

Exploratory Analysis of Recommending Urban Parks for Health-Promoting Activities

Linus W. Dietz

linus.dietz@kcl.ac.uk

Department of Informatics, King's College London
London, United Kingdom

Ke Zhou

ke.zhou@nokia-bell-labs.com

Nokia Bell Labs

Cambridge, United Kingdom

Sanja Šćepanović

sanja.scepanovic@nokia-bell-labs.com

Nokia Bell Labs

Cambridge, United Kingdom

Daniele Quercia

daniele.quercia@nokia-bell-labs.com

Nokia Bell Labs

Cambridge, United Kingdom

Abstract

Parks are essential spaces for promoting urban health, and recommender systems could assist individuals in discovering parks for leisure and health-promoting activities. This is particularly important in large cities like London, which has over 1,500 named parks, making it challenging to understand what each park offers. Due to the lack of datasets and the diverse health-promoting activities parks can support (e.g., physical, social, nature-appreciation), it is unclear which recommendation algorithms are best suited for this task. To explore the dynamics of recommending parks for specific activities, we created two datasets: one from a survey of over 250 London residents, and another by inferring visits from over 1 million geotagged Flickr images taken in London parks. Analyzing the geographic patterns of these visits revealed that recommending nearby parks is ineffective, suggesting that this recommendation task is distinct from Point of Interest recommendation. We then tested various recommendation models, identifying a significant popularity bias in the results. Additionally, we found that personalized models have advantages in recommending parks beyond the most popular ones. The data and findings from this study provide a foundation for future research on park recommendations.

1 Introduction and Background

With increasing urbanization [17], populations grow, and so does the need for spaces where people can relax and play. Parks are urban spaces for such activities, and they support public health and well-being [30, 32], particularly for elderly and individuals from socioeconomically disadvantaged backgrounds [10, 20]. Given the numerous advantages of parks, it is important to make them accessible to all citizens. Understanding what each park offers can be overwhelming, especially in large cities such as London. A way of helping citizens navigate this challenge is through recommender systems that can aid in understanding which parks are available for their needs. However, it is unclear what kind of recommendation algorithms are well-suited for park recommendations, so that each park may be recommended for various activities.

To tackle this gap, we introduce and study the intricacies of park recommendations, with the main purpose of supporting the promotion of urban health. We treat parks as items, and state the problem as follows: a user wants to do a certain activity in a park, but it is unknown which park would be suitable for this activity.

The recommender system should infer the user's preferences from previous visits and recommend parks that are most suitable. Our main contributions are as follows:

1. Creating datasets for park recommendation. Given the absence of datasets for our task, our first contribution was to create datasets of park visitations by profiling London parks. Cities or other urban areas have been characterized in several ways, for example, by capturing the vibe of neighborhoods in different cities [19], comparing data sources to capture the touristic experience [8] or the smell and soundscape of cities [2, 23]. However, when it comes to parks, specifically, we did not identify research to characterize them in meaningful ways that would allow for using them in a recommender system. To advance beyond previous work [22] and support the goal of promoting urban health, we decided to characterize parks by the health-promoting activities they offer to citizens. We found the taxonomy of activities in parks introduced by previous work [9] suitable for park recommender systems. This taxonomy, based on a literature survey on health-promoting activities and an expert panel, introduced five distinct activity categories: cultural, environmental, nature-appreciation, physical, and social. We then gathered and released¹ two distinct sets of user preference data (section 2), one from a survey, and one from geo-located Flickr images.

2. Evaluating baseline recommendation models. Analyzing these data sets, we found that the geographic patterns of park visits differ from those of POIs in a way that park visits are not highly local. This limits the adoption of POI recommendation algorithms. In response to this finding, we designed a series of offline recommendation experiments to understand the factors that influence the quality of park recommendations, especially regarding recommending different activities. Thus, the second contribution of this paper is an exploratory analysis and evaluation of existing methods (section 3). We found that recommending nearby parks as done by techniques for recommending POIs is not effective. Our follow-up analysis of five widely used recommendation models revealed a significant popularity bias in the recommendations, requiring managing popularity bias to make for personalized models competitive [1].

