\mathbf{C}) = \frac{\sum_{r \in \mathbf{C}^e} A_r}{\sum_{r \in ec(\mathbf{C})} A_r} \in (0, 1] . \quad (13)$$

Industries with a relatively high density of clusters in their e -containments might be classified as “locally dense” industries [illustrated by the industries in Figures 2(a,b)]. Similarly, industries with a substantially lower density of clusters in their e -containments might be classified as “locally sparse” industries [illustrated by those in Figures 2(c,d)].

More generally, Figure 2 is intended to summarize the main features of this classification system. First, the concept of the e -containment is designed to capture the region of most significant agglomeration for an industry. This is illustrated in each of the stylized figure panels by filled regions containing the largest clusters within the cluster schemes shown. In each case, the “outlier” clusters excluded from this region are implicitly assumed to be less significant in terms of their contributions to BIC .

Each of the four panels in this figure depicts a type of extreme case in the present classification system. However, it should be emphasized that there is no unambiguous ordering among these extremes. Indeed, it is a fundamental tenet of this paper that the types of concentration/dispersion continua implied by scalar measures of concentration are simply too limiting. In contrast, Figure 2 can be said to represent the extremes of a *two-dimensional* ordering: For any given level of Local Density, higher values of Global Extent tend to reflect industrial patterns that are more dispersed throughout the country. Similarly, for any given level of Global Extent, higher values of Local Density tend to reflect

²³One might consider more exact classifications, such as $GE < 1/2$ for “globally confined” and $GE \geq 1/2$ for “globally dispersed.” But in our view, the appropriate ranges of GE may often be context dependent.

industrial location patterns that are more dispersed throughout their essential containments. More detailed examples of these extremes will be developed in Section 4 below.²⁴

4 Detection of Industrial Clusters in Japan

In this section, we apply the above set of cluster-analytic tools to study the agglomeration patterns of manufacturing industries in Japan. We begin in Section 4.1 with a description of the relevant data for analysis. This is followed in Section 4.2 with a summary of results for the spurious-cluster test described in Section 2.4. The classification scheme developed in Section 3 is then given an operational form for the present application. Finally, this classification scheme is illustrated by means of a number of selected examples in Section 4.3.

4.1 Data for Analysis

The data required for this application includes both quantitative descriptions of the relevant system of regions and the class of industries to be studied. We consider each of these data types in turn.

4.1.1 Basic Regions

The relevant notion of a “basic region” for this analysis is taken to be the *shi-ku-cho-son*, which is a municipality category equivalent to a city-ward-town-village in Japan. The specific municipality boundaries are taken to be those of October 1, 2001.²⁵ While there are a total of 3363 municipalities in Japan, we only consider 3207 of these (as shown in Figure 3), namely those that are *geographically connected to the major islands of Japan (Honshu, Hokkaido, Kyushu and Shikoku)*

²⁴It should also be noted that the extremes in Figure 2 have differing implications for the overall *size* of the industries involved. In particular, only industries with many establishments can possibly exhibit dense patterns of significant clusters over large areas [such as Figure 2(b)], and only industries with small numbers of establishments can exhibit sparse patterns of significant agglomeration in confined areas. [such as Figure 2(c)]. This contrast can also be seen by comparing Figures 7 and 13 in Section 4 below.

²⁵The data source for these municipality boundaries is the Statistical Information Institute for Consulting and Analysis [51, 52].

via a road network. This avoids the need for ad-hoc assumptions regarding the effective distance between non-connected regions.

The only exception here is Hokkaido, which is one of the four major islands (refer to Figure 3), but is disconnected from the road network covering the other three. Given its size (217 municipalities), as seen in Figure 3, we still include Hokkaido as a potential location for establishments. Aside from this exceptional case, we adopt the following conventions. First, while we allow establishments to locate freely within the 3207 municipalities, we do not allow the formation of any clusters including municipalities in both Hokkaido and other major islands.²⁶ Second, e -containments for each industry are obtained as the union of the two convex solidified subsets of essential clusters within and without Hokkaido [see, e.g., the cases of “sliding doors and screens”, “livestock products”, and “manufactured ice” shown in Figures 7(b), 8(b) and 9(b), respectively, in Section 4.3 below].

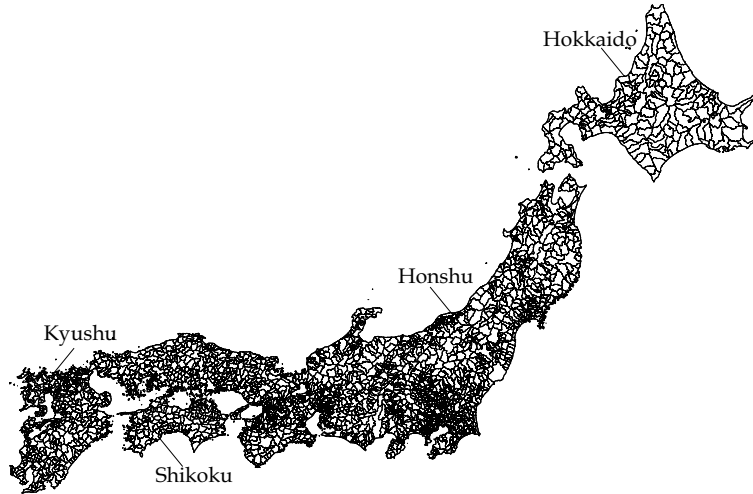


Figure 3: Municipalities in Japan

4.1.2 Economic Area

To represent the areal extent of each basic region we adopt the notion of “economic area”, obtained by subtracting forests, lakes, marshes and undeveloped area from the total area of the region (available from the Statistical Information

²⁶In terms of our δ -neighborhood definition in Mori and Smith [45, §4.2.2], the distances between Hokkaido regions and those of the major islands are implicitly assumed to exceed δ .

Institute for Consulting and Analysis [51, 52]).²⁷ The economic area of Japan as a whole (120,205km²) amounts to only 31.8% of total area in Japan. Among individual municipalities this percentage ranges from 2.1% to 100%, with a mean of 48.5%. Not surprisingly, those municipalities with highest proportions of economic area are concentrated in urban regions. In this respect, our present approach is relatively more sensitive to clustering in rural areas.²⁸

4.1.3 Interregional Distances

The travel distance between each pair of neighboring municipalities is computed as the length of the shortest route between their municipality offices along the road network.²⁹ From the computed pairwise distances between neighboring (contiguous) municipalities, the *shortest-path distances* (and associated sequences of neighboring municipalities) are computed in terms of Mori and Smith [45, eq.(15)].³⁰ While there is of course some degree of interdependency between industrial locations and the road network, the spatial structure of this network is mainly determined by topographical factors.

4.1.4 Industry and Establishments Data

Finally, the industry and establishments data used for this analysis is based on the Japanese Standard Industry Classification (JSIC) in 2001. Here we focus on three-digit manufacturing industries, of which 163 industrial types are present

²⁷There is of course a certain degree of interdependence between the size of economic areas and the presence of industries in those areas. In particular, industrial growth in a region may well lead to a gradual increase in the economic area of that region (say by land fills or deforestation). But to capture agglomeration patterns at a given point in time, we believe that it is more reasonable to adopt economic area than total area as the potential location space for establishments. In Japan, for example, it is doubtful that mountainous forested regions (which account for 98% of non-economic areas) can be easily be made available for industrial location in the short run.

²⁸In other words, for any given number of firms, n_r , in a basic region r , our clustering algorithm implicitly regards n_r as a more significant concentration in regions with smaller economic areas (other things being equal).

²⁹This road network data is taken from Hokkaido-chizu Co. Lit. [24], and includes both prefectural and municipal roads. However, if a given municipality office is not on one of these roads, then minor roads are also included.

³⁰As noted in Mori and Smith [45, §3.1], shortest-path distances are always at least as large as shortest-route distances. But in the present case, shortest-path distance appears to approximate shortest-route distance quite well. For the distribution of ratios of short-path over shortest-route distances across all 4,491,991 relevant pairs of municipalities, the median and mean are both equal to 1.14. In fact, the 99.5 percentile point of this distribution is only 1.28.

in the set of basic regions chosen for this analysis.³¹ The establishment counts (n) across these 163 industries is taken from the Establishment and Enterprise Census of Japan [30] in 2001. The mean and median establishment counts per industry are respectively 3958 and 1825. In addition, 147 (90%) of these industries have more than 100 establishments, and 125 (77%) have more than 500 establishments.

4.2 Tests of Spuriousness of Cluster Schemes

Using the cluster-detection procedure developed in Section 2.3, optimal cluster schemes, $\mathbf{C}_i^*$, were identified for each industry, $i = 1, \dots, 163$. Each cluster scheme, $\mathbf{C}_i^*$, was then tested for spuriousness using the testing procedure developed in Section 2.3.³² Among the 163 industries studied, the null hypothesis of complete spatial randomness (Section 2.4 above) was strongly rejected for 155 of these industries. For the remaining eight industries, this null hypothesis could not be rejected at the .01 level. The main reason for non-rejection in these cases (which include seven arms-related industries, together with “coke”), appears to be the small size of these industries, with $n < 40$ in all cases.³³ In view of these findings, we chose to drop the eight industries in question and focus our subsequent analyses on the 155 industries exhibiting significant clustering.³⁴

For these 155 industries, Figure 4 shows the frequency distribution of the share of establishments for each industry i that are included in the clusters of its cluster scheme, $\mathbf{C}_i^*$. These shares range from 39.1% to 100% with a median (mean) share of 95.2% (93.6%). The industries with the smallest shares of establishments in clusters are typically those which exhibit the weakest tendency for clustering. For instance, “paving materials” industry and “sawing, planning

³¹More precisely, out of the 164 industrial types in Japan, all but one have establishments in at least one of our basic regions.

³²These tests of spuriousness were based on the *BIC* values for a sample of 10,000 completely random location patterns for each industry.

³³These industries are also rather special in other ways. Arms-related industries are highly regulated industries, so that their location patterns are not determined by market forces, while “coke” is a typical declining industry in Japan (steel industries have gradually replaced coke production by less expensive powder coal after the 1970s).

³⁴In the application in Mori and Smith [45, §5] using the same data, 154 instead of 155 industries exhibited significant clustering based on 1000 samples for random establishment locations, instead of 10,000 samples in the present study. Specifically, “tobacco manufacturing” industry turned out to exhibit significant clustering under the present larger Monte Carlo simulation.

mills and wood products” industry have 39.1% and 54.0% of their establishments in the clusters, respectively. Since both of these industries are typically sensitive to transport costs, their establishment locations tend to reflect population density.

