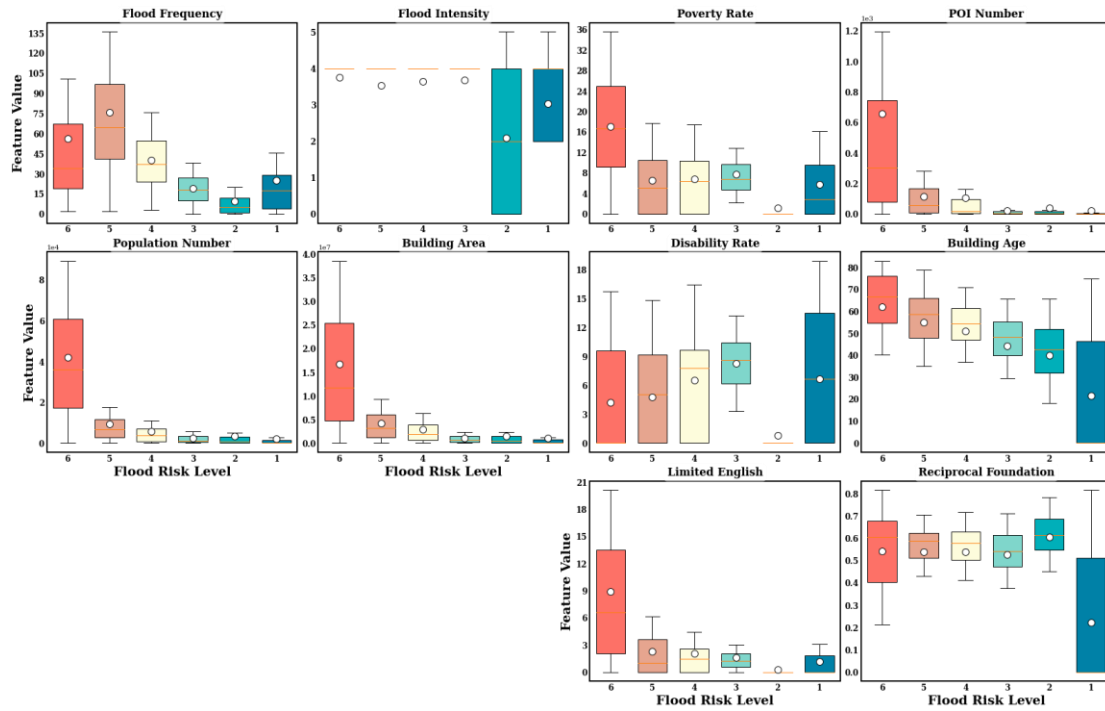
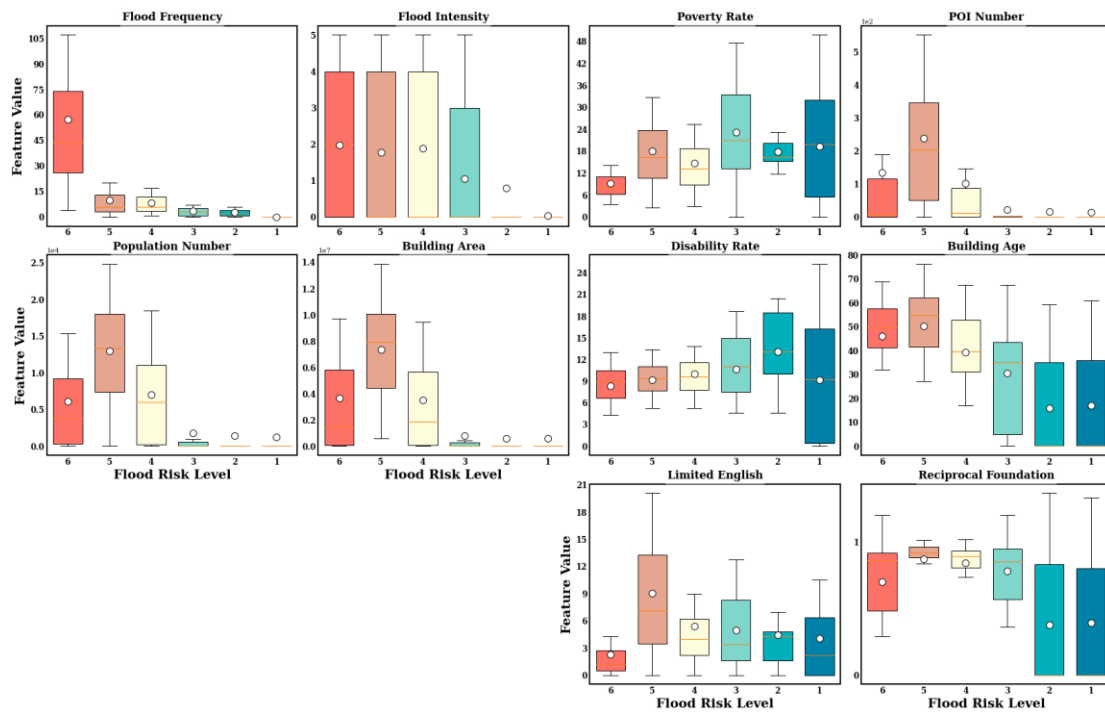
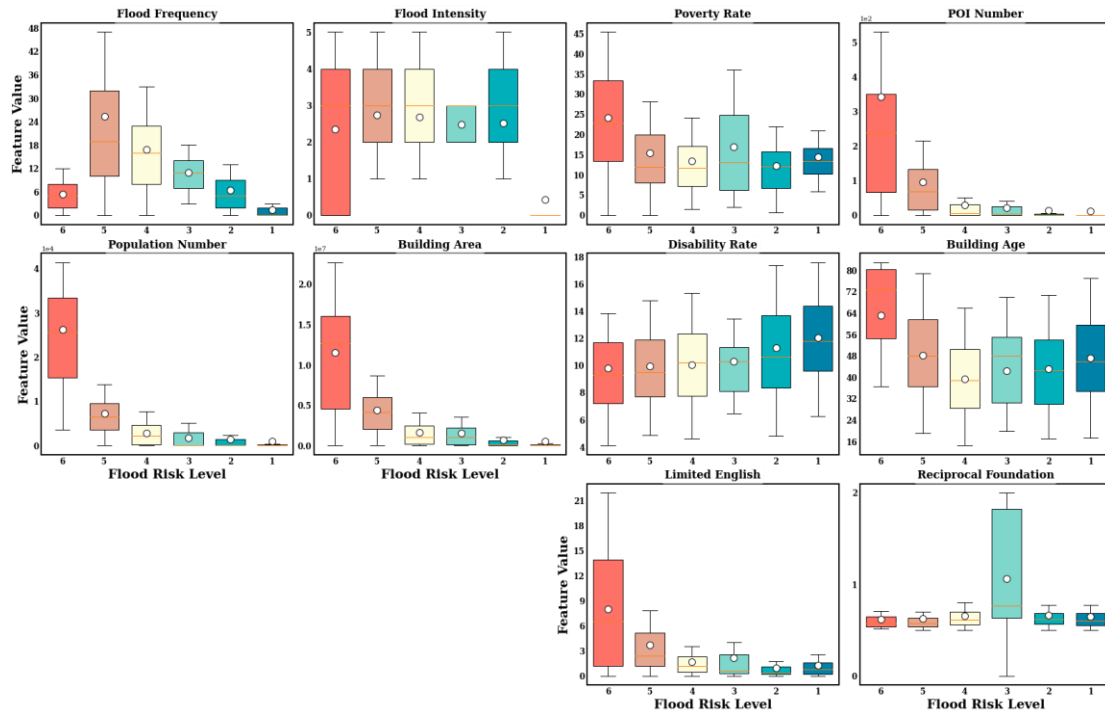




