

NORTHWESTERN UNIVERSITY

Operations Research Techniques in Data Limited Environments: Applications in
Public Services

A DISSERTATION

SUBMITTED TO THE GRADUATE SCHOOL
IN PARTIAL FULFILLMENT OF THE REQUIREMENTS

for the degree

DOCTOR OF PHILOSOPHY

Field of Industrial Engineering and Management Sciences

By

Kezban Yagci Sokat

EVANSTON, ILLINOIS

December 2017

© Copyright by Kezban Yagci Sokat 2017
All Rights Reserved

ABSTRACT

Operations Research Techniques in Data Limited Environments: Applications in Public Services

Kezban Yagci Sokat

Industrial engineering and management sciences (IEMS) tools have been shown to improve the efficiency of many systems from manufacturing to emergency rooms in a hospital. These tools can offer much to public officials in decision making. These decisions impact thousands of people; IEMS tools have the potential to provide guidance in decision making. However, as such decision making tools emerge in the IEMS literature, most models assume the availability of data while there is often limited and unreliable data for many public services. This work explores the impact of data availability in decision making tools in public service. The first part of the thesis, Chapter 2 and Chapter 3, focuses on controlling diseases: human immunodeficiency virus (HIV) and measles. HIV is a chronic disease and we use ordinary differential equations (ODEs) to study the population level impact of contraceptive use on births and HIV spread. We focus on HIV prevalence, the number of new HIV infections, the number of HIV-related deaths, the number of cases of vertical transmission, and the number of births over a 15-year period. We next use stochastic models to assess low-level tactics to control measles outbreaks. We analyze the impact of the community connectivity on spread of measles. We also investigate the impact of interventions (better case detection and a faster public health response) that have not been studied before. The second part of the thesis, Chapter 4 and Chapter 5, looks at the integration of data into humanitarian logistics models. We first provide a framework for assessing near real-time data. This study highlights the implications of real time data collection, processing and analysis from a first-hand experience. We provide a logistical content analysis of data based on a recent disaster. Inspired by Chapter 4, Chapter 5 presents a framework for dealing with incomplete information of road networks using machine and statistical learning methods. We develop a model in ArcGIS to automate the data gathering and processing steps as much as possible. We present an application of this framework to a past disaster where we estimate the status of a road: open, partially blocked or damaged. The third part of the thesis, Chapter 6, looks at the logistics of rural healthcare in Africa where supplies as well as the data are limited.

Limited historical data and limited on hand inventory bring additional difficulty to supply chain management. We develop a forecasting tool to estimate demand by combining alternative resources when the historical consumption data are not available. We also develop a method to calculate shortage costs in healthcare by linking health outcomes to inventory optimization under limited data. In collaboration with a non-governmental organization in Liberia, we apply these methods to obtain stocking levels for different medicines.

Acknowledgements

Ph.D. is a long journey. It is only possible to go through with passion, patience, able advising and lots of encouragement. This thesis would not be possible without endless support of many wonderful people. I would like to express my gratitude to these amazing people, and I apologize in advance to the ones whose names are not mentioned in this limited space.

First and foremost, I would like to thank my advisers Benjamin Armbruster, Irina Dolinskaya and Karen Smilowitz for not only supervising this thesis but also guiding my research, teaching, and academic life goals. Thank you for the opportunity to mentor undergraduate students. More importantly, I have been very lucky to have advisers who supported my family and personal goals along with my academic goals, encouraged me to go to INFORMS with my two months old son, and dedicated their time, experience and wisdom despite the time and space difference. I believe that each Ph.D. student is unique and I appreciate my advisers for treating me as a unique intellectual. I would like to thank my thesis committee member, Burcu Balcik for her support and advice on research, academia and future goals.

Collaborations are essential part of research. During my Ph.D., I had the opportunity to collaborate with a number people from different areas. I would like to thank to practitioners Ryan Bank and Jennifer Chan for their input in the research papers, Kelsey Rydland for his ArcGIS support, Dr. Rebecca Wurtz for her insight in infectious diseases, Nan Chen for the unique opportunity with work a non-governmental organization and James Kaufman for his able advisory during IBM research internship.

I want to express my deep gratitude to the faculty and staff of the Northwestern IEMS for providing a high quality research and teaching environment and the opportunity to interact with a diverse group of researchers. Apart from my advisers, I have been blessed with amazing mentors in the IEMS Department.

I would like to particularly thank Prof. Barry Nelson and Prof. Jill Wilson for their support on many different topics during my Ph.D. I would like to acknowledge Prof. Morton who has kindly shared his knowledge and provided advice when needed.

My dream to come to USA for higher education was a long life dream. I would like to thank to several educators and people who helped me make this dream reality. I would like to thank Prof. Nihat Berker of Massachusetts Institute of Technology for creating the Istanbul Technical University - Massachusetts Institute of Technology Freshman Physics Scholar Program. With this scholarship, he opened the doors of US higher education as a sophomore. I am grateful for the Fulbright for not only supporting my masters degree at Georgia Tech but also giving the chance to meet with full and bright people from all over the world. I would like to you Sevil Sezen for believing in me and making me feel proud.

I am in debt to Georgia Tech Faculty and their collaborators for preparing me for Ph.D. I am especially honored to have met with my lovely Fulbright adviser Dave Goldsman and my humble and always encouraging mentor Paul Griffin who provided continuous support during my higher education, believed in me and let me develop a personal relationship in addition to the academic one. I would like to thank Mark Ferguson, Julie Swann and Sarang Deao for the research opportunities as well as supporting my Ph.D. application.

I am also grateful to fellow students at Northwestern IEMS where I have been able to meet with wonderful people from all over the world. I am far most grateful for the opportunity to meet Ekkehard Beck who has been the best academic sibling, represents the Ph.D. - love of wisdom- and with whom I have shared many sleepless study sessions and most enjoyable and sophisticated memories about almost anything from politics to art. I would like to thank Kenan Arifoglu for his help in settling to Evanston and IEMS and endless mentorship throughout my Ph.D. I have appreciated my officemates Gokce Kahvecioglu, Nastaran Shojaee and Mariline Cherklesly for their encouragement and support at much needed times.

My Turkish family in Evanston and Chicago have been a huge part of my support system. Among many, I would like to thank my friends, Taner Aytun, Tuncay Ozel, Armagan Bayram, Kaan Kanad, Koc Family, Irem Unlu, who showed enormous support before, during and after my son's arrival. As from my Chicago family, I would like to express my gratitude to one special family, Nisa, Ilhan and Pelin, for being

our family in Chicago and being our side for our happiest and hardest times.

I am also thankful to my Georgia Tech family for their continuous encouragement and support during my Ph.D. I have turned to my good friend to Melih Celik for research, academic life and future goals. My dear friends Alicia and Amisha expressed their continuous belief and encouragement for me to maintain a strong attitude. My dear friends Hilal and Yavuz are another good example that distance does not matter to show your support.

Finally, but not the least, I would like to express my deepest gratitude to my beloved family. My parents, Hayriye and Sadik Yagci, always believed and continuously expressed that I have no limits and can overcome any challenge I may face. Thank you for emphasizing the importance of education, putting your daughters' education as the highest priority, and working continuously to provide us with good future. My sisters, Ulviye Yagci Karaca and Yasemin Hoke, have been my true best friends over thirty years. My dear nephew, Kerem, and niece, Zeynep, are the proof that distance is only as far as you make it. Lastly, I would like to thank to two important men in my life, my husband, Gurkan Sokat, and my son, Efe Sokat. First of all, thank you for your patience and literally following me across the country so that I can follow my dream of Ph.D. Thank you for sharing this journey with me and your constant support and encouragement. Your confidence and love have uplifted me and helped me to see the finish line. This thesis is dedicated to you.

List of abbreviations

ACT Artemisinin-based Combination Therapy

APAN All Partners Access Network

ART Antiretroviral Therapy

CaART Classification and Regression Trees

CCAA Clustering Combined with Adjacent Arc

CCMMM Clustering Combined with Modified Mean-Mode

CDC Centers for Disease Control and Prevention

CEMS Copernicus Emergency Management Service

COC Combined Oral Contraceptives (COCs)

CHW Community Health Worker

CSV Comma Separated Value

CT CT Classification Trees

DSWD the Department of Social Welfare and Development

DHN Digital Humanitarian Network

DALY Disability Adjusted Life Years

DMPA Depot-medroxyprogesterone acetate

ESRI Environmental Systems Research Institute

FC Female Condoms

FEMA Federal Emergency Management Agency

GIS Geographic Information System

HDX Humanitarian Data Exchange

HXL Humanitarian Exchange Language

HOT Humanitarian OpenStreetMap Team

HIV Human Immunodeficiency Virus

IASC Inter-Agency Standing Committee

iCCM Integrated Community Case Management

ICT Information and Communication Technology

IM Information Management

IUD Intrauterine Device

KML Keyhole Markup Language

k-NN k-Nearest Neighbors

LogIK Logistics Information about In-Kind Relief

LPI Logistics Performance Index

MC Male Condom

MMR Measles-Mumps-Rubella

NDRRMC The National Disaster Risk Reduction & Management Council

NET-EN Norethisterone Enanthate

NGO Nongovernmental Organization

NON No Contraception

OCHA The United Nations Office of Coordination for Humanitarian Affairs

OCP Oral Contraceptive Pills

ODE Ordinary Differential Equations

OSM OpenStreetMap

OTH Other Forms of Contraception

PDF Portable Data Formats

POIHC Progestogen-only Injectable Hormonal Contraception

QALY Quality Adjusted Life Years

SEIR Susceptible-Exposed-Infectious-Recovered

SHP Shapefile

UN United Nations

UNITAR United Nations Institute for Training and Research

UNOSAT United Nations Operational Satellite Applications Programme

US United States

VM Vaginal Microbicides

VISOV Volontaires Internationaux en Soutien aux Opérations Virtuelles

WFP World Food Programme

WHO World Health Organization

*This dissertation is dedicated to
my family, my husband and my little man.*

Contents

Acknowledgements	5
List of abbreviations	8
1 Introduction	19
2 Implications of Switching Away From Injectable Hormonal Contraceptives on the HIV epidemic	22
2.1 Introduction	22
2.2 Methods	23
2.2.1 Model	23
2.2.2 Scenarios of future contraceptive use	25
2.2.3 Parameter values	26
2.2.4 Analysis	27
2.3 Results	29
2.4 Discussion	32
2.5 Conclusion	35
3 Modeling Measles Outbreaks: The Role of Community Structure and Non-Vaccination Interventions	60
3.1 Introduction	60
3.2 Methods	61
3.2.1 Model	62
3.2.2 Outcome Measures of Interest	63
3.3 Theoretical Results	64
3.3.1 Convergence As Population Size Increases	64
3.3.2 Outbreaks Are Longer With Two Subgroups	66
3.3.3 Heterogeneity In Infection Risk	66
3.4 Numerical Results	67
3.4.1 The Impact of Population Size	68
3.4.2 The Impact of Heterogeneity	68
3.4.3 The Impact of Public Health Interventions	69
3.5 Discussion	70

Acknowledgement	73
4 Capturing Real-Time Data in Disaster Response Logistics	74
4.1 Introduction	74
4.2 Disaster Response Data Stakeholders	75
4.2.1 Data-gathering Communities	75
4.2.2 Disciplines Using Data	77
4.3 Framework for Analyzing Real-time Logistics Data for Modeling	80
4.3.1 Measures of Data Quality and Applicability	81
4.4 Framework Application: Typhoon Haiyan Case Study	88
4.4.1 Information and Data Retrieval	89
4.4.2 Data and Data Sources Classification	91
4.4.3 Logistical Content Data Availability Analysis	91
4.4.4 Lessons Learned: Challenges Faced During the Collection	94
4.5 Conclusion	98
5 Incomplete Information Imputation In Limited Data Environments	104
5.1 Introduction	104
5.2 Related Literature	105
5.2.1 Transportation Network Information in Humanitarian Logistics Studies	105
5.2.2 Estimating Incomplete Data	106
5.3 Methods	108
5.3.1 Attribute Selection	108
5.3.2 Data Collection and Processing	111
5.3.3 Data Imputation	112
5.3.4 Validation	117
5.4 Application to 2010 Haiti Earthquake Case Study	117
5.4.1 Data Used in the Case Study	118
5.4.2 Results	120
5.5 Modeling and Policy Implications and Lessons Learned	128
5.5.1 Making the Most Out of Limited Data	128
5.5.2 Applicability to Other Disasters	128
5.5.3 Applicability Beyond Disaster Response and Humanitarian Transportation Planning	129
5.6 Conclusion	130
Acknowledgement	131
6 Supply Chain Management in Limited Data Environments: An Application to Community Health	135
6.1 Introduction	135
6.2 Literature Review	136
6.2.1 Community Health and Data	136
6.2.2 Supply Chain Management/ Inventory Management	138
6.2.3 Public Health and Health Economics	138

	15
6.3 Problem Setting and Modeling	139
6.3.1 1 CHW 1 item: Single Item Stocking Level	140
6.3.2 Extensions	145
6.4 Discussion and Future Work	147
7 Conclusion	149

List of Tables

2.1	Parameters Acronyms: Progestogen-only injectable hormonal contraception (POIHC), IUD, female condoms (FC), vaginal microbicides (VM), male condoms (MC), oral contraceptive pills (OCP), other forms of contraception (OTH), and no contraception (NON). b Even though VM is not a contraceptive method, it is included in the study for its protective effect on HIV transmission.	28
2.2	Marginal Analysis for Kenya a) Change in New Infections b) Increase in Births. Change in outcomes based on the change in the usage of each contraceptive type and their effectiveness per unit of usage in the population on reducing births and HIV infections. BC refers to birth control while BC/HIVp refers to any form of contraception that also prevents HIV (FC and MC).	31
3.1	Notation	64
3.2	Parameters	68
3.3	Additional Sensitivity Analysis	71
3.4	Sensitivity Analysis for Unequal Sized Groups	71
4.1	Assessment of Framework Attributes	88
4.2	Classification of Information Data Files Collected by Our Team (see Section 4.3.1 for details on classification categories and Conclusion for description of sources.)	92
4.3	Road and Port Status Information Update after Typhoon Haiyan between 11/8/2013 and 11/13/2013 in NDRMMC Situational Reports.	93
4.4	Demand Information Update after Typhoon Haiyan between 11/8/2013 and 11/13/2013 from NDRMMC Situational Reports.	94
4.5	Demand Information Update after Typhoon Haiyan between 11/8/2013 and 11/20/2013 in OpenStreetMap and CEMS.	95
4.6	Vehicle Information Update after Typhoon Haiyan between 11/8/2013 and 11/13/2013 from LogIK Entries.	96
4.7	Medical Team Information Update after Typhoon Haiyan between 11/15/2013 and 11/30/2013 from WHO Health Cluster Reports.	96
5.1	List of potential attributes for road status estimation	110
5.2	Types of Imputation Techniques Used in This Study	112

5.3	Attributes used in the case study to estimate the road status	120
5.4	Accuracy Calculation Step 1: Example for Downtown Carrefour with Geographical Holdout	122
5.5	Accuracy Calculation Step 2: Example for Downtown Carrefour with Geographical Holdout	123
5.6	Summary of p-values on Impact of Percent of Missing Data Tests for Downtown Carrefour	124
5.7	Summary of p-values on Impact of Missing Data Spread Type for Downtown Carrefour . .	124
5.8	Summary of p-values on Impact of Imputation Method Comparison for Downtown Carrefour	124
5.9	Summary of p-values on Impact of Mapping Focus Area (Geographical holdout 1 km for Downtown Carrefour)	125
5.10	Impact of Mapping Focus Area (Random holdout 1 km for Downtown Carrefour)	126
5.11	Impact of Grid Size (Geographical holdout 1 km to 2 km comparison)	127
5.12	Impact of Grid Size (Random holdout 1 km to 2 km comparison)	127
A1	Summary of p-values on Impact of Percent of Missing Data Tests for Entire Carrefour . . .	133
A2	Summary of p-values on Impact of Missing Data Spread Type for Entire Carrefour	133
A3	Summary of p-values on Impact of Imputation Method Comparison for Entire Carrefour .	133
A4	Summary of p-values on Impact of Percent of Missing Data Tests for Entire Carrefour(2 km Grid)	134
A5	Summary of p-values on Impact of Missing Data Spread Type for Entire Carrefour(2km Grid)	134
A6	Summary of p-values on Impact of Imputation Method Comparison for Entire Carrefour(2km)	134

List of Figures

2.1	Compartmental model	24
2.2	State in 15 Years: New infections averted, decrease in prevalence, increase births and change in vertical transmission per 1000. POIHC, NON, OTH, IUD, MC, FC and VM stands for progestogen-only injectable hormonal contraception, no contraception, other contracep- tion methods, intrauterine device, male condom, female condom and vaginal microbio- cides, respectively.	30

2.3	Sensitivity Analysis of Outcomes for Kenya. Case 1: 100% increase in HIV acquisition risk, 100% increase in female-to-male transition transmission risk, Case 2: 50% increase in HIV acquisition risk, 50% increase in female-to-male transition transmission risk, Case 3: 50% increase in HIV acquisition risk, 50% increase in female-to-male transition transmission risk, Case 4: 50% increase in HIV acquisition risk, 50% increase in female-to-male transition transmission risk, Case 5: 0% increase in HIV acquisition risk, 0% increase in female-to-male transmission risk. For the panels showing the new infections averted and the increase in births, we show for comparison, σ , the standard deviation of the baseline number of new infections and births, respectively, from the probabilistic sensitivity analysis.	33
3.1	The average number of infected cases ($E[C]$, left, blue) and the probability of a large outbreak ($P[C \geq 5]$, right, green) as we vary the size of the population. The dashed lines are for the case of a homogenous population while the solid lines are for the case where the population is divided into two equal-sized subgroups. There were 1000 replications and the standard error is less than 0.11 for $E[C]$ and 0.02 for $P[C \geq 5]$	69
3.2	Histogram of total number of infected cases in a homogenous population of 1000 or a population with two subgroups each of 500 individuals. There were 1000 replications.	70
3.3	The average number of infected cases ($E[C]$, left, blue) and the probability of a large outbreak ($P[C \geq 5]$, right, green) as the mixing parameter changes in a population with two subgroups each of 500 individuals. The mixing parameter is the ratio of the probability of infecting a susceptible in the other subgroup to the probability of infecting a susceptible in the same subgroup. There were 10000 replications and the standard error is less than 0.04 for $E[C]$ and 0.02 for $P[C \geq 5]$	72
3.4	The average number of infected cases ($E[C]$, left, blue) and the probability of a large outbreak ($P[C \geq 5]$, right, green) as we vary the fraction vaccinated. The dashed lines are for the case of a homogenous population of 1000 while the solid lines are for the case where the population is divided into two equal-sized subgroups of 500. The red line indicates the deterministic herd immunity threshold. There were 1000 replications and the standard error is less than 0.22 for $E[C]$ and 0.02 for $P[C \geq 5]$	73
4.1	Data Analysis Framework: Applicability Measures and Related Attributes	81
4.2	Information and Data Retrieval Epochs	89
5.1	Data Collection and Processing Steps.	111
5.2	Missing Value Imputation with Clustering.	113
5.3	Missing Value Imputation Process with Classification Tree Adapted from [202].	117
5.4	Timeline After Haiti Earthquake.	118
5.5	1 km Grid of Carrefour.	119

Chapter 1

Introduction

Industrial engineering and management sciences (IEMS) tools have been shown to improve the efficiency of many systems from manufacturing to emergency rooms in a hospital. These tools can offer much to public officials in decision making. For example, public officials face difficult decisions from vaccination policies to allocation of relief supplies after a disaster. These decisions impact thousands of people; IEMS tools have the potential to provide guidance in decision making. However, as such decision making tools emerge in the IEMS literature, most models assume the availability of data while there is often limited and unreliable data for many public services. This work explores the impact of data availability in decision making tools in public service. Chapter 2 and 3 include applications in the public health domain, Chapter 4 and 5 represent the humanitarian logistics domain and Chapter 6 provides example on community health.

Chapter 2 studies the population level impact of contraceptive use on births and HIV spread. A study by Heffron et al. published in *Lancet Infectious Diseases* in 2011 [1] revealed that women using hormonal contraceptives, especially injectable depot-medroxyprogesterone acetate (DMPA), have double the risk of acquiring HIV and transmitting it to their partners. Over 12 million women in sub-Saharan Africa use hormonal contraceptive to avoid unintended pregnancies [2]. Moreover, United Nations' 2011 World Contraceptive Use Report [3] states that it is the preferred form of contraception in many sub-Saharan countries. We develop a compartmental model to analyze the population-level effects for various scenarios of future contraceptive use for five countries in sub-Saharan Africa. We focus on HIV prevalence, the number of new HIV infections, the number of HIV-related deaths, the number of cases of vertical transmission, and the number of births over a 15-year period.

We show that if all hormonal contraceptive users discontinue use, there is a large increase in births and vertical transmission. However, if a significant fraction of these users are successfully transferred to other forms of contraception, prevalence, new infections, and HIV-related deaths decrease considerably while the increase in births are limited. Also, the amount of vertical transmission only changes slightly, with the direction of the change depending on the country. These results can inform the public discussion on contraceptives in the wake of the DMPA study by Heffron et al. in 2011. This model can guide policy makers who are reconsidering marketing strategies for contraceptives by allowing them to consider the impact for both the fight against unintended pregnancies and HIV.

Chapter 3 presents our collaborative work with Northwestern University Public Health Department on measles. We propose a simulation model of an anticipated measles outbreak in the United States. De-

creased levels of Measles-Mumps-Rubella (MMR) vaccination, combined with extensive international travel, resulted in measles outbreaks in 2013. According to the Centers for Disease Control and Prevention (CDC), the United States has observed the highest levels of measles cases since 1996 [4]. We develop a mathematical model to study a measles outbreak. We analyze the impact of the community connectivity on spread of measles. We prove that the infection process converges to a finite limit as the population size increases. We also prove why outbreaks are longer with two subgroups and that heterogeneity increases the number of new infections, and the outbreak has a greater chance of dying out in the homogenous case. We complement our theoretical analyses with numerical results of this stochastic model. We find that community structure has little impact on the results. Models with a heterogeneous population find slightly greater outbreaks but that the differences are minimal. Even the size of the modeled population ceases to have an impact after the first few hundred individuals. We also investigate the impact of interventions (better case detection and a faster public health response) that have not been studied before. In addition to the expected impact of increased vaccination, we also find that better case detection and to a lesser extent a faster response have a significant impact on outcomes. In the wake of recent measles outbreaks in the United States and in Europe, our results can inform the public discussion about potential interventions to decrease the size of future outbreaks.

Chapter 4 initiates the second area of application in public services: humanitarian logistics. There has been a significant increase in disaster and humanitarian logistics research. Yet, there is still need to proactively improve the disaster cycle. One opportunity in this area is the enhancement in the information and communications technology and involvement of digital humanitarians that enable access to data after disasters. The amount of available data is rising; however, researchers are still not incorporating these data into decision making as much as possible. This chapter highlights the implications of real time data collection, processing and analysis from a first-hand experience. Examples of data include initial post disaster assessment in portable data format (PDF) maps, list of relief aid in Excel spreadsheet (XLS) format and user-mapped damage in comma separated value (CSV) format. We assess multiple properties of the data in addition to their implications and limitations. We provide logistical content analysis based on a recent disaster for demand, infrastructure and supply. To the best of our knowledge, this work is the first of its kind to focus on humanitarian logistics researchers' perspective on data collection process. We also follow an interdisciplinary approach of combining humanitarian practitioner's perspective with modeler's view. This work establishes the ground for an ongoing taxonomy of new and emerging data sources focusing on their applicability to humanitarian logistics operations, which is discussed in Chapter 5.

Inspired by the study on the availability of logistics data after a disaster in Chapter 4, Chapter 5 focuses on developing methods to estimate incomplete information on infrastructural damage in limited data environments. Utilizing the new data sources such as OpenStreetMap along with currently available data, we develop a framework to estimate incomplete information in limited data environments. In this framework, we benefit from similarities between road segments with known status, named attributes. We provide guidance for identification of selected attributes that enable successful imputation. The framework next explains data collection and processing and various imputation techniques such as naive methods (optimistic, pessimistic, neutral and popularistic), clustering combined models (clustering combine with modified mean-and-mode method and adjacent arc), and decision tree based methods (classification tree). Furthermore, we develop a model in ArcGIS to automate the data gathering and processing steps as much as possible. We present an application of this framework to a recent disaster where we estimate the status of a road: open, partially blocked or damaged. Results show high levels of success in accurately identifying damage levels for infrastructure. This study aids field opera-

tions managers to deploy help effectively and efficiently given limited information about road damage immediately after disaster so that they can save as many lives as possible.

Chapter 6 presents supply chain management in limited data environment with an application to community health setting. This study is in collaboration with a non-governmental organization (NGO) in Liberia that aims to increase access to healthcare in remote areas with through community health worker (CHW) program. The NGO trains and deploys CHWs to diagnose and treat the most common causes of morbidity and mortality. These CHWs use a backpack to store the medicinal items. Composition of the backpack and the stocking levels of the items in the backpack is critical to successful implementation of the CHW program. In order to expand and seek for donor funding, the NGO needs to know the demand for each medicinal item in the backpack. On the other hand, there is limited historical consumption data. We first develop a forecasting model to estimate the demand for medicinal items requirements by combining alternative resources. We provide guidance of these alternative resources for future applications. We next estimate the inventory management costs; specifically overage and underage costs, which are hard to calculate in the healthcare setting, by linking the population health measures to operational outcomes. With the proposed forecasting method and the estimated costs, we develop stocking levels for the items in the backpack for the case of the NGO.

The final chapter (Chapter 7) concludes the thesis with the summary and future directions.

Chapter 2

Implications of Switching Away From Injectable Hormonal Contraceptives on the HIV epidemic

2.1 Introduction

Injectable hormonal contraception is the preferred form of contraception in many sub-Saharan countries representing 8.1%- 28.4% of the contraceptive use in the countries we consider [3]. A shot provides two to three months of protection depending on the type. There are two types of progestogen-only injectable hormonal contraception (POIHC): depot-medroxy progesterone acetate (DMPA) and norethisterone enanthate (NET-EN). DMPA is the most widely used progestin-only injectable [1]. More than 12 million women in Africa use DMPA for pregnancy prevention [2]. POIHC is highly effective in preventing pregnancies compared to other methods, and it has a .3% risk of failure with perfect use and a 3% risk with typical use [5].

However, some observational studies (the earliest from 1991 [5]) link the use of certain contraceptives with an increased risk of HIV acquisition [1, 6, 7, 8, 9, 10, 11, 12, 13]. Most studies focus on combined oral contraceptives (COCs) and/or POIHC (including DMPA and NET-EN). There is limited data on the potential relationship between HIV risks and other hormonal contraceptive methods such as implants, vaginal rings, or intrauterine devices (IUD). Studies involving NETEN did not conclude any significant relationship between NET-EN use and HIV risk [13, 14].

On the other hand, there is evidence that DMPA increases HIV acquisition and transmission risk. Clinical and laboratory studies suggest several possible biological reasons including vaginal structural changes, higher cervicovaginal HIV shedding and higher number of inflammatory cells in cervicovaginal fluid [15]. A recent study of HIV-1-serodiscordant couples in Africa (Botswana, Kenya, Rwanda, South Africa, Tanzania, Uganda, and Zambia) by Heffron et al. suggests that DMPA may double the risk of HIV infection for women [1]. Subsequent meta-analyses also find increased HIV acquisition risk for women using DMPA [16, 17]. In addition, Heffron et al. is the only study directly looking at the relationship between POIHC and HIV transmission risk, and find that this rate is doubled. They observed increased concentrations of HIV-1 RNA in endocervical secretions from HIV-1 infected women using injectable contraceptives as the potential cause for the increased risk of HIV transmission [1].

Results from this study caused debate among healthcare providers and policy makers. The previous

recommendation of the World Health Organization (WHO) did not have any restrictions on the use of hormonal contraceptives [18]. Even though the WHO kept its policy recommendation, it recommended that women using POIHC should get dual protection for HIV and pregnancy by using female or male condoms in addition to POIHC [19, 20].

While there is no consensus about the exact effect of POIHC on HIV risk, these studies may lead to changes in public health programs for family planning. Especially, in African countries with high prevalence of both POIHC use and HIV, governments may advise women to quit using POIHC or to switch to other methods. Switching from POIHC to other forms of contraception may reduce HIV infections. On the other hand, decreased use of these contraceptives may cause an increase in unintended pregnancies and higher mother-to-child transmission (vertical transmission) of HIV. This of course depends on the type of contraception because male condoms, for example, are also recognized as a way of controlling the HIV epidemic, preventing HIV infection among adults, and preventing mother-to-child HIV transmission [21].

In this article, we model the population level impact of the potential association of POIHC with increased HIV risk. We predict the effect of potential changes in DMPA use on childbirths, vertical transmission, HIV infections and prevalence in different countries in sub-Saharan Africa for a variety of scenarios of changes in contraceptive use. The remainder of this paper is organized as follows. Section 2 presents the model. Section 3 provides numerical results, and section 4 discusses these results. Section 5 concludes the paper with final remarks.

2.2 Methods

2.2.1 Model

The population we consider is adults aged 15-49. We use a deterministic compartmental model of HIV spread. We select Kenya, Zambia, South Africa, Rwanda and Botswana for our study. These are the same countries that Heffron et al. studied except that we omit Uganda and Tanzania as their prevalence is similar to Kenya's (6.5%, 5.6%, and 6.3%, respectively), which we do include. Contraceptive use information was not available for Tanzania and overall contraceptive use was lower in Uganda (23.7%) than in Kenya (45.5%).

The WHO states that the risk of HIV transmission due to POIHC is more important for countries where women have a high risk of acquiring HIV; where hormonal contraceptives (especially POIHC) count for a significant portion of all modern methods used; and where the maternal mortality rate (MMR) is high [18]. These countries are in sub-Saharan Africa where the majority of women with HIV in the world reside [22]; where POIHC represents from 20% to almost 60% of all modern methods of contraception [3]; and where the MMR is mostly higher than the global average. The MMRs of Kenya, Zambia, South Africa, Rwanda, and Botswana are 413, 603, 237, 383 and 513 per 100,000 live births respectively while the global average is 251 per 100,000 live births [23]. In addition, Kenya, Zambia, South Africa, Rwanda, and Botswana display variety in the levels of HIV prevalence and in the usage of different contraceptives.

We assume all transmission is heterosexual and divide the adult population into four compartments by gender and HIV status. We assume that the ratio of women to men is one, since the actual ratio is quite close. We also assume that contraceptive use does not affect the progression of HIV in an HIV-infected female [18]. Fig. 2.1 illustrates the compartmental model and Table A.1 in the Appendix summarizes the notation used.

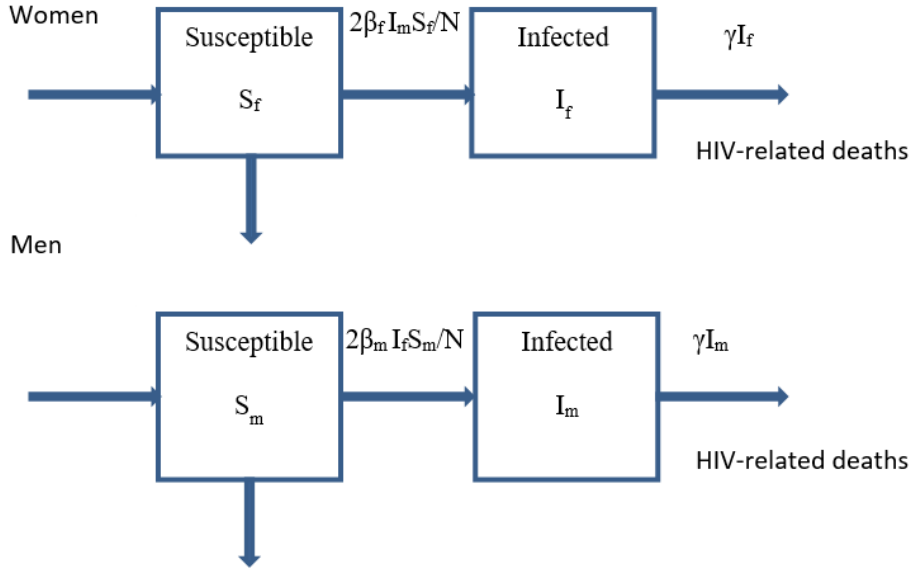


Figure 2.1: Compartmental model

Below are the differential equations (Eqs. (1)-(7)) that fully specify the model. Here N is the size of our population; I_f and I_m are the number of infected females and males in the population; and S_f and S_m are the number of susceptible females and males in the population. Note that $I_f + S_f = I_m + S_m = N/2$, $(1 - \frac{I_m}{N/2} = S_f/N$ and $(1 - \frac{I_m}{N/2} = S_m/N$. The first two equations describe the rate of new infections; the third equation describes the overall population growth; and Eqs. (4)-(6) describe outcomes of interest:

$$\dot{I}_f = -\gamma I_f + \beta_f I_m (1 - \frac{I_f}{N/2}) \quad (2.1)$$

$$\dot{I}_m = -\gamma I_m + \beta_m I_f (1 - \frac{I_m}{N/2}) \quad (2.2)$$

$$\dot{J} = \alpha N \quad (2.3)$$

$$\dot{N} = \beta_f I_m (1 - \frac{I_f}{N/2}) + \beta_m I_f (1 - \frac{I_m}{N/2}) \quad (2.4)$$

$$\dot{V} = \delta_f I_f \quad (2.5)$$

$$\dot{B} = \phi N \quad (2.6)$$

The four outcomes we track are the cumulative number of infections since the start of the simulation, J , (not the current number of infected); the cumulative number of cases of vertical transmission (i.e., mother-to-child transmission), V ; the cumulative number of births, B ; and the HIV prevalence at the end of the simulation. The time horizon we look at, 15 years, is short enough that infants infected vertically do not enter the adult population in our model, allowing us to assume that individuals enter the population uninfected.

2.2.2 Scenarios of future contraceptive use

To compare different levels of contraceptive use, we track the fraction using POIHC, IUD, female condoms (FC), vaginal microbicides (VM), male contraceptives (condoms) (MC), oral contraceptive pills (OCP), no contraception (NON), and other forms of contraception (such as vaginal barrier methods) (OTH). We include vaginal microbicides even though they do not prevent pregnancy, because they do prevent HIV infection. We assume that OTH do not affect HIV transmission. We do not consider couples using multiple forms of contraception simultaneously because that is not common [24]. Overall, we consider methods only preventing pregnancy (POIHC, IUD, OCP and OTH), methods only preventing HIV (VM), and dual protection methods which refers to methods that prevent both HIV and pregnancy (MC and FC).

We consider various scenarios of future contraceptive use. A scenario $P = (p_{POIHC}, p_{IUD}, p_{FC}, p_{VM}, p_{MC}, p_{OCP}, p_{OTH}, p_{NON})$ is a vector that tracks the fraction using each type. For example, the scenario (10%, 2%, 1%, 3%, 1%, 5%, 20%, 55%) has 10% of females using POIHC, 2% of females using IUD, 1% of females using FC, 3% of females using VM, 1% of males using MC, 5% of females using OCP, 20% of females using other forms of contraception, and 55% of couples using no contraception. Our sources do not directly give the fraction using other forms of contraception (OTH) so we calculate this from the fact that the components of P sum to 100%.

We consider six different scenarios in addition to a baseline (i.e., status quo) scenario and then look at the various outcomes over a 15-year horizon. We assume that sexual behavior does not change beyond these explicit changes in contraceptive use detailed in the following scenarios. Thus, our analysis does not consider self selection of one type of contraceptive over another or why people switch from one form to another. The baseline scenario (scenario 0) represents the current situation and corresponds to vector P_0 of contraceptive use and its values are given in Table 1 for various countries. Scenarios 2-6 consider five alternative futures: in each, the population switches to a different distribution of contraceptive use. Since changing behavior takes time [25], whether due to a public health campaign or not, we assume that in scenarios 1-5, the contraceptive use switches from the current levels after one year. Note that in scenarios 2, 5, and 6, the total fraction of the population using any kind of protection does not change. In the other scenarios, the total fraction using protection may be significantly less than in the baseline except for Botswana where the total contraceptive use increases slightly.

Scenario 0 (Baseline): Current contraceptive use. $P_0 = (p_0, p_{0,POIHC}, p_{0,IUD}, p_{0,FC}, p_{0,VM}, p_{0,MC}, p_{0,OCP}, p_{0,OTH}, p_{0,NON})$,

Scenario 1 (POIHC to NON): After one year, all POIHC users stop using POIHC and switch to NON. $P_1 = (0, p_0, p_{0,IUD}, p_{0,FC}, p_{0,VM}, p_{0,MC}, p_{0,OCP}, p_{0,OTH}, p_{0,NON} + p_{0,POIHC})$,

Scenario 2 (POIHC to OTH): After one year, all POIHC users stop using POIHC and switch to OTH. $P_2 = (0, p_0, p_{0,IUD}, p_{0,FC}, p_{0,VM}, p_{0,MC}, p_{0,OCP}, p_{0,OTH} + p_{0,POIHC}, p_{0,NON})$,

Scenario 3 (POIHC to IUD & MC): After one year, POIHC use drops 50%, IUD use increases 25%, and MC use increases 25%. The remaining individuals switching from POIHC will not use any form of contraception. $P_3 = (0.5p_0, p_{0,POIHC}, 1.25p_0, p_{0,IUD}, p_{0,FC}, p_{0,VM}, 1.25p_0, p_{0,MC}, p_{0,OCP}, p_{0,OTH}, p_{0,NON} + (0.5p_0, p_{0,POIHC} - 0.25p_0, p_{0,IUD} - 0.25p_0, p_{0,MC}))$,

Scenario 4 (POIHC to MC): After one year, POIHC use drops 50% and MC use increases 50%. The remaining individuals switching from POIHC will not use any form of contraception. $P_4 = (0.5p_0, p_{0,POIHC}, p_0, p_{0,IUD}, p_{0,FC}, p_{0,VM}, 1.5p_0, p_{0,MC}, p_{0,OCP}, p_{0,OTH}, p_{0,NON} + (0.5p_0, p_{0,POIHC} - 0.5p_0, p_{0,MC}))$,

Scenario 5 (POIHC to FC): After one year, 25% of POIHC users switch to FC. $P_5 = (0.75p_0, p_{0,POIHC}, p_0, p_{0,IUD}, p_{0,FC} + 0.25p_0, p_{0,POIHC}, p_{0,VM}, p_{0,MC}, p_{0,OCP}, p_{0,OTH}, p_{0,NON})$,

Scenario 6 (POIHC to VM): After five years, 25% of POIHC users switch to VM. $P_6 = (0.75p_{0,POIHC}, p_{0,IUD}, p_{0,FC}, p_{0,VM} + 0.25p_{0,POIHC}, p_{0,MC}, p_{0,OCp}, p_{0,OTH}, p_{0,NON})$.

We have no basis upon which to predict the future prevention behavior in these countries. We choose these scenarios because they cover a variety of possibilities of what might happen. They allow for the continued use of POIHC; users ceasing to use any contraceptives; and users switching to other contraceptives, including a new HIV prevention method (VM). They also cover the possibility of a net decrease in the number of people using contraception. We should note that these are scenarios rather than explicit interventions. We do not consider costs; the feasibility of achieving a particular scenario using a public health campaign; or try to determine the optimal distribution of contraceptive that such a campaign should aim for. Changing behavior on a national level has all kinds of difficulties that are hard to model [26].

2.2.3 Parameter values

For the HIV acquisition risk, Heffron et al. find that DMPA doubles it [1], and subsequent meta-analyses find a 1.5-fold [16] and a 1.4-fold [17] increase. We use 1.5 because [16] surveys 18 studies while [17] considers only 10 studies. For HIV transmission risk, Heffron et al. is the only study directly measuring the impact of POIHC on it [27], and so we use its finding that the risk is doubled.

We assume that 65% of eligible individuals receive antiretroviral therapy (ART) [28]; that the life expectancy without ART is 11.6 years after HIV infection [29]; and that the life expectancy with ART is 37 years after HIV infection [30]. We obtain the HIV-related mortality rate of the population, $\gamma = 35\%(1/11.6\text{year}) + 65\%(1/37\text{year})$, by taking a weighted average of the rates with and without ART. Equations 8-16 below describe how to calculate for a scenario P the four model parameters $\beta_f, \beta_m, \varphi$, and δ , which are not given directly in Table 1. Here β_f and β_m are the infection rates for females and males, φ is the total birth rate, and δ is the vertical transmission rate.

$$\xi_{ff} = \xi_{IHC}^f p_{IHC} + \xi_{VM}^f p_{VM} + \xi_{FC}^f p_{FC} + 1(p_{IUD} + p_{OTH} + p_{NON}) \quad (2.7)$$

$$\xi_{mf} = \xi_{MC}^f p_{MC} + 1(1 - p_{MC}) \quad (2.8)$$

$$\beta_f = \beta_{0f} \xi_{ff} \xi_{mf} \quad (2.9)$$

$$\xi_{mm} = \xi_{MC}^m p_{MC} + 1(1 - p_{MC}) \quad (2.10)$$

$$\xi_{fm} = \xi_{IHC}^m p_{IHC} + \xi_{FC}^m p_{FC} + 1(1 - p_{IHC} - p_{FC}) \quad (2.11)$$

$$\beta_m = \beta_{0m} \xi_{mm} \xi_{fm} \quad (2.12)$$

$$\chi = \sum_i \chi p_i \quad (2.13)$$

$$\varphi = \chi \varphi_0 \quad (2.14)$$

$$v = p_{ART} \pi + (1 - p_{ART}) \pi' \quad (2.15)$$

$$\delta = v \varphi \quad (2.16)$$

We calculate β_f , β_m , and φ by multiplying risk-adjustment factors with the values these parameters take with no contraception, β_{0f} , β_{0m} , and φ_0 . Specifically, we calculate the HIV infection rate of women, β_f , in Eqs. (7)-(9) by multiplying β_{0f} by ξ_{ff} and ξ_{mf} , and the risk adjustment factors for female contraception (POIHC, VM, and FC) and male condoms, respectively. These risk adjustment factors are weighted sums of the risk-reduction for female (male) infection due to each form of contraception, for

example the risk reduction of MC for female (male) infection, ξ_{MC}^f (ξ_{MC}^m), weighted by the use of MC, p_{MC} . Similarly, we calculate in Eq. 10-12 the HIV infection rate of men, β_m . In Eq. 13 and 14, we calculate the birth rate, φ , by multiplying φ_0 with the relative risk of pregnancy, c . This relative risk is again a weighted sum of the effectiveness of each form of contraception [5]. We calculate the rate of vertical transmission, d , by multiplying the risk of vertical transmission per birth, v , by the birth rate φ (Eq. 15 and 16). The risk of vertical transmission is π if the mother is on ART and π' without any treatment or intervention. Thus, we calculate the risk of vertical transmission, v , as a weighted sum of π and π' weighted by the percentage of the HIV + pregnant females who are on antiretroviral drugs, p_{ART} [31].

We choose the remaining parameters, β_{0f} , β_{0m} , and φ_0 , so that the above equations give the current values of β_f , β_m , and φ for the status quo scenario, P_0 . We fit the base contact rates β_{0m} and β_{0f} for each country such that (a) the infection rate for females is twice that of males, $\beta_{0f} = 2\beta_{0m}$ [32]; and (b) that the ratio of the simulated prevalence at year 5 to the starting prevalence matches the ratio of the prevalence in 2012 to the prevalence in 2007. To determine φ_0 , we first calculate the risk of being pregnant, χ , under. We then look up the current birth rate, τ [33], and using Eq. 13 and 14 set. Table 1 gives the parameter values and Table A.2 in the Appendix shows the calculated values of the parameters we discussed above for the baseline scenario. To validate the model, Appendix Fig. A.1 compares the relative change in the simulated prevalence over time to the UNAIDS prevalence estimates from 2007 to 2012.

2.2.4 Analysis

To better understand the simulation results in the various scenarios, we also conducted a marginal analysis that decomposed the increase in the births and the change in new infections into the change in usage of each contraceptive type and their effectiveness per unit of usage in the population on reducing births and HIV infections. These results are then compared to the simulation results. We should also note that these are linear approximations while the simulation follows the disease dynamics over 15 years. As discussed in the previous subsection, the degree to which POIHC increases HIV acquisition is uncertain, and so far, Heffron et al. is the only study directly looking at the degree to which POIHC increases HIV transmission [1]. We perform sensitivity analysis focusing on this key factor by varying the risk of HIV male-to-female and female-to-male transmission when using POIHC. Specifically, we investigate the following cases for the impact on POIHC use:

Case 0 (Baseline): 50% increase in HIV acquisition risk, 100% increase in female-to-male transmission risk,

Case 1: 100% increase in HIV acquisition risk, 100% increase in female-to-male transmission risk,

Case 2: 50% increase in HIV acquisition risk, 50% increase in female-to-male transmission risk,

Case 3: 50% increase in HIV acquisition risk, 0% increase in female-to-male transmission risk,

Case 4: 0% increase in HIV acquisition risk, 50% increase in female-to-male transmission risk,

Case 5: 0% increase in HIV acquisition risk, 0% increase in female-to-male transmission risk.

In addition, we conduct a probabilistic sensitivity analysis for five parameters (HIV acquisition and transmission rate with POIHC, birth rate, and initial contraceptive use of POIHC, initial contraceptive use of MC). For each replication in the probabilistic sensitivity analysis, we simultaneously draw each parameter from a uniform distribution ranging from -10% to $+10\%$ of its baseline value. We chose this distribution for its simplicity.

Table 2.1: Parameters Acronyms: Progestogen-only injectable hormonal contraception (POIHC), IUD, female condoms (FC), vaginal microbicides (VM), male condoms (MC), oral contraceptive pills (OCP), other forms of contraception (OTH), and no contraception (NON). b Even though VM is not a contraceptive method, it is included in the study for its protective effect on HIV transmission.

Parameter	Value					Source
	Kenya	Zambia	South Africa	Rwanda	Botswana	
Initial female prevalence, $2If/N$ (%)	7.15	13.1	21.35	3.22	25.09	[38]
Initial male prevalence, $2Im/N$ (%)	5.1	12.3	14.4	2.6	20.09	[38]
Annual population growth rate, α (%)	2	2	2	2	2	[31]
Annual mortality rate of HIV+ (deaths per year), γ	0.0477	0.0477	0.0477	0.0477	0.0477	[28, 29, 30]
Contraceptive use in status quo, P_0 (%) *						[3]
POIHC	21.6	8.5	28.4	15.2	8.1	
IUD	1.6	0.1	1	0.2	1.7	
FC	0	0	0	0	0	
VM	0	0	0	0	0	
MC	1.8	4.7	4.6	1.9	15.5	
OCP	7.2	11	0.9	6.4	14.3	
OTH	13.3	16.5	15	12.7	4.8	
NON	54.5	59.2	40.1	63.6	55.6	
Relative risk of contraceptives on female infection rate, **						[20, 34, 5, 35]
POIHC	1.5	1.5	1.5	1.5	1.5	
IUD	1	1	1	1	1	
FC	0.24	0.24	0.24	0.24	0.24	
VM	0.46	0.46	0.46	0.46	0.46	
MC	0.2	0.2	0.2	0.2	0.2	
OCP	1	1	1	1	1	
OTH	1	1	1	1	1	
NON	1	1	1	1	1	
Relative risk of contraceptives on male infection,						[1, 34, 5, 35]
POIHC	2	2	2	2	2	
IUD	1	1	1	1	1	
FC	0.24	0.24	0.24	0.24	0.24	
VM	1	1	1	1	1	
MC	0.2	0.2	0.2	0.2	0.2	
OCP	1	1	1	1	1	
OTH	1	1	1	1	1	
NON	1	1	1	1	1	
Risk reduction of a birth control method for pregnancy, $(1 - \chi)$ (%)						[[5]
POIHC	97	97	97	97	97	
IUD	99.2	99.2	99.2	99.2	99.2	
FC	79	79	79	79	79	
VM	15	15	15	15	15	
MC	85	85	85	85	85	
OCP	92	92	92	92	92	
OTH	70	70	70	70	70	
NON	15	15	15	15	15	
Percentage of pregnant females on ART, pART (%)	73	69	88	65	95	[28]
Vertical transmission probability (%)						
When HIV+ mother is on ART, pi	5	5	5	5	5	[36]
When HIV+ mother is not on ART, pi'	26	26	26	26	26	[37]
Current annual birth rate (per 1000), τ	31.93	43.51	19.32	36.14	22.02	[33]

2.3 Results

Fig. 2 compares the simulation outcomes of the various scenarios to the baseline scenario. The absolute magnitude of the outcomes is shown in Table A.4 in the Appendix, and the change in prevalence over time is shown in Fig. A.2 in the Appendix. Compared with the baseline, all scenarios had fewer new infections and lower prevalence. In most scenario-country combinations, the births increased compared with the baseline, while the change in vertical transmission had no clear trend.

Scenarios POIHC to NON and POIHC to OTH provide the most reduction in terms of new infections and prevalence for all countries except Botswana. These reductions are larger in countries where both the current prevalence and POIHC use are high, such as Kenya and South Africa. Since the initial HIV prevalence was significantly higher in South Africa than in Kenya, these scenarios lead to correspondingly larger reductions.

In Botswana, POIHC to MC followed by POIHC to IUD & MC provide the largest reduction in terms of new infections and prevalence. This result is driven by the significant increase in MC use and the fact that the initial contraceptive use of MC is higher in Botswana than in other countries. Thus, as mentioned before, the 25% or 50% increase in MC use outweighs the users switching away from POIHC, resulting in a net increase in the total contraceptive use whereas in all other scenario-country combinations, there is a net decrease in contraceptive use. Additionally, MC has the most protection against HIV. Table A.5 in the Appendix uses marginal analysis to explain in detail such changes for each country.

In Table A.3 of the Appendix we also see the relative reduction in new infections by sex. We find that the benefit for men and women is almost the same except in scenario POIHC to VM. Females have a greater decrease in new infections than males in scenario POIHC to VM because VM only reduces the risk of HIV infection for women.

We see that births increase for almost all scenarios in all countries except in Botswana where the births decrease in scenario POIHC to MC. As before, Botswana is an exception due to the high initial MC use. In all countries, scenario POIHC to NON results in the largest increase in births and vertical transmission.

We observe that scenarios with a method preventing both pregnancy and HIV (POIHC to IUD&MC, POIHC to MC, and POIHC to FC) perform better than scenarios focusing on a method preventing only HIV (POIHC to VM) for most of the outcomes. This effect is most obvious when comparing scenario POIHC to FC to scenario POIHC to VM where the exact same number of the women switch to FC in the former and VM in the latter. We see that scenario POIHC to FC outperforms scenario POIHC to VM in all outcome measures. This is due to the fact that VM only prevents HIV infection and is not a form of contraception while FC does both.

The marginal analysis is shown in Table 2 below for Kenya and in Table A.4 of the Appendix for all the countries of births and new infections in Table 2 below and Table A.4 in the Appendix. It is reassuring that the marginal analysis, which used a linear approximation, gives results with relative differences that are similar to those of the simulation results. The POIHC to VM scenario seems to show larger differences, which can be explained by the fact that in this scenario VM was not in place until five years after other contraceptive changes.

Fig. 3 shows the sensitivity analysis for Kenya. Sensitivity analysis results for all countries are given in Fig. A.3 and Table A.5 in the Appendix. For the new infections averted and the increase in births, Fig. 3 also includes the standard deviation of those outcomes in the probabilistic sensitivity analysis, which can be found in Fig. A.4 of the Appendix. Births (though not cases of vertical transmissions) remain the same for all sensitivity cases since we only consider changes in HIV acquisition and transmission risk.