¹<https://github.com/LinusDietz/park-recommendation-datasets/tree/recsys-lbr>

2 Datasets for Park Recommendation

Parks are urban spaces with a wide array of offerings to different users and even the same user in different situations depending on their activities, such as sports or enjoying nature. To evaluate activity-aware park recommendations, we first needed a dataset for these recommendations. As such data does not exist, we established two complementary data sets for our evaluations. The first is a survey-based dataset, where we asked citizens to name parks they find suitable for doing certain activities. The second is based on a large-scale image data set from Flickr, which we processed into an implicit feedback recommendation dataset of park visits.

2.1 Survey on Parks for Different Activities

In an online survey, we asked Londoners about suitable parks for performing each of the five activity categories from a taxonomy of health-promoting activities [9]. The main questions were phrased as: “Can you name several parks suitable for *physical activities* (e.g., sports)?”, with equivalent phrasing for other the activities.

We recruited the participants using the first author’s institutional research recruitment portal ($n = 81$), as well as through Prolific ($n = 178$). The participants (F: 48%, M: 41%, O: 11%) were informed about the purpose of the survey, the voluntary nature of their participation, and the scope of the data collection. The average age was 36.2 ± 10.94 and the home locations of the respondents were uniform across Greater London, with the exception that only a few were from the centrally located City of London. Finally, we asked people for a park close to their homes, which we could use as an instructional manipulation check in conjunction with the reported postal area and an attention check, where in one of the steps, we asked to select ‘Hyde Park’ instead of naming parks suitable for a certain activity. The data collection was registered as a minimal-risk study at the first author’s institutional review board. After having removed the answers of users who failed any attention check, we interpreted each park of the survey as a signal that the user finds it suitable for the respective activity. To establish a uniform terminology, we use the term *visit* for each mention by a participant of a park in the rest of the paper.

2.2 Geotagged Images from Flickr

Since its inception in 2004, Flickr has gained considerable popularity for sharing photography, accumulating billions of images. Notably, many of these images have been precisely geo-located, thanks to the utilization of the (phone) camera’s GPS module. We utilized a substantial dataset, collected using Flickr’s API, comprising geo-located images posted between 2004 and 2015 to record park visits. By intersecting the 12 million images with the park outlines as defined in OpenStreetMap (OSM), we identified individual visits to parks in London. Note that we only analyzed named parks, as unnamed parks are quite small in size, are typically not under active maintenance, and without a name are cumbersome to refer to. Following several pre-processing steps, we converted these visits into a useful and realistic dataset for park recommendations.

First, we selected all images of a user within London and computed the centroid from these to estimate the user’s center of life within the city, which has been shown to approximate the home location [21]. In the next step, we used OSM to gather the shapes of

all parks in London and discarded all images that were not within the park boundaries. This left us with 1,065,197 individual park visits. As we are interested in specific activities of the user, we analyzed two types of tags the images were partially annotated with: user-generated tags and computer vision labels from a computer vision algorithm [28]. To match these tags to activities, we employed Sentence-BERT [24] for text embeddings. For this, we utilized a lexicon of amenities and spaces found in parks [9] to match them with the Flickr tags. After embedding the OSM tags using the `all-mpnet-base-v2` model, we matched each Flickr label to the closest OSM tag in the embedding space using cosine distance as the similarity measure and a similarity threshold of at least 0.7. This approach helped us avoid matching Flickr labels that did not have meaningful OSM counterparts. We manually confirmed the agreement of the 20 most frequent Flickr tags to activities using three expert annotations aggregated using majority voting. The agreement was 82%, which is highly accurate given that the matchings are only based on individual tags.

The outcome was that we annotated all labels attached to an image with an activity. For example, if the labels were [‘sunshine’, ‘volleyball’, ‘pond’], we would discard sunshine, as it is unrelated to any activities, and count the other two tags towards physical and nature-appreciation, respectively, using a relative frequency count of 0.5 for physical and 0.5 for nature-appreciation. Doing so for all images, we could then subdivide them into visits to parks and infer the activity that the user was involved in. To prevent recording spurious activities, we only assigned an image to an activity if the relative frequency count is at least 0.5.

In this way, we obtained an implicit feedback dataset from Flickr, which records visits of a user to parks and also contains information about which activity the user was interested in when capturing the photo. As we were interested in recommending novel items to a user, we removed all duplicate visits to a park for one activity.

