

**Are perceived neighborhoods palimpsests? Analyzing
self-defined New York City Neighborhoods in the context of historical redlining**

By
Kai L. Kresek

An honors thesis submitted in partial
fulfillment of the requirements for the degree of
Bachelor of Arts with honors designation in the Department of Geography

Examining Committee:

Dr. Colleen E. Reid, Primary Thesis Advisor
Department of Geography

Dr. William Travis, Secondary Thesis Advisor
Department of Geography

Dr. Daniel C. Jones, Honors Committee Advisor
College of Arts and Science Honors Program

UNIVERSITY OF COLORADO AT BOULDER
DEFENSE DATE: APRIL 2, 2018

“What else is a nation but a patchwork of cities and towns; cities and towns a patchwork of neighborhoods; and neighborhoods a patchwork of homes?”

- Matthew Desmond, *Evicted: Property and Profit in the American City*

To my family,
Thank you for always being beside me,
from Anhui Province to Colorado.
Your love has made this journey possible.
我爱你

To my friends,
Thank you for your unwavering patience and support –
I wouldn't have made it without you.

To Dr. Colleen Reid,
Thank you for your constant engagement, generosity
and dedication throughout this entire process.
You have truly helped me grow as a student and a researcher.

To Dr. Daniel Jones,
Thank you for your mentorship throughout these four years of study.
Your guidance has opened the door to so many opportunities.

To Dr. Bill Travis and the Department of Geography,
Thank you for fostering within me an interdisciplinary perspective of the world,
as well as the tools and the critical thinking skills I will need for the future.

ABSTRACT

Neighborhood analysis is the inquiry into how neighborhoods are constructed and implemented for applications in human-centered studies. This is a relevant field of research given the impact that these units have on the analysis of spatial data. As neighborhood-based research continues to develop, it is necessary to examine how perceived neighborhoods, which are neighborhoods defined by individuals, are being constructed over time and what factors are contributing to their delineation.

There is a lack of inquiry into why perceived neighborhoods take on specific dimensions. This study proposes that the dimensions assumed by these neighborhoods are influenced by the demographic characteristics of individuals, and the historical housing policies known as redlining which have shaped neighborhoods. The data used in this study are perceived neighborhoods and demographic characteristics from survey respondents in New York City and historical Home Owner's Loan Corporation (HOLC) mortgage risk-assessment grades. Our methods involved developing a procedure for geocoding perceived neighborhood data from a telephone survey, integrating perceived neighborhood data from both telephone and online surveys, and conducting a statistical analysis to look for trends in the neighborhood polygon characteristics of area, perimeter and compactness in relation to survey methodology, demographic data and HOLC grades.

Our results include a novel methodology for geocoding perceived neighborhoods, as well as salient relationships between neighborhood polygon characteristics (area, perimeter and compactness) and demographic variables (age, gender, Hispanic identity, racial identity, employment status, education attainment, marriage status, and length of residency in neighborhood). Mean neighborhood perimeters were significantly different between male and female survey participants. Neighborhood area and perimeter were significantly different between categories within our Hispanic and employment status demographic variables. Area and compactness showed statistically significant differences between our racial identity and education variables. Groups within our income variable showed differences in neighborhood compactness. Lastly, we found that categories within demographic variables of race, income and education showed statistically significant associations with our HOLC data.

Keywords: Urban geography, geographic information systems, social environment, New York City, neighborhoods, historical redlining

TABLE OF CONTENTS

I. Introduction

i. Background.....	6
ii. Literature Review.....	8
a. Neighborhood Analysis.....	8
b. Palimpsest Conceptual Model.....	10
c. Historical Housing Policies I: Home Owner's Loan Corporation.....	10
d. Historical Housing Policies II: Federal Housing Administration.....	12

II. Methodology

i. Constructing the Data.....	16
a. The Random Digit Dial Dataset.....	16
b. The Online Dataset.....	25
c. The HOLC Dataset.....	25
ii. Statistical Analysis of Data.....	27

III. Results.....32

IV. Discussion.....42

V. Works Cited.....50

I. INTRODUCTION

i. BACKGROUND

In the summer of 2017, a real-estate industry led effort to rename Harlem “SoHa,” incited protest and outrage from community members who saw the move as obscuring the neighborhood’s notable legacy as the center of African American art, culture and business in New York City. This renaming accompanied rising rent costs, the demolition of historical buildings such as the church where Malcolm X’s funeral was held, and the installation of new fixtures such as Whole Foods and clothing boutiques. These most recent developments are being done with a wealthier target demographic in mind than current Harlem residents, whose median income amounts to less than \$37,000 per year (Adams, 2016). These acts of gentrification not only stripped the neighborhood of its past, but threatened the ability of Harlem’s current minority residents to continue living in this area. In response to public outrage, lawmakers drafted a bill that proposed an intensive vetting process in order to rename a neighborhood, and penalties to real-estate brokers who advertised properties in neighborhoods whose boundaries and names they had essentially fabricated (Bellafonte 2017). The incident highlighted the ways in which contemporary conflict over neighborhood boundaries, their names, and the sense of personal identity tied up in these neighborhood characteristics are largely reminiscent of a past in which mortgage lending and housing restrictions segregated minority groups and those of lower socioeconomic status. In “SoHa” and elsewhere, there remains a sense of injustice perpetrated against minority populations living in neighborhoods which are constantly being redefined and re-configured by the real estate industry.

The direct interference into the composition of neighborhoods by stakeholders in the field of urban real-estate and development is hardly a new idea. One could even consider the practice of segregating neighborhoods, both racially and socioeconomically, through selective mortgage lending

to be an American tradition. This practice arose in the New Deal era and was perpetuated throughout the 20th century. Some of the earliest iterations of racially and socioeconomically biased mortgage lending can be seen in the 1933 Home Owner's Loan Corporation maps that grade neighborhoods on a scale of A – D, from the “best” neighborhoods to the most “hazardous,” as determined in part by the racial and socioeconomic composition of these areas (Nelson, Winling, Marciano, Connolly 2016). In fact, the practice of representing the most hazardous neighborhoods in red on these maps, is thought to be the origin of the term “redlining” which is still used to describe these types of lending strategies today. In 1934, the U.S. Federal Housing Administration released their first Underwriting Manual, which codified racial and socioeconomic prejudices and made segregation into a standard practice (Federal Housing Administration, 1934). While these practices were officially outlawed in 1986 with the Fair Housing Act, they have not vanished. As recently as 2017, the U.S. Justice Department sued KleinBank for redlining minority neighborhoods in Minnesota (Department of Justice, 2017). Later that year, a report by Colorado Public Radio asserted that discriminatory lending in Denver [was] fueling present day gentrification (Sakas, 2017), which echoes the sentiments of those fighting gentrification in Harlem. While the legacy of redlining is being felt in cities across the United States, it's difficult to directly measure how the lending practices of the past are impacting the neighborhoods of present day residents. Being able to identify and quantify some of the long-term effects of discriminatory lending is important because concrete evidence of this phenomenon is needed in order to better inform the actions and strategies of policy makers, researchers, activists and other stakeholders who are seeking solutions.

This study addresses the need to quantify the impact of historical redlining by analyzing perceived neighborhoods, neighborhoods which are self-defined by individuals. We have chosen the study area of New York City for the purpose of understanding how historical housing policies alter the way that present day residents view the area that they consider to be their neighborhood. In

addition to gathering information about the boundaries of an individual's neighborhood, we have also calculated important, and measurable, characteristics about these perceived neighborhoods, as well as demographic data regarding those we have surveyed. We aim to contribute to the present day conversation which is currently establishing links between historical housing policies and contemporary urban inequality. Additionally, we believe that the methods and procedures we use in this research can support the field of neighborhood analysis, which is an important area of inquiry for social science researchers who seek to examine human-centered data on a spatial scale. In order to understand how historical context can shape the perception that residents have of their neighborhood, we introduce the idea of understanding place through the palimpsest conceptual framework. These terms and concepts will be expanded upon in the upcoming paragraphs.

ii. LITERATURE REVIEW

Neighborhood analysis is the inquiry into how neighborhood units are constructed and implemented for applications in human-centered studies. This is an increasingly relevant field of research given the reliance that different social science fields have on neighborhood units. Neighborhood units, and how they are delineated, have an impact on the aggregation of spatial data. This concept is related to the modifiable areal unit problem (MAUP), which is defined in this context as the variation in the outcome of an analysis based on alternative groupings of neighborhood units at a constant spatial scale (Openshaw 1984). Theoretical considerations and approaches to neighborhoods have defined many different characteristics of neighborhoods. This study employs perceived neighborhoods, which are the cognitive constructs of an individual's environmental exposures (Coulton, Korbin, Chan, Su 2001). Perceived neighborhoods have been found to represent environmental exposures more comprehensively than territorial neighborhoods,

an alternative neighborhood definition that bases neighborhood units on higher level characteristics such as administrative boundaries or spatial homogeneity (Colabianchi, Coulton, Hibbert, McClure, Ievers-Landis & Davis 2014; Weeks, Hill & Stoler 2013).

Neighborhood units are integral to the many fields of social science which study how human-centered variables interact on a spatial scale. They are the spaces in which people live their everyday lives – an area that encompasses their place of residence, the shops they frequent, the parks in which they spend their free time, and the people with whom they interact with on a daily basis. Research into place-related identity finds that places such as neighborhoods, and the meaning and significance that they hold to individuals, are capable of influencing conceptualizations of self (Smith & Bender 2001). Neighborhoods also lend important insight into the environmental and social exposures of an individual (Basta, Richmond & Wiebe 2010), and they are frequently used to measure health and social variables on a spatial scale.