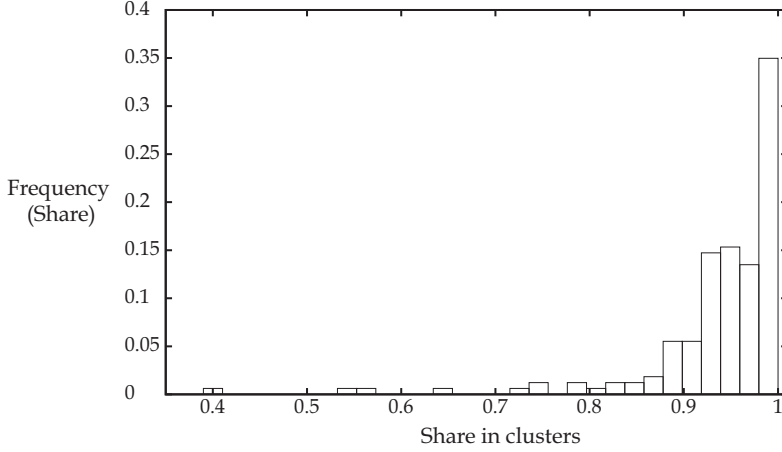


Figure 4: Share of establishment counts in clusters

4.3 Classification of Cluster Patterns

To apply our two measures (GE, LD) for classification purposes, we begin by recalling that the key parameter defining e -containments for industries with cluster schemes, $\mathbf{C}$, is the *share*, λ , of the total $BIC_{\mathbf{C}}$ values accounted for by clusters in these e -containments. So it is necessary to specify an appropriate value of λ . In terms of classification, it is useful to consider the consequences of λ for possible correlations between GE and LD . For if these measures are too highly correlated (either positive or negative), then it is doubtful that they can both provide distinct information useful for classification purposes. With this in mind, we first observe if λ is very small, then e -containments will include only a few highly significant clusters. If these clusters are concentrated in a small region for a given industry, then GE will be small and LD is likely to be large. Conversely, if these clusters are widely separated for a given industry, then GE will be large and LD is likely to be smaller. So for small λ it seems clear that GE and LD should be strongly negatively correlated across industries. On the other hand, if λ is very large, then e -containments will tend to include almost all of an industry’s clusters. So the question is whether industries that are more spread out (i.e.,

with higher GE values) also tend to have denser cluster patterns (i.e., higher LD values). In our data this appears to be the case, so that GE and LD are in fact positively correlated at high values of λ .

These observations are quantified in Figure 5, where we have plotted the (Pearson) correlations, ρ , between GE and LD across our 155 industries (with non-spurious cluster schemes) for a the full range of λ values. Here the solid red curve shows correlation values, ρ , and the dashed blue curve shows the corresponding p -values (for a two-sided test of ρ significance). As is seen in the figure, ρ is significantly negative (at the 0.05 level) for λ less than about 0.67, and is significantly positive for λ above 0.92. Moreover, since p -values rise sharply between these two extremes, it can be concluded that GE and LD are essentially uncorrelated within the range, $\lambda \in [0.67, 0.92]$, so that industries are most differentiated in terms of their agglomeration patterns within this range of λ .

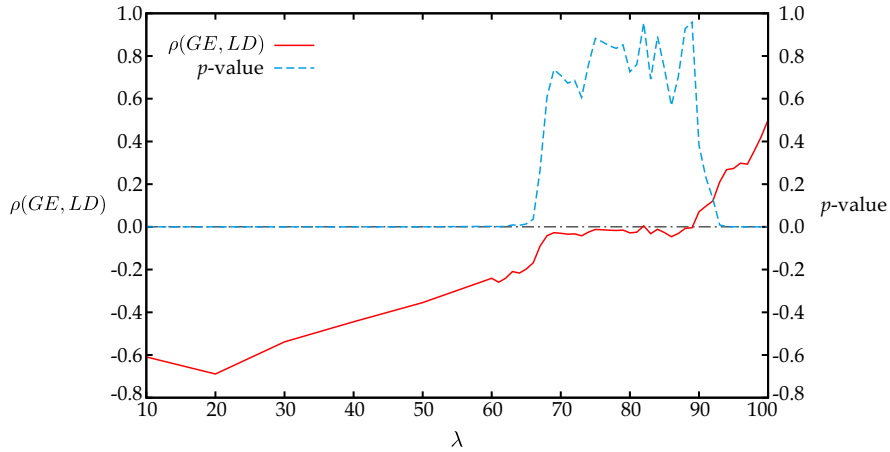


Figure 5: Correlation between global extent and local density of clusters

In particular, the correlation between GE and LD is seen to be least significant at approximately $\lambda = 0.88$. For this value of λ , we have plotted the pairs (GE, LD) for each of the 155 industries in Figure 6. Here it is seen that GE and LD are essentially unrelated, so that all four extremes in Figure 2 do in fact occur simultaneously. For convenience, the relative positions of panels (a) through (d) in Figure 2 are arranged to match the relative positions in Figure 6. For example, the types of globally confined patterns illustrated in the left panels (a,c) of Figure 2 are typical of industries with (GE, LD) pairs in the left portion of Figure 6. Similarly, the locally dense patterns in the top panels (a,b) of Figure 2 are

typical of industries with (GE, LD) pairs toward the top of Figure 6.

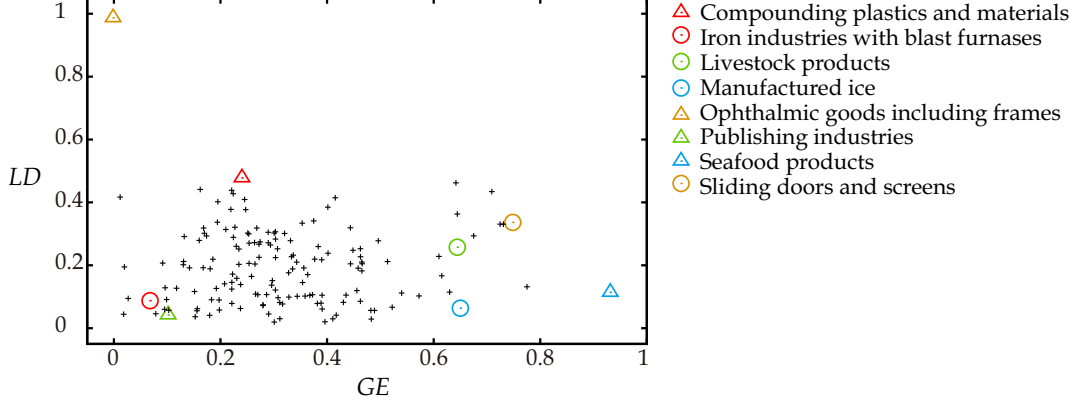


Figure 6: LD versus GE under $\lambda = 0.88$

Within this general framework, it is of interest to consider more detailed examples of industries with cluster schemes exhibiting a variety of (GE, LD) combinations. Here we focus on the case, $\lambda = 0.88$, in Figure 6 which exhibits the widest variation of GE and LD values.³⁵ Figures 7 through 12 focus on different industries. For each industry i , panel (a) of the figure shows the density of i establishments across municipalities (where darker colors denote municipalities with higher densities). Panel (b) of the figure shows both the spatial pattern of clusters and their e -containment for industry i . Here individual clusters are represented by the enclosed red areas,³⁶ and the corresponding e -containment (for $\lambda = 0.88$) is shown in yellow.

4.3.1 Globally Dispersed and Locally Dense Patterns

Industries with high values of both GE and LD (located in the upper-right portion of Fig. 6) can be described as exhibiting patterns of agglomeration that are “globally dispersed and locally dense”. Such industries are by definition present almost everywhere, and can equivalently be described as *ubiquitous industries*. As discussed in Section 3.2, this pattern is evaluated as the “maximally dispersed” in terms of scalar indices of agglomeration. A typical example is the

³⁵The basic results remain the same for different values of $\lambda \in [0.67, 0.92]$ in which the correlation between GE and LD are insignificant.

³⁶The portion of each cluster in pink shows those basic regions which contain no establishments (but are included in the cluster by the process of convex solidification).

“sliding doors and screens” (with $GE = 0.749$, $LD = 0.336$; $\odot$ in Fig. 6). As indicated by Figure 7(a), establishments are present almost all municipalities, and the clusters are found to be densely distributed throughout the country. Their products are often custom made and require face-to-face contact with customers, and hence, there are strong market-attraction forces that contribute to the ubiquity of this industry.

It is also of interest to note (as mentioned in footnote 24) that such ubiquitous industries are by their very nature quite large in terms of establishment numbers. In the present case, “sliding doors and screens” industry has 15,363 establishments, which is well above the mean of 4189 for all industries. In terms of establishments in clusters, this industry has 13,565 establishments relative to the mean of only 3896 for all industries.

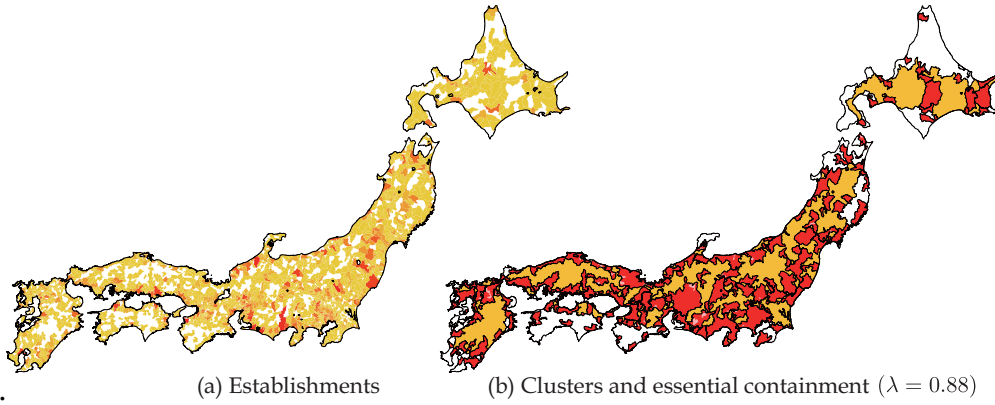


Figure 7: Location pattern of “sliding doors and screens” industry

Figure 8 shows the location patterns of another ubiquitous industry, “live-stock products”. The clusters of this industry exhibit slightly smaller global extent and local density ($GE = 0.645$ and $LD = 0.258$; $\odot$ in Fig. 6) than does “sliding doors and screens” industry above. But, they are still relatively globally dispersed and locally dense. The reason for ubiquity of clusters in this industry is straightforward, since freshness is critical for most of its products so that market proximity is a major determinant of establishment locations.