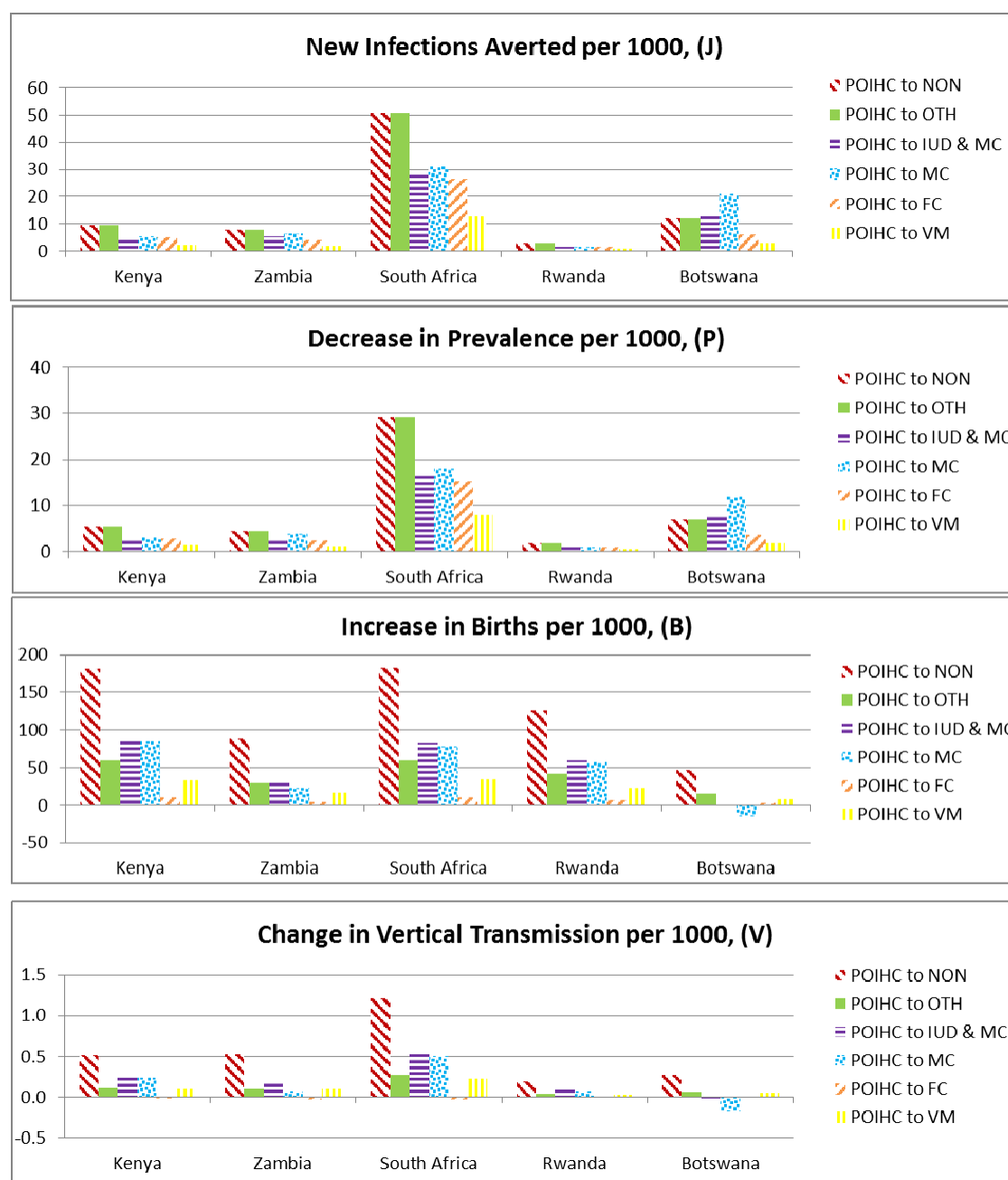


Figure 2.2: State in 15 Years: New infections averted, decrease in prevalence, increase births and change in vertical transmission per 1000. POIHC, NON, OTH, IUD, MC, FC and VM stands for progestogen-only injectable hormonal contraception, no contraception, other contraception methods, intrauterine device, male condom, female condom and vaginal microbicides, respectively.

Contraceptive type	Change in each contraceptive type per each scenario							Change in relative risk							Change in new infections (per 1000)						
	Baseline (%)							FF							NON						
	NON	OTH	IUD & MC	MC	FC	VM		MF	MM	FM		NON	OTH	IUD	MC	FC	VM				
POIHC	21.6	-21.6	-10.8	-10.8	-10.8	-5.4	-5.4	0.39	0.01	0.01	0.78	2.3	2.3	1.1	1.1	0.6	0.6				
IUD	1.6	-	0.4	-	-	-	-	-0.11	0.01	0.01	-0.22	-	-	0.0	-	-	-				
FC	-	-	-	-	5.4	-	-	-0.87	0.01	0.01	-0.98	-	-	-	-	0.9	-				
VM	-	-	-	-	-	5.4	-	-0.65	0.01	0.01	-0.22	-	-	-	-	-	0.5				
MC	1.8	-	0.5	0.9	-	-	-	-0.11	-0.79	-0.79	-0.22	-	-	0.1	0.2	-	-				
OCp	7.2	-	-	-	-	-	-	-0.11	0.01	0.01	-0.22	-	-	-	-	-	-				
OTH	13.3	-	21.6	-	-	-	-	-0.11	0.01	0.01	-0.22	-	-	-	-	-	-				
NON	54.5	21.6	-	10.0	9.9	-	-	-0.11	0.01	0.01	-0.22	0.5	-	0.2	0.2	-	-				
Any BC	45.5	-21.6	-	-9.9	-9.9	-	-5.4	-	-	-	-	Total Change	2.8	2.8	1.5	1.5	1.1				
Any BC/HIVp	1.8	-	0.5	0.9	5.4	-	-	-	-	-	-	Simulation	9.4	9.4	5.0	5.3	4.9	2.3			
Contraceptive type	Contraceptive usage as fraction of the population							Difference in pregnancy risk compared to average							Increase in births per 1000						
	Baseline (%)							Change in percentage points							NON						
	NON	OTH	IUD & MC	MC	FC	VM		NON	OTH	IUD & MC	MC	FC	VM		NON	OTH	IUD & MC	MC	FC	VM	
POIHC	21.6	-21.6	-10.8	-10.8	-10.8	-5.4	-5.4	-0.49	-	-	-	-	-	-	105.5	105.5	52.7	52.7	26.4	26.4	
IUD	1.6	-	0.4	-	-	-	-	-0.51	-	-	-	-	-	-	-	-	-2.0	-	-	-	
FC	-	-	-	-	5.4	-	-	-0.31	-	-	-	-	-	-	-	-	-	-	-16.6	-	
VM	-	-	-	-	-	5.4	-	0.33	-	-	-	-	-	-	-	-	-	-	-	17.9	
MC	1.8	-	0.5	0.9	-	-	-	-0.37	-	-	-	-	-	-	-	-	-1.8	-	-	-	
OCp	7.2	-	-	-	-	-	-	-0.44	-	-	-	-	-	-	-	-	-	-	-	-	
OTH	13.3	-	21.6	-	-	-	-	-0.22	-	-	-	-	-	-	-	-	-	-	-	-	
NON	54.5	21.6	-	10.0	9.9	-	-	0.33	-	-	-	-	-	-	-	-	0.0	33.2	32.8	-	
Any BC	45.5	-21.6	-	-9.9	-9.9	-	-5.4	Total change	-	-	-	-	-	-	177.1	58.3	82.0	82.3	9.7	44.3	
Any BC/HIVp	1.8	-	0.5	0.9	5.4	-	-	Simulation Result	-	-	-	-	-	-	181.4	59.7	84.0	84.2	10.0	33.6	

Table 2.2: Marginal Analysis for Kenya a) Change in New Infections b) Increase in Births. Change in outcomes based on the change in the usage of each contraceptive type and their effectiveness per unit of usage in the population on reducing births and HIV infections. BC refers to birth control while BC/HIVp refers to any form of contraception that also prevents HIV (FC and MC).

For HIV-related outcomes, all cases show a smaller decrease (some by up to a factor of 2.5) compared to the baseline case where we used the risk numbers provided by Morrison [17] and Heffron et al. [1]. When we compare the cases to Case 1 (rather than Case 0), in which we take risk numbers from Heffron et al., the other sensitivity cases show a smaller decrease for the HIV-related outcomes (some by up to a factor of 3.5). The impact depends on the country. For example, when these parameters decrease, Kenya, South Africa and Rwanda show similar behavior (where POIHC to FC is most favorable), different from Zambia and Botswana (where POIHC to MC is most favorable).

Case 0 and Case 1 have the most dramatic results. Case 1, which uses the HIV risk parameters from Heffron et al. gives the largest decrease in new infections and prevalence and the lowest increase in vertical transmission. Since births are unaffected and since Case 1 has the highest transmission and acquisition risk, these results are expected. As the case number increases, the decrease in the new infections and prevalence slows down while the increase in vertical transmission increases slightly. Case 3 and 4 where the acquisition and transmission risks are increased by 50%, respectively, show similar results with Case 3 being more favorable for HIV outcomes.

Comparing for the new infections averted in Fig. 3, the size of the standard deviation from the probabilistic sensitivity analysis to the size of the differences of the Cases to the baseline case 0, confirms that a key parameter is the HIV acquisition and transmission risk when using POIHC.

2.4 Discussion

The large increase in births and cases of vertical transmission make scenario POIHC to NON, where all POIHC users stop using any form of protection, undesirable. In all other scenarios, HIV-related outcome measures (new infections and prevalence) improve up to 21% while births increase less than 15% with the exception being South Africa where births increase up to 24%. In most scenarios and countries, the change in the level of vertical transmission is small and can go in either direction. In contrast, the POIHC to OTH scenario leads to similarly good HIV outcomes, while having a substantially smaller increase in births. However, it is the only scenario that keeps all POIHC users on some form of contraception, making it unfair to compare it to the other scenarios and emphasizing the importance of keeping women that discontinue POIHC on some form of contraception. In the following we compare the remaining four scenarios.

In general, dual protection methods perform the best as expected. Aside from POIHC to NON, POIHC to MC provides the largest decrease in new infections and prevalence followed by POIHC to IUD&MC and POIHC to FC. POIHC to FC results in the lowest increase in births followed by POIHC to VM and POIHC to MC. While these different scenarios have similar effects in most of the countries, the magnitudes of the effects are different in each country due to the initial distribution of the contraceptive use. Similar to births, POIHC to FC provides the lowest increase in vertical transmission. However, since vertical transmission depends on both birth and new infections, it is hard to identify similar trends for other scenarios as we did for the other metrics. We can also conclude that for Botswana, scenario POIHC to MC, which emphasizes male condoms, is the preferred scenario because it provides the greatest reduction in all outcomes (provided of course that the significant increase in MC use is possible in Botswana). Even the births are expected to decrease by 4% for Botswana with this scenario since the current MC use is quite high, and the increase in MC use compensates for their decreased birth-control efficacy as compared to POIHC.

In many cases, scenario POIHC to FC, which increases female condom use, may be the preferred outcome because it not only decreases HIV-related outcomes but also decreases vertical transmission. In

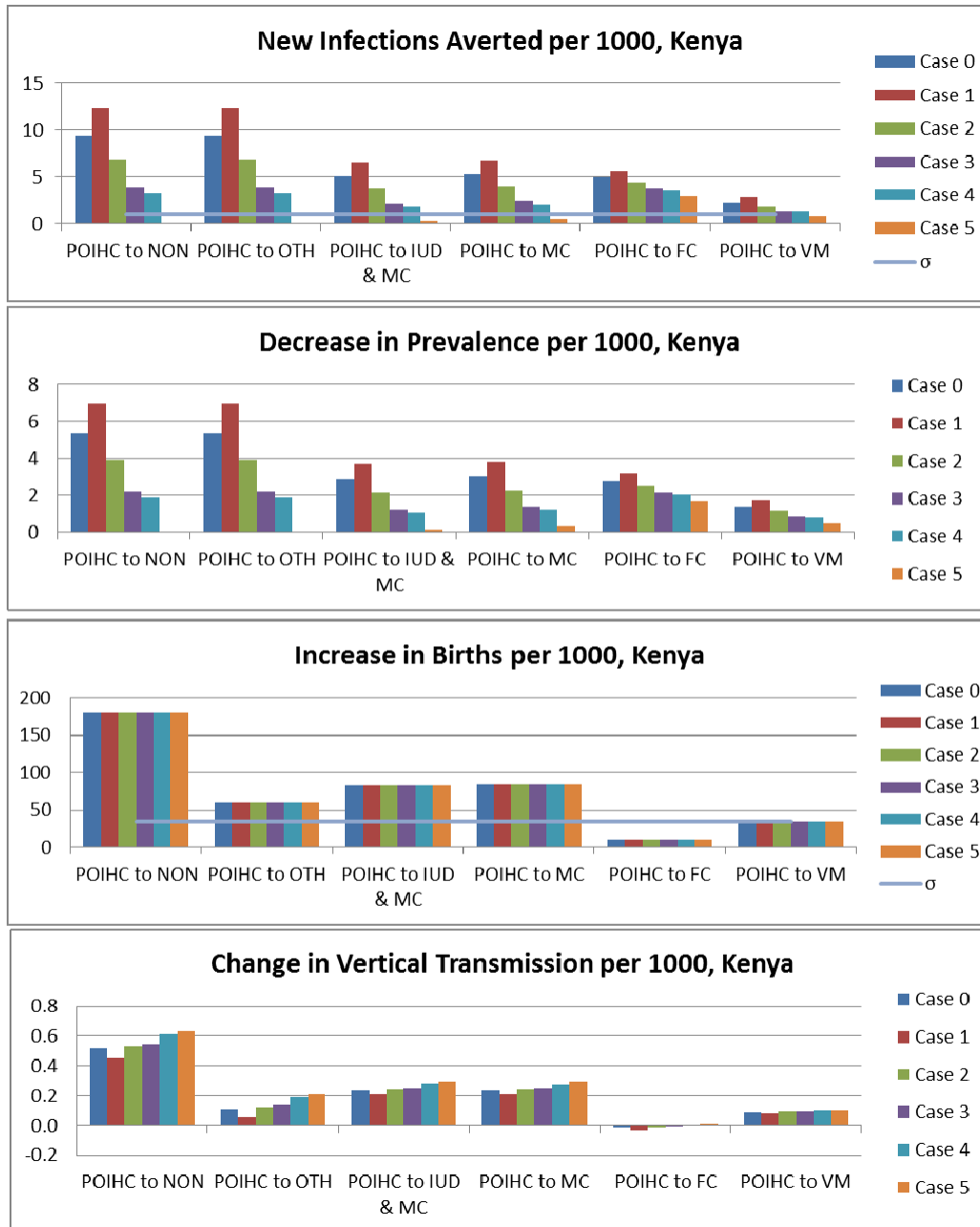


Figure 2.3: Sensitivity Analysis of Outcomes for Kenya. Case 1: 100% increase in HIV acquisition risk, 100% increase in female-to-male transition transmission risk, Case 2: 50% increase in HIV acquisition risk, 50% increase in female-to-male transition transmission risk, Case 3: 50% increase in HIV acquisition risk, 50% increase in female-to-male transition transmission risk, Case 4: 50% increase in HIV acquisition risk, 50% increase in female-to-male transition transmission risk, Case 5: 0% increase in HIV acquisition risk, 0% increase in female-to-male transmission risk. For the panels showing the new infections averted and the increase in births, we show for comparison, σ , the standard deviation of the baseline number of new infections and births, respectively, from the probabilistic sensitivity analysis.

addition, births increase less than 3% in that scenario. However, female condom use is rare despite being recommended by public health agencies [3]. For that reason we did not consider the even less common scenario involving the simultaneous use of POIHC and either male or female condoms, providing dual protection for birth control and HIV transmission [24].

Some scenarios might be more feasible than others. In addition to FC being very rare, the differences in a specific country's contraceptive use behavior (both the prevalence of all forms of contraception and the distribution among different contraceptive choices) might make some scenario more practical than others. For example, public officials might need to more marketing effort to change behavior in condom use in Kenya compared to Botswana where prevalence of MC is already high.

There is no consensus about the relationship between POIHC use and increased HIV risk. However, the meta-analysis by Morrison et al. is strong evidence for a relationship between POIHC use and specifically, male-to-female HIV transmission [16]. For female-to-male HIV transmission, Heffron et al. is the only study that finds an association with POIHC [1]. While [27] identified 16 studies that indirectly looked for an association, most of which did not find any, Heffron et al. was the only study identified that directly looked for an association [1]. Sensitivity analysis (comparison of Case 1 with the highest transmission risk values to comparison of Case 5 with the lowest transmission risk values) shows that different assumptions about HIV transmission for POIHC users can have up to a 3.5-fold difference in the magnitude of the results (i.e., the effect of a contraceptive use scenario compared to the baseline). This stresses the need for a more conclusive study of the relationship between POIHC use and HIV transmission. Fig. 3 and Appendix Fig. A.3 also show that the uncertainty in the other parameters as described in the probabilistic sensitivity analysis is of a similar magnitude as the uncertainty explored by the Cases for the POIHC-linked HIV transmission risk.

The limitations of our model are as follows. The rate at which the population will change its contraceptive usage is unknown. We assumed that changes would occur after one year. Slower changes would lessen the differences to the status quo. We made two modeling assumptions that are not true but still reasonably close to reality: we assumed that the ratio of females to males is one and we based our base contact rates on the epidemic dynamics of 2007-2012. We used a simple compartmental model instead of a detailed simulation model. However, this is appropriate since we are studying population-level outcomes over 15 years and have included details (contraception and new HIV infections) important to the factors being studied. We excluded other details such as the different stages of HIV progression or changes to treatment coverage since they affect the role of contraceptives on new infections only very indirectly. In addition, we do not know the likelihood of the different scenarios occurring or the feasibility of using public health campaigns to achieve them. Currently DMPA is much more common in these countries than NET-EN [1, 2]. We assumed that POIHC users would continue to prefer DMPA over NET-EN in the future. If this is not true then the magnitude of changes may be different since there is no reported association between HIV risk and NET-EN use [21] (unlike the case for DMPA [1]). We also assumed that the sexual behavior in a country will not otherwise change when one form of birth control is replaced by another. Currently, differences in fertility characteristics such as birth spacing in various countries and their relation to the use of various contraceptive methods are not well understood. However, we must build our model and base our recommendations on the data available. The most sensible assumption is that POIHC use does not cause short birth intervals but that these fertility characteristics are due to behavioral and cultural factors that merely correlate with the use of POIHC.

2.5 Conclusion

We observe that switching from POIHC to other types of protection will be beneficial for HIV related outcome measures. Especially for countries where both the POIHC use and the HIV prevalence are high, the HIV-related benefits of switching from POIHC to other protection options are great. For countries with low birth rates, the negative impact of switching from POIHC on births and vertical transmission will be less. Overall, the outcomes depend on the countries and models such as these are useful for tailoring any potential public health intervention to a specific country or population of interest. Especially when combined with analyses of feasibility and costs, our simulations of the various scenarios can form the basis of future public health interventions.

Our results depend highly on the value of the HIV acquisition and transmission risk parameters for those using POIHC, which are currently uncertain. Our analysis explored three major sources of uncertainty: (1) the unknown future sexual behavior of the population in different scenarios; (2) the biological parameter values such as transmission probabilities and risk reductions in different cases and in a probabilistic sensitivity analysis; and (3) potential limitations and sources of model error in the discussion. The female-to-male and male-to-female transmission risks impact the HIV outcomes but they do not impact the births. Despite the uncertainty of these parameters, even low percentages of dual protection method use can help balancing between HIV and birth related population level outcomes. The simulations in this study show that stopping POIHC use, with those individuals not switching to any other form of contraception, results in the worst outcomes of all the scenarios considered, an important fact public policy decision makers should keep in mind when designing interventions and preparing for the potential population-level changes in sexual behavior due to the link between POIHC and HIV.

Acknowledgments We have no conflicts of interest to declare. We thank Stéphane Helleringer and Rebecca Wurtz for insightful discussions. We also thank the editor-in-chief, Vedat Verter, and the reviewers for their insightful comments.

Supplementary Appendix**Table of Contents**

Fig. A.1 Validation

Fig. A.2 Prevalence over 15 years in various countries

Fig. A.3 Sensitivity analysis for each country

Fig. A.4 Probabilistic sensitivity analysis for each country

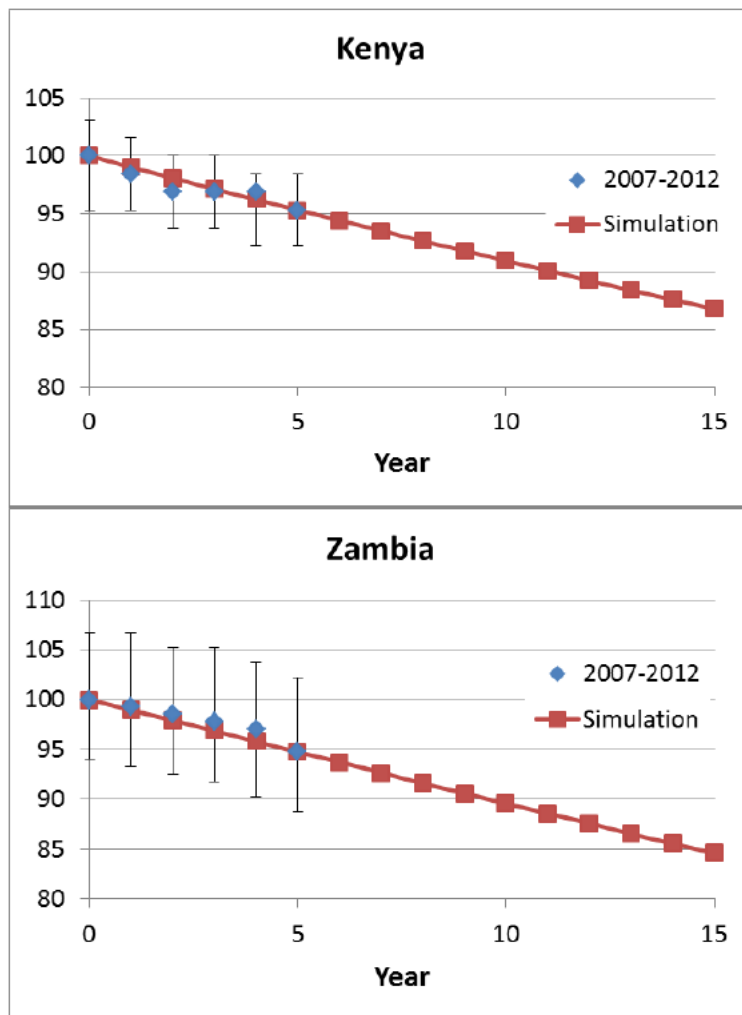
Table A.1 Notation used in the model

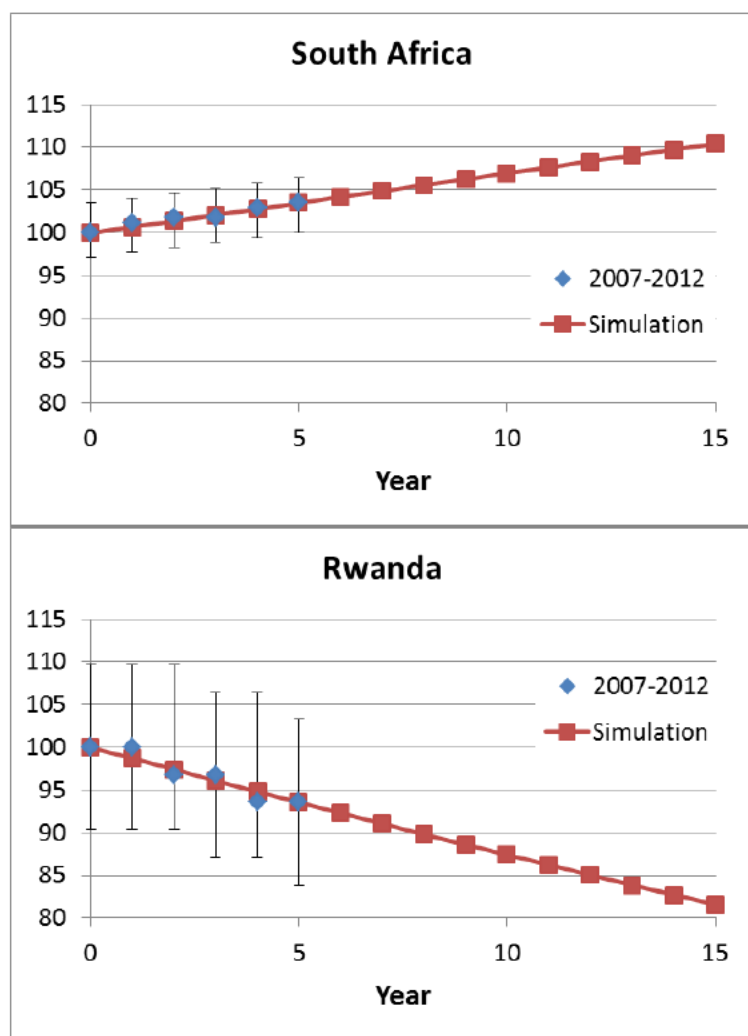
Table A.2 Values of calculated parameters for baseline scenario

Table A.3 Cumulative number of new infections over 15 years by sex

Table A.4 Marginal Analysis

Table A.5 Sensitivity Analysis





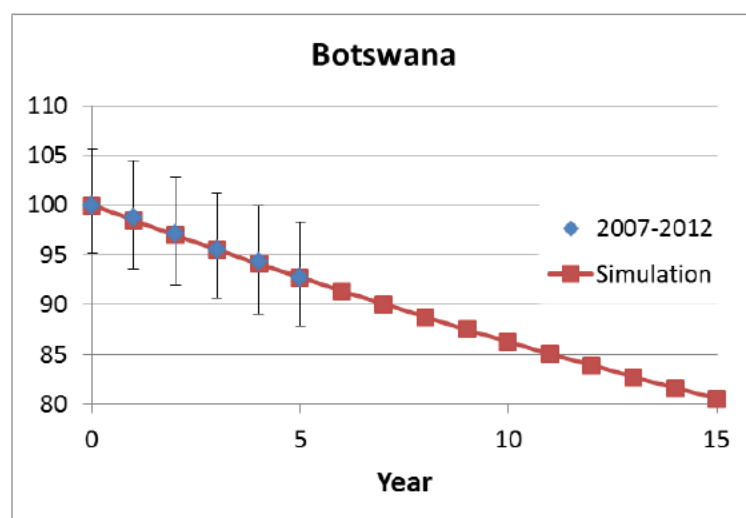
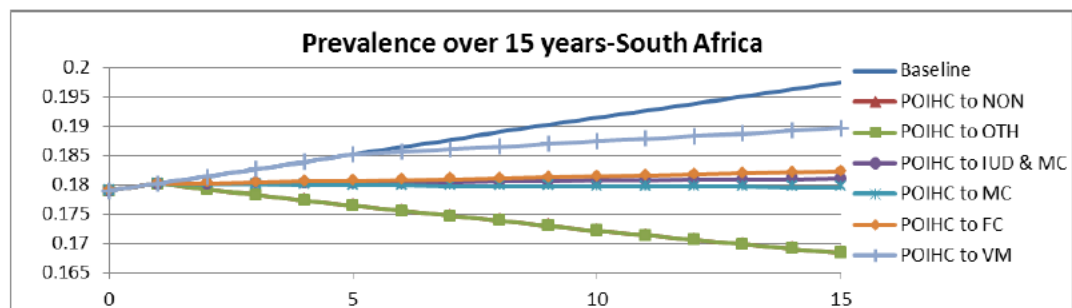
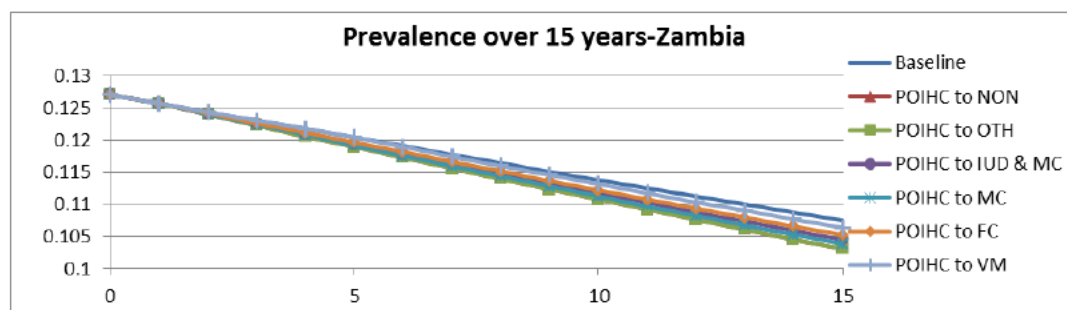
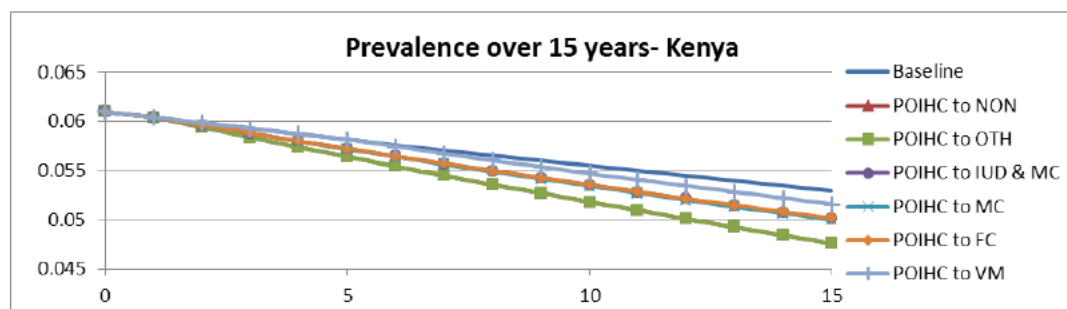


Figure A.1 Validation. The blue time series are the UNAIDS estimates for adult (age 15-49) prevalence from 2007 to 2012, along with the low and high estimates. The red time series is the simulation. Both have been normalized to start at 100.



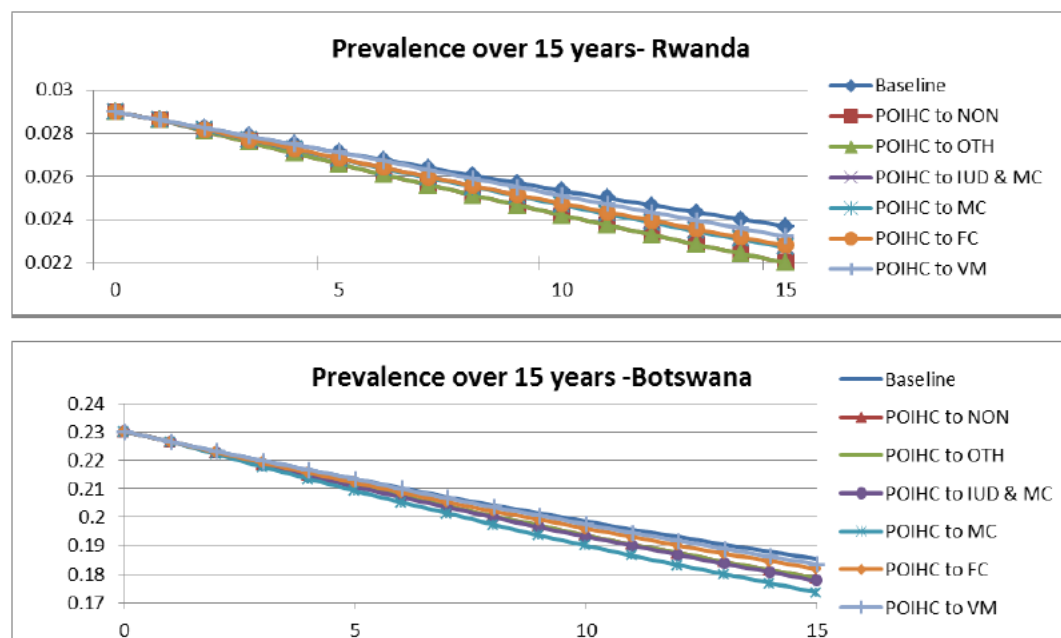
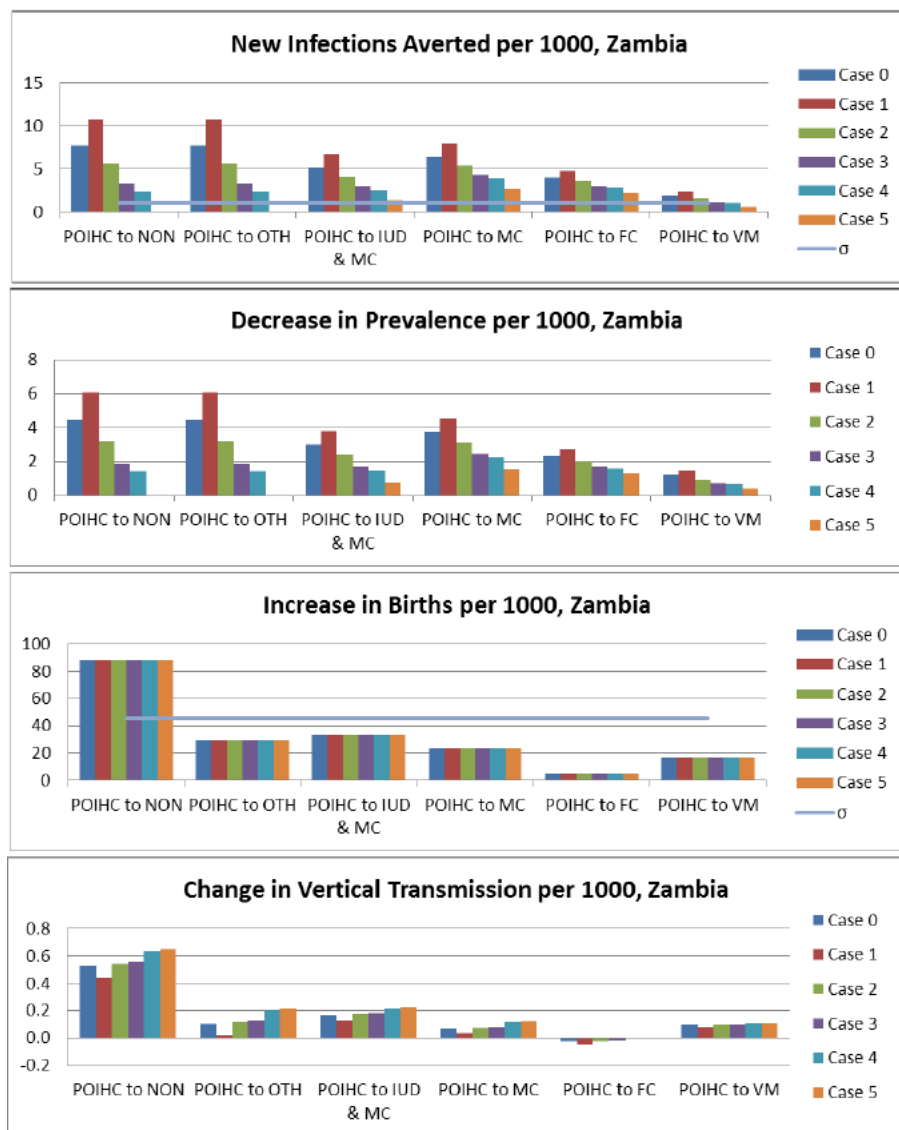


Figure A.2 Prevalence over 15 years in various countries



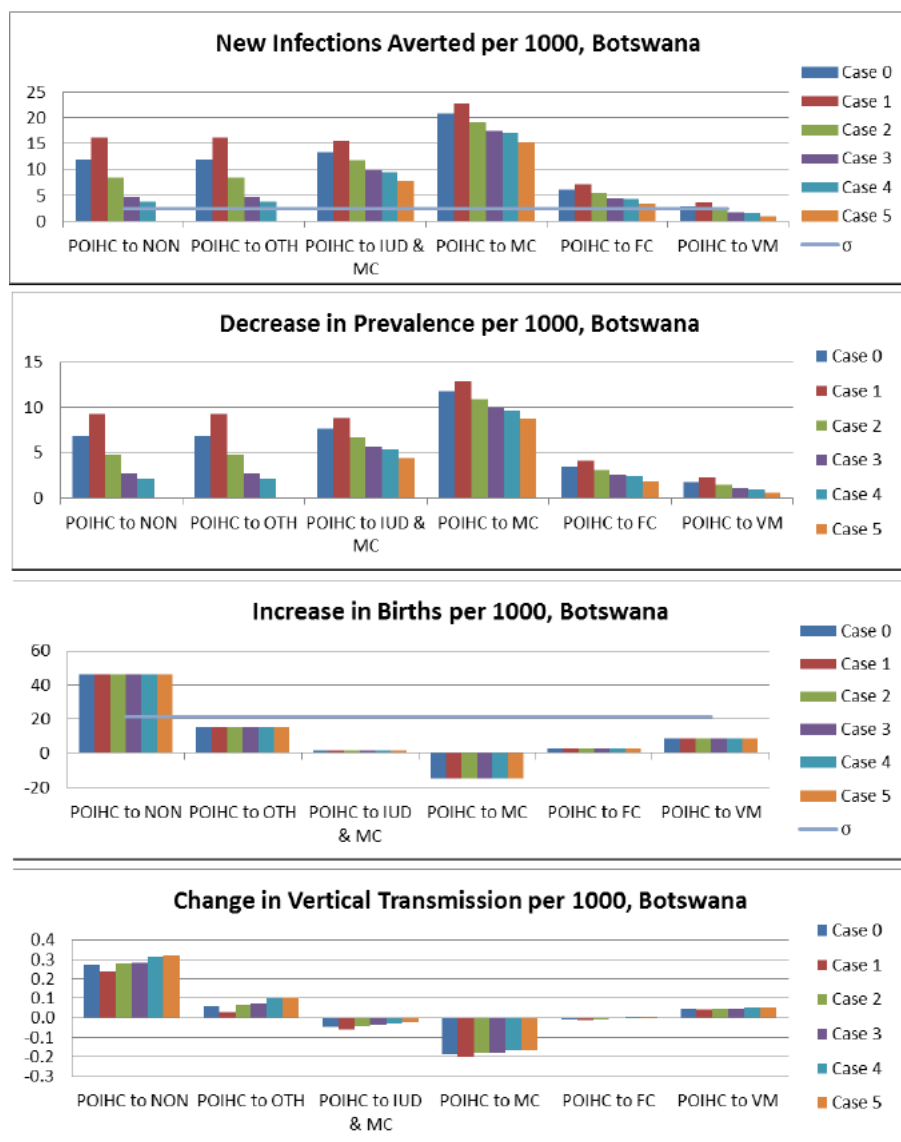
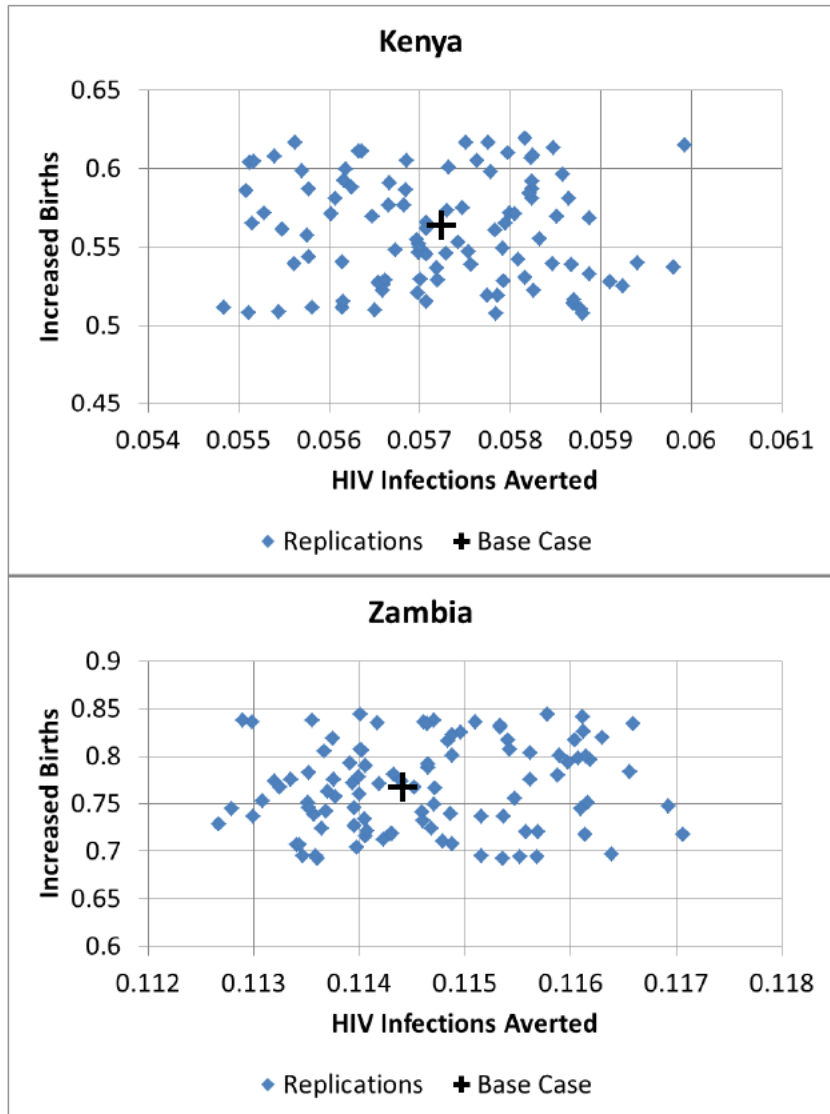
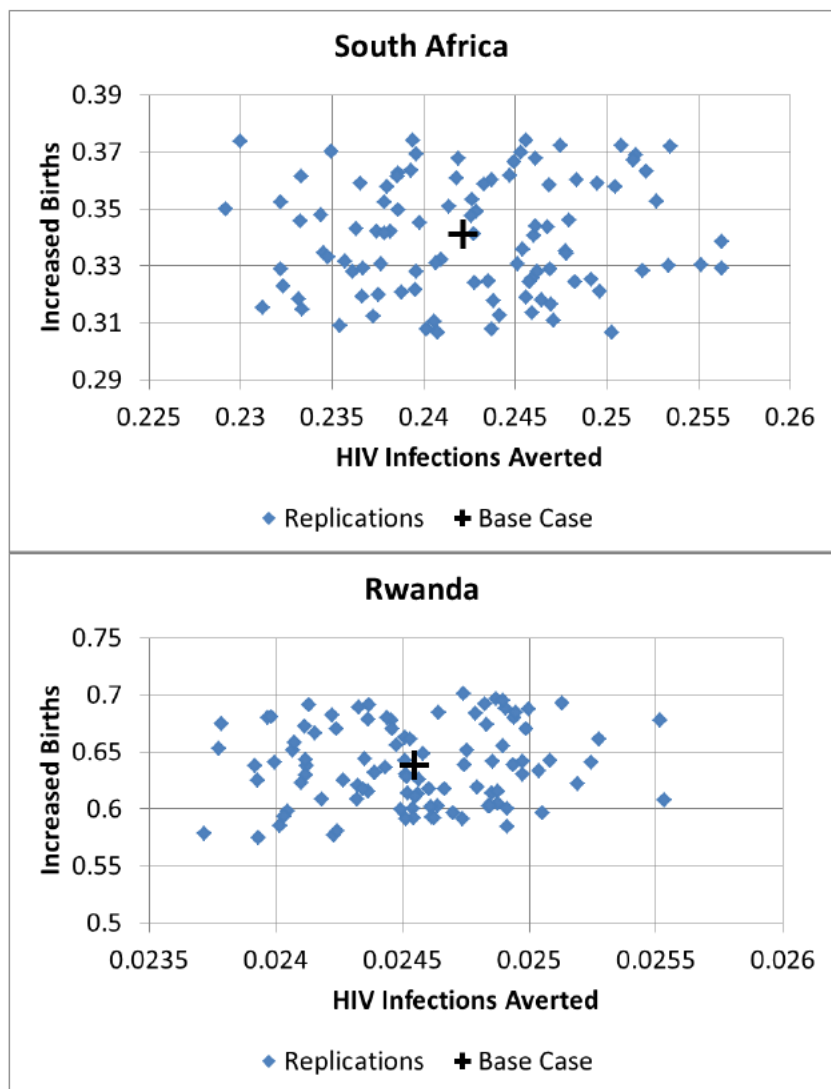


Figure A.3 Sensitivity analysis for each country





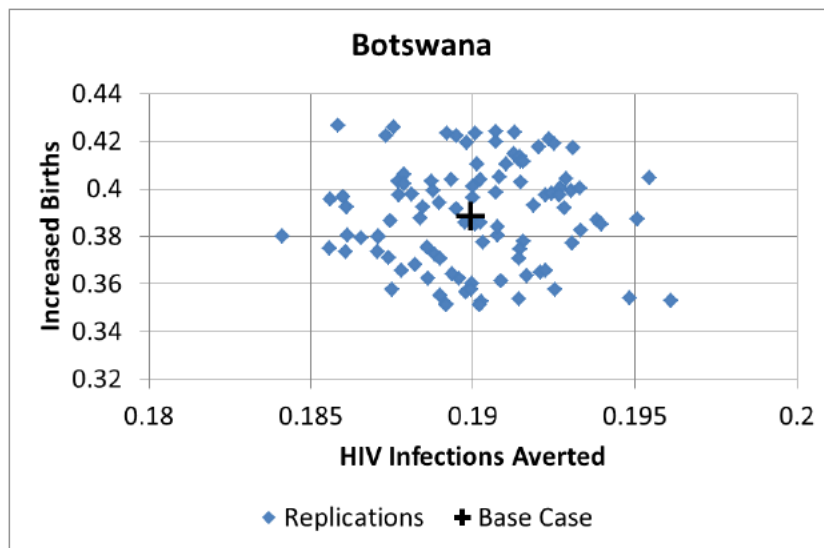


Figure A.4 Probabilistic sensitivity analysis. The figure shows for each country the base case and the 100 replications of a probabilistic sensitivity analysis where some parameters (HIV acquisition and transmission rate with POIHC, birth rate, and initial contraceptive use of POIHC, initial contraceptive use of MC) were independently drawn from a uniform distribution ranging from -10% to $+10\%$ of the baseline value.

Notation	
Total Population (ages 15-49), $N(t)$	
Female and male susceptible population, $S_f(t), S_m(t)$	
Female and male infected population, $I_f(t), I_m(t)$	
Male-to-female and female-to-male and infection rate, β_f, β_m	
Annual population growth rate, α	
Annual mortality rate of HIV+ (deaths per year), γ	
Annual rate of vertical transmission per infected female, δ	
Annual rate of births per capita, ϕ	
Cumulative number of new infections, J	
Cumulative number of cases of vertical transmission, V	
Cumulative number of births, B	
Female, male, and total prevalence, $2I_f/N, 2I_m/N, (I_f+I_m)/N$	
Portfolio of contraceptive use, $P=(p_{PIHC}, p_{IUD}, p_{FEC}, p_{VM}, p_{MC}, p_{OTH}, p_{NON})$	
Female and male infection rates assuming no protection, β_{0f}, β_{0m}	
Relative risk of female (male) infection using contraceptive x , $\xi_x^f, (\xi_x^m)$	
Total relative risk of $j=\{m,f\}$ infection using $i=\{m,f\}$ contraceptives in current portfolio, ξ_{ij}	
Per capita birth rate assuming no protection, ϕ_0	
Relative risk of pregnancy (compared to no protection), χ	
Probability of mother-to-child (vertical) transmission	
Overall, with mother on ART, without mother on ART: v, π, π'	
Percentage of pregnant females on ART, p_{ART}	
Current portfolio of contraceptive use, P_0	
Current birth rate, τ	

Table A.1. Notation used in the model

Parameter	Value *
$N(0)$	1, 1, 1, 1, 1
β_f	0.0873, 0.0885, 0.1304, 0.0776, 0.0939
β_m	0.0437, 0.0442, 0.0652, 0.0388, 0.0470
δ	0.0034, 0.0050, 0.0015, 0.0045, 0.0013
φ	0.0319, 0.0435, 0.0193, 0.0361, 0.0220
β_{of}	0.0799, 0.0882, 0.1185, 0.0732, 0.1030
β_{om}	0.0364, 0.0424, 0.0527, 0.0342, 0.0496
ζ_{ff}	1.108, 1.043, 1.142, 1.076, 1.041
ζ_{mf}	.9856, .9624, .9632, .9848, .8760
ζ_{mm}	.9856, .9624, .9632, .9848, .8760
ζ_{fm}	1.216, 1.085, 1.284, 1.152, 1.081
φ_0	0.0616, 0.0762, 0.0471, 0.0611, 0.0420
χ	.5182, .5711, .4101, .5912, .5243
ν	.1067, .1151, .0752, .1235, .0605

* Values are in order: Kenya, Zambia, South Africa, Rwanda and Botswana

Table A.2. Values of calculated parameters for baseline scenario.

	Kenya		Zambia		South Africa		Rwanda		Botswana	
	Male	Female	Male	Female	Male	Female	Male	Female	Male	Female
Baseline*	0.02	0.03	0.04	0.07	0.10	0.14	0.01	0.01	0.08	0.11
POIHC to NON**	-19.58	-14.10	-8.72	-5.58	-24.30	-18.38	-14.64	-9.96	-7.96	-5.18
POIHC to OTH	-19.58	-14.10	-8.72	-5.58	-24.30	-18.38	-14.64	-9.96	-7.96	-5.18
POIHC to IUD & MC	-10.38	-7.63	-5.57	-3.91	-13.45	-10.50	-7.87	-5.50	-8.02	-6.46
POIHC to MC	-10.79	-8.04	-6.75	-5.01	-14.50	-11.58	-8.31	-5.94	11.98	-10.27
POIHC to FC	-9.28	-8.12	-4.10	-3.20	-11.62	-10.45	-6.90	-5.74	-3.71	-2.96
POIHC to VM	-3.85	-4.19	-1.71	-1.65	-4.98	-5.49	-2.84	-2.94	-1.56	-1.55

* Values are a fraction of the initial population.

** Values for scenarios 1-6 are percent change compared to baseline.

Table A.3a. Cumulative number of new infections over 15 years by sex

	Kenya		Zambia		South Africa		Rwanda		Botswana	
	Male	Female	Male	Female	Male	Female	Male	Female	Male	Female
Baseline*	0.02	0.03	0.04	0.07	0.10	0.14	0.01	0.01	0.08	0.11
POIHC to NON**	0.47	0.47	0.38	0.39	2.54	2.53	0.14	0.15	0.61	0.58
POIHC to OTH	0.47	0.47	0.38	0.39	2.54	2.53	0.14	0.15	0.61	0.58
POIHC to IUD & MC	0.25	0.25	0.25	0.28	1.41	1.45	0.08	0.08	0.62	0.73
POIHC to MC	0.26	0.27	0.30	0.35	1.51	1.59	0.08	0.09	0.92	1.16
POIHC to FC	0.22	0.27	0.18	0.22	1.21	1.44	0.07	0.08	0.29	0.33
POIHC to VM	0.09	0.14	0.08	0.12	0.52	0.76	0.03	0.04	0.12	0.18

* Values are a fraction of the initial population.