The way that neighborhoods have been defined play a central role in numerous studies in the fields of public health. For example, a study which investigated youth exposure to tobacco retailers found that measures of tobacco retailer density and proximity measures varied drastically across the two territorial and four perceived neighborhood types that were employed by researchers (Vallée & Shareck 2017). Another study found that when studying neighborhood influences on physical activity, overweight, and obesity in children and youths, perceived neighborhoods yielded less biased results than territorial neighborhoods, and were able to more accurately represent an individuals' environmental exposures (Colabianchi et al. 2014). If perceived neighborhoods are capable of producing information regarding important health and social outcomes, then it is necessary to evaluate how they are being defined, and then critically examine the kinds of assumptions that are being made about their delineations. These precautions are imperative because if researchers seek to serve the public in any formal capacity, such as through the influencing of policy-making decisions,

then researchers must be confident that any results that are produced are truly representative of the population and the process that they are seeking to describe.

In approaching the factors that contribute to neighborhood delineation, this article employs the term ‘palimpsest’. The original definition of palimpsest describes sixth century Western European manuscripts that have been reused, but upon which remnants of prior writings can be observed (Lyons 2011). However, this term has transcended its original usage, and has been used in the field of historical geography to describe the concept of landscapes and spaces as being comprised of different temporal layers (Vervloet 1986). Later applications of the palimpsest model evaluate the different ways that place can be interpreted by a variety of social groups and individuals who apply a unique combination of factors – such as cultural identity, lifestyle, and lived experiences – to their definition of place (Schein 1997). Place as a palimpsest can be summarized as a conceptual model that acknowledges how historical context, cultural differences, and a diversity of interpretations coalesce to form a multilayered landscape (Mitin 2010). We assert that this conceptual framework is applicable to neighborhood analysis and urban geography, especially as a lens with which to view how perceived neighborhoods are defined by individuals.

Our study seeks to examine how the historical housing policy known as mortgage redlining effects the way that individuals in the present day develop a mental map of their neighborhoods. Redlining is thought to have originated with the Home Owner Loan Corporation (HOLC) in 1933. This government-sponsored corporation was an element of President Franklin Roosevelt’s New Deal legislation and played a prominent role in restructuring home finance, reshaping the mortgage market, and recovering the American financial market after the Great Depression (Crossney & Bartelt 2005). From the corporation’s inception to the late 1940s when the HOLC focused their operations on the maintenance of loans, it had refinanced millions of homes and turned \$14 million in profit to the U.S. treasury. The consequences of redlining were not brought to light until 1938,

and it took until 1968 for segregation in housing to become illegal with the Fair Housing Act (Nelson, Winling, Marciano, Connolly 2016). Redlining describes the denial of services, in this case home loans, as well as access to different properties, based on a risk assessment that was constructed using the racial and ethnic composition of different neighborhoods, in addition to many other valuation factors such as housing age and style. The HOLC used maps, called Residential Security Maps, that marked neighborhoods with a high proportion of minority residents in red. Other HOLC documentation included neighborhood surveys which included racist sentiment and language (Nelson et al 2016). An example of an HOLC Residential Security Map is shown in Figure 1.

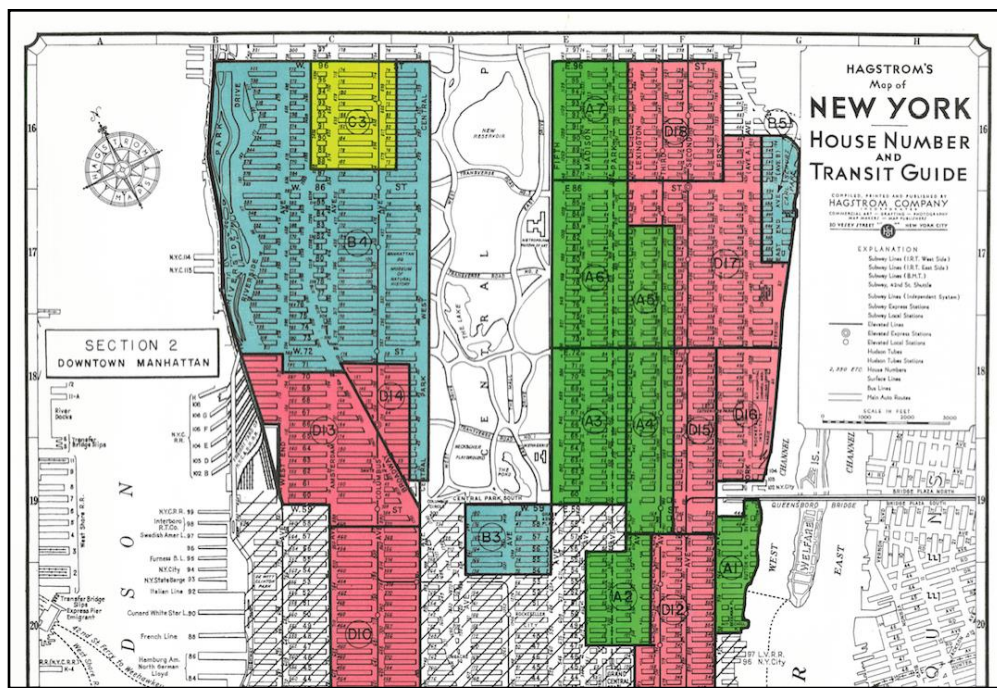


Figure 1: Excerpt from an HOLC Residential Security Map of Downtown Manhattan. Where green regions (Grade A) were the "Best" neighborhoods, blue (Grade B) were "Still Desirable" neighborhoods, yellow were "Definitely Declining" neighborhoods, and red were "Hazardous" neighborhoods in terms of the HOLC's risk assessment methods. These designations assisted property valuers and mortgage lenders in their determination of property value and mortgage lending risks (Mapping Inequality, 2016).

HOLC maps and surveys were re-discovered in the 1980s, and new inquiries into the corporation's role in applying racial and socioeconomic prejudices to real-estate appraising are

possible, especially through the application of geographic information science. While there is an ongoing debate surrounding whether or not the HOLC's policies implied that the United States Government was willfully complicit in the segregation of American neighborhoods, there is outstanding evidence that redlining and HOLC policies played a role in supporting existing racist biases and preventing different ethnic groups from having equal access to opportunities and wealth accumulation (Nelson et al 2016; Crossney & Bartnett 2005). Additionally, it is important to acknowledge that redlining is not the sole mechanism for segregation, and that despite the HOLC designation that certain regions received, there is some evidence that the actual distribution of mortgages had a limited racial bias (Hillier 2003; Crossney et al 2005). Given HOLC's ability to accurately represent historical racial biases, as well as influence the mortgage lending decisions of several private lenders during the formative New Deal era, HOLC designations are maintained as a fair measure of how historical neighborhoods were shaped by housing policies that reinforced and perpetuated ethnic and racial prejudices.

The HOLC maps also helped to inform, and were informed by, the Underwriting Manual for the Federal Housing Administration, which was first published in 1934 and revised several times over the years (Hillier, 2003). This guide was intended for real estate valuers of the Home Owner's Loan Corporation who were tasked with the appraisal of neighborhoods in 63 cities across the United States. Their ultimate goal was to establish a mortgage risk rating for different neighborhoods, through the valuation of dwellings, as well as through the rating

The Property:
Structural Soundness
Resistance to Elements
Resistance to Use
Livability and Functional Plan
Mechanical and Convenience Equipment
Natural Light and Ventilation
Architectural Attractiveness
Adjustment for Nonconformity
The Location:
Relative Economic Stability
Protection from Adverse Influences
Adequacy of Transportation
Need for Housing
Appeal
Sufficiency of Utilities and Conveniences
Adequacy of Civic, Social, and Commercial Centers
Level of Taxes and Special Assessments
Topography and Special Hazards
The Borrower:
Reputation
Attitude Toward Obligations
Ability to Pay
Prospects for Future
Past Record
The Mortgage Pattern:
Ratio of Loan to Value
Ratio of Debt Service to Rental Value
Ratio of Life of Mortgage to Economic Life of Building
Lowest Category Rating
Intermediate Category Rating
Highest Category Rating

Figure 2: The main features of the risk-rating system for mortgage lending as outlined by the FHA Underwriting Manual (1936)

of properties and locations (Federal Housing Administration, 1936). This calculation of mortgage risk took into account not only the architectural design and style of a home, but also a wide variety of neighborhood and locational characteristics. For the entire list of risk-rating factors, see Figure 2.

The FHA believed that “a most important group of factors which affect mortgage risk is the one which embraces the relationship between the physical property and the neighborhood in which it is located,” (Federal Housing Administration, 1936). This meant that valuers were encouraged to preserve the homogeneity, stability and “character” of a neighborhood and to incorporate the presence of “adverse influences” into their rating of risk (Federal Housing Administration, 1936). For example, geographical characteristics that were favorable for a neighborhood were natural or artificial barriers that could “prevent the infiltration of business and industrial uses, lower-class occupancy, and inharmonious racial groups.” (Federal Housing Administration, 1936). Additionally, valuers were encouraged to investigate the surrounding areas of a certain location in order to determine “whether or not incompatible racial and social groups [were] present” and to estimate “the probability of the location being invaded by such groups.” This was because “a change in social or racial occupancy often leads to instability and a reduction in values.” (Federal Housing Administration, 1936). Given these factors, zoning regulations were developed by the FHA to prevent such “instability” and “adverse influences” through the institution of zoning regulation recommendations that included the “prohibition of the occupancy of properties except by the race they were intended for,” and the homogenization of centers for community such as schools which “should not be attended in large numbers by inharmonious groups” (Federal Housing Administration, 1936).

HOLC grades and FHA regulations, which were both used to inform one another, as well as influence the practices of other private and government housing lenders throughout the 1900s, perpetuated racial segregation in urban areas. This was done through the way they valued property

and neighborhoods, which reinforced racial prejudices and led to the codification of these beliefs into the mortgage lending policies of the real-estate and appraisal industries. While the actual distribution of mortgages may not have not been less in areas with higher risk-ratings (Hillier 2003; Crossney et al 2005), it is undeniable that different racial and socioeconomic groups were barred from accessing property in different neighborhoods. These instances of neighborhood segregation greatly influenced the spatial and social mobility of minority groups and those with low-socioeconomic status, and greatly reduced the access that these people had to educational attainment and employment options that could have increased the possibility for their children and grandchildren to escape the confines of these housing policies. Additionally, because these practices officially took place up until 1968, it is possible that the patterns of neighborhood segregation have been perpetuated over time, and may still be reflected in the composition of neighborhoods in the present day.