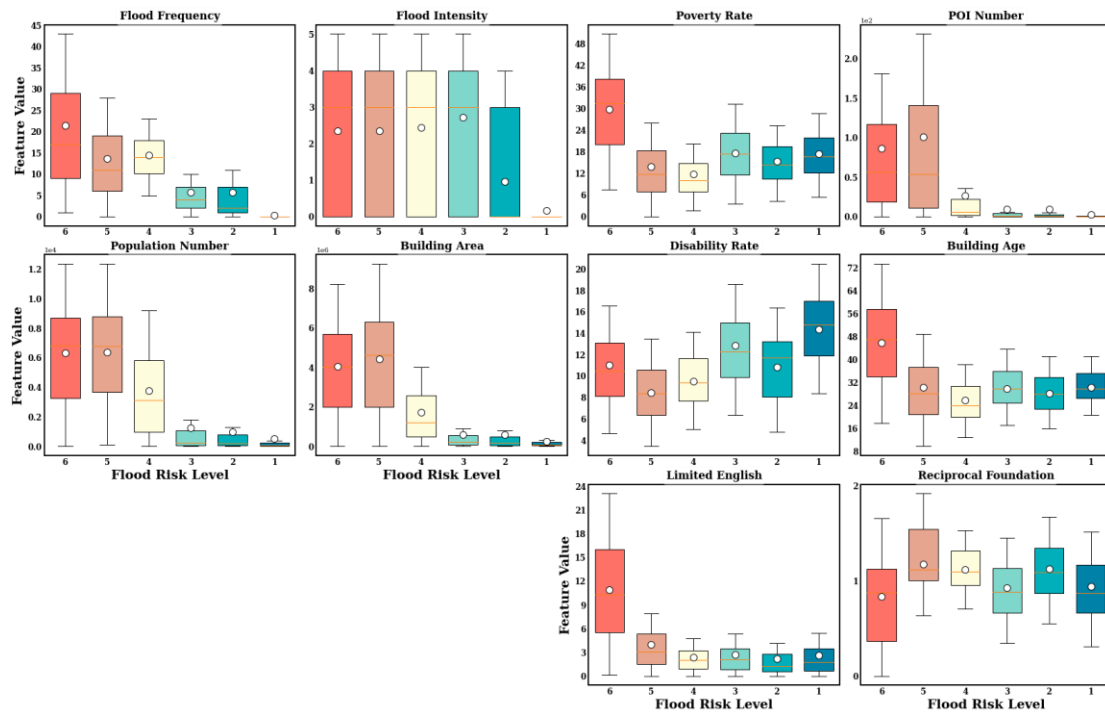
Supplementary Figure 1 Characteristics of clusters in Greater New York



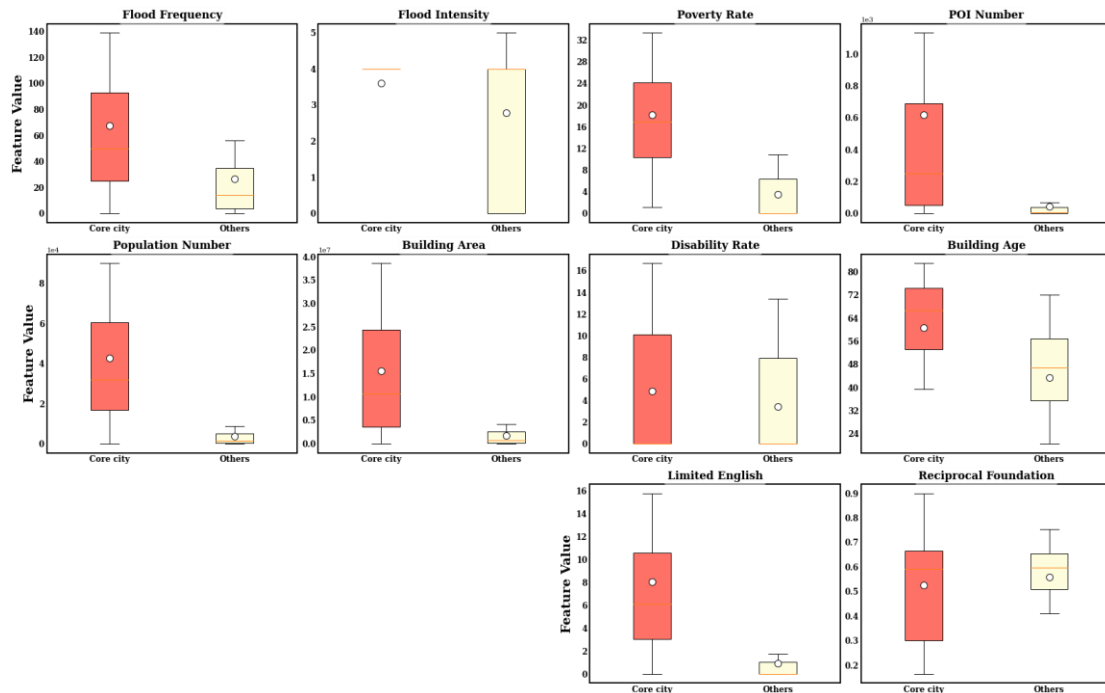
Supplementary Figure 2 Characteristics of clusters in Greater Los Angeles



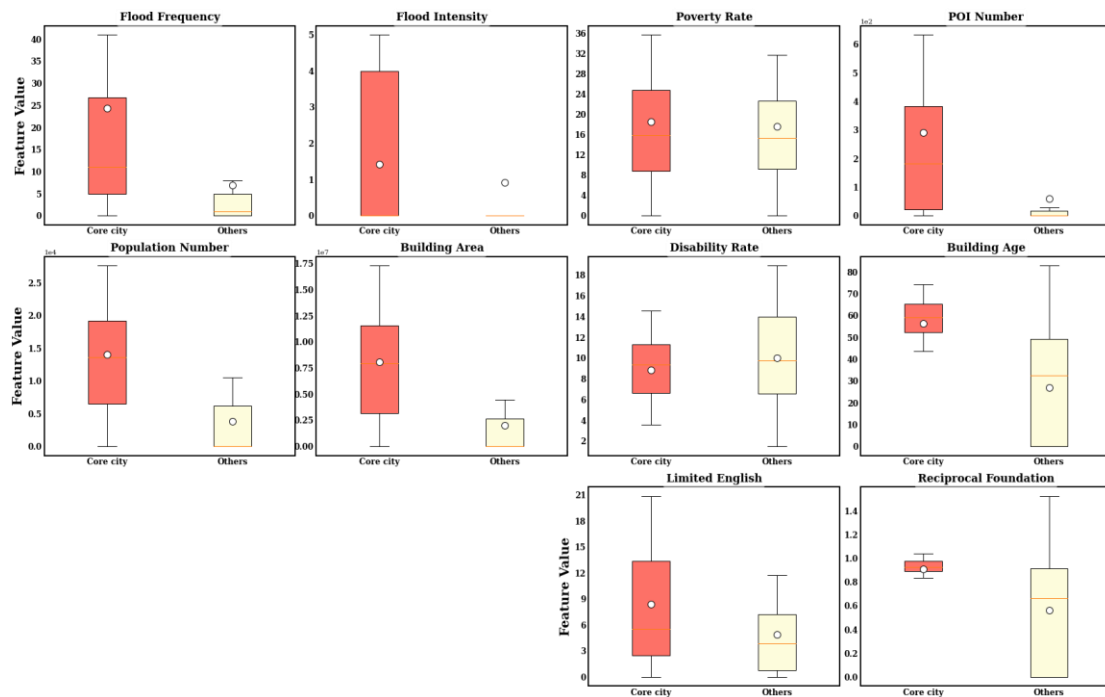
Supplementary Figure 3 Characteristics of clusters in Greater Chicago