** Values for scenarios 1-6 are change compared to baseline (10^{-3}).

Table A.3b. Cumulative number of new infections over 15 years by sex

Marginal Analysis of New Infections in Zambia																	
Contraceptive Type	Change in Each Contraceptive Type per Each Scenario							Change in Relative Risk				Change in New Infections (per 1000)					
	Baseline (%)	NON	OTH	IUD & MC	MC	FC	VM	FF	MF	MM	FM	NON	OTH	IUD	MC	FC	VM
POIHC	8.5	-8.5	-8.5	-4.2	-4.2	-2.1	-2.1	0.46	0.04	0.04	0.92	5.2	5.2	2.6	2.6	1.3	1.3
IUD	1.0	-	-	-	-	-	-	-0.04	0.04	0.04	-0.09	-	-	0.0	-	-	-
FC	-	-	-	-	-	2.1	-	-0.80	0.04	0.04	-0.85	-	-	-	-	1.5	-
VM	-	-	-	-	-	-	2.1	-0.58	0.04	0.04	-0.09	-	-	-	-	-	0.7
MC	4.7	-	-	1.2	2.4	-	-	-0.04	-0.76	-0.76	-0.09	-	-	1.0	2.0	-	-
OCP	11.0	-	-	-	-	-	-	-0.04	0.04	0.04	-0.09	-	-	-	-	-	-
OTH	16.5	-	8.5	-	-	-	-	-0.04	0.04	0.04	-0.09	-	0.1	-	-	-	-
NON	58.3	8.5	-	2.8	1.9	-	-	-0.04	0.04	0.04	-0.09	0.1	-	0.0	0.0	-	-
Any BC	41.7	-8.5	-	-2.8	-1.9	-	-2.1				Total Change	5.3	5.3	3.6	4.6	2.8	2.0
Any BC/HIVp	4.7	-	-	1.2	2.4	2.1	-				Simulation	7.8	7.8	5.2	6.5	4.1	1.9

Marginal Analysis of Births in Zambia														
Contraceptive Type	Contraceptive usage as fraction of the population							Difference in Pregnancy Risk Compared to Average	Increase in Births per 1000					
	Baseline (%)	Change in Percentage Points							NON	OTH	IUD & MC	MC	FC	VM
		NON	OTH	IUD & MC	MC	FC	VM							
POIHC	8.5	-8.5	-8.5	-4.2	-4.2	-2.1	-2.1	-0.53	45.4	45.4	22.4	22.4	11.2	11.2
IUD	1.0	-	-	-	-	-	-	-0.56	-	-	-1.7	-	-	-
FC	-	-	-	-	-	2.1	-	-0.35	-	-	-	-	-7.4	-
VM	-	-	-	-	-	-	2.1	0.29	-	-	-	-	-	6.0
MC	4.7	-	-	1.2	2.4	-	-	-0.41	-	-	-5.0	-9.9	-	-
OCP	11.0	-	-	-	-	-	-	-0.48	-	-	-	-	-	-
OTH	16.5	-	8.5	-	-	-	-	-0.26	-	-22.4	-	-	-	-
NON	58.3	8.5	-	2.8	1.9	-	-	0.29	24.3	-	8.0	5.4	-	-
Any BC	41.7	-8.5	-	-2.8	-1.9	-	-2.1	Total change	69.7	23.0	23.8	17.9	3.8	17.2
Any BC/HIVp	4.7	-	-	1.2	2.4	2.1	-	Simulation Result	88.3	29.1	33.4	23.3	4.8	16.4

Marginal Analysis of New Infections in South Africa																	
Contracepti ve Type	Change in Each Contraceptive Type per Each Scenario							Change in Relative Risk				Change in New Infections (per 1000)					
	Baseline (%)	NON	OTH	IUD & MC	MC	FC	VM	FF	MF	MM	FM	NON	OTH	IUD	MC	FC	VM
POIHC	28.4	-28.4	-28.4	-14.2	14.2	-7.1	-7.1	0.36	0.04	0.04	0.72	27.6	27.6	13.7	13.7	6.9	6.9
IUD	1.0	-	-	-	-	-	-	-0.14	0.04	0.04	-0.28	-	-	0.1	-	-	-
FC	-	-	-	-	-	7.1	-	-0.90	0.04	0.04	-1.04	-	-	-	-	11.8	-
VM	-	-	-	-	-	-	7.1	-0.68	0.04	0.04	-0.28	-	-	-	-	-	6.6
MC	4.6	-	-	1.2	2.3	-	-	-0.14	-0.76	-0.76	-0.28	-	-	2.5	4.8	-	-
OCP	0.9	-	-	-	-	-	-	-0.14	0.04	0.04	-0.28	-	-	-	-	-	-
OTH	15.0	-	28.4	-	-	-	-	-0.14	0.04	0.04	-0.28	-	7.4	-	-	-	-
NON	50.1	28.4	-	12.8	11.9	-	-	-0.14	0.04	0.04	-0.28	7.4	-	3.4	3.1	-	-
Any BC	49.9	-28.4	-	-12.8	11.9	-	-7.1				Total Change	35.0	35.0	19.7	21.7	18.7	13.5
Any BC/HIVp	4.6	-	-	1.2	2.3	7.1	-				Simulation	50.7	50.7	28.5	31.1	26.5	12.8

Marginal Analysis of Births in South Africa														
Contraceptive Type	Contraceptive usage as fraction of the population							Difference in Pregnancy Risk Compared to Average	Increase in Births per 1000					
	Baseline (%)	Change in Percentage Points							NON	OTH	IUD & MC	MC	FC	VM
		NON	OTH	IUD & MC	MC	FC	VM							
POIHC	28.4	-28.4	-28.4	-14.2	-14.2	-7.1	-7.1	-0.46	129.8	129.8	64.9	64.9	32.5	32.5
IUD	1.0	-	-	-	-	-	-	-0.48	-	-	-1.4	-	-	-
FC	-	-	-	-	-	7.1	-	-0.28	-	-	-	-	-19.7	-
VM	-	-	-	-	-	-	7.1	0.36	-	-	-	-	-	25.8
MC	4.6	-	-	1.2	2.3	-	-	-0.34	-	-	-4.0	-7.8	-	-
OCP	0.9	-	-	-	-	-	-	-0.41	-	-	-	-	-	-
OTH	15.0	-	28.4	-	-	-	-	-0.19	-	-53.1	-	-	-	-
NON	50.1	28.4	-	12.8	11.9	-	-	0.36	103.1	-	46.5	43.2	-	-
Any BC	49.9	-28.4	-	-12.8	-11.9	-	-7.1	Total change	232.9	76.7	105.9	100.3	12.8	58.2
Any BC/HIVp	4.6	-	-	1.2	2.3	7.1	-	Simulation Result	182.3	60.0	83.2	78.6	10.0	33.8

Marginal Analysis of New Infections in Rwanda																	
Contraceptive Type	Change in Each Contraceptive Type per Each Scenario							Change in Relative Risk				Change in New Infections (per 1000)					
	Baseline (%)	NON	OTH	IUD & MC	MC	FC	VM	FF	MF	MM	FM	NON	OTH	IUD	MC	FC	VM
POIHC	15.2	-15.2	-15.2	-7.6	-7.6	-3.8	-3.8	0.42	0.02	0.02	0.85	0.3	0.3	0.2	0.2	0.1	0.1
IUD	0.2	-	-	-	-	-	-	-0.08	0.02	0.02	-0.15	-	-	0.0	-	-	-
FC	-	-	-	-	-	3.8	-	-0.84	0.02	0.02	-0.91	-	-	-	-	0.1	-
VM	-	-	-	-	-	-	3.8	-0.62	0.02	0.02	-0.15	-	-	-	-	-	0.1
MC	1.9	-	-	0.5	1.0	-	-	-0.08	-0.78	-0.78	-0.15	-	-	0.0	0.0	-	-
OCP	6.4	-	-	-	-	-	-	-0.08	0.02	0.02	-0.15	-	-	-	-	-	-
OTH	12.7	-	15.2	-	-	-	-	-0.08	0.02	0.02	-0.15	-	0.0	-	-	-	-
NON	63.6	15.2	-	7.1	6.7	-	-	-0.08	0.02	0.02	-0.15	0.0	-	0.0	0.0	-	-
Any BC	36.4	-15.2	-	-7.1	-6.6	-	-3.8				Total Change	0.4	0.4	0.2	0.2	0.2	0.1
Any BC/HIVp	1.9	-	-	0.5	1.0	3.8	-				Simulation	2.9	2.9	1.6	1.7	1.5	0.7

Marginal Analysis of Births in Rwanda														
Contraceptive Type	Contraceptive usage as fraction of the population							Difference in Pregnancy Risk Compared to Average	Increase in Births per 1000					
	Baseline (%)	Change in Percentage Points							NON	OTH	IUD & MC	MC	FC	VM
		NON	OTH	IUD & MC	MC	FC	VM							
POIHC	15.2	-15.2	-15.2	-7.6	-7.6	-3.8	-3.8	-0.56	85.3	85.3	42.7	42.7	21.3	21.3
IUD	0.2	-	-	-	-	-	-	-0.58	-	-	-0.6	-	-	-
FC	-	-	-	-	-	3.8	-	-0.38	-	-	-	-	-14.5	-
VM	-	-	-	-	-	-	3.8	0.26	-	-	-	-	-	9.8
MC	1.9	-	-	0.5	1.0	-	-	-0.44	-	-	-2.2	-4.4	-	-
OCP	6.4	-	-	-	-	-	-	-0.51	-	-	-	-	-	-
OTH	12.7	-	15.2	-	-	-	-	-0.29	-	-44.3	-	-	-	-
NON	63.6	15.2	-	7.1	6.7	-	-	0.26	39.3	-	18.4	17.3	-	-
Any BC	36.4	-15.2	-	-7.1	-6.6	-	-3.8	Total change	124.6	41.0	58.2	55.6	6.8	31.2
Any BC/HIVp	1.9	-	-	0.5	1.0	3.8	-	Simulation Result	126.6	41.7	59.5	56.6	6.9	23.5

Marginal Analysis of New Infections in Botswana																	
Contraceptive Type	Change in Each Contraceptive Type per Each Scenario							Change in Relative Risk				Change in New Infections (per 1000)					
	Baseline (%)	NON	OTH	IUD & MC	MC	FC	VM	FF	MF	MM	FM	NON	OTH	IUD	MC	FC	VM
POIHC	8.1	-8.1	-8.1	-4.0	-4.0	-2.0	-2.0	0.46	0.12	0.12	0.92	17.7	17.7	8.7	8.7	4.3	4.3
IUD	1.7	-	-	-	-	-	-	-0.04	0.12	0.12	-0.08	-	-	-0.1	-	-	-
FC	-	-	-	-	-	2.0	-	-0.80	0.12	0.12	-0.84	-	-	-	-	3.8	-
VM	-	-	-	-	-	-	2.0	-0.58	0.12	0.12	-0.08	-	-	-	-	-	1.5
MC	15.5	-	-	3.9	7.8	-	-	-0.04	-0.68	-0.68	-0.08	-	-	9.7	19.4	-	-
OCP	14.3	-	-	-	-	-	-	-0.04	0.12	0.12	-0.08	-	-	-	-	-	-
OTH	4.8	-	8.1	-	-	-	-	-0.04	0.12	0.12	-0.08	-	-2.2	-	-	-	-
NON	55.6	8.1	-	-	-3.7	-	-	-0.04	0.12	0.12	-0.08	-2.2	-	0.1	1.0	-	-
Any BC	44.4	-8.1	-	-	3.7	-	-2.0				Total Change	15.5	15.5	18.3	29.1	8.1	5.9
Any BC/HIVp	15.5	-	-	3.9	7.8	2.0	-				Simulation	12.0	12.0	13.5	20.8	6.2	3.0

Marginal Analysis of Births in Botswana														
Contraceptive Type	Contraceptive usage as fraction of the population							Difference in Pregnancy Risk Compared to Average	Increase in Births per 1000					
	Baseline (%)	Change in Percentage Points							NON	OTH	IUD & MC	MC	FC	VM
		NON	OTH	IUD & MC	MC	FC	VM							
POIHC	8.1	-8.1	-8.1	-4.0	-4.0	-2.0	-2.0	-0.49	40.0	40.0	19.8	19.8	9.9	9.9
IUD	1.7	-	-	-	-	-	-	-0.52	-	-	-2.1	-	-	-
FC	-	-	-	-	-	2.0	-	-0.31	-	-	-	-	-6.3	-
VM	-	-	-	-	-	-	2.0	0.33	-	-	-	-	-	6.5
MC	15.5	-	-	3.9	7.8	-	-	-0.37	-	-	-14.6	-29.2	-	-
OCP	14.3	-	-	-	-	-	-	-0.44	-	-	-	-	-	-
OTH	4.8	-	8.1	-	-	-	-	-0.22	-	-18.2	-	-	-	-
NON	55.6	8.1	-	-	-3.7	-	-	0.33	26.4	0.0	-0.7	-12.1	-	-
Any BC	44.4	-8.1	-	-	3.7	-	-2.0	Total change	66.4	21.9	2.5	-21.5	3.6	16.4
Any BC/HIVp	15.5	-	-	3.9	7.8	2.0	-	Simulation Result	46.4	15.3	1.7	-14.7	2.5	8.6

Table A.4 Marginal Analysis

a) New Infections Averted								
Country	Case	Baseline	POIHC to NON	POIHC to OTH	POIHC to IUD	POIHC to MC	POIHC to FC	POIHC to VM
Kenya	Case 0	0.06	9.38	9.38	5.02	5.26	4.92	2.32
	Case 1	0.06	12.25	12.25	6.52	6.75	5.56	2.78
	Case 2	0.06	6.89	6.89	3.74	3.99	4.39	1.89
	Case 3	0.06	3.84	3.84	2.18	2.43	3.74	1.37
	Case 4	0.06	3.23	3.23	1.88	2.13	3.61	1.31
	Case 5	0.06	0.00	0.00	0.27	0.53	2.95	0.79
Zambia	Case 0	0.11	7.77	7.77	5.21	6.50	4.06	1.92
	Case 1	0.11	10.68	10.68	6.67	7.94	4.75	2.40
	Case 2	0.11	5.61	5.61	4.13	5.44	3.55	1.55
	Case 3	0.11	3.26	3.26	2.96	4.29	2.99	1.14
	Case 4	0.11	2.40	2.40	2.54	3.87	2.79	1.02
	Case 5	0.11	0.00	0.00	1.35	2.70	2.24	0.62
South Africa	Case 0	0.24	50.70	50.70	28.51	31.08	26.53	12.76
	Case 1	0.24	64.70	64.70	35.85	38.32	29.47	15.06
	Case 2	0.24	37.38	37.38	21.70	24.37	23.85	10.42
	Case 3	0.24	20.30	20.30	12.95	15.76	20.40	7.41
	Case 4	0.24	18.44	18.44	12.04	14.85	20.04	7.44
	Case 5	0.24	0.00	0.00	2.96	5.91	16.54	4.39
Rwanda	Case 0	0.02	2.90	2.90	1.58	1.69	1.52	0.71
	Case 1	0.02	3.89	3.89	2.09	2.20	1.75	0.87
	Case 2	0.02	2.12	2.12	1.18	1.30	1.35	0.58
	Case 3	0.02	1.22	1.22	0.72	0.84	1.14	0.42
	Case 4	0.02	0.95	0.95	0.59	0.70	1.08	0.39
	Case 5	0.02	0.00	0.00	0.12	0.24	0.88	0.24
Botswana	Case 0	0.19	11.98	11.98	13.47	20.83	6.20	2.95
	Case 1	0.19	16.28	16.28	15.55	22.81	7.22	3.68
	Case 2	0.19	8.52	8.52	11.82	19.26	5.38	2.36
	Case 3	0.19	4.78	4.78	10.03	17.56	4.50	1.71
	Case 4	0.19	3.81	3.81	9.56	17.10	4.27	1.57
	Case 5	0.19	0.00	0.00	7.75	15.39	3.38	0.93

b) Decrease in Prevalence								
Country	Case	Baseline	POIHC to NON	POIHC to OTH	POIHC to IUD	POIHC to MC	POIHC to FC	POIHC to VM
Kenya	Case 0	0.05	5.37	5.37	2.88	3.01	2.82	1.42
	Case 1	0.05	6.99	6.99	3.73	3.86	3.18	1.71
	Case 2	0.05	3.94	3.94	2.14	2.28	2.51	1.16
	Case 3	0.05	2.19	2.19	1.24	1.39	2.14	0.84
	Case 4	0.05	1.86	1.86	1.08	1.23	2.07	0.81
	Case 5	0.05	0.00	0.00	0.15	0.30	1.69	0.48
Zambia	Case 0	0.11	4.44	4.44	2.97	3.71	2.31	1.18
	Case 1	0.11	6.07	6.07	3.80	4.52	2.71	1.47
	Case 2	0.11	3.19	3.19	2.35	3.10	2.02	0.95
	Case 3	0.11	1.84	1.84	1.68	2.43	1.70	0.70
	Case 4	0.11	1.39	1.39	1.45	2.21	1.60	0.63
	Case 5	0.11	0.00	0.00	0.77	1.54	1.27	0.38
South Africa	Case 0	0.20	29.24	29.24	16.47	17.95	15.32	7.88
	Case 1	0.20	37.23	37.23	20.68	22.09	17.01	9.30
	Case 2	0.20	21.55	21.55	12.53	14.07	13.77	6.43
	Case 3	0.20	11.69	11.69	7.46	9.08	11.77	4.57
	Case 4	0.20	10.71	10.71	6.98	8.60	11.58	4.59
	Case 5	0.20	0.00	0.00	1.71	3.41	9.55	2.71
Rwanda	Case 0	0.02	1.66	1.66	0.90	0.97	0.87	0.44
	Case 1	0.02	2.21	2.21	1.19	1.25	1.00	0.54
	Case 2	0.02	1.21	1.21	0.67	0.74	0.77	0.35
	Case 3	0.02	0.69	0.69	0.41	0.48	0.65	0.26
	Case 4	0.02	0.54	0.54	0.34	0.40	0.62	0.24
	Case 5	0.02	0.00	0.00	0.07	0.13	0.50	0.15
Botswana	Case 0	0.19	6.81	6.81	7.65	11.81	3.52	1.81
	Case 1	0.19	9.23	9.23	8.82	12.92	4.10	2.25
	Case 2	0.19	4.83	4.83	6.70	10.91	3.05	1.44
	Case 3	0.19	2.69	2.69	5.68	9.95	2.55	1.04
	Case 4	0.19	2.18	2.18	5.43	9.71	2.43	0.96
	Case 5	0.19	0.00	0.00	4.40	8.73	1.92	0.56

c) Increase in Births								
Country	Case	Baseline	POIHC to NON	POIHC to OTH	POIHC to IUD	POIHC to MC	POIHC to FC	POIHC to VM
Kenya	Case 0	0.56	181.37	59.72	84.01	84.23	9.95	33.64
	Case 1	0.56	181.37	59.72	84.01	84.23	9.95	33.64
	Case 2	0.56	181.37	59.72	84.01	84.23	9.95	33.64
	Case 3	0.56	181.37	59.72	84.01	84.23	9.95	33.64
	Case 4	0.56	181.37	59.72	84.01	84.23	9.95	33.64
	Case 5	0.56	181.37	59.72	84.01	84.23	9.95	33.64
Zambia	Case 0	0.77	88.25	29.06	33.44	23.30	4.84	16.37
	Case 1	0.77	88.25	29.06	33.44	23.30	4.84	16.37
	Case 2	0.77	88.25	29.06	33.44	23.30	4.84	16.37
	Case 3	0.77	88.25	29.06	33.44	23.30	4.84	16.37
	Case 4	0.77	88.25	29.06	33.44	23.30	4.84	16.37
	Case 5	0.77	88.25	29.06	33.44	23.30	4.84	16.37
South Africa	Case 0	0.34	182.34	60.04	83.22	78.57	10.01	33.82
	Case 1	0.34	182.34	60.04	83.22	78.57	10.01	33.82
	Case 2	0.34	182.34	60.04	83.22	78.57	10.01	33.82
	Case 3	0.34	182.34	60.04	83.22	78.57	10.01	33.82
	Case 4	0.34	182.34	60.04	83.22	78.57	10.01	33.82
	Case 5	0.34	182.34	60.04	83.22	78.57	10.01	33.82
Rwanda	Case 0	0.64	126.62	41.69	59.50	56.55	6.95	23.49
	Case 1	0.64	126.62	41.69	59.50	56.55	6.95	23.49
	Case 2	0.64	126.62	41.69	59.50	56.55	6.95	23.49
	Case 3	0.64	126.62	41.69	59.50	56.55	6.95	23.49
	Case 4	0.64	126.62	41.69	59.50	56.55	6.95	23.49
	Case 5	0.64	126.62	41.69	59.50	56.55	6.95	23.49
Botswana	Case 0	0.39	46.36	15.27	1.75	-14.69	2.54	8.60
	Case 1	0.39	46.36	15.27	1.75	-14.69	2.54	8.60
	Case 2	0.39	46.36	15.27	1.75	-14.69	2.54	8.60
	Case 3	0.39	46.36	15.27	1.75	-14.69	2.54	8.60
	Case 4	0.39	46.36	15.27	1.75	-14.69	2.54	8.60
	Case 5	0.39	46.36	15.27	1.75	-14.69	2.54	8.60

d) Increase in Vertical Transmission								
Country	Case	Baseline	POIHC to NON	POIHC to OTH	POIHC to IUD	POIHC to MC	POIHC to FC	POIHC to VM
Kenya	Case 0	0.00	0.52	0.11	0.24	0.24	-0.02	0.09
	Case 1	0.00	0.45	0.06	0.21	0.21	-0.03	0.09
	Case 2	0.00	0.53	0.12	0.25	0.24	-0.01	0.09
	Case 3	0.00	0.55	0.14	0.25	0.25	-0.01	0.10
	Case 4	0.00	0.62	0.19	0.28	0.28	0.00	0.10
	Case 5	0.00	0.63	0.21	0.29	0.29	0.00	0.10
Zambia	Case 0	0.01	0.53	0.10	0.17	0.07	-0.03	0.09
	Case 1	0.01	0.45	0.02	0.13	0.03	-0.05	0.08
	Case 2	0.01	0.54	0.11	0.17	0.08	-0.03	0.09
	Case 3	0.01	0.56	0.13	0.18	0.08	-0.02	0.10
	Case 4	0.01	0.63	0.20	0.22	0.12	-0.01	0.10
	Case 5	0.01	0.65	0.21	0.22	0.12	0.00	0.11
South Africa	Case 0	0.00	1.21	0.27	0.55	0.50	-0.03	0.23
	Case 1	0.00	1.06	0.15	0.49	0.44	-0.05	0.21
	Case 2	0.00	1.25	0.30	0.57	0.52	-0.03	0.23
	Case 3	0.00	1.31	0.34	0.59	0.54	-0.02	0.24
	Case 4	0.00	1.46	0.45	0.65	0.60	0.00	0.25
	Case 5	0.00	1.52	0.50	0.68	0.62	0.01	0.26
Rwanda	Case 0	0.00	0.19	0.04	0.09	0.08	-0.01	0.03
	Case 1	0.00	0.16	0.01	0.07	0.07	-0.01	0.03
	Case 2	0.00	0.19	0.04	0.09	0.08	-0.01	0.03
	Case 3	0.00	0.20	0.05	0.09	0.08	-0.01	0.03
	Case 4	0.00	0.22	0.07	0.10	0.10	0.00	0.04
	Case 5	0.00	0.23	0.08	0.11	0.10	0.00	0.04
Botswana	Case 0	0.00	0.27	0.06	-0.04	-0.19	-0.01	0.05
	Case 1	0.00	0.24	0.03	-0.06	-0.20	-0.01	0.04
	Case 2	0.00	0.28	0.07	-0.04	-0.18	-0.01	0.05
	Case 3	0.00	0.28	0.07	-0.04	-0.18	-0.01	0.05
	Case 4	0.00	0.31	0.10	-0.03	-0.17	0.00	0.05
	Case 5	0.00	0.32	0.11	-0.02	-0.17	0.00	0.05

Table A.5 Sensitivity Analysis a) New Infections, b) Prevalence, c) Births d) Vertical Transmission: Case 1: 100% increase in HIV acquisition risk, 100% increase in female-to-male transmission risk, Case 2: 50% increase in HIV acquisition risk, 50% increase in female-to-male transmission risk, Case 3: 50% increase in HIV acquisition risk, 0% increase in female-to-male transmission risk, Case 4: 0% increase in HIV acquisition risk, 50% increase in female-to-male transmission risk, Case 5: 0% increase in HIV acquisition risk, 0% increase in female-to-male transmission risk.

Chapter 3

Modeling Measles Outbreaks: The Role of Community Structure and Non-Vaccination Interventions

3.1 Introduction

Epidemiologists distinguish between endemic diseases such as HIV or hepatitis C and endemic outbreaks, for which measles are a prototypical example. Measles is a highly contagious disease, which can have severe consequences, especially in young infants and pregnant women. The introduction of the measles, mumps, and rubella (MMR) vaccine decreased the incidence of measles in the 1990s and measles was declared eliminated in the United States in 2000 [39]. However, large scale measles outbreaks in Europe increasingly expose susceptible American travelers, reintroducing measles to the United States dozens of times each year by infected travelers and leading to large outbreaks in the US [40, 41, 42].

This paper analyzes three issues, two modeling issues (the necessary model size and complexity), and one public health issue (the benefit of a better detection and faster response to epidemic outbreaks). To describe the issue of model size, we note that the median outbreak size is five individuals [41], though the largest outbreaks in the US infected a couple hundred [4, 43]. This raises the question whether an outbreak simulation needs to include 100, 1000, or 100,000 individuals. For a more detailed agent based simulation (e.g., [44]), the related question is what to include in the simulation (i.e., where to draw the boundary of the study population). For example, when simulating an outbreak initiated by an elementary school child, the question would be whether to only simulate the individuals in the child's school or whether to include the elementary school of the initial child's older teenage sister in the simulation¹. For our stochastic model (Section 3.2) of epidemic outbreaks, we prove (Section 3.3) that the outbreak size converges in distribution as the model population increases and then use simulation (Section 3.4) to determine that a population size of 1000 is more than enough for measles.

The most related stream of research on model size is on the critical community size, the minimum size within which measles can persist. Early research suggested that the critical community size was 250000-500000 [45, 46, 47]. Keeling and Grenfell (1997) [48] examine outbreaks in 60 towns in Eng-

¹The issue whether to simulate the older sister's high school also touches on the issue of heterogeneity which we address later.

land and Wales from 1944 to 1968 to arrive at a critical community size for measles of 250000-400000. However, this research is about the pre-vaccination era where significant fractions of children became infected each year [49] instead of the isolated outbreaks we see today.

We now turn to the issue of model complexity or heterogeneity. Measles is a well-studied disease and most models incorporate some form of heterogeneity such as age-structure, spatial heterogeneity, vaccination coverage heterogeneity, etc. for the transmission probability between people [50, 51, 44, 52, 53, 54, 55]. The cost of including many and detailed forms of heterogeneity are additional model complexity and simulation run-time. To capture the range of potential heterogeneity we use a stylized model that divides the population into two subgroups and varies the mixing rates between the subgroups and within each subgroup, an approach similar to that taken by one of the earliest measles model with heterogeneity [50] and a recent study on the role of heterogeneity [44]. We compare this to a model with a homogeneous, undivided population. We prove that outbreaks are longer when the population is divided into two subgroups and that heterogeneity increases infection risk. However, using simulation, we find that the actual effect of varying the model heterogeneity is quite small.

Turning to the public health interventions for measles, the effectiveness of measles vaccination in school and university settings is well studied [56, 57, 58, 59]. The increased rates of religious or philosophical exemptions to vaccination requirements [60] make it important to consider alternative interventions. However, studies of other interventions such as a better case detection strategy or a faster public health response are missing. Our study includes three intervention variables: the vaccinated fraction, the probability of detecting and reporting a case, and the number of days' delay in the public health response. Our numerical simulations show that a better detection strategy can be just as important in increasing vaccination coverage.

The rest of the manuscript is outlined as follows. In Section 3.2, we describe stochastic discrete-time model that keeps track of the number of infected, exposed, etc. in each of the two population subgroups and the outcomes of interest: the expected size of the outbreak, the probability of the outbreak spreading between groups, and the size of the outbreak in the second community group, the one that was not initially infected. This model forms the framework for the theoretical results in Section 3.3 and the simulation results in Section 3.4. We conclude in Section 3.5.

3.2 Methods

We build a model to determine the size of the outbreak, when it is detected, and how various interventions or model characteristics affect it. We model the progression of measles using an SEIR (Susceptible-Exposed-Infectious-Recovered) model with a 10-day exposed (or incubation) period and a subsequent 8-day infectious period [61, 62]. We model measles transmission using a stochastic simulation model. Every day there is a chance that the outbreak will be detected and contained shortly thereafter with a vaccination program (or closure of the school). At the start of the simulation there is a single infectious individual, and we focus on the final size C of the outbreak. We assume a constant population of size n , of which a fraction v is unvaccinated.

We examine the effect of community structure on measles transmission, how people connect and communicate. To study the extremes of community structure we compare the case of homogeneous mixing (where everybody is part of the same community group) and the case where the population is divided into two equal-sized community groups. In the latter case, an individual is more likely to infect another individual in the same community group than one in the other group, where this relative likelihood is the mixing parameter, ϕ . When $\phi = 0$, there is no mixing between the two community groups and when

$\phi = 1$, there is no preferential mixing between the groups (there is effectively only one big community group). However, we keep the basic reproductive number, R_0 , the same in all cases. We also consider the less stylized case of transmission between two schools, a smaller elementary school of 500 students and a high school with 2000 students.

We assume each infectious individual has a daily probability q of being detected and reported. Throughout the paper we use detecting and reporting a case interchangeably. Note that sometimes the outbreak is never detected due to underreporting [63] and not detecting the imported cases [64]. Under homogenous mixing (a single community group), the simulation is halted ϵ days after the first case is detected and reported. With two community groups, we assume that public health officials separately stop the spread in each group. Specifically, the spread in each group is halted ϵ days after the first detected case in that group and $\sigma = 5$ days after the first detected case in the other group, whichever comes first. The details of the model and outcome measures of interest are given next.

3.2.1 Model

Here we describe the equations that describe our model. Below is also a table of notation (Table 3.2). We start by describing the equations for the case with a homogenous population. We let n be the size of our population and ν the fraction unvaccinated. Thus before the outbreak, there are νn susceptible individuals in the population. On day t , this population is divided into $S(t)$ susceptible individuals; $E_k(t)$ individuals in day $k \in \{1, \dots, 10\}$ of the incubation period; $I_k(t)$ individuals in day $k \in \{1, \dots, 8\}$ of the infectious period; and $R(t)$ individuals who have recovered from an infection. Note that $S(t) + \sum_{k=1}^{10} E_k(t) + \sum_{k=1}^8 I_k(t) + R(t) = \nu n$. We denote the total number of infectious individuals on day t as $T(t)$. On day 1 we have exactly one infected individual, who is in the first day of the incubation period. Then the system dynamics are given by equations (1-9), where $N(t)$ is the number of new infections on day t .

$$T(t) = \sum_{k=1}^8 I_k(t) \quad (3.17)$$

$$S(t) + \sum_{k=1}^{10} E_k(t) + T(t) + R(t) = \nu n \quad (3.18)$$

$$S(1) = \nu n - 1, E_1(1) = 1 \quad (3.19)$$

$$S(t) = S(t-1) - N(t) \quad (3.20)$$

$$E_1(t) = N(t) \quad (3.21)$$

$$E_j(t) = E_{j-1}(t), \text{ for } j \in \{2, \dots, 10\} \quad (3.22)$$

$$I_1(t) = E_1(t) \quad (3.23)$$

$$I_k(t) = E_{k-1}(t), \text{ for } k \in \{2, \dots, 8\} \quad (3.24)$$

$$R(t) = R(t-1) + I_8(t-1) \quad (3.25)$$

Equations (10-11) describe the number of new infections $N(t)$ on day t . These are distributed binomially with each susceptible individual on day $t-1$ having a probability $p(t)$ of being infected. We calculate $p(t)$ as the total number of infectious individuals, $T(t-1)$ times the basic reproductive number R_0 and divided by the length of the infectious period, 8, and the size of the population. This way the first infectious individual will create R_0 secondary cases in expectation, if the entire population were

susceptible.

$$N(t) \sim B(S(t-1), p(t)) \quad (3.26)$$

$$p(t) = T(t-1)(R_0/8)/n \quad (3.27)$$

Equations (12-14) describe the final size of the outbreak, C . Here $Z(t)$ is the number of cases reported on day t . Thus on day τ , ϵ days after the first reported case, public health authorities stop the spread of the outbreak. Then the outbreak size C is the number of infected and removed individuals on that day. Here the number of reported cases per day, $Z(t)$, is distributed binomially with each infectious individual having a probability q of being detected and reported each day.

$$Z(t) \sim B(T(t), q) \quad (3.28)$$

$$\tau = \epsilon + \min(t : Z(t) \geq 1) \quad (3.29)$$

$$C = nv - S(t) \quad (3.30)$$

Now we describe the small modifications to the model we need to make for the case where the population is divided into two subgroups. These are summarized in equations (15-22). We let $n^x, S^x, E_k^x, I_k^x, R^x, T^x, N^x, Z^x, \tau^x, C^x$, and p^x be the respective variables for subgroup $x = 1$ or $x = 2$. Initially, the only infected individual is in subgroup 1 and is in the first day of the incubation period. Recall that the mixing parameter ϕ is the ratio of the probability of infecting a susceptible in the other subgroup to the probability of infecting a susceptible in the same subgroup. Thus we adjust the infection probabilities $p_x(t)$ by considering both infections from the same subgroup and the other subgroup. In addition, we not only have a public health intervention stopping the spread in a subgroup ϵ days after the first reported case in that subgroup but also σ days after the first reported case in the other subgroup.

$$S^1 = n^1 v - 1, E_1^1(1) = 1 \quad (3.31)$$

$$S^2 = n^2 v \quad (3.32)$$

$$n = n^1 + n^2 \quad (3.33)$$

$$p^1(t) = T^1(t-1)\left(\frac{R_0}{8}\right)\frac{1}{n_1}\frac{1}{(1+\phi)} + T^2(t-1)\left(\frac{R_0}{8}\right)\frac{1}{n_1}\frac{\phi}{(1+\phi)} \quad (3.34)$$

$$p^2(t) = T^2(t-1)\left(\frac{R_0}{8}\right)\frac{1}{n_2}\frac{1}{(1+\phi)} + T^1(t-1)\left(\frac{R_0}{8}\right)\frac{1}{n_2}\frac{\phi}{(1+\phi)} \quad (3.35)$$

$$\tau^1 = \min(\epsilon + \min(t : Z^1(t) \geq 1), \sigma + \min(t : Z^2(t) \geq 1)) \quad (3.36)$$

$$\tau^2 = \min(\epsilon + \min(t : Z^2(t) \geq 1), \sigma + \min(t : Z^1(t) \geq 1)) \quad (3.37)$$

$$C = C^1 + C^2 \quad (3.38)$$

3.2.2 Outcome Measures of Interest

The outcome measures of interest are the final number of infected cases, C , and the number of infected cases in the second community group, the one that was initially **not** infected, C_2 , when we assume two

Table 3.1: Notation

Population, n
Fraction unvaccinated, ν
Susceptibles, $S(t)$
Individuals in day k of the incubation period, $E_k(t)$
Individuals in day k of the infectious period, $I_k(t)$
Recovered individuals, $R(t)$
Total number of infectious individuals, $T(t)$
New infections, $N(t)$
Probability a susceptible individuals is infected, $p(t)$
Basic reproduction number, R_0
Number of reported individuals, $Z(t)$
Prob. an infectious case is reported each day, q
Lag between first reported case in same subgroup and stopping, ϵ
Stopping time, τ
Mixing parameter, ϕ
Lag between first reported case in other subgroup and stopping, σ

community groups. Specifically, we will examine the expected size of the outbreak, $E[C]$, the probability the outbreak is at least five, $P[C \geq 5]$, and the probability that the outbreak spreads beyond the initial community group, $P[C_2 \geq 1]$. In the post elimination era, the median number of outbreak cases is 5; thus five measles cases are chosen as a metric of interest [41].

3.3 Theoretical Results

In this section, we analyze the impact of population size of the model and the model heterogeneity. For the population size, we show that the epidemic outcome measures converge to finite limits as the population size increases. For the model heterogeneity, we show that the outbreaks are larger and longer in the inhomogeneous case.

Recall, that $T = T(t, n)$ denotes the infectious population and $N(t, n)$ the number of new infections, where we may drop the dependence on t and n for convenience. We let $C(t, n)$ denote the size of the outbreak at time t , the number who are infected or recovered at this time. Also, let $Y(i)$ denote whether individual i becomes infected at time t , so that the number of new infections at time t is $N = \sum_{i=1}^{n\nu-C} Y(i)$. Then for given T , $Y(i)$ are independent and identically distributed with $Y(1) \sim \text{Bern}(T\beta/n)$ where β is effective contact rate.

3.3.1 Convergence As Population Size Increases

Our goal in this subsection is to prove for the homogenous case that the infection process converges to a finite limit as the population size becomes large. For notational convenience, let $C' = C(t+1, n)$ and note that $C(1, n) = 1$ and $C' = C + N$.

Theorem 1 *For all t , $T(t, n)$, $N(t, n)$, and $C(t, n)$ converge in distribution as $n \rightarrow \infty$.*

Lemma 1 For all t , $E[C(t, n)] = O(1)$ as $n \rightarrow \infty$. If $E[C] = O(1)$ then $E[C'] = O(1)$ as $n \rightarrow \infty$.

Proof Since $C(1, n) = 1$, it suffices to show that if $E[C] = O(1)$ then $E[C'] = O(1)$ as $n \rightarrow \infty$ and induct on t . Since $N \leq \sum_{i=1}^n \nu Y(i)$ a.s., $E[N|T] \leq n\nu E[Y(1)|T] = \beta\nu T$ a.s. Hence, $E[N] = E[E[N|T]] \leq \beta\nu E[T]$. Therefore, $E[C'] = E[C] + E[N] \leq E[C] + \beta\nu E[T]$. Now since $T \leq C$ a.s., $E[T] \leq E[C]$, and $E[C'] \leq E[C](1 + \beta\nu)$, proving the induction hypothesis. $\square$

Lemma 2 Suppose that $X(n) \geq 0$ a.s. and $E[X(n)] \rightarrow 0$. Then, $X(n) \rightarrow 0$ in distribution.

Proof Since $X(n) \geq 0$ a.s., $X(n) = |X(n) - 0|$ a.s. Hence, $E[X(n)] = E[|X(n) - 0|] \rightarrow 0$. Finally, convergence in mean implies convergence in distribution. $\square$

Lemma 3 If $w(n) \rightarrow w(\infty)$ and $y(n) \rightarrow 0$, then $(1 + w(n)/n)^{n(v-y(n))} = e^{\nu w(\infty)}$.

Proof Note that, $\log((1 + w(n)/n)^{n(v-y(n))}) = (v - y(n))n \log(1 + w(n)/n)$ equals $(v - y(n))nw(n)/n(1 + o(1)) = (v - y(n))w(n)(1 + o(1))$. Since this is a product of limits, it converges to $\nu w(\infty)$. The claim follows from the continuity of $\exp()$. $\square$

Lemma 4 Suppose $T(t, n) \rightarrow T(t, \infty)$ in distribution. Then, $N(t, n) \rightarrow \text{Poisson}(T(t, \infty)\beta\nu)$ in distribution.

Proof The characteristic function of a random variable X is $E[e^{isX}]$ as a function of s . Let $\phi_n(s) = E[e^{isN(t, n)}]$ be the characteristic function of $N(t, n)$. Define the characteristic function of $\text{Poisson}(T(t, \infty)\beta\nu)$ as

$$\begin{aligned} (\phi(s)) &= E[e^{is\text{Poisson}(T(t, \infty)\beta\nu)}] \\ &= E[E[e^{is\text{Poisson}(T(t, \infty)\beta\nu)} | T(t, \infty)]] \\ &= E[e^{(T(t, \infty)\beta\nu)(e^{is} - 1)}] \end{aligned}$$

where we use the fact that the characteristic function of $\text{Poisson}(\lambda)$ is $e^{\lambda(e^{is} - 1)}$. By Levy's continuity theorem, it suffices to show for all $s > 0$, that $\phi_n(s) \rightarrow \phi(s)$ as $n \rightarrow \infty$.

Note that, $e^{isN} = \prod_{i=1}^{n\nu-C} e^{isY(i)}$ a.s. Thus,

$$\begin{aligned} (E[e^{isN} | T, C]) &= E[e^{isY(1)} | T, C]^{n\nu-C} = (e^{isT\beta/n} + (1 - T\beta/n))^{n(v-C/n)} \\ &= (1 + \frac{T\beta(e^{is} - 1)}{n})^{n(v-C/n)} \end{aligned}$$

where we use the fact that the characteristic function of $\text{Bern}(p)$ is $1 - p + pe^{is}$. By Lemmas 1 and 2, $C/n \rightarrow 0$ in distribution. We now construct a joint probability space where $T(t, n) \rightarrow T(t, \infty)$ a.s. and

$C(t, n)/n \rightarrow 0$ a.s. Then, by Lemma 3, $E[e^{isN}|T, C] \rightarrow e^{((T(t, \infty)\beta v(e^i s - 1)))}$ a.s. Since e^{isN} is bounded, we can apply the bounded convergence theorem to take the expectation of both sides and prove the claim. $\square$

Proof of Theorem 1: Since $C(1, n) = 1$ and T and C are the sum of previous new infections, we can apply induction using Lemma 4.

3.3.2 Outbreaks Are Longer With Two Subgroups

We now turn our attention model heterogeneity. There are two reasons why heterogeneity increases the expected outbreak size. The first reason is when a community is divided into two subgroups then the outbreak lasts longer when the community is divided into two subgroups. This is because if a case is detected in one subgroup, then public health officials do not intervene in both subgroups at the same time. Instead, they intervene later in the subgroup that did not detect the first case. In the homogenous case, the intervention when the first case is detected affects more people immediately (i.e., the entire population) than in the two-subgroup case. To formalize this, let $\tau^{(1)} = \max(\tau^1, \tau^2)$ and $\tau^{(2)} = \min(\tau^1, \tau^2)$. Also, define the case detection times in the homogenous case and the case with two subgroups, $\tilde{\tau} = \min(t : Z(t) \geq 1)$ and $\tilde{\tau}^x = \min(t : Z^x(t) \geq 1)$, respectively.

Theorem 2 *Suppose that delay between detection and intervention is larger if the case is detected in the other subgroup, $\sigma \geq \epsilon$, and that everything else, the total number of infections and detected cases is the same as in the homogenous case, $Z^1(t) + Z^2(t) = Z(t)$ for all t . Then the stopping time in the homogenous case is the same as in one of the subgroups but before the stopping time in the other subgroup, $\tau = \tau^{(2)} \leq \tau^{(1)} \leq \tau + \sigma - \epsilon$.*

Proof Note that $\tilde{\tau} = \min(\tilde{\tau}^1, \tilde{\tau}^2)$. Hence, $\tau = \tilde{\tau} + \epsilon = \min(\tilde{\tau}^1, \tilde{\tau}^2) + \epsilon = \tau^2$. Obviously, $\tau^{(2)} \leq \tau^{(1)}$. The last inequality, $\tau^{(1)} \leq \tau + \sigma - \epsilon$ follows from the definition of the stopping time (20)-(21). $\square$

Additional intuition for $\tilde{\tau} \leq \tilde{\tau}^1, \tilde{\tau}^2$ (a consequence of $\tilde{\tau} = \min(\tilde{\tau}^1, \tilde{\tau}^2)$) is the following. The number of person-days of infected individuals required until a case is detected is a geometric random variable with mean $1/q$. In populations with more infected individuals (which larger populations tend to have), one has a higher probability of detecting a case, and detection occurs more quickly.

3.3.3 Heterogeneity In Infection Risk

We now focus on the second way that model heterogeneity leads to larger outbreaks: the variation in the infection risk between individuals in the case with two subgroups. We will abstract beyond two subgroups and consider the general case where the infection probability varies among individuals, $Y(i) \sim \text{Bern}(p_i)$. To ensure a valid comparison, we use the average infection probability in the homogenous case, $p^* = \frac{1}{n} \sum_{i=1}^n p_i = T\beta/n$, ensuring that $E[N] = \sum_{i=1}^{nv-C} p_i = T\beta(v - C/n)$ in any case. Formally, we again assume that the $Y(i)$ are independent given T . Our argument that outbreaks grow more slowly in the homogenous case has two parts (later we show that outbreaks are also more likely to fizzle out in the homogenous case). First, we will show that the variance in the number of new infections, $\text{Var}[N]$, is maximized in the homogenous case. Second, we go on to show that increases in the variance of the number of infectious individuals, $\text{Var}[T]$, will decrease the expected number of new infections. While we prove both parts below, we are unable to combine them to argue about the outbreak size at any particular time, $C(t)$, due to difficulty with the induction step.

Lemma 5 *The homogenous case maximizes $\text{Var}[N|T, C]$.*

Proof Let $f(x) = x(1 - x)$. Note that, $\text{Var}[N|T, C] = \text{Var}[\sum_{i=1}^{nv-C} Y(i)|T, C] = \sum_{i=1}^{nv-C} f(p_i)$. Since f is concave, $(1/k)\sum_{i=1}^k f(p_i) \leq f(p^*)$ where p^* is the average infection probability as defined above. Thus, $\sum_{i=1}^{nv-C} f(p_i) \leq (nv - C)f(p^*) = \text{Var}[B(n, p^*)|T, C]$, proving the claim. $\square$

Lemma 6 *Let $C=T+R$. Then, $E[N|R]$ decreases in $\text{Var}[T|R]$.*

Proof By definition,

$$\begin{aligned} (E[N|R]) &= E\left[\sum_{i=1}^{nv-T-R} Y(i)|R\right] = E\left[E\left[\sum_{i=1}^{nv-T-R} Y(i)|T\right]|R\right] \\ &= E\left[\sum_{i=1}^{nv-T-R} p(i)|R\right] \end{aligned}$$

Since $Y(i)$ are iid, $1/n \sum_{i=1}^n p_i = T\beta/n$, $E[N|R] = E[(nv - T - R)(T\beta/n)|R]$, and thus

$$\begin{aligned} (E[N|R]) &= E[T|R]\beta\left(v - \frac{R}{n}\right) - E[T^2|R]\beta/n \\ &= E[T|R]\beta\left(v - \frac{R}{n}\right) - (E[T|R]^2 + \text{Var}[T|R])\beta/n \end{aligned}$$

proving the claim. $\square$

In addition to the expected number of new infections being smaller in the homogenous case, we can also show that the probability of no new infections is larger in that case. This is particularly important at the beginning of an outbreak, which can fizzle out if the first few infected individuals do not manage to infect any others.

Lemma 7 *The homogenous case maximizes $P[N = 0|T, C]$.*

Proof The proof is similar to that of Lemma 5. Note that, $P[N = 0|T, C] = \prod_{i=1}^{nv-C} P[Y(i) = 0|T] = \prod_{i=1}^{nv-C} (1 - p_i)$. Since $\log(1 - x)$ is a concave function, $(1/k)\sum_{i=1}^k \log(1 - p_i) \leq \log(1 - p^*)$. Hence, $\log P[N = 0|T, C] \leq (nv - C)\log(1 - p^*) = \log P[B(n, p^*) = 0]$, proving the claim. $\square$

3.4 Numerical Results

In this section, we conduct numerical experiments to verify our analytical results and provide additional insights. We also examine the impact of public health interventions such as increasing the vaccinated population, increasing daily reporting probability, and a faster public health response. All parameters not varied in a figure or table are at their baseline value as given in Table 3.2.

Table 3.2: Parameters

Parameter	Value	Source
Population, n	1000	
Basic reproduction number, R_0	18	[49]
Fraction unvaccinated, ν	10%	[4]
Incubation period (days)	10	[61]
Infectious period (days)	8	[62]
Mixing parameter, ϕ	0.5	
Prob. of case being reported each contagious day, q	0.1	
Lag between first reported case in subgroup stopping, ϵ (days)	3	
Lag between first reported case in other subgroup and stopping, σ (days)	5	

3.4.1 The Impact of Population Size

In this section, we not only confirm the convergence proof in Section 3.3.1 but also determine the population size at which we are “close enough” for public health purposes. Figure 3.1 examines the effect of changing the population size on the average outbreak size and the probability of a large outbreak, one where five or more become infected. After increasing initially, both outcome measures vary little as the population increases beyond a thousand. This implies that there is no gain to simulating measles outbreaks with larger populations.

3.4.2 The Impact of Heterogeneity

In this section, we confirm that heterogeneity results in more infections and numerical results allow us to determine whether the increase is significant. Figure 3.1 also shows how the average outbreak size and the probability of a large outbreak are greater for a community divided into two subgroups than a homogenous community, while keeping the total population size fixed. Figure 3.2 focuses on this aspect of the community structure by comparing the distribution of the outbreak size in the two cases. It also shows that large outbreaks are slightly more frequent in the inhomogeneous case. This may be for multiple reasons and our theoretical results suggest the difference in stopping times between subgroups and the heterogeneity of the force of infection.

Figure 3.3 further explores the community structure in the case of two subgroups by varying the mixing parameter, which specifies how separate the two subgroups are. The mixing parameter is the ratio of the likelihood that an individual infects someone in the other subgroup rather than someone in the same subgroup. When the mixing parameter is zero there is no mixing between subgroups. In that case, we can focus on the subgroup with the initial infection (since it won’t ever spread to the other subgroup) and view it as a single homogenous population of half the original size. This partly explains why the expected number of infected cases is lower with no mixing. When the mixing parameter is one, there is complete mixing. This is like a single homogenous population except that as discussed in the previous figure, there is a lag until the outbreak is contained. Figure 3.3 shows that as the mixing parameter increases, the expected size of the outbreak increases slightly.

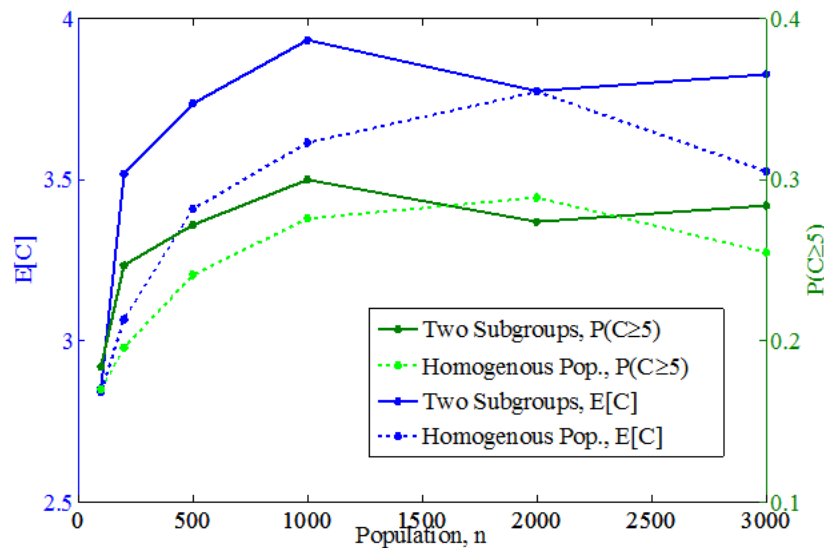


Figure 3.1: The average number of infected cases ($E[C]$, left, blue) and the probability of a large outbreak ($P[C \geq 5]$, right, green) as we vary the size of the population. The dashed lines are for the case of a homogenous population while the solid lines are for the case where the population is divided into two equal-sized subgroups. There were 1000 replications and the standard error is less than 0.11 for $E[C]$ and 0.02 for $P[C \geq 5]$.

3.4.3 The Impact of Public Health Interventions

We now turn our sensitivity analysis to important parameters for public health interventions. Figure 3.4 focuses on the effect of changing the vaccinated fraction of the population. For each percentage point that the vaccinated fraction increases, the expected size of the outbreak decreases by approximately 0.35 and the probability of a large outbreak decreases by approximately 3.3 percentage points.

Two other important parameters for public health interventions are the daily probability of detecting an infected case and the lag between the first detected case and a public health response stopping the spread of the outbreak. Table 3.3 examines the sensitivity to these parameters for the average outbreak size ($E[C]$), the probability of a large outbreak with five or more infected ($P[C \geq 5]$), and for the case of a population with two subgroups, the probability the infection spreads from one subgroup to the other ($P[C_2 \geq 1]$). While the benefits of increasing the detection probability are significant, we see diminishing returns for greater increases. On average, for each percentage point that the detection probability increases, the expected size of the outbreak decreases by approximately 0.21 and the probability of a large outbreak decreases by approximately 2.3 percentage points. Decreasing the lag between detection and response also decreases the outcome measures (except for the probability of the second subgroup being infected, which changes little) but only by a small amount. By the time a case is detected, the measles has already spread to other subgroup thus a few days of delay does not make much difference and change in the probability of the second subgroup being infected was less than other outcomes. Thus, a few days of delay does not make much difference and the change in $P[C_2 \geq 1]$ is even less.

We simulated several scenarios to conduct a sensitivity analysis for the less stylized case of transmission between two schools, a smaller elementary school of 500 students and a high school with 2000 students. We assume that initially there is a single infected case in the elementary school. In Table 3.4,

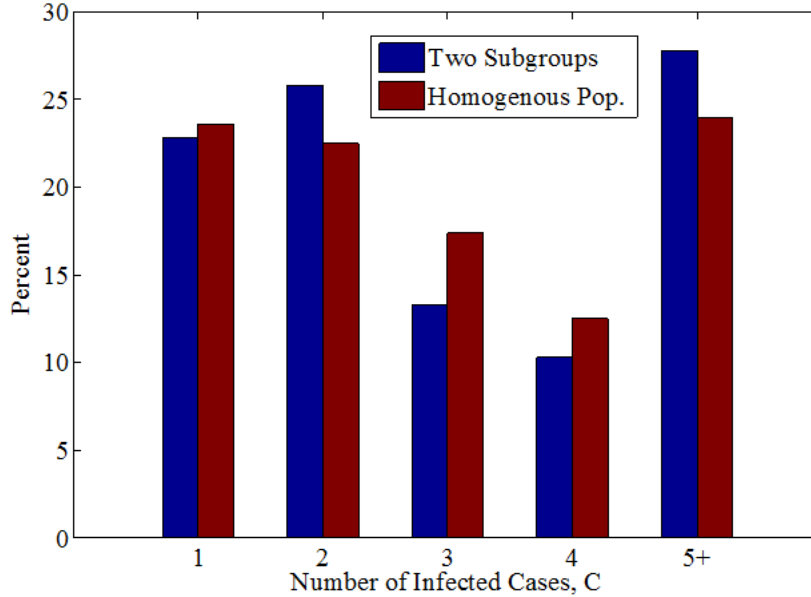


Figure 3.2: Histogram of total number of infected cases in a homogenous population of 1000 or a population with two subgroups each of 500 individuals. There were 1000 replications.

we examine different choices for the fraction vaccinated, the mixing parameter, the daily probability of detecting a case, and the lag between the first reported case and stopping the outbreak. The outcome measures are the same as in Table 3.3, with $P[C_2 \geq 1]$ now denoting the probability the outbreak spreads from the elementary school to the high school. The results have the same trend as those in the case of a population with two equal-sized subgroups we discussed previously.