Our locational focus is on New York City, which is subdivided into five boroughs, and contains over 250 neighborhoods whose boundaries are malleable, in that residents frequently define their exact boundaries in different ways, including in area, shape, extent and composition. Additionally, individual residents may define their own neighborhood differently from the 250 ‘defined’ neighborhoods of New York City. Currently, there are a number of New York City neighborhoods which are in flux, and changing in character and composition. Given the relevance of neighborhood definition, and the demonstrated importance of perceived neighborhoods in modern day policy-making, studying the influence of the historical housing policies that shaped earlier iterations of New York City neighborhoods is a way of understanding their ongoing legacy. This study examines this legacy by connecting two interrelated surveys regarding perceived neighborhood with 1930s redlining maps.

We are interested in placing perceived neighborhoods in the purview of the palimpsest conceptual model by investigating how the dimensions of neighborhoods may be impacted by both an individual's demographic traits and historical redlining. We will do this by constructing datasets using the results of a random digit dial survey and an online survey of perceived neighborhoods that were performed together in New York City. Additionally, we will calculate the area, perimeter, and compactness of these neighborhoods and perform a statistical analysis which we will examine whether or not demographic characteristics of an individual (age, gender, race, income, educational attainment, marriage status, and length of residency) result in different values for each of these variables. Next, we will assign each of these present-day perceived neighborhoods a HOLC grade based on what these neighborhoods were assigned in the past, and determine whether or not area, perimeter or compactness values are varying in relation to a specific HOLC grade. Lastly, we will assess whether or there is a relationship between which HOLC grade a perceived neighborhood is assigned, and what the race, income level and education attainment of an survey respondent living in that neighborhood. We hypothesize that there will be significant differences in the area, perimeter and compactness of an perceived neighborhood based on their individual demographic characteristics. Additionally, we hypothesize that there will be significant differences between the area, perimeter and compactness of neighborhoods which fall into different HOLC grades. Finally, we hypothesize that different HOLC grades contain significantly different levels of individuals, based on their race, income level and educational attainment. These hypotheses align with our palimpsest conceptual model, and if true, would support the theory that individuals from different walks of life perceive their neighborhoods in measurably different ways, as well as the theory that historical redlining practices impact the characteristics of contemporary perceived neighborhoods.

II. METHODOLOGY

i. CONSTRUCTING THE DATA

This section describes the data used in this analysis. The sources, data handling methods, and challenges that these data pose to the research findings will be covered. This project uses three different datasets: perceived neighborhoods of New York City obtained using a random digit dial (RDD) telephone survey, perceived neighborhoods of New York City obtained using an online Geographic Information System (GIS) tool, and HOLC neighborhoods digitized and geocoded by the Mapping Inequality Project. These datasets will be referred to respectively as the RDD perceived neighborhoods dataset, the online perceived neighborhoods dataset, and the HOLC dataset.

The RDD perceived neighborhood data was collected by the University of Pittsburg, and geocoded by myself under the advisory of Dr. Colleen Reid at the University of Colorado Boulder. This dataset originally consisted of 703 participants, surveyed throughout the summer season (June and September 2012), and the winter season (December and March 2012-2013). The process for acquiring survey participants consisted of RDD of both landlines and cell phones. Participation in this survey was possible after eligibility was determined and informed consent was given by the survey participants. The RDD surveys were led by trained administrators originating from the Survey Research Program at the University of Pittsburg. Survey participants consisted of residents from all five boroughs of NYC, and their distribution was approximately proportional to the populations of each borough.

Participants were asked to provide the names of four streets that formed the outer boundaries of their perceived neighborhoods. Additional locational information provided by survey

participants included the name of their neighborhood and their nearest cross-streets. The procedure for geocoding neighborhoods involved searching for bounding streets in the MyMaps (mymaps.google.com) application and using the line tool to draw a polygon. Polygons were stored in MyMaps, and downloaded in the keyhole markup language file format (.kml), which were then converted into shape files in ArcGIS 10.5. Hypothetical examples of perceived neighborhood data collected and geocoded for this project is shown below in Figure 3.

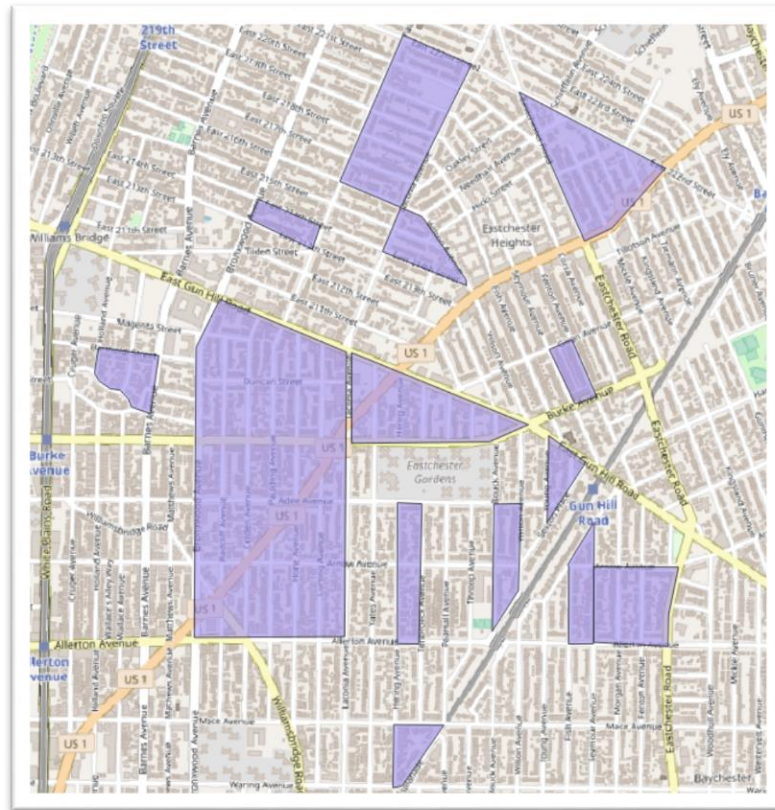


Figure 3. Hypothetical perceived neighborhood polygons geocoded using MyMaps and converted into polygons for analysis in ArcGIS 10.5.

The RDD survey format introduced many different opportunities for error and uncertainty in both the relaying and the interpretation of locational information over the phone. One main challenge that was common was stating as one of four neighborhood boundaries as landmarks, which represent as points or polygons on a map, as opposed to streets which represent as lines on a

map. This introduced uncertainty in the geocoding process, because it was difficult to decide how to connect from lines to points to create a neighborhood polygon as there are multiple ways to do so. A hypothetical example of this specific geocoding challenge is shown below in Figure 4.

Other challenges to geocoding were inaccurate street names or unknown street names, refusal to give street names, or giving one or more boundaries that did not intersect. Errors that were common on the part of the survey administrator, and which were a source of uncertainty when geocoding, were interpretation-based, and included interpretation of misspellings and typos, difficulty understanding local terminology for boundary names, and the phonetic spelling of unfamiliar boundary names.

Dealing with these different types of challenges and uncertainties in the raw data was an important part of our preliminary analysis when geocoding RDD perceived neighborhood polygons. We created a framework for acknowledging these, based on the error type and the level of uncertainty, to ensure that our methodology was consistent. This framework was developed based on an initial subset of polygons, based on the different error types that we encountered. After fully developing our system for dealing with different error types and uncertainty levels, we started over, editing and updating previously geocoded neighborhoods as necessary, in order to ensure that all neighborhoods conformed to the same set of geocoding standards.

An important factor to address when geocoding polygons, was how to draw polygons that included one or more boundaries that was a point or polygon location, instead of a street. Examples of a point boundary would be an intersection, or a local landmark such as a statue. Examples of polygon locations would be a park, a school, a business or a neighborhood. Survey participants were asked to provide streets as boundaries, but it was common to receive boundaries in either of these point or polygon forms. When geocoding polygons that included one or more of these boundary types, we adhered to the following rules. When using point locations to form one or more of a

polygon's boundaries, we drew a line from the point, directly to the connecting boundary lines, as the crow flies, as opposed to following street networks. This decision was based on our inability to make an assumption about which street the participant would have used to connect a point location to the other boundaries that they defined. When dealing with a polygon boundary, such as a park (e.g., Central Park), a body of water (i.e. Grasmere Lake) or a business (e.g., IKEA Brooklyn), we chose not to encompass the polygon within our neighborhood, but instead selected the edge of the polygon that was nearest to the other given boundaries as the definitive boundary. An example of this concept is shown below, in Figure 4.

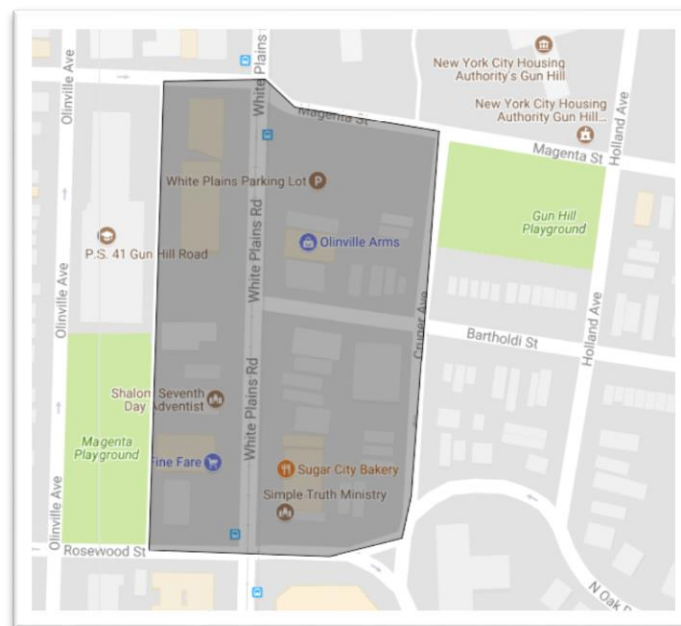


Figure 4. This figure shows an example of how a neighborhood was geocoded if a survey participant provided researchers with two points and two lines in order to define their neighborhood boundaries. The two points in this example are Magenta Playground and Gun Hill Playground, while the two line boundaries are Magenta St. and Rosewood Street.