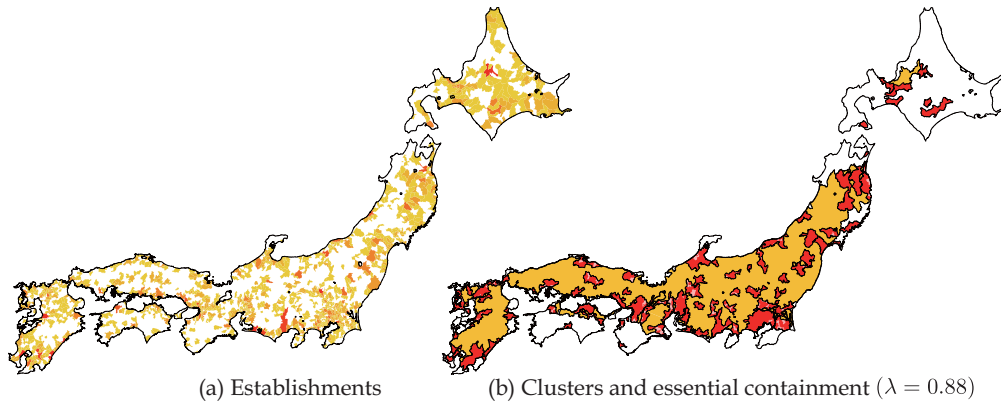


Figure 8: Location pattern of “livestock products” industry

4.3.2 Globally Dispersed and Locally Sparse Patterns

Industries with relatively high values of GE and low values of LD (near the lower-right portion of Fig. 6) can be described as exhibiting patterns of agglomeration that are “globally dispersed and locally sparse”. A clear example is provided by the “manufactured ice” industry shown in Figure 9 (with $GE = 0.589$ and $LD = 0.133$; $\odot$ in Fig. 6). Global dispersion here reflects the high cost of shipping ice, while local sparseness suggests that there are scale economies in production. In fact, the number of establishments in this industry is only 387 which is about one tenth of the mean establishment counts of all the three-digit manufacturing industries.

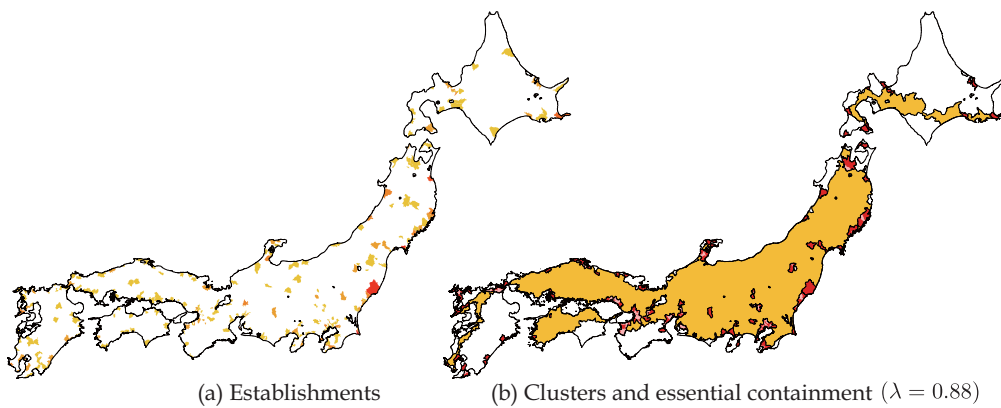


Figure 9: Location pattern of “manufactured ice” industry

Another extreme example is provided by the “seafood products” industry depicted in Figure 10 (with $GE = 0.931$ and $LD = 0.116$; $\triangle$ in Fig. 6). The primary location determinant for this industry is obviously proximity to the coast, so that establishment locations are dense along the coast but sparse inland.

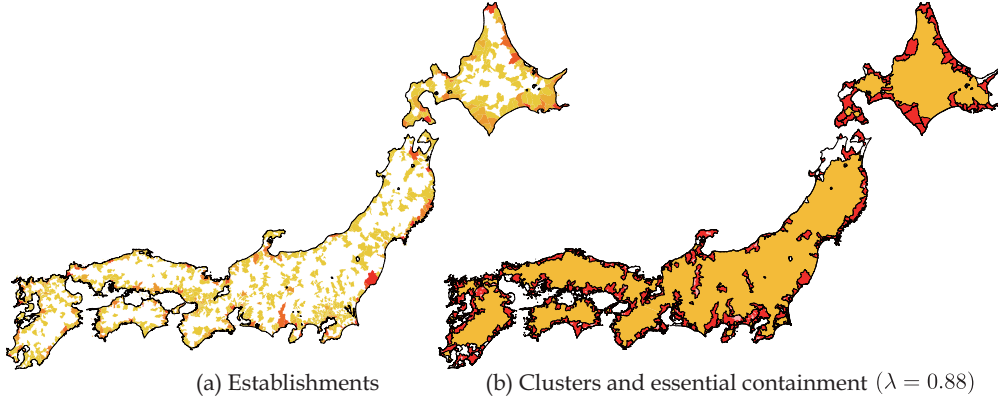


Figure 10: Location pattern of “seafood products” industry

4.3.3 Globally Confined and Locally Dense Patterns

Industries with relatively low values of GE and high values of LD (in the upper-left portion of Fig. 6) can be described as exhibiting patterns of agglomeration that are “globally confined and locally dense”. An extreme example of such industries is provided by the “ophthalmic goods, including frames” industry in Figure 11 (with $GE = 0.009$ and $LD = 0.988$; $\triangle$ in Fig. 6). This industry is strongly concentrated in a single town, Sabae, with a population of only 65,000. In fact, this one small town accounts for more than 90% of the national market share in ophthalmic goods. Not surprisingly, the e -containment for this industry consists only of this single town, as shown in Figure 11(b). As with many specialized industries, the location pattern of this industry is governed more by historical circumstances than economic factors. In terms of establishment counts, such industries are necessarily small in size. In the present case, there are only 1139 establishments, which is well below the mean of 4188 for all industries. So even though all of its 1139 establishments are in clusters, this number is still well below the mean of 3896 for all industries.

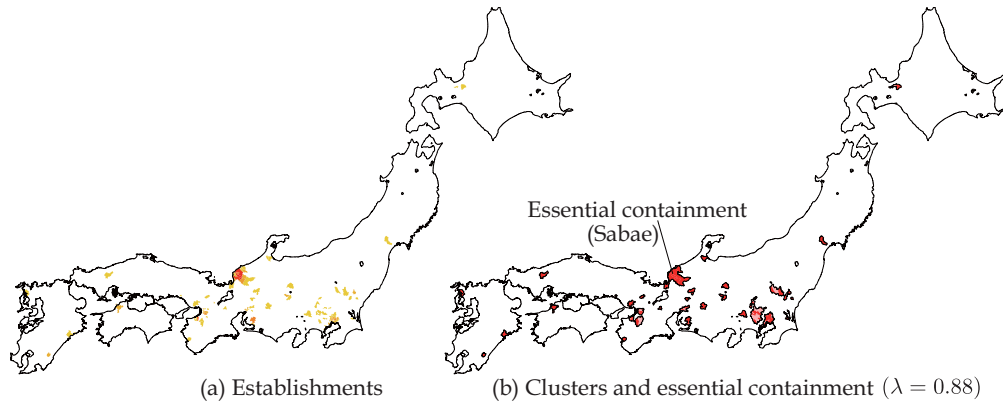


Figure 11: Location pattern of “ophthalmic goods including frames” industry

A second example is provided by the “compounding plastics and reclaimed plastics” industry (with $GE = 0.240$ and $LD = 0.478$; $\triangle$ in Fig. 6). From Figure 12, it is clear that most clusters for this industry, and indeed most of its establishments, lie in the Industrial Belt. The outputs of this industry are primarily intermediate inputs for a variety of manufactured goods produced along the Belt, particularly home electronics appliances and motor vehicles (such as the molded plastic parts for seats, fenders, and instrument panels). Thus the intermediate locations between these manufacturers constitute natural market-oriented locations for this industry. In fact, many industries with (GE, LD) values similar to this industry also exhibit Industrial-Belt type agglomerations.

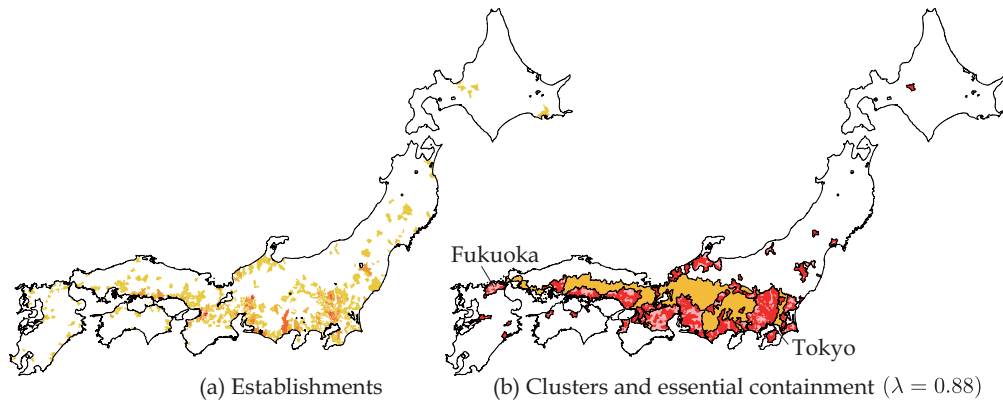


Figure 12: Location pattern of “compounding plastic materials” industry

4.3.4 Globally Confined and Locally Sparse Patterns

Finally, “globally confined and locally sparse” agglomeration patterns (in the lower-left portion of Fig. 6) are mostly exhibited (in Japan) by those industries with establishments concentrated in the major cities along the Pacific coast. A representative case is provided by the “iron industry with blast furnaces” industry (with $GE = 0.068$ and $LD = 0.090$; $\odot$ in Fig. 6), where plant-level scale economies are so large that the entire industry consists of only 38 establishments. As seen in Figure 13(a), most establishments are concentrated around the major ports along the Pacific coast (in order to gain access to both their imported inputs and largest output markets). Since these major ports are widely spaced along the coast from Tokyo to Oita (more than 1000km apart), clustering also appears to be locally sparse in this region.