3.5 Discussion

We present a stochastic simulation model of a measles outbreak to investigate the impact of the simulated population size, model complexity and public health interventions. We analyze this for both the homogeneous mixing (i.e., a single community) and the inhomogeneous case where the population is divided into two community groups. We observe the final number of infected cases, the number of infected cases in the second community group where initial number of infected people is zero, and the probability of the infection spreading from the initially infected community group to the second community group.

Our theoretical analysis shows that results converge to a finite limit as the population size increases and give multiple reasons (from the way that outbreaks are stopped to the role of heterogeneity) why the expected outbreak size for a community divided into two subgroups is greater than in the case of a single homogenous community. We also performed sensitivity analysis on various model parameters as well as suggested public health interventions such as increasing the probability of reporting a case and decreasing the lag between detecting a case and the public health authority stopping transmission.

We find that the size of the modeled population ceases to matter after exceeding 1000 individuals, matching the theoretical result. This makes sense because with 10% unvaccinated, the unvaccinated

Table 3.3: Additional Sensitivity Analysis

		Two Community Groups			One Community	
		$E[C]$	$P(C \geq 5)$	$P(C_2 \geq 1)$	$E[C]$	$P(C \geq 5)$
Daily reporting probability, q	0.05	5.9	0.465	0.65	5.2	0.403
	0.1	3.9	0.287	0.554	3.4	0.257
	0.15	3.2	0.199	0.533	2.8	0.159
	0.2	2.7	0.116	0.499	2.6	0.108
Lag between first case reported and simulation stopped, ϵ	0	3.1	0.212	0.558	2.8	0.185
	1	3.4	0.247	0.531	3	0.192
	2	3.7	0.262	0.567	3.4	0.239
	3	3.7	0.251	0.551	3.5	0.242

Here C_2 is the number of infected cases in the other subgroup (not the one with the initial infection). The number of replications is 1000. The standard error is less than 0.18 for $E[C]$ and 0.02 for $P[C \geq 5]$ and $P[C_2 \geq 1]$.

Table 3.4: Sensitivity Analysis for Unequal Sized Groups

		$E[C]$	$P(C \geq 5)$	$P(C_2 \geq 1)$
Fraction Vaccinated, $\%, (1 - \nu)$	80	8.3	0.621	0.823
	85	5.9	0.464	0.717
	90	4	0.308	0.579
	95	2.1	0.076	0.306
Mixing Parameter, ψ	0	3.4	0.237	0
	0.25	3.8	0.285	0.451
	0.5	3.9	0.297	0.57
	0.75	3.9	0.279	0.659
	1	4.1	0.315	0.677
Daily Reporting Probability, q	0.05	5.9	0.46	0.639
	0.1	3.8	0.276	0.552
	0.15	3.1	0.2	0.525
	0.2	2.8	0.128	0.492
Lag Between First Case Reported and Simulation Stopped, ϵ	0	3.1	0.211	0.537
	1	3.6	0.261	0.572
	2	3.8	0.286	0.556
	3	3.9	0.307	0.572

Scenario with two unequal sized subgroups, an elementary school of 500 and a high school of 2000 individuals. The initial infection is in the elementary school and C_2 is the number of cases in the high school. There were 1000 replications. The standard error is less than 0.23 for $E[C]$ and 0.02 for $P[C \geq 5]$ and $P[C_2 \geq 1]$.

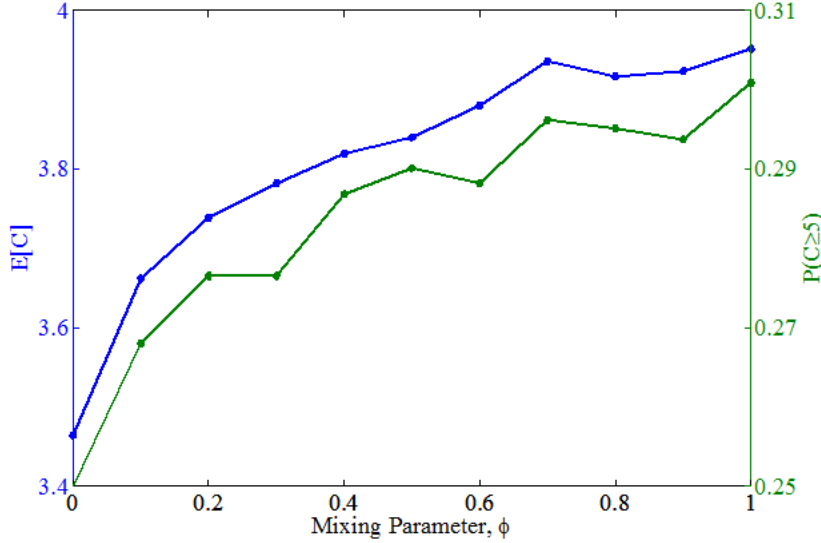


Figure 3.3: The average number of infected cases ($E[C]$, left, blue) and the probability of a large outbreak ($P[C \geq 5]$, right, green) as the mixing parameter changes in a population with two subgroups each of 500 individuals. The mixing parameter is the ratio of the probability of infecting a susceptible in the other subgroup to the probability of infecting a susceptible in the same subgroup. There were 10000 replications and the standard error is less than 0.04 for $E[C]$ and 0.02 for $P[C \geq 5]$.

population is 100 individuals, large enough so that the stochastic effects of small numbers are small. For smaller populations, the outbreaks are smaller as they are limited by the population size. Overall, the two-subgroup model has slightly larger outbreaks compared to the homogenous model, also matching the theoretical results. Increasing the mixing parameter in the two-subgroup model increases the average outbreak size but the effects are small. In addition, the probability of an outbreak spreading from one subgroup to another is around 50% and is somewhat insensitive to the parameters.

Increasing the probability of detecting a case is a good alternative to vaccination strategy. A 1.5 percentage point in probability of detecting a case can achieve the equivalent of a 1 percentage point increase in vaccination coverage. This is especially important in communities where the vaccination refusal rates are higher. However, decreasing the delay between case detection and stopping the outbreak had only a minimal effect. The asymmetric case with a bigger high school and a smaller elementary school has similar results. Larger outbreaks (i.e., due to lower vaccination rates or smaller detection probabilities) are also associated with higher probabilities of the infection spreading from one subgroup to another.

One limitation of our study is that we only examined two community structures (one with homogenous mixing and one with two subgroups). The hope is that these represent the extremes and other community structures will have outcomes between the two extremes. Our model does not capture the possibility that the unvaccinated individuals are all clustered around the index case in the social network; our model assumed that the unvaccinated are spread randomly within each subgroup. Social network data could evaluate the likelihood of this possibility. This could also be captured in our model by restricting the study population to the smaller set of social contacts of the index case (e.g., maybe 20 individuals with 40% vaccinated). We are not aware of such social network data. Another limitation is that there

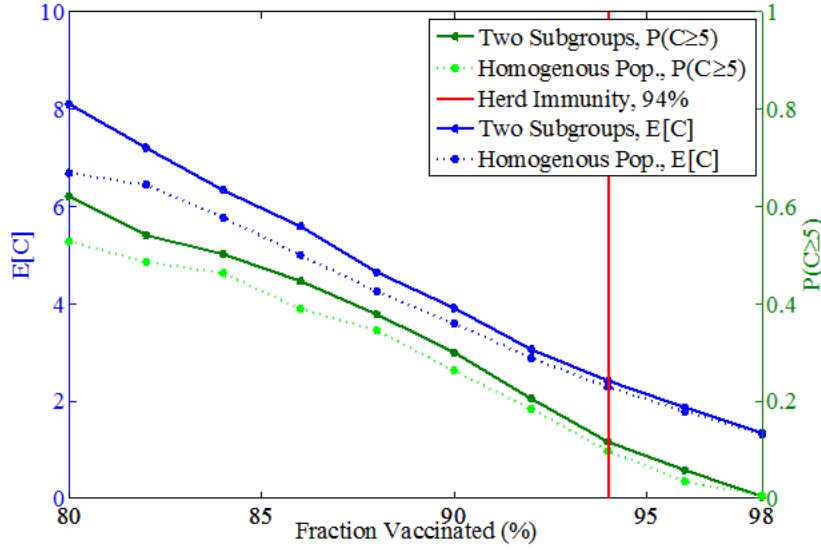


Figure 3.4: The average number of infected cases ($E[C]$, left, blue) and the probability of a large outbreak ($P[C \geq 5]$, right, green) as we vary the fraction vaccinated. The dashed lines are for the case of a homogenous population of 1000 while the solid lines are for the case where the population is divided into two equal-sized subgroups of 500. The red line indicates the deterministic herd immunity threshold. There were 1000 replications and the standard error is less than 0.22 for $E[C]$ and 0.02 for $P[C \geq 5]$.

are no studies to justify a value for the case detection probability in our model. Such a study would be valuable, especially since the true extent of measles outbreaks is uncertain (i.e., underreporting [63]) and since we find that increasing the detection probability would be an effective way of decreasing the size of measles outbreaks. In addition, we also did not consider details such as variability in the measles progression across different individuals. However, this was not the focus of our study.

There are three conclusions we can draw. First, unless the population of interest is actually small, modelers can save a lot of computational effort by limiting their simulated population to 1000 individuals, with little cost in accuracy. Second, this is irrespective of the heterogeneity in the population. As heterogeneity increases, the metrics of outbreak severity increase slightly. Thus, unless focusing on the impact of specific interactions, one can exclude complex heterogeneity from the simulation models. This would also speed up validation, calibration, and sensitivity analysis. Third, in contexts where it is difficult to raise the vaccination coverage (which are many), alternative interventions such as better detection of infections and a faster response to any detected cases should be considered.

Acknowledgement

We thank Rebecca Wurtz for her insights.

Chapter 4

Capturing Real-Time Data in Disaster Response Logistics

4.1 Introduction

Over the past two decades, the field of humanitarian logistics has progressed significantly, with a growing number of researchers and practitioners studying problems, such as relief distribution, post-disaster debris removal, and evacuation of affected populations. Much work within the academic community has focused on the development and application of operations research tools for humanitarian logistics (e.g., see recent surveys: [65, 66, 67, 68]). As the ultimate goals and benefits of these efforts are to improve real world applications, integration of in-field data is an important, yet often overlooked, component of such humanitarian logistics models. For example, in their recent review, Sangiamkul & Hillegrersberg (2011) [69] identify only two papers [70, 71] out of 30 surveyed that use realtime data in logistical modeling. In another survey, Ortuño et al. (2013) [72] describe only two papers among 87 that integrate dynamically updated data. Finally, Özdamar & Ertem (2015) [73] acknowledge three papers [70, 71, 179] out of 110 studies mentioned in their review of humanitarian logistic models, solutions and technologies that capture such data. Outside these academic disciplines, extensive efforts have been made in information communication technology, especially in regard to the use of social media and crowdsourcing in disaster management. At the same time, the volume, accuracy, accessibility and level of detail of near real-time data emerging from disaster-affected regions continue to significantly improve. Considerable efforts are currently focused on the collection, aggregation and dissemination of field data, which, together with the help of the humanitarian logistics decision tools, have the potential to considerably impact relief efforts. In this paper, we present a structure for analyzing humanitarian logistics data, explore the process of retrieving real post-disaster relief data from sources available online, and examine the data for the purpose of integrating data streams into response logistics models to facilitate future modeling.

We present a framework for evaluating real-time humanitarian logistics data focused on use in mathematical modeling. The framework reflects the integration of our recent experience of near real-time data collection, a survey of different communities producing data and disciplines using data, and a development of measures to evaluate the quality of data and applicability to other disasters for logistical modeling. We also discuss how to measure the attributes of the framework and describe the application of this framework to a case study of near real-time data collection in the days following the landfall of Typhoon Haiyan. We detail our first-hand experience of capturing data as the post-disaster response

unfolded, starting November 10, 2013 until March 31, 2014 and assess the characteristics and evolution of data pertaining to humanitarian logistics modeling. The case study, illustrating our information retrieval process, presents an example of the classification of data and data sources using the proposed framework. The logistical content analysis, using the available information following Typhoon Haiyan, examines the availability of data and informs modelers about the current state of near real-time data. This analysis illustrates what data is available, how early it is available, and how the data changes after the disaster. The study describes how our humanitarian logistics team approached the emergence of dynamic online data after the disaster and the challenges faced during the collection process, as well as recommendations to address these challenges in the future (when possible) from an academic humanitarian logistics perspective.

This study signifies the importance of an interdisciplinary team approach when exploring real-time humanitarian logistics data, its value and challenges. The retrieval of information needed for humanitarian logistics models and knowledge of in field data collection and dissemination come from a unique collaboration between logistical researchers and humanitarian practitioners. This research shows that well-formed and growing relationships allow for parties to gain insights into each other's respective use of terminology and broader domains. Such insights may enable each party to alert the other about potential opportunities for exploration, such as the uniqueness of Typhoon Haiyan with regards to public data, while the modelers can inform the practitioners and the broader humanitarian response community about the needs for accessible field-appropriate data for the on-ground-personnel or agencies to aid in their operations.

The rest of the paper is organized as follows. The following section provides some background information, which includes a literature review of different communities involved in data generation, processing and dissemination, and a description of disciplines using humanitarian logistics data. Next, section a proposed framework for humanitarian logistics data with respect to humanitarian logistics modeling is presented with the focus on data quality and applicability measures, and the framework implications for mathematical modeling. An application of this framework to the recent Typhoon Haiyan is illustrated with a logistical content analysis and describes the lessons learned. This paper concludes with final remarks on scarcity of data and points up the need for humanitarian logistics models that integrates multidisciplinary work to validate the limited data.

4.2 Disaster Response Data Stakeholders

Multiple entities play a role in the evolution of post-disaster data via collection, processing or dissemination. Furthermore, various communities are the intended users and beneficiaries of this data. Understanding the roles and motivation of the key stakeholders is essential to analyzing the emergence of near real-time data following a disaster. In this section we describe the data-gathering communities and the disciplines using the data.

4.2.1 Data-gathering Communities

Altay and Labonte (2014) [65] discuss the challenges of information management and coordination amongst intergovernmental organizations, such as OCHA, government and non-governmental organizations in the case of the Haiti response in 2010. The United Nations Foundation Disaster (2011) [75] examines the future of information-sharing in humanitarian emergencies with a focus on volunteer and

technical communities, such as OSM and Crisis Mappers (2014) [76]. The response to the Haiti earthquake highlights the disaster response operations where thousands of volunteers around the world collaborated around various information communication technologies to inform the public about the affected population. Grünewald and Binder (2010) [77] analyze the humanitarian response following the Haiti earthquake and discuss the inter-agency real-time evaluation. They also provide the list of people consulted from different organizations from government, donor representatives, international NGOs, national NGOs and UN agencies. These studies provide a good starting point for the main stakeholders in data gathering.

We also survey a number of practitioners and responders involved with Typhoon Haiyan and other major responses (e.g., the earthquake in Haiti) through the pre-existing relationships of our team members with other humanitarian practitioners and responders. These surveys and relationships brought to our attention considerable amount of disaster relief operations relevant data and their sources. In addition, the on-ground personnel involved with major responses, including Typhoon Haiyan, assisted the team with assessing reliability and understanding the nature of the datasets and the data sources.

We first present the key players we observed in the data collection, processing and dissemination, motivated by the literature, survey of practitioners and our observations of publicly available online data and information sources between November 10, 2013 and March 31, 2014 following Typhoon Haiyan.

Large International Humanitarian Response Organizations

The United Nations Office of Coordination for Humanitarian Affairs (OCHA) aims to play a critical role in "mobiliz[ing] and coordinat[ing] effective and principled humanitarian action in partnership with national and international actors in order to alleviate human suffering in disasters and emergencies" [78]. In the immediate aftermath of a disaster, when OCHA receives an international call for providing assistance, it often sends United Nations Disaster Assessment and Coordination teams to provide an initial assessment of the situation [79]. The information management activities within OCHA aim to support information collection and sharing needs of humanitarian actors to support coordination [80]. In large disasters, responding organizations often coordinate in "clusters" based upon major sectoral activities, such as logistics, health and shelter. In such cases, the information collection, management and sharing are often targeted to the specific cluster's activities. For example, in Typhoon Haiyan, specifically, OCHA facilitated the activities of 11 clusters.

Information management activities at OCHA take advantage of different types of information in different phases of the response. According to OCHA's description of its services, they aim to create and share information in media that are simple to understand and easily accessible. Datasets including common operational datasets, contact lists, and "who, what, where" data are also maintained and shared by this group. Shared documents in the form of portable data formats (PDF) reports and maps are frequently used.

Geographic information system (GIS) data plays a key role in OCHA's information management (IM) services. Recent collaborations with external organizations, including MapAction, Humanitarian OpenStreetMap Team (HOT) and GISCorps, have advanced the timely processing of map generation. These organizations sometimes leverage microtasking and crowdsourcing methods to process large amounts of geographic information from imagery datasets and non-traditional geographic sources (e.g., satellite imagery, photos).

National Government

After a disaster, the host government coordinates governmental departments and agencies, such as

the department of health, department of defense and emergency management authority, for its disaster response. For example, in the case of Typhoon Haiyan, The National Disaster Risk Reduction & Management Council (NDRRMC), which functions under the Department of National Defense, managed gathering and reporting data [81]. In addition to NDRRMC, the Department of Social Welfare and Development (DSWD) played a crucial role in providing information about affected citizens [82]. A detailed description of these two government offices' role in Typhoon Haiyan with respect to data efforts is provided in the Appendix.

Digital Humanitarians

The Digital Humanitarian Network (DHN) is a consortium of volunteer and technical communities that “provide an interface between formal, professional humanitarian organizations and informal yet skilled-and-agile volunteer and technical networks” with the purpose “to leverage digital networks in support of 21st century humanitarian response” [83]. In the case of Typhoon Haiyan, OCHA activated DHN immediately after the disaster. According to media reference, this was the first time officials were appointed to coordinate the crowdsourced mapping efforts with volunteer groups [84] during the early stage of the response. Some of these volunteer groups include HOT, Standby Task Force and MapAction. The general role of HOT is to serve as a bridge between the OSM community and the traditional humanitarian relief organizations. In the Philippines, there were more than 1000 OSM volunteers from 82 countries who provided maps to nongovernmental organizations, including Doctors without Borders [84] and the American Red Cross [85]. The Standby Task Force analyzed more than one million texts, tweets and other social media posts with the help of MicroMappers software, which uses machine-learning techniques to filter potentially relevant messages [84]. MapAction, a longtime partner of OCHA, worked in the Philippines to generate more than a hundred files per day to be shared with the disaster relief community. With all these efforts, the data from Typhoon Haiyan is a notable example of the evolution of collaboration between digital humanitarian and response agencies, where access to information, collaboration and the next steps of information-sharing were pushed forward.

Operationally-focused Humanitarian Practitioners and Responders

Humanitarian practitioners and responders (both local and international) in affected regions are often among the most knowledgeable people about the changing post-disaster environment. They are frequently aware of information sources and datasets, sometimes generating data themselves, which may not only reflect the current context but also represent information and data used by organizations for planning and executing response activities. In our experience, the pre-existing relationship of our practitioner team member with other practitioners and responders has brought tremendous value to identifying and better understanding various information outlets and how they can be integrated in future logistical models using real-time data.

As information communication technology improves, connecting with responding humanitarian organizations, theoretically, is more feasible. However, developing trusting relationships and personal networks still requires years of engagement in working with people from various backgrounds, often having different short-term goals but with common overarching missions.

4.2.2 Disciplines Using Data

Current humanitarian logistics models aim to capture real-time data in order to improve their decision support tools. Here, we discuss how the data is used in these models and the assumptions the researchers

make, especially in relation to the data availability. Various disciplines utilize humanitarian logistics data and impact the changes in data collection, processing and dissemination. Therefore, the role and purpose of each discipline as it relates to real-time data should be taken into account by the logistics modelers in order to better understand the data characteristics.

As modelers studying the real-time humanitarian logistics data focused on use in mathematical modeling ourselves, we naturally investigate the academic humanitarian logistics as one of the disciplines benefiting from data. The efforts on understanding data gathering communities (such as literature review and surveys) enlighten us about the other disciplines benefiting from humanitarian logistics data. Similarly, the work on digital humanitarians enables us to recognize the importance of information and communication technology for understanding the current status of the real-time data, its implications and challenges. Thus, in addition to the academic humanitarian logistics discipline, we also discuss the role humanitarian data plays for practitioners, the intended primary users of this data, and information and communication technology (ICT), the data collection, processing and communication facilitators.

Academic Humanitarian Logistics

We first describe the current efforts of humanitarian logistics models with real-time data. As mentioned previously, although the number of models related to humanitarian logistics is growing, models with real-time data are limited. We review those papers below.

Liu and Ye (2014) [86] present a decision model for the allocation of relief resources in natural disasters using information updates. These updates predominantly contain information on disaster states (population transfer rates) and traffic conditions (road affected level). Authors suggest that this information can be obtained from the disaster database of governmental agencies, such as the National Oceanic and Atmospheric Administration (2013) [87], the National Climatic Data Center, and the National Geophysical Data Center, among others. Liu and Ye [86] apply their model to the Wenchuan earthquake in China with data provided by China's National Committee on Disaster Reduction.

Sheu (2010) [70] develops a dynamic relief-demand management model that forecasts the demand in real-time and dynamically allocates supplies based on those forecasts, as well as urgency and population vulnerability measures. The main components of the information used in the model are 1) timevarying ratio of the estimated number of trapped survivors relative to the local population; 2) population density associated with a given affected region; 3) proportion of frail population (e.g., children and the elderly) relative to the total number of population trapped; 4) time elapsed since the most recent relief arrival; and 5) level of building damage. Sheu [70] uses the official statistics from the 921 earthquake (also known as the Jiji earthquake) special report from Taiwan to demonstrate the application of the developed model. The model contains the most detailed amount of data in comparison to other academic studies we have surveyed and is valuable for estimating regional level demand under dynamic information updates. The author also generates simulation data to replace the missing data points in an effort to tackle incomplete information.

Yi and Özdamar (2007)[71] study a dynamic coordination problem of supply distribution and transfer of injured people. They apply their model to an earthquake scenario for Istanbul. The demand distribution (number of wounded people) and supply distribution (people, fleet composition, and total capacity transport) are provided for each time period. Researchers employ the widely used data from the Earthquake Engineering Department of Bogazici University (2002) [88] for attrition numbers and possible structural damage to Istanbul, which are used to calculate the number of affected people. Information about permanent emergency units is gathered from local municipalities and the Turkish Medical Doctors Association. However, the information about the number and capacity of vehicles, the capacity

of temporary emergency units, as well as how this information is updated are not explicitly provided.

Huang et al. (2013) [179] study the impact of incorporating real-time data into disaster relief routing for search and rescue operations in the aftermath of the 2010 Haiti earthquake. They use OpenStreetMap (OSM) to obtain road and building data. They also extract demand information on collapsed structures and trapped persons using Mission 4636, a text-message communication initiative. This research provides insights into incorporating crowdsourced data into humanitarian logistics models.

In addition to the integration of data into logistical models, some researchers also study classification frameworks for humanitarian data. This work is discussed in more detail later (see Measures of Data Quality and Applicability section) as it closely relates to our developed measures of data quality and applicability.

Humanitarian Practitioners

Humanitarian practitioners often rely on situational awareness to make critical decisions in difficult situations with limited resources and time. Endsley (1988) [89] defines situational awareness as “the perception of the elements in the environment within a volume of time and space, the comprehension of their meaning and the projection of their status in the near future” The availability of timely and accurate data is critical to personnel making operational decisions. Therefore, there have been numerous and ongoing efforts to improve the collection, management and sharing of humanitarian data for humanitarian practitioners such as Humanitarian Data Exchange (HDX) (See Information and communication technology section below for more details).

There are also ongoing efforts among practitioners to build vocabulary standards for crisis management [90]. This is of particular interest to researchers, as we observe not only differences between researchers and practitioners in the terminology used, but even among various handbooks and guidelines intended for practitioners [91, 92, 93, 94, 95, 96, 97, 98, 99]. For example, while humanitarian logistics researchers extensively use the term “supply”, practitioners use “supply”, “resources”, “capacity”, “stockpile” and “availability”. Until these vocabulary standards are developed and implemented in the field, researchers should be aware of the various terms different data sources might use in the same context.

The role that data plays in humanitarian operations continues to change as data gathering, processing and sharing technologies evolve. Many humanitarian agencies actively acknowledge, assess and forecast the effects of the corresponding changes. The annual World Disaster Report 2013 from the International Federation of Red Cross and Red Crescent Societies (2013) “examines the profound impact of technological innovations on humanitarian action, how humanitarians employ technology in new and creative ways, and what risks and opportunities may emerge as a result of technological innovations” [100]. These and other similar reports can provide logistics models with insights into how data is perceived by the humanitarian practitioners discipline.

Information and Communication Technology (ICT)

Information and communication technology plays a critical role in facilitating data collection, processing and communication. From the academic perspective, with the ongoing evolution of these technologies, the studies that analyze their application to crisis and emergency management have also significantly expanded [101, 239, 103, 104, 105, 106]. The majority of this research focuses on crowdsourcing and social media applications in disaster responses [107, 108, 109, 110, 111, 112, 113, 114, 115] with particular interest in Twitter [111]. For example, Ashktorab et al. (2014) [107] present a Twitter-mining tool to classify, cluster and extract tweets. The authors include the keywords processed in their study, such as “bridge”, “intersection”, “evacuation”, “impact”, “injured” and “damage”, among others. The authors

implement their algorithm to tweets collected from 12 different crises in the United States. Purohit et al. (2013) [113] present machine-learning methods developed for social media specifically to identify needs (demands) and offers (supplies) to facilitate relief coordination, by matching the needs with offers, encompassing shelter, money, clothing, volunteer work, etc.

DeLone and McLean (1992) [116] develop an information system success model with six interdependent success variables: system quality, information quality, use, user satisfaction, individual impact, and organizational impact. Over the years, this model has been extensively studied and improved [118]. DeLone and McLean (2003) [117] later update their model with the following success variables: system quality, information quality, service quality, system use, user satisfaction, and net benefits. Each variable also has numerous dimensions. For example, relevance, understandability, accuracy, conciseness, completeness, currency, timeliness, and usability are provided for information quality, which are defined as the desirable characteristics of the system outputs such as reports and web pages. System quality is defined as the desirable characteristics of an information system and ease of use system flexibility, system reliability, and ease of learning, as well as system features of intuitiveness, sophistication, flexibility, and response times. Bharosa, Appelman, Zanten, and Zuurmond, (2009) [119] examine information and system quality as requirements for information system success during disaster management. The researchers state that, although information quality requirements are very relevant for information system success during disaster management, they are very hard to measure. In the case of systems quality measurements, much of the effort is focused on the inter-operability and ease of use.

From the practitioner's perspective, ICT has improved data collection, processing and communication in recent decades. UN OCHA's (2002) Symposium on Best Practices in Humanitarian Information Exchange resulted in humanitarian information management principles as: accessibility, inclusiveness, inter-operability, accountability, verifiability, relevance, impartiality, humanity, timeliness and sustainability. In a later symposium, reliability, reciprocity, and confidentiality were added to the list [120]. The ongoing Humanitarian Data Exchange (HDX) project, led by OCHA, aims to "make humanitarian data easily available and useful for decision-making," by bringing together multi-country, multi-sourced, curated data for analytical use through a single platform [121]. As part of HDX project, Humanitarian Exchange Language (HXL) is intended to offer the standardization of humanitarian data [122]. In order to further facilitate standardization, the HDX Quality Assurance Framework identifies five dimensions of quality as accuracy, timeliness, accessibility, interpretability and comparability [123]. Logistics modelers can benefit and often directly utilize the data collection and processing tools developed by the ICT discipline.

4.3 Framework for Analyzing Real-time Logistics Data for Modeling

This section presents the framework for analyzing real-time post-disaster data, specifically focusing on the measures that describe the quality of data and data sources, as well as their applicability to different disasters for logistics modeling to learn from past disasters. Information during the aftermath of a disaster can more frequently be found on websites and is often shared via listservs and emails. Each disaster context will vary in the degree of online information access for several reasons, such as: the type of disaster (e.g., disasters with predictable timing and location, such as hurricanes, versus disasters with unpredictable timing, such as earthquakes), availability of information technology, and level of involvement of host nations and their national and local governments. In order to better assess the quality of numerous information sources that emerge after a given disaster and their applicability to other disasters, we classify the outlets and data provided from these outlets based on a number of measures

relevant to the focus of this study. We first determine broad areas of quality and applicability measures, which help us understand humanitarian logistics focused real-time data and their indication for modeling. We then introduce attributes of data and data sources to explain each measure in detail and discuss the implications of the presented attributes and measures for disaster response logistics models. Finally, we address how to measure the attributes of the proposed framework.

4.3.1 Measures of Data Quality and Applicability

In order to develop the framework, we first review the data standards and literature from the disciplines described earlier in the disciplines using data, as well as survey our practitioner contacts for their input. Based on our previous modeling experiences and observations from the data and data sources over time, we identify four fundamental questions about the data and data sources, which lead us to the measures in the framework. These questions are: what data is available (relevance), when data is available (timeliness), where data is available and to what degree data represents the surrounding environment (generalizability), and the degree to which the data reflects the true environment (accuracy). We believe these four categories emphasize the critical characteristics that describe the quality of data and data sources and their applicability to different disasters for logistical modeling. As a result, we propose the following four measures: 1) relevance, 2) timeliness, 3) generalizability, and 4) accuracy. Similar to measures, we survey the literature and practitioners to identify attributes that sufficiently represent each measure and their implications for logistical modeling. The attributes are described in this section in detail. Figure 1 outlines our developed framework for analyzing real-time logistics data for the purpose of use in mathematical modeling. Following a disaster, as the data starts to become available, the data user evaluates it based on his or her purpose (e.g., research, situational awareness and decision-making on the ground). As a result, the importance of overall quality measures and the level of their applicability depend on the purpose of the data user; thus we span these measures around the user's purpose. The data format is highly correlated with the purpose of the research team and refers to the format of the files that data or the information relevant to our humanitarian logistics modeling focus is represented. Day, Junglas, and Silva, (2009) [124] and Altay and Labonte (2014) [65] stress inconsistent information and data formats as one of the information flow impediments that impact decision-making and coordination in the humanitarian response. The data is available in many formats from portable data formats (PDF) to keyhole markup language (KML) and Microsoft Word Documents (.doc).

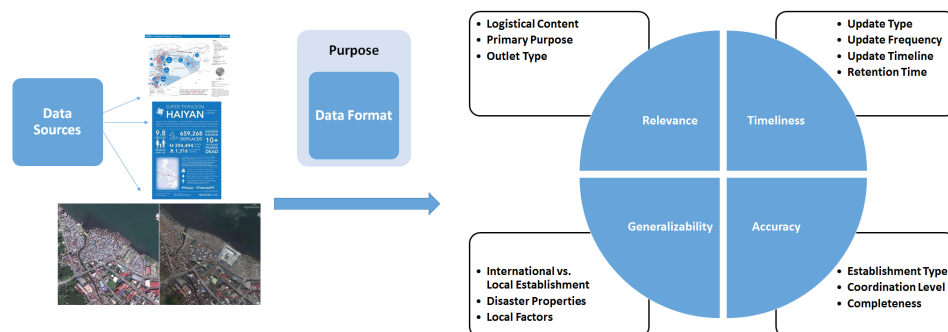


Figure 4.1: Data Analysis Framework: Applicability Measures and Related Attributes

The data format plays a critical role in the accessibility of data and smooth integration of data into models. Many of the available files identified in our case study (see section Framework Application: Typhoon Haiyan Case Study) depict humanitarian logistics information, such as roads and hospitals, yet frequently only in static formats. Optimal data formats for integration with logistics models are editable documents and dataset file types. For example, PDF maps often contain a great amount of information about the severity of damage in an area; however, due to their format limitations, it is hard to access information about road structure. Many released reports do not contain easily or quickly transferable numerical data. KML files and shape (SHP) files might contain the relevant data for testing models, however not in an immediately accessible manner, often requiring file conversions and format manipulations. On the other hand, even after these files are converted, discrepancies may exist between the content. Two different sources converted through the same online SHP to comma separated value (CSV) formatter [125] can return different marking systems, e.g., XY coordinates vs. “osm_id,” as location indicators. This variation may be problematic for analysis, since it might require multiple infrastructure information sources. On the other hand, some sources contain Microsoft Excel data spreadsheets, and classification in this data simplifies the identification process of sources with promising data for the purpose of logistics models.

Next, we describe the measures and attributes used in our framework based on their utility for logistical modeling and potential challenges. We should note that, while some attributes might refer to multiple measures (e.g., the classification category primary purpose might infer about the relevance and generalizability), some attributes can be related (e.g., local factors and disaster properties).

Relevance

Relevance is determined by whether the data meets the needs of its users. The relevance of the data obtained from the data streams refer to the degree to which the data meets the current and future needs of the data required for logistical modeling. For our framework, we identify the following measures that represent relevance.

Logistical Content

Humanitarian and relief organizations constantly collect, process and disseminate immense amounts of information in a broad range of settings and applications. Our work and this specific paper focuses on humanitarian logistic operations in post-disaster settings, such as on-the-ground operations immediately following natural or man-made disasters (e.g., immediate medical assistance or search and rescue operations). Thus, our first step consists of identifying specific types of information relevant to humanitarian logistics. The data is divided into the following categories: demand, supply and infrastructure. These categories are similar to those used in the literature [66, 126].

- Demand: Effective and efficient relief efforts require the identification of the location, quantity and types of supplies needed within the affected region. Demand in these settings can correspond to needed physical goods, such as food, medication or shelter, as well as needed services, such as medical assistance, rescue, and telecommunication.
- Supply: Information about relief supplies that are pre-existing or gathered after a disaster, transportation vehicles and expert or key personnel (e.g., see [127] for examples specific to Typhoon Haiyan) is another important component for the relief efforts.
- Infrastructure: In order to facilitate distribution of supplies to demand, we need to have knowledge of the infrastructure (e.g., roads, airports, seaports and their post-disaster conditions). While

this critical component can be highly uncertain, there is great potential benefit from accurate and timely data from the field.

Primary Purpose

Primary purpose indicates the focus of the information posted in the outlet and/or role of the organization providing the information. Some of the examples are initial assessment, evacuation and providing maps. The primary purpose of the data outlet or organization might help researchers search for relevant information for their models. For example, if a researcher is working on search and rescue operations, it might be easier for him or her to start from an outlet that is focused on initial damage assessments. More specifically, in the case of Typhoon Haiyan, a researcher can begin his or her analysis from CEMS, NDRRMC and OpenStreetMap (see Table 2 for the primary purpose of surveyed information outlets in this study).

Outlet Type

Outlet type identifies the ownership of information, as observed in our study. We distinguish between two types of information outlets: original information outlets and aggregator information outlets. Original information outlets refer to organizational websites that primarily provide information and data, either collected by that organization or transform the data for their specific response purposes. We use the term aggregator for outlets that predominantly disseminate information collected by various other sources.

Outlet type classification can be helpful for practitioners and researchers for the following reasons. Original information outlets might be a good choice to look at when a researcher or practitioner is searching for a certain type and/or format of information, such as maps or reports about damaged areas. In this case, researchers or practitioners might need to search for several original information outlets to find particular data or accumulate a series of different information pieces. Aggregator information outlets generally compile information about supply, demand and/or infrastructure from assorted outlets. Thus, a researcher might want to start his or her initial search from these sources.

Timeliness

The timeliness of data refers to time dimensions of the collected, analyzed and disseminated data. Timeliness is deemed by many users as generally the most important characteristic of data [128]. Humanitarian data should be collected, analyzed and disseminated efficiently, and must be kept up to date [129]. Timeliness is also defined as the delay between when the data is collected and when data becomes available and accessible. Timeliness is a crucial measure for modelers as it highly impacts the types of the models they can feasibly implement.

Data Update Type (new update/incremental vs. overwrite)

Data update type represents the method by which the status updates are provided after the initial file upload. Incremental updates suggest that the new information is described in a new file or new field. On the other hand, overwrite updates indicate that the additional information is appended to the existing file containing the original information.

Practitioners and researchers may have different needs with respect to data updates. While practitioners may seek the most current data with cumulative statistics for operational decision-making during a certain disaster, logistics modelers often prefer to see the evolution of the post-disaster data. As a result, while an overwrite update may be preferable for practitioners, it is not desirable for humanitarian

logistics researchers who focus on adaptive modeling. Modelers generally find the piecewise information about additional available roads or estimated needs (e.g., refugee camp population requiring food assistance from WFP) at a location more useful. Update type also provides important implications for the prioritization of the data collection process for research purposes. Sources such as OpenStreetMap may have openly available data that archives the prospective near real-time updates of roads, which can be retrieved later on, while other map sources that have overwrite updates should be monitored frequently to capture the evolution of data over time.

Update Frequency

Update frequency represents the frequency with which data is updated. Update frequency can be by minute, hour, day or other. Due to the nature of humanitarian operations and impact of time in the output, the humanitarian community benefits from data being updated timely and frequently.

Update frequency may have a large influence on modeling decisions. In our case study (see Table 2), many organizations that share data sources appear to update and upload their datasets daily and frequently post new content data related to logistics. This may influence the model type and the inclusion of the dynamic information into the models. For example, as the frequency of the information increases, a researcher might prefer dynamic programming to multi-stage optimization for modeling.

Data Update Timeline and Retention Time

While both data update timeline and retention time are associated with the time perspective of data, update timeline refers to the timeline from the initial time data starts to be uploaded until the last time data is uploaded. On the other hand, retention time captures the last time when the data will be available for public use. In other words, update timeline describes when files are updated on the website, e.g., from November 1, 2014 to December 15, 2014. Retention time is for how long that data will stay up on the website, e.g., the data can be accessed for another year after it has been uploaded. Retention time is also often associated with the establishment type of the information outlet, which is discussed later.

As stated before, timeliness is generally one of the most important measures of data and it has several implications for modelers. Different update timelines might express different values to various types of modeling purposes. The first available post-disaster data is crucial, especially when one wants to find as much information as possible to understand the immediate context and link information with high-priority humanitarian logistics activities within a certain period after the disaster. For example, the data available during the golden time (first few days after the disaster) is vital for models focused on search and rescue operations [179]. As the initial few days pass after the disaster, the information about the supplies (which team is where with how much medical or other supplies) becomes widely available. This information provides a basis for the relief-distribution modeling. Additionally, longer timelines imply more data for the modelers and this allows them to conduct more comprehensive analysis for test case generations of past disasters.

Similar to update type, retention time might also impact and aid in enhancing the data collection process. Longer retention times enable researchers to access time-sensitive critical information. Postponing the collection of data that remains well after its posting date may allow for greater time spent on more volatile sources.

Generalizability

The generalizability measure in this framework indicates the applicability of the information obtained from the data resources of a particular disaster to other disasters for preparedness, analysis,

lessons learned and evaluation. We determine an establishment's local versus global designation, disaster properties and local factors as key indicators of generalizability. Generalizability facilitates modelers to learn from previous disasters to improve the preparedness and response to future disasters.

Local/ National vs. Global/International

This classification addresses whether the information source is administered by an international organization or a local government/organization where the disaster occurs. It signifies the level of involvement from local government in the disaster response operations.

Local/national versus international data ownership might inform about the generalizability of this data and analysis for future disasters. For example, the level of involvement from the local government for the post-disaster response can be included in the discussion of different disaster comparisons. Similarly, depending on the disaster type, pre-disaster evacuation efforts of the local/national government can be a critical factor when comparing different disasters and making inference from them.

Disaster Properties

As the name suggests, this attribute describes the main characteristics of a disaster. Tatham et al. (2013) [130] develop a 13-parameter framework that captures the factors impacting logistical preparedness and response. A significant part of the classification categories from Tatham L'Hermitte, Spens, and Kovács's framework, such as the time available for action (disaster onset), disaster size, magnitude of impact, duration of time and environmental factors (such as the topography or weather conditions of the affected area) are related to disaster characteristics used in our framework.

Disaster characteristics can educate modelers about decision-making in different stages of the disaster cycle. For example, a sudden-onset disaster with predictable timing, such as Typhoon Haiyan, can help modelers and practitioners on the ground to improve the prepositioning strategy to save as many lives as possible. In addition, disasters with predictable timing impact the level of information available in the response phase by advance notice, which can impact the specific characteristics of the model and model validation.

Local Factors

The World Bank Logistics Performance Index (LPI) measures the "friendliness" of a country based on six factors: customs, infrastructure, services quality, timeliness, international shipments and tracking/ tracing [244]. Tatham et al. (2013) [130] suggest the Logistics Performance Index as one of the four factors impacting the logistical preparation and response. While LPI focuses on factors that impact logistical performance directly such as infrastructure, local factors refer to metrics for local environment. L'Hermitte, Bowles, and Tatham (2013) [132] present a classification model of disasters from a humanitarian logistics perspective. Their model composes the time and space components of the disasters and five external situational factors of the disaster environment. The external factors stated in the paper are the government situational factors, the socioeconomic situational factors, the infrastructure situational factors, the environmental situational factors, and the conflict environment. Five external situational factors of the disaster environment of L'Hermitte et al.'s work (2013) [132] and some of the categories from Tatham et al.'s 13-parameter framework [130] such as the geographic context, population density, per capital GDP, and potential for the reoccurrence of the disaster in the same area are examples of local factors. Another local factor is the language of the local environment, as in the case of the 2010 Haiti earthquake [75].

According to existing studies (e.g., [245]), number of affected people generally decreases. Logistics

modelers can benefit from this information when estimating service demand for a given region.

Using a country's LPI value, modelers can evaluate applicability of the available data from one country to another. Furthermore, higher LPI scores usually correspond to less restriction from the government [134] on relief response operations, corresponding to another generalizability measure of the data.

Geographical context of the local area where a disaster takes place might also provide multiple insights for modeling. For example, the Philippines being a combination of islands implies routing decisions using different modes of transportation, as well as the coordination of relief items among islands. Furthermore, whether or not a given island is the hub for relief operations can also impact the routing decisions. Another local factor, disaster reoccurrence probability, can provide useful information to logistics modelers at various stages of the disaster management cycle. For example, Ergun, Stamm, Keskincok, and Swann (2010) [135] describe numerous efforts of Waffle House Restaurants to effectively respond to hurricanes in southeast US. These efforts include equipment prepositioning, special menus and advanced personnel scheduling. Similar strategies can improve disaster management in areas that are prone to storms such as the Philippines.

Accuracy

Multiple researchers discuss data accuracy directly or using related terms, such as reliability, verifiability and accountability [129, 120, 122, 127, 136, 124, 65, 117]. According to the Humanitarian Data Exchange Quality Assurance Framework, the accuracy of the data is defined as “the degree to which the information correctly describes the phenomenon it was designed to measure” Synthesizing the definitions from these resources, we define accuracy as the measure that represents the reliability and the credibility of the information obtained from the data streams and data sources.

DeLone and McLean (2003) [117] emphasize accuracy in their information success model. For example, accuracy-related terms appear in systems quality as “system reliability”, in information quality as “accuracy”, “conciseness” and “completeness”, and in service quality as “accuracy” and “reliability”. Day et al. (2009) [124] state that self-reported information like shelter registration is generally unreliable.

Altay and Labonte (2014) [65] assess unreliability as an information impediment for decision-making and provide several sources of unreliable information examples from the case of the earthquake in Haiti in 2010, such as crowdsourced data. They also note that, rather than waiting for the perfect information, practitioners should utilize the available information and make sure coordination is stressed in the information-sharing. From a modeler's perspective, accuracy implies various assumptions about the available information, as well as the unknown data.

Establishment Type

Establishment type denotes whether the placement of the data (e.g., website or repository) was established specifically and solely for the purpose of a specific disaster or maintains data across multiple disasters. We use establishment type to separate data sources into two categories: multiple disaster and disaster specific. Multiple disaster sources correspond to organizations and websites involved with data on disasters prior to a given disaster, often retaining data from such disasters. Such sources or repositories usually retain data for multiple disasters after the relief operations are completed, e.g., ReliefWeb. Disaster specific sources generally provide information during relief operations of a specific disaster and may become unavailable shortly afterwards.

Establishment type can often inform researchers and practitioners about reliability, completeness and accuracy of information since it often relates to the structure of the organization that maintains the information outlet. Multiple disaster data sources may possess additional verification processes, which

may increase their reliability. However, these data sources (e.g., PDF maps) rarely contain the raw pre-processed data. In contrast, based upon our specific case study experience, we observe that disaster specific sources may directly post datasets without progressing through a verification process, or may not fully share the verification process with the public. They may also lack the level of reliability and trust in comparison to sources established for previous disasters. Regardless of the situation, connecting with practitioners at appropriate times might assist in further understanding how and whether data is verified. It may also open up opportunities to explore datasets in pre-filtered formats, reliable short-term sources, or data sources more relevant to practitioners. However, datasets may not necessarily be complete or entirely accurate, requiring data fusion or synthesis of various datasets together to achieve logistics modeling requirements. The resultant fused datasets will need to be reassessed for applicability for effective on-the-ground operations.

Coordination Level

This attribute represents the level of coordination among different communities during a disaster response. It might indicate various types of collaboration, such as coordination between different relief agencies and coordination between local and international governing bodies. The concept of coordination through subgroups introduced by Jahre and Jensen (2010) [137] aims to organize humanitarian help in a number of different areas by predefined management. Jahre and Jensen (2010) [137] discuss the importance of the balance between horizontal and vertical coordination in any period of disaster management. Additionally, the authors mention the role of logistics cluster, one of the 11 OCHA formed clusters, on information management and the challenges of coordinating the information.

One key component to a successful coordination is the exchange of information, which in return results in additional information generation (e.g., cluster reports to be shared with participants). In addition, the involvement of various parties in the combined mission improves data accuracy as information is validated by the distinct participants. Moreover, similar to the establishment's international versus local designation, the level of cooperation of the local government with international communities in the disaster response efforts impacts the evolution of available data.

Completeness

Completeness represents whether there is missing information or not. Examples of incomplete information might be in regards to the status of certain roads or damage level of buildings.

Missing data raises questions regarding the accuracy of information to be used in the logistical models and forces modelers to account for the incomplete information. For example, when road status information is not provided, the modeler needs to make assumptions about that information. In order to account for the inaccurate and incomplete information, uncertainty factors should be included in the models.

Assessment of Framework Attributes

Table 1 provides an example of a metric that might be used to assess and categorize each framework attribute developed above, as well as sample metric categories for illustration purposes. Some of these attributes and the metrics have been previously introduced in existing literature, in which case we include the appropriate references in the last column of Table 1. As discussed by Bharosa et al. (2009) [119], some parameters, such as relevance, accuracy, and timeliness that signify information quality, are often hard to measure during disaster management. Thus, what we provide in Table 1 is just one set of examples to measure the attributes, and there might be many other ways to measure each of them. For

example, in addition to the sample metric provided in Table 1, it is possible to measure the coordination level using the number of agencies sharing resources or presence and the role of a single agency during a disaster such as local government.

Framework Measure	Attribute	Metric Example	Sample Metric Categories	Reference
Relevance	Logistical Content	number of logistical content categories that the data source cover	a) infrastructure b) demand c) supply d) all	Ergun et al. (2010a)
	Primary Purpose	main focus as described in mission statement	a) (initial) post disaster assessment, damaged infrastructure b) relief aid c) coordination d) all	
	Outlet Type	percentage of the content that is originally prepared by the data source	a) < 20% (aggregator) b) 20-50% (mixed) c) >50% (original)	
Timeliness	Update Type	number of previous versions of the file stored	a) 1 (overwrite) b) >1 (incremental)	
	Update Frequency	average time between two consecutive updates	a) hourly b) daily c) weekly d) monthly e) one time	OCHA (2002)
	Update Timeline	average time between the first and last update	a) > 1 year b) > 3 months c) >1 month d) > 1 week	
	Retention Time	average time the information is kept by the data source	a) > 5 years b) > 3 years c) > 1 year d) < 1 year	
Generalizability	International vs. Local Establishment	percentage of funding available from local government	a) < 50 (local) b) > 50 (international)	
	Disaster Properties	number of people impacted	a) > 1,000,000 b) > 10,000 c) > 1000 d) > 100	Tatham et al. (2013)
	Local Factors	LPI score	a) > 4 b) >3.5 c) >3 d) <3	Tatham et al. (2013)
Accuracy	Establishment Type	total number of disasters involved	a) 1 (disaster-specific) b) >1 (multiple - disaster)	
	Coordination Level	average number of agencies listed in the reports	a) 1 b) >5 c) >10 d) > 20	Jahre and Jensen (2010)
	Completeness	percentage of the attributes tagged or have information about	a) > 20% b) > 10% c) > 5% d) < 5%	DeLone and McLean (2003)

Table 4.1: Assessment of Framework Attributes

4.4 Framework Application: Typhoon Haiyan Case Study

To illustrate the application of the framework presented above, we focus on a specific disaster response, Typhoon Haiyan, as the case study for investigating the role, value and limitations of integrating new information streams for logistical modeling during the aftermath of a disaster. On November 8, 2013, Typhoon Haiyan, named as Typhoon Yolanda, the strongest storm recorded at landfall [85]. and one of the strongest tropical cyclones in recorded history [87], hit the Philippines and resulted in catastrophic damage throughout the country. As of April 7, 2014, 6,300 individuals were reported dead, 28,689 injured, and 1,061 are still missing according to the National Disaster Risk Reduction and Management Council (NDRRMC) [81]. Both the Philippines government and international humanitarian organizations started their preparedness activities as early as November 7, 2013 and began response activities immediately following the fall.

Typhoon Haiyan represents an evolution of disaster response in which the emergence and growth of data during relief operations brought new opportunities for addressing humanitarian challenges. Multiple factors played a role in the generation of this outstanding level of information, including local factors, the nature of the disaster and efforts of OCHA and the availability of data from the Philippines government responding entities. More specifically, the Philippines benefited from advanced information communications technology and widespread media and organizational coverage within the country in the disaster response efforts.

As this research is carried out by an English-speaking team, the synergy between information and data predominantly shared in the English language provides us an opportunity to pursue this case study. In addition, the Typhoon Haiyan post-disaster environment was permissive with respect to information sharing across stakeholders; unlike conflict-driven, complex humanitarian crises where repressive environments often restrict information-sharing, especially through public online sources. Typhoon Haiyan is a sudden-onset disaster with relatively predictable timing and location, which made it possible for the advance staging of volunteers. OCHA's call for digital volunteer support through the Digital Humanitarian Network prior to the typhoon activated volunteers around the world to participate in collecting and processing information [138]. In addition, the growth of digital humanitarians or "crisis mappers" has expanded nontraditional data streams during recent crises (Haiti earthquake, Pakistan, Chile, Christchurch, Bhopa, super-storm Sandy, etc.), often in online formats and frequently available to the public.

In this section, we first present the timeline of events following Typhoon Haiyan for the case study. Next, we examine the data sources and provide the classification of the data and data sources based on the framework developed in Section 3. We provide a brief logistical content analysis to assess the available information for logistical models. Finally, we present lessons learned from this case study.

4.4.1 Information and Data Retrieval

This section highlights the chain of events describing our information retrieval process. Figure 2 illustrates the information and data retrieval epochs of this study.



Figure 4.2: Information and Data Retrieval Epochs

Initialization

Shortly after November 8, 2013, when Typhoon Haiyan moved into the central Philippines region, one of our team members, who is an active researcher and practitioner in the humanitarian technologies community, began receiving numerous email communications from the crisis mappers network. These emails contained publicly available website links to information sources pertaining to Typhoon Haiyan. This network is also linked to other digital humanitarian groups that were being activated by the Digital Humanitarian Network, including the Standby Task Force, Humanitarian OpenStreetMap Team, GIS-Corps, HumanityRoad, Info4Disaster, MapAction, Translators without Borders, Statistics without Bor-

ders and other members of the Digital Humanitarian Network [139].