In addition to standardizing our geocoding methods when dealing with point and polygon boundaries, we established four different types of errors based on the most common and consistent errors that we saw throughout our dataset, assigned these specific numbers which corresponded to

the degree of uncertainty about the polygon associated with that error type, and took detailed notes regarding why these numbers were assigned for each polygon (Table 1). Within each type of error, we assigned numbers 1-4 to denote the level of uncertainty. Polygons could possess more than one type of error but only one level of uncertainty for a given error type (i.e., polygons could fall into more than one row in Table 1 but not in more than one column for any given row). Levels of uncertainty were also established to serve as guidelines for geocoding, and to describe and standardize different types of errors that we observed throughout the RDD perceived neighborhood dataset. Error type and uncertainty levels were used as a guideline to inform the geocoding process, and set limits to how much uncertainty was acceptable in a polygon, as opposed to when a polygon was unable to be created based on a violation of these limits.

Polygons were assigned type one if there were no errors that introduced uncertainty into the geocoding process. Polygons were assigned type two errors based on spelling, prefix and naming issues with the boundaries provided. Type three polygons were polygons that were able to be created but only through the omission of one out of four of the given boundaries, due to the fourth boundary not intersecting the other three, or the fourth boundary being contained within the polygon. Type four polygons were assigned their designation based on the reliance of extensions when connecting one or more of the boundaries to the others in order to complete the polygon. This entailed extending streets boundaries beyond where they terminated, in order to ensure that the resulting polygon had contiguous boundaries that intersected with one another.

In order to visualize the different uncertainty types and levels, Table 1 (below) was created as a summary and reference.

Table 1. Levels of Uncertainty and Corresponding Error Types

<i>Error Types</i>	<i>Level of Uncertainty</i>				
	1	2	3	4	5
1. No uncertainty	<i>Boundaries given create a polygon without uncertainty.</i>	<i>NA</i>	<i>NA</i>	<i>NA</i>	<i>NA</i>
2. Boundary names	<i>Minor spelling error in street name.</i>	<i>Directional prefix of street changed to geocode a complete polygon.</i>	<i>Street suffix changed on boundary to geocode a complete polygon.</i>	<i>Given street continues to make a polygon but name of street changes.</i>	<i>Name of one street is different but is inferred given context clues to geocode a complete polygon.</i>
3. Three out of four boundaries	<i>Three bounds form a polygon, fourth does not. Made polygon out of three sides.</i>	<i>Three bounds form polygon, fourth is contained within polygon. Made polygon out of three sides.</i>	<i>Three bounds form polygon, one extends from polygon. Extension not incorporated into polygon.</i>	<i>NA</i>	<i>NA</i>
4. Extensions	<i>Given streets extended by less than 0.5 miles to form polygon.</i>	<i>Streets extended by more than 0.5 miles and less than 1 mile to form polygon.</i>	<i>Streets extended by more than 1 mile but less than 1.5 miles to form polygon.</i>	<i>NA</i>	<i>NA</i>

Assigning a neighborhood polygon to a type one error was a way for us to classify polygons that had no errors in their boundary data. These polygons conformed to all of our survey guidelines and possessed boundaries that intersected logically.

Type two errors were prominent in the raw data for our RDD perceived neighborhoods, and introduced several different kinds of uncertainty into our resulting polygons depending on which level of uncertainty they fell into. We believe that the main reason for the high prevalence of Type two errors was that this error type encompassed spelling errors, lack of specificity in terms of street cardinality, general typing errors, as well as issues in communication between researcher and survey participant. Within type two errors, there were five different categories of uncertainties.

Level one uncertainty within the type two error category included RDD perceived neighborhood data in which there was a minor spelling error. An example of this would be “Centrel Park” instead of “Central Park.” These uncertainties were insignificant and did not hinder the ability for a polygon boundary to be successfully interpreted.

Within type two errors, level two uncertainties were when a cardinal direction associated with a street boundary had to be changed or added. For example, a boundary would be listed as “14th Street” but when viewed in context with the rest of the given boundaries, it became clear that “West 14th Street” was the only possible street that one could use to enclose a neighborhood polygon. These provided more uncertainty than type two, level one in terms of geocoding because in some cases, multiple boundary names with different cardinalities would be present in close spatial proximity with one another. In these cases, other factors were taken into account when geocoding the polygon, such as assessing which of these streets intersected with the rest of the given boundaries while introducing the least amount of additional uncertainty.

Type two, level three uncertainty involved boundary suffixes, and encompassed any street boundary in which a suffix had to be added or changed. An example of this is “Riverside Parkway” vs “Riverside Drive.” In this example, while “Riverside Parkway” was given as a boundary, when making the polygon, it was most logical to use “Riverside Drive” instead, because this was the only boundary that completed the polygon in the context of the other given boundaries. For these cases, the boundary definition was updated with the appropriate suffix. In other cases, “Riverside Parkway” would not exist whereas in others, both existed, but one was too far away to incorporate into the polygon. Similar to level two uncertainties, there were scenarios in level three uncertainties in which multiple streets with the same base name but different suffixes would exist in close proximity, which then required an overall assessment of which street would make sense in context

of the other boundaries and we created a neighborhood polygon that adhered to the other rules that we established.

Level four uncertainties were cases within the type two error group in which a named boundary intersected with the rest of the polygon but changed name while remaining a continuous street. An example of this uncertainty level is Rockaway Avenue in Brooklyn, which merges into Rockaway Parkway as it extends south. If we were given the boundaries Rockaway Avenue and Farragut Road, which lies south of Rockaway Avenue, we could not intersect these two boundaries without Rockaway Avenue merging into Rockaway Parkway. In these cases, uncertainty could be considered higher than the previous uncertainty levels because the streets that the given boundaries merged into were not explicitly stated as boundaries by the survey participant, though their inclusion was necessary in order to connect the rest of the boundaries.

Level five uncertainties, the final uncertainty level for type two errors, consisted of instances when a boundary that was used to create a polygon was inferred based on information from a survey participant that was incorrectly spelled or interpreted by the survey administrator. There were many instances in which the name of a boundary was difficult to ascertain as written, and could not be found using the Google Maps search. However, when the other given boundaries were entered, the boundary could be inferred to be a certain location based on a street that completed the polygon, and either sounded like or was spelled similarly to what had been originally entered in the data. Some of these may have been different spellings with the same pronunciation, but for others, it was a similar sounding word and the one given does not exist. The prevalence of these errors could be attributed to the fact that accents and unfamiliar-sounding place names could make the interpretation of a stated boundary difficult to understand on the part of the survey administrator. An example of this error would be a given boundary of “Urban Avenue,” a location that could not be found in the context of the other boundaries. However, when scanning the area, the boundary

“Irving Avenue” could be identified and used to complete a neighborhood polygon with the other given boundaries. Thus, it was concluded that the survey administrator misheard the boundary name, and the participant had intended to include “Irving Avenue” as one of their boundaries. This uncertainty level was highest of all the type two uncertainty levels given that we were selecting boundaries that were not explicitly stated by the survey participants, but that made sense in place of boundary names that could not be identified.

Where type two errors encompass boundary naming, type three errors include the situation in which three out of four given boundaries form a polygon, and the fourth is excluded due to different circumstances. Within the type three errors, three different uncertainty levels were established, with the same amount of severity, meaning that while there were three different cases of type three errors, they all had relatively the same impact on the uncertainty of the resulting polygons. Level one uncertainty described a situation in which three out of the four boundaries formed a polygon, but the fourth boundary did not connect to the rest of the polygon in any way. In this case, the disconnected boundary was not incorporated into the polygon. Level two uncertainty described a situation in which three out of the four boundaries formed a polygon, but the fourth boundary lay inside of the resulting polygon. In this case, we kept the in-tact polygon that encompassed that fourth boundary. Level three uncertainty described a situation in which three out of the four boundaries formed a polygon, with the fourth boundary intersecting, but extending beyond the limits of the polygon. In this case, the fourth boundary was not factored into the polygon as a boundary.

Type four errors indicated that, in order to make a complete polygon, one or more of the boundaries provided had to be extended a specific distance in order to intersect with the other given boundaries. For example, a street would be given that would end abruptly, while the other three

boundaries connected. In order to form a polygon, the street boundary would be extended by no more than 1.5 miles in order to create a complete neighborhood polygon.

We were unable to geocode polygons if boundaries did not connect to form a polygon despite considerations made by our error types and uncertainty levels. In addition to providing insight into how the geocoding of these polygons took place, levels of uncertainty were also factored into the data analysis portion of our study in order to better inform our results.

The online perceived neighborhood dataset was obtained with permission from the University of Pittsburgh's Survey Research Program, and consisted of neighborhood polygons collected using an online survey. 575 participants were successfully surveyed using this method. The online survey was self-administered among participants of a voluntary, standing survey panel of NYC adults (Survey Sampling International, <http://www.surveysampling.com>, MyOpinions Pty Ltd., Shelton, CT, USA). This survey implemented a novel GIS tool with which participants could define/draw a polygon that encompassed their perceived neighborhood. The GIS tool used a base map modelled on Google Maps due to the widespread familiarity and utilization of this platform. The GIS tool also included detailed instructions intended to guide survey participants and ensure that all participants, regardless of technical ability, could access the tool. This dataset was already created before the start of this project, and we did not have to do any geocoding of these data.