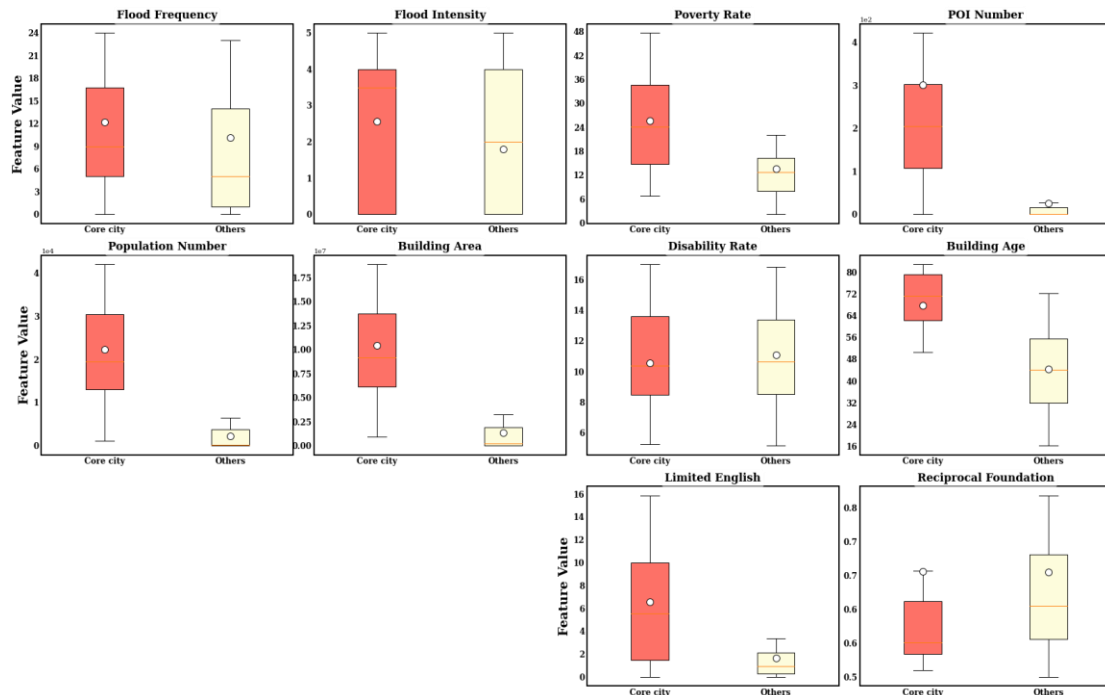
Supplementary Figure 4 Characteristics of clusters in Greater Dallas



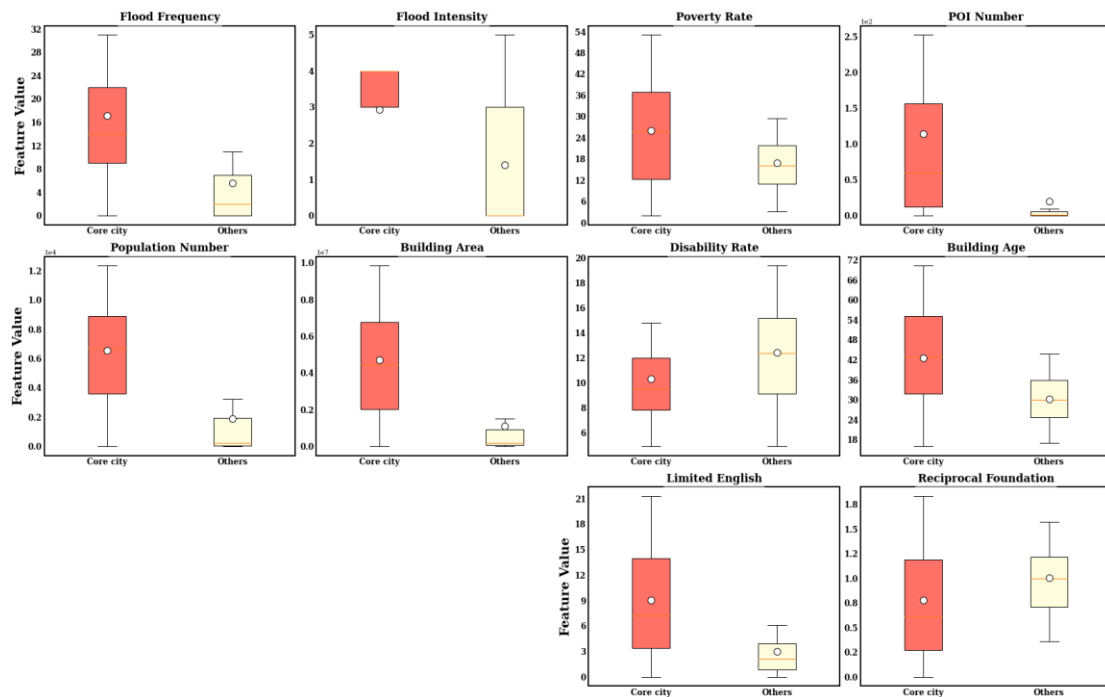
Supplementary Figure 5 Characteristics of core city and other areas in Greater New York



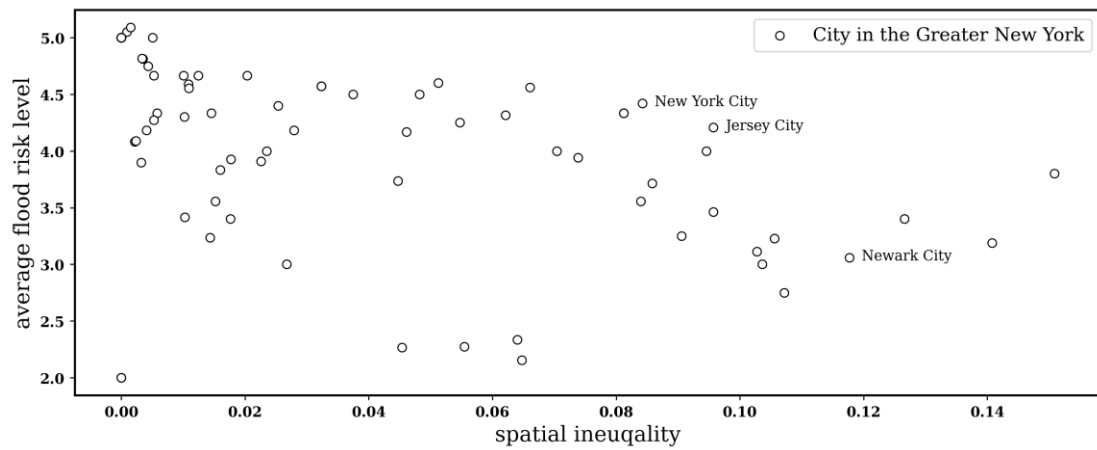
Supplementary Figure 6 Characteristics of core city and other areas in Greater Los Angeles



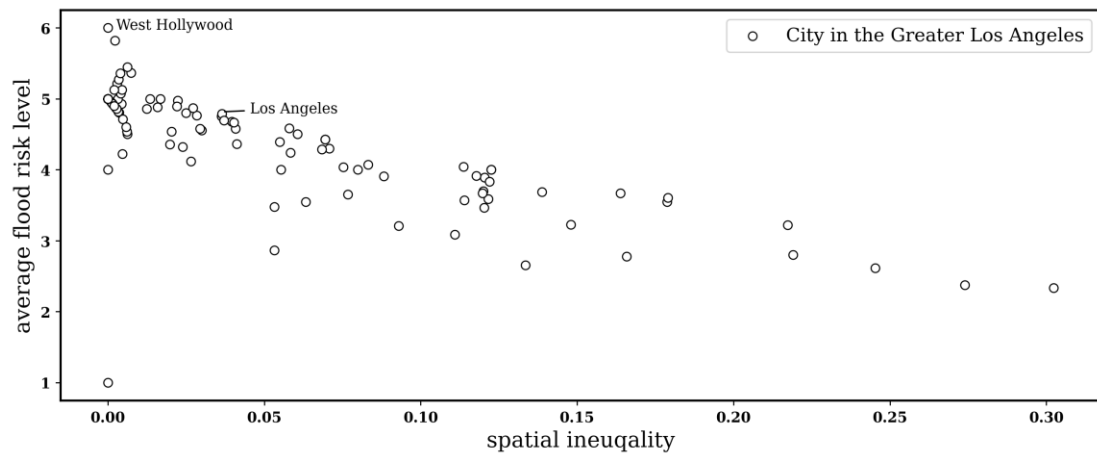
Supplementary Figure 7 Characteristics of core city and other areas in Greater Chicago