2.3 Dataset Characteristics

Table 1 shows the data characteristics of the two datasets. The survey dataset has a drawback in terms of the number of users and visits but has a manageable density of around 2%. The Flickr dataset has an order of magnitude larger size, but its density is very low, ranging from 0.21% (nature-appreciation) to 0.62% in the environmental category, albeit with only 189 distinct park visits recorded for this activity. Still, the city-level density of the Flickr dataset is comparable to that of POI recommendation datasets [27].

Analyzing the geographical aspects of visiting parks reveals different patterns of visits for distinct activity categories. We fit a scaling factor α for the distances between the users’ home/center location and the visited parks [26]:

$$p_{\text{close}} = \frac{1}{d_{u,p}^\alpha},$$

where $d_{u,p}$ is the distance between the user’s location (home post-code in the survey, center of geographic interest in Flickr) and the centroid of the visited park. To derive the probability distribution, we used a bin width of 1km. We then fit the scaling factor α , tabulated in the last columns of Table 1, for the probability of traveling a certain distance for a certain activity. The results reveal that using the center of geographic interest is likely not a good proxy for

Table 1: Characteristics of the datasets (Survey on the left, Flickr on the right). VPU – visits per user, VPP – visits per park.

Survey	Users	Parks	Visits	Density %	VPU	VPP	α
Cultural	208	97	501	2.4831	2.4087	5.1649	1.90
Env.	178	125	347	1.5596	1.9494	2.	2.03
Nature	249	146	630	1.7330	2.5301	4.3151	2.02
Physical	252	203	947	1.8512	3.7579	4.6650	2.04
Social	248	156	882	2.2798	3.5565	5.6538	1.94

Flickr	Users	Parks	Visits	Density %	VPU	VPP	α
Cultural	6132	655	10426	0.2596	1.7003	15.9176	2.33
Env.	420	189	490	0.6173	1.1667	2.5926	2.06
Nature	10203	821	17692	0.2112	1.7340	21.5493	2.22
Physical	4743	628	7202	0.2418	1.5184	11.4682	1.99
Social	3919	510	5626	0.2815	1.4356	11.0314	2.1

home location as the resulting values were smaller compared to the survey and inconclusive regarding the distances traveled to visit a park for a certain activity. Hence, we focused on the results in the survey dataset. The higher the value of α , the less significant distance becomes in deciding to visit a park for this activity [26]. In the survey, we found that α is highest for cultural activities. This result is logical, as fewer parks offer high-quality cultural experiences like museums, artwork, and concerts, making people willing to travel longer distances for these activities. Moreover, specific cultural events in these parks might be occasional rather than continuous. Conversely, α was lowest for environmental and physical activities, which also makes sense since people prefer cultivating city plots nearby for frequent visits and find it easier to engage in sports in nearby parks.

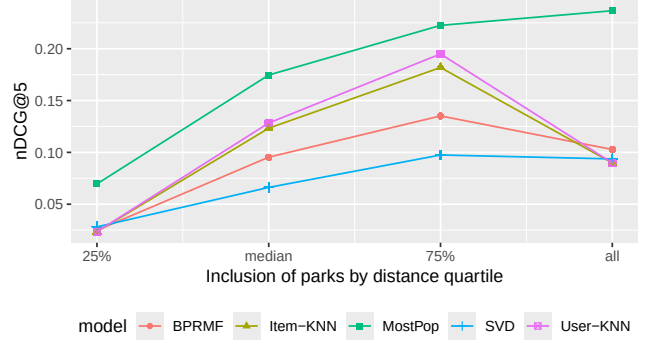
These findings revealed the advantages and shortcomings of the two park recommendation datasets we created. Specifically, while our survey dataset is highly accurate, the Flickr dataset is about 12 times larger. However, besides having less accurate user home locations, it also has a more imbalanced distribution of activities, with very few images depicting environmental activities.

3 Evaluation

The goal was to find out which algorithms and strategies lead to the best recommendation quality for recommending parks for a specific activity. We treat this problem as an implicit feedback recommendation problem and our experiments aim to uncover what are the most influential factors to compute high-quality park recommendations. The following experiments are conducted using Elliot [4] to enable reproducible evaluation in terms of preprocessing steps, recommendation models, and evaluation metrics.