This study also used digitized HOLC maps and survey information from the Mapping Inequality project (Nelson et al 2016). Data related to their activities – including color-coded security maps and data regarding the ethnic composition of the neighborhoods they surveyed – has been provided as open data by the Mapping Inequality project. This project recovered HOLC data from the National Archives and released it in the form of digitized maps and geocoded neighborhood polygons. Four designations, assigned grades A-D, were given to subdivisions of surveyed regions. Grade A described an area with the least amount of monetary risk for the lending agency, which

accounted for a low amount of ethnic minority groups, and D described an area with the most amount of risk, taking into account the high proportion of ethnic minority groups, with B and C falling in between these two. HOLC neighborhoods are available for all five boroughs of NYC, and the polygons have been geo-rectified in order to conform to modern day maps of NYC. The geocoded neighborhoods for the HOLC dataset attributes included HOLC designations, and the meta-datum included digitized versions of the original surveys from the year 1933, which detail the ethnic composition of each surveyed neighborhood based on estimations performed by the original HOLC valuers.

In order to assign HOLC Grades to our perceived neighborhood data, a GIS analysis was performed using the spatial analyst tools available in ArcGIS 10.5. The aim of this analysis was to establish which HOLC neighborhood had the highest percent area overlap with a perceived neighborhood from our RDD and online perceived neighborhoods datasets. For perceived polygons that were completely contained within one HOLC neighborhood, the assignment of HOLC neighborhood was clear. In the event that more than one HOLC neighborhood was overlapping, the HOLC grade of the HOLC neighborhood with the highest percent overlap was assigned to that polygon. In order to obtain accurate measurements of polygon area, all Online/RDD perceived neighborhoods and HOLC neighborhoods were projected to NAD 1983, UTM zone 18N. The ArcGIS tool ‘Spatial Join’ was used to join the attributes of the all perceived neighborhood datasets (online, winter RDD and summer RDD) with the HOLC neighborhood data. When using the join tool, it was important to use ‘join’ as the merge rule for the HOLC grade attribute in order to output a joined feature class with an HOLC grade attribute table field that displays all the HOLC grades that overlap with each Online/RDD perceived neighborhood.

Using the output from ‘Spatial Join’, a shape file that included Online/RDD perceived neighborhoods and all HOLC neighborhoods which overlap them, the tool ‘Tabulate Intersection’

was used to assess the percent area of HOLC perceived neighborhoods that overlapped with the Online/RDD perceived neighborhoods along with their corresponding HOLC grades. With this information, an HOLC grade was assigned to each Online/RDD perceived neighborhood based on the HOLC neighborhood which overlapped the most in terms of area.

ii. STATISTICAL ANALYSIS OF DATA

The four main aims of this data analysis were to determine whether or not there are differences in neighborhood polygon characteristics between the Winter and Summer RDD data frames, assess if there are differences in polygon characteristics between the RDD combined dataset and the online panel, investigate whether or not there are differences in perceived neighborhood characteristics by demographic variables, and determine if there are differences in perceived neighborhood characteristics by historical HOLC designation.

The first goal, to compare our Summer RDD perceived neighborhood dataset with our Winter RDD dataset, was performed in order to determine if these datasets could be merged into one RDD dataset, as well as to determine if our methodology yielded data that was consistent when replicated. The second goal of our project, to examine differences between the RDD and Online perceived neighborhoods, was performed in order to gain insight into whether these two methods produced similar perceived neighborhoods. We calculated perimeter, area, and compactness of the neighborhoods to compare the dimensions of all perceived neighborhood data frames.

We were also interested in whether there were differences in polygon characteristics by demographic variables. We wanted to know whether or not the age, race, employment, income, education level, marriage status, and gender of a survey participant affected the dimensions of their perceived neighborhood. The last aim of our project was to determine whether or not historical

HOLC grades impacted the dimensions (area, perimeter and compactness) of perceived neighborhoods, and whether present day perceived neighborhoods showed a demographic (age, race, employment, income, education level, marriage status, and gender) relationship with their historical HOLC grade. All analysis was done using the R Statistical Software, specifically the R Graphics Package, R base functions, and the ‘pander’ Package (R Core Team, 2016). The GIS software used in this research was ArcGIS 10.5 (Esri, 2016) and QGIS 2.18.2 (QGIS Development Team, 2018).

The first aspect of the data analysis was to determine the dimensions of all of our perceived neighborhood polygons. Our three neighborhood dimension variables – area, perimeter and compactness – were calculated in QGIS, after first projecting all of the neighborhoods into NAD 1983, UTM zone 18N. Perimeter and area were determined using the field calculator capability, and was calculated in km and km², respectively. Compactness is defined as,

$$Compactness = \sqrt{area/area_2}, \quad (Eq. 1)$$

where *area* is the area of the perceived neighborhood polygon, *area₂* is defined as the area of a circle with the same perimeter as the neighborhood polygon (O’Sullivan & Unwin., 2003). This ratio compares the area of a neighborhood polygon to the area of a circle (the most compact shape) with the same perimeter. The closer the perceived polygon is to a circle, the closer this value will be to 1. The less compact, the closer it would be to zero.

After calculating the neighborhood dimensions, the second aspect of data analysis was to review the summary statistics of our Summer RDD, Winter RDD and Online perceived neighborhood datasets in order to assess how many neighborhoods were successfully geocoded, what the dimensions of these neighborhoods were, and how these neighborhood dimensions

compared with one another both within the different datasets and between the datasets. The procedure for doing this involved calculating the mean, median, standard deviation, minimum value, maximum value and interquartile ranges for each neighborhood dimension.

Additionally, we visualized the data using density plots in order to gain insight into their distributions. Statistical transformations, both log and cubic functions, were performed on skewed neighborhood dimension data, which were chosen for their ability to normalize right and left skewed data, respectively. The log function was applied to the area and perimeter datasets, due to their right skewedness, and the cubic function was applied to the compactness data as a result of the left skewed distribution.

To assess the possibility of combining the summer and winter RDD datasets into one dataset for further analysis, we performed unpaired t-tests on normalized dimensional variables to evaluate whether there were significant differences between the datasets. If there were no statistically significant differences between the dimensional datasets of the Summer and Winter RDD perceived neighborhoods, we would proceed to merge the datasets for further analysis, as opposed to dealing with each dataset separately.

After comparing the Summer RDD and Winter RDD neighborhoods, we also performed a similar statistical analysis in order to compare our combined RDD datasets to the online dataset.

After analyzing our datasets individually, we sought to also gain insight into the demographic characteristics that could be related to how people reported their perceived neighborhoods, as a function of our three dimensional variables. The individual-level demographic variables that we were interested in for this portion of the analysis were gender, age, education level, income level, race, marriage status, employment status, and length of residency. Before analyzing these demographic variables, we chose to create groups within each variable. These bins were treated as categorical variables, and an ANOVA test was performed in order to compare the means of these levels of the

categorical variables as a function of area, perimeter, and compactness separately. In order to determine which specific variables showed significant relationships with one another, we performed a Tukey Honest Significant Difference (HSD) test.

The individual level demographic variables were categorical, and were grouped into bins. Our statistical analysis, described below, employed these groupings, which are elaborated upon below. In some cases, we grouped the original bins established by the survey into a smaller number of bins.

For the demographic variable of gender, participants were asked to identify themselves as either male or female. Group 1 consists of survey participants who self-identify as male, and Group 2 consists of survey participants who self-identify as female.

Age was categorized into four age groups: Group 1 (18-30yrs), Group 2 (30-50yrs), Group 3 (50-65yrs), and Group 4 (65+yrs).

The original survey contained two questions regarding racial identification, and asked survey participants to respond to whether or not they identified as Hispanic, as well as what race they identified as. Hispanic respondents were assigned Group 1, and non-Hispanic respondents were assigned Group 2. The variable of race was determined based on the responses to survey questions that allowed participants to self-identify as White, Black, Asian/Pacific Islander, Native American and other. Additionally, participants that identified solely as 'Hispanic' are also included in this demographic variable as Hispanic. Participants were allowed to select all races that they identified with, resulting in numerous participants identifying as mixed race. Based on their self-identification, five race categories were established: Group 1 (White), Group 2 (Black), Group 3 (Asian/Pacific Islander), Group 4 (Hispanic), Group 5 (Mixed Race, Native American, Other).

The variable of education was determined based on the responses to survey questions that allowed survey participants to select their highest level of educational attainment, which originally

fell into 10 different categories. For our analysis, these education groups were further divided into four groups. These groups were: Group 1 (Eighth grade or less, some high school), Group 2 (High School/GED, Trade/Vocational/Nursing School, Associate's degree), Group 3 (Bachelor's Degree), Group 4 (Master's degree, professional degree, Doctoral degree).

The variable of income was determined based on the responses to survey questions that allowed participants to select the income bracket that applied to themselves. Based on their self-identification, three income categories were established: Group 1 ($25000 \leq \text{income} < 46000$), Group 2 ($45000 \leq \text{income} < 93000$), and Group 3 ($93000 \leq \text{income}$).

The variable of employment was determined based on the responses to survey questions that allowed participants to select the employment type that applied to themselves. Based on their self-identification, five employment categories were established: Group 1 (Employed full-time, employed part-time, self-employed), Group 2 (Out of work and looking for work, out of work and not looking for work, unable to work), Group 3 (Student), Group 4 (Retired), Group 5 (Homemaker).

The variable of marriage was determined based on the responses to survey questions that allowed participants to select the marriage status that applied to themselves. Based on their self-identification, four marriage status categories were established: Group 1 (Married), Group 2 (Separated/Divorced), Group 3 (Widowed), Group 4 (Never Married/Single).

The variable of length of residency was determined based on the responses to survey questions that allowed participants to select their length of residency based on four groups containing time-interval brackets in years. These groups were: Group 1 (Less than a year), Group 2 (1 year to just under 5 years), Group 3 (5 years to just under 10 years), Group 4 (10 years or more).

Using the HOLC grades that were assigned to each of our perceived neighborhoods in the Combined and Online datasets, we then turned our analysis towards looking for trends in area, perimeter and compactness within our HOLC neighborhood categories by performing an ANOVA

and a Tukey HSD test. Additionally, if no HOLC Grade polygons overlapped with our perceived neighborhood polygons, we did not assign a grade to the neighborhood, and they were not included in the HOLC analysis.