Collection and sorting

The team recruited an undergraduate research assistant to collect data beginning on November 12, 2013, starting with the resources identified in initial emails. As a first step, we iteratively devised a sorting scheme for data intake. After teasing out the emails relevant to the typhoon, our initial approach in the first days, between November 12, 2013 and November 19, 2013 was to continuously download from the data sources identified on those websites, such as Logistics Cluster and Copernicus Emergency Management Service (CEMS), with the goal of capturing the temporal aspect of the available data. The potential relevance for each data source was assessed using the accompanying description of the data. However, since time constraints and our advancing knowledge of humanitarian information sources did not allow for a complete understanding of each dataset at the time of retrieval, we continued to download from as many of these sources as we could find, in the hopes that each of those sources might contain desired information.

Prioritization

After one week of sorting data, around November 19, 2013, we realized that the volume, both in number and size, of data was overwhelming. For example, the Humanitarian Response had over a hundred new files [140] and OpenStreetMap had several gigabits of data (OSM, “Index of Haiyan”). The magnitude of the data created a need to balance the trade-off between processing the known data sources and searching for new sources. After the initial few days of collecting a breadth of data, we attempted to focus on understanding the content of the resources, especially with respect to their logistical content. For almost a week, we focused on differentiating between logistical content, storing and archiving current data according to its content type.

Searching for new sources

As time passed, around November 25, 2013, we continued to search for new data sources. The sources from the email lists served as a starting point, and allowed for retrieval from some important sources earlier than otherwise possible. However, finding new sources required other methods, such as subscriptions to appropriate newsletters and mailing lists, and manual internet searches. Yet, even with these methods as aids to supplement the manual search for sources, we were not able to identify all additional sources, particularly those not openly shared on the internet. This illustrates the challenge to discover the relevant sources far enough to reduce lost data and the limits of remote research activities that explore field operations. Within this particular case study, many relevant sources, e.g., the OSM repository specific to Typhoon Haiyan, were further researched by our team over one month later, despite HOT and the American Red Cross activities commencing very early in the response, around November 8, 2013.

Reorganization

With the addition of new resources, the team acquired significant information and continued content differentiation, storing and archiving. On December 4, 2013, the team noticed the duplication of resources and decided to reorganize the list of them. At this point, the team discontinued downloading from repeating outlets.

Analysis

The team started the analysis of downloaded data on November 21, 2013, which consisted of the

initial assessment of the data content, especially the logistical content discussed in measures of data quality and applicability section. While the team worked on analysis and data collection in parallel after this point, the detailed analysis of data sources classifications began on December 5, 2013.

4.4.2 Data and Data Sources Classification

In this section, we implement the framework proposed in the study to the data obtained from Typhoon Haiyan. We begin with the detailed description of the studied data outlets (see Appendix for details). The information posted at the outlets, generally in their “about us” and data pages are analyzed based on their relevance to humanitarian logistics models during the timeline of the study, between November 10, 2013 and March 31, 2014. In developing measures of data quality and applicability we identify two types of information outlets: information aggregation sites and organization or sector specific data repositories (original, in other words, primary outlets). However, this dichotomy may not necessarily be entirely distinct, as some organization-focused sources will also include some data collected from other groups; for example, MapAction takes advantage of data from NDRRMC to generate some of the maps. Thus, one might expect to categorize MapAction as an aggregator. Also, there might be more information outlets providing information after Typhoon Haiyan. This classification includes analysis of a subset of available resources. Table 2 comprises a subset of the proposed classification categories such as outlet type, primary purpose and logistical content as they apply to the studied data sources.

4.4.3 Logistical Content Data Availability Analysis

In this section, we examine the availability of logistical content information following the days after Typhoon Haiyan. As discussed in disciplines using data section there are differences in the terminology used by researchers and practitioners. In order to find the relevant information (e.g., demand and infrastructure) to integrate into models, academic researchers should first learn about these differences. For example, in the case of Typhoon Haiyan, while searching for relevant information, we establish a set of key search terms that are used to identify the potential data sources within our compiled database that contain information about service demand and infrastructure. We use sources from different disciplines described in Section to help us build the list of these search terms. Then, using this list, we examine the data gathered from sources such as OpenStreetMap, Logistics Cluster, NDRRMC, and Copernicus Emergency Management Service (CEMS) for availability of the logistical content (see Table 2 for logistical contents and the Appendix for description of these data sources). The filtered data is then analyzed to identify the specific content they contain in relation to the kind of information needed for logistical models. For example, we want to know the level of damage each road link sustained, its post-disaster status and location of potential beneficiaries of medical and rescue services.

In the analysis of data availability for road damage, we look specifically at the data for Tacloban City, which is expected to have more data in comparison to more rural areas [139]. The number of roads labeled with some level of damage and the total number of roads recorded are compiled for each day between November 7, 2013 and November 28, 2013 using OSM and CEMS. However, even by November 28, 2013, only approximately 5 links among these data sources are labeled to contain at least some level of damage. This result demonstrates how limited logistics data is for modeling following a disaster.

As an example, Table 3 shows the details about the information provided by NDRRMC regarding the road and port status between November 7, 2013 and November 13, 2013 in its situation reports. Regarding road status, the reports include information about the name of the road, status (not passable or passable) and additional comments (why it is not passable or status of ongoing clearance operations).

Table 4.2: Classification of Information Data Files Collected by Our Team (see Section 4.3.1 for details on classification categories and Conclusion for description of sources.)

Source Name	Outlet Type	Primary Purpose Posted Related to HL Models	Logistical Content	Main Downloaded File Type(s)	Update Frequency Observed in Timeline	Update Type Observed	Establishment	Update Timeline	International / Local
APAN	aggregator	user-uploaded files	all	mixed	varies	new	disaster-specific	11/12/13 <i>present</i> *	- international
Copernicus EMS	primary	initial post disaster assessment	infrastructure	SHP	daily	new	multiple disasters	11/12/13 11/18/13	- international
DSWD	primary	evacuation	supply, demand	PDF	every few days	new	multiple disasters	11/09/13 12/12/13	- local
Google Crisis Maps	aggregator	facility placement and damage assessment	infrastructure	KML	unknown	overwrite	multiple disasters	11/08/13 11/26/13	- international
Humanitarian Response	aggregator	satellite imagery	all	mixed	daily	new	multiple disasters	11/12/13 <i>present</i> *	- international
IFRC	aggregator	displacement	all	PDF	daily	new	disaster-specific	11/10/13 <i>present</i> *	- international
Logistics Cluster	primary	coordination and resource aggregation	infrastructure	PDF	almost daily	new	multiple disasters	11/12/13 <i>present</i> *	- international
LogIK	primary	relief aid	supply	XLSX	daily	overwrite	multiple disasters	unknown	international
MapAction	primary	disaster response maps	all	PDF	daily	new	multiple disasters	11/12/13 <i>present</i> *	- international
NDRMC	primary	facility Operations and Damage Assessment	All	PDF	daily	new	multiple disasters	11/08/13 <i>present</i> *	- local
OSM (repository)	primary	user-mapped damage, damaged infrastructure	infrastructure	SHP, OSM	daily	overwrite	multiple disasters	11/08/13 <i>present</i> *	- international
Reliefweb	aggregator	aggregates various sources	all	mixed	daily	new	multiple disasters	11/08/13 <i>present</i> *	- international
UNITAR - OSAT	primary	initial post disaster assessment	infrastructure	PDF	daily	new	multiple disasters	11/12/13 11/20/13	- international
VISOW	primary	user-mapped damage	infrastructure	CSV	daily	new	disaster-specific	11/09/13 <i>present</i> *	- international

The reports before November 8, 2013 12:00 pm do not include road status information. The information provided in the reports is cumulative; in other words, a road that was previously identified as impassable is kept in the following reports until November 13, 2013. As of November 13, 2013 6:00 pm all roads that were previously affected are stated to be passable and road status information is not included in the following situational reports. In the entire Philippines area, the reports include at most only 18 roads. The airport information consists of the name of the suspended airports as well as cancelled flight information. Data related to sea ports include the name of the port and number of strandeers by type (passengers, vessels, rolling cargoes and motor banca boats).

Update Date	Update Time	Number of Roads with Status Information	Number of Roads Reported Not Passable	Number of Suspended Airports	Number of Sea Ports with Strandeers
11/8/2013	12:00 PM	5	5	5	12
11/8/2013	6:00 PM	5	5	13	15
11/9/2013	6:00 AM	5	5	13	20
11/9/2013	6:00 PM	13	13	4	4
11/10/2013	6:00 AM	18	18	4	3
11/10/2013	7:00 PM	18	3	4	0
11/11/2013	6:00 AM	18	3	4	0
11/12/2013	10:00 AM	18	3	4	0
11/12/2013	10:00 PM	18	3	4	0
11/13/2013	7:00 AM	18	2	0	0
11/13/2013	10:00 PM	18	0	0	0

Table 4.3: Road and Port Status Information Update after Typhoon Haiyan between 11/8/2013 and 11/13/2013 in NDRMMC Situational Reports.

In order to obtain service demand information for logistical models, we also examine the available building data, which might provide information about the individuals either trapped under collapsed structures or displaced persons due to loss of property. The building damage reports vary by different data sources. Table 4 contains a sample of information from the NDRRMC reports on the number of damaged buildings (totally or partially damaged), deaths, injured, and missing individuals, affected families and persons, and stranded individuals between November 8, 2013 and November 13, 2013. These numbers are aggregated for the entire Philippines.

For more detailed building data analysis, we next focus on five cities: Tacloban City, Guiuan, Palo, Ormoc and Cebu. Notably within three of these cities (Tacloban City, Guiuan and Palo), the percentage of buildings with “collapse” or “damage” indicators range between 40% and 60% by November 20, 2014. This analysis is conducted by using OpenStreetMap and CEMS. Table 5 shows the total number of buildings and number of damaged or collapsed buildings for each city between November 8, 2013 and November 20, 2013. The information about building damage appears to be delayed in these sources in comparison to NDRRMC reports. Even in Tacloban City, the damage information starts on November 14, 2013. This is another example of limited information immediately after disaster, especially for search and rescue purposes.

We next examine the available data for supply information. For illustration purposes, we examine two pieces of supply information. Table 6 depicts the number of vehicles available each day throughout November, beginning on November 11, 2013 from UN OCHA's Logistics Information About In-Kind Relief Aid (LogIK) records. The table includes the number of vehicles decided by certain dates and their status categories (i.e., dispatched, committed, delivered) as of December 8, 2013.

Update Date	Number of Damaged Houses	Number of Total Damage	Number of Partial Damage	Number of Deaths	Number of Injured	Number of Missing	Number of Affected/ Pre-emptively Evacuated Families	Number of Affected/ Pre-emptively Evacuated Persons
11/8/2013							26675	125604
11/8/2013				3	7		151910	748572
11/9/2013				4	7	4	161973	792018
11/9/2013	3438	2055	1383	138	14	4	944597	4282636
11/10/2013	3480	2071	1409	151	23	5	982252	4459468
11/10/2013	19651	13191	6360	229	48	28	2055630	9497847
11/11/2013	23190	13473	9717	255	71	38	2095262	9679059
11/12/2013	41176	21230	19946	1774	2487	82	1387446	6937229
11/12/2013	149015	79726	69289	1798	2582	82	1387446	6937229
11/13/2013	149756	80047	69709	1833	2623	84	1387446	6937229
11/13/2013	188225	95359	92886	2344	3804	79	1730005	8012671

Table 4.4: Demand Information Update after Typhoon Haiyan between 11/8/2013 and 11/13/2013 from NDRMMC Situational Reports.

In addition to vehicle information, we also explore data about the medical teams from WHO Health Cluster Reports. The PDF maps show the total number of foreign medical teams divided by city and origin of the medical team, starting November 15, 2013. Table 7 shows the number of foreign medical teams and their operational status. In this context, operational teams refer to medical teams that are actually seeing patients. Similar to infrastructure information, updates on the medical teams are very scarce.

4.4.4 Lessons Learned: Challenges Faced During the Collection

Process and Potential Solutions

This case study enables us to better understand the current situation of real-time data, how data evolves, and to what extent real-time data is available. We believe that this valuable experience will inform and aid modelers in building improved models. This section describes the challenges faced during our study, and recommendations for how these challenges can be addressed in the future (when possible) from an academic humanitarian logistics perspective.

Time-sensitive information

As information evolves after a disaster, some sources that commonly recur during separate disaster relief efforts do not retain their data for long periods of time. This generally depends on the establishment type and update type of the information outlet. Investigating retention time and update type of information outlets before the onset of a disaster can aid in gathering time-sensitive information. For instance, collection and analysis of data from resources that have the tendency to retain their files longer can be postponed to later times depending on the ultimate goals of the data. If a researcher focuses on modeling the initial few days of the disaster, this approach might not work. However, if a researcher wants to model a later period, postponing collection of retained files can be beneficial, thus avoiding the trade-off between collection of time-sensitive information and the prioritization of analysis. Additionally, the format of the time-sensitive information provided by an outlet tends to be similar for each disaster. Familiarity with the data format can ease the data collection process. Moreover, some out-

Date	Tacloban City		Palo		Guiuan		Ormoc		Cebu	
	Total Number of Buildings	Total Number of Collapse / Damage	Total Number of Buildings	Total Number of Collapse / Damage	Total Number of Buildings	Total Number of Collapse / Damage	Total Number of Buildings	Total Number of Collapse / Damage	Total Number of Buildings	Total Number of Collapse / Damage
11/8/2013	2967	0	65	0	0	0	327	0	327	0
11/9/2013	15456	0	469	0	2	0	327	0	327	0
11/10/2013	24468	0	1649	0	28	0	327	0	327	0
11/11/2013	30278	0	2942	0	1585	0	327	0	327	0
11/12/2013	35786	0	3579	0	2262	0	328	0	327	0
11/13/2013	38722	0	3746	0	2277	0	328	0	327	0
11/14/2013	55374	14808	5432	1604	2294	0	329	0	2087	0
11/15/2013	65579	23242	8422	3934	2308	0	337	0	5999	0
11/16/2013	71338	28409	10465	5975	2308	0	337	0	6111	0
11/17/2013	75781	32738	11009	6519	2309	0	350	0	6111	0
11/18/2013	76045	32937	11010	6520	2309	0	350	0	6111	0
11/19/2013	76132	32992	11010	6520	4248	1890	350	0	6112	0
11/20/2013	76133	32993	11010	6520	4343	1963	350	0	6112	0

Table 4.5: Demand Information Update after Typhoon Haiyan between 11/8/2013 and 11/20/2013 in OpenStreetMap and CEMS.

lets, such as OSM, benefit from specifically searching for a separate repository that might be linked to their wiki webpage. Accessing those repositories in the early days of the disaster response supports the smooth collection of time-sensitive data.

Delays

The data on collapsed and damaged buildings did not begin to appear in Palo and Tacloban City until November 14, 2013, and in Guiuan until November 19, 2013. Depending on the main objectives of the humanitarian agencies using the data, those dates might be too late. For example, from the perspective of search and rescue operations, receiving information three days after a disaster may seem too late to assist most of the people trapped under buildings. However, people who were stranded but not directly impacted by the collapsed structures would likely still be alive during that particular time frame and could benefit from relief supplies.

Data explosion and information overload

With the emergence of technology and increasing number of humanitarian organizations, the amount of information available after a disaster is accelerating. From the perspective of a humanitarian logistics research team, it is challenging to account for this information load in a timely manner. One major factor is a limited research workforce. The limited human resource capacity for information processing is a shared challenge across the humanitarian ecosystem and acutely experienced by field practitioners. The evolution and success of many digital humanitarian efforts is the harnessing of remote workforces, often through crowdsourcing and microtasking efforts. Based on this case study and our team's experience, a possible solution for this issue is to recruit additional help in the data collection process whenever possible. Exploring collaborative opportunities with research teams and potentially practitioners to build a feasible and appropriate workforce to identify, filter, assign and prioritize humanitarian logistics datasets for modeling purposes may be a long-term goal. In the short term, future efforts might consist of two groups working in tandem on the data retrieval: one searching for new sources and initializing their retrieval, while the other investigates whether to continue retrieving data from the identified sources or not.

Date	Number of Vehicles Dispatched	Number of Vehicles Committed	Number of Vehicles Delivered by	Total Number of Vehicles
11/11/2013	1	4	13	18
11/12/2013	2	6	18	26
11/13/2013	3	9	33	45
11/14/2013	4	10	48	62
11/15/2013	4	11	54	69
11/16/2013	4	13	55	72
11/17/2013	4	15	60	79
11/18/2013	4	15	63	82
11/19/2013	5	15	68	88
11/20/2013	5	15	80	100
11/21/2013	5	16	83	104
11/22/2013	5	20	85	110
11/23/2013	5	20	85	110
11/24/2013	5	20	86	111
11/25/2013	5	20	86	111
11/26/2013	5	20	88	113
11/27/2013	5	20	90	115
11/28/2013	5	20	90	115
11/29/2013	5	20	90	115
11/30/2013	5	20	90	115

Table 4.6: Vehicle Information Update after Typhoon Haiyan between 11/8/2013 and 11/13/2013 from LogIK Entries.

Date	Standby in country / without destination (out-of country)	At Destination (not registered)	Operational teams (not registered)	Left Deployment (not registered)
11/15/13	6(0)	12(0)	10(0)	0
11/17/13	1(4)	14(0)	9 (0)	0
11/20/13	5(3)	10(0)	22 (0)	0
11/22/13	1(6)	2(0)	42 (0)	0
11/26/13	2(6)	0(2)	41(10)	1(1)
11/30/13	0(5)	0(2)	39(11)	6(3)

Table 4.7: Medical Team Information Update after Typhoon Haiyan between 11/15/2013 and 11/30/2013 from WHO Health Cluster Reports.

Duplication

Further complications arose when we observed that several data sources were reposting data from other sources; for example, the NDRRMC situation reports were placed on both ReliefWeb.int and the

NDRRMC site, in the case of Typhoon Haiyan. Additionally, many updating files with recent timestamps appeared to be identical to older versions of the files. Recognizing the overlapping data segments between primary information outlets and aggregator information outlets early is a key point for researchers. Furthermore, prior information about the expected update timeline and primary purpose of the information outlets can help resolve these issues. For example, if an information outlet is only focused on initial damage assessment, a researcher might stop downloading from this outlet a few weeks after a disaster to prevent any possible duplication.

Relevance

Some of the data retrieved was not as relevant to humanitarian logistics models as originally hoped. Identifying logistical content of an information outlet and prioritizing these outlets might be helpful in organizing the data collection process in the future. For example, seeking out resources, such as OSM and the Logistics Cluster, earlier in future disasters would be valuable, since the coordinates used in OSM provide information about the damaged structures. Additionally, collaborating with and supporting these organizations before disasters strike would help researchers to understand data relevance more clearly.

Compatibility

Even assuming file compatibility, problems might exist between perceptions of the sources and how the sources are developed. For example, when one data source is developed by people on the ground, and another source is developed by digital humanitarian mappers, conflicting information would have to have a system for prioritization. Such a system would also require the differentiation of information obtained from mappers versus on field personnel. Furthermore, the particular mapping techniques of various sources may differ. One information source may mark a singular section as damaged or impassable, whereas another source might tag the entire street. This discrepancy can cause major differences in routing decisions and make it challenging to combine multiple resources to build a larger database.

Availability

Our data analysis shows that there is only information for 5% of the roads that indicate some sort of damage to the road structure. This low level of information in the case of a high-impact disaster, with a high level of media coverage and a large amount of data tracking efforts within the first few weeks of the typhoon, shows that the available information is not enough to integrate real-time data into models without putting efforts into adding accuracy measures and finding missing data. On the other hand, 60% of the available building damage information also requires validity checks for source data.

The availability of data across regions change. Some locations receive more attention than others. In particular with the building data, Palo, Tacloban City, and Guiuan appear to receive more mapping than Ormoc or Cebu City. While some of this can be attributed to proximity to the storm at the peak intensity, looking at displacements suggests that more people were affected in Cebu and Iloilo than Guiuan, yet these regions were less mapped [141]. Some of these variations in coverage may be due to the focus and collaboration of mapping activities with the online communities such as OSM, and need to be further explored. We also observe that most of the damage indicated by the data was for roads along the coast. A possible explanation of this phenomena might be the presence of multiple medical teams close to the coast [142], as well as UN On-Site Operations Coordination Centers having predominant locations along the coast [143]. In addition, digital humanitarians assisting with updating maps might more easily distinguish damage to a larger coastal road than to more crowded neighborhood streets.

We recognize that numerous pieces of data from aggregator outlets have citations to their primary information sources. However, availability of the underlying raw data is frequently limited for a number of reasons. In addition, multiple primary information outlets share PDF maps, yet the detailed information about the infrastructure damage is challenging to obtain since the original core datasets from the primary source are infrequently cited or made available. This results in information loss.

Technological Status of Disaster-Affected Regions

The pre- and post-disaster states of the communication system play a significant role in the opportunities, limitations and gaps of the available data. No matter how technologically advanced a particular geographic area may be, gaps in telecommunication coverage in the post-disaster setting are often present. Significant communication problems arose due to the destruction of power and communication lines in the Philippines soon after Typhoon Haiyan [144]. Over a month after the onset, connecting with field teams within specific regions on a daily basis presented significant challenges, as exemplified by “a survey undertaken in the affected community in Guiuan, which reconfirmed the need for clearer and more frequent communication between aid partners and affected communities” [145]. Recognizing the damage sustained by regional communication systems can help researchers understand the information flow and better explain missing data (completeness) for specific geographic regions and time periods. Absence of information flow from an area can also serve as a signal of significant damage and imply increased needs for humanitarian relief (i.e., demand).

4.5 Conclusion

This study introduces a framework for real-time humanitarian logistics data focused on use in mathematical models. We define a set of measures to assess the quality of the data and their applicability to different disasters. Additionally, we provide modeling implications of data based on the proposed framework and discuss how to measure the attributes listed in the framework. We then apply this framework to the data collected from Typhoon Haiyan and present an example of data sources classification based on proposed measures. We also provide an analysis of the data focused on the logistical content to inform modelers about the availability of logistical data, at least in the case of Typhoon Haiyan. The study describes how our humanitarian logistics team approached the emergence of data after the disaster and the challenges faced during the collection process, as well as our observations.

The study shows that, even with accumulating information from different resources, real-time logistical information is very scarce. The analysis demonstrates that only 5% of the infrastructure in Tacloban City has damage information. The number would be much less for other cities that did not receive as much attention as Tacloban City. We encourage researchers to design appropriate models that consider this issue. The framework and its application illustrate what data is available to the team, when data is available, and how data changes after the disaster. It also provides direction about which data sources to search for a particular purpose after a disaster which would be beneficial in future disasters.

The information and observations included in this study are based only on one disaster, Typhoon Haiyan. Future experiences might differ based on multiple factors, such as the disaster type (e.g., complex emergency, man-made disaster), ICT environment, and involvement of organizations and affected populations. The information outlets described and analyzed in this work constitute only a subset of the available resources and focus on those with an English content and online availability. The description of organizations and digital humanitarian groups involved in information management and data sharing is based upon a growing knowledge of our research team and one that is a work in progress. Furthermore,

the classification provided in this paper is only one of the many possible ways, where other researchers might approach the same data differently. The development of parameters to measure the attributes of the framework is in its early stages. More work needs to be done to improve the measurement structure and customize it for specific purposes.

To the best of our knowledge, this is the first study conducted by humanitarian logistics researchers focusing on the real-time data collection process in a post-disaster setting. It also presents a unique team approach that combines the expertise of both humanitarian logistics researchers and a researcher with humanitarian practitioner experience. The data retrieval and aggregation process described in this paper would not have been possible to carry out in a timely fashion without the pre-existing relationships between researchers and humanitarian practitioners. Through comprehensive mathematical models built specifically for the emerging data sources, researchers can identify the most valuable and promising data for the purpose of more efficient humanitarian logistics operations, and ways to integrate this data into a decision-making process. Ideally, validated humanitarian logistic models developed based on near real-time data shared by humanitarian agencies should undergo a series of iterative processes with practitioners to translate logistic models into relevant tools for field logisticians and agencies to assist in their operational activities. The study enlightens researchers about the availability of real-time data and its challenges. Additionally, it provides a ground work for the integration of real-time data into logistical models.

Acknowledgements: This work has been in part funded by the National Science Foundation, Grant CMMI-1265786: “Advancing Dynamic Relief Response: Integration of New Data Streams and Routing Models” and accompanying REU supplement, CMMI-1265786/001.

APPENDIX

A.1 Original Information Outlets

Logistics Cluster

Logistics Cluster, created by OCHA, aids the cooperation of groups of humanitarian organizations. World Food Programme [91] is the lead agency as appointed by the IASC [146]. Common types of datasets are maps, meeting minutes, and situation updates. Logistics Cluster maintains maps that detail infrastructure data such as operations access constraints and general operations. The meeting minutes from the Coordination - Roads Transport - Sea and Rivers Transport Group include information about the road conditions. These almost daily meeting notes start from November 11, 2013. The situation updates also provide information about various types of transportation channel availability including the overland transport. These updates start on November 14, 2013, with limited data about the road conditions. Most of Logistics Cluster’s files found during the case study comprise of portable data formats (PDF).

LogIK

LogIK, or Logistics Information about In-Kind Relief, is a global online database maintained by OCHA [147]. LogIK provides detailed reports in PDF or XLS format of supply operations, with three categories of data: relief items, transports and contributions. The database reflects reported international/regional humanitarian contributions of relief items. Within these categories, LogIK offers information such as supplier information, destination, and quantity [148]. Specifically, the relief items section includes data about donated item types (e.g., blankets, tents and emergency kits), quantity, senders and other. The contributions section provides the decision date and dollar value of the contributions. The transport section contains information about vehicles provided from different organizations by air, road and sea. This information source may be of higher reliability because its source data originates from donors affil-

iated with the United Nations Office (UN). The information in LogIK is updated daily.

HOT

The Humanitarian OpenStreetMap Team (HOT) is a volunteer-based community formed within the larger OpenStreetMap community that has emerged as a pivotal provider and platform for data in humanitarian operations by providing open source data. The data placed into OSM by volunteers continues to increase its scope and accuracy with rising numbers of users verifying information in more locations. OSM furthermore contributes to the large growth in information in post-disaster operations by frequently updating geographic data, sometimes every minute [174]. OpenStreetMap globe data takes several gigabytes, so specific repositories exist for disasters such as Haiyan [150]. While the “history” feature of the OSM helps to see the previous actions, space limitations tend to prompt these repositories to update at longer intervals and not retain many previous updates. The initial OSM updates for Typhoon Haiyan date back to November 7, 2013 since the HOT team was called by OCHA to start mapping the region a day before the typhoon touchdown. The basic street maps of the cities of Port-Au-Prince and Carrefour provided by OpenStreetMap in about 48 hours following the previous crises were claimed to be the best available maps [151].

MapAction

MapAction is a non-governmental organization that produces maps for the humanitarian crisis. From November 13, 2013 to January 17, 2014, they provided maps in JPEG and PDF formats of affected areas, along with information on the populations, road conditions, coordination (cluster activities by location), shelters and others. The main type of MapAction maps seem to be “Who, What, Where” maps that exhibit the locations of organizations. MapAction compiles data from several sources such as OCHA and NDRRMC for different cities, and each map is accompanied by a summary. A lag appears to exist between the report date and the update time; however, MapAction directs its users to Humanitarian Response [152] Philippines portal for primary MapAction maps [153]. Similar to OSM, MapAction was also present in the Philippines before the typhoon hit [154].

Copernicus EMS

CEMS [155], maintained by the European Commission, “monitors and forecasts the state of the environment on land, sea and in the atmosphere, in order to support climate change mitigation and adaptation strategies, the efficient management of emergency situations and the improvement of the security of every citizen.” The website appears to present these maps, which seem to run between November 9, 2013 to November 18, 2013, in numerous formats and resolutions for over a year after the disaster [156]. Moreover, while CEMS is claimed to have published some of the best pre- and post-event analysis images in the first 36 hours of the Haiyan’s landfall, the website appears to make only a few updates publicly once the initial assessment occurs.

ESRI

Environmental Systems Research Institute (ESRI) holds data from the US Government on infrastructural damage. The ESRI Disaster Response Program supports organizations responding to disaster. They provide “software, data coordination, technical support, and other GIS assistance to organizations” [157]. In this case study, similar to CEMS, these files appear to consist of initial damage assessments. They supported the Typhoon Haiyan response by providing an ESRI platform for publicly licensed imagery after the event, and have supported other disasters [158]. ESRI was one of the organizations that responded

to the OCHA call for volunteers as part of the DHN. After Typhoon Haiyan, ESRI collaborated with the digital volunteer mapping groups such as Standby Taskforce and GISCorps to process social media reports and provide interactive maps [159]. The website also provides maps from other groups such as MapAction and OSM. The Haiyan maps start from November 8, 2013 and were last updated on November 25, 2013 (as of March 2014).

UNITAR - UNOSAT

United Nations Institute for Training and Research's (UNITAR) Operational Satellite Applications Programme (UNOSAT) is a satellite program that provides "solutions to relief and development organizations within and outside the UN system to help make a difference in critical areas such as humanitarian relief" [160]. The satellite images appear to allow digital mapping volunteers to contribute to changing sources, such as OSM, and often remain available for several years. Daily maps illustrating brief overviews of satellite-detected areas of destroyed and possibly damaged structures of different areas of Philippines are provided from November 11, 2013 to November 20, 2013. While the first few are presented only in PDF format, the rest are also offered in Shapefile and ESRI's geodatabase format.

VISOV

The goal of *Volontaires Internationaux en Soutien aux Opérations Virtuelles* [161] or International Volunteers in Support of Virtual Operations (VISOV) "is to help coordinate disaster responses with those of emergency organizations (formal or humanitarian) via digital spaces on which they organize and communicate" [162]. VISOV appears to openly share and maintain its datasets on the website, possibly due to its intention to "become a tool in the hands of local communities" These datasets contain relevant tweets and map tags to estimate the road damage and relief progression VISOV datasets in particular include information such as the type of damage, description of the damage, geographical location, and time of notification. The data is available in the comma separated value (CSV) and keyhole markup language (KML) format from November 11, 2013 to December 3, 2013.

NDRRMC

NDRRMC, a governmental agency of the Philippines, develops detailed situation reports used by many mapping efforts and other situational reports [81]. These PDF reports include information about situation overviews, casualties, affected populations, damaged houses, status of roads and bridges, standees, prepositioned and deployed assets/resources, cost of assistance, cost of damages, status of lifelines (both power and network outage), and emergency management. The status of the roads and bridges demonstrates the damaged areas, declares if the roads are passable and adds remarks such as closing reasons or efforts made to make the roads passable. The level of detail includes even missing persons' names, as well as the coordination efforts (involvement of different governmental and international agencies and humanitarian groups). These reports were initiated immediately after Haiyan on November 8, 2013. NDRRMC retains the situational reports during the recovery operations and appears to archive a large number of files.

DSWD

The Department of Social Welfare and Development (DSWD) appears to play a similar role to NDRRMC. However, it seems to focus on breaking down the information on citizens by geographic regions, as well as statuses within each region such as the number of families in each evacuation center [163]. For the Typhoon Haiyan, DSWD frequently publishes effect, service and intervention reports from November 8, 2013 to December 12, 2013. Viewing previous disasters suggests that DSWD also retains its reports

for several months after the onset of the disaster.

A.2 Information Aggregation Outlets

Humanitarian Response

Humanitarian Response, maintained by OCHA, “aims to be the central website for Information Management tools and services” It appears to compile files from OCHA sectors (Logistics Cluster, Education Cluster, Protection Cluster, etc.), and other groups such as the Canadian Red Cross, Logistics Cluster, MapAction and OSM. The Humanitarian Response website possesses a large number of files, retaining information from several past operations. The outlet provides numerous file filters such as content and data source, and within each filter, multiple items may be selected. Humanitarian Response also maintains a registry of common operational datasets and fundamental operational datasets that often contains files with numerical data, which it claims “should represent the best available datasets for each theme” [164]. The relevant data starts from as early as the moment Typhoon Haiyan hit, and new information is still being uploaded months after the event.

ReliefWeb

As with Humanitarian Response, OCHA maintains the ReliefWeb website. ReliefWeb appears to differ from Humanitarian Response in that it provides files, from situation reports to maps, from a broad range of sources and topics, not focusing on information management to the extent that Humanitarian Response does. ReliefWeb may be effective for identifying primary sources, since it “collects, updates and analyzes from more than 4,000 global information sources” [227]. Alternatively, ReliefWeb may help narrow which sources’ files do not need to be captured right away since it appears to contain most of the files from each source and retains them long after relief operations. The OCHA sourced information about Typhoon Haiyan is directly linked to the ReliefWeb website on the OCHA website. The ReliefWeb page for Typhoon Haiyan was activated on November 8, 2013 and different updates are still being uploaded, as of March 2014.

APAN

All Partners Access Network (APAN) functions similarly to ReliefWeb and Humanitarian Response, but differs mainly in that users upload the files themselves and that the specific page for Typhoon Haiyan is reactionary [166]. User uploaded files allow for the identification of reactionary sources that may be overlooked in the expansive collection of ReliefWeb and Humanitarian Response sources. However, user uploading tends to lack consistency in uploading files from any given primary source, so using APAN as a data retrieval site may be problematic. In contrast, users may sometimes upload files not on a given website but derived from nonpublic datasets, e.g., insurance industry datasets. APAN amalgamates maps, briefs, reports from a variety of different organizations, agencies and groups from November 10, 2013 to January 7, 2014. The community for Typhoon Haiyan provides a link to an ESRI map [167].

Red Crescent Societies (BRC, ARC)

The Red Crescent Societies do not appear to put out files as an overarching system of organizations; rather some individual Red Cross societies may choose to do so on their own. In this case, a collaboration between the American Red Cross and British Red Cross [168] makes available numerous maps throughout the disaster recovery efforts using various data sources. Since the map files specify what data sources each map employs, they may be used to locate the primary sources that contain the desired raw data. Moreover, the maps seem to specify the exact file, e.g., report number, from the source, allowing for direct retrieval of specific data. The Red Cross also provides reports about damage assessments, affected peo-

ple, shelter, etc. They collect information from a variety of resources, such as OSM, UNITAR-UNOSAT, and ReliefWeb.

Google Crisis Maps

Google Crisis Maps, one of the tools of Google Crises Response Group crowdsources data not only within its self-produced facility locations files, but also provides options to access files from sources such as Waze, a traffic mapping application, and CNES/Astrium, which provides satellite imagery [169]. In particular, the self-produced map from Google Crisis Maps plays a role as infrastructure data. However, the facilities that Google Crisis Maps display appear to remain relatively constant at each map, so frequent downloads may not be necessary depending on the goals. The map shows damaged areas, their severity, evacuation centers, and relief drop zone areas. When color coding the damaged areas, the map shows the data as aggregated chunks. However, it is not always clear if this means that the roads to those areas or the roads within that area are closed or not; and more detailed explanation about classification of damages might be useful.

Chapter 5

Incomplete Information Imputation In Limited Data Environments

5.1 Introduction

With the increasing number of large scale disasters, there has been more attention in the operations research community to develop humanitarian logistics models over the past decades. Often, this work presents deterministic models with the assumption of known data [170]. However, complete data is generally not available or hard to gather and integrate in a disaster setting [171]. It is even harder to obtain relevant data for transportation purposes despite the fact that transportation plays a critical role in humanitarian aid [172]. In Chapter 4, discuss the availability of the infrastructural damage and quantify the available road damage information as 8% in the case of Typhoon Haiyan. There is uncertainty about the status of the remaining road segments, which implies that the majority of the information is incomplete.

The first update map provided by by United Nations Institute for Training and Research's Operational Satellite Applications Programme (UNITAR/UNOSAT) three days after the disaster [173] shows road status information about some road segments, highlighted as partially blocked or blocked. However, the UNITAR/UNOSAT map does not provide information about which road segments are open and it is not clear, whether the remaining unmarked roads are open or not. At the same time, field operations managers have to make urgent deployment decisions given this limited available information. This chapter explores how we, as operation researchers, can assist field managers quickly and efficiently make decisions in such settings.

To utilize every available piece of information to the greatest extent possible, we introduce an innovative approach that uncovers similarities between road segments with known road status and utilizes these similarities to fill in the knowledge gaps. We explore whether certain characteristics, defined as *attributes*, such as road type or distance to epicenter, can be used to understand the damage impact across the network. We conjecture that roads with similar attributes have higher correlation in their damage level. For example, if two road segments have similar number of damaged buildings around each of them, and one of the road segments is known to be blocked, it is likely that the other road segment is also blocked. As part of this study, we develop a list of potential attributes that can be used to estimate road status information.

Utilizing new data sources (such as OpenStreetMap (OSM) [174]) along with more traditional data sources for road damage status and relevant attributes (such as portable data format (PDF) maps and

shapefiles (SHP)- a geospatial vector data format for geographical information system), we develop a unifying framework to estimate incomplete information in limited data environments. Since much of the current publicly available transportation network data is disseminated in PDF formats, to make this information useful to decision models, the road status information has to be migrated to a digital geographical information platform such as ArcGIS [175]. This migration process requires extensive georectification (i.e., matching the PDF map to pre-existing road network), identifying the damage and manually entering this information into ArcGIS. ArcGIS is also critical for collecting the attribute information for individual road segments and preparing them for imputation algorithms. We build an ArcGIS model to automate these data gathering and processing steps to the extent possible. This model eliminates a significant amount of the manual work and enables faster data processing for use in future disasters.

Once the available data are captured and processed, we propose various imputation techniques to estimate the missing data. We utilize the properties of the available data structure and the unique characteristic of the post disaster environment to develop new algorithms such as clustering combined with modified mean-and-mode and clustering combined with adjacent arcs. In addition to these unsupervised clustering methods, we also test supervised learning based methods such as classification tree. We also employ global constant methods (optimistic, pessimistic, neutral and popularistic) for benchmarking as they represent reasonable and easy to implement strategies.

We validate our framework with a case study based on the 2010 Haiti Earthquake. We estimate the status of roads after the disaster and demonstrate results with a high level of accurately predicted road damage. We provide key insights from this case study and discuss the applicability of this framework to other disasters. Our test cases are made publicly available for the broader community's use.

The remainder of this paper is organized as follows. Section 2 provides an overview of relevant literature focusing on the use of infrastructural data in humanitarian logistics models. Section 3 presents our framework for imputing incomplete information, while Section 4 presents the results from application of this framework to a past disaster case study. Section 5 discusses the implications of this framework for humanitarian response and broader areas. Finally, Section 6 summarizes the study with concluding remarks.

5.2 Related Literature

5.2.1 Transportation Network Information in Humanitarian Logistics Studies

Existing studies recognize the lack of infrastructure information following a disaster and importance of this information to humanitarian logistics decisions [176]. However, to our knowledge, none of the prior work addresses incomplete information systematically. Below, we summarize the type of transportation network information discussed in the humanitarian logistics literature and examples of different types of data used. We should note that some studies assume Euclidean distances to overcome lack of data; however, we do not include such papers in our review. We also do not discuss the decision problem modeling, but rather focus on the modeling of the transportation networks.

To tackle the incomplete information in transportation networks, different elements of uncertainty are included in the humanitarian logistics models, such as travel time, availability, link capacity, and reliability [177]. One approach to address uncertain travel times is the use of distributions. Shen et al. [178] use lognormal distribution to generate realizations of travel times. Huang et al. [179] model the travel times with a uniform distribution with a lower bound based on the free-flow speed and an upper bound based on a percent of noise added to the lower boundary. Another method is to alter travel time

by multiplying the travel times with a specific coefficient per scenario. Mete and Zabinsky [180] create six scenarios based on the location of the earthquake, an expected earthquake in Seattle, USA triggered by Seattle fault or Cascadia fault, and the time of the day: working hours, rush hours and nonworking hours.

Link availability is also a popular approach to account for the uncertainty in the transportation network [182, 183, 184, 185, 186]. Link availability is generally modeled in two ways. The first option is to randomly destroy a certain percent of the roads [184, 186]. For example, Celik et al. [184] randomly determine the blocked roads and vary the percent of blocked roads in the grid and ring networks. Another method is to assign probability of survival or failure for each link based on different scenarios or instances with respect to intensity and location of the disaster [182, 183]. Gunnec and Salman [182] adjust the link survival probabilities (the complement of link failure probability) with respect to disaster scenarios developed by experts for a possible earthquake in Istanbul, Turkey based on the magnitude and location of the earthquake [187].

Link capacity and reliability are also introduced to capture uncertainty in the transportation network [188, 189, 190, 191, 192]. Barbarosoglu and Arda [188] employ random arc capacities that are generated based on the building damage scenarios developed by Erdik et al. [193] based on the 1999 Marmara earthquake in Istanbul, Turkey. Rawls and Turnquist [190] also use scenario dependent arc capacities to account for the damage in the transportation links for a possible hurricane in North Carolina, USA with the hurricane scenarios coming from HAZUS (geographic information system-based natural hazard developed by the Federal Emergency Management Agency (FEMA) [194]) and their previous research. Since the authors focus on models rather than the data, there is no discussion about how the arc capacities are obtained and scenarios are generated. Yazici and Ozbay [189] develop a procedure to assign capacity loss probability for the specifically chosen links with application to Cape May network in New Jersey, USA. The procedure runs through all arcs to identify one out of every three arcs and assign capacity loss probability to these arcs. Yazici and Ozbay [189] follow a cell based approach, which captures the impact of geography in the capacity loss by assigning higher loss probabilities to cells that are at higher risk. For example, under a flooding case, authors allocate higher capacity loss probabilities to cells near a shore. Noltz et al. [191] use the elementary catastrophe hazard of each arc calculated by catastrophe models for specific regions as risk of an arc to capture link reliability. In addition to reliability, Vitoriano et al. [192] model link security with the probability of a vehicle to be ransacked when traveling through an arc following the Haiti earthquake in 2010.

In summary, researchers use four common elements to model infrastructural damage information: traversing time, link availability, link capacity, and link reliability. However, most of these methods do not incorporate the available information from the field. In addition, there is no systematic scheme to account for the unknown information in the network. This research aims to fill these gaps by developing a framework for estimating the incomplete information in these settings.

5.2.2 Estimating Incomplete Data

In this study, we present various imputation techniques to treat missing values in a more general setting. Imputation methods estimate the missing value through identification of relationships among *attributes* (also known as features or independent variables) [195]. The type of an attribute depends on the set of possible values it can take. Examples of attribute types include: *numeric*, *binary*, *nominal* and *ordinal*. *Numeric* or continuous variables may take any value within a finite or infinite interval. *Binary* refers to the attributes that can take only two values. *Nominal*, also known as categorical, attributes refer to at-

tributes where the value is a symbol or name of an entity [195] and does not have any intrinsic ordering to the category. For example, road type is a nominal attribute with categories such as primary road, highway and pathway. *Ordinal* attributes are similar to categorical variables; however, there is an implicit ordering of the variables. For example, distance to epicenter can be viewed as a continuous attribute. However, when it is categorized into threshold bins, it can be viewed as an ordinal variable. The specific types of attribute and *outcome variable* (also known as dependent variable or response variable) impact the appropriate imputation method to be used for estimating missing data [196]. In this chapter we consider the road status to be the outcome variable that we want to estimate. We assume the road status can only take on categorical values such as open, closed and partially blocked, which impacts our selection of appropriate imputation methods.

A variety of techniques have been proposed for imputing missing data in the literature. Following a disaster, complicated methods might not be appropriate for large data sets due to high computational requirements and urgency to deploy help as quickly as possible. Thus, we tend to focus on simpler methods to estimate missing information. A common unsupervised imputation approach is filling in the missing values with a global constant or mean-and-mode method calculated from the available instances of the corresponding attribute [197, 198]. Mean and mode are used when the attribute is numerical and categorical, respectively. In our setting, since the road status is categorical, a common method used is to find the mode for the available data and apply it to all road segments with unknown road status. Mean-and-mode method is a simple and widely used approach in the literature [199]. However, since this approach does not utilize the underlying correlation among data objects, it can result in poor performance [200]. On the other hand, it can be efficient and effective for the large data sets [201].

Various data mining, machine learning and statistical methods have also been utilized for imputing the missing data. Simple learning methods include the neighbor-based supervised imputation methods such as k-Nearest Neighbor (k-NN) and weighted k-nearest neighbor. k-NN computes the k nearest neighbors of the data objects and finds a value for the missing data based on these neighbors [202]. Mean and mode are used for numerical and nominal values, respectively. Weighted k-nearest neighbors is similar to k-NN, except the estimated value considers the different instances from the neighbors using a weighted average or mode [203]. In our setting, k-nearest neighbors would correspond to adjacent arcs for a given road segment.

Another widely used group of machine learning technique is clustering based methods. As an unsupervised learning method, clustering partitions objects into subsets, known as *clusters*, such that similar objects based on (pretermined criteria) are grouped together to minimize intra-cluster dissimilarity [204, 205, 196]. K-means clustering, where intra-cluster dissimilarity is measured by the sum of the distances among the objects and centroid of the cluster that they are assigned to, divides the data into k groups [196, 199, 205]. The centroid of the cluster refers to the mean value of the objects in the cluster. In terms of clustering, we are interested in algorithms that can efficiently cluster large data sets with categorical variables. While K-means is efficient for clustering numeric data, it does not work with categorical data. The traditional approach of converting categorical data into numeric values does not provide meaningful results when the values of categorical data are not ordered [196]. An alternative approach for utilizing k-means for categorical data is to convert multiple categorical attributes into binary attributes [206]. This increases the computational demand when the number of attribute categories is large and the cluster means do not represent the characteristics of the clusters. K-modes overcomes these issues by using a simple matching dissimilarity measure for categorical objects, replacing means of clusters by modes, and using a frequency-based method to find the modes [196] (see details in Section 5.3.3.2). Once the appropriate clustering method is selected, a rule to fill in the missing value needs to be determined.

Common methods include k-means algorithm with mean-and-mode, weighted distance [199, 198], and inverse distance weighted [205]. Inverse distance weighted method uses similar distance functions as weighted distance method, however it has an inverse relationship with the weight.

Decision tree is another widely used supervised learning method for completing missing information and prediction [202]. A decision tree classifies data objects through recursive partitioning of the instance space [202, 207]. The major advantage of decision trees is the class focused visualization, which allows users to understand the data structure and see the most influential attribute at the root node. Thus, they are more interpretable compared to other classification schemes, such as neural networks and support vector machines [207]. Popular decision tree methods include C4.5 and classification and regression trees (CART), which differ in the measurement of class separations [201, 208]. CART produces either regression or classification of data objects depending on whether the outcome variable is numeric or categorical, respectively. Note that the methods discussed in this section are generally used for data sets with missing values in attributes or in attributes combined with outcome variables. In this chapter, we adopt the classification trees approach to predict the incomplete outcome variable (i.e., road status).

5.3 Methods

This section presents our framework for imputing missing road status information. Road status (our outcome value) refers to the operational status of a road after the disaster, defined here as *unrestricted*, *closed* and *partially blocked*. The first step is to identify the set of attributes (i.e. independent variables) that characterize the road, disaster or the environment that can potentially explain the road damage status in the data. Next, we collect and process information relevant for our analysis: available road status information and attribute data. In addition to unrestricted, closed and partially blocked, we develop a new aggregate road status category called *likely restricted* to improve estimation accuracy. The likely restricted category is an aggregation of partially blocked or closed, i.e., the road is likely restricted to traffic. When limited data about the true state of the damage is available, likely restricted status can be useful in informing operations manager of the possible network damage.

Once the data available are processed, we evaluate several imputation techniques, covering simple methods that follow current practices in the field, as well as other less computationally expensive yet efficient methods. We then validate the framework with a case study based on the 2010 Haiti earthquake.

5.3.1 Attribute Selection

Attributes, such as information related to roads, disaster properties, or geography, help us identify the similarities between road segments to infer unknown information. We utilize existing prior work, domain knowledge and practical availability of data to identify a good representative set of attributes for imputing road status.

We first investigate damage estimation models in the literature. The Federal Emergency Management Agency (FEMA)'s HAZUS multi-hazard models provide damage estimation for flood, hurricane and earthquake events. The models estimate the impact of physical damage to residential and commercial buildings, schools, critical facilities, and infrastructure [194]. The models and characteristics of the damage estimation functions change based on the hazard and its properties. For example, while hurricane models focus on wind characteristics, such as wind speed and direction, flood models focus on the hydrological characteristics, such as water depth, velocity and flood duration [209, 210]. The main focus of HAZUS multi-hazard models is damage to buildings and type of critical characteristics of the building

to estimate the damage changes based on the disaster type. For earthquake and hurricane models, the building type, level of design, and quality of construction play a critical role in estimating damage [211] while for flood models, building materials, building age and configuration are more crucial [212].

Damage modeling for transportation networks is not systemically present in HAZUS. The hurricane model does not include damage estimation to transportation network [213] while the flood model only addresses damage to bridges [214]. The earthquake model provides damage estimation to roads, bridges and tunnels in terms of their classification and level of ground motion or ground failure [215]. The roads are classified as either major or urban roads. The bridge classification uses age, angle of skew, seismic design, bridge type, bridge material, number of spans, maximum span length, and total bridge length. Finally, the tunnels are classified based on the construction type such as bored/drilled or cut & cover tunnels. One limitation of the HAZUS model is that it does not include interdependence of transportation network components [215]. In order to account for interdependence between the road segments, we consider administrative boundaries and information prevalence (i.e., the level of information available) within geographic region in our attribute selection process.

Additional natural disaster damage estimation models exist. Examples include RMS [216], AIR [217], EQECAT [218] for hurricanes, ELER [219], SELENA [220], Openquake [221] for earthquakes and FLEMO [222], RAM [223], and European Commission/ HKV [224] for floods. To our knowledge, most of these models do not include infrastructural damage estimation while Huizinga et al. [225] suggest global infrastructural damage function based on water depth for flooding.

In addition to reviewing existing damage prediction models, we interview practitioners from humanitarian logistics, on-ground disaster responders, digital humanitarians, and public officials. Discussions with the domain experts help us translate damage models into compact metrics such as floodplain, geology, elevation and land use.

Public availability of data is critical for practical use of imputation algorithms. For example, while water depth can be helpful in estimating whether a bridge or road is blocked, this information is often hard to obtain immediately after a disaster at the individual road segment level.

Table 5.1 presents a list of attributes that can be practically collected before and after the disaster. To develop a general framework that can be applied to future disasters, we include numerous attributes. The list of attributes to include in algorithms might change per disaster type, disaster location and availability of data. Next, we explain each of the proposed attributes.

Pre-disaster attributes: These attributes represent the properties of road segments or geography that can be collected prior to a disaster. Since the impact of a disaster depends on disaster and location specific properties [226], these attributes can increase the accuracy of our road status predictions. These attributes are as follows:

- Road type: Road type is a functional classification of a road and is similar to the road classification used in HAZUS [215]. The categories used to classify the road type might change depending on the available data source.
- Administrative boundaries: We use administrative boundaries (e.g. county or city limits) to obtain geographical information for individual regions.
- Grid: In order to accompany the high level information provided by administrative boundaries, we also use ArcGIS's grid feature, which creates geographic location indicators, to have more details about the geography.

Table 5.1: List of potential attributes for road status estimation

Pre-disaster attributes	Post-disaster attributes
Road type	Building damage level
Administrative boundaries	Distance to focal point of the disaster
Grid	Information prevalence
Building density	
Land use	
Flood plain	
Geology	
Elevation	
Age	
Construction material	
Soil type	

- Building density: Building density represents information about the surrounding buildings for each road segment. When the information regarding a building's height or the number of stories is available, we can find the weighted building density per road segment. One potential way to compute this is to divide the sum of adjacent buildings multiplied by their heights by the length of the arc. When the height information is not available, we calculate the surrounding building density by dividing the total number of adjacent buildings by length of the arc.
- Land use: Motivated by the building types used in HAZUS, land use identifies the management and development of land. The main categories include residential, industrial, and commercial.
- Flood plain: We include the comparative flood inundation probability for each given location, corresponding to their flood hazard: low, medium or high.
- Geology: Following the HAZUS earthquake model, we include the geology data that identifies the local rock types.
- Elevation: Road elevation corresponds to the height of the ground where the road segment is built.
- Age: Age of the road refers to the year the road was built.
- Construction material: Motivated by the bridge damage modeling in HAZUS, this attribute accounts for the construction material of the roads, such as asphalt or concrete.
- Soil type: Soil type is the layer that covers the rock. It can be an important attribute in the case of earthquakes accompanied by landslides or floods.

Post-disaster attributes: These attributes represent road segment properties that can be collected after the disaster strikes. Disaster size, magnitude of destruction and extent of the geographical impact area drive logistical needs [226].

- Building damage level: We consider the total building damage and severity of the damage per road segment as potential attributes to capture this.