3.1 Selecting Recommendation Models

To understand park recommendations, we needed models that are: (i) suitable for small datasets (since the number of parks in a city limits the dataset size), (ii) interpretable (enabling straightforward explanations for recommendation performance), and (iii) perform competitively [11]. This led to the choice of the following models: MostPop [6], User/Item-KNN [3], SVD [16], and BPRMF [25]. We also considered the latest deep learning-based recommendation algorithms and established tensor factorization techniques [13, 14]. However, these were excluded due to their black-box nature or unsuitability for smaller datasets, which need to scale with model complexity [31]. Additionally, we disregarded POI algorithms, as their geographic assumptions did not align with our task.

**Figure 1: Recommendation accuracy (nDCG@5) depends on the percentage (x-axis) of distant parks included (survey).**

3.2 Training-test Split

Due to the small size of the datasets, we performed a leave-one-out split on a user basis, where one random visit per activity was reserved for the test set. We acknowledge that a temporal split of a certain number of interactions per user would have been preferable, but considering the data characteristics (Table 1), particularly the low number of visits per user, this was not viable.

3.3 Results

(1) *Recommending nearby parks (as done by techniques for recommending POIs) is not effective.* In the first step, we expanded on the analysis of geographic influences in §2.3 by examining how recommendation accuracy is affected by including parks at increasingly greater distances from the user’s home. For each user, we subdivided parks into 4 quartiles based on the distance to the users’ home location. To ensure an accurate home location was used in this experiment, we only used the survey dataset, where this information was available. We then ran the recommendation experiments four times, each time with one more quartile of parks being available for recommendation. To eliminate the effects of different sizes of possible items, we randomly subsampled the number of interactions to be equal to the first quartile. We repeated the experiments 100× to average out the influences from the subsampling.

With increasingly more distant parks in the candidate pool, we observe an increase in the average recommendation accuracy for the five activity categories until the 75% quartile, and smaller gains or even worsening afterward when including more distant items (Figure 1). Especially, the first quartile of parks seems to include very few relevant parks that can be recommended with abysmal recommendation accuracy of an nDCG@5 < 0.05. Overall, MostPop

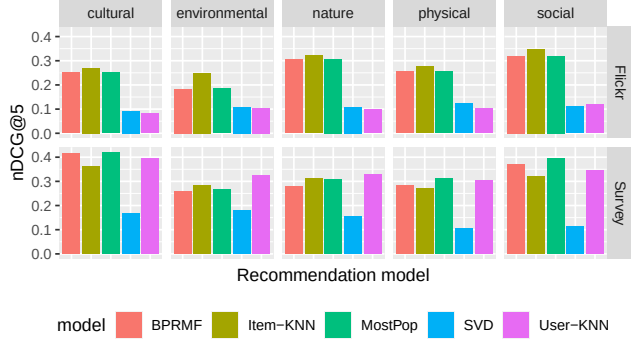


Figure 2: Recommendation accuracy (nDCG@5) is highest for recommendation models influenced by popularity.

performs best, while User-KNN slightly outperforms Item-KNN. The matrix-factorization algorithms BPRMF and SVD follow a similar pattern but seem to suffer from the small number of interactions compared to the other models. This finding is highly relevant for the upcoming experiments, as it shows that including distant parks up to Q3 improves the accuracy of the recommendations, which is in violation of the basic assumption of POI recommendation models that incorporate geographic proximity into their scoring [27].

(2) *Recommending default (popular) parks is effective.* We now turn our attention to recommending parks for individual activities, by subdividing the test set interactions by activity (Figure 2). In the Flickr dataset, the Item-based KNN, MostPop, and BPRMF methods generally performed best, with SVD and User-KNN consistently being in the range of 0.08 – 0.1. Activity-wise, the highest accuracy was achieved in nature-appreciation and social activities. In the survey dataset, MostPop was best for cultural, physical, and social activities, but is beaten in the environmental and nature-appreciation categories by User-KNN. BPRMF and Item-KNN were also competitive leaving SVD as the only model that fails to produce high-quality recommendations. In this dataset, the highest NDCG@5 was achieved in the recommendations for cultural activities, followed by social.