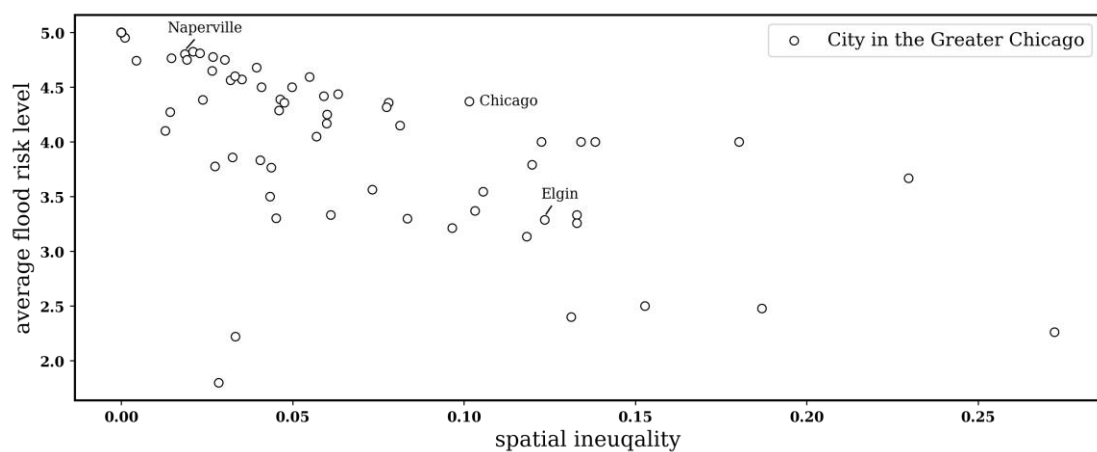
Supplementary Figure 8 Characteristics of core city and other areas in Greater Dallas



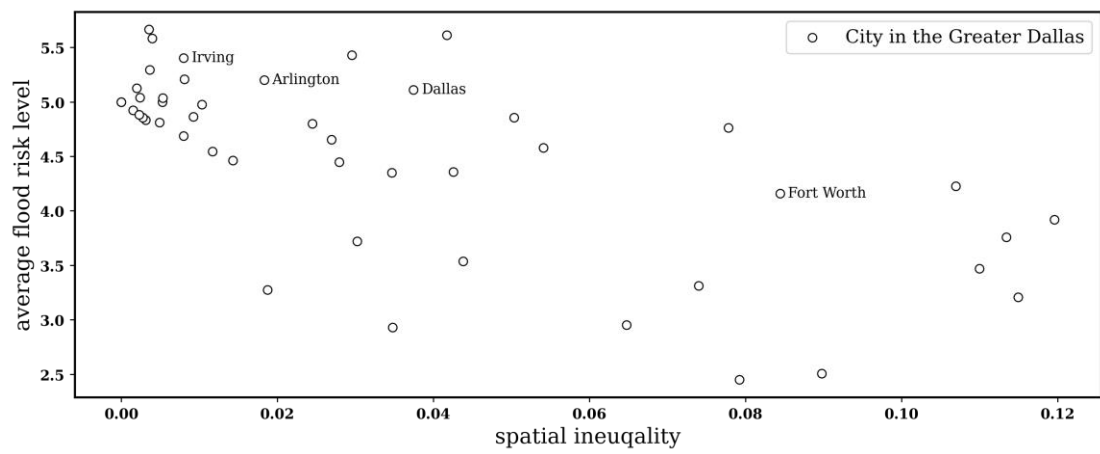
Supplementary Figure 9 Relations between average flood risk and spatial inequity of cities within Greater New York



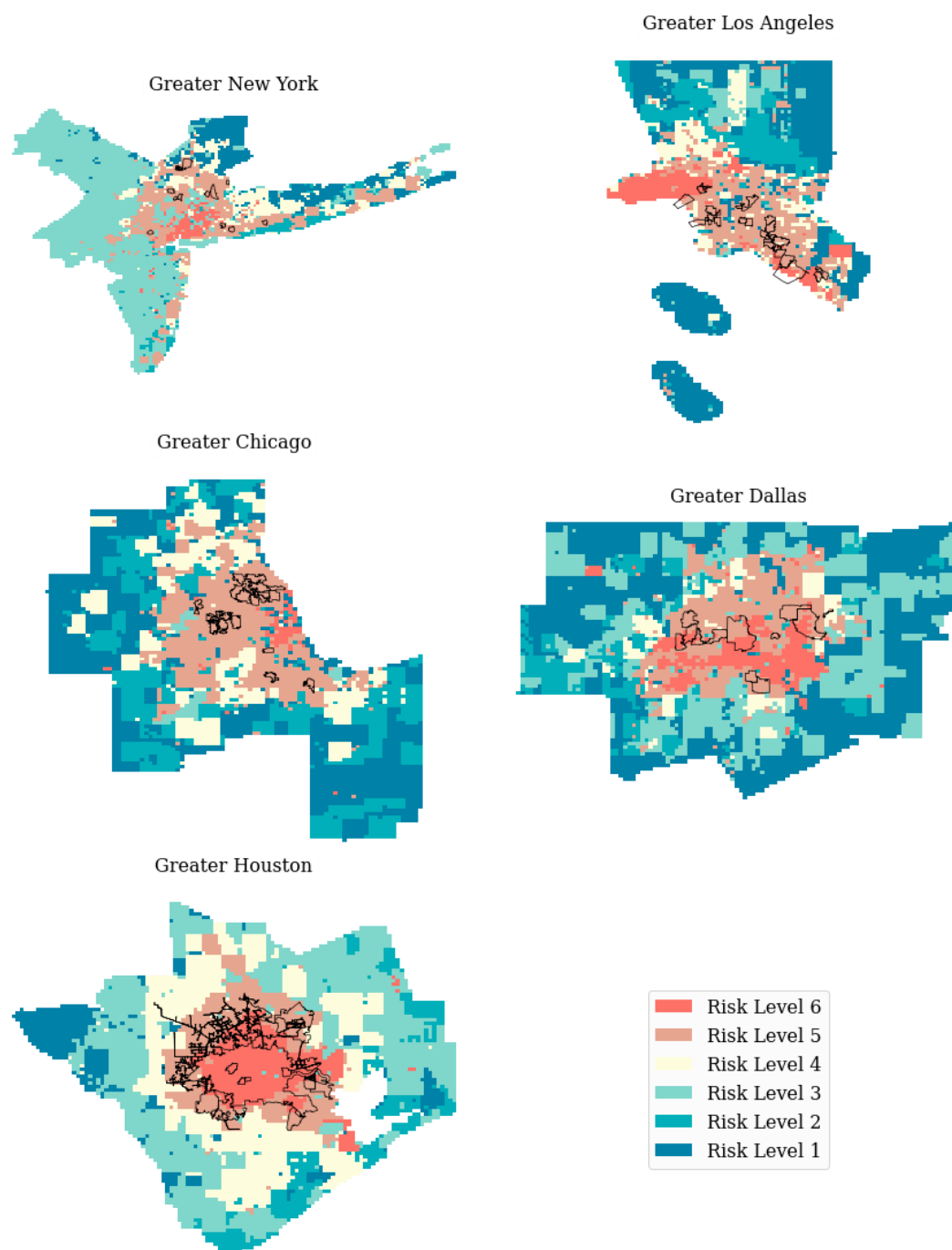
Supplementary Figure 10 Relations between average flood risk and spatial inequity of cities within Greater Los Angeles



Supplementary Figure 11 Relations between average flood risk and spatial inequity of cities within Greater Chicago



Supplementary Figure 12 Relations between average flood risk and spatial inequality of cities within Greater Dallas



Supplementary Figure 13 Cities showing large average flood risk levels and low spatial inequality. These cities whose population numbers are larger than 100,000 are listed in Table S2.