- Distance to focal point of the disaster: We calculate the distance from the midpoint of each road segment to the focal point of the disaster. For an earthquake, this metric measures the distance to the epicenter. For a hurricane, it can be measured as distance to hurricane path or storm surge.
- Information prevalence: We utilize the level of available information, such as percent of missing data and percent of road segments that are open, closed and partially blocked. We study this information at the grid level, as well as the entire operation area.

5.3.2 Data Collection and Processing

Figure 5.1 presents the data gathering and processing steps. Two primary types of data are collected and processed: road status information and attribute data. The first step is to gather road status data from information outlets that release data immediately following a disaster, such as OpenStreetMap (OSM) and ReliefWeb [174, 227]. The majority of such data is in the form of PDF maps and PDF reports rather than in the desirable shapefiles (SHP), see [228]. Attribute data can be gathered from a variety of resources such as OSM and UNITAR/UNOSAT, generally in the form of shapefiles.

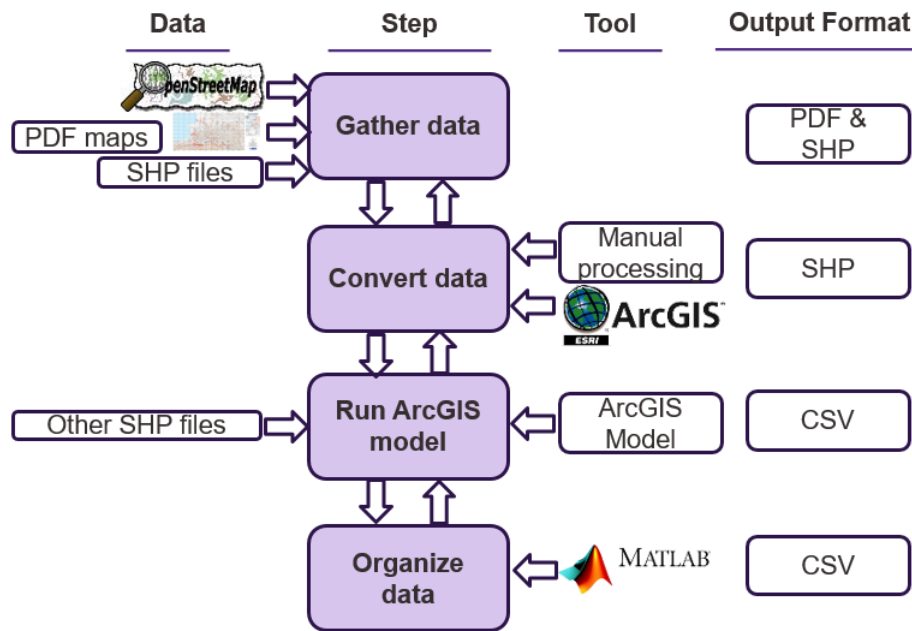


Figure 5.1: Data Collection and Processing Steps.

Next, the collected data are transformed into a format that can be integrated into humanitarian logistics models. We first geo-rectify the PDF maps (i.e., match the road network from PDF map to the digital environment), identify the damage, and manually enter the damage information into the ArcGIS where it can be easily extracted to a comma separated value format. We segment roads by intersections to create smaller links called *road segments*. This is particularly important for about where exactly identifying the damage location along a road. Following this data conversion step, we have the available road status information in ArcGIS (in SHP format) by segments, which enables the attribute information collection per road segment as well.

In order to collect the attribute information (gathered in a CSV format) efficiently, we build a model in ArcGIS that provides an automated approach to merge road status information data with the attribute data construct for each road segment. Our unit of analysis, road segment, is represented in line data, however, much of the available information is presented at different units, both larger (e.g., polygon data such as administrative boundaries) and smaller (e.g., point data such as building damage). The ArcGIS model transforms these data streams to the road segment unit definition. For example, buildings are attached to the closest road segment thus transforming the point data into line data. The benefit of our GIS model (available at http://users.iems.northwestern.edu/kezbani16/ArcGISModelandData/Paper_Submission.zip along with the underlying data) is that once all the data are collected we can link the new data sets and then run the model. Thus, the model eliminates most of the manual steps in integrating data and improves reproducibility, as well as gives ability to others to make changes. As a result, it reduces the data processing time for future disasters and facilitates real-time information integration. Details of the model are provided in the Appendix 5.6.

After the attribute information for each road segment is extracted using our ArcGIS model, some attributes might need to be classified and organized into categorical data desired for specific imputation algorithms (see Section 5.3.3). For example, the distance to epicenter, initially given in continuous values, needs to be transformed into categorical values. After completing this step, a CSV file with all road status and categorical attribute information is created for use with the proposed imputation methods.

5.3.3 Data Imputation

In this section we present data imputation techniques to estimate the missing road segment status. Table 5.2 summarizes the methods employed in this study.

Table 5.2: Types of Imputation Techniques Used in This Study

Imputation Technique Group	Imputation Method
Global Constant [Benchmark]	Optimistic
	Pessimistic
	Neutral
	Popularistic
Cluster Based	Modified Mean-and-Mode
	Adjacent Arc
Decision Tree Based	Classification Tree

5.3.3.1 Global Constant Based Imputation Rules

We first consider four simple imputation approaches based on the global constant method and mean-and-mode method as they are reasonable for benchmark [198]. In this set of approaches, all unknown values take the same constant value.

Optimistic: An “optimistic” view assumes that each unknown entry is realized at its best value.

Pessimist: A “pessimistic” view assumes that each unknown entry is realized at its worst value.

Neutral: A “neutral” view assumes that each piece of unknown information takes on its intermediate value. For example, each road segment with unknown information is assigned to be accessible with some damage or debris that slows the traffic flow.

Popularistic: A “popularistic” view assumes that each piece of unknown information is set to its most common value, i.e., the popularistic view assigns the missing value to its mode.

5.3.3.2 Cluster (Unsupervised Learning) Based Imputation Rules

Embedding global constant methods within more sophisticated approaches has been shown to improve their performance [201]. One such technique is the cluster based method. Here, we first cluster the road segments based on similarities in attributes and then compute the missing values in each cluster with imputation rules [199, 201, 203]. For example, following an earthquake, one can cluster the road segments based on attributes such as road type and distance to epicenter. Then, imputation rules, such as modified mean-and-mode method or adjacent arc method, can be used to estimate the unknown road segments’ status by utilizing the information for the known road segments within the same cluster. Figure 5.2 shows the general steps to impute the missing information with the clustering based methods.

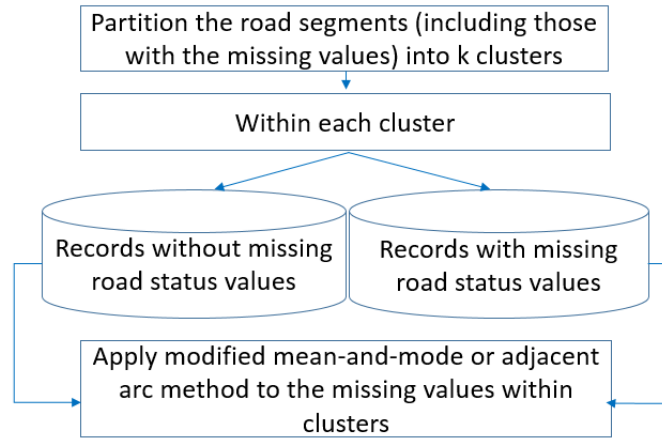


Figure 5.2: Missing Value Imputation with Clustering.

The attributes we use in clustering are categorical; thus we employ k-modes algorithm by Huang [196], see Algorithm 1. The solutions of the k-modes algorithm are sensitive to the choice of initial cluster centers [229]. To overcome local optimality in k-modes clustering, we use the cluster center initialization algorithm discussed in [230]. After forming the clusters, we employ one of the following cluster based imputation rules to complete the missing data.

Clustering Combined with Modified Mean-and-Mode Imputation (CCMMM)

We develop a simple mean-and-mode consolidated method to assign values to unknown information that integrates cluster level information, as well as information from all road segments. Among the road segments with known status, majority of the road segments are unrestricted, causing the classical mean-and-mode method to consistently assign unrestricted status to all road segments within a cluster. To overcome this, in Algorithm 2, we define a new parameter called “critical road status value” that represents the type of road status that has the lowest percentage of appearance in the data. (In our example, the critical road status value is “closed” for overall data and majority of the clusters.) Then, the CCMMM

Algorithm 1 K-modes Algorithm, Huang [196]

- 1: **Input:** All data
 - 2: $k \leftarrow$ Number of clusters
 - 3: **Output:** Clusters
 - 4: Select k initial modes, one for each cluster
 - 5: Allocate an object to the cluster whose mode is the nearest to it according to dissimilarity measure. Update the mode of the cluster after each allocation.
 - 6: After all objects have been allocated to clusters, retest the dissimilarity of objects against the current modes. If an object is found such that its nearest mode belongs to another cluster rather than its current one, reallocate the object to that cluster and update the modes of both clusters.
 - 7: Repeat line 6 until no object has changed clusters after a full cycle test of the whole data set.
-

algorithm adjusts the percent of roads segments with the critical status for each cluster to the global critical value road segment status. Thus, this method takes advantage of the relative information about the road status values that are not dominant in the data, yet critical to be accurately captured for post-disaster operations.

Clustering Combined with Arc Adjacent Method (CCAA)

Adjacent arc method is another cluster based method to impute unknown information. This method is inspired by the $k - NN$ method however, instead of having a specified k nearest neighbors, we utilize the information from the adjacent arcs (i.e., the neighboring arcs sharing a common endpoint), including the percent of unknown arcs adjacent to that road segment (see Algorithm 3). This modification allows us to integrate specific domain knowledge into the imputation methods, which can improve the estimation accuracy [200]. Similar to the CCMMM method, CCAA method also compares the percentage of partially blocked and closed roads within the cluster to these statistics for the entire data set.

5.3.3.3 Decision Tree Based Imputation

We next consider a decision tree based approach to impute the unknown transportation network information.

Classification Trees (CT) (Supervised Learning)

Classification trees employ a set of independent variables (attributes in our case) to predict class to which an object belongs [231]. (For details about classification trees, we refer readers to [232].) Recursive partitioning is used to construct a tree that has more homogeneous subsets [233]. Similar to clustering, classification trees take advantage of the attributes and categorize data based on the decision classes. However, in classification trees, we consider the entire data set with no missing values to inform the missing values.

The data set is partitioned into two sets: a set with no missing road status values (known as training set) and a set of records with missing road status values (test set). We train the classification tree with training data and select the road status to be predicted by the classification tree. Next, we feed the test set to the model and impute the predicted value for the missing field, the road status. Following these steps, the data objects with missing values are imputed by the classification tree. Figure 5.3 shows the missing value imputation process using classification trees (adapted from [202]).

Algorithm 2 Missing value imputation process with clustering combined with modified mean-and-mode method

```

1: Notation Used:
2:  $p(r, i)$  percent of road status type  $r = \{u, l, p, c\}$  in cluster  $i \in \{0..k\}$  where  $\{u, l, p, c\}$  denote unre-
   restricted, likely restricted, partially blocked and closed, respectively,  $i = 0$  represents the entire data
   and  $i = \{1..k\}$  denote the clusters
3:  $m$  as mode of entire known road status value
4:  $c(i)$ : critical road status value for cluster  $i \in \{0..k\}$ 
5: Input: Clusters
6:  $k \leftarrow$  Number of clusters
7: Output: Imputed value of missing road status data
8: Calculate  $p(r, 0)$  for all  $r \in \{u, l, p, c\}$ 
9: Identify  $c(0)$ 
10: for  $i \in 1 \dots k$  do
11:   Calculate  $p(r, i)$  for all  $r \in \{u, l, p, c\}$ 
12:   Identify  $c(i)$ 
13: end for
14: for  $i \in 1 \dots k$  do
15:   for  $j \in 1 \dots |cluster_i|$  do
16:     if data object( $i, j$ ) is unknown then
17:       if  $p(u, i) = 0$  then
18:         if  $p(c, i) \geq p(c, 0)$  then
19:           Set unknown data object( $i, j$ ) as closed
20:         else
21:           Set unknown data object( $i, j$ ) as  $m$ 
22:         end if
23:       else if  $p(l, i) \geq p(l, 0)$  then
24:         if  $p(p, i) \geq p(c, i)$  then
25:           Set unknown data object( $i, j$ ) as partially blocked
26:         else
27:           Set unknown data object( $i, j$ ) as closed
28:         end if
29:       else if  $p(c(i), i) \geq p(c(i), 0)$  then
30:         Set unknown data object( $i, j$ ) as value of critical road status value
31:       else
32:         Set unknown data object( $i, j$ ) as  $m$ 
33:       end if
34:     end if
35:   end for
36: end for

```

Algorithm 3 Missing value imputation process with adjacent arc

```

1: Notation Used:
2:  $p(r, i)$  percent of road status type  $r = \{u, l, p, c\}$  in group  $i \in \{0..k\}$  where  $\{u, l, p, c\}$  denote unrestricted, likely restricted, partially blocked
   and closed respectively;  $i = 0$  represents the entire data and  $i = \{1..k\}$  denote the clusters
3:  $a(r, i, j)$  percent of road status type  $r = \{u, l, p, c, h\}$  adjacent to road segment  $j$  in group  $i \in \{1..k\}$  where  $r = h$  represents the unknown road
   status
4:  $m$  : mode of the entire known road status value
5: Input: Clusters
6:  $k \leftarrow$  Number of clusters
7: Output: Imputed value of missing road status data
8: Calculate  $p(r, 0)$  for all  $r \in \{u, l, p, c\}$ 
9: for  $i \in 1 \dots k$  do
10:   Calculate  $p(r, i)$  for all  $r \in \{u, l, p, c\}$ 
11: end for
12: for  $i \in 1 \dots k$  do
13:   for  $j \in 1 \dots |cluster_i|$  do
14:     if  $data\ object(i, j)$  is unknown then Calculate  $a(r, i, j)$  for all  $r \in \{u, l, p, c\} \forall i \in \{1..k\}$ 
15:       if  $a(u, i, j) + a(h, i, j) = 1$  then
16:         if  $p(l, i) \geq p(l, 0)$  then
17:           if  $p(p, i) \geq p(c, i)$  then
18:             Set unknown  $data\ object(i, j)$  as partially blocked
19:           else
20:             Set unknown  $data\ object(i, j)$  as closed
21:           end if
22:         else
23:           Set unknown  $data\ object(i, j)$  as  $m$ 
24:         end if
25:       end if
26:       if  $a(p, i, j) + a(c, i, j) > p(l, 0)$  then
27:         if  $a(p, i, j) \geq a(c, i, j)$  then
28:           Set unknown  $data\ object(i, j)$  as partially blocked
29:         else if  $a(p, i, j) = a(c, i, j)$  then
30:           if  $p(p, i) \geq p(c, i)$  then
31:             Set unknown  $data\ object(i, j)$  as partially blocked
32:           else
33:             Set unknown  $data\ object(i, j)$  as closed
34:           end if
35:         else
36:           Set unknown  $data\ object(i, j)$  as closed
37:         end if
38:       end if
39:     end if
40:   end for
41: end for

```

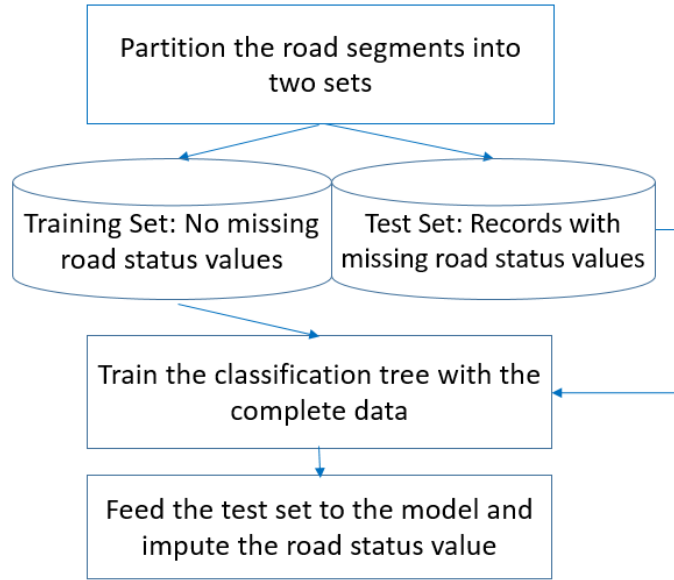


Figure 5.3: Missing Value Imputation Process with Classification Tree Adapted from [202].

5.3.4 Validation

The final step in this framework is validation. The goals of validation are to evaluate the accuracy of the model and to determine whether certain characteristics of the missing data impact the accuracy. For example, one could imagine missing data are grouped geographically, corresponding to disaster damage assessment resources directed to specific geographic regions.

During the validation, one might need to revisit the framework multiple times to check feasibility of assumptions and improve the estimates of unknown information. The nature of the problems we work with are generally computationally intensive, thus, in the validation step, we initially test a smaller data set to learn about the algorithm. Once the framework is tested on the small instances, we apply the algorithms to the entire data set.

5.4 Application to 2010 Haiti Earthquake Case Study

Among the publicly available relevant post-disaster data sets, the information stemming from the 2010 earthquake in Haiti proved to be the most complete. These publicly available data include the initial damage assessment of buildings in and around Port-au-Prince and Carrefour, and maps showing the road status (whether they are partially or completely blocked by damage or debris). Figure 5.4 shows the timeline of damage maps released in the days following the Haiti earthquake. We first describe the data used in this case study, both road damage and attribute data (Section 5.4.1). Then, Section 5.4.2 presents the results of the case study for Carrefour.

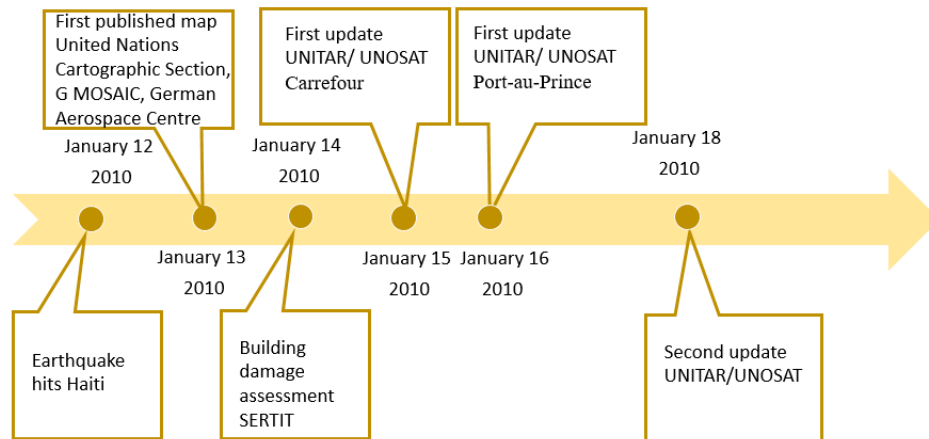


Figure 5.4: Timeline After Haiti Earthquake.

5.4.1 Data Used in the Case Study

5.4.1.1 Road Damage Data

The earthquake struck Haiti on January 12, 2010 at 4:53 pm local time. The first map showing damage was published on January 13, 2010 by the United Nations Cartographic Section, G MOSAIC, and the German Aerospace Centre [173] (see Figure 5.4). This Port-au-Prince map showed both building damage and road conditions. Roads were identified as either “limited access” or “blocked”, while building damage was not categorized (i.e., only identified as “damaged building”).

On January 14, 2010, SERTIT (Service Régional de Traitement d’Image et de Télédétection - a regional image processing and remote sensing service) provided a map of building damage assessment for each urban block of Port-au-Prince [234]. The map did not specifically identify road damage however, it showed the highly impacted areas, which can be used to infer the road damage areas. The building damage was categorized into three groups: obvious/widespread damage (> 40%), evident but sporadic damage (11–40%), and scarce or nonvisible damage (0–10%). It should be noted that even with relatively high resolution of the map, the details on the provided PDF map are difficult to see.

The first updated maps, provided by UNITAR / UNOSAT (United Nations Institute for Training and Research UNITAR’s Operational Satellite Applications Programme) on January 15th for Carrefour and January 16th for Port-au-Prince, contained information about road status as either “likely restricted by debris” or “likely closed by debris” [235]. On January 18, 2010, a second map update was provided by UNITAR/ UNOSAT in the form of damage density maps, showing the color coded relative road damage density by geographic areas [236].

There is no consensus on the details of the damage classification in the maps. The first released map [173] uses “limited access” and “blocked” categories for the road status information, while the first update map [235] and the second updated map [236] show “likely closed by debris” or “likely restricted by debris” road status. Furthermore, while the first map shows the individual building damage without categorizing the damage level, the SERTIT map [234] shows the building damage by block. A nice feature of this map is that it distinctly states the areas that are not analyzed. This is important for identifying the unknown areas and estimating and preparing logistical needs accordingly. The first updated maps for Carrefour and Port-au-Prince show that approximately 4.2 % roads are analyzed, where approximately

2.8% of the roads are identified as partially blocked, 1.4% of the analyzed roads are closed and the remaining 95.8% is unrestricted.

Since the products released to the public and international community were published as PDF maps, we first geo-rectify the maps over the base map of Haiti in ArcGIS software to convert the data as discussed in Section 5.3.2. Among the provided maps, we utilized the first updated map provided by UNITAR/UNOSAT focusing on Carrefour rather than Port-au-Prince since it provides a good level of details for road conditions while still being visible enough to geo-rectify manually. We systematically processed square grid by grid and manually digitized all mapped damage locations (partially blocked or closed) in ArcGIS. Figure 5.5 shows an example of 1 km grids in Carrefour.

We should note that the time needed for georectification depends on the size of the impacted area, number of the impacted roads, and the quality of available maps. In the study area, there are 6952 road segments and 290 of them have some kind of damage. It took our team almost 4 hours to migrate this information into ArcGIS. For future disasters, if the damage information is provided in PDF format, one might need to spend similar time or even more (depending the area affected) to integrate this information. On the other hand, with the developed ArcGIS model, once all the data are collected, one can link the new datasets and then run the model.

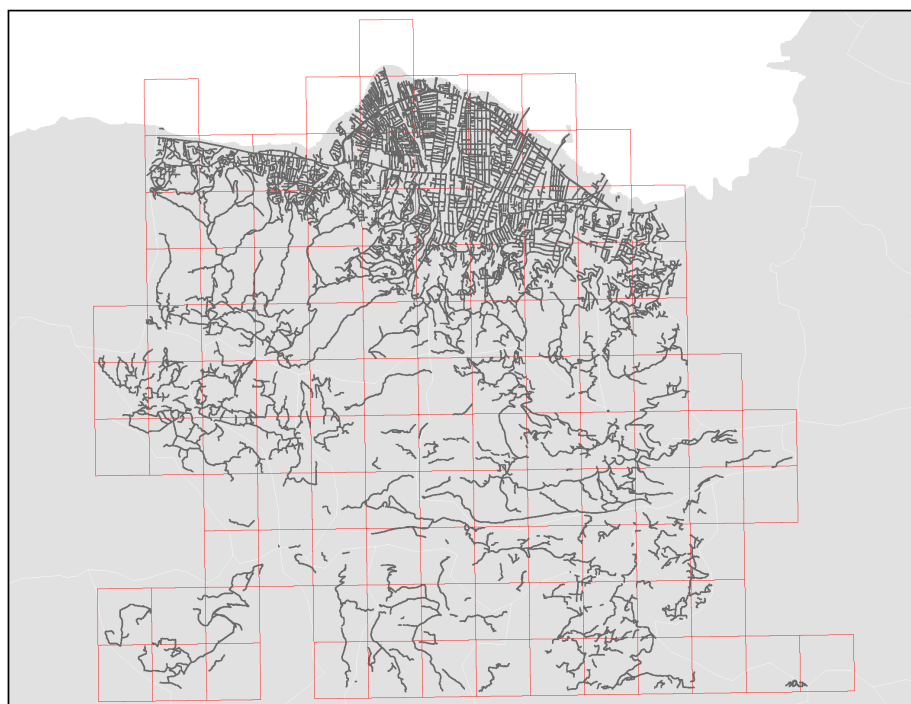


Figure 5.5: 1 km Grid of Carrefour.

5.4.1.2 Attributes

In this section, we discuss how the attributes data are collected and how the attributes are organized for imputation. Table 5.3 lists a selection of attributes for this case study. Limited by the data availability, we

use a subset of the attributes presented in Table 5.1. The details about the attributes and the data used for each attribute is provided in the Appendix 5.6.

Table 5.3: Attributes used in the case study to estimate the road status

Pre-disaster attributes	Post-disaster attributes
Road type	Building damage level
Administrative boundaries	Distance to epicenter
Grid	Information prevalence
Building density	
Land use	
Flood plain	
Geology	

We use the ArcGIS model discussed in Section 5.3.2 to convert the attribute information into line data. We develop two attributes to transform the building damage information to the road segment level. For the first attribute, we find the percentage of damaged building along each road segment. Next, among the damaged buildings, we compute at the percentage of heavily damaged buildings, which includes very heavy damage (grade 4) and destruction (grade 5) since according to the EMS 98 these two categories correspond to failures and collapses of the structures [237]. We should note that before we developed the ArcGIS model this step took our team two days to complete. Once the model was built, the processing time reduced to 15 minutes for the same operations. We then classify the attribute information to obtain categorical variables when required for the imputation methods, such as building density, building damage level and distance to epicenter. For each of these categories, we use Matlab to form five equal sized/length bins to be consistent with number of bins used for the building damage.

5.4.2 Results

To test our imputation methods, we first focus on downtown Carrefour, a smaller area with approximately 14 % of the road segments' status information provided as either partially blocked or closed. We study the impact of percentage of missing data, dispersion of missing data and imputation techniques. We then extend the analysis to the entire Carrefour area, where we include our study of the impact of the level of geographical detail by altering the grid size.

As discussed in Section 5.3.4, we first need to find the true states of the road segments with missing road status information, however no detailed road status analysis following the Haiti earthquake was conducted. Therefore, based on the interviews with experts involved in the field operations following the earthquake and a review of maps posted later based on the same satellite images such as [236], we make the assumption that the road segments with missing data are unrestricted. Then, in order to test the performance of our imputation methods, we introduce a certain percentage of missing value as holdout data varying from 10% to 50% in 10% increments (similar to work in [200, 205]) under random holdout or geographic holdout scheme. A *random holdout scheme* assumes that the probability of a data object having missing value is independent of both the known values and the missing data (this is analogous to the missing completely at random assumption from imputation literature [200]). Thus, under the random holdout scheme, the data are removed randomly from the entire study region. A *geographical holdout scheme* assumes that the probability of a data object having missing value may depend on the

known values but not the value of the missing data itself (this is equivalent to missing completely at random assumption from imputation literature [200]). Under the geographical holdout scheme, the data are removed randomly from each geographical unit (i.e, grid). Holdout type also can be interpreted as the missing data spread. We use the remaining data as the training data set and predict the road damage status of the test data using the various imputation techniques discussed in Section 5.3.3.

We use Matlab to run our experiments. For each experiment type (each missing data percentage and holdout combination), we simulate 50 runs. Since some grids contain a relatively small number of road segments, repeated sampling from these grids to obtain the missing data results in the same underlying network with missing data as we increase the number of runs. Thus, we set the limit at 50 runs. We set the number of cluster, k , as 10 to balance the trade off between coarseness of our estimation and ensuring the sufficient number of arcs within each cluster [203]. We use “*fitctree*” function in Matlab to build classifications that utilize the standard classification and regression trees introduced by Breiman et al. [232, 238].

Tables 5.4 and 5.5 show how we calculate the success rate for one of the runs. Table 5.4 tracks the number of road segments identified as unrestricted, partially blocked and closed for each imputation method. The second and third columns show the distribution of the missing data amongst the road status types in the original data, while columns 4-6 show the distribution of the imputed values. The boxed values on the diagonals in Table 5.4 represent the number of correctly identified road segments. Then using Table 5.4, we calculate the allocation percent of imputed roads for each road status type in the original data (illustrated in Table 5.5). Rows in Table 5.5 demonstrate the conditional probabilities of an imputed road status type (columns 3-6) given a road status type in the original data and the imputation method. We also track the probabilities for likely restricted road status type in this analysis. Recall, likely restricted refers to the union of partially blocked and closed road segments. We should note that the likely restricted road status is calculated a posteriori, after predicting the other three possible road statuses (unrestricted, partially blocked and closed). The diagonals in Table 5.5 represent the accuracy of an imputation method for each road status type. Accuracy is calculated by the correctly projected number of road segments for a road status type divided by the number of road segments in the holdout for that road status type.

We take the average of accuracy results over 50 runs to find the average success rate of imputation methods. In the remainder of this section, we analyze the results from clustering based and decision tree based imputation methods in more detail.

5.4.2.1 Application to Downtown Carrefour

We first provide results from Downtown Carrefour area where approximately 9%, 5% and 86% of the road segments have partially blocked, closed and unrestricted status in the original data, respectively. We track the average success rate of each method as a function of the percentage of missing data for two types of missing data spread (geographical and random holdout). We observe an average success rate of up to 81% for correctly identifying likely restricted road segments. We should remind that likely restricted is calculated based on the partially blocked and closed road status and it measures the probability of damage in the road segment (either partially blocked or closed) rather than the specific road status. For the majority of cases, when the percent of missing data increases, corresponding to less information being available to make a decision, the success rate decreases.

In order to test the effects of multiple factors on average success level, Table 5.6-5.8 show the results of the t -tests (p-value) analyzing the impact of the percent of missing data, spread of missing data and

Table 5.4: Accuracy Calculation Step 1: Example for Downtown Carrefour with Geographical Holdout

Imputation Type	Road Type in the Original Data	Holdout	Road Type in the Imputation Results		
			Unrestricted	Partially Blocked	Closed
Global Constant: Optimistic	unrestricted	201	201	0	0
	partially blocked	23	23	0	0
	closed	11	11	0	0
Global Constant: Pessimistic	unrestricted	201	0	0	201
	partially blocked	23	0	0	23
	closed	11	0	0	11
Global Constant: Neutral	unrestricted	201	0	201	0
	partially blocked	23	0	23	0
	closed	11	0	11	0
Global Constant: Popularistic	unrestricted	201	201	0	0
	partially blocked	23	23	0	0
	closed	11	11	0	0
Clustering Combined with Modified Mean-and-Mode	unrestricted	201	92	53	56
	partially blocked	23	2	15	6
	closed	11	0	5	6
Clustering Combined with Adjacent Arc	unrestricted	201	123	67	11
	partially blocked	23	2	20	1
	closed	11	1	8	2
Classification Tree	unrestricted	201	172	19	10
	partially blocked	23	17	5	1
	closed	11	8	1	2

imputation method type, consecutively. In Table 5.6, we investigate the percent of missing information where the rows show the imputation method type and the columns represent the spread of missing data. Each cell tracks the p -value per road status type. Table 5.6 shows that for almost all cases, the percent of missing data seems to play a significant role on the average success level. We next investigate the impact of missing data spread. Since the percent of missing information was significant, we conduct a two-way ANOVA with interaction. In Table 5.7, rows correspond to the imputation method types while columns 3-5 represent the percent of missing data, spread of missing data and their interaction, respectively. Similarly, each cell tracks the p -value per road status type. As seen in Table 5.7, for nearly all cases, the missing data spread type does not have a significant impact on the data estimation success rate, especially for clustering based methods. The results show that there is no interaction between level of available information and spread of missing data. In Table 5.8, we illustrate the impact of imputation method, where the rows correspond to the spread of missing data while columns 3-5 represent the percent of missing data and imputation method. Table 5.8 demonstrates that selected imputation method impacts the average success rate significantly, however there is no interaction between percent of missing information level and missing data spread type.

Observation: *As the percent of missing information increases, the average success level tends to decrease. However, this trend is not a steep trend.*

Insight: Due to the low presence of certain road status types, obtaining more information does not have a dramatic impact. Beyond a certain percent of missing data in some grids, the remaining road status information will have even fewer road segments of closed and/or partially blocked road segments or none at all. Thus, obtaining more information will not impact the success dramatically.

Table 5.5: Accuracy Calculation Step 2: Example for Downtown Carrefour with Geographical Holdout

Imputation Type	Road Type in the Original Data	Road Type in the Imputation Results			
		Unrestricted	Partially Blocked	Closed	Likely Restricted
Global Constant: Optimistic	$P(\cdot U)$	1	0	0	0
	$P(\cdot PB)$	1	0	0	0
	$P(\cdot C)$	1	0	0	0
	$P(\cdot LR)$	1	0	0	0
Global Constant: Pessimistic	$P(\cdot U)$	0	0	1	0
	$P(\cdot PB)$	0	0	1	0
	$P(\cdot C)$	0	0	1	0
	$P(\cdot LR)$	0	0	1	0
Global Constant: Neutral	$P(\cdot U)$	0	1	0	0
	$P(\cdot PB)$	0	1	0	0
	$P(\cdot C)$	0	1	0	0
	$P(\cdot LR)$	0	1	0	0
Global Constant: Popularistic	$P(\cdot U)$	1	0	0	0
	$P(\cdot PB)$	1	0	0	0
	$P(\cdot C)$	1	0	0	0
	$P(\cdot LR)$	1	0	0	0
Clustering Combined with Modified Mean-and-Mode	$P(\cdot U)$	.46	.26	.28	.54
	$P(\cdot PB)$	.09	.65	.26	.91
	$P(\cdot C)$	.00	0.45	.55	1.00
	$P(\cdot LR)$	.06	.59	.35	.94
Clustering Combined with Adjacent Arc	$P(\cdot U)$	.61	.33	.05	.39
	$P(\cdot PB)$	.09	.87	.04	.91
	$P(\cdot C)$	.09	.73	.18	.91
	$P(\cdot LR)$	.09	.82	.09	.91
Classification Tree	$P(\cdot U)$	.86	.09	.05	.14
	$P(\cdot PB)$	.74	.22	.04	.26
	$P(\cdot C)$	.73	.09	.18	.27
	$P(\cdot LR)$	.74	.18	.09	.26

Observation: Spread of missing data does not have a significant impact on the average success level.

Insight: There is high variance in the distribution of various road statuses: 86% unrestricted, 9% partially blocked and 5% closed. Similar to the insight above, with the low presence of damage, obtaining more information either randomly or geographically does not have a significant difference on the average success.

Observation: Average success rate depends on the imputation method and road status type.

Insight: If an operations manager is interested in identifying the damage (i.e., likely restricted roads), he or she should use clustering based methods. On the other hand, the results show that if an operations

Table 5.6: Summary of p-values on Impact of Percent of Missing Data Tests for Downtown Carrefour

Imputation Technique	Road Status Type	Missing Data Spread Type	
		random	geographic
Clustering Combined with Modified Mean-and-Mode	unrestricted	0.04	0.00
	partially blocked	0.09	0.01
	closed	0.00	0.00
	likely restricted	0.08	0.00
Clustering Combined with Adjacent Arc	unrestricted	0.02	0.28
	partially blocked	0.88	0.00
	closed	0.18	0.50
	likely restricted	0.03	0.01
Classification Tree	unrestricted	0.03	0.00
	partially blocked	0.01	0.12
	closed	0.01	0.00
	likely restricted	0.00	0.00

Table 5.7: Summary of p-values on Impact of Missing Data Spread Type for Downtown Carrefour

Imputation Technique	Road Status Type	Percent of Missing Data	Missing Data Spread Type	Interaction
Clustering Combined with Modified Mean-and-Mode	unrestricted	0.00	0.84	0.77
	partially blocked	0.10	0.37	0.09
	closed	0.00	0.85	0.64
	likely restricted	0.00	0.48	0.30
Clustering Combined with Adjacent Arc	unrestricted	0.25	0.57	0.11
	partially blocked	0.01	0.56	0.33
	closed	0.13	0.75	0.61
	likely restricted	0.01	0.25	0.27
Classification Tree	unrestricted	0.00	0.68	0.99
	partially blocked	0.00	0.02	0.01
	closed	0.00	0.77	0.61
	likely restricted	0.00	0.12	0.13

Table 5.8: Summary of p-values on Impact of Imputation Method Comparison for Downtown Carrefour

Missing Data Spread Type	Road Status Type	Percent of Missing Data	Imputation Technique	Interaction
Random	unrestricted	0.25	0.00	0.95
	partially blocked	0.06	0.00	0.71
	closed	0.01	0.00	0.00
	likely restricted	0.00	0.00	0.98
Geographic	unrestricted	0.36	0.00	1.00
	partially blocked	0.00	0.00	0.99
	closed	0.00	0.00	0.00
	likely restricted	0.00	0.00	1.00

manager is interested in identifying the unrestricted roads more accurately, classification trees should be used. This makes sense since clustering based methods utilize critical road status value in the outcome and adjacent arc information. Additionally, classification trees utilize entire training data for predictions where unrestricted road segments have high presence.

5.4.2.2 Application to Entire Carrefour

The percent of road segments with partially blocked or closed road status decreases to around 4.2% in the entire Carrefour area. We still keep the grid size in ArcGIS as one km in our analysis.

Similar to the steps described in Section 5.4.2.1, we vary the percent of missing data for both spread types and impute the missing values using the seven imputation methods for 50 runs for each experiment. We obtain an average success rate of up to 87% in estimating the road segments with likely restricted road status. We also conduct the t -tests to investigate the significance of the factors. Similar to the results for the downtown area, Tables A1 and A3 show that the percent of missing data and the imputation method play a significant role while the missing data spread type does not have a significant role in accurately estimating the missing data (Table A2).

We also investigate the impact of focus area (see Tables 5.9 and 5.10.) Mapping focus area refers to choosing either downtown Carrefour or entire Carrefour area for results analysis. The rows represent the imputation method types while columns 3-5 show the percent of missing data, mapping focus area and their interaction, respectively. Tables 5.9 and 5.10 demonstrate that the geographical focus area has a significant impact on the success of road status estimation.

Following close analysis of our results, we observe that employing clustering based imputation techniques in both random and geographic holdout schemes using entire Carrefour area result in significantly better average road status estimation success rates for unrestricted, partially blocked and likely restricted road statuses. On the other hand, focusing on the downtown area provides a higher success rate for estimating the road segments with closed road status in clustering combined with adjacent arc, while both geographical focus areas are indifferent in clustering combined with mean-and-mode technique. When utilizing classification trees in both random and geographic holdout schemes, using entire Carrefour area results in significantly better average success rates for estimating the road segments with unrestricted road status, while focusing on the downtown area yields a higher success rate for partially blocked, closed and likely restricted road status.

Table 5.9: Summary of p-values on Impact of Mapping Focus Area (Geographical holdout 1 km for Downtown Carrefour)

Imputation Technique	Road Status Type	Percent of Missing Data	Mapping Focus Area	Interaction
Clustering Combined with Modified Mean-and-Mode	unrestricted	0.02	0.00	0.54
	partially blocked	0.01	0.00	0.18
	closed	0.60	0.22	0.68
	likely restricted	0.01	0.00	0.25
Clustering Combined with Adjacent Arc	unrestricted	0.00	0.00	0.36
	partially blocked	0.84	0.00	0.31
	closed	0.00	0.00	0.00
	likely restricted	0.04	0.00	0.16
Classification Tree	unrestricted	0.00	0.00	0.95
	partially blocked	0.01	0.00	0.01
	closed	0.00	0.00	0.00
	likely restricted	0.00	0.00	0.12

Observation: Our results suggest that a larger geographical area (using the entire Carrefour area) yields better estimation of the unknown information when higher levels of limited information of a certain attribute or outcome value types are present.

Insight: Above a certain threshold level, as the geographical focus area expands, the average success level for the imputation value categories with limited information increases.

Table 5.10: Impact of Mapping Focus Area (Random holdout 1 km for Downtown Carrefour)

Imputation Technique	Road Status Type	Percent of Missing Data	Mapping Focus Area	Interaction
Clustering Combined with Modified Mean-and-Mode	unrestricted	0.19	0.00	0.09
	partially blocked	0.03	0.00	0.50
	closed	0.10	0.29	0.29
	likely restricted	0.22	0.00	0.56
Clustering Combined with Adjacent Arc	unrestricted	0.00	0.00	0.74
	partially blocked	0.10	0.00	0.05
	closed	0.00	0.00	0.00
	likely restricted	0.31	0.00	0.58
Classification Tree	unrestricted	0.01	0.00	0.99
	partially blocked	0.00	0.00	0.01
	closed	0.00	0.00	0.01
	likely restricted	0.00	0.00	0.03

5.4.2.3 Changing Geographical Detail: Impact of Granularity

While the application of algorithms to the entire Carrefour area results in up to 87% average success for estimating majority of road status categories, it requires significant computational time due to geographical detail of one km grid size. In order to decrease the run time, we consider at coarser geography granularity or in other words, lower level of geographical information detail. Thus, we aggregate to a grid size of two km instead of one km. Larger grid size decreases the total number of grids, which in turn reduces the processing time necessary for integrating information prevalence in the geography attribute. Tables A4 - A6 show similar results for percent of missing data, missing data spread type and imputation method analogous to previous tests for downtown Carrefour and entire Carrefour area.

We observe up to 96% average success in estimating likely restricted road segments, which is a significant increase compared to the previous 87% average success. This success rate is obtained from 50 runs for 10% of data withheld using the CCMMM method. Random holdout and geographic holdout result in a similar success rate (only .01% difference). We should note that the high success in likely restricted road segments depends partially on how this status is calculated. The proposed methods have the least success in estimating closed road status, up to 22%, due to the low occurrence of this type in data (1.4%) and the high occurrence of unrestricted road status.

We further investigate the impact of grid size, one km versus two km, in Tables 5.11 and 5.12. The rows represent the imputation method type while columns 3-5 show the percent of missing data, grid size and their interaction, respectively. In Tables 5.11 and 5.12, we observe that selected grid size has an important impact on the average success rates for almost all cases, except when using clustering combined with modified mean-and-mode in random holdout scheme while estimating the closed road segments. When estimating the unrestricted road status, one km grid size yields significantly higher success rate than two km grid size for all imputation techniques under both holdout schemes. On the other hand, two km grid size is preferred for estimating the road segments with partially blocked, closed or likely restricted road status.

Observation: Our results suggest that lower granularity (using two km grid size) yields better estimates of the unknown information when limited information of a certain outcome value type is present.

Insight: Below a certain threshold level of available information for road status types, lower granularity yields better estimation of unknown information for the road status types with limited information.

One would expect that finer granularity yields better results, however we show that it is not the case when there is a limited number of a certain output type. This is true due to higher variance in the data.

Table 5.11: Impact of Grid Size (Geographical holdout 1 km to 2 km comparison)

Imputation Technique	Road Status Type	Percent of Missing Data	Grid Size	Interaction
Clustering Combined with Modified Mean-and-Mode	unrestricted	0.00	0.00	0.52
	partially blocked	0.00	0.00	0.35
	closed	0.00	0.00	0.01
	likely restricted	0.00	0.00	0.83
Clustering Combined with Adjacent Arc	unrestricted	0.02	0.00	0.20
	partially blocked	0.61	0.00	0.62
	closed	0.00	0.00	0.00
	likely restricted	0.03	0.00	0.31
Classification Tree	unrestricted	0.00	0.00	0.18
	partially blocked	0.00	0.00	0.04
	closed	0.78	0.00	0.60
	likely restricted	0.00	0.00	0.52

Table 5.12: Impact of Grid Size (Random holdout 1 km to 2 km comparison)

Imputation Technique	Road Status Type	Percent of Missing Data	Grid Size	Interaction
Clustering Combined with Modified Mean-and-Mode	unrestricted	0.19	0.00	0.09
	partially blocked	0.03	0.00	0.50
	closed	0.10	0.29	0.29
	likely restricted	0.22	0.00	0.56
Clustering Combined with Adjacent Arc	unrestricted	0.00	0.00	0.74
	partially blocked	0.10	0.00	0.74
	closed	0.00	0.00	0.00
	likely restricted	0.31	0.00	0.58
Classification Tree	unrestricted	0.01	0.00	0.99
	partially blocked	0.00	0.00	0.01
	closed	0.00	0.00	0.01
	likely restricted	0.00	0.00	0.03

When we have coarser granularity, on average 5.4 road segments are partially blocked and 2.7 are closed road per grid. When we have finer granularity, there are on average 1.7 partially blocked and 0.8 closed road segments per grid. As a result, there are some grids with only one closed road segment. When that road segment is selected to be held, there are not enough data points for closed road segments in the grid, thus reducing the success of imputation for closed road segments. A similar argument applies to partially blocked and in return likely restricted road statuses.

A good threshold for grid size to be used depends on the specific data, the amount of data available and the distribution of data (i.e., how much is allocated to each road status type). In our case study, we can say that the success of partially blocked and closed (and also likely restricted), which represent approximately 2.8% and 1.4% (and 4.2%) of data, increases as the granularity decreases. Thus, the threshold for this setting is above 4.2%. The default granularity in ArcGIS is one km grid size. We suggest that the researchers or the operations managers to run a sample of coarser granularity prior to making a decision.

5.5 Modeling and Policy Implications and Lessons Learned

5.5.1 Making the Most Out of Limited Data

While infrastructural data are critical for successful humanitarian operations management, it is often hard to obtain information about damaged roads [239, 228]. We exploit new data resources, such as OSM, and incorporate GIS to increase and utilize scarcely available road damage data after a disaster. Utilizing ArcGIS can save significant time in data collection and conversion efforts, which is key to successful post-disaster decision making.

When only 4.2% of data is available, rather than identifying each damage category explicitly, it is often better to identify for whether there is any type of damage or closure in the road segment with the aggregated category “likely restricted”. This information can help the deployment teams prepare for possible damage. If needed, the decision makers can then include reliability/error estimates on these forecasts in their logistical models or update imputation methods to tune their damage estimation. Some imputation methods can overestimate the damage (labeling the unrestricted roads as damaged) while others underestimate (labeling the damaged roads as unrestricted). In many settings, it is expected that overestimating damage is more preferable than underestimating damage due to high risk involved in humanitarian operations. This is analogous to estimating overage and underage costs in healthcare setting, where the cost of underage is much higher than the cost of overage in medical supplies demand estimation. Algorithms, especially CCMMM, can be tuned by changing the impact of the critical outcome status. This will serve as the first step in road status imputation. Hopefully, as more post-disaster damage data are collected, the imputation accuracy will be improved.

One additional approach to take advantage of the available information is to check for the granularity of the available information. Testing for granularity can improve the computational time while positively impacting the quality of the missing data estimation. This issue is more pressing when the level of available information for a specific data type decreases significantly. Future research should establish the generalizability and boundary conditions of these findings.

5.5.2 Applicability to Other Disasters

In this section, we discuss the applicability of this framework to other disasters. We utilize the measures developed in Chapter 4 to evaluate the quality of data and applicability to other disasters for logistical modeling. We use disaster properties, local factors and completeness of data from this study to assess the applicability of this framework.

In Chapter 4, we find that “time available for action (disaster onset), disaster size, magnitude of impact, duration of time and environmental factors (such as the topography or weather conditions of the affected area)” are good indicators of disaster properties. We suggest using Logistics Performance Indicator (LPI), geographical context of local area and disaster reoccurrence probability for local factors. In addition to disaster properties and local factors, we also include completeness measure, percent of available information, from the framework in Chapter 4 to characterize a disaster. As a minimum requirement, we need to have road status information and some of the attributes listed in this chapter if not all. We also believe that distribution of the damage information among the available information is key for successful implementation of the proposed methods.

Considering these factors, we believe that the proposed methods have potential to work better in the reoccurring disasters, disasters with advanced warning stages, and disasters in developed countries and urban areas due to availability of data. For example, hurricanes in Southeast US have higher occur-

rence probability and are well-studied by researchers [240]. Disasters with advanced warning such as hurricanes and typhoons have potential for more data during the disaster response due to better preparedness for mapping efforts, as it was in the case of Typhoon Haiyan. Disaster response maps need the pre-disaster conditions: transportation network and building information. Additionally, the majority of the attributes proposed such as flood plain, geology and administrative boundaries need GIS information. In general, developed countries such as the United States and Europe have a well mapped geography [241]. Thus, in addition to road status data, we also expect to find the necessary attribute data more in the US and Europe. Additionally, amount and quality of geographical information is higher in urban areas [242, 243].

As discussed in Chapter 4, LPI measures the “friendliness” of a country based on six factors: customs, infrastructure, services quality, timeliness, international shipments and tracking/tracing [244]. As LPI increases, the expected number of people impacted generally decreases [245]. Higher LPI scores indicate supply chain reliability, predictability, and timeliness [246]. Top 10 positions of the LPI ranking is occupied by high income countries, majority of them are from Europe, except Singapore, Hong Kong SAR, and United States. Thus, in addition to US and Europe, the framework has potential to work well in disasters based in these developed countries as well.

On the other hand, the proposed framework is less likely to produce successful results in areas where the given data do not meet minimum necessary requirements discussed above and due to geographical conditions, the network is disconnected. For example, in Cyclone Pam that hit South Pacific Ocean islands in 2015, there was no publicly available road damage data. In 2013, Typhoon Haiyan hit a large area in the Philippines. The framework might not work for the entire affected area since some islands were mapped less than others [247]. On the other hand, the framework can work better in selected areas of Philippines, such as Tacloban City, which had extensive post-disaster data. In the 2015 earthquake in Nepal, a large area was affected. Yet, there was limited coverage of estimated damage regarding road status. Moreover, the outcome of interest was the landslides, which were distributed far from each other. In this setting, we can not easily take advantage of adjacency and we can not utilize clustering combined with adjacent arcs method. Similarly, when there is extremely scarce data, classification trees might not work well since they depend on the availability of sufficient training data.

In Chapter 4, we also discuss the impact of technological status of the disaster-affected area on the available data. To the best of our knowledge and connections with practitioners as well as private industry, US has one of the leading disaster damage recording technology, such as 3D imaging and drones. In addition, there is an increasing effort in the US for better ways of identifying post-disaster damage information. One of our team members is currently leading the efforts in post-disaster damage identification in the US after disasters in collaboration with the local authorities and private companies. We believe that these and similar efforts are going to increase the availability of data in the near future.

5.5.3 Applicability Beyond Disaster Response and Humanitarian Transportation Planning

The framework presented in this paper is relevant to limited data environments that require actions before extensive data can be obtained. Such problems are encountered under highly constrained settings, such as funding, time, and manpower to name a few. In these cases, decisions must be made before sufficient amount of data are acquired to find an optimal solution.

Apart from finding the infrastructural information, this framework could be used in other mapping efforts, such as identifying the types of buildings. One of the common tasks in the Humanitarian OpenStreetMap Team (HOT) Tasking Manager, the platform where HOT reaches out to the volunteers for com-

pleting desired tasks, is the request to identify and tag building types [248]. Generally, the satellite images that volunteers base their work on are hard to read, and it takes significant time to identify and verify each task. A volunteer can utilize the framework presented here by finding attributes, such as number of houses in the neighborhood, soil type and type of road in the near distance, in order to identify the building types. Building types are important for disaster management and urban planning. In addition to disaster response, building type information can enhance preparedness and recovery in disaster management cycle, as well as emergency management such as fire protection and response. Similarly, building type information facilitates better urban planning and decision making for sustainable development [249].

5.6 Conclusion

We develop a framework to estimate the incomplete information in limited data environments. We identify a set of attributes based on the prior research, domain knowledge and availability of data. We demonstrate an application of this framework to the 2010 Haiti earthquake where the road status of unknown road segments amongst unrestricted, closed and partially blocked status is estimated. We define a new category of road damage status called “likely restricted” to better estimate the unknown information under the limited resources. Our results show up to 96% average success in estimating the unknown information. We investigate the impact of the percent of missing data, spread of missing data and imputation techniques used in approximating the incomplete information. Our findings suggest that more granularity is not always better for decision making and imply a threshold policy for granularity depending on the percent of missing data for certain road status types. We develop a general application, which includes a variety of attributes from flood plain to soil type so that it can be replicated for other disasters. We also investigate the applicability of this framework to other disasters.

In this study, we present a variety of imputation methods to emulate the current practices, such as global constant methods, as well as more sophisticated methods inspired by literature on missing data imputation. We also incorporate the nature of the problem to come up with new imputation methods: clustering combined with modified mean-and-mode method and clustering combined with adjacent arc.

Information is critical for any logistical operations. Through an in-depth study of humanitarian logistics models employing infrastructural information in humanitarian community, we highlight the need for publicly available data and data preservation. We found multiple instances where the available public information about the response to a disaster is “situational awareness via PDF” – a phrase often heard in the response community. We see a need in creating a common, public response and humanitarian data repository not only for final GIS products, but also the raw data that allowed the initial map creation. This will facilitate the international communities contributing to and analyzing the current digital response processes. In addition, access to private data can increase the success of the methods described here, as well as the disaster response process. Finally, the attribute selection also depends on the availability of data. For example, if the additional information such as building height is provided publicly, the algorithms might be improved by integrating it as another attribute.