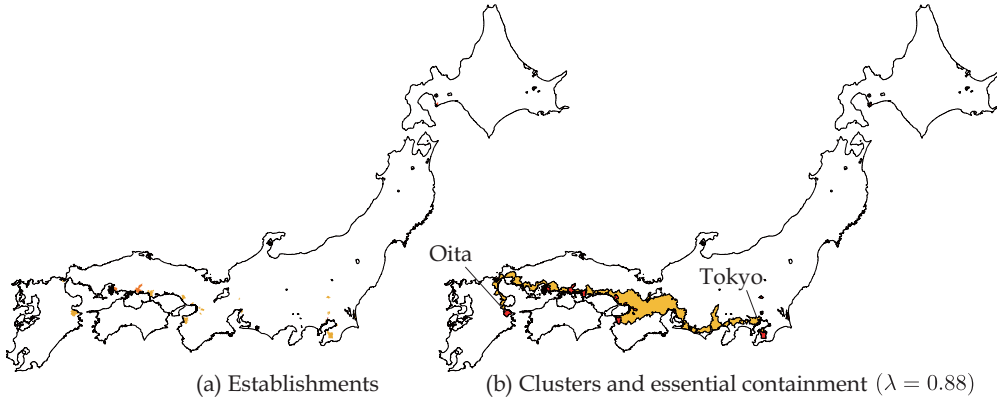


Figure 13: Location pattern of “iron industry with blast furnaces” industry

A final example is provided by the “publishing industry” depicted in Figure 14 (with $GE = 0.354$ and $LD = 0.145$; $\triangle$ in Fig. 6). Publishing is typical of “urban-oriented” industries with location patterns tending to reflect urban density. While both the establishments and clusters of this industry are spread throughout the country, Figure 14 shows that there is relatively more concentration in the Pacific coast area between Tokyo and Osaka, with a narrow band

stretching beyond Osaka to include the major metro areas further west (Kobe, Okayama, Hiroshima, and Fukuoka).

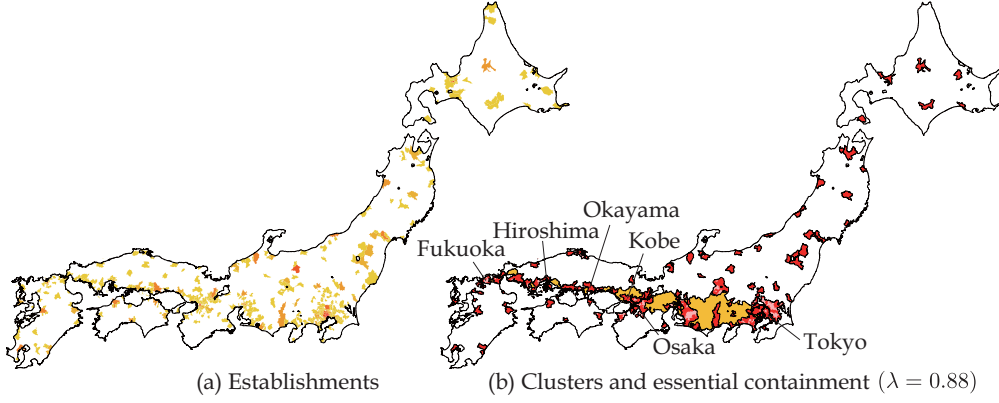


Figure 14: Location pattern of “publishing industries”

5 Comparisons with Scalar Indices

The most dominant approach to agglomeration comparisons between industries has been in terms of scalar measures of the overall *degree* of industrial agglomeration (see, e.g., Rosenthal and Strange [49] for a survey). These indices are computed by measuring the discrepancy between the spatial distribution of establishments within an industry and a given reference distribution representing “complete dispersion” of establishments.³⁷ But, not surprisingly, such scalar measures often yield similar values for industries with very different spatial patterns of agglomeration.

As will be seen below, the scalar index which is most closely related to our cluster detection approach is the *D-index* developed by Mori et al. [42]. This *D-index* for a given industry i is defined by the Kullback-Leibler [34] divergence of its establishment location probability distribution, $P_i \equiv [P_i(r) : r \in R]$, from a purely random establishment location patterns, $P_0 \equiv [P_0(r) : r \in R]$, as defined in Section 2.3 above. By using the sample estimate of P_i , namely, $\hat{P}_i = [\hat{P}_i(r) : r \in R]$

³⁷Refer to footnote 1 for the choice of reference distributions in the existing indices.

with $\widehat{P}_i(r) \equiv n_r/n$, a corresponding estimate of this D -index is given by

$$D(\widehat{P}_i|P_0) = \sum_{r \in R} \widehat{P}_i(r) \ln \left(\frac{\widehat{P}_i(r)}{P_0(r)} \right). \quad (14)$$

The intuition behind this particular index is that it provides a natural measure of distance between probability distributions. So if uniformity is taken to represent the complete absence of clustering, then it is reasonable to assume that those distributions “more distant” from the uniform distribution should involve more agglomeration. Note that since both D and BIC given by (2) are based on similar log-likelihood measures of “distance from uniformity”, our cluster identification procedure is closer in spirit to this scalar measure than other possible choices.

In terms of popularity, the primary index has been the γ -index developed by Ellison and Glaeser [13]. For a given industry $i \in I$, the γ -index is defined by

$$\gamma_i = \frac{G_i - (1 - \sum_{r \in R} x_r^2) H_i}{(1 - \sum_{r \in R} x_r^2) (1 - H_i)}. \quad (15)$$

In (15), G_i represents the Herfindahl-Hirschman index of employment concentration of industry i given by $\sum_{r \in R} (x_{ir} - x_r)^2$, where x_{ir} and x_r are the shares of region $r \in R$ in the total employment of industry i and that of the aggregate industry, respectively³⁸; H_i is the Herfindahl-Hirschman index of employment distribution across all the establishments in industry i given by $\sum_{j \in E_i} h_j^2$, where E_i is the set of all establishments in the industry, and h_j is the share of establishment $j \in E_i$ in the total employment of industry i . Notice that the definition of “complete dispersion” for this index is different from the D -index above as well as our present cluster detection. Specifically, γ measures the *squared deviation* of the employment distribution of industry in question from that of the aggregate industry (with certain adjustments for heterogeneity in establishment sizes), which means that industries whose establishments are either more spatially concentrated or dispersed than that of the aggregate industry are evaluated as *equally* more concentrated than the aggregate industry. As pointed out below, this raises certain questions about the interpretation of γ .

Another popular index is the one proposed by Duranton and Overman [10]. They start by computing the Euclidean distance between each pair of establish-

³⁸The “aggregate industry” in the present case is all manufacturing.

ments in a given industry $i \in I$. For given n establishments in this industry, the estimator of the density of bilateral distances, called K -density, at each distance level, d , is defined by

$$\widehat{K}_i(d) = \frac{1}{n_i(n_i - 1)h} \sum_{j=1}^{n_i-1} \sum_{k=j+1}^{n_i} f\left(\frac{d - d_{jk}}{h}\right), \quad (16)$$

where d_{jk} is the distance between establishments j and k , f the Gaussian kernel function, and h the bandwidth set according to Silverman [50, §3.4.2]. Roughly speaking, $\widehat{K}_i(d)$ is larger when the distances between many establishment pairs in industry i are approximately d . For each industry i , this K_i -density is then compared with the *counterfactual K -density* estimated from 1000 simulations of bilateral distances between n_i randomly sampled (distinct) establishments in the aggregate industry.

To identify the distance levels at which the industry in question exhibit significant concentration (or dispersion), Duranton and Overman [10] distinguish between “local” and “global” confidence bands. Our interest focuses only on global confidence bands, which are defined in the following way. First, one defines *local $p\%$ confidence bands* by identifying the p -percentiles of the simulated counterfactual distributions of $K(d)$ values at each distance $d = 0, 1, \dots, 296$, and then interpolating these percentile points into continuous bands, where $d = 296\text{km}$ is the median bilateral distance of all the establishments. Given these local bands, one then defines the *upper* [resp., *lower*] 5% *global confidence* $\overline{K}_i(d)$ [resp., $\underline{K}_i(d)$] for this industry to be the highest [resp., lowest] local confidence band that is hit by at least 5% of the simulated counterfactual K -densities. In these terms, industry i is said to be *localized* if $\widehat{K}_i(d) > \overline{K}_i(d)$ for at least one distance $d \in [0, 296]$, and similarly, is said to be *dispersed* if it is not localized, and $\widehat{K}_i(d) < \underline{K}_i(d)$ for at least one $d \in [0, 296]$.³⁹ In these terms, the *degree of localization* at each distance, d , is defined by

$$\Gamma_i(d) \equiv \max \left\{ \widehat{K}_i(d) - \overline{K}_i(d), 0 \right\}, \quad (17)$$

³⁹Duranton and Overman [10] use the respective terms “globally localized” and “globally dispersed”.

and the corresponding *degree of dispersion* is defined by

$$\Psi_i(d) \equiv \begin{cases} \max \{ \underline{K}_i(d) - \widehat{K}_i(d), 0 \} & , \text{ if } \sum_{d=0}^{296} \Gamma_i(d) = 0 , \\ 0 & , \text{ otherwise .} \end{cases} \quad (18)$$

While the overall degrees of localization and dispersion for a given industry i are defined separately in Duranton and Overman [10] by $\Gamma_i \equiv \sum_{d=0}^{296} \Gamma_i(d)$ and $\Psi_i \equiv \sum_{d=0}^{296} \Psi_i(d)$, respectively, these can be combined to define a single *localization index* as follows:

$$\Gamma_i^* \equiv \Gamma_i - \Psi_i , \quad (19)$$

where industry i is a *localized* (*dispersed*) industry (relative to the aggregate industry) if Γ_i^* is positive (negative).

To relate these indices to GE and LD , the most direct approach is simply to plot their pairwise relations (for $\lambda = 0.88$) as in Figure 15, where these relations are seen most clearly in terms of $\log(GE)$ and $\log(LD)$ [and where γ is has also been transformed to $\log \gamma$].⁴⁰

⁴⁰ Γ^* values were computed by using the R-package, *dbmss*, developed by Marcon et al. [38]. We thank Kohei Takeda for his research assistance.