Assessing the Popularity Bias of the recommended items (cf. Figure 3) using the Average Ranking Popularity [1], we observed a clear trend that the recommended items by MostPop and BPRMF are highly popular, whereas Item-KNN provided good recommendations, especially in the Flickr dataset without an over-reliance on popular items. This phenomenon of popularity bias is likewise prevalent in POI recommendation [1, 27], where a common strategy is to remove a certain portion of the most popular items to increase the value of the resulting recommendations to the user [7, 27].

(3) *After disregarding highly popular parks, personalized recommendation models consistently outperform non-personalized models.* Removing the most popular parks, i.e., those visited by the highest number of different users, led to an even more difficult recommendation problem as the sparsity is further increased. However, the recommendation problems become more realistic and the usefulness of the recommendations for the user typically increases [12].

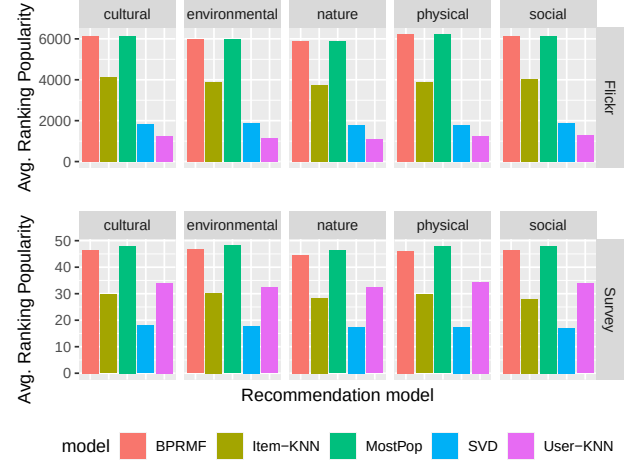


Figure 3: Popularity Bias (y-axis) in the recommendations of different recommendation models. BPRMF and MostPop rely on recommending highly-popular parks questioning the usefulness of these recommendations.

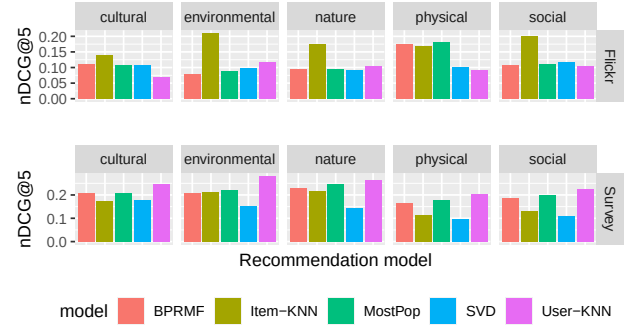


Figure 4: Recommendation accuracy (nDCG@5) after dropping the top 0.5% most popular parks. The popularity-influenced methods (BPRMF and MostPop) are outperformed by neighborhood-based models for most activities.

Experimenting with dropping different amounts of the most popular parks, we chose 0.5% as the threshold. Even with this tiny portion removed, we observed a change in the ranking of the results, with neighborhood-based models outperforming popularity-influenced methods such as MostPop and BPRMF.

We interpret this finding as follows: removing the *obvious* parks leads to a more realistic recommendation problem, and the user and item-based neighborhood models are best suited for these small but sparse cold-start recommendation problems [18]. The success of the Item-KNN model in the Flickr dataset and the User-KNN in the survey can be explained by the high number of visits per park and visits per user, respectively (cf. Table 1). Surprisingly, SVD was consistently outperformed, which appears to be a consequence of the small size of the datasets coupled with their sparsity.

4 Conclusions and Future Work

In this paper, we did an initial investigation into park recommendation, which seems to differ from POI recommendations in two key aspects. Each park supports various health-promoting activities and the geographic factors play a different role in park recommendations compared to standard POIs recommendation. After creating two datasets (survey and Flickr), we analyzed park recommender system requirements, noting the significance of distant parks and challenges with sparsity and popularity biases of park visits. After mitigating popularity bias, neighborhood models performed well in our recommendation experiments, with User-KNN and Item-KNN outperforming alternatives in the survey and Flickr datasets, respectively. Building on our initial findings, it would be worthwhile to analyze the park recommendation task with larger datasets in various cities and explore potentially further data sources beyond a survey and geotagged images. Larger datasets would provide chances to test more sophisticated models, such as deep or reinforcement learning, while more cities and data sources would improve the generalizability of the findings. Furthermore, user studies need to be conducted to understand human aspects like presentation and reception of recommendations, which will be crucial for real-world effectiveness [15]. Park recommendations can play a pivotal role in enhancing urban health by decreasing information barriers and, thus, improving the accessibility of parks [5] especially for vulnerable communities, which is a central pillar of the UN Sustainability Development Goals [29].