Supplementary Table 1 The relation between flood risk level and entered content
 (Adapted from <https://www.fema.gov/openfema-data-page/fima-nfip-redacted-claims-v1>)

Level	Entered content
1	Basement/Enclosure/Crawlspace/Subgrade Crawlspace only
2	Basement/Enclosure/Crawlspace/Subgrade Crawlspace and above
3	Lowest floor only above ground level (No basement/enclosure/crawlspace/subgrade crawlspace)
4	Lowest floor above ground level and higher floors (No basement/enclosure/crawlspace/subgrade crawlspace)
5	Above ground level more than one full floor

Supplementary Table 2 Cities with a population larger than 100,000 and show high flood risk levels while low spatial inequality

City name	Population number (in the unit of thousand)
Houston city, Texas	228.8
Irving city, Texas	25.42
Frisco city, Texas	21.07
Huntington Beach city, California	19.67
Garden Grove city, California	17.05
East Los Angeles CDP, California	12.05
Lewisville city, Texas	11.29
Sugar Land city, Texas	10.94
Allen city, Texas	10.69
Norwalk city, California	10.04

Supplementary Table 3 Summarize of flood risk reduction strategies based on features of each risk level

Flood risk aspects		Strategies
High flood hazard	/	Engineering solutions to control the flow of stormwater runoff
High flood exposure	/	Land use planning to restrict development in flood-prone areas, or the government's buyout of flood-prone areas
High social vulnerability	High poverty rate	Financial aid to help them buy flood insurance, reinforce their homes, rebuild flood-damaged homes, migrate to lower flood risk areas
	High limited English-speaking	Provide flood-related information in different languages
	High disability rate	Policy to ensure accessibility of shelters and financial aid to rebuild homes after floods
High physical vulnerability	Large POI number	Engineering solutions to protect POIs providing essential service
	Large building age	Engineering solutions to reinforce the building
	Low foundation height	Engineering solutions to raise foundation heights to prevent floods from entering the building

Supplementary Table 4 Notation for core parameters used in the urban flood risk rating model

<i>Notation</i>	<i>Definition</i>
$\tilde{\mathbf{A}} \in R^{m*m}$	The learned adjacency matrix of the spatial flood dependence network with unknown parameters, where m is the number of grid cells in each MSA
$\tilde{\mathbf{A}}^* \in R^{m*m}$	Adjacency matrix of spatial flood dependence network after post-processing
$\mathbf{A}^* \in R^{m*m}$	The final optimal adjacency matrix of spatial flood dependence network
$\overline{\mathbf{A}}^* \in R^{m*m}$	$\overline{\mathbf{A}}^* = \mathbf{A}^* + \mathbf{I}$, and $\mathbf{I}$ is the identity diagonal matrix of $\mathbf{A}^*$
$\mathbf{BF} \in R^{m*d_{BF}}$	Binary flood occurrence matrix, where d_{BF} is the number of features
$\mathbf{FR} \in R^{m*10}$	Flood risk-related feature matrix
$\mathbf{G}_K = (\mathbf{K}, \mathbf{BF})$	The constructed KNN graph view, where $\mathbf{K}$ is the KNN graph structure
$\mathbf{G}_L = (\tilde{\mathbf{A}}^*, \mathbf{BF})$	The constructed learned graph view
$\overline{\mathbf{BF}}, \overline{\mathbf{G}}_K, \overline{\mathbf{G}}_L$	The augmented binary flood occurrence matrix, KNN graph view and learned graph view
$\mathbf{Z}^{(L_{GE})}$	The output representation learned by GCN (with unknown parameters) and L_{GE} is the layer number of GCN
$\mathbf{H}^{(L_{GE})}$	The output representation learned by DNN (without unknown parameters)
L_G	The loss function of the contrastive learning
L_{res}	The reconstruction loss
L_{ta}	The KL-divergence loss between ancillary cluster assignment Q distribution and the target distribution P
L_{clu}	The KL-divergence loss between cluster assignment distribution Z and the target distribution P
L	The loss function of the clustering model

Supplementary Notes

S1 Binary Flood occurrence matrix construction

FEMA historical flood claim dataset¹ which contains the field of location and time of the flood claim is adopted to generate a binary flood occurrence matrix ($\mathbf{BF} \in R^{m \times d_{BF}}$). Following Brunner et al²'s work, flood claims in the same week are recognized as the same flood event. Flood claims from 8/31/1970 through 4/14/2022 are adopted and this time period is split into week intervals. The element (i, j) in $\mathbf{BF}$ describes whether there is a flood event in week j for grid cell i .

The location information of the flood claim includes census tract, Zip code, longitude, and latitude¹. However, the precision of the longitude and latitude information of the flood claim is reduced to 1 decimal for privacy consideration. To set the location of the flood claim as accurately as possible, the location polygon of each flood claim is generated using the information of census tract, Zip code, and (longitude, latitude). A square is first constructed using the (longitude, latitude) coordinates of the flood claim as the center with an edge length of 0.1 degrees. The precise location of the flood claim should be within this square. This square is further intersected with polygons of the census tract and Zip code of this flood claim. The intersection among them is treated as the location polygon. For grid cell i , if there is a location polygon with flood claim in week j intersecting with it, the value of (i, j) in the $\mathbf{BF}$ is 1, otherwise it is 0.

S2 Flood risk feature matrix construction

Flood hazard. Flood hazard is measured by flood frequency and flood intensity to describe the probability and intensity of possible flood events following the work of de Moel et al³., Mohanty et al⁴., and Wing et al⁵. The flood frequency is obtained by summing up each row of the $\mathbf{BF}$, which is the value of the total flood events of each grid cell during the study period. Flood intensity is approximated using the field of "Location of Contents" in the FEMA flood insurance claim dataset. This field indicates where the contents within the building are impacted by the flood. Its value ranges from level-1 basement only to level-5 more than one full floor¹ (See Table S.1 for detailed information.) The larger value of Location of Contents indicates the deeper flood inundation depth, which represents the more intensive the flood event is. The flood intensity of a grid cell is calculated as the average value of "Location of Content" for all location polygons intersecting with it.

Flood exposure. Flood exposure is measured by the population number and building area, which represent the valued social components (e.g., people and buildings) exposed to flood hazards following De Moel et al⁶. and Koks et al⁷.'s definition of flood exposure.

The population number is calculated using the 2020 Decennial Census data⁸ at the block level. This dataset provides the total population (P1_001N) in each block. The population density in each block is calculated by dividing the population number by the block area. The population density of each grid is determined by the average population density for all blocks that intersect with it. The total population number of each grid cell is calculated by multiplying its population density and area.

The building area is calculated using the National Structure Inventory (NSI) dataset⁹ which contains the field of location and estimated square footage of the structure. The highly detailed geographic coordinates of the structure are at the resolution of seven decimals. The

building area of a grid cell is determined by the total area of all structures inside it.

Social vulnerability. Poverty rate, disability rate, and limited English-speaking rate are selected to measure grid cell's social vulnerability. People in poverty have limited economic capacity to buy flood insurance to protect their houses, and repair their homes after floods^{10,11}. Those with disability have limited physical ability to escape from floods, and repair their damaged homes¹¹⁻¹³. Floods would also increase the mortality of these population groups^{11,12}. The limited English-speaking population has language barriers to finding needed flood information, which hampers their familiarity with local flood hazards, ability to make a timely response to floods, and availability of rescue resources after floods^{11,14}. These socially vulnerable populations have a limited ability to cope with and recover from floods^{15,16}, which makes the floods cause greater impacts on these populations and amplify flood risk on them.