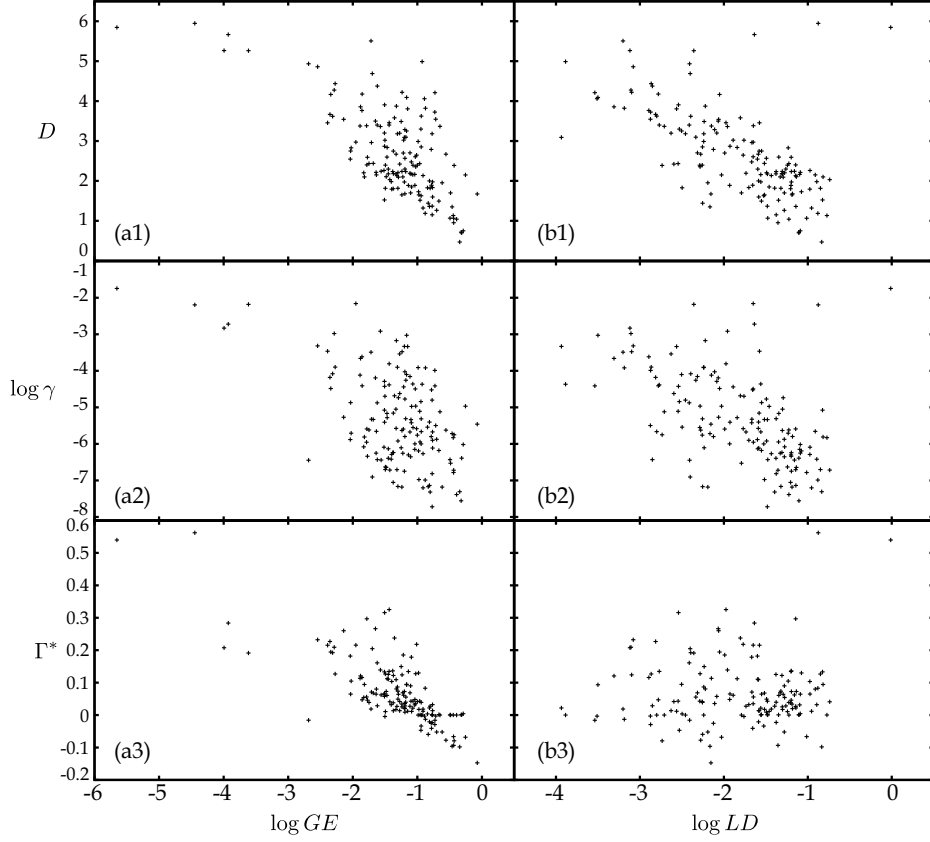


Figure 15: Scalar indices versus GE and LP

Here it is clear that both D and γ are significantly (negatively) correlated with both GE and LD , while Γ^* is correlated only with GE .⁴¹ But recall that since GE and LD are uncorrelated (for $\lambda = 0.88$), these visual relations can best be quantified in terms of the following multiple regression model:

$$Y = a + b \log GE + c \log LD + \varepsilon \quad (20)$$

where $Y = D, \log \gamma$ or Γ^* and where a, b and c are parameters to be estimated (assuming normal errors, ε). The results of these regressions are shown in Table 1, where all the visual observations above are confirmed.

⁴¹Spearman's rank correlations between GE and (D, γ, Γ^*) and are respectively $(-.0574, -0.375, -0.773)$, and between LD and (D, γ, Γ^*) are respectively $(-0.681, -0.598, 0.060)$, where only the correlation between LD and Γ^* is not significant.

Index	D	$\log \gamma$	Γ^*
Intercept	-0.425 (-3.80)	-8.212 (-36.63)	-0.060 (-3.46)
$\log GE$	-1.035 (-22.68)	-0.914 (-10.07)	-0.106 (-15.03)
$\log LD$	-0.937 (-19.89)	-0.915 (-9.59)	0.007 (0.91)
adj. R^2	0.853	0.561	0.594
#Obs.	155	153	155

(The numbers in parentheses are t -values.)

Table 1: Regression results for the scalar indices of agglomeration

Moreover, the adjusted R^2 results suggest that these scalar indices ($D, \log \gamma, \Gamma^*$) are reasonably well approximated by their predicted values ($\hat{D}, \widehat{\log \gamma}, \hat{\Gamma}^*$) as linear combinations of $\log(GE)$ and $\log(LD)$. This is confirmed by the regression plots shown in Figure 16.

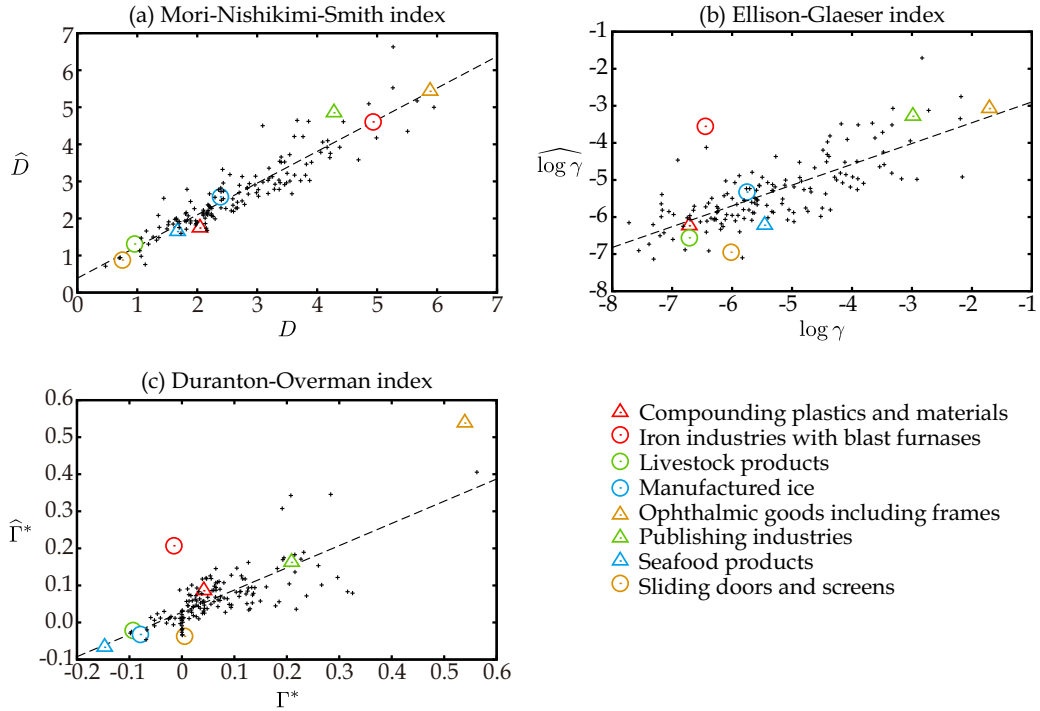


Figure 16: Scalar indices of agglomeration

These regressions help to illustrate the more important similarities and distinctions between the three indices in terms GE and LD . With respect to similarities, it should be clear that all these indices tend to agree when industries are

unabiguously concentrated in space, i.e., when GE is extremely small. This is well illustrated by the “ophthalmic goods” industry (Fig. 11), which corresponds to the symbol, $\triangle$, at the extreme end of the clustering spectrum on all three indices. A less extreme example is provided by “publishing industries” (Fig. 14) with GE again quite small and with symbol, $\triangle$, located toward the extreme clustering end for all three indices.

But aside from these extreme cases, the interpretations of such scalar indices can often be quite ambiguous. In particular, it is difficult for these indices to differentiate between “globally confined and locally dense” and “globally dispersed and locally sparse” patterns – which can be very different. Such differences are most often related to the *spatial scale* of agglomeration in the relevant industries. A good example of the first type of industry is provided by “compounding plastics materials” (refer to Fig. 12) with e -containment confined to the Industrial Belt stretching for more than 1000km along the Pacific coast area, but with e -clusters quite densely packed inside this area. The spatial scale of agglomeration for this industry is thus best described by the Industrial Belt itself. More generally, industries with relatively large LD compared to GE can be said to exhibit agglomeration at *larger* spatial scales. The converse is true for industrial patterns that are globally dispersed but locally sparse, i.e., with relatively large GE compared to LD . A good example is provided by the “manufactured ice” industry (refer to Fig. 9) where agglomeration is seen to occur at a *smaller* spatial scale, in this case extending only over a few widely dispersed municipalities. But in spite of the differences between these two types of industrial patterns, such industries are often grouped closely together by scalar indices. For “compounding plastics materials” and “manufactured ice” in particular, this is seen to be true for all three indices (as indicated by the closeness of their respective positions, $\triangle$ and $\odot$ on the horizontal axes in Figure 16.⁴²

Turning now to a more detailed consideration of these three indices themselves, note first from the adjusted R^2 values in Table 1 (as well as an inspection of Figure 16) that the D -index is most fully captured by model (20). Note in particular that since the estimated coefficients of both $\log(GE)$ and $\log(LD)$ for D are close to one, the relative values of D are well approximated by $-\log(GE \times LD)$. Moreover, since the product, $GE \times LD$, is seen from (12) and (13) to be simply the

⁴²Note however that “closeness” between Γ^* values on either side of zero is somewhat more difficult to gauge.

areal share of an industry’s e -clusters within the nation as a whole, it follows that D itself is essentially a decreasing function of this areal share. This simple relation is due largely to the fact that the uniform reference measure on which D is based is essentially area itself.

Turning next to the γ -index, note from Table 1 that the estimated coefficients of both $\log(GE)$ and $\log(LD)$ for $\log(\gamma)$ are almost identical, so that γ is again seen to be essentially decreasing in areal share. This accounts for much of the similarity in the behavior of D and γ . But note also that there are important differences, as seen by the much larger degree of unexplained variation in $\log(\gamma)$ [i.e., lower adjusted R^2]. As mentioned above, this is largely due to the *unsigned* nature of squared deviations implicit in γ , which can in principle equate very different types of patterns. This is well illustrated by a comparison of the spatially concentrated “iron industries with blast furnaces” (Fig. 13) with the much more ubiquitous “sliding doors and screens” industry (Fig. 7), as denoted respectively by $\odot$ and $\circ$ in Fig. 16. This difference is strongly reflected by D in panel (a) where the “sliding doors and screens” industry is seen to be much more uniformly distributed (i.e., smaller D). But the γ -index essentially equates the two, reflecting the fact that these two industries are deviating in opposite directions from the aggregate industry.

Turning finally to Γ^* , we begin by observing from Table 1 that this index is far more sensitive to GE than to LD . This asymmetry is partly explained by the fact that Γ^* focuses entirely on bilateral distances, whose magnitudes are far more sensitive to GE than to LD . But a more subtle factor contributing to this difference is the relation of bilateral distances for individual industries to those of the aggregate industry. As documented by Mori et al. [43] and Mori and Smith [44], clustering tends to be spatially coordinated *across* industries, so that clusters of many of industries tend to coincide in larger cities.⁴³ (As an extreme case, Tokyo contains clusters of all industries.) So when sampling counterfactuals from the aggregate industry, there tend to be larger numbers of small distances than would be expected for individual industries. The result is that such frequency comparisons tend to *understate* the significance of local concentrations for individual industries relative to the aggregate industry. In fact, if one considers all 41 industries, i , that are “globally dispersed and locally

⁴³This result holds also for the US manufacturing industries as evidenced Akamatsu et al. [1].

sparse” in the sense that GE_i is above the median and LD_i is below the median, then it turns out that *none* of these 41 exhibit significant localization at distances below 100km relative to the aggregate industry.