Because HOLC grades were created based on race, income and other socio-economic indicators from the 1930s, we wanted to know if our survey respondents had similar characteristics as those designations would have implied in the 1930s. We focused on the variables of race, income, and education level for this analysis. These variables remained grouped in the same manner that they were for our previous analysis. We then created two-way tables for each demographic variable and HOLC grade pairing, and performed chi-square tests on each separate HOLC grade and demographic variable category pair. Chi-square tests were performed in order to investigate whether or not there were statistically significant differences between the demographic variable groups in each HOLC grade. In addition to chi-square tests, the data were also analyzed using ANOVA and Tukey HSD tests in order to establish whether or not there was a significant difference between the means of area, perimeter and compactness as a function of their HOLC class.

III. RESULTS

The results for our summary of the Summer and Winter RDD datasets, as well as our online dataset are reported below, in Tables 3, 4 and 5. Table 3 shows the geocoding success rate for all of the neighborhood data received for both the Summer and Winter datasets. This table shows the number of individuals for which neighborhood boundary data were originally collected by the survey administrators for the Summer and Winter RDD datasets, as well as how many neighborhoods were able to be geocoded using our methodology.

Table 3. Summer RDD, Winter RDD, and Online Geocoding Success Rates

<i>Dataset</i>	<i>Total No. of Neighborhoods</i>	<i>No. of Successfully Geocoded Neighborhoods</i>	<i>Percent Geocoded</i>
<i>Summer RDD</i>	350	211	60.29%
<i>Winter RDD</i>	353	238	67.42%
<i>Online*</i>	856	575	67.17%

*The online panel was geocoded by other researchers.

Table 4 details the number of perceived neighborhood polygons that fell into each uncertainty level by dataset from our geocoding of the two RDD perceived neighborhoods datasets.

Table 4. Combined RDD Error Type Counts

<i>Dataset</i>	<i>Level 1</i>	<i>Level 2</i>	<i>Level 3</i>	<i>Level 4</i>	<i>Multiple Levels</i>
<i>Summer RDD</i>	91	40	33	30	17
<i>Winter RDD</i>	67	77	19	29	46

Table 5 shows the summary statistics for the three neighborhood dimension variables for all datasets. These variables were area (km²), perimeter (km) and compactness.

Table 5. Summary Statistics for Combined RDD and Online Dimensional Variables

<i>Dataset</i>	<i>Variable</i>	<i>Min</i>	<i>Median</i>	<i>Mean</i>	<i>Max</i>	<i>Std. Deviation</i>
Summer	<i>Area (km²)</i>	0.005	0.417	1.370	16.860	2.208
	<i>Perimeter (km)</i>	0.301	3.217	4.051	17.450	3.358
	<i>Compactness</i>	0.465	0.797	0.7832	0.8968	0.088
Winter	<i>Area (km²)</i>	0.007	0.530	1.342	14.590	2.076
	<i>Perimeter (km)</i>	0.392	3.377	4.222	20.050	3.432
	<i>Compactness</i>	0.263	0.804	0.782	0.899	0.098
Online	<i>Area (km²)</i>	0.00014	0.85350	1.54600	33.95000	2.387663
	<i>Perimeter (km)</i>	0.05804	4.12600	4.60700	23.09000	3.14996
	<i>Compactness</i>	0.0298	0.8352	0.7782	0.8728	0.9409

The results for the comparison between our Summer and Winter RDD datasets, are shown in Table 6, and contain the results of our Welch Two Sample t-tests for each dimensional variable. Before performing the t-tests, it was necessary to transform the Summer and Winter area and perimeter data using a natural logarithmic function due to the positively skewed distributions of the pre-transformed data. This transformation was chosen to better conform to the assumption of a normal distribution when conducting parametric tests. The compactness data for the Summer and Winter datasets showed a negative skew, and was normalized using a cubic function in order to produce a distribution that was closer to normal before executing any t-tests.

Table 6. T-Test Results Comparing Summer RDD, Winter RDD and Online Datasets by Neighborhood Dimension Variables

	<i>Summer RDD and Winter RDD</i>		<i>Combined RDD and Online</i>	
<i>Variable</i>	<i>t-statistic</i>	<i>p-value¹</i>	<i>t-statistic</i>	<i>p-value¹</i>
Area	0.7995	0.4245	-1.346	0.1785
Perimeter	0.8502	0.3957	2.246	0.0249*
Compactness	0.0552	0.956	-0.5374	0.5911

1. Where (*) denotes significance when $\alpha < 0.05$

Because there were no significant differences by area, perimeter, or compactness between the Summer and Winter RDD datasets, these two datasets were combined into one dataset to be compared with the Online datasets.

The Online dataset showed differently skewed distributions for the dimensional variables than the Summer and Winter RDD, therefore, the Welch's t-tests to compare the Combined RDD dataset to the online dataset were done on non-transformed data in order to ensure that the comparisons made were the most accurate. Table 6 shows that the Online perceived neighborhoods were significantly different from the RDD perceived neighborhoods for perimeter, but not for area and compactness.

After the datasets for the Combined RDD Dataset and the Online Dataset were compared with one another, these datasets were merged for the following analysis. For future reference, the merged Summer and Winter RDD Dataset will be referred to as the Combined RDD Dataset. In this analysis, demographic characteristics of our study population were evaluated for any possible relationships with our neighborhood dimensional variables. These demographic characteristics were gender, age, length of residency, marriage status, education status, income, employment status, and race. Additionally, we used the transformed versions of our neighborhood dimension datasets. For

neighborhood area and perimeter, which showed an overall right skewed distribution in the Combined RDD/Online Dataset, we used a natural log transformation, whereas for neighborhood compactness, we used a cubic transformation to transform the distribution of the dataset to better conform to the assumption of a normal distribution when conducting parametric tests.

In order to compare the mean values of our dimensional variables of area, perimeter, and compactness in terms of our demographic variables we used both an ANOVA test and a Tukey HSD? test on our Combined RDD and Online datasets. Exception to using ANOVA tests to compare means were when demographic variables had only two groups – such as Gender (Male/Female) and Hispanic Identity (Yes/No). The association of these variables with our dimensional variables were evaluated using a t-test. Additionally, due to the significant difference in perimeter values between the Combined RDD dataset and the Online Datasets, perimeter was evaluated with demographic data separately for the Combined RDD datasets and the Online Datasets. The results of these tests are summarized in Table 7. Additionally, the results of the Tukey HSD? tests are explained when significant results were seen.

Table 7: Resulting P-Values of ANOVA Tests Comparing Mean Dimensional Variables with Categorical Demographic Variables

	<i>P-Values of ANOVA and T-Tests Comparing Dimensional Variables of Demographic Variable Group Means¹</i>								
<i>Dimensional Variables</i>	<i>Gender</i>	<i>Age</i>	<i>Hispanic Identity</i>	<i>Racial Identity</i>	<i>Income</i>	<i>Employment Status</i>	<i>Marriage Status</i>	<i>Length of Residency</i>	<i>Education</i>
<i>Area (RDD/Online)</i>	0.0967	0.5800	1.52e-05*	0.0056*	0.1230	0.0069*	0.4510	0.3020	0.0415*
<i>Perimeter (RDD)</i>	0.2519	0.8990	0.0659	0.6520	0.6520	0.3420	0.3960	0.4120	0.412
<i>Perimeter (Online)</i>	0.0087*	0.2280	2.01e-05*	0.2280	0.2360	0.0143*	0.9750	0.3740	0.142
<i>Compactness (RDD/Online)</i>	0.0523	0.763	0.3238	0.0016*	0.0058*	0.7320	0.6160	0.0750	0.0004*

1. Where (*) denotes significance when $\alpha < 0.05$

2. RDD/Online indicates that all perceived neighborhoods were combined together, whereas for perimeter RDD and online were evaluated separately.

There were several instances in which our p-values denoted that there was a significant difference in dimensional means between demographic groups.

For the ‘Racial Identity’ demographic variable, there were significant differences in mean area and mean compactness. The Tukey HSD test showed that differences in mean area existed specifically between those who identified their race as ‘White’ and ‘Black’ (p-value = 0.038) as well as ‘White’ and ‘Mixed Race’ (p-value = 0.053). For mean compactness, the Tukey HSD test showed that differences in mean compactness existed only between ‘White’ and ‘Black’ participants (p-value = 0.0025).

Our ‘Income’ demographic demonstrated a statistically significant difference in mean compactness between income groups. These differences existed specifically between the income bracket (25000 ≤ income < 46000), and income bracket (\$93000 ≤ income) (p-value = 0.004).

For the 'Employment Status' demographic variable, there were significant differences in mean area, mean perimeter for the online dataset specifically, and mean compactness. The Tukey HSD test showed that differences in mean area existed between the 'Employed full-time, employed part-time and self-employed' and the 'Out of work and looking for work, out of work and not looking for work, and unable to work' groups ($p\text{-value} = 0.006$). Additional differences in mean area existed between the 'Student' and 'Out of work and looking for work, out of work and not looking for work, and unable to work' groups ($p\text{-value} = 0.0213$). Mean perimeter for the online dataset differed between the 'Employed full-time, employed part-time and self-employed' and the 'Out of work and looking for work, out of work and not looking for work, and unable to work' groups ($p\text{-value} = 0.0280$). Additional differences in mean perimeter existed between the 'Student' and 'Out of work and looking for work, out of work and not looking for work, and unable to work' groups ($p\text{-value} = 0.0489$). Lastly, the mean compactness variable differed between the 'Employed full-time, employed part-time and self-employed' and the 'Student' groups ($p\text{-value} = 0.004$).