These social vulnerability-related features are calculated using the CDC/ATSDR Social Vulnerability Index (SVI) dataset¹⁷. It provides the rate of poverty, disability, and limited English-speaking at the census tract level¹⁷. The population in poverty is calculated by multiplying the poverty rate and the total population number. The population density of poverty in each census tract is then calculated by dividing the population number in poverty and the census tract area. The population density of impoverished people for each grid cell is determined by the average population density of poverty across all census tracts intersecting with it. By multiplying the poverty density by the grid cell area, the population number of impoverished individuals for each grid cell can be computed. Finally, after the population number in poverty divided by the total population number, the poverty rate could be obtained. Disability rate and limited English-speaking rate are calculated similarly. This process can be expressed using Equation (S1) as follows:

$$R_{gc-j} = \frac{\sum_{i=1}^{N_{ct}} \frac{R_i \times P_{ct-i}}{S_{ct-i}}}{N_{ct}} \times S_{gc-j} \quad (S1)$$

where, R_i is the rate of poverty, disability, or limited English-speaking in the census tract i , P_{ct-i} is the total population number in the census tract i , S_{ct-i} is the area of the census tract i , N_{ct} is the number census tract intersected with the grid cell j , S_{gc-j} is the area of grid cell j , P_{gc-j} is the total population number in grid cell j , and R_{gc-j} is the rate of poverty, disability, or limited English-speaking in grid cell j .

Physical vulnerability. Physical vulnerability is measured by POI number, building age, and the foundation height of the building. When floods impact communities with a large number of POIs, they will severely disrupt the most fundamental service provided by these communities, such as access to health and food services, schools, and child care centers¹¹. Structure materials, construction practices, and building codes for older buildings are less advanced than the newer buildings. Older buildings tend to have a higher risk of structure damage under floods¹⁸⁻²⁰. These conditions make the old buildings less resistant to floods than the newly built ones. Buildings with higher foundation heights are more capable of preventing floodwater entry^{18,21}. Flood events will intensify their impacts on these physically vulnerable communities^{5,22}.

Building age and the reciprocal foundation height are calculated using the NSI dataset. The location information of the foundation height is the same as that of the building area in the

NSI dataset. The foundation height in each grid cell is calculated by taking the average value of foundation height across all structures inside it. The reciprocal foundation height value is further calculated by taking the reciprocal of the foundation height. The location information of building age is at the census tract level. The building age of each grid cell is calculated by taking the average value of building age across all census tracts intersecting with it.

Point-of-interest information is from SafeGraph²³. Fields of latitude, longitude, and naics_code of the POI dataset are adopted. The latitude and longitude are in the resolution of seven decimal places, which is accurate enough to calculate the number of POIs inside the grid cell. By filtering POI categories by NAICS codes starting with 11 and 21, facilities services POIs were included in this study. POIs categorized as facility services are included in this study. The total number of POIs within each grid cell is counted.

S3 Construction of KNN graph view

For each data sample in $\mathbf{BF}$, its similarity with other data points is calculated, and edges are set to its $top - K$ similar neighbors, following the work of Bo et al²⁴.

We first calculate the similarity matrix using the Heat Kernel as shown in the following equation:

$$S_{ij} = e^{-\frac{\|x_i - x_j\|^2}{t}} \quad (S2)$$

where S_{ij} is the similarity between grid cell i and j , $x_i \in R^{d_{BF}}$ is the row vector represents the i -th data sample in $\mathbf{BF}$, and t is the time parameter in the heat conduction equation.

With this similarity matrix, for each data sample, nodes whose similarity values are among the $top-K$ are chosen to be this node's neighbors. Edges are set to connect it with its neighbors. In this way, the KNN related adjacency matrix $\mathbf{K}$ is constructed. $\mathbf{K}$ is integrated with $\mathbf{BF}$, then the KNN graph view $\mathbf{G}_K = (\mathbf{K}, \mathbf{BF})$ is constructed. Where $\mathbf{K}$ provides graph structure information, $\mathbf{BF}$ provides feature matrix information.

S4 Data augmentation in the first layer of FloodRisk-Net

Attribute masking. As summarized by Wu et al²⁵, in attribute masking, researchers usually randomly select a fraction of feature dimensions to be masked. In our study, we randomly select some columns in $\mathbf{BF} \in R^{m \times n}$ to be masked, the value of selected columns will be set to be zeros while those not selected will not be changed. The column in $\mathbf{BF}$ represents whether there is flood occurrence during this week, and the row represents the data sample (grid cell). The practical meaning of this attribute masking is that we randomly select some weeks' flood information to be masked for all data samples. The probability distribution of the column selection is generated by Bernoulli distribution, which is frequently used in attribute masking as summarized by Wu et al²⁵. This attribute masking could be expressed using the following formula:

$$\overline{\mathbf{BF}} = \tau_{Am}(\mathbf{BF}) = \mathbf{BF}^T \odot \mathbf{Am} = [x_1^T \odot \mathbf{Am}, \dots, x_n^T \odot \mathbf{Am}]^T \quad (S3)$$

where, $\overline{\mathbf{BF}} \in R^{m \times n_{BF}}$ is the augmented binary flood occurrence matrix, $\tau_{Am}(\cdot)$ is the attribute masking transformation, $\mathbf{BF}^T$ is the transpose of the $\mathbf{BF}$, $\odot$ is the Hadamard operation, $\mathbf{Am} \in \{0,1\}^m$ is the masking column selection vector whose value is sampled from the Bernoulli distribution with probability $p^{(x)}$, x_i^T is the transpose of the x_i .

Edge perturbation. Wu et al²⁵ summarized that graph structure would be perturbed by randomly adding or removing a certain fraction of connections in edge perturbation

augmentation. In this work, a portion of edges in the graph structures in both views are randomly selected to be masked. The practical meaning is that the spatial flood dependence will be corrupted by adding or removing edges. The edge perturbation can be expressed as the following equation:

$$\bar{\mathbf{A}} = \tau_{ep}(\mathbf{A}) = \mathbf{A} \odot \mathbf{M} \quad (\text{S4})$$

where, $\mathbf{A} \in R^{m \times m}$ is the KNN graph structure $\mathbf{K}$ in KNN graph view, or the learned graph structure $\tilde{\mathbf{A}}^*$ in the learned graph view. $\tau_{ep}(\cdot)$ is the edge perturbation transformation. $\bar{\mathbf{A}}$ is the augmented graph structure. $\mathbf{M} \in \{0,1\}^{n_{BF} \times n_{BF}}$ is the graph masking matrix whose value is sampled from the Bernoulli distribution with probability $p^{(a)}$.

In this model, we conduct these two data augmentation strategies on both views as the following:

$$\overline{\mathbf{G}}_K = (\tau_{ep}(\mathbf{K}), \tau_{Am}(\mathbf{BF})), \overline{\mathbf{G}}_L = (\tau_{ep}(\tilde{\mathbf{A}}), \tau_{Am}(\mathbf{BF})) \quad (\text{S5})$$

where, $\overline{\mathbf{G}}_K$ and $\overline{\mathbf{G}}_L$ are augmented KNN graph view and learned graph view, respectively. We set the probability $p^{(x)}$ and $p^{(a)}$ to be different in both views, which means $p_K^{(x)} \neq p_L^{(x)}$, $p_K^{(a)} \neq p_L^{(a)}$.

S5 Eligibility for setting target distribution P as ground-truth labels

The target distribution P can be viewed as the square and normalization of each assignment in Q .^{24,26} Nodes close to the cluster center are regarded as high-confidence assignments. By raising the distribution Q to the second power in P , it aims to emphasize the role of these high-confidence assignments so as to optimize the node representation learning. From the other hand, the soft cluster frequencies normalization can prevent the solution that all instances assigned to the same cluster²⁶. These two properties have been analyzed as the key points of the success of the self-training strategy²⁶ and have been demonstrated powerful in clustering tasks^{24,27–30}.