This can be illustrated in more detail by comparing two globally dispersed industries “Livestock products” (Fig. 8) and “manufactured ice” (Fig. 9) with similar global extents ($GE = 0.645$ and 0.589) but with very different local densities ($LD = 0.258$ and 0.133) reflecting the more locally concentrated nature of “manufactured ice”. While both D and γ reflect this difference, and evaluate “manufactured ice” to be more concentrated, these two industries are essentially indistinguishable in terms of Γ^* [compare the relative locations of $\odot$ and $\odot$ in panels (a) and (b) with those in panel (c) of Fig. 16]. For as seen by their respective K -densities in Figure 17, all differences between these two patterns are completely overwhelmed by the lack of any discernible localization at small distances under such K -density tests. More generally, tests of localization at small distances tend to be much more conservative under Γ^* than under D or γ .⁴⁴

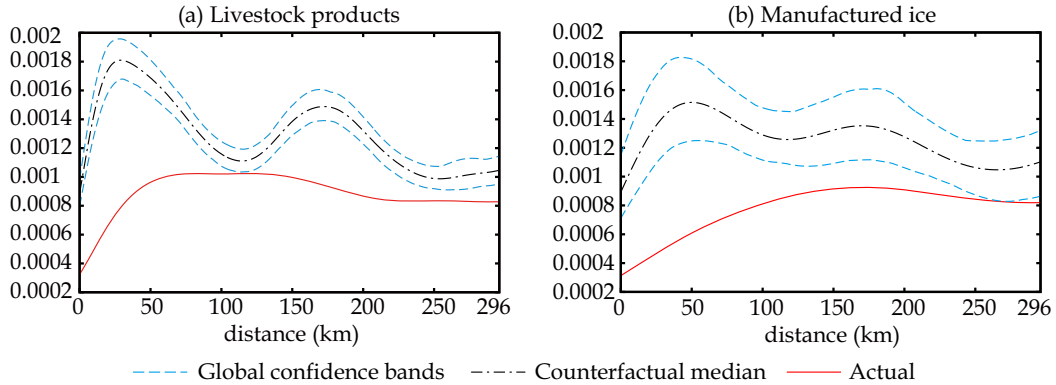


Figure 17: K -densities, global confidence bands, and median counterfactual K -densities of two illustrative industries

⁴⁴See Marcon and Puech [37] for more detailed discussions on the interpretation of K -densities by Duranton and Overman [10].

6 Shipment Distances and Spatial Patterns of Agglomeration

Theoretical relationships between transport costs and the spatial patterns of industrial agglomeration have been studied extensively in the context of NEG (see, e.g., Fujita et al. [17]). Here it has been shown that the influence of transport costs on spatial patterns of industrial location is complex, and in particular that interactions between global and local dispersion forces play a key role (see Fujita and Mori [20] for a survey). On the one hand, the spatial dispersion of consumers (driven mainly by land-intensive production together with a general scarcity of usable land) generates a *global dispersion force* in which industries with higher transport costs tend to spread over spatially dispersed local markets (both cities and towns) in order to minimize their transport costs to final markets. In our terminology, industries with more dispersed cluster patterns (higher *GE*) should thus be those in which firms tend to ship more locally.

On the other hand, there are two types of *local dispersion forces* affecting industries. One is a *crowding-out force* due to congestion and local scarcity of land that motivates some firms (and residents) to expand existing clusters, rather than form new clusters (as in the case of global dispersion above).⁴⁵ The other is a *filling-in force* that tends to transform collections of distinct clusters into a continuum of clusters (as in the formation of “industrial belts”). This happens for example when firms in footloose industries with relatively lower transport costs are attracted to locations between existing clusters to gain access to markets in more than one cluster.⁴⁶ Under both crowding-out and filling-in forces, local dispersion takes place that tends to leave the degree of global dispersion relatively unaffected. In our terminology, one thus expects industries with relatively lower transport costs to exhibit more locally dispersed cluster patterns (higher *LD*)

⁴⁵In continuous-location NEG models (such as Fujita and Krugman [15]), where land is neither consumed nor used as inputs in cities, each city initially occupies only a single point in space. But as populations grow and congestion externalities increase, mobile agents in such cities have incentives to relocate just outside the city, where they can avoid congestion costs while still enjoying proximity to the city market. In this sense, cities can be said to expand spatially in equilibrium. Similar crowding-out effects are found in the many-region extension of the model by Helpman [22] in which land scarcity is the primary dispersion force (see Akamatsu and Takayama [2] for more detail).

⁴⁶See Mori [41] for the theoretical mechanism underlying the formation of a continuum of cities.

for any given level of GE . Such dispersion in turn implies that these industries should also exhibit longer shipment distances to their extended markets.

In this section, we investigate the theoretical predictions above by focusing on the shipment distances for individual establishments obtained from the 2000 Net Freight Flow Census [40] for Japanese two-digit manufacturing industries. Although our ultimate goal is to relate transport costs directly to spatial patterns of industries,⁴⁷ we believe that shipment distances often reflect sensitivity to transport costs more directly than observed transport costs themselves.⁴⁸ In particular, while transport costs of outputs may appear to be large for an industry, the importance of these costs can only be gauged relative to sales revenues and transport costs of inputs. Moreover, even when such costs are of major importance for an industry, there is often a strong interdependence between the observed transport costs and realized location patterns of individual firms. So the importance of these costs as locational determinants may be better reflected by observed shipment distances.

Under the assumption that higher transport cost sensitivities do translate into shorter shipment distances, we seek to determine whether the observed shipment distances of different industries can provide empirical support for the theoretical predictions above.⁴⁹ But since industrial shipment data in Japan is only available at the two-digit level of classification, we must here analyze spatial patterns in terms of average GE and LD values within each two-digit category. In particular, our 155 three-digit industries are grouped into 22 categories at the two-digit level. So by letting I_i denote the set of three-digit industries in each two-digit category, i , we can summarize the spatial pattern for each category $i = 1, \dots, 22$ by its *average global extent*, $\overline{GE}_i \equiv \frac{1}{|I_i|} \sum_{j \in I_i} GE_j$, and *average local density*, $\overline{LD}_i \equiv \frac{1}{|I_i|} \sum_{j \in I_i} LD_j$. While these average spatial patterns, $(\overline{GE}_i, \overline{LD}_i : i = 1, \dots, 22)$, for categories are far fewer in number than our original 155 spatial patterns for industries, it is seen in Figure 18 that they continue to

⁴⁷See, e.g., Combes and Lafourcade [8] and Konishi et al. [32] for attempts to estimate transport costs within a given country.

⁴⁸There have been very few empirical studies relating transport costs to spatial patterns of industries in a regional context (although there have been some efforts to quantify the effects of transport costs on international trade, as for example in Head and Mayer [23] and Limão and Venables [35]).

⁴⁹Since the origins and destinations of the shipments can be identified in terms of municipalities, the shipment distances are computed as the shortest-route distances along the road network as in Section 2.3.

be uncorrelated in a manner similar to Figure 6 [with $\rho(\overline{GE}, \overline{LD}) = 0.10$ and a p -value of 0.64 for a two-sided test of ρ significance].⁵⁰ So these two average indices continue to provide distinct spatial information.

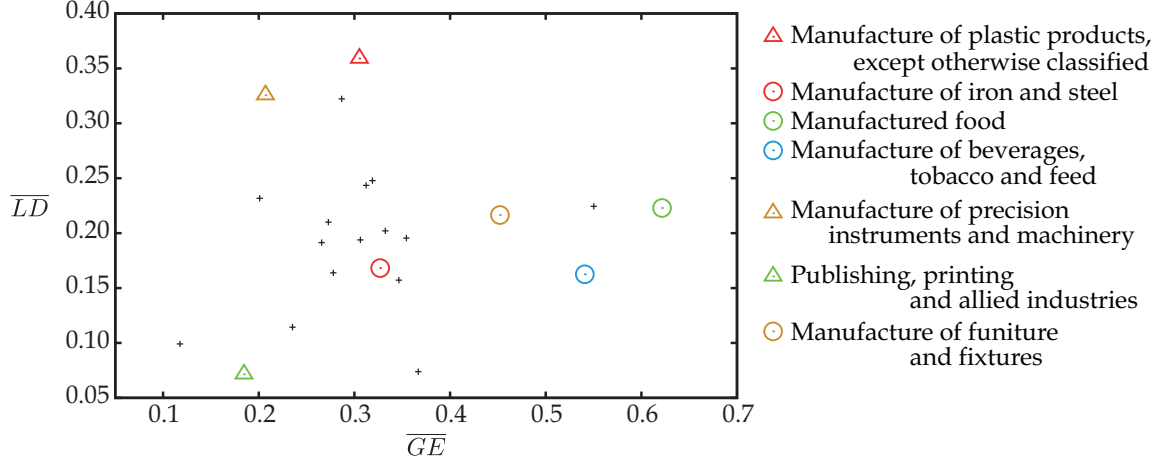


Figure 18: Relationship between average values of GE and LD at two-digit industrial categories

Given these average indices, if we now let the average shipment distance for establishments in each industry $j = 1, \dots, 155$, be denoted by SD_j , then our objective is to relate these indices to the *average shipment distance*, $\overline{SD}_i \equiv \frac{1}{|I_i|} \sum_{j \in I_i} SD_j$, for each two-digit category i by employing a multiple regression approach paralleling expression (20) above. The results of this regression are shown below:

$$\log \overline{SD}_i = 5.812 - 0.491 \log \overline{GE}_i + 0.529 \log \overline{LD}_i, \quad \text{adj. } R^2 = 0.496, \quad (21)$$

(23.86) (-3.40) (4.11)

where the numbers in the parentheses are t -values, and all the estimated coefficients are significant at the 1% level. In view of the low correlation between the two explanatory variables, $\overline{GE}$ and $\overline{LD}$, these relations are well approximated by their corresponding simple regressions, which can be shown graphically as in panels (a) and (b) of Figure 19, respectively.⁵¹

⁵⁰The three-digit industries indicated in Figure 6 belong to the two-digit categories indicated by using the same symbols, except that “livestock products” and “seafood products” both belong to “manufactured food” category.

⁵¹The simple-regression coefficients are naturally somewhat different, and in this case -0.334 for $\log \overline{GE}$ in panel (a) and 0.412 for $\log \overline{LD}$ in panel (b).