For the 'Education' variable, we found that there were statistically significant differences in mean area and compactness. For mean area, our Tukey HSD Test showed that there was statistically significant difference between the 'bachelor's degree' group and the 'eighth grade or less/some high school' group. ($p\text{-value} = 0.035$) as well as between the 'master's degree/professional degree/doctors degree' group and the 'eighth grade or less/some high school' group ($p\text{-value} = 0.015$). For compactness, our mean differences existed between the 'master's degree/professional degree/doctors degree' group and the 'high School/GED/trade/vocational/nursing associate's degree' group ($p\text{-value} = 0.026$) and the 'bachelor's degree' group and the 'high school/GED/trade/vocational/nursing/associate's degree' group ($p\text{-value} = 0.0007$).

The last analysis that was performed was our comparison of neighborhood dimensions between perceived neighborhoods with a specific HOLC designation. The distribution of HOLC grades is shown in Table 8.

Table 8. HOLC Neighborhood Grades for the Perceived Neighborhood Datasets

	<i>Total Perceived Neighborhood HOLC Grade Distribution</i>	
<i>HOLC Grade</i>	Count	Percentage of Total
A	21	2.15%
B	198	20.3%
C	418	42.8%
D	339	34.7%
Totals	976	100%

The results of the ANOVA tests for HOLC designation of perceived neighborhoods by dimensional variables are shown below, in Table 9. There were no significant results, and thus, the results of the Tukey HSD test are not expanded on.

Table 9: Resulting P-Values of ANOVA Tests Comparing Mean Dimensional Variables with Perceived Neighborhood HOLC Grades

<i>Dimensional Variables</i>	<i>p-value</i>
<i>Area RDD/Online</i>	0.412
<i>Perimeter RDD</i>	0.495
<i>Compactness RDD/Online</i>	0.757

Lastly, in order to analyze the distributions of different demographic variables within the HOLC neighborhoods, we placed the four HOLC neighborhoods and a corresponding demographic variable into a two-way table, and ran a chi-square test on the data contained in these tables to determine whether there are differences in the distribution of these variables between the different neighborhood classes. The chi-square tests were performed on each separate column within the two-way tables in order to assess whether there were significant differences in the distribution of different demographic variable groups between HOLC grades. The two-way tables for three demographic variables used in this analysis comprise Table 10, Table 11 and Table 12. The results of these chi-square tests performed on these tables are shown in Table 13.

Table 10. Two-Way Table of ‘Race’ Demographic Variable Groups and Counts of HOLC Class Within Each Group

	‘Race’ Demographic Variable Categories And HOLC Grade Counts Within Perceived Neighborhood Dataset				
HOLC Grade	White	Black	Asian/Pacific Islander	Hispanic	Mixed Race, Native American, Other
A	18	0	1	0	2
B	119	47	16	3	13
C	221	88	45	10	52
D	161	105	24	8	41

Table 11. Two-Way Table of ‘Income’ Demographic Variable Groups and Counts of HOLC Class Within Each Group

	‘Income’ Demographic Variable Categories And HOLC Grade Counts Within Perceived Neighborhood Dataset		
HOLC Grade	(\$25000 <= income <\$46000)	(\$45000 <= income <\$93000)	(\$93000 <= income)
A	3	7	11
B	72	67	51
C	185	148	72
D	182	91	57

Table 12. Two-Way Table of ‘Education’ Demographic Variable Groups and Counts of HOLC Class Within Each Group

	‘Education’ Demographic Variable Categories And HOLC Grade Counts Within Perceived Neighborhood Dataset			
HOLC Grade	Group 1	Group 2	Group 3	Group 4
A	1	7	4	9
B	6	68	71	52
C	13	206	129	68
D	21	178	73	66

1. **Group 1:** Eighth grade or less, some high school
2. **Group 2:** High School/ GED, Trade/ Vocational/ Nursing School, Associate’s degree)
3. **Group 3:** Bachelor’s degree
4. **Group 4:** Master’s degree, professional degree, doctoral degree

Table 13. Chi-Squared Test Results for Two-Way Tables of HOLC Grades and Demographic Variables

Demographic Variable Categories	<i>Results Of Chi-Squared Test On Two-Way Tables Of Demographic Variables By HOLC Grade¹</i>		
	Race p-values	Income p-values	Education p-values
1	2.2e-16*	2.2e-16*	6.2e-05*
2	2.2e-16*	2.2e-16*	2.2e-16*
3	3.6e-10*	2.2e-16*	2.2e-16*
4	7.55e-3*	NA	4.8e-10*
5	3.9e-13*	NA	NA

1. Where (*) denotes significance when $\alpha < 0.05$

2. NA values denote lack of this category for the corresponding demographic variable.

Table 13 shows that we saw significant relationships between the distribution of categories within our demographic variables and the HOLC grade each neighborhood was assigned. Essentially, these results show that within each HOLC grade, the variation of people falling into different categories of race, income and education were not random, and that there was a relationship between an HOLC grade, and the demographic traits of an individual.

IV. DISCUSSION

Our investigation into the possible differences between our Summer and Winter RDD datasets, the variation between the Random Digit Dial and Online survey methodology, the distinctions between our demographic variables and their corresponding neighborhood dimensional variables, and the ways that HOLC grade designations may affect both the dimensional variables of our perceived neighborhoods, and their demographic characteristics, yielded an extremely interesting

array of information that did much to expand our understanding of how to collect data about perceived neighborhoods, and how demographic characteristics may impact the way that perceived neighborhoods are defined by individuals.

Our first analysis, which involved the comparison of Summer and Winter RDD datasets yielded results that indicated that the geocoding methods we used and the survey methods between the Summer and Winter data collection periods did not vary, and that the individuals we surveyed defined consistent distributions of area, perimeter and compactness when describing their perceived neighborhoods. This was also demonstrated through the fact that the dimensional variables between the two datasets showed similar distributions in the data (Table 5). Additionally, the geocoding success rate was similar for these two datasets (Table 3). This similarity indicates that we were able to apply our methods consistently and that this procedure could be used to assemble a combined perceived neighborhood dataset in New York City, for RDD data, even when the time frame for data collection was not constant. Given that we consistently produced informative perceived neighborhood data using this methodology, it would be possible to use the same methodology to obtain perceived neighborhoods in future research with perceived neighborhoods in New York City, as well as consider collecting perceived neighborhood data in similar regions using our methodology. However, given the unique qualities of an urban environment, it would be difficult to know if a similar analysis could be performed in a rural or suburban environment.

The results of our comparison between our combined Combined RDD dataset and our Online Survey data also yielded similar geocoding success rates and similar distributions in neighborhood dimensional data (Table 3, Table 5, Table 6). The only statistically significant difference seen between these two datasets was is the perimeter dimensional variable (p-value 0.0249) (Table 6). Despite this discrepancy, we saw the distributions for area and compactness data remain consistent between the two datasets, and the geocoding success rate of the Online dataset

was extremely similar to the RDD dataset (Table 3). Because of the variation in perimeter data distribution, we cannot say that these two survey methodologies would yield data that is entirely analogous with one another, but due to the similarity between the other dimensional variables and the geocoding success rate, the two methodologies are comparable. However, because the online datasets enabled a participant to draw his or her own polygon, this methodology may yield more accurate neighborhood data than the RDD survey methodology, which relied on our interpretation of how a person described their neighborhood in the geocoding of transcribed neighborhood boundaries. More research will need to be done into the effectiveness of one methodology compared to the other, but overall, because we were able to yield a similar number of viable perceived neighborhoods for analysis, which showed (with one exception) homologous dimensional characteristics, our methods for geocoding polygons using different survey methods appear to be valid.

The next analysis that was performed sought to explore the hypothesis that perceived neighborhood dimensions were related to different demographic variables. We used the demographic variables of gender, age, length of residency, marriage status, education status, income, employment status, and race. We did not see any statistically significant relationship between demographic variables and neighborhood dimensions for age, marriage status and length of residency. However, we did see statistically significant differences in neighborhood dimensions for gender, Hispanic identity, racial identity, income, and employment status (Table 7).

Mean neighborhood perimeter was greater for our female survey participants than our male participants, despite similarities in terms of area and compactness. This indicated that female survey participants perceived their neighborhoods to be more elongated with a greater spatial spread as compared to male participants. This trend was only seen in our online perimeter results, which

indicates that something about the online survey methodology likely influenced the way that survey participants perceived their neighborhood perimeter.

In addition to gender identity influencing one aspect of neighborhood dimension, we saw Hispanic identity as affecting multiple dimensional aspects of perceived neighborhoods. According to our results (Table 7), people who identified as Hispanic had statistically significant differences between the way that they defined their neighborhood's area and perimeter (for the online dataset only). Non-Hispanic survey participants, on average, described larger areas than Hispanic survey participants, as well as larger perimeters. This indicates that those who identify as Hispanic perceive their neighborhoods as being smaller, overall, with a more irregular shape, and covering a wider spatial range than those who don't identify as Hispanic. Again, we saw the perimeter trend in our online perimeter results only, which indicates that something about the online survey methodology influences the way that survey participants reported their neighborhood perimeter.

Additionally, there were statistically significant differences between area and compactness definitions when examined across a variety of different racial groups. Specifically, our results showed that our survey participants who identified as White perceived their neighborhoods as being, on average, larger in area than our survey participants who identified as Black and Mixed Race. For our compactness variable, the average compactness value for our White participants was larger than for our Black participants (Table 7). These results indicate that White participants defined neighborhoods as being larger than those defined by Black and Mixed raced participants, as well as with a wider and less condensed spatial range (as indicated by compactness) than their Black counterparts.

Income bracket also impacted our mean neighborhood compactness. Our results show that those in a higher income bracket ($\$93000 \leq \text{income}$) had a larger value for compactness than those

in our lowest income bracket ($\$25000 \leq \text{income} < \46000). This means that neighborhoods of higher income survey participants tended to be more compact, and less irregularly shaped.

Education attainment also impacted the mean area and mean compactness of perceived neighborhoods. For mean area, we found that participants with an eighth grade or less education or had some high school education perceived their neighborhoods as being smaller in area than those with a bachelor's degree or those with a master's, professional or doctoral degree. In terms of mean compactness, our survey participants who attained a high school/GED degree, a Trade/Nursing/Vocational school degree or an Associate's degree tended to perceive their neighborhoods as having a smaller compactness than those with a bachelor's degree or those with a master's, professional or doctoral degree.