References

- [1] Himan Abdollahpour, Robin Burke, and Bamshad Mobasher. 2019. Managing Popularity Bias in Recommender Systems with Personalized Re-ranking. In *FLAIRS*. 413–418.
- [2] Luca Maria Aiello, Rossano Schifanella, Daniele Quercia, and Francesco Aletta. 2016. Chatty maps: constructing sound maps of urban areas from social media data. *Royal Society Open Science* 3, 3 (2016). <https://doi.org/10.1098/rsos.150690>
- [3] Fabio Aioli. 2013. Efficient top-n recommendation for very large scale binary rated datasets. In *RecSys*. ACM, 273–280.
- [4] Vito Walter Anelli, Alejandro Bellogin, Antonio Ferrara, Daniele Malatesta, Felice Antonio Merla, Claudio Pomo, Francesco Maria Donini, and Tommaso Di Noia. 2021. Elliot: A Comprehensive and Rigorous Framework for Reproducible Recommender Systems Evaluation. In *SIGIR*. ACM, 2405–2414.
- [5] Alice Battiston and Rossano Schifanella. 2024. On the need for a multi-dimensional framework to measure accessibility to urban green. *npj Urban Sustainability* 4, 1 (March 2024). <https://doi.org/10.1038/s42949-024-00147-y>
- [6] Alejandro Bellogin, Pablo Castells, and Iván Cantador. 2017. Statistical biases in Information Retrieval metrics for recommender systems. *Inf. Ret. J.* 20, 6 (July 2017), 606–634. <https://doi.org/10.1007/s10791-017-9312-z>
- [7] Linus W. Dietz, Pablo Sánchez, and Alejandro Bellogin. 2023. Understanding the Influence of Data Characteristics on the Performance of Point-of-Interest Recommendation Algorithms. (2023). [arXiv:2311.07229](https://arxiv.org/abs/2311.07229) [cs.LG]
- [8] Linus W. Dietz, Mete Sertkan, Saadi Myftija, Sameera Thimbiri Palage, Julia Neidhardt, and Wolfgang Wörndl. 2022. A Comparative Study of Data-driven Models for Travel Destination Characterization. *Front. Big Data* 5 (April 2022).
- [9] Linus W. Dietz, Sanja Šćepanović, Ke Zhou, André Felipe Zanella, and Daniele Quercia. 2024. Examining Inequality in Park Quality for Promoting Health Across 35 Global Cities. (2024). [arXiv:2407.15770](https://arxiv.org/abs/2407.15770) [cs.LG] <https://arxiv.org/abs/2407.15770>
- [10] Angel Mario Dzhambov and Donka Dimitrova. 2014. Elderly visitors of an urban park, health anxiety and individual awareness of nature experiences. *Urban For. Urban Gree.* 13, 4 (2014), 806–813.
- [11] Maurizio Ferrari Dacrema, Paolo Cremonesi, and Dietmar Jannach. 2019. Are we really making much progress? A worrying analysis of recent neural recommendation approaches. In *RecSys*. ACM, NY, USA, 101–109.
- [12] Mouzhi Ge, Carla Delgado-Battenfeld, and Dietmar Jannach. 2010. Beyond accuracy: evaluating recommender systems by coverage and serendipity. In *RecSys*. ACM, NY, USA, 257–260.
- [13] Yuchin Juan, Yong Zhuang, Wei-Sheng Chin, and Chih-Jen Lin. 2016. Field-aware Factorization Machines for CTR Prediction. In *RecSys*. ACM, NY, USA.
- [14] Alexandros Karatzoglou, Xavier Amatriain, Linas Baltrunas, and Nuria Oliver. 2010. Multiverse recommendation: n-dimensional tensor factorization for context-aware collaborative filtering. In *RecSys*. ACM, NY, USA, 79–86.
- [15] Bart P. Knijnenburg and Martijn C. Willemsen. 2015. *Evaluating Recommender Systems with User Experiments*. Springer, Boston, MA, 309–352.