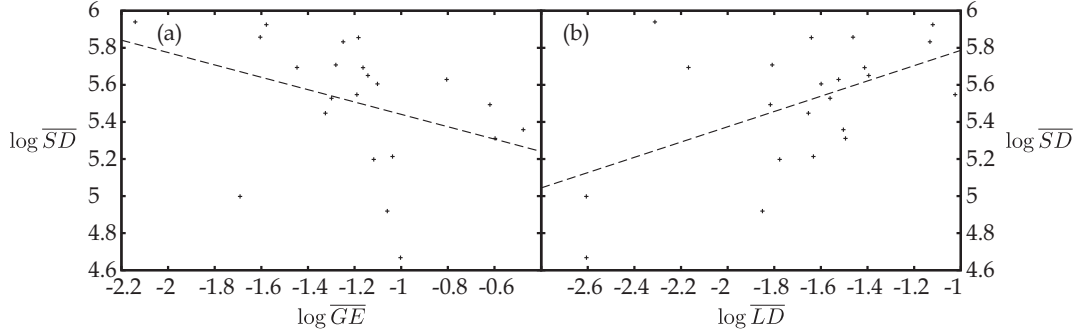


Figure 19: Average shipment distance and the spatial scales of agglomeration

Here we see that larger values of average global extent, $\overline{GE}$, correspond to lower average shipment distances, $\overline{SD}$ – which is consistent with the *global dispersion force* prediction of NEG above. Similarly, larger values of average local density, $\overline{LD}$, correspond to larger average shipment distances, $\overline{SD}$ – which is consistent with the *local dispersion force* prediction of NEG. Thus, this regression provides perhaps the first empirical support for these theoretical predictions of NEG. However, there are at least two caveats in interpreting eq.(21). One is that these relations involve only average values across rather broad two-digit industry categories. Second, even if this same relation were to hold for individual industries, the average shipment distance, SD , for each industry would only be associated with those values of GE and LD realized in equilibrium. So no causal inferences could be drawn from these relations. Thus, the relative magnitudes of these estimated coefficients must be interpreted with caution.

7 Concluding Remarks

In this paper we have applied the cluster-detection procedure developed by Mori and Smith [45] to study the agglomeration patterns of manufacturing industries in Japan. In particular, we have proposed a simple classification of pattern types based on a pair of quantitative measures, global extent (GE) and local density (LD), for distinguishing both the scale and degree of industrial agglomeration derived from the cluster schemes. But the ultimate utility of this approach will of course depend on how it can be applied in practical situations.

As alluded to in the Introduction, these measures can already help to sharpen certain concepts in the literature. For example, the differences between spatial

dispersion of manufacturing at high versus low levels of transport costs, as derived in general NEG models, can be characterized in terms of these measures. In particular, the type of dispersion associated with high levels of transport costs (“first-phase” dispersion) can in principle be quantified empirically in terms of large GE values and small LD values.⁵² In contrast, dispersion patterns associated with low levels of transport costs (“second-phase” dispersion) might be quantified in terms of small GE values and large LD values. Hence, such differences between dispersion patterns might be quantified in terms of directed distances within a given GE - LD space. In fact, given appropriate historical data on industrial location patterns at various stages of transportation technology, one might even be able to test the significance of such differences.

As another illustration, the Japanese Industrial Belt discussed in Section 4.3.3 can be considered as an instance of the more general notion of a “megapolopolis,” first proposed by Gottman [21] to describe the continuum of cities along the US Atlantic seaboard (stretching from Boston to Washington, DC, via New York). But to date, no formal methods have been developed for identifying such agglomeration structures statistically. In this light, the analysis of Section 4.3.3 shows that such structures can also be regarded as natural instances of “globally confined and locally dense” agglomeration patterns. Hence, the emergence of such large scale structures might in principle be characterized in term of urban agglomeration pattern shifts within an appropriate GE - LD space.

But it should also be emphasized that these two measures are by no means the only relevant properties of agglomeration patterns that can be quantified. Indeed, our present construction of such patterns in terms of cluster schemes provides a potentially rich spatial data set for studying a wide range of problems. Along these lines, it is appropriate to mention three possible research directions involving, respectively, the spacing of clusters within industries and the coordination of clusters between industries.

⁵²Here it should be noted that since firms have no “area” in such continuous models, our present notion of local density is somewhat ambiguous. But given fixed employment levels for industries, the essence of this type of dispersion is that individual clusters become smaller and more scattered throughout the spatial continuum. So local “employment” density decreases under this type of dispersion.

7.1 Agglomeration Spacing within Industries

Within the NEG, a number of models have been developed to explain the spacing between individual clusters for a given industry (e.g., Krugman [33], Fujita and Krugman [15], Fujita and Mori [18], Fujita et al. [17, Ch.6], Tabuchi et al. [55], Ikeda et al. [27], Akamatsu et al. [3]). From the view point of general equilibrium theory, these models predict whether an agglomeration of industrial firms will be viable at a given location, depending on how other clusters of the same industry (as well as population) are distributed over the location space. In these models, industrial agglomeration is typically induced by demand externalities arising from the interactions between product differentiation, plant-level scale economies and transport costs. In particular, Fujita and Krugman [15] have shown that each agglomeration casts a so-called *agglomeration shadow* in which firms have no incentive to relocate from the existing clusters, since within this “shadow” firms are too close to existing clusters (i.e., competitors) to realize sufficient local monopoly advantages. Hence the presence of such shadows serves to limit the number of viable clusters within each industry. Note also that since the level of internal competition differs between industries (depending on their degree of product differentiation and transport costs), the size of agglomeration shadows should also be industry specific.

But while there has been empirical work to study the spacing between urban centers (e.g., Marshall [39, Ch.7], Ioannides and Overman [28], Hsu et al. [26]), to our knowledge there have been no systematic efforts to study the spacing between industrial clusters – and in particular, no efforts to identify the presence of actual agglomeration shadows. However, it should be clear that our present approach to cluster identification offers a promising method for doing so. In particular, since our cluster-detection procedure enables one to identify individual clusters for each industry, it is a simple matter to construct explicit measures of the spacing between them. For example, one natural measure of spacing between clusters in our present framework would be the shortest path distance between their closest basic regions. Agglomeration spacing for cluster schemes as a whole might then be summarized by the mean nearest-neighbor distance between their constituent clusters. To test whether such spacing is larger (or more uniform) than would be expected by chance alone, one could in principle generate appropriate random versions of cluster schemes to serve as counterfac-

tuals. While such random collections of disjoint sets are of course more difficult to construct than random point patterns, initial investigations with a variety of rejection-sampling techniques suggest that this is certainly possible. Hence by constructing mean nearest-neighbor distances for each random version sampled, one could use this sampling distribution to test a variety of agglomeration spacing properties in terms of cluster schemes. Such spacing analyses will be reported in subsequent work.⁵³

7.2 Agglomeration Coordination between Industries

Within the context of Christaller’s [7] celebrated theory of *Central Places*, a topic of major interest has long been the spatial coordination of locations across industries. In particular, the “Hierarchy Principle” underlying this theory asserts that the set of industries found in smaller metro areas is always a subset of those found in larger metro areas.⁵⁴ Theoretical efforts to explain this phenomenon have focused mainly on the role of demand externalities in determining industrial locations (see Quinzii and Thisse [48], Fujita et al. [16], Tabuchi and Thisse [54] and Hsu [25]).⁵⁵ In particular, the types of demand externalities which induce industrial agglomeration are often shared by many different industries, so that their spatial markets overlap. In such cases, it is natural for these industries to co-locate. Moreover, in terms of market sizes, it is also natural for clusters in more concentrated industries (with larger markets) to coincide with those of less concentrated industries (with smaller markets), thus leading to the type of synchronization predicted by the Hierarchy Principle.

But while these theoretical arguments are quite plausible, there has been surprisingly little work done to actually test the empirical validity of the Hierar-

⁵³Here it is of interest to note that initial investigations of such spacing properties suggest that further restrictions need to be imposed. In particular, for those industries with small *e*-containments, it is clear that random versions located throughout all of Japan will necessarily tend to exhibit larger mean spacing for rather spurious reasons. One possibility here is to preserve the *e*-containment of each industry, and to restrict random versions to these *e*-containments. This should provide more meaningful tests of the presence of agglomeration shadows in which the overall spatial scale of each industry is preserved.

⁵⁴Obviously, this principle implicitly assumes a certain degree of industry aggregation, since it could not hold if industries are fully disaggregated, i.e., where each industry consists of one establishment.

⁵⁵There were earlier attempts by, e.g., Christaller [7], Lösch[36], Beckmann [4] and Eaton and Lipsey [12]. But, all lacked formal microeconomic foundations leading to the Hierarchy Principle.

chy Principle itself. One approach proposed by Mori and Smith [44] focuses on the hierarchical industrial structure of cities implied by this principle. In particular, the present cluster-detection procedure was used to identify those cities containing establishments that are actually part of clusters for the industry. By restricting the classical Hierarchy Principle to these “cluster-based choice cities” for each industry, it was shown that this *cluster-based Hierarchy Principle* holds even more strongly than the classical version for our Japanese data.⁵⁶

However, the detailed spatial structure of cluster schemes also permits more direct comparisons of spatial coordination between individual industries. In particular, by associating larger market sizes with smaller numbers of clusters for an industry,⁵⁷ one may ask whether industries with larger market sizes do in fact tend to coordinate their spatial locations with industries having smaller market sizes. More specifically one may ask whether their cluster schemes are closer to those of industries with smaller market sizes than would be expected by chance alone. By again measuring “closeness” in terms of (shortest-path) nearest-neighbor distances, one could test this hypothesis in a manner similar to Section 7.1 above. Note that this test can also be interpreted as the test of *co-localization* among different industries, which could in principle provide an alternative approach to those of Ellison et al. [14] and Duranton and Overman [11]. Such investigations will be reported in subsequent work.

7.3 Refining Essential Containments

Because our present spatial measures, GE and LD , are defined solely in terms of area, there of course remains a certain degree of ambiguity regarding spatial *patterns* of agglomeration. This is particularly evident when analyzing the nature of “local dispersion” within e -containments, as illustrated by the two e -containment patterns in Figure 20. While both GE and LD are identical for each pattern, it is evident that “local dispersion” is far more ubiquitous in panel (b) than panel (a). In fact panel (a) might be better described as two major agglomerations of clusters concentrated at opposite ends of this e -containment. So it is important to ask how our present set of measures might be extended to capture

⁵⁶Using the same method, Akamatsu et al. [1] have established the same result for the US for the 4-, 5- and 6-digit manufacturing industries in North American Industry Classification System in 2007.

⁵⁷In fact this relationship underlies the results in the theoretical papers above.

such distinctions.