Lastly, employment status was also associated with the neighborhood dimensions of survey respondents, in terms of both area and perimeter. In terms of neighborhood area, our results showed that our participants who identified themselves as being employed full time, employed part-time or self-employed had larger average perceived neighborhood areas than those who were out of work and looking for work, out of work and not looking for work and unable to work. Additionally, in terms of area, our survey participants who were students also had larger average perceived neighborhood areas than those who were out of work and looking for work. For the online neighborhood dataset only, participants who identified themselves as being employed full time, employed part-time or self-employed had larger average perceived neighborhood perimeters than those who were out of work and looking for work, out of work and not looking for work and unable to work. These findings imply that those who are employed full or part time, as well as those who are self-employed tend to perceive their neighborhoods as being larger in both area and perimeter than those who are unemployed. Students also have larger neighborhoods, in terms of both area and perimeter than our unemployed participants.

The next part of our analysis involved looking for differences in neighborhood dimensions between HOLC grades. Our hypothesis asserted that we would observe differences in area, perimeter and compactness between neighborhoods that were assigned different HOLC grades. However, the results of our analysis showed that there were no statistically significant differences. This means that the HOLC assignments for our perceived neighborhoods are independent from any neighborhood dimension variable. This refutes our hypothesis that... and means that HOLC grades, as established by the original HOLC, do not contain any trends in neighborhood dimensions in modern day perceived neighborhoods.

The last part of our analysis involved creating two-way tables for HOLC grades and demographic variable groups that we believed were most salient in terms of their relationship with the original demographic variables incorporated into HOLC's lending policy. These variables were race, income and education. According to our chi-square test results, there were statistically significant associations between HOLC grade and race, income, and education, as identified by a survey participant. These results indicate that differences between the distribution of the demographic variables of race, income and education are not due to random variation in specific cases outlined above.

These results imply that within our present day perceived neighborhoods, there is a relationship between historical HOLC grade, and the current demographic variables of race, income and education. However, there are differing trends that exist between each specific demographic variable. Looking at our variable for race identification (Table 10), there is a statistically significant association between each racial group and the HOLC grade in which their neighborhood lies. Within Group 1, or those who identified as White, have a majority of neighborhoods concentrated in the C and D neighborhood range. Additionally, while the proportion of HOLC grade A is small within Group 1, outside of Group 1, we see that 18 perceived neighborhoods out of 21 total

neighborhoods that are assigned grade A, are those associated with our White survey participants. Within Groups 2 (Black), 3 (Asian/Pacific Islander) and 4 (Hispanic), our survey population also tended towards a majority within the C and D grades. Within our income variable, the trends within groups were similar to our race variable, in that we see a concentration of neighborhoods within each group falling into the C and D HOLC grades. However, the trends across HOLC grades show that, as income increases, the likelihood that someone whose perceived neighborhood lies within HOLC grade A and also falls within a wealthier income bracket is higher. Conversely, as income decreases, the likelihood that someone whose perceived neighborhood lies with HOLC Grades B, C and D and also fall within a lower income bracket is higher. The trends within education are similar to the other two variables in that there was a statistically significant association between HOLC grades and education groups. However, when looking at the trends of education levels by HOLC grade, most perceived neighborhoods, regardless of education level, fell within the B and C grade range. Overall, in terms of race and income, our sample of perceived neighborhoods consisted largely of those falling within B and C neighborhood grades. Our sample of neighborhoods falling into HOLC grade A was small, but of those neighborhoods, most belonged to White-identified survey participants of a higher income bracket. Operating within our palimpsest conceptual model, this outcome indicates that vestiges of historical housing policies and restrictions, which sought to segregate neighborhoods by race and socioeconomic status, may exist in neighborhoods which were historically given a HOLC grade of A.

The potential for future research into using perceived neighborhoods for different kinds of spatial analysis, how to better design surveys and methodology for the purpose of acquiring data about perceived neighborhoods, and why individuals choose to define their perceived neighborhoods in different ways is immense. This particular study was unique in that it explores the possibility of demographic variables playing a part in perceived neighborhood definition, and while

there were many salient results that we uncovered through our analysis, it would be interesting to explore additional causes for neighborhood dimensions that are behavioral and geared towards activity spaces, due to our survey not differentiating between the neighborhood one calls “home” and the neighborhood in which one works, attends school, shops...etc. Accounting for time away from home is important as it allows researchers to get a better impression of the environmental exposures encountered by an individual which do not lie within the bounds of where one lives. Looking for demographic trends in how activity spaces are defined would also be an interesting subject of research as well, as it would lend a deeper insight into our findings regarding differences in neighborhood size, perimeter and compactness. If an individual feels limited in the region that they call home, perhaps understanding the dimensions of their activity space could shed further light on why these differences exist.

WORKS CITED

- Basta, Luke A., Therese S. Richmond, and Douglas J. Wiebe. 2010. "Neighborhoods, Daily Activities, and Measuring Health Risks Experienced in Urban Environments." *Social Science & Medicine* (1982) 71(11): 1943–50.
- Chaix, B., Merlo, J., Evans, D., Leal, C., & Havard, S. (2009). Neighbourhoods in eco-epidemiologic research: delimiting personal exposure areas. A response to Riva, Gauvin, Apparicio and Brodeur. *Social Science & Medicine* (1982), 69(9), 1306–1310.
- Chaix, B., Kestens, Y., Perchoux, C., Karusisi, N., Merlo, J., & Labadi, K. (2012). An interactive mapping tool to assess individual mobility patterns in neighborhood studies. *American Journal of Preventive Medicine*, 43(4), 440–450.
- Colabianchi, N., Coulton, C. J., Hibbert, J. D., McClure, S. M., Ievers-Landis, C. E., & Davis, E. M. (2014). Adolescent self-defined neighborhoods and activity spaces: spatial overlap and relations to physical activity and obesity. *Health & Place*, 27, 22–29.
<https://doi.org/10.1016/j.healthplace.2014.01.004>
- Coulton, C. J., Jennings, M. Z., & Chan, T. (2013). How big is my neighborhood? Individual and contextual effects on perceptions of neighborhood scale. *American Journal of Community Psychology*, 51(1-2), 140–150. Retrieved from <http://link.springer.com/article/10.1007/s10464-012-9550-6>
- Crossney, K. B., & Bartelt, D. W. (2005). Residential Security, Risk, and Race: The Home Owners Loan Corporation and Mortgage Access in Two Cities. *Urban Geography*, 26(8), 707-736.
- Cummins, S. (2007). Commentary: investigating neighbourhood effects on health--avoiding the "local trap." *International Journal of Epidemiology*, 36(2), 355–357.
<https://doi.org/10.1093/ije/dym033>
- Department of Justice. (2017). Justice Department Sues KleinBank for Redlining Minority Neighborhoods in Minnesota. (2017, January 13). Retrieved from <https://www.justice.gov/opa/pr/justice-department-sues-kleinbank-redlining-minority-neighborhoods-minnesota>.
- Desmond, M. (2017). *Evicted: Poverty and profit in the American city*. New York: Broadway Books.
- ESRI. (2016). ArcGIS 10.1; Redlands, CA.
- Federal Housing Administration. (1934-6). *Underwriting Manual: Underwriting and Valuation Procedure Under Title II of the National Housing Act With Revisions to April 1, 1936. Part II, Section 2, Rating of Location*. (Washington, D.C.),
- Hillier, A. E. (2003). Redlining and the Home Owners Loan Corporation. *Journal of Urban History*, 29(4), 394-420.

- Nelson, R., Winling, L., Marciano, R., Connolly, N. (2016) Mapping Inequality. *American Panorama*, ed. Robert K. Nelson and Edward L. Ayers.
- Oliver, L. N., Schuurman, N., & Hall, A. W. (2007). Comparing circular and network buffers to examine the influence of land use on walking for leisure and errands. *International Journal of Health Geographics*, 6, 41. <https://doi.org/10.1186/1476-072X-6-41>
- O’Sullivan, D., & Unwin, D. J. (2003) Geographic Information Analysis. John Wiley & Sons, Hoboken NJ, 178.
- QGIS Development Team. (2018). QGIS Geographic Information System. Open Source Geospatial Foundation.
- R Core Team R: A Language and Environment for Statistical Computing. (2016). R Foundation for Statistical Computing: Vienna, Austria.
- Sakas, M. (2017). Discriminatory Lending That Once Stifled Denver Neighborhoods Now Fuels Gentrification. Retrieved from <http://www.cpr.org/news/story/denver-neighborhoods-with-high-minority-populations-were-deemed-too-hazardous-for>.
- Schmool, J., Johnson, I., Dodson, Z., Keene, R., Gradeck, R., Beach, S., Clougherty, J. (2017) “Developing a GIS-based online survey instrument to elicit perceived neighborhood geographies to address the Uncertain Geographic Context Problem.” *Professional Geographer* (2017).
- Spielman, S. E., & Yoo, E. H. (2009). The spatial dimensions of neighborhood effects. *Social Science & Medicine* (1982), 68(6), 1098–1105. <https://doi.org/10.1016/j.socscimed.2008.12.048>
- Stewart, T., Duncan, S., Chaix, B., Kestens, Y., Schipperijn, J., & Schofield, G. (2015). A novel assessment of adolescent mobility: a pilot study. *The International Journal of Behavioral Nutrition and Physical Activity*, 12, 18. <https://doi.org/10.1186/s12966-015-0176-6>
- Vallée, J., Le Roux, G., Chaix, B., Kestens, Y., & Chauvin, P. (2015). The “constant size neighbourhood trap” in accessibility and health studies. *Urban Studies*, 52(2), 338–357. Retrieved from <http://journals.sagepub.com/doi/abs/10.1177/0042098014528393>
- Weeks, J. R., Hill, A. G., & Stoler, J. (2013). *Spatial inequalities. Health, poverty, and place in Accra, Ghana* (Vol. 110, GeoJournal Library). Springer Science and Business Media