S6 Model training

Hyperparameters tuning. Hyperparameters in *FloodRisk-Net* are summarized in Table S5. We set the searching space for each hyperparameter and use grid search to select the optimal set of hyperparameters. In the first layer of the model, the feature masking rate is searched in $\{0.1, 0.3, 0.6, 0.9\}$, edge dropping probability is tuned among $\{0, 0.25, 0.5, 0.75\}$, the embedding dimension of representation and projection in contrastive learning are both searched among $\{16, 32, 64, 128\}$, and the coefficient (τ) controlling the update of K is searched among $\{0.99, 0.999, 0.9999\}$. In the second layer, the embedding dimension of the autoencoder and GCN model is searched among $\{10, 16, 32, 64\}$. In the last layer, the clustering number which represents the number of flood risk levels is tuned in $\{5, 6, 7, 8, 9\}$, and coefficients α and β are both searched in $\{0.01, 0.1, 1\}$. The rest of these hyperparameters are set to be deterministic. The layer number of the graph learner embedding network (L_{en}) is set as 2. The layer number of encoding, projection and embedding in the contrastive learning part are all set as 2. Temperature parameter t in $NT - Xent$ is set as 0.2 following the work, of Chen et al.³¹. We set L_{GE} as 5 and set the dimension of the GCN or encoding network of DNN as $d_G - 500 - 500 - 2000 - 10$, where d_G is the dimension of the input data. Set the dimension of decoding network of DNN as $10 - 500 - 500 - 2000 - d_G$ following the work of Xie et al.³², Y. Chen

et al.³³, Bo et al²⁴. The DNN module is first pre-trained with 100 epochs and the learning rate is 10^{-3} following the work of Bo et al²⁴.

For each combination of hyperparameters, the model is tested and the performance of the model is evaluated using the Calinski and Harabasz (C&H) score ³⁴.

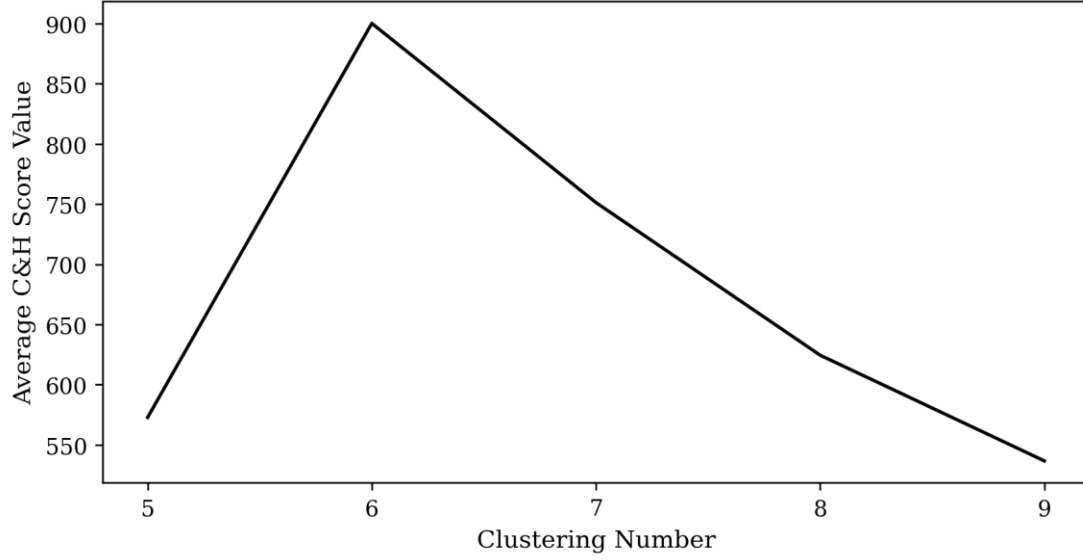
Supplementary Table 4. Hyperparameters and their value range

Layer	Hyperparameters	Value Range
1 st layer— Spatial Flood dependence Network Learner	Dimension of representation and projection (d_1 and d_2)	{16, 32, 64, 128}
	Layer number of the graph learner embedding network (L_{en})	2
	Layer number of encoding, projection and embedding in the contrastive learning part	2, 2, 2
	Feature masking probability	{0.1, 0.3, 0.6, 0.9}
	Edge dropping probability	{0, 0.25, 0.5, 0.75}
	Temperature for contrastive loss	0.2
	K on KNN graph view	10
	Learning rate	10^{-3}
	Training epoch	2000
	τ	{0.99, 0.999, 0.9999}
2 nd layer— Node representation learning	Dimension of the GCN	$d_G - 500 - 500 - 2000 - 10$
	Dimension of the encoding part of DNN	$d_G - 500 - 500 - 2000 - 10$
	Dimension of decoding part of DNN	$10 - 500 - 500 - 2000 - d_G$
	Balance coefficient (ϵ)	0.5
3 rd layer— Flood risk level determination	Clustering number	{5, 6, 7, 8, 9}
	Clustering coefficient in loss function (α)	{0.01, 0.1, 1}
	Coefficient β in the loss function	{0.01, 0.1, 1}
	Pretrain epoch and learning rate	100, 10^{-3}
	Model training epoch and learning rate	1000, 10^{-3}

Optimal shared clustering number. The optimal set of hyperparameters for each MSA is obtained after hypermeter tuning. The optimal value of the same type of hyperparameter in different MSAs would be different. In our model, the clustering number equals the number of flood risk levels. In order to compare the flood risk spatial distribution in each MSA better, we try to set the clustering number in each MSA to be the same through the following steps:

1. For each optional clustering number, search for the optimal set of hyperparameters other than the clustering number that maximizes the C&H score of each MSA.
2. Then, calculate the average C&H score across all five MSAs at each clustering number.
3. The optimal shared clustering number of all five MSAs is determined by selecting the clustering number that yields the highest average C&H score across all MSAs.

The average C&H score across all five MSAs for each clustering number is calculated and shown in Fig. S14 following the above steps. It shows that the average C&H score reaches its maximum when the clustering number is 6. Therefore, the optimal clustering number across five MSAs is set to be 6, which means we set the possible flood risk from level 1 (the lowest flood risk) to level 6 (the highest flood risk).



Supplementary Figure 14. The relations between the average Calinski and Harabasz score value across all five MSAs and the clustering number

S7 Calculation of Moran's I taking spatial dependence network as weight matrix

Following Li et al.³⁵ and Moran³⁶, global Moran's I can be calculated as:

$$I = \frac{n}{\sum_i \sum_j w_{ij}} \frac{\sum_i \sum_j w_{ij} z_i z_j}{\sum_i z_i^2} \quad (S6)$$

Where n is the number of grid cells in each MSA, z_i is the value of flood risk level in grid cell i , and w_{ij} is the strength of spatial dependence between grid cell i and grid cell j .

S8 Calculation of average flood risk level and spatial inequality

The average flood risk level of each city is calculated as the average flood risk value of all grid cells within the city boundary.

Spatial inequality is measured by Theil index. Following Manz and Mansmann³⁷; Rey et al.³⁸, and Theil³⁹, the Theil index can be calculated and decomposed as follows:

$$T = \sum_{i=1}^m \left(\frac{y_i}{\sum_{i=1}^m y_i} \ln \left[m \frac{y_i}{\sum_{i=1}^m y_i} \right] \right) \quad (S7)$$

$$= \left[\sum_{g=1}^w s_g \ln \left(\frac{m}{m_g} s_g \right) \right] + \left[\sum_{g=1}^w s_g \sum_{i \in g} s_{i,g} \ln(m_g s_{i,g}) \right] \quad (S8)$$

$$= B + W$$

where m is the total number of grid cells in each MSA, y_i is the flood risk level in grid cell i , w is the number of mutually exclusive groups (in this study it equals 2, representing the core city and other areas), m_g is the number of grid cells assigned to region g , $s_g = \frac{\sum_{i \in g} y_i}{\sum_i y_i}$, and

$$s_{i,g} = \frac{y_i}{\sum_{i \in g} y_i}.$$

The first item (B) measures the spatial inequality originates from between groups. It is

calculated by treating each group as a whole and then calculating their spatial inequality value similar to the normal Theil index equation (Equation S7). The second term represents the within-group spatial inequality. It is the weighted sum of spatial inequality in each group and each group's spatial inequality is calculated in the same way as the normal Theil index equation (Equation S7).