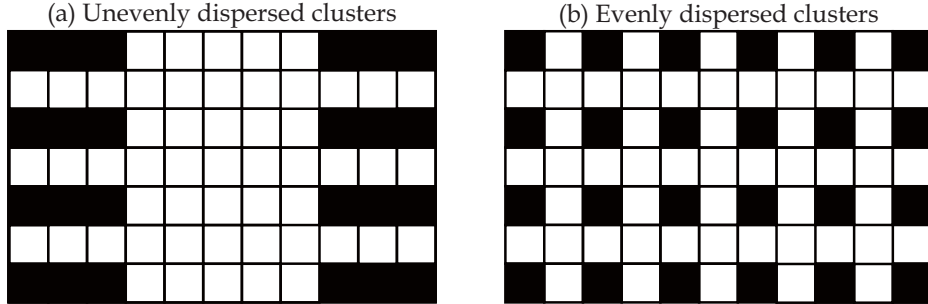


Figure 20: Two cluster patterns with the same GE and LD

One possibility is suggested by our procedure of building e -containments, where clusters are added until some appropriate BIC threshold is achieved. Having done so, one may continue to combine e -clusters in a pairwise manner that “least detracts” from this threshold value, and then analyze the resulting sequence of decreasing values. For example one would expect that an application of this procedure to panel (a) of Figure 20 would first combine clusters on either end of the e -containment until at some point these two agglomerations of clusters would be joined. At this point, a much larger drop in BIC might be expected, reflecting the “loss of fit” resulting when these two agglomerations are combined. In contrast, panel (b) should be expected to yield a more even sequence of decreases, with no major drops. So by studying these respective patterns of decrease, one might be able to detect major changes in spatial patterns that represent important intermediate levels of agglomeration structure.

Initial experimentations with this type of *decrement analysis* on empirical examples similar in nature to Figure 20 suggest that such intermediate structures can indeed be identified. However, these experiments also show that decrement sequences are highly sensitive to the order in which clusters are joined. So more robust types of procedures (such as resampling and model averaging) are evidently needed to overcome such path-dependencies. Further investigations along these lines will be reported in subsequent work.

References

- [1] Akamatsu, Takashi, Tomoya Mori and Yuki Takayama. 2014. “Spatial coordination among industries and the common power law for city size distribu-

- tions.” In progress. Institute of Economic Research, Kyoto University.
- [2] ——— and Yuki Takayama. 2014. “Do polycentric patterns emerge in NEG models?” Unpublished manuscript. Graduate School of Information Sciences, Tohoku University.
 - [3] ———, Yuki Takayama and Kiyohiro Ikeda. 2012. “Spatial discounting, fourier, and racetrack economy: A recipe for the analysis of spatial agglomeration models.” *Journal of Economic Dynamics & Control* 36: 1729-1759.
 - [4] Beckmann, Martin J. 1958. “City hierarchies and the distribution of city size,” *Economic Development and Cultural Change* 6: 243-248.
 - [5] Behrens, Kristian. 2007. “On the location and lock-in of cities: Geography vs transport technology.” *Journal of Urban Economics* 37(1): 22-45.
 - [6] Brühlhart, Marius and Rolf Traeger. 2005. “An account of geographic concentration patterns in Europe.” *Regional Science and Urban Economics* 35(6): 597-624.
 - [7] Christaller, William. 1933. *Die Zentralen Orte in Suddeutschland* (Jena, Germany: Gustav Fischer). English translation by Baskin, Carlisle W. (1966), *Central Places in Southern Germany* (London: Prentice Hall).
 - [8] Combes, Pierre-Phillippe and Miren Lafourcade. 2005. “Transport costs: measures, determinants, and regional policy implications for France.” *Journal of Economic Geography* 5(3): 319-349.
 - [9] Combes, Pierre-Phillippe, Mayer, Thierry, Thisse, Jaques-François. 2008. *Economic Geography: The Integration of Regions and Nations* (Princeton, NJ: Princeton University Press).
 - [10] Duranton, Gilles and Henry G. Overman. 2005. “Testing for localization using micro-geographic data.” *Review of Economic Studies* 72: 1077-1106.
 - [11] ——— and ———. 2008. “Exploring the detailed location patterns of U.K. manufacturing industries using microgeographic data.” *Journal of Regional Science* 48(1): 213-243.
 - [12] Eaton, B. Curtis, Richard G. Lipsey. 1982. “An economic theory of central places.” *Economic Journal* 92(365): 56-72.
 - [13] Ellison, Glenn and Edward L. Glaeser. 1997. “Geographic concentration in US manufacturing industries: A dartboard approach,” *Journal of Political Economy* 105(5): 889-927.

- [14] ———, ——— and William R. Kerr. 2010. "What causes industry agglomeration? Evidence from coagglomeration patterns." *American Economic Review* 100(3): 1195-1213.
- [15] Fujita, Masahisa and Paul R. Krugman. 1995. "When is the economy monocentric?: von Thünen and Chamberlin unified." *Regional Science and Urban Economics* 25: 505-528.
- [16] ———, Paul R. Krugman and Tomoya Mori. 1999. "On the evolution of hierarchical urban systems," *European Economic Review* 43: 209-251.
- [17] ———, ——— and Anthony J. Venables. 1999. *The Spatial Economy: Cities, Regions and International Trade* (Cambridge, MA: The MIT Press).
- [18] ——— and Tomoya Mori. 1997. "Structural stability and evolution of urban systems." *Regional Science and Urban Economics* 27: 399-442.
- [19] ———, ———. 2005. "Frontiers of the new economic geography." *Papers in Regional Science* 84(3): 307-405.
- [20] ———, ———. 2005. "Transport development and the evolution of economic geography." *Portuguese Economic Journal* 4(2): 129-156.
- [21] Gottman, Jean. 1961. *Megalopolis: The Urbanized Northern Seaboard of the United States* (New York: The Twentieth Century Fund).
- [22] Helpman, Elhanan. 1998. "The size of regions." In: *Topics in Public Economics: Theoretical and Applied Analysis* (Cambridge: Cambridge University Press): 33-54.
- [23] Head, Keith and Thierry Mayer. 2004. "The empirics of agglomeration and trade" in J.V. Henderson and J.-F. Thisse (eds.) *Handbook of Regional and Urban Economics* volume 4, Amsterdam: North-Holland, 2609-2669.
- [24] Hokkaido-chizu, Co. Ltd. 2002. *GIS Map for Road*.
- [25] Hsu, Wen-Tai. 2012. "Central place theory and the city size distribution." *Economic Journal* 122: 903-932.
- [26] ———, Tomoya Mori and Tony E. Smith. 2014. "Spatial patterns and size distributions of cities." Discussion paper No.882. Institute of Economic Research, Kyoto University.
- [27] Ikeda, Kiyohiro, Takashi Akamatsu, and Tatsuhito Kono. 2012. "Spatial period doubling agglomeration of a core-periphery model with a system of cities." *Journal of Economic Dynamics & Control* 36: 743-778.

- [28] Ioannides, Yannis M. and Henry G. Overman. 2004. "Spatial evolution of the US urban system." *Journal of Economic Geography* 4(2): 131-156.
- [29] Japan Statistics Bureau. 2000. *Population Census of Japan* (in Japanese).
- [30] ———. 2001. *Establishments and Enterprise Census of Japan* (in Japanese).
- [31] Kerr, William R. and Scott Duke Kominers. 2014. "Agglomerative forces and cluster shapes." *Review of Economics and Statistics*, forthcoming.
- [32] Konishi, Yoko, Se-il Mun, Yoshihiko Nishiyama, Ji Eun Sung. 2014. "Measuring the value of time in freight transportation." Discussion paper 14-E-004, Research Institute of Economy, Trade and Industry.
- [33] Krugman, Paul R. 1993. "On the number and location of cities." *European Economic Review* 37: 293-298.
- [34] Kullback, Solomon and Richard A. Leibler. 1951. "On information and sufficiency." *Annals of Mathematical Statistics* 22: 79-86.
- [35] Limão, Nuno and Anthony J. Venables. 2001. "Infrastructure, geographical disadvantage, transport costs and trade." *World Bank Economic Review* 15: 451-479.
- [36] Lösch, August. 1940. *The Economics of Location* (Jena, Germany: Fischer). English translation by Stolper, Wolfgang F. and Woglom, William H. (1954) (New Haven, CT: Yale University Press).
- [37] Marcon, Eric and Florence Puech. 2010. "Measures of the geographic concentration of industries: improving distance-based methods." *Journal of Economic Geography* 10(5): 745-762.
- [38] ———, Stéphane Traissac, Florence Puech and Gabriel Lang. 2013. "dbmss: Distance-based measures of spatial structures." <http://cran.r-project.org/web/packages/dbmss/index.html>.
- [39] Marshall, John U. 1989. *The Structure of Urban Systems* (Toronto: University of Toronto Press). polis formatin: the maturing of city systems," *Journal of Urban Economics* 42: 133-157.
- [40] Ministry of Land, Infrastructure, Transport and Tourism. 2000. *Net Freight Flow Census*.
- [41] Mori, Tomoya. 1997. "A modeling of megalopolis formation: The maturing of city systems." *Journal of Urban Economics* 42: 133-157.
- [42] ———, Koji Nishikimi and Tony E. Smith. 2005. "A divergence statistic for industrial localization." *Review of Economics and Statistics* 87(4): 635-651.

- [43] ———, Koji Nishikimi and Tony E. Smith. 2008. “The number-average size rule: A new empirical relationship between industrial location and city size.” *Journal of Regional Science* 48: 165-211.
- [44] ——— and Tony E. Smith. 2011. “An industrial agglomeration approach to central place and city size regularities.” *Journal of Regional Science* 51(4): 694-731.
- [45] ———, ———. 2014. “A probabilistic modeling approach to the detection of industrial agglomeration.” *Journal of Economic Geography* 14(3): 547-588.
- [46] Murata, Yasusada and Jaques-François Thisse. 2005. “A simple model of economic geography à la Helpman-Tabuchi.” *Journal of Urban Economics* 58(1): 137-155.
- [47] Özak, Ömer. 2010. “The voyage of homo-œconomicus: Some economic measure of distance.” Unpublished manuscript. Department of Economics, Brown University.
- [48] Quinzii, Martine and Jacques-François Thisse. 1990. “On the optimality of central places.” *Econometrica* 58: 1101-1119.
- [49] Rosenthal, Stuart S. and William C. Strangelliam. 2004. “Evidence on the nature and sources of agglomeration economies.” In Henderson, J. Vernon and Thisse, Jacques-François (eds.), *Handbook of Regional and Urban Economics*, Vol.4 (Amsterdam: North-Holland), Ch.49.
- [50] Silverman, Bernard W. 1986. *Density Estimation for Statistics and Data Analysis*. New York: Chapman and Hall.
- [51] Statistical Information Institute for Consulting and Analysis. 2002. *Toukei de Miru Shi-Ku-Cho-Son no Sugata* (in Japanese).
- [52] ———. 2003. *Toukei de Miru Shi-Ku-Cho-Son no Sugata* (in Japanese).
- [53] Tabuchi, Takatoshi. 1998. “Urban agglomeration and dispersion: A synthesis of Alonso and Krugman.” *Journal of Urban Economics* 44(3): 333-351.
- [54] ——— and Jacques-François Thisse. 2006. “Regional specialization, urban hierarchy, and commuting costs.” *International Economic Review* 47: 1295-1317.
- [55] ———, ——— and Dao-Zhi Zeng. 2005. “On the number and size of cities.” *Journal of Economic Geography* 5(4): 423-448.