The data from this study and all the test cases developed through different geography are publicly available at http://users.iems.northwestern.edu/kezban16/ArcGISModelandData/Paper_Submission.zip. The current process requires significant time for manual collection of data and converting data into workable format. Although employing a model in ArcGIS (also shared in the aforementioned link) to automate the process provides a substantial work and time reduction, this study points out the

need for publicly available data post-disaster data, especially in formats other than PDF maps.

Acknowledgement

This work has been in part funded by the National Science Foundation, Grant CMMI-1265786: “Advancing Dynamic Relief Response: Integration of New Data Streams and Routing Models” and accompanying REU supplements. We thank Kelsey Rydland of Northwestern University Library for technical support.

Appendix

ArcGIS Model Explanation

We use ArcGIS Desktop 10.x and ArcMap ModelBuilder to develop our model. Below we list the input layers used for this study and their corresponding type as well the steps used in building the model.

Input layers used in the model

- Damaged buildings - Points
- Damaged road segments - Lines
- Administrative areas - Polygon
- Floodplain - Polygon
- Surface geology - Polygon
- Epicenter - Point
- Data grids - Polygon

Model Steps

- Spatial join within a certain distances assigning number of buildings to the road segment
- Calculate density of damaged buildings per road segment
- Calculate number of damaged buildings within 25 meters
- Calculate percent of damaged buildings
- Add in administrative areas using polygon features.
- Assign distance to each segment from the earthquake epicenter (point feature) where each segment has distance from epicenter joined.
- Add to the segment if it is in the floodplain (polygon feature)
- Add the surface geology (polygon feature)
- Add centroids and assign the segment values to data grids (polygon) one at 1 km and one at 2 km

The resultant grids carry the road damage value and can be used for predictive modeling.

Case Study Attributes

Pre-disaster attributes

- Road type: We collect the road type data from OSM [174]. The possible categories are pathway, primary, primary link, tertiary, footway, residential, secondary, service, step, track, trunk, unclassified, road, pier, living street, and trunk link.
- Administrative boundaries: We use the administrative boundaries shapefile from OSM [174] for our case study region. This information is provided at the city level.
- Grid: In order to accompany the aggregate information provided by administrative boundaries, we break the study area into grids and use this for obtaining more detailed geographical information.
- Building density: Based on the OSM data, we calculate the building density for each road segment. Due to unavailability of building height data, we calculate the surrounding building density by dividing the total number of adjacent buildings by arc length.
- Land use: We consider land cover identification number data from OSM to study the impact of land use. The main categories include residential, industrial, and commercial and originate from OSM's internal classification, as accessed at the time of initial data retrieval from OSM.
- Flood plain: We include the comparative flood inundation probability for any given location in Haiti from Haiti Data, which provides Haiti-related geo-spatial information [250].
- Geology: We use the geology data from Haiti Data [250].

Post-disaster attributes

- Building damage level: For the building damage data, we use the joint damage assessment data from UNITAR/UNOSAT, World Bank and EC Joint Research Centre [251]. The damage assessment level ranges among five grades where grade 1 represents no visible damage and grade 5 represents destruction.
- Distance to epicenter: We find the epicenter of Haiti Earthquake data from UNITAR/UNOSAT, World Bank and EC Joint Research Centre [251]. We calculate the distance from each midpoint of the road segment to the epicenter.
- Information prevalence: We calculate the level of information available in the geography based on the available road status information per grid, as well as for the entire Carrefour.

Table A1: Summary of p-values on Impact of Percent of Missing Data Tests for Entire Carrefour

Imputation Technique	Road Status Type	Missing Data Spread Type	
		random	geographic
Clustering Combined with Modified Mean-and-Mode	unrestricted	0.01	0.01
	partially blocked	0.19	0.41
	closed	0.15	0.95
	likely restricted	0.22	0.03
Clustering Combined with Adjacent Arc	unrestricted	0.01	0.04
	partially blocked	0.54	0.11
	closed	0.36	0.38
	likely restricted	0.42	0.09
Classification Tree	unrestricted	0.00	0.00
	partially blocked	0.04	0.01
	closed	0.91	0.87
	likely restricted	0.03	0.01

Table A2: Summary of p-values on Impact of Missing Data Spread Type for Entire Carrefour

Imputation Technique	Road Status Type	Percent of Missing Data	Missing Data Spread Type	Interaction
Clustering Combined with Modified Mean-and-Mode	unrestricted	0.30	0.87	0.79
	partially blocked	0.03	0.77	0.64
	closed	0.20	0.46	0.30
	likely restricted	0.40	0.55	0.42
Clustering Combined with Adjacent Arc	unrestricted	0.00	0.31	0.34
	partially blocked	0.07	0.73	0.43
	closed	0.11	0.72	0.91
	likely restricted	0.30	0.87	0.79
Classification Tree	unrestricted	0.00	0.30	0.80
	partially blocked	0.00	0.50	0.00
	closed	0.96	0.48	0.81
	likely restricted	0.00	0.37	0.38

Table A3: Summary of p-values on Impact of Imputation Method Comparison for Entire Carrefour

Missing Data Spread Type	Road Status Type	Percent of Missing Data	Imputation Technique	Interaction
Random	unrestricted	0.01	0.00	1.00
	partially blocked	0.72	0.00	1.00
	closed	0.39	0.00	0.05
	likely restricted	0.79	0.00	1.00
Geographic	unrestricted	0.05	0.00	1.00
	partially blocked	0.45	0.00	1.00
	closed	0.88	0.00	0.99
	likely restricted	0.33	0.00	1.00

Table A4: Summary of p-values on Impact of Percent of Missing Data Tests for Entire Carrefour(2 km Grid)

Imputation Technique	Road Status Type	Missing Data Spread Type	
		random	geographic
Clustering Combined with Modified Mean-and-Mode	unrestricted	0.13	0.00
	partial blocked	0.02	0.00
	closed	0.00	0.00
	likely restricted	0.33	0.00
Clustering Combined with Adjacent Arc	unrestricted	0.00	0.00
	partially blocked	0.22	0.03
	closed	0.05	0.00
	likely restricted	0.02	0.00
Classification Tree	unrestricted	0.21	0.00
	partially blocked	0.00	0.00
	closed	0.60	0.67
	likely restricted	0.03	0.00

Table A5: Summary of p-values on Impact of Missing Data Spread Type for Entire Carrefour(2km Grid)

Imputation Technique	Road Status Type	Percent of Missing Data	Missing Data Spread Type	Interaction
Clustering Combined with Modified Mean-and-Mode	unrestricted	0.00	0.17	0.73
	partially blocked	0.00	0.84	0.79
	closed	0.00	0.57	0.44
	likely restricted	0.03	0.06	0.83
Clustering Combined with Adjacent Arc	unrestricted	0.00	0.68	0.70
	partially blocked	0.03	0.93	0.93
	closed	0.00	0.01	0.13
	likely restricted	0.00	0.40	0.73
Classification Tree	unrestricted	0.00	0.63	0.62
	partially blocked	0.00	0.27	0.57
	closed	0.50	0.03	0.78
	likely restricted	0.00	0.80	0.47

Table A6: Summary of p-values on Impact of Imputation Method Comparison for Entire Carrefour(2km)

Missing Data Spread Type	Road Status Type	Percent of Missing Data	Imputation Technique	Interaction
Random	unrestricted	0.00	0.00	1.00
	partially blocked	0.00	0.00	1.00
	closed	0.01	0.00	0.00
	likely restricted	0.05	0.00	1.00
Geographic	unrestricted	0.00	0.00	1.00
	partially blocked	0.00	0.00	1.00
	closed	0.15	0.00	0.00
	likely restricted	0.00	0.00	1.00

Chapter 6

Supply Chain Management in Limited Data Environments: An Application to Community Health

6.1 Introduction

At present, 5.9 million children under five years of age die per year worldwide, and the majority of these children die from preventable and treatable diseases [252]. The success of reducing this disease burden highly depends on the availability of treatment and prevention strategies for people at risk. Many low-income countries have limited access to simple, affordable and effective disease control measures due to resource constraints in areas such as healthcare work force and transportation [253].

In order to increase access to healthcare, community-based healthcare delivery models, which decentralize healthcare by moving care out of hospitals and into care in the community, have been implemented. Community-based programs have been shown to be effective in improving health outcomes in addition to extending healthcare delivery [254]. For this reason, many countries including Ethiopia, Malawi and Rwanda have undertaken national community healthcare programs [255, 256]. In this study, we collaborate with a nongovernmental organization (NGO) that aims to provide access to healthcare in remote areas through community-based healthcare in Liberia.

Community-based healthcare programs are executed with community health workers (CHWs). CHWs are lay people who live in the communities they serve and are the first point of contact for community members seeking medical care. Each CHW is responsible for a predetermined catchment area. The most common services that CHWs provide are reproductive health, maternal and neonatal health, and child health and nutrition. The CHWs perform integrated community case management (iCCM) for diarrhea, pneumonia and malaria for child health.

The success of the community health worker program highly depends on effective supply chain management. Our partner's supply chain starts with the procurement of internationally sourced medicinal supplies. These medicinal supplies move through a country depot and reach community health worker supervisors (CHWL), who are located at the supervisor hubs and distribute supplies to CHWs, via restock trucks. The supply chain ends with the distribution of service (care or medicine) from CHW to a patient. Often CHWs are provided with a backpack to carry supplies for critical services. Each CHW is given monthly stocks. When they are out-of-stock, they go to supervisor hubs for replenishment.

Medical supply stockouts are a significant concern as insufficient supplies can negatively impact services provided. There are two critical decisions that can impact stockout rates: 1) allocation decisions (i.e., which supplies to place in the backpack) and 2) restock policies (i.e., when to trigger a replenishment). Optimal allocation of supplies in the backpack and restock policies can decrease stockouts substantially and improve the CHW program performance. For example, from the operational management perspective, fewer stockouts lead to fewer transportation needs for the CHWs, which facilitates lower CHW program costs. From the health outcomes perspective, fewer stockouts enable better treatment, which increases the success of key performance indicators and the CHW impact.

Due to budget and space constraints in the backpack, a CHW can not stock unlimited supplies in the backpack. Therefore, the decision maker has to find a way to allocate different items in the backpack, which may likely cause understocking of items compared to a desired level or exclusion of some items from the backpack entirely. Changes in the composition and stocking levels of the backpacks can have significant consequences in the diagnosis of life threatening conditions, as well as long-term diseases. Thus, the decision maker has to carefully assess the demand, potential health benefit, and cost of each product in the backpack. Due to limited historical data, it is hard to employ a statistical forecasting model to estimate the demands. It is difficult to quantify the health benefits of medicinal items or services since intangible benefits are usually ignored [257]. Similarly, it is hard to measure the inventory cost of each product, especially the understocking cost, since a variety of outcomes can happen such as loss of patient's goodwill and death of a patient [258].

As in the work presented in earlier chapters, we use techniques from operations research in parallel with public health domain knowledge to tackle resource-limited environment challenges for inventory management models. In order to better understand the demand for each item, we first develop a demand and cost forecasting tool to carefully estimate the medical supply requirements and their corresponding costs. We propose alternative resources for obtaining data for the forecasting tool when historical consumption and/or cost data are limited. The forecasting tool assists in identifying the needs and costs of required commodities based on the established community health services guidelines [259].

Motivated by the backpack problem, in this study we investigate the impact of different items in the backpack. We develop a simple method to approximate cost measures for inventory management, such as underage and overage costs. This method employs population health measures; thus, it links cost to medical outcomes. To the best of our knowledge, this method is the first of its kind to provide direction on how to calculate shortage costs in healthcare, and it results in reasonable solutions for inventory management. Using the understanding of the costs, we develop models for stocking levels and more general inventory management policies.

The remainder of this chapter is structured as follows. In Section 2, we discuss the relevant literature on community health, supply chain management and public health. In Section 3, we present a series of models to represent the knapsack problem and solution methodology. Section 4 concludes the study with discussion and future directions.

6.2 Literature Review

6.2.1 Community Health and Data

During the past decade, there have been numerous studies analyzing the potential of CHW programs to improve health for populations where there is a limited healthcare workforce and limited access due to road conditions [256]. Similar to other healthcare programs, accurate program data plays a critical role

in informing decision-making, especially scaling up CHW programs in other areas. For example, data are needed to identify the gaps in areas where CHWs are needed most, the amount of resources that need to be allocated to these CHWs, etc. However, one of the biggest challenge in the development of CHW programs remains data availability and quality of data [260].

In order to overcome the data availability challenges, researchers put extensive efforts on data collection and processing. Mehta et al. [261] study fleet management to increase access to healthcare delivery in Zambia. Zambia's severe shortage of healthcare workers combined with poor transportation conditions (i.e. poor road conditions, unknown roads and limited budget for transportation) prevent healthcare workers from accessing rural areas. In order to assess the impact of different fleet management strategies, a group of researchers from Stanford University in collaboration with Riders for Health first map unmapped roads in Zambia. [261]. The research team combines GPS data (odometer readings) from Stanford tracked vehicles with other data sources, such as health worker surveys, to create color-coded heat maps. Von Achen et al. [262] also study the operational challenges such as limited resources and poor infrastructure in community healthcare networks in Liberia. In collaboration with Last Mile Health and Northwestern, researchers present extensive efforts to collect road network data (i.e. road conditions, road features, community locations) in the field through census workers and process these data to use in decision-making. Researchers describe the common data issues, such as disconnection between roads and repeated road segments, and potential solutions to create a detailed transportation network.

There are many concerns regarding the quality of the CHW collected data. Otieno et al. [263] observe that the reliability of the CHW collected data in Kenya changes from the reliability of technically trained team collected data depending on the different health indicator of interest. For example, while child health reports have over 10% disparities, maternal health reports have less than 1% difference. Admon et al. [264] observe significant data aggregation problems such as entry mistakes and calculation errors where chart data were migrated incorrectly in CHW reports in Malawi. Researchers observe that the majority of the areas where CHWs operate are operating with low quality, i.e., 70% or worse agreement between CHW and household reports. For example, Admon et al. [264] conclude that 41.8% of CHW reports in Malawi are poor quality for measuring the number of children aged < 5 years. Similarly, Mahmood and Ayub [265] find that only 47.5% of the reports administered by community based health workers in Pakistan are accurate while a significant percentage of the reports contain missed entries, misreported data and false information. Additionally, underreporting is also a common problem among CHWs [266].

Researchers employ quality control methods to assess and improve the quality of the CHW reports. Admon et al. [264] develop a quality assessment tool to identify areas with high and low quality reporting using CHW reports in Malawi. They also develop interventions such as reallocation of data aggregation responsibilities to address poor data quality. Mitsunaga et al. [267] design a data quality assessment system and integrate the system into the routine activities to respond to programmatic and training related quality problems in Rwanda's CHW program. Researchers suggest the use of data quality checklists for training and supervision.

In this study, we focus on the roll-out of a relatively new CHW program in Liberia. Since this is a new program, only two months of historical data were available for medicinal items for child health provided in the backpack at the time of the study. The two months of data include data from CHW reports and consumption data from the NGO reports. We observe significant discrepancies between these two reports. Thus, in this study we seek for alternative options for obtaining consumption data to develop forecasting strategies when there is limited data.

6.2.2 Supply Chain Management/ Inventory Management

The problem studied in this paper is an application of the knapsack problem where a set of items, each with a weight and value are considered to be allocated in the knapsack to maximize the total value of items while not exceeding the knapsack capacity [268]. Various stochastic knapsack problems, where either the item weights and/or the item values are random variables, have been studied. For more details, we refer the reader to [269, 270]. In the stochastic knapsack problem we are studying, the item weights and the values are known but the item demands are unknown a priori; thus, the total value and the total weights are unknown.

The class of stochastic knapsack problems considered in this paper is related to the newsvendor problem and its variants. In the single item newsvendor problem, a newsvendor sells a certain product whose demand is unknown but generally follows a certain probability distribution. The newsvendor needs to place an order of x units before observing the demand, considering the unit ordering cost, unit inventory holding cost and the unit stock-out cost. The newsvendor problem is solved by the critical ratio, which quantifies the relative cost of understocking to overstocking. For details on the newsvendor problem, we refer the reader to [271]. The unbounded knapsack with stochastic demand where there are no upper bounds on the number of copies of each kind of item in the knapsack can be viewed as classical, single period, multi-item newsvendor problem (MPNP) [272].

Several solutions have been proposed to solve the multi-product newsvendor problem (MPNP). While solving the unconstrained case, since the costs are separable for different items, the optimal inventory level for each item is found using the individual critical ratio [273]. The problem becomes challenging when there is a constraint, such as space or budget that links the products. One approach to solving this problem would be to start with the solution to the unconstrained multi-product newsvendor problem/knapsack problem and insert this solution into the constraint. If the constraint is not violated, then we have an optimal solution. Since we work in resource-constrained settings, this constraint is generally violated.

There have been multiple approaches to solve this problem. Hadley and Whitin [313] developed a Lagrange multiplier technique and dynamic programming procedure for finding the optimal order quantities. Nahmias and Schmidt [275] proposed several heuristic methods to solve the MPNP with a single constraint. Zhang et al. [276] developed a binary search method to obtain the optimal solution using marginal benefit function of each product. The authors note that the algorithm might only provide approximate solutions when the demand distribution is discrete and at the approximate solution, the budget constraint might be violated or there may be leftover budget. We modify Zhang's algorithm to overcome these challenges. We also utilize the methods developed by [275] in our study for comparison.

These supply chain models have been used in the industry for a long time. However, applying these models for medical supply chain in the developing world remains challenging [277]. In this study, we investigate the applications of these models in community health, an area where in addition to the historical consumption data, cost data is also hard to obtain.

6.2.3 Public Health and Health Economics

In this study, we explore other domains to learn about demand and cost information. In order to obtain the demand information, we need to understand the basic public health terms such as incidence and prevalence. Incidence refers to the new cases of a medical condition during a period in a specified population while prevalence indicates the number of people affected at a point in time [278, 279]. When

to use incidence or prevalence depends on what one knows, the type of condition, and purpose of the study [278, 280, 279].

In this study, we aim to link the operational outcomes such as stocking cost to medical outcomes. Thus, we also investigate public health and health economics literature. We examine the population health summary measures that quantify the health of a population by combining mortality and non-fatal health outcomes data into a single metric. The most widely used measure, Disability Adjusted Life Years (DALY), measures the health gap to the standard life of expectancy by summing the number of years of life lost due to premature mortality and years of life lived with a disability [281]. We utilize DALY averted to find the shortage cost for a medicinal item. We discuss details of the conversion in Section 3.2. We refer the interested readers to the following studies to examine the benefits and challenges of different population health measures [282, 283, 284, 285, 286].

6.3 Problem Setting and Modeling

In the unbounded knapsack problem studied here, there is a set of n medicinal items to be allocated to a backpack of capacity B . Capacity can be in terms of volume or budget. Each item i has a reward $r(i)$ and weight $w(i)$. Similarly, weight can be in terms of budget or cost per item. The demand for each item i is unknown and $D(i)$ represents the random variable for demand of item i . Let $F_i(d) = P[D(i) \leq d]$ be the distribution function of the demand. There is no upper bound on the number of copies of each kind of medicinal item. The objective is to determine the number of each item to include in the backpack, $x(i)$, to maximize the total value while staying in the capacity limit. We should note that in this setting there might be a collection of drugs needed to treat a disease yet we treat each collection as a single medicinal item. For example, diarrhea is suggested to be treated with a combination of oral rehydration salt sachet and zinc sulphate. For simplicity, we refer to these as “antidiarrheal drug” and we assume that the collection is one medicinal item.

As mentioned in the literature review, the unbounded knapsack problem with random demands can be studied as multi-product newsvendor problem. In the multi-product newsvendor problem, each item has an overage cost, $c_o(i)$, the cost incurred per item i for surplus at the end of the specified period and an underage cost, $c_u(i)$, for the cost incurred per item i for not being able to satisfy the demand.

The main challenge in this study is the limited availability of data, both historical consumption data and cost data. We develop a forecasting tool to estimate the demand per medicinal item under limited availability of consumption data. In particular, shortage costs are hard to obtain in healthcare [258]. Thus, we first start with a simpler case of the problem with one of the relatively simple items (i.e., items for which population health measure summaries exist) and then build on it. Here is the order of the cases we study:

- Case 1: 1 CHW 1 item
- Case 2: 1 CHW 2 item
- Case 3: 2 CHWs 1 item
- Case 4: 2 CHWs 2 item

Case 1 shows us how the stocking decisions are made when there is only one CHW who only stocks one medicinal item. The purpose of this case is to understand how the demand and different costs are

calculated given limited data. Next, using this case we study the impact of having additional CHWs and/or items in the backpack. In the next subsections, we discuss about these different cases.

6.3.1 1 CHW 1 item: Single Item Stocking Level

This is a classical single item newsvendor problem, which can be solved by using the critical ratio:

$$F_i(x(i)^*) = \frac{c_u(i)}{c_u(i) + c_o(i)}, i = 1, 2, \dots, n. \quad (6.39)$$

We include the index i to build for the upcoming cases. In this case, we assume unlimited budget. We next outline our steps for addressing limited data and application.

6.3.1.1 Demand Information

In this section, we explain the details of the forecasting tool developed to estimate the demand for various medicinal items and the data sources used in the process. We first determine the list of services to be provided by a CHW, medicinal items and required dosage information needed to perform these services using the general guidelines on community health, as well as previous community health applications [287, 288]. We should note that a service can be diagnosing and treating a disease such as malaria or providing supplies for educational programs like family planning.

After determining the list of services and the corresponding commodities, we next explore the data necessary for forecasting. We examine accurate data from sources including the NGO, Demographic Health Survey (DHS), which disseminates detailed national data on health and population, and World Health Organization (WHO) [287, 289]. Unfortunately, data on both historical medicinal item consumption and disease progressions are sparse.

Consequently, we seek alternative domains for gathering and estimating historical consumption data to develop forecasting strategies when there is limited data. For this reason, we investigate and amalgamate data from local resources, international data bases and reports, grey literature and academic literature. We utilize demographic data, morbidity data and program targets from these resources to estimate quantities of medicinal items needed. Demographic data include population level details regarding specific population target group residing in rural area per county (i.e., remote population). Examples of relevant populations can be number of children under five for iCCM and number of pregnant women for neonatal care. In order to reach these population numbers, demographic data include total population, percent of population residing in rural area per county, growth rate, distribution of male or female, percent of children under five, percent of pregnant women, etc. Morbidity data refers to disease or health condition prevalence or incidence rates per specific population target group [290]. Program targets include the coverage of the CHW per each county and the coverage of a service by CHWs. For example, the NGO targets to serve 50% of the remote population in the first year of CHW program application, 75% in the second year and 100% after third year of the CHW program application in the Gbarpolu county. Table 6.1 shows the types of resources used, examples, description of examples and data type for our forecasting tool.

Using these data, we formulate the average demand for item i per CHW per month:

Resource Type	Resource Type Description	Examples	Example Description	Data Type for Our Forecasting Tool Usage
Local Resources	Local resources include surveys and reports that provide detailed statistic on the population and its health, guidelines and targets for future operations.	Demographic Health Survey	This survey presents a wide range of country based population and health topics, such as population-based (age, sex, urban/rural) data on fertility, contraception, maternal and child health and nutrition.	demographic information, morbidity data
		Liberia - Core Welfare Indicators Questionnaire Survey 2010	The 2010 Core Welfare Indicators Questionnaire Survey cover household demographic characteristics, health, education and literacy, household assets and amenities, employment, access to amenities, household perceptions of well-being, subjective poverty, displacement and food	
		Ministry of Health and Social Welfare Republic of Liberia Annual Report	This report provides extensive information about multiple diseases, rates, percent of relevant population information, etc.	
International Databases and Reports	International resources offer reports and databases for country and region based statistics such as disease prevalence numbers for different regions of the world.	National Health and Social Welfare Policy	This policy covers a variety of statistics and goals related Liberian population, its health and welfare such as health and social welfare financing, infrastructure, human resources, basic package of health services, social welfare, pharmaceuticals and health commodities.	program targets
		WHO Resources (e.g., African Health Observatory)	AHO provides regional, sub-regional and country level statistical tables, factsheets, and overviews describing the health situation and trends on health outcomes, health systems, programmes.	
		UNICEF Global Databases United Nation (World Population Prospects)	UNICEF Global Databases present data collected through Multiple Indicator Cluster Surveys (MICS), an international household survey programme that is designed to collect data on more than 100 indicators related to children and women's education, health, gender equality, World Population Prospects demonstrates United Nations population estimates and projections.	
Grey Literature	Grey literature refers to the policies and guides developed by international organizations as well as program applications to other countries.	WHO Guidelines: The Community Health Worker	This document provides guidelines for training and adaptation of community health workers.	
		DELIVER Project Jarrah et al., Jarrah et al (2), The Ghana One Million Community Health Workers (1mCHW) Campaign	This guide is designed to assist users in applying a systematic, step-by-step approach to quantifying health commodity requirements and costs. These reports show the results of the costing analysis of an iCCM program and community health program applications to other countries in sub-Saharan Africa.	epidemiological factor for morbidity data
		Lukacik et al., Watson et al.	These papers provide epidemiological factors such as average duration for diarrhea and seasonality index for malaria.	epidemiological factor for morbidity data
Academic Literature	Academic literature presents published research on various parameters, such as epidemiological factors, cost of medicines as well as cost estimates from previous applications.	Walker et al., Rudan et al., Walker et al. (2010)	These studies provide global incidence rates and total population living with pneumonia and diarrhoea.	
		McCord et al., Oliphant et al., Nefdt et al. and Legesse et al.	These papers provide cost analysis of integrated community case management and greater community healthcare programs in sub-Saharan Africa.	cost

Table 6.1 Data Sources Used in the Forecasting Model

$$\bar{d}(i) = \kappa(i) \frac{\alpha \beta \theta(i) \phi(i)}{n_{CHW}}, i = 1, 2, \dots, n. \quad (6.40)$$

In this formulation, α and $\theta(i)$ are the data related to demographic data where α represents the remote population in the county studied and $\theta(i)$ represents the percent of relevant population for a medicinal item i . $\phi(i)$ refers to prevalence or incidence of a disease or health condition that is treated by medicinal item i . β is the target coverage of the CHW per each county per item. Additionally, we use the monthly coefficient of item, $\kappa(i)$, to find the average demand for CHW per month and the number of CHWs, n_{CHW} to find the average yearly demand for item i per CHW.

In the CHW program literature, demand for items with limited historical data are calculated as a function of prevalence and number of episodes [291], just incidence [292], or a combination of prevalence or incidence [293, 254] in addition to target coverage and target population. As discussed in the literature review, the use of incidence or prevalence should depend on the type of health condition and services provided. For example, for antimalarial drugs, we should use the incidence rate to estimate the demand for antimalarial drugs to serve all the new cases. For providing deworming, it might be better to use the prevalence of infestations [254]. Similarly, for providing antiretroviral treatment (ARV) for HIV, we should use the HIV prevalence rates to account for both new and old cases of HIV since HIV needs a continuing care. However, finding incidence rate estimates is hard. We benefit from the prevalence rates reported in 2013 DHS for Liberia for rural area to obtain incidence rates [289]. For each disease, the DHS reports the percentage of children who had symptoms in the two weeks prior to the survey. We use the methodology provided by Jarrah et al. [290] for their application in Senegal to convert these numbers to incidence rates. We first find the total number of two-week periods in a year including the duration of a disease episode. We then multiply this number with the prevalence rate stated in DHS and obtain the incidence rate per year.

We include monthly coefficient of item, $\kappa(i)$, to account for different characteristics of diseases and services provided. For example, while demand for family planning services is steady over the year, malaria is known to be changing from month to month [294]. In order to calculate average monthly consumption for seasonal items, using last six months, last 12 month, previous month and last year's consumption data for the same three months are suggested [294]. Watson et al. [295] develop look-ahead seasonality index (LSI) approach to adjust monthly consumption rates and present applications to Zambia, Zimbabwe, and Burkina Faso. Researchers also employ time series models to investigate seasonality patterns in the malaria incidence [296, 297, 298, 299]. On the other hand, due to unavailability of historical data, many researchers and guidelines only provide yearly estimates [300]. We use the NGO's guidelines as well as crude seasonality index approach developed by Watson et al. [295] to explore demand estimation.

Using the average demand per item i per CHW calculated in equation 6.40, we next find the appropriate distribution for selected medicinal items with the help of literature and historical data. The details are provided in the application section below, 6.3.1.3.

6.3.1.2 Cost Calculation

Cost calculations are a challenge when applying the newsvendor model to healthcare setting. Often, healthcare shortage costs are calculated subjectively [258]. We develop a formula to estimate shortage costs based on the value of the item i , $v(i)$, and the cost of the item, $c(i)$.

It is conventional to write the underage cost in terms of unit price sold and unit cost [301]. Rather than the price of item sold, we use the value of item. In the public health literature, in order to analyze the clinical and economical evaluation of healthcare, health state valuations (such as healthy, dead and ill) are practiced [302]. WHO recommends use of DALY to find global burden of a disease because the health state evaluations combine mortality and non-fatal outcomes for an intervention for DALY [303, 304]. Additionally, in order to compare the impact of different interventions, WHO values cost of a year of life. To find the value of item i , we multiply the DALY averted by using item i , $\gamma(i)$, with the cost of a year lost of life, c_l . We then deduct the cost of item i to find the underage cost. Equation 6.41 presents our shortage cost formulation.

$$c_u(i) = v(i) - c(i) = \gamma(i)c_l - c(i), i = 1, 2, \dots, n. \quad (6.41)$$

What makes this method attractive is that it utilizes population health measures; as a result, it links medical outcomes to operational outcomes. Moreover, the application Section 6.3.1.3 shows that it produces reasonable results. Note that we assume no additional cost such as loss-of-good-will in this method.

We next calculate the overage cost. In general, overage cost can be calculated in terms of holding cost $h(i)$, cost of item, $c(i)$ and salvage value, $s(i)$, as $c_o(i) = h(i) + c(i) - s(i)$ [305]. In this setting, we assume that the medicinal items can be salvaged fully. We find the overage cost using the holding cost per item which is equal to the percent of unit cost as carrying cost, h_p , multiplied by the cost of item. Equation 6.42 shows our overage cost calculation.

$$c_o(i) = h(i) = h_p c(i), i = 1, 2, \dots, n. \quad (6.42)$$

6.3.1.3 Application

We next present the application of the classical newsvendor model for the case of the NGO. Even with the proposed shortage cost calculation method, it might be hard to calculate the shortage cost for certain items due to unavailability of the underlying value of the item data, $v(i)$. In order to overcome the data challenge, we investigate the items in the backpack and categorize them as relatively easy or hard items based on their availability of value of item data. Commodities for iCCM; malaria, diarrhea and pneumonia are grouped as relatively easy items since there are studies regarding population health measures for these items in public health literature. On the other hand, it is hard to find data that estimate the value of nutrition, maternal health and family planning items. Thus, we group them as hard items. For this application, we focus on the top two diseases with highest estimated demands: malaria and diarrhea.

We first find the stock levels for antimalarial drugs. In order to calculate the critical ratio for antimalarial drugs, $F_1(x(1)^*)$, we first estimate the underage and overage costs, $c_u(1)$ and $c_o(1)$. Arrow et al. [306] state that each malaria treatment has .262 DALYs as an effect per person. According to WHO, a year lost of life is worth 150\$/DALY [307]. By multiplying these two, we find the value of antimalarial drugs, $v(1)$ as \$39.3. The cost of antimalarial drugs to treat one case of malaria, $c(1)$ is \$1.38 according to our partner. Then, by subtracting this cost from the value of the item, we find the shortage cost of antimalarial drug, $c_u(1)$ as \$37.92. According to WHO, 30-40% is a reasonable percent of holding cost for medicinal items, h_p [308]. Using 35% as an average percent of unit cost for inventory carrying cost, we find the overage cost for an antimalarial drug, $c_o(1)$ as \$.483. Plugging $c_u(1)$ and $c_o(1)$ into the equation

(6.39), we find the critical ratio, $F_1(x(1)^*)$, to be .988 for antimalarial drugs.

We now estimate the demand distribution for antimalarial drugs. Using the forecasting formula in equation (6.40), we find the average demand for antimalarial drugs per CHW per month as 9. Note that after discussions with our partner, as an initial step, we use the average monthly coefficient as (1/12) for all items because they want to have higher stock levels overall to account for changing road conditions for different seasons. In other words, the average demand is same for each month. Leung et al. [309] model the demand for antimalarial drugs with lognormal distribution with 50% coefficient of variation (CV) for calculating stockouts in Zambia. The CV for lognormal distribution is calculated by:

$$CV[X] = \sqrt{e^{\sigma^2} - 1} \quad (6.43)$$

Solving for this, we find the standard deviation, σ , to be .47. The solution to newsvendor problem under the lognormal distribution can be found by:

$$x^* = F^{-1}(c_u / (c_u + c_o)) = e^{\mu + (Z^{-1}(c_u / (c_u + c_o)) * \sigma)} \quad (6.44)$$

The mean for the lognormal distribution is $e^{\mu + \frac{\sigma^2}{2}}$. Notice here rather than the mean, we use the mean divided by $e^{\frac{\sigma^2}{2}}$. Plugging the mean and the standard deviation into equation (6.44), we find the optimal stocking level for antimalarial drugs, $x^*(1)$, to be 24. Note that this optimal stocking level is much higher than the average demand per month due to log normal distribution. However, it is not much different than the stocking level used by our partner.

Our partner NGO uses the following stocking level policy for their items. They assume that the stocking levels equate to a 99% probability of always having stock available for each medicinal item plus a month's average monthly consumption (AMC). This is designed to overcome challenges with bad roads, broken bridges and problems with motorbikes, seasonality, etc. The NGO assumes average monthly consumption rate of 7 for antimalarial drugs and uses a stocking level of 20. This is slightly lower than the stocking level we obtain using critical ratio. The NGO's average stockout report shows 1-3 stockouts for antimalarial products per month for Rivercess county. This shows that our stocking level of 24 is a good estimate for the NGO operations.

We also employ the crude seasonality indices provided by Watson et al. [295] to adjust AMC for seasonality under no historical data. Researchers provide an example of using coefficient of 1 in low season and coefficient of 2.5 for the peak season. This indicates that the consumption in the peak season is 2.5 times of the consumption in the low season. Watson et al. [295] also assumes that peak season lasts for half of the year. Using these indices, we adjust the monthly coefficient of item, $\kappa(i)$, accordingly and find the average demand per month per low season as 5 and the average demand per month per peak season as 12. Plugging these numbers in the equation 6.44, we find the stocking level for antimalarial drugs for low and peak season as 13 and 30. We should note that these estimates should be approached with caution and seen as a comparison since the peak season might change based on the geography and country.

We now find the stocking levels for antidiarrheal drugs. Using the equations (6.40)-(6.42), we find the average demand rate for antidiarrheal drugs as 23, underage cost as \$ 4.62, and overage cost as \$.116 [310]. We should note that the item cost for diarrhea treatment is \$.33. As a result, the critical ration for the antidiarrheal drugs is $c_u / (c_u + c_o) = .977$. We next look at the demand distribution of anti-diarrhoeal

drugs. A study in Brazil detects that the diarrhea related deaths and hospitalization are seasonal [311]. Xu et al. [312] also discover that the demand for diarrhea commodities were seasonal. Researchers state that diarrhea in children < 5 years appeared to peak in fall-winter seasons, which is in accord with findings in Brazil [311], and this may partially be attributable to the fact that rotavirus is the predominant aetiology of diarrhea in infants and young children and rotavirus favors low temperature [312]. Additionally, according to our partner NGO activities, diarrhea follows a similar trend to malaria. Thus, we choose to use log normal distribution with 50% CV for antidiarrheal drug demand as well [310]. Then, we find the optimal stocking level to be around 53 for antidiarrheal drugs. According to our partner NGO, the average monthly consumption is 10 and the stocking levels seem to be 30 for antidiarrheal drugs. These numbers are very different than our calculated levels. If you look at the average demand calculated based on equation 6.40 and the NGO's crude assumption, there is almost two fold difference. According to DHS, in the rural areas 63.1% of the diarrhea cases were administered antidiarrheal drugs [289]. On the other hand, our model assumes 100% coverage for all diarrhea cases. If we use the target numbers from DHS to calculate the amount of antidiarrheal drugs, then we observe the mean to be around 15 cases and the top off level should be 33 cases which is similar to the NGO's stocking level.

Applying similar crude seasonality indices, we find the average demand for antidiarrheal drugs to 13 for the low season and 33 for the peak season. Consecutively, we find the stocking levels as 29 and 72 for low and peak season, respectively.

6.3.2 Extensions

In this section, we discuss the extensions to the single item single CHW case. We also include budget constraint in these extensions.

6.3.2.1 1 CHW 2 items

We first study the case where there is an additional item in the knapsack. In this case, the problem becomes multi-product newsvendor problem (MPNP). It can be formulated as:

$$\min \sum_{i \in I} \max((x(i) - D(i), 0)c_o(i) + \max(D(i) - x(i), 0)c_u(i)) \quad (6.45)$$

s.t.

$$\sum_{i \in I} x_i w_i \leq B \quad (6.46)$$

$$x(i) \geq 0 \forall i \in I \quad (6.47)$$

While solving the unconstrained case, since the costs are separable for different items, the optimal order quantity can be found by Turken et al. [273]:

$$F_i(x^*(i)) = \frac{c_u(i)}{c_u(i) + c_o(i)} \quad (6.48)$$

In other words, each item is set to its optimal value using the critical ratio. Researchers suggest multiple strategies for addressing the constrained problem [273, 276, 313, 275]. When the budget is not binding, each item is set to unconstrained solution value. Nahmias and Schmit[275] propose a simple

heuristic, which scales down the solution of the unconstrained problem, $x(i)^{unc}$, to constrained problem by using a scale, k to fit the budget and find the constrained solution, $x(i)^{cons}$ as $x(i)^* = kx(i)^{unc}$, when the budget is binding. Alternatively, Zhang et al. [276] develop a binary search method that uses marginal benefit of each item to find the optimal solution. In order to accommodate for the under budget or over budget solutions, we use the Zhang's algorithm for obtaining initial results and then use a search algorithm to add and/or drop items in the *Modified Zhang's Algorithm*. The details of the Modified Zhang's Algorithm are:

- Use Zhang's algorithm and find the budget difference.
- If the budget difference is zero, output the solution.
- If the budget difference is less than zero, find the item with the minimum critical ratio and reduce the number of items one by one until the difference is positive.
- If the budget difference is greater than zero, find the item with the maximum critical ratio. While the budget difference is greater than the minimum cost item, check to add item with the highest critical ratio if there is enough budget to cover the cost of that item. If there is not enough budget, move on the list of critical ratio order until an item can fit in the budget.

We now apply these two methods to our case study. We want to stock two items: anti-malarial drugs and antidiarrheal drugs. Let's assume that the budget is \$40. The unconstrained solution is found as $x(1) = 24$ and $x(2) = 53$, with a required budget of \$50.61, which is over budget. Using Nahmias and Schmidt's heuristic, we find $k = .79$ and the corresponding constrained solution as $x(1) = 19$ and $x(2) = 41$. We obtain the same solution with the modified Zhang's algorithm. We observe that these values are approximately scaled down version of the unconstrained case in proportion with the available budget. The solution is rounded around the scaled numbers based on the available budget.

6.3.2.2 2 CHW 1 item

We are particularly interested in this case to see how the solution changes when multiple CHWs work together. This is especially important to build the ground work for future work such as assessing pooling strategies later. We assume that the underage and overage costs to be the same for each CHW; however, the demand consequently the order quantity might change per CHW. Thus, we define $D(j)$ as demand for CHW $j \in J$ and $x(j)$ as the quantity ordered for CHW $j \in J$. The formulation changes to:

$$\min \sum_{j \in J} E[\max(x(j) - D(j), 0)]c_o + E[\max(D(j) - x(j), 0)]c_u \quad (6.49)$$

s.t.

$$\sum_{j \in J} x(j)w \leq B \quad (6.50)$$

$$x(j) \geq 0 \forall j \in J \quad (6.51)$$

Currently, our partner assumes that each CHW serves a population of 350 people. They also assume that demand for a medicinal item is independent of the region in which the CHW serves. Thus, in this

setting, we assume that the demand for the item is independent and identically distributed (i.i.d.). In reality, this might not be true because some areas have more children population so the incidence rates of childhood diseases are expected to be higher. Additionally, in some areas the population is more dispersed and a CHW spends more time in walking. This leads to a lower ratio. However, currently there is not enough data to analyze these factors.

The optimal solution for this problem is same as the 1 CHW 1 item for both CHWs assuming we are operating under the proportional budget. When the budget is not proportional, the allocation for both CHWs are same, just the number is different than the 1 CHW 1 item case. This is due to the i.i.d. assumption where we assume that each region in which the CHW serves has same demand. If this assumption did not hold, then the stock should be allocated based on the demand of the item per CHW. For example, if a CHW serving a community has higher demand for an item compared to another CHW serving another community, s/he should be allocated more stock for that item.

6.3.2.3 n CHWs m items

We now study n CHWs m items case to motivate for multiple CHWs multiple items for the larger scale. We assume the underage and overage costs to be same for item i for different CHWs. However, the demand and consecutively stock levels can change per CHW. Thus, we define $D(i, j)$ as demand of an item $i \in I$ for CHW $j \in J$ and $x(i, j)$ as the quantity ordered for item $i \in I$ by CHW $j \in J$. The formulation changes to:

$$\min \sum_{i \in I} \sum_{j \in J} E[\max(x(i, j) - D(i, j), 0)] c_o(i) + E[\max(D(i, j) - x(i, j), 0)] c_u(i) \quad (6.52)$$

s.t.

$$\sum_{i \in I} \sum_{j \in J} x(i, j) w(i) \leq B \quad (6.53)$$

$$x(i, j) \geq 0 \forall i \in I, j \in J \quad (6.54)$$

For our setting, similar to 2 CHWs 1 item case, we assume the demand to be i.i.d. Assuming the budget is proportional to the 1 CHW two items case, the solution should be same as 1 CHW two items. The allocation should be same for all CHWs. If the i.i.d. assumption does not hold, similarly to 2 CHWs 1 item case, the stock should be allocated among CHWs based on the demand for the item i per CHW j .

6.4 Discussion and Future Work

In this chapter, we study supply chain management in limited data environments with an application to community health. We work in collaboration with an NGO providing access to healthcare in the world's most remote areas. They operate through CHWs who carry a backpack with medicine to serve their community. Due to limited access to training and limited budget, the country specific data about CHWs are hard to obtain. However, organizations need to make stocking decisions to serve the population. This chapter provides steps on how to tackle scarce data. We investigate the impact of different medicinal items and backpack composition when there is limited historical data, budget and costing data.

Our contributions in this chapter are following. We develop a forecasting tool to estimate the demand and cost of medicinal items under limited data using morbidity data. We combine alternative resources to estimate the demand information and provide guidance of different resources for future applications. We also develop a simple methodology to approximate underage and overage costs for healthcare supply chain. We use DALY averted with treatment and cost of life of a year to find the value of an item and deduct the cost of item to find the shortage cost. To best of our knowledge, this is the first method in the literature to address the shortage costs in healthcare supply chain. Using the results from the costing tool and the supply chain cost estimation methodology, we calculate stocking levels for medicinal items. Application to our partner shows that the methodology provides reasonable results.

These resources and methods provided in this study can not only be used for community health but for other healthcare applications. For example, if there is need for calculating the demand of a new medicinal item provided in hospitals or items that do not have reliable historical data, practitioners can use the suggested demand estimation tool with adjustment to appropriate population targeted for the medicinal item. Similarly, the shortage cost calculation developed here can be used for other medicinal items in addition to the drugs to serve community health patients.

The current application of monthly coefficient for average monthly demand calculations needs to be tailored for disease characteristics. Although current formula accounts for some of the disease characteristics with the crude seasonal indices, these indices are generic and might need to be updated for Liberia. As additional county based consumption data are gathered, the forecasting model can be tuned using bootstrapping and fitting the data to find seasonality indices. This would alleviate the dependence to the lognormal distribution assumption as well. These data are also valuable to understand the dependence among counties for different diseases. Another possible method to tackle this problem is to utilize our knowledge about infectious disease modeling from Chapter 2 and 3 and develop an epidemic model of the diseases studied here. Then, we can use the output of these epidemic models in the resource allocation problem. This is plausible because the items in the backpack do not interact much so we can have a separate model for each disease.

As mentioned before, the stocking levels should be approached cautionary considering the seasonality and availability of infrastructure to support restocking in peak seasons. For this reason, lead time should be incorporated and stocks for peak season should be allocated much in advance compared to low season.

The methods developed provides the ground work for other applications, such as resource pooling. Currently, the NGO assumes that each CHW is responsible for their own community. Due to the resource and space constraints, CHWs might not always carry all items. In the future, we plan to investigate the impact of different chain configurations for transshipment between CHWs. Due to the nature of our CHW setting, full flexibility is not viable. Thus, CHWs can pool with their neighboring CHWs. We also aim to study different characteristics of the system such as correlations.

Acknowledgement

This work has been in part funded by Northwestern University Transportation Center Fellowship and Royal E. Cabell Tearminal Year Fellowship.

Chapter 7

Conclusion

This thesis presents industrial engineering and management sciences applications to public services. We used a variety of methods to explore the role of operations research in a variety of settings with different levels of data availability. In the public health domain, we used ordinary differential equations and stochastic simulation to study impact of two diseases: human immunodeficiency virus (HIV) and measles. We investigate the population level impact of switching away from progestogen-only injectable hormonal contraceptives (POIHCs) on births and HIV spread: prevalence, new infections, HIV-related deaths, and vertical transmission in sub-Saharan Africa. We show that if a significant portion of the POIHC users undertook other forms of contraceptions, prevalence, new infections and HIV-related deaths decrease and the vertical transmission changes slightly. Additionally, the increase in births is limited compared to the case where all POIHC users stop. In a second setting, we study the impact of community connectivity for the case of measles analytically as well as numerically. We show that community structure had little impact on the results and the infection process converges to a finite limit as the population size increases. We develop a simulation model to assess these impacts numerically. The simulation showed that population size does not have a significant impact after a few hundred. We also investigate the impact of new public health interventions such as better case detection and faster response and show that better case detection can improve outbreak outcomes significantly.

In the humanitarian logistics domain, we develop frameworks to analyze and impute humanitarian logistics data. in quantity and quality. Using Typhoon Haiyan as a case study we study the implications of real-time collection for humanitarian logistics models. We first hand collect near real-time data and develop a framework to analyze data after disasters. We perform logistical content analysis to quantify amount of data. To best of our knowledge, this is the study focuses on humanitarian logistics researchers' perspective on data collection and provide estimates and examples on the logistical data. Our analysis from multiple disasters show that only 5-8% of damage information is available after disasters while the rest is missing. Inspired by this analysis, we employ statistical and machine learning techniques to impute the missing information using similarities in the attributes. We develop an ArcGIS model to automatize the data collection and processing efforts to the extent possible and we are in the process of providing the model and sharing it publicly to reduce the data collection and processing for future disasters. We provide an application of these to a past disaster, 2010 Haiti Earthquake, to estimate the missing road network information. The results show high level success of identifying damage in the damage in the infrastructure. We also discuss the applicability of these methods to other disasters.

In the final chapter, we combine our experience in public health, humanitarian operations and limited data environment to address supply chain management in limited data for community health in

collaboration with a non-governmental organization (NGO). The NGO trains and deploys community health workers to increase access to healthcare and reduce the mortality and morbidity rate of common diseases in the community. Each CHW carries a backpack equipped with medicinal items while serving the community. We develop a forecasting model to estimate the demand and cost required for different medicinal items to serve the community by utilizing alternative resources under limited historical consumption data. We also build a simple methodology to estimate shortage and overage costs for medicinal items. To best of our knowledge, this is the first formulation of shortage cost estimation in medicinal items and this method links the population health measures to the cost. Using these methods, we determine the stocking levels for different medicinal items for CHWs to effectively serve their communities. This work can be useful not only in other community health settings but also other healthcare settings. Additionally, this work provides ground work for resource sharing in community health settings.

Available data informs modelers about real life conditions. Much of the operations research models assume availability of data. On the other hand, especially in the public health and humanitarian logistics we see that data can be limited. We present different techniques to address missing data such as using assumptions, forecasting, imputation and linking medical outcomes to costs. We hope that these different methods can inform modelers and practitioner for decision making as well as data collection and processing. As the available information increases, operations researchers can potentially aid public officials and on-field teams better.

Bibliography

- [1] Heffron R, Donnell D, Rees H, Celum C, Mugo N, Were E, et al. Use of hormonal contraceptives and risk of HIV-1 transmission: a prospective cohort study. *Lancet Infect Dis* 2012;12:19-26.
- [2] Pelluck P. The New York Times. 2012; 3 October. Contraceptive used in Africa may double risk of H.I.V. Available from: <http://www.nytimes.com/2011/10/04/health/04hiv.html?pagewanted=all> [accessed 05.10.11].
- [3] United Nations, DESA Population 2011 World Contraceptive Use. Available from: <http://www.un.org/esa/population/publications/wcu2010/Main.html> [accessed 12.12.11]
- [4] CDC (2014) Measles - United States, January 1 - May 23, 2014 *MMWR*. 63(22):496-499.
- [5] Agboghroma CO. Contraception in the context of HIV/AIDS: a review. *Afr J Reprod Health* 2011;15(3):15-24.
- [6] Plummer FA, Simonsen JN, Cameron DW, Ndinya-Achola JO, Kreiss JK, Gakinya MN, et al. Co-factors in male-female sexual transmission of human immunodeficiency virus type 1. *J Infect Dis* 1991;163:233e9.
- [7] Baeten JM, Benki S, Chohan V, Lavreys L, McClelland RS, Mandaliya K, et al. Hormonal contraceptive use, herpes simplex virus infection, and risk of HIV-1 acquisition among Kenyan women. *AIDS* 2007;21:1771e7.
- [8] Morrison CS, Turner AN, Jones LB. Highly effective contraception and acquisition of HIV and other sexually transmitted infections. *Best Pract Res Clin Obstet Gynaecol* 2009;23:263e84.
- [9] Morrison CS, Chen PL, Kwok C, Richardson BA, Chipato T, Mugerwa R, et al. Hormonal contraception and HIV acquisition: reanalysis using marginal structural modeling. *AIDS* 2010;24:1778e81.
- [10] Watson-Jones D, Baisley K, Weiss HA, Tanton C, Changalucha J, Everett D, et al. Risk factors for HIV incidence in women participating in an HSV suppressive treatment trial in Tanzania. *AIDS* 2009;23:415e22.
- [11] Lavreys L, Baeten JM, Kreiss JK, Richardson BA, Chohan BH, Hassan W, et al. Injectable contraceptive use and genital ulcer disease during the early phase of HIV-1 infection increase plasma virus load in women. *J Infect Dis* 2004;189: 303e11.
- [12] Stringer EM, Kaseba C, Levy J, Sinkala M, Goldenberg RL, Chi BH, et al. A randomized trial of the intrauterine contraceptive device vs hormonal contraception in women who are infected with the human immunodeficiency virus. *Am J Obstet Gynecol* 2007;197(144):e1e8.