Reference

1. FEMA(Federal Emergency Management Agency). OpenFEMA Dataset : FIMA NFIP Redacted Claims - v1. 1–18 <https://www.fema.gov/openfema-data-page/fima-nfip-redacted-claims-v1> (2023).
2. Brunner, M. I., Furrer, R. & Favre, A. C. Modeling the spatial dependence of floods using the Fisher copula. *Hydrology and Earth System Sciences* **23**, 107–124 (2019).
3. de Moel, H. *et al.* Flood risk assessments at different spatial scales. *Mitigation and Adaptation Strategies for Global Change* **20**, 865–890 (2015).
4. Mohanty, M. P. *et al.* A new bivariate risk classifier for flood management considering hazard and socio-economic dimensions. *Journal of Environmental Management* **255**, 109733 (2020).
5. Wing, O. E. J. *et al.* Inequitable patterns of US flood risk in the Anthropocene. *Nature Climate Change* **12**, 156–162 (2022).
6. De Moel, H., Aerts, J. C. J. H. & Koomen, E. Development of flood exposure in the Netherlands during the 20th and 21st century. *Global Environmental Change* **21**, 620–627 (2011).
7. Koks, E. E., Jongman, B., Husby, T. G. & Botzen, W. J. W. Combining hazard, exposure and social vulnerability to provide lessons for flood risk management. *Environmental Science and Policy* **47**, 42–52 (2015).
8. USCB(United States Census Bureau). Decennial Census of Population and Housing Data. <https://www.census.gov/programs-surveys/decennial-census/data.html> (2020).
9. US Army(US Army Corps of Engineers Hydrologic Engineering Center). NSI Technical Documentation. <https://www.hec.usace.army.mil/confluence/nsi/technicalreferences/latest/technical-documentation> (2022).
10. Green, R., Bates, L. K. & Smyth, A. Impediments to recovery in New Orleans’ upper and lower ninth ward: One year after Hurricane Katrina. *Disasters* **31**, 311–335 (2007).
11. NASEM(National Academies of Sciences, Engineering, and M. *Framing the Challenge of Urban Flooding in the United States*. THE NATIONAL ACADEMIES PRESS (2019) doi:10.17226/25381.
12. Hemingway, L. & Priestley, M. Natural Hazards, Human Vulnerability and Disabling Societies: A Disaster for Disabled People? *The Review of Disability Studies* **2**, (2006).
13. Stough, L. M., Sharp, A. N., Resch, J. A., Decker, C. & Wilker, N. Barriers to the long-term recovery of individuals with disabilities following a disaster. *Disasters* **40**, 387–410 (2016).
14. Hamel, L., Wu, B., Brodie, M., Sim, S.-C. & Marks, E. *Early Assessment of Hurricane Harvey’s Impact on Vulnerable Texans in the Gulf Coast Region: Their Voices and Priorities to Inform Rebuilding Efforts*. Kaiser Family Foundation and Shao-Chee Sim and Elena Marks Episcopal Health Foundation <http://files.kff.org/attachment/Report-An-Early-Assessment-of-Hurricane-Harveys-Impact-on-Vulnerable-Texans-in-the-Gulf> (2017).
15. Hallegatte, S., Vogt-Schilb, A., Bangalore, M. & Rozenberg, J. *Unbreakable: Building the Resilience of the Poor in the Face of Natural Disasters*. World Bank, 2017 (2017).

16. Liu, Z., Liu, C. & Mostafavi, A. Beyond Residence : A Mobility-based Approach for Improved Evaluation of Human Exposure to Environmental Hazards. (2023).
17. CDC(Centers for Disease Control and Prevention). CDC/ATSDR Social Vulnerability Index 2020. 15–16 <https://www.atsdr.cdc.gov/placeandhealth/svi/index.html> (2023).
18. Fernandez, P., Mourato, S., Moreira, M. & Pereira, L. A new approach for computing a flood vulnerability index using cluster analysis. *Physics and Chemistry of the Earth* **94**, 47–55 (2016).
19. M. Neubert, T. Naumann, J. H. and J. N. The Geographic Information System-based flood damage simulation model HOWAD. *Flood Risk Management* (2016) doi:10.1111/jfr3.12109.
20. Simonovic, S. P. *et al. Floods : Mapping Vulnerability in the Upper Thames Watershed under a Changing Climate*. (2007).
21. Godfrey, A., Ciurean, R. L., Westen, C. J. Van, Kingma, N. C. & Glade, T. Assessing vulnerability of buildings to hydro-meteorological hazards using an expert based approach – An application in Nehoiu Valley , Romania. *International Journal of Disaster Risk Reduction* **13**, 229–241 (2015).
22. Malgwi, M. B., Fuchs, S. & Keiler, M. A generic physical vulnerability model for floods: Review and concept for data-scarce regions. *Natural Hazards and Earth System Sciences* **20**, 2067–2090 (2020).
23. SafeGraph. SafeGraph Technical Documentation. 1–13 <https://docs.safegraph.com/docs/core-places> (2023).
24. Bo, D. *et al.* Structural Deep Clustering Network. *The Web Conference 2020 - Proceedings of the World Wide Web Conference, WWW 2020* 1400–1410 (2020) doi:10.1145/3366423.3380214.
25. Wu, L., Lin, H., Tan, C., Gao, Z. & Li, S. Z. Self-supervised Learning on Graphs: Contrastive, Generative, or Predictive. *IEEE Transactions on Knowledge and Data Engineering* 1–1 (2021) doi:10.1109/tkde.2021.3131584.
26. Zhou, S. *et al.* A Comprehensive Survey on Deep Clustering: Taxonomy, Challenges, and Future Directions. *Proceedings of ACM Computing Survey* **1**, (2022).
27. Guo, X., Gao, L., Liu, X. & Yin, J. Improved deep embedded clustering with local structure preservation. *IJCAI International Joint Conference on Artificial Intelligence* **0**, 1753–1759 (2017).
28. Wang, C. *et al.* Attributed graph clustering: A deep attentional embedding approach. *IJCAI International Joint Conference on Artificial Intelligence* **2019-Augus**, 3670–3676 (2019).
29. Fan, S. *et al.* One2Multi Graph Autoencoder for Multi-view Graph Clustering. *The Web Conference 2020 - Proceedings of the World Wide Web Conference, WWW 2020* 3070–3076 (2020) doi:10.1145/3366423.3380079.
30. Kulatilleke, G. K., Portmann, M. & Chandra, S. S. *SCGC : Self-Supervised Contrastive Graph Clustering. Proceedings of ACM Conference (Conference '17)* vol. 1 (Association for Computing Machinery, 2022).
31. Chen, T., Kornblith, S., Norouzi, M. & Hinton, G. A simple framework for contrastive learning of visual representations. *37th International Conference on Machine Learning, ICML 2020 Part F* **16814**, 1575–1585 (2020).

32. Xie, J., Girshick, R. & Farhadi, A. Unsupervised deep embedding for clustering analysis. *33rd International Conference on Machine Learning, ICML 2016* **1**, 740–749 (2016).
33. Chen, Y., Wu, L. & Zaki, M. J. Iterative deep graph learning for graph neural networks: Better and robust node embeddings. *Advances in Neural Information Processing Systems* **2020-Decem**, (2020).
34. Caliński, T. & Harabasz, J. A Dendrite Method For Cluster Analysis. *Communications in Statistics* **3**, 1–27 (1974).
35. Li, H., Calder, C. A. & Cressie, N. Beyond Moran's I: Testing for spatial dependence based on the spatial autoregressive model. *Geographical Analysis* **39**, 357–375 (2007).
36. Moran, P. Notes on Continuous Stochastic Phenomena Published by : Biometrika Trust Stable URL : <http://www.jstor.org/stable/2332142>. *Biometrika* **37**, 17–23 (1950).
37. Manz, K. M. & Mansmann, U. Inequality indices to monitor geographic differences in incidence, mortality and fatality rates over time during the COVID-19 pandemic. *PLoS ONE* **16**, 1–17 (2021).
38. Rey, S. J., Arribas-Bel, D. & Wolf, L. J. Spatial Inequality Dynamics. 1–24 (2022).
39. Theil, H. *ECONOMICS AND INFORMATION THEORY*. (1969).