- [16] Yehuda Koren and Robert Bell. 2011. *Advances in Collaborative Filtering*. Springer, Boston, MA, 145–186.
- [17] Debolina Kundu and Arvind Kumar Pandey. 2020. World urbanisation: trends and patterns. *Developing national urban policies: Ways forward to green and smart cities* (2020), 13–49.
- [18] Urszula Kuźelewska. 2020. *Effect of Dataset Size on Efficiency of Collaborative Filtering Recommender Systems with Multi-clustering as a Neighbourhood Identification Strategy*. Springer, 342–354.
- [19] Géraud Le Falher, Aristides Gionis, and Michael Mathioudakis. 2015. Where is the Soho of Rome? Measures and algorithms for finding similar neighborhoods in cities. In *ICWSM*, Vol. 9. 228–237.
- [20] Jolanda Maas, Robert A Verheij, Peter P Groenewegen, Sierp de Vries, and Peter Spreeuwenberg. 2006. Green space, urbanity, and health: how strong is the relation? *J. of Epidemiology & Community Health* 60, 7 (July 2006), 587–592.
- [21] Adam Poulston, Mark Stevenson, and Kalina Bontcheva. 2017. Hyperlocal Home Location Identification of Twitter Profiles. In *HT*. ACM, NY, USA, 45–54.
- [22] Kostis Proustouris, Harry Nakos, Yannis Stavrakas, Konstantinos I Kotsopoulos, Theofanis Alexandridis, Myrto S Barda, and Konstantinos P Ferentinos. 2021. An integrated system for urban parks touring and management. *Urban Science* 5, 4 (2021), 91.
- [23] Daniele Quercia, Rossano Schifanella, Luca Maria Aiello, and Kate McLean. 2015. Smelly maps: the digital life of urban smells. In *ICWSM*, Vol. 9. 327–336.
- [24] Nils Reimers and Iryna Gurevych. 2019. Sentence-BERT: Sentence Embeddings using Siamese BERT-Networks. In *EMNLP*. ACL.
- [25] Steffen Rendle, Christoph Freudenthaler, Zeno Gantner, and Lars Schmidt-Thieme. 2009. BPR: Bayesian Personalized Ranking from Implicit Feedback. In *UAI*. AUAI Press, 452–461.
- [26] Diego Saez-Trumper, Daniele Quercia, and Jon Crowcroft. 2012. Ads and the City: Considering Geographic Distance Goes a Long Way. In *RecSys*. ACM, NY, USA, 187–194.
- [27] Pablo Sánchez and Alejandro Bellogin. 2022. Point-of-interest Recommender Systems Based on Location-based Social Networks: A Survey from an Experimental Perspective. *Comput. Surveys* (Jan. 2022).
- [28] Bart Thomee, David A. Shamma, Gerald Friedland, Benjamin Elizalde, Karl Ni, Douglas Poland, Damian Borth, and Li-Jia Li. 2016. YFCC100M: The New Data in Multimedia Research. *Commun. ACM* 59, 2 (Jan. 2016), 64–73.
- [29] Department of Economic United Nations and Social Affairs Sustainable Development. 2015. Transforming our World: the 2030 Agenda for Sustainable Development. , 16301 pages. <https://sdgs.un.org/2030agenda>
- [30] David Vlahov, Nicholas Freudenberger, Fernando Proietti, Danielle Ompad, Andrew Quinn, Vijay Nandi, and Sandro Galea. 2007. Urban as a Determinant of Health. *J. Urban Health* 84, S1 (March 2007), 16–26.
- [31] Carole-Jean Wu, Robin Burke, Ed H. Chi, Joseph Konstan, Julian McAuley, Yves Raimond, and Hao Zhang. 2020. Developing a Recommendation Benchmark for MLPerf Training and Inference. <https://doi.org/10.48550/ARXIV.2003.07336>
- [32] Shengbiao Wu, Bin Chen, Chris Webster, Bing Xu, and Peng Gong. 2023. Improved human greenspace exposure equality during 21st century urbanization. *Nat. Commun.* 14, 1 (Oct. 2023).