- [13] Morrison CS, Skoler-Karpoff S, Kwok C, Chen PL, van de Wijgert J, Ehret- Plagianos MG, et al. Hormonal contraception and the risk of HIV acquisition among women in South Africa. *AIDS* 2012;26:497e504.
- [14] Myer L, Denny L, Wright T, Kuhn L. Prospective study of hormonal contraception and women's risk of HIV infection in South Africa. *Int J Epidemiol* 2007;36:166e74.
- [15] Hel Z, Stringer E, Mestecky J. Sex steroid hormones, hormonal contraception, and the immunobiology of human immunodeficiency virus-1 infection. *Endocr Rev* 2010;31:79e97.
- [16] Morrison C, et al. Presented at AIDS 2014. Melbourne, Australia: Jul 20e25, 2014. Hormonal contraception and HIV infection: results from a large individual participant data meta-analysis.
- [17] Ralph LJ, McCoy SI, Shiu K, Padian NS. Hormonal contraceptive use and women's risk of HIV acquisition: a meta-analysis of observational studies. *Lancet Infect Dis* 2015;15:181e9.
- [18] WHO. Medical eligibility criteria for contraceptive use. 4th ed. A WHO Family Planning Cornerstone; 2009 Available from: http://whqlibdoc.who.int/publications/2010/9789241563888_eng.pdf [accessed at 16.02.12].
- [19] WHO. Women need access to dual protection effective contraceptives and HIV prevention options. 16 February 2012. Available from: <http://www.unaids.org/en/resources/presscentre/pressreleaseandstatementarchive/2012/february/20120216pshormonal/> [accessed 20.02.12].
- [20] WHO. Hormonal Contraception and HIV. Technical Statement. Available from: http://whqlibdoc.who.int/hq/2012/WHO_RHR_12.08_eng.pdf [accessed at 20.02.12].
- [21] Padian NS, Buve A, Balkus J, Serwadda D, Cates W. Biomedical interventions to prevent HIV infection: evidence, challenges, and way forward. *Lancet* 2008;372:585e99.
- [22] Abdool Karim Q, Sibeko S, Baxter C. Preventing HIV infection in women: a global health imperative. *Clin Infect Dis* 2010;50(S3):S122e9.
- [23] Hogan MC, Foreman KJ, Naghavi M, Ahn SY, Wang M, Makela SM, et al. Maternal mortality for 181 countries, 1980e2008: a systematic analysis of progress towards millennium development goal 5. *Lancet* 2010;375(9726): 1609e23.
- [24] Sales JM, Whiteman MK, Kottke MJ, Madden T, DiClemente RJ. Dual Protection use to prevent STIs and unintended pregnancy. *Infect Dis Obstetrics Gynecol* Volume 2012, Article ID 972689, 2 pp.
- [25] UNAIDS. Guidelines for Behavior Change Interventions to Prevent HIV: Sharing Lessons from an Experience in Bangladesh Based on the Application of Lessons from Sonagachi, Kolkata. Available from: <http://www.hivpolicy.org/Library/HPP001312.pdf> [accessed 20.12.14].
- [26] The Global HIV Prevention Working Group. Behavior change and HIV prevention: (Re)Considerations for the 21st Century. Available from: http://www.malecircumcision.org/advocacy/documents/PWG_behavior_report_FINAL.pdf [accessed 20.12.2014].
- [27] Haddad LB, Polis CB, Sheth AN, Brown J, Kourti AP, King C, et al. Contraceptive methods and risk of HIV acquisition or female-to-male transmission. *Curr HIV/AIDS Rep* 2014 December;11(4):447e58.

- [28] UNAIDS. Report on the global AIDS epidemics. 2010. Available from: http://www.unaids.org/en/media/unaids/contentassets/documents/unaidspublication/2010/20101123_globalreport_en.pdf [accessed 09.01.12].
- [29] Glynn JR, Sonnenberg P, Nelson G, Bester A, Shearer S, Murray J. Survival from the HIV-1 seroconversion in Southern Africa: a retrospective cohort study in nearly 2000 gold miners over 10 years follow up. *AIDS* 2007;21:625e32.
- [30] Mills EJ, Bakanda C, Birungi J, Chan K, Ford N, Cooper CL, et al. Life expectancy of persons receiving combination antiretroviral therapy in low-income countries: a cohort analysis from Uganda. *Ann Intern Med* 2011;155:209e16.
- [31] World Bank. Available from: <http://www.worldbank.org/depweb/english/modules/social/pgr/datasubs.html> [accessed 20.02.12].
- [32] Boily MC, Baggaley RF, Wang L, Masse B, White RG, Hayes RJ, et al. Heterosexual risk of HIV-1 infection per sexual act: systematic review and metaanalysis of observational studies. *Lancet Infect Dis* 2009;9:118e29.
- [33] CIA World Factbook 2011. Available from: <https://www.cia.gov/library/publications/the-world-factbook/rankorder/2054rank.html> [accessed at 16.01.12].
- [34] Abdool Karim Q, Abdul Karim SS, Frohlich JA, Grobler AC, Baxter C, Mansoor LE, et al. Effectiveness and safety of tenofovir gel, an antiretroviral microbicide, for the prevention of HIV infection in women. *Science* 2010;329: 1168e74.
- [35] Weller S, Davis K. Condom effectiveness in reducing heterosexual HIV transmission. *Cochrane Database Syst Rev* 2002;1:CD003255.
- [36] PEPFAR Prevention of mother-to-child transmission of HIV: expert panel report and recommendations to the U.S. Congress and U.S. Global AIDS coordinator. 2010. Available from: <http://www.pepfar.gov/documents/organization/135465.pdf> [accessed at 16.01.12].
- [37] Soderlund N, Zwi K, Kinghorn A, Gray G. Prevention of vertical transmission of HIC: analysis of cost effectiveness options available in South Africa. *BMJ* 1999;318:1651e6.
- [38] UNAIDS. UNAIDS report on the global AIDS epidemic. 2013. HIV estimates with uncertainty bounds. Available from: www.unaids.org/en/media/unaids/contentassets/documents/epidemiology/2013/gr2013/GR2013_HIV_Estimates_AnnexTable.xls.
- [39] Blumberg, S., Enaroria, W.T.A., Lyoyd-Smith, J.O, Lietman, T.M., Porco, T.C. (2014) Identifying Postelimination Trends for the Introduction and Transmissibility of Measles in the United States. *American Journal of Epidemiology*, 79(11):1375-82
- [40] Health Protection Agency (HPA). (2012) Laboratory confirmed cases of measles, mumps and rubella in England and Wales: update to end June 2012. *Health Protection Report*. 6(34). 24 August 2012. Available from: <http://www.hpa.org.uk/hpr/archives/2012/hpr3412.pdf>.
- [41] Fiebelkorn, A.P, Redd, S.B., Gallagher, K., Rota, P.A., Rota, J., Bellini, W., Seward, J. (2010) Measles in the United States during the postelimination era. *Journal of Infectious Diseases*, 202:1520-8.

- [42] New York City Department of Health and Mental Hygiene (NYCDHMH). (2013) ALERT 9: Measles in New York City. Available from [https://a816-health29ssl.nyc.gov/sites/NYCHAN/Lists/AlertUp-date AdvisoryDocuments/HAN_Measles_2013-04-12.pdf](https://a816-health29ssl.nyc.gov/sites/NYCHAN/Lists/AlertUp-date%20AdvisoryDocuments/HAN_Measles_2013-04-12.pdf) Accessed at August 12, 2013.
- [43] CDC (2016) Measles Cases. Available from <http://www.cdc.gov/measles/cases-outbreaks.html> Accessed at 12 April 2016.
- [44] Glass, K., Kappey, J. and Grenfell, B.T. (2004) The Effect of Heterogeneity in Measles Vaccination on Population Immunity. *Epidemiology and Infections*, 132:675-683.
- [45] Bartlett, M.S. (1957) Measles Periodicity and Community Size. *Journal of Royal Statistical Society A*, 120(1):48-70.
- [46] Bartlett, M.S. (1960) The Critical Community Size for Measles in the United States. *Journal of Royal Statistical Society A*, 123:37-44.
- [47] Black, F. (1966) Measles Endemicity in Insular Populations: Critical Community size and Its Evolutionary Implication. *Journal of Theoretical Biology*, 11(2):207-211.
- [48] Keeling, M.J. and Grenfell, B.T. (1997) Disease Extinction and Community Size: Modeling the Persistence of Measles. *Science*, 275:65-67.
- [49] Anderson, R.M., May, R.M. (1992) *Infectious diseases in humans. Dynamics and control*. Oxford: Oxford University Press.
- [50] Murray, G.D. and Cliff, A.D. (1977) A Stochastic Model for Measles Epidemics in a Multi-Region Setting. *Transactions of the Institute of British Geographers, New Series*, 2(2):158-174.
- [51] Bolker, B. and Grenfell B.T. (1995) Space, persistence and dynamics of measles epidemics. *Philosophical Transactions: Biological Sciences*, 348:309-320.
- [52] Hethcote, H.W. and Van Ark, J. W. (1987) Epidemiological Models for Heterogeneous Populations: Proportionate Mixing, Parameter Estimation and Immunization Programs. *Mathematical Biosciences*, 75:205-227.
- [53] Castillo-Chavez, C., Hethcote, H.W., Andreasen, V., Levin, S.A., Liu, W.M. (1989) Epidemiological Models with Age Structure, Proportionate Mixing, and Cross-immunity. *Journal of Mathematical Biology*, 27:233-258.
- [54] Hethcote, H.W. Modeling heterogeneous mixing in infectious disease dynamics. V. Isham, G. Medley (Eds.), *Models for Infectious Human Diseases: Their Structure and Relation to Data*, Cambridge University Press, Cambridge, UK (1996), pp. 215-238.
- [55] Llyod, A.L. and May, R.M. (1996) Spatial Heterogeneity in Epidemic Models. *Journal of Theoretical Biology*, 179:1-11.
- [56] Allen, L.J.S., Jones, M.A., Martin C.F. (1991) A discrete-time model with vaccination for a measles epidemic. *Mathematical Biosciences*, 105:111-131.

- [57] Bonačić Marinović, A.A., Swaan, C., Wichmann, O., Steenbergen, J.V., Kretzschmar, M. (2009) Effectiveness and timing of vaccination during school measles outbreak. *Emerging Infectious Diseases*, 18(9):1405-1413.
- [58] Ejima, K., Omori, R., Aihara, K., Nishiura, H. (2012) Real-time investigation of measles epidemics with estimate of vaccine efficacy. *International Journal of Biological Sciences*, 8:620.
- [59] Babad, H.R., Nokes, D.J., Gay, N.J., Miller, E., Morgan-Capner, P., Anderson, R.M. (1995) Predicting the impact of measles vaccination in England and Wales: Model validation and analysis of policy options. *Epidemiology and Infection*, 114:319-344.
- [60] CDC. (2012) Measles - United States, August 24, 2012. *MMWR*. 61:647-652.
- [61] Public Health Agency of Canada (2013) Measles. Available from <http://www.phac-aspc.gc.ca/im/vpd-mev/measles-rougeole-eng.php> Accessed at January 12, 2013.
- [62] Perez, L. and Dragicevic, S. (2009) An Agent-based Approach for Modeling Dynamics of Contagious Disease Spread. *International Journal of Health Geographics*, 8:50.
- [63] Wallinga, J., Heijne, J.C.M., Kretzschmar, M. (2005) A Measles Epidemic Threshold in a Highly Vaccinated Population. *PLoS Med*, 2(11):e316.
- [64] Kutty, P., Rota, J., Bellini, W. (2013). Manual for the Surveillance of Vaccine Preventable Diseases, 6th Edition. In: CDC (editor). Chapter 7 Measles Available at <http://www.cdc.gov/vaccines/pubs/survmanual/chpt07-measles.html> Accessed at December 12, 2015.
- [65] Altay, N., & Labonte, M. (2014). Challenges in humanitarian information management and exchange: Evidence from Haiti. *Disasters*, 38(s1), S50-S72.
- [66] Ergun, Ö., Karakus, G., Keskinocak, P., Swann, J., & Villarreal, M. (2010). Operations research to improve disaster supply chain management. *Wiley Encyclopedia of Operations Research and Management Science*, John Wiley & Sons, Hoboken, NJ.
- [67] Kovács, G., & Spens, K. M. (2007). Humanitarian logistics in disaster relief operations. *International Journal of Physical Distribution & Logistics Management*, 37(2), 99-114.
- [68] De la Torre, L.E., Dolinskaya, I. S., & Smilowitz, K. R. (2012). Disaster relief routing: Integrating research and practice. *Socioeconomic Planning Sciences*, 46(1), 88-97.
- [69] Sangiamkul, E., & Hillegersberg, J. van. (2011, May 8-11). Research directions in information systems for humanitarian logistics. 8th International Conference on Information Systems for Crisis Response and Management, ISCRAM, Lisbon, Portugal.
- [70] Sheu, J. B. (2010). Dynamic relief-demand management for emergency logistics operations under large-scale disasters. *Transportation Research Part E: Logistics and Transportation Review*, 46(1), 1-17.
- [71] Yi, W., & Özdamar, L. (2007). A dynamic logistics coordination model for evacuation and support in disaster response activities. *European Journal of Operational Research*, 179(3),

- [72] Ortuño, M. T., Cristóbal, P., Ferrer, J. M., Martín-Campo, F. J., Muñoz, S., Tirado, G., & Vitoriano, B. (2013). Decision aid models and systems for humanitarian logistics. A survey. In: Vitoriano B, Montero J, Ruan D (eds) Decision aid models for disaster management and emergencies (pp. 17-44). Atlantis Press, Springer Business + Science Media, New York, NY
- [73] Özdamar, L., & Ertem, M. A. (2015). Models, solutions and enabling technologies in humanitarian logistics. *European Journal of Operational Research*, 244(1), 55-65.
- [74] Huang, A., Ma, A., Schmidt, S., Xu, N., Zhang, B., Meineke, L., Shi, Z. E., Chan, J. & Dolinskaya, I. (2013). Integration of real time data in urban search and rescue center for the commercialization of innovative transportation technology. Transportation Center, Northwestern University. Retrieved from http://www.ccitt.northwestern.edu/documents/Integration_of_real_time_data_in_urban_search_and_rescue.pdf
- [75] United Nations Foundation. (2011). Disaster relief 2.0: The future of information sharing in humanitarian emergency. Retrieved from <http://www.unfoundation.org/what-we-do/legacy-of-impact/technology/disaster-report.html>
- [76] Crisis Mappers. (2014). Home. Retrieved from <http://crisismappers.net/>
- [77] Grünewald, F., & Binder, A. (2010) Inter-agency real time evaluation in Haiti: 3 Months after the earthquake. Plaisians, France: Groupe Urgence Développement and Global Public Policy Institute.
- [78] United Nations Office of Coordination of Humanitarian Affairs. (2014e). Who we are. Retrieved from <http://www.unocha.org/about-us/who-we-are>
- [79] United Nations Office of Coordination of Humanitarian Affairs. (2014d). What is UNDAC? Retrieved from <http://www.unocha.org/what-we-do/coordination-tools/undac/overview>
- [80] United Nations Office of Coordination of Humanitarian Affairs. (2014c). Information management. Retrieved from <http://www.unocha.org/what-we-do/information-management/overview>
- [81] National Disaster Risk Reduction and Management Council. (2014). NDRRMC home page. Retrieved from <http://www.ndrrmc.gov.ph>
- [82] Presidential Management Staff, Presidential Communications Development & Strategic Planning Office. (2014, April 24). Infographic: Timeline of government actions in response to Typhoon Haiyan. Official Gazette. Retrieved from <http://www.gov.ph/2014/04/24/infographic-timeline-of-government-actions-in-response-to-typhoon-haiyan/>
- [83] Digital Humanitarian Network (2014). DHNetwork about. Retrieved from <http://digitalhumanitarians.com/about>
- [84] Butler, D. (2013). Crowdsourcing goes mainstream in typhoon response. *Nature*, 20. Retrieved from <http://www.nature.com/index.html>
- [85] Open Street Map (2014c). Typhoon Haiyan. Retrieved from http://wiki.osm.org/wiki/Typhoon_Haiyan
- [86] Liu, N., & Ye, Y. (2014). Humanitarian logistics planning for natural disaster response with Bayesian information updates. *Journal of Industrial and Management Optimization*, 10(3), 665-689.

- [87] The National Oceanic and Atmospheric Administration. (2013, November 21). NWS Director Uccellini speaks on CNN about Typhoon Haiyan and climate change. Retrieved from http://www.nws.noaa.gov/com/weatherreadynation/news/131121_louis.html
- [88] Bogazici University, Earthquake Eng. Dept. (2002). Earthquake risk analysis of Istanbul metropolitan area (in Turkish). Bogazici Universitesi Kandilli Rasathanesi ve Deprem Arastirma Enstitusu, Istanbul, Turkey.
- [89] Endsley, M. R. (1988, October). Design and evaluation for situation awareness enhancement. Proceedings of the human factors and ergonomics society annual meeting (Vol. 32, No. 2, pp. 97-101). SAGE Publications Santa Monica, CA.
- [90] Limbu, M., Wang, D., Kauppinen, T., & Ortmann, J. (2012). Management of a crisis (MOAC) vocabulary specification. Retrieved from <http://observedchange.com/moac/ns>
- [91] World Food Programme. (2014a). About. Retrieved from <http://www.wfp.org/about>
- [92] World Food Programme. (2014b). Food aid information system. Retrieved from <http://www.wfp.org/fais/>
- [93] Logistics Cluster. About the Logistics Cluster. Retrieved from <http://www.logcluster.org/logistics-cluster>
- [94] Federal Emergency Management Agency. (2010). The Federal Emergency Management Agency. Retrieved from <https://www.fema.gov/pdf/about/pub1.pdf>
- [95] Federal Emergency Management Agency. (2014a). Data feeds. Retrieved from: <http://www.fema.gov/data-feeds> on 12 January 2014.
- [96] Federal Emergency Management Agency. (2014b). FEMA's International Programs & Activities. Retrieved from <http://www.fema.gov/femas-international-programs-activities>
- [97] Inter-Agency Standing Committee. (2010a). IASC guidelines Common Operational Datasets (CODs) in disaster preparedness and response. Retrieved from https://www.humanitarianresponse.info/system/files/documents/files/iasc_guidelines_on_common_operational_datasets_in_disaster_preparedness_and_response_2010-11-01.pdf
- [98] Inter-Agency Standing Committee. (2010b). IASC strategy, meeting humanitarian challenges in urban areas. Retrieved from <https://interagencystandingcommittee.org/meeting-humanitarian-challenges-urban-areas/documents-public/iasc-strategy-meeting-humanitarian>
- [99] Inter-Agency Standing Committee. (2012). Reference module for cluster coordination at the country level. Retrieved from https://www.humanitarianresponse.info/system/files/documents/files/iasc-coordination-referencemodule-en_0.pdf.
- [100] International Federation of Red Cross and Red Crescent Societies. (2013). World disasters report. Retrieved from <http://www.ifrc.org/en/publications-and-reports/world-disastersreport/world-disasters-report-2013/>

- [101] Cordeiro, K. de F., Campos, M. L. M., & Borges, M. R. da S. (2014). Adaptive integration of information supporting decision making: A case on humanitarian logistic. Proceedings of ISCRAM 2014., Pennsylvania, USA.
- [102] Howden, M. (2009), How humanitarian logistics information systems can improve humanitarian supply chains: A view from the field. Presented at 6th International ISCRAM Conference, Gothenburg, Sweden.
- [103] Scott, N., & Batchelor, S. (2013). Real time monitoring in disasters. IDS Bulletin, 44(2), 122-134.
- [104] Li, J., Li, Q., Liu, C., Khan, S. U., & Ghani, N. (2014). Communitybased collaborative information system for emergency management. Computers & Operations Research, 42, 116-124.
- [105] Dorasamy, M., Ramen, M., & Kaliannan, M. (2013). Knowledge management systems in support of disasters management: A two decade review. Technological Forecasting and Social Change, 80(9), 1834-1853.
- [106] Palen, L., Anderson, K. M., Mark, G., Martin, J., Sicker, D., Palmer, M., & Grunwald, D. (2010). A vision for technology-mediated support for public participation and assistance in mass emergencies and disasters. Proceedings of ACM-BCS Visions of Computer Science 2010, Swinton, UK. In: Proceedings of the 2010 ACM-BCS Visions of Computer Science Conference (Edinburgh, United Kingdom, April 14-16, 2010). ACM-BCS Visions of Computer Science, British Computer Society, Swinton, UK, 1-12.
- [107] Ashktorab, Z., Brown, C., Nandi, M., & Culotta, A. (2014). Tweedr: Mining twitter to inform disaster response. Proceedings of ISCRAM 2014, Pennsylvania, USA.
- [108] Hester, V., Shaw, A., & Biewald, L. (2010, December). Scalable crisis relief: Crowdsourced SMS translation and categorization with Mission 4636. Proceedings of the first ACM Symposium on Computing for Development (p. 15), ACM, London, United Kingdom.
- [109] Imran, M., Elbassuoni, S. M., Castillo, C., Diaz, F., & Meier, P. (2013). Extracting information nuggets from disaster-related messages in social media. Proceedings of ISCRAM 2013, Baden-Baden, Germany.
- [110] Manso, M., & Manso, B. (2012). The role of social media in crisis: A European holistic approach to the adoption of online and mobile communications in crisis response and search and rescue efforts. Proceedings of the 17th International Command & Control Research & Technology Symposium, Virginia, USA. Retrieved from http://www.dodccrp.org/events/17th_iccrts_2012/papers/007.pdf
- [111] Munro, R. (2013). Crowdsourcing and the crisis-affected community. Info Retrieval, 16(2), 210-266.
- [112] Ortmann, J., Limbu, M., Wang, D., & Kauppinen, T. (2011, October). Crowdsourcing linked open data for disaster management. Proceedings of the Terra Cognita Workshop on Foundations, Technologies and Applications of the Geospatial Web in conjunction with the ISWC (pp. 11-22), Bonn, Germany.
- [113] Purohit, H., Castillo, C., Diaz, F., Sheth, A., & Meier, P. (2013). Emergency-relief coordination on social media: Automatically matching resource requests and offers. First Monday, 19(1).

- [114] Sarcevic, A., Palen, L., White, J., Starbird, K., Bagdouri, M., & Anderson, K. (2012, February). Beacons of hope in decentralized coordination: Learning from on-the-ground medical twitterers during the 2010 Haiti earthquake. *Proceedings of the ACM 2012 Conference on Computer Supported Cooperative Work* (pp. 47-56), ACM, New York, USA.
- [115] Velev, D., & Zlateva, P (2012). Use of social media in natural disaster management. *International Proceedings of Economic Development and Research*, 39, 41-45.
- [116] DeLone, W. H., & McLean, E. R. (1992). Information systems success: The quest for the dependent variable. *Information Systems Research*, 3(1), 60-95.
- [117] DeLone, W. H., & McLean, E. R. (2003). The DeLone and McLean model of information systems success: A ten-year update. *Journal of Management Information Systems*, 19(4), 9-30.
- [118] Petter, S., DeLone, W., & McLean, E. (2008). Measuring information systems success: Models, dimensions, measures, and interrelationships. *European Journal of Information Systems*, 17(3), 236-263.
- [119] Bharosa, N., Appelman, J. A., Zanten, B. Van, & Zuurmond, A. (2009, May 10-13). Identifying and confirming information and system quality requirements for multi-agency disaster management. *Proceedings of the 6th International Conference on Information Systems for Crisis Response and Management*, Gothenborg, Sweden, ISCRAM.
- [120] Haggarty, A., & Naidoo, S. (2008). Global symposium+5 final report, information for humanitarian action. UNOCHA Global Symposium Secretariat. Palais des Nations 1211, Geneva, Switzerland.
- [121] United Nations Office of Coordination of Humanitarian Affairs. (2013c). Humanitarian data exchange. Retrieved from [http:// docs.hdx.rwlab.org/wp-content/uploads/HDX-Project-Document-abridged.pdf](http://docs.hdx.rwlab.org/wp-content/uploads/HDX-Project-Document-abridged.pdf)
- [122] United Nations Office of Coordination of Humanitarian Affairs. (2013a). Developing humanitarian data standards: An introduction and plan for 2014. Retrieved from http://docs.hdx.rwlab.org/wp-content/uploads/HXL_Paper-forsite.pdf
- [123] United Nations Office of Coordination of Humanitarian Affairs. (2014b). HDX quality assurance framework. Retrieved from http://docs.hdx.rwlab.org/wp-content/uploads/HDX_Quality_Assurance_Framework_Draft.pdf
- [124] Day, J. M., Junglas, I., & Silva, L. (2009). Information flow impediments in disaster relief supply chains. *Journal of the Association for Information Systems*, 10(8), 637-660.
- [125] Geodata Converter. (2014). Vector. Retrieved from <http://converter.mygeodata.eu/vector>
- [126] Tatham, P., & Spens, K. (2011). Towards a humanitarian logistics knowledge management system. *Disaster Prevention and Management*, 20(1), 6-26.
- [127] United Nations Office for the Coordination of Humanitarian Affairs. (2013b). Governmental contact list: Region VI (Western Visayas). Retrieved from <https://philippines.humanitarianresponse.info/document/government-contact-list-region-vi-western-visayas>

- [128] Teran, J. (2014). Measuring the quality of humanitarian data: An emerging framework. Retrieved from <http://docs.hdx.rwllabs.org/measuring-the-quality-of-humanitarian-data-anemerging-framework/#comments>
- [129] United Nations Office for Coordination of Humanitarian Affairs. (2002). Symposium on best practices in humanitarian information exchange. Palais des Nations, Geneva, Switzerland.
- [130] Tatham, P., L'Hermitte, C., Spens, K., & Kovács, G. (2013). Humanitarian logistics: Development of an improved disaster classification framework. The 11th ANZAM Operations, Supply Chain and Services Management Symposium (pp. 1-10), Brisbane, QLD, Australia.
- [131] Arvis, J. E., Saslavsky, D., Ojala, L., Shepherd, B., & Busch, C. (2014). Connecting to compete 2014, trade logistics in the global economy: The logistics performance index and its indicators. Retrieved from <http://lpi.worldbank.org/>
- [132] L'Hermitte, C., Bowles, M., & Tatham, P. (2013). A new classification model of disasters based on their logistics implications. 11th ANZAM Operations, Supply Chain and Services Management Symposium (pp. 1-19), Brisbane, QLD, Australia.
- [133] Vaillancourt, A. (2013). Government decentralization and disaster impact, an exploratory study. Retrieved from <http://www.buildresilience.org/2013/proceedings/files/papers/352.pdf>
- [134] Haavisto, I. (2014). Performance in humanitarian supply chains. *Ekonomi och Samhälle/Economics and Society*, No. 275, Hanken School of Economics, Helsinki, Finland.
- [135] Ergun, Ö., Stamm, J. L. H., Keskinocak, P., & Swann, J. L. (2010). Waffle House Restaurants hurricane response: A case study. *International Journal of Production Economics*, 126(1), 111-120.
- [136] Galton, A., & Worboys, M. (2011, May). An ontology of information for emergency management. *Proceedings of 8th International Conference on Information Systems for Crisis Response and Management*, Lisbon, Portugal.
- [137] Jahre, M., & Jensen, L. M. (2010). Coordination in humanitarian logistics through clusters. *International Journal of Physical Distribution & Logistics Management*, 40(8/9), 657-674.
- [138] United Nations Office of Coordination of Humanitarian Affairs. (2013d, November 8). Philippines: Super Typhoon Haiyan makes landfall, UN and humanitarian community on high alert. Retrieved from <http://www.unocha.org/roap/top-stories/philippines-super-typhoon-haiyan-makes-landfall-unand-humanitarian-community-high-alerton>
- [139] Digital Humanitarian Network. (2014b). Super Typhoon Haiyan. Retrieved from <http://digitalhumanitarians.com/content/super-typhoon-yolanda>
- [140] Humanitarian Response. (2014c). Humanitarian Response Philippines. Retrieved from <https://philippines.humanitarianresponse.info/search/>
- [141] Protection Cluster (2013). Protection Cluster displacement and 3W map. Retrieved from <http://www.humanitarianresponse.info/operations/philippines/infographic/protection-clusterdisplacement-and-3w-map>

- [142] World Health Organization. (2013). Foreign medical teams in Tacloban City. Retrieved from <http://www.humanitarianresponse.info/operations/philippines/infographic/foreignmedical-teams-tacloban-city>
- [143] MapAction. (2013a). MA029 Philippines Typhoon Haiyan (Yolanda) 3W Overview (21-Nov-2013). Retrieved from <http://www.humanitarianresponse.info/operations/philippines/infographic/ma029-philippines-typhoon-haiyan-yolanda-3w-overview-21-nov-2013>
- [144] Palatino, M. (2013, November 13). Lessons from the Haiyan Typhoon Tragedy. The Diplomat. Retrieved from <http://thediplomat.com/2013/11/lessons-from-the-haiyan-typhoon-tragedy/>
- [145] United Nations Office of Coordination of Humanitarian Affairs. (2014f). Philippines: Typhoon Haiyan Situation Report No. 30. Retrieved from https://reliefweb.int/sites/reliefweb.int/files/resources/OCHAPhilippinesTyphoonHaiyan-SitrepNo30_06January2014.pdf
- [146] World Food Programme Office of Evaluation, United Nations Children's Fund Evaluation Office, & the Ministry of Foreign Affairs, Policy and Operations Evaluation Department, Netherlands. (2012). Joint evaluation of the global logistics cluster. Retrieved from <http://documents.wfp.org/stellent/groups/public/documents/reports/wfp251775.pdf>
- [147] Logistics Information about In-Kind. (2014a). About LogIK. Retrieved from <http://logik.unocha.org/SitePages/about.aspx>
- [148] Logistics Information about In-Kind. (2014b). Logistics Information about In-Kind relief detailed report. Retrieved from <http://logik.unocha.org/SitePages/DetailedReportAdmin.aspx?soid=2,3>
- [149] Open Street Map. Humanitarian OSM Team. Retrieved from http://wiki.openstreetmap.org/wiki/Humanitarian_OSM_Team
- [150] Open Street Map. (2014a). Index of Haiyan. Retrieved from <http://labs.geofabrik.de/haiyan/>
- [151] Open Street Map. (2014b). Some editing stats from the Typhoon Haiyan response. Retrieved from <http://hot.openstreetmap.org/updates/2014-01-14-some-editing-stats-from-the-typhoon-haiyan-response>
- [152] Humanitarian Response. (2014a). About Humanitarian Response. Retrieved from <https://philippines.humanitarianresponse.info/about>
- [153] MapAction. (2014). MapAction home page. Retrieved from <http://www.mapaction.org/>
- [154] MapAction. (2013b). Typhoon Haiyan update. Retrieved from <http://www.mapaction.org/more-news/385-typhoon-haiyan-update.html>
- [155] Copernicus Emergency Management System. (2014b). What is Copernicus. Retrieved from <http://emergency.copernicus.eu/mapping/ems/what-copernicus>
- [156] Copernicus Emergency Management System. (2014a). EMSR058: Typhoon in Philippines. Retrieved from <http://emergency.copernicus.eu/mapping/list-of-components/EMSR058>

- [157] Digital Humanitarian Network. (2014a). Esri disaster response program. Retrieved from <http://digitalhumanitarians.com/content/esri-disaster-response-program>
- [158] Environmental Systems Research Institute. (2014a). Cyclone Haiyan impact (MapServer). Retrieved from http://fema-dev.esri.com/arcgis/rest/services/DisasterResponse/Cyclone_Haiyan_Impact/MapServer
- [159] Environmental Systems Research Institute. (2014b). Typhoon Haiyan/Yolanda maps. Retrieved from <http://www.esri.com/services/disaster-response>
- [160] United Nations Institute for Training and Research. (2013). UNOSAT. Retrieved from <http://www.unitar.org/unosat/>
- [161] Volontaires Internationaux en Soutien aux Opérations Virtuelles. (2014a). About us. Retrieved from <https://haiyan.crowdmap.com/page/index/1>
- [162] Volontaires Internationaux en Soutien aux Opérations Virtuelles. (2014c). Retrieved from <http://visov.org>
- [163] Department of Social Welfare and Development. (2014). Reports and updates. Retrieved from <http://disaster.dswd.gov.ph/reports-and-updates/>
- [164] Humanitarian Response. (2014b). Humanitarian Response COD-FOD Registry Philippines. Retrieved from https://cod.humanitarianresponse.info/search/field_country_region/164?search_api_views_fulltext=
- [165] ReliefWeb. (2014). ReliefWeb about page. Retrieved from <http://reliefweb.int/about>
- [166] All Partners Access Network. (2013). Typhoon Haiyan response 2013. Retrieved from <https://community.apan.org/typhoonhaiyan/p/map.aspx>
- [167] All Partners Access Network. (2014). Typhoon Haiyan community. Retrieved from <https://community.apan.org/typhoonhaiyan/b/updates/archive/2013/11/09/typhoon-haiyan-community.aspx>
- [168] Red Cross. (2014). Typhoon Haiyan mapfolio. Retrieved from http://americanredcross.github.io/haiyan_mapfolio/
- [169] Google Crisis Maps. (2014). Typhoon Yolanda relief map. Retrieved from <http://google.org/crisismap/2013-yolanda>
- [170] Jiang, Y., Yuan, Y., Huang, K., and Zhao, L. (2012). Logistics for Large-Scale Disaster Response: Achievements and Challenges. 45th Hawaii International Conference on System Sciences (HICSS), 2012, pp. 1277 - 1285.
- [171] Argollo, S., Bandeira, R., and Campos, V. (2013). Operations Research in Humanitarian Logistics Decisions. 13th World Conference on Transportation Research, July 15-18, 2013, Rio de Janeiro, Brazil.
- [172] Van Wassenhove, L. N. (2006) Humanitarian aid logistics: supply chain management in high gear. *Journal of the Operational Research Society*, 57(5), 475-489.

- [173] United Nations Cartographic Section. (2010). Haiti: Port-au-Prince Damage Assessment as of 13/01/2010. Retrieved from <http://reliefweb.int/sites/reliefweb.int/files/resources/F3A04278B9989ED1C12576B00036CD87-map.pdf>
- [174] Open Street Map. Humanitarian OSM Team. Retrieved from http://wiki.openstreetmap.org/wiki/Humanitarian_OSM_Team
- [175] Environmental Systems Research Institute. (2016). ModelBuilder Tutorial. Retrieved from <https://www.arcgis.com/>
- [176] Chen, L. and Miller-Hooks, E. (2012). Optimal team deployment in urban search and rescue. *Transportation Research Part B*, 46(8), 984-999.
- [177] Liberatore, F., Pizarro, C., Blas, C., Ortuño, M., and Vitoriano, B. (2013). Uncertainty in humanitarian logistics for disaster management. A review, in: B. Vitoriano, J. Montero, D. Ruan (Eds.), *Decision aid models for disaster management and emergencies*, Atlantis computational intelligence systems, Vol. 7, Atlantis Press (2013), pp. 45-74.
- [178] Shen, Z., Dessouky, M. M., and Ordoñez, F. (2009). A two-stage vehicle routing model for large-scale bioterrorism emergencies. *Networks*, 54(4), 255-269.
- [179] Huang, Y., Fan, Y., and Cheu, R. L. (2007). Optimal Allocation of Multiple Emergency Service Resources for Protection of Critical Transportation Infrastructure. *Transportation Research Board*, 2022, pp. 1-8.
- [180] Mete, H. O. and Zabinsky, Z. B. (2010). Stochastic optimization of medical supply location and distribution in disaster management. *International Journal of Production Economics*, 126(1), 76-84.
- [181] Hentenryck, P. V., Bent, R., and Coffrin, C. (2010). Strategic Planning for Disaster Recovery with Stochastic Last Mile Distribution. *Proceedings of the Seventh International Conference on Integration of Artificial Intelligence and Operations Research Techniques in Constraint Programming (CPAIOR 2010)*, pp. 318-333.
- [182] Günneç, D. and Salman, F. S. (2011). Assessing the reliability and the expected performance of a network under disaster risk, *OR Spectrum*, 33, 499-523.
- [183] Akbari, V. and Salman, F. S. (2016) Multi-vehicle synchronized arc routing problem to restore post-disaster network connectivity. *European Journal of Operations Research*.
- [184] Celik, M., Ergun, O., and Keskinocak, P. (2015) The post-disaster debris clearance problem under incomplete information. *Operations Research*, 63(1), 65-85.
- [185] Clark, A. and Culkin, B. (2007). A network transshipment model for planning humanitarian relief operations after a natural disaster. Paper presented at EURO XXII- 22nd European Conference on Operational Research, Prague.
- [186] Ahmadi, M., Seifi, A., and Tootooni, B. (2015) A humanitarian logistics model for disaster relief operation considering network failure and standard relief time: a case study on San Francisco district, *Transportation Research Part E: Logistics and Transportation Review*, 75, 145-163.

- [187] JICA IMM. (2002). The study on a disaster prevention/mitigation basic plan in Istanbul including microzonation in the Republic of Turkey. Technical report. Japanese International Cooperation Agency, Local Municipality of Istanbul.
- [188] Barbarosoglu, G. and Arda, Y. (2004). A Two-Stage Stochastic Programming Framework for Transportation Planning in Disaster Response. *Journal of the Operational Research Society*, 55, 43-53.
- [189] Yazici, M. A. and Ozbay K. (2007). Impact of probabilistic road capacity constraints on the spatial distribution of hurricane evacuation shelter capacities. *Transportation Research Record: Journal of the Transportation Research Board*, 2022, 55-62.
- [190] Rawls, C. G. and M. A. Turnquist. (2012). Pre-positioning and dynamic delivery planning for short-term response following a natural disaster. *Socio-Economic Planning Sciences*, 46(1), 46-54.
- [191] Nolz, P. C., Semet, F., and Doerner, K. F. (2011). Risk approaches for delivering disaster relief supplies. *OR Spectrum*, 33, 543-569.
- [192] Vitoriano, B., Ortuño, M. T., Tirado, G., and Montero, J. (2011). A multi-criteria optimization model for humanitarian aid distribution. *Journal of Global Optimization*, 51, 189-208.
- [193] Erdik, M. and Aydinoglu, N. (2002). Earthquakes risk to buildings in Istanbul and a proposal for mitigation. KOERI Report no 2001/16 Bogazici Univeristy, Istanbul, Turkey.
- [194] FEMA (2012). Hazus: FEMA's methodology for estimating potential losses from disasters. Federal Emergency Management Agency. Retrieved from <http://www.fema.gov/hazus>
- [195] Han, J., Kamber, M., and Pei, J. (2012). Data mining : concepts and techniques. (3rd ed). San Francisco, CA: Mogan Kaufmann.
- [196] Huang, Z. (1998). Extensions to the k-Means Algorithm for Clustering Large Data Sets with Categorical Values. *Data Mining and Knowledge Discovery*, 2(3), 283-304.
- [197] Michie, D., Spiegelhalter, D. J., and Taylor, C. C. (1994). Machine Learning, Neural and Statistical Classification. Ellis Hardwood Limited.
- [198] Purward, A. and Singh, S. K. (2014). Empirical Evaluation of Algorithms to impute Missing Values for Financial Dataset 2014 International Conference on Issues and Challenges in Intelligent Computing Techniques (ICICT), pp. 652-656.
- [199] Patil, B., Joshi, R., and Toshniwal, D. (2010). Missing value imputation based on k-mean clustering with weighted distance. *Contemporary Computing, Communications in Computer and Information Science*, 94, 600-609.
- [200] Schmitt, P., Mandel, J., and Guedj, M. (2015). A Comparison of Six Methods for Missing Data Imputation. *J Biomet Biostat*, 6, 224.
- [201] Fujikawa, Y. and Ho, T. B. (2002). Cluster-based algorithms for filling missing values. *Lecture Notes in Artificial Intelligence*, 2336, 549-554.
- [202] Rahman, M. M. and Davis, D. N. (2013). Machine learning-based missing value imputation method for clinical datasets. *IAENG transactions on engineering technologies*, 229, 245-257.

- [203] Luengo, J., Garca, S., and Herrera, F. (2012). On the choice of the best imputation methods for missing values considering three groups of classification methods. *Journal of Knowledge and Information Systems*, 32, 77-108.
- [204] Rokach, L. and Maimon, O. (2005). Clustering. In Maimon, O., Rokach L. (eds), *The Data Mining and Knowledge Discovery Handbook* (pp. 37-57). Springer.
- [205] Li, D., Deogun, J., Spaulding, W., and Shuart, B. (2005). Dealing with missing data: Algorithms based on fuzzy sets and rough sets theories. Springer-Verlag, 37-57.
- [206] Ralambondrainy, H. 1995. A conceptual version of the k-means algorithm. *Pattern Recognition Letters*, 16, 1147-1157.
- [207] Kingsford, C. and Salzberg, S. L. (2008). What are decision trees? *Nat Biotechnol*, 26(9), 1011-1013.
- [208] Loh, W-Y. (2011). Classification and regression trees. *WIREs Data Mining Knowl Discov*, 1, 14-23.
- [209] Vickery, P. J., Lin, J. X., Skerlj, P. F., Twisdale, L. A., and Huang, K. 2006. HAZUS-MH hurricane model methodology. I: Hurricane hazard, terrain and wind load modeling. *Nat. Hazards Rev.*, 7(2), 82-93.
- [210] Scawthorn, C., Flores, P., Blais, N., Seligson, H., Tate, E., Chang, S., Mifflin, E., Thomas, W. Murphy, J., Jones, C., and Lawrence, M. 2006. HAZUS-MH flood loss estimation methodology. II: Damage and Loss Assessment. *Nat. Hazards Rev.*, 7(2), 72-81
- [211] Kircher, C.A., Robert V. Whitman, R.V., and Holmes, W.T. 2006. HAZUS Earthquake Loss Estimation Methods. *Nat. Hazards Rev.*, 7(2), 45-59.
- [212] Scawthorn, C., Flores, P., Blais, N., Seligson, H., Tate, E., Chang, S., Mifflin, E., Thomas, W. Murphy, J., and Jones, C. 2006. HAZUS-MH flood loss estimation methodology. I: Overview and flood hazard characterization. *Nat. Hazards Rev.*, 7(2), 60-71
- [213] Federal Emergency Management Agency FEMA. 2003. Multi-hazard loss estimation methodology, hurricane model, HAZUS, technical manual, developed by the Department of Homeland Security, Emergency Preparedness and Response Directorate, FEMA, Mitigation Division, Washington, D.C., under a contract with the National Institute of
- [214] Federal Emergency Management Agency FEMA. 2003. Multi-hazard loss estimation methodology, flood model, HAZUS, technical manual, developed by the Department of Homeland Security, Emergency Preparedness and Response Directorate, FEMA, Mitigation Division, Washington, D.C., under a contract with the National Institute of Building Sciences, Washington, D.C.
- [215] Federal Emergency Management Agency FEMA. 2003. Multi-hazard loss estimation methodology, earthquake model, HAZUS, technical manual, developed by the Department of Homeland Security, Emergency Preparedness and Response Directorate, FEMA, Mitigation Division, Washington, D.C., under a contract with the National Institute of Building Sciences, Washington, D.C.
Building Sciences, Washington, D.C.
- [216] Olsen A.H., Porter K.A. What We Know about Demand Surge: Brief Summary *Natural Hazards Review*, Vol. 12, No. 2, May 1, 2011, 62-71.

- [217] Cole C.R., Macpherson, D.A., Mccullough, K.A. A Comparison of Hurricane Loss Models. *Journal of Insurance Issues*, 33(1):31-53 2010
- [218] Clark, Karen, Near Term Hurricane Models: How Have They Performed? Boston, Mass.: Karen Clark & Company, December 2008. As of September 28, 2010: http://www.karenclarkandco.com/pdf/KCC_NearTermHurricanes.pdf
- [219] Hancilar U, Tuzun C, Yenidogan C, Erdik M (2010) ELER software: a new tool for urban earthquake loss assessment. *Nat Hazards Earth Syst Sci*, 10:2677-02696
- [220] Molina S, Lindholm C (2005) A logic tree extension of the capacity spectrum method developed to estimate seismic risk in Oslo, Norway. *J Earthq Eng*, 9(6):877-897
- [221] Crowley H, Silva V (2013) OpenQuake engine book: risk. GEM Foundation, Pavia, July 2013
- [222] Kreibich, H., Seifert, I., Merz, B., and Thieken, A. H. Development of FLEMOcs: A new model for the estimation of flood losses in companies. Hydrological Sciences Journal, *J. Sci. Hydrol.*, 55, 1302-1314, 2010.
- [223] ICPR: Atlas of flood danger and potential damage due to extreme floods of the Rhine, International Commission for the Protection of the Rhine, Koblenz, 2001.
- [224] Huizinga, H. J.: Flood damage functions for EU member states, HKV Consultants, Implemented in the framework of the contract 382442-F1SC awarded by the European Commission Joint Research Centre, 2007.
- [225] Huizinga, J., Moel, H. de, Szewczyk, W. (2017). Global flood depth-damage functions. Methodology and the database with guidelines. EUR 28552 EN. doi: 10.2760/16510
- [226] Tatham, P., L'Hermitte, C., and Spens, K., & Kovács, G. (2013). Humanitarian logistics: Development of an improved disaster classification framework. The 11th ANZAM Operations, Supply Chain and Services Management Symposium, (pp. 1-10), Brisbane, QLD, Australia.
- [227] ReliefWeb. (2014). ReliefWeb about page. Retrieved from <http://reliefweb.int/about>
- [228] Yagci Sokat, K., Zhou, R., Dolinskaya, I., Smilowitz, K., and Chan, J. (2016). Capturing Real-Time Data in Disaster Response Logistic. *Journal of Operations and Supply Chain Management*, 9(1), 23-54.
- [229] He, Z. (2006). Approximation Algorithms for K-Modes Clustering. Harbin Institute of Technology, China.
- [230] Khan, S. S. and Ahmad, A. (2013). Cluster center initialization algorithm for k-modes clustering. *Expert Systems with Applications*, 40, 7444-7456.
- [231] Breiman, L., Friedman, J., Olshen, R., and Stone, C. (1984). Classification and Regression Tree. New York: Chapman & Hall.
- [232] Breiman, L., Friedman, J. H., Olshen, R. A., and Stone, C. J. (1993). Classification and Regression Trees. Chapman & Hall, Boca Raton.

- [233] Hansen, M., Dubayah, R., and Defries, R. (1996). Classification trees: an alternative to traditional land cover classifiers, *International Journal of Remote Sensing*, 17(5),1075-1081.
- [234] Service Régional de Traitement d'Image et de Télédétection - a regional image processing and remote sensing service. (2010). Haiti Port-au-Prince Building damage, assessment per urban block. Retrieved from http://www.esa.int/spaceinimages/Images/2010/01/Damage_around_Port-au-Prince_Haiti
- [235] United Nations Institute for Training and Research. (2010a). Satellite-Identified IDP Concentrations, Road & Bridge Obstacles in Carrefour, Haiti (Update 1). Retrieved from http://unosat-maps.web.cern.ch/unosat-maps/HT/EQ20100114HTI/UNOSAT_HTI_EQ2010_IDP_PP_v2_LR.pdf (Accessed 2012-05-23)
- [236] United Nations Institute for Training and Research. (2010b). Density of Bridge & Road Obstacles in Port-au-Prince and Carrefour, Haiti (Update 2). http://unosat-maps.web.cern.ch/unosat-maps/HT/EQ20100114HTI/UNOSAT_HTI_EQ2010_ObstaclesOverview_v2_LR.pdf
- [237] European Commission - EC, Joint Research Centre - JRC, United Nations Institute for Training and Research - UNITAR, Operational Satellite Applications Programme - UNOSAT, World Bank Global Facility for Disaster Reduction and Recovery - GFDRR, and Centre National d'Information Géo-Spatial - CNIGS (2010). Building Damage Assessment Report Haiti earthquake 12 January 2010 Post Disaster Needs Assessment and Recovery Framework (PDNA), Report to the Haitian Government
- [238] Mathworks. (2015). Global Optimization Toolbox: User's Guide (r2014b). Retrieved May 10, 2015 from www.mathworks.com/help/pdf_doc/gads/gads_tb.pdf
- [239] Howden, M. (2009). How humanitarian logistics information systems can improve humanitarian supply chains: a view from the field. Proceedings of the 6th International ISCRAM Conference, Gothenburg, Sweden, May 2009. J. Landgren and S. Jul, eds.
- [240] Ergun, Ö., Stamm, J. L. H., Keskinocak, P., & Swann, J. L. (2010). Waffle House Restaurants hurricane response: A case study. *International Journal of Production Economics*, 126(1), 111-120.
- [241] Holdeman, E. How GIS can aid emergency management. 2014. Emergency Management. Available from <http://www.govtech.com/em/disaster/How-GIS-Can-Aid-Emergency-Management.html>
- [242] Haklay, M. 2010a. How good is volunteered geographical information? A comparative study of OpenStreetMap & Ordnance Survey datasets. *Env & Planning B: Planning & Design*, 37.4: 682.
- [243] Mashhadi, A., Quattrone, G., & Capra, L. 2013. Putting ubiquitous crowd-sourcing into context. Proceedings of CSCW, 611-622.
- [244] Arvis, J. F., Saslavsky, D., Ojala, L., Shepherd, B., & Busch, C. (2014). Connecting to compete 2014, trade logistics in the global economy: The logistics performance index and its indicators. Retrieved from <http://lpi.worldbank.org/>
- [245] Vaillancourt, A. (2013). Government decentralization and disaster impact, an exploratory study. Retrieved from <http://www.buildresilience.org/2013/proceedings/files/papers/352.pdf>

- [246] Connecting To Complete 2016. Trade Logistics in the Global Economy; The Logistics Performance Index and Its Indicators. Washington: The World Bank
- [247] Protection Cluster (2013). Protection Cluster displacement and 3W map. Retrieved from <http://www.humanitarianresponse.info/operations/philippines/infographic/protection-clusterdisplacement-and-3w-map>
- [248] OpenStreet Map (2016). Humanitarian OpenstreetMap Team. Retrieved from https://wiki.openstreetmap.org/wiki/Humanitarian_OSM_Team
- [249] Pauleit, S. and Duhme, F. (2000). Assessing the environmental performance of land cover types for urban planning. *Landscape and Urban Planning*, 52(1), 1-20.
- [250] Haiti Data. (2016). Maps. Retrieved from haitidata.org
- [251] United Nations Institute for Training and Research. (2010c). Haiti Earthquake 2010: Remote Sensing based Building Damage Assessment Data. Retrieved from <http://www.unitar.org/unosat/haiti-earthquake-2010-remote-sensing-based-building-damage-assessment-data>
- [252] UNICEF 2016. State of the World's Children 2016: A fair chance for every child. New York: UNICEF https://www.unicef.org/publications/files/UNICEF_SOWC_2016.pdf
- [253] Odhiambo, G.O, Musuva, R. M., Odiere, M.R., and Mwinzi, P.N. Experiences and perspectives of community health workers from implementing treatment for schistosomiasis using the community directed intervention strategy in an informal settlement in Kisumu City, western Kenya BMC Public Health BMC series 2016, 16:986.
- [254] McCord, G.C., Liu, A., and Singh, P. Deployment of community health workers across rural sub-Saharan Africa: financial considerations and operational assumptions. *Bulletin of the World Health Organization* 2012;91:244-253B.
- [255] Collins, D.H., Jarrah, Z., Gilmartin, C., and Saya, U. The costs of integrated community case management (iCCM) programs: A multi-country analysis. *Journal of Global Health* December 2014, Vol. 4, No. 2. 020407
- [256] Perry, H.B., Zulliger, R., and Rogers, M. Community health workers in low-, middle-, and high-income countries: an overview of their history, recent evolution, and current effectiveness. *Annu Rev Public Health*. 2014;35:399-421. Medline:24387091 doi:10.1146/annurev-publhealth-032013-182354
- [257] Anuj C. High altitude medicine. In: Munjal YP, Sharma SK, Agarwal AK, eds. *API Textbook of Medicine*. 9th ed. Mumbai, India: Jaypee Brothers Medical Publishers; 2012
- [258] Ozcan, YA (2009) *Quantitative methods in health care management: techniques and applications*, 2nd edn. Jossey-Bass, San Francisco, CA
- [259] Ministry of Health, Liberia. Revised National CommunityHealth Services Policy 2016-2021.
- [260] One Million Community Health Workers. (2015) Data for Decision Making Series: Diana Frymus. Available from <http://1millionhealthworkers.org/2015/05/04/data-for-decision-making-series-the-importance-of-chw-data-collection-with-diana-frymus/>

- [261] Mehta KM, Rerolle F, Rammohan SV, Albohm DC, Muwowo G, Moseson H, Sept L, Lee HL, Bendauid E (2016) Systematic motorcycle management and health care delivery: A field trial. *American Journal of Public Health*, 106(1), 87-94.
- [262] Von Achen, P., Smilowitz, K., Raghavan, M., Feehan, R. (2016) Optimizing community healthcare coverage in remote Liberia. *Journal of Humanitarian Logistics and Supply Chain Management*, Vol. 6 Issue: 3, pp.352-371
- [263] Otieno, C. E., Kaseje, D., Ochieng, B. M., & Githae, M. N. (2012). Reliability of community health worker collected data for planning and policy in a peri-urban area of Kisumu, Kenya. *Journal of Community Health*, 37(1), 48-53.
- [264] Admon, A. J., Bazile, J., Makungwa, H., et al. (2013). Assessing and improving data quality from community health workers: A successful intervention in Neno, Malawi. *Public Health Action*, 3(1), 56-59.
- [265] Mahmood, S., & Ayub, M. (2010). Accuracy of primary health care statistics reported by community based lady health workers in district Lahore. *Journal of the Pakistan Medical Association*, 60(8), 649-653.
- [266] Mitsunaga, T., Hedt-Gauthier, B. L., Ngizwenayo, E., Farmer, D. B., Gaju, E., Drobac, P., Basinga, P., Hirschhorn, L., Rich, M.L., Winch, P. J., Ngabo, F., Mugeni, C. Data for Program Management: An Accuracy Assessment of Data Collected in Household Registers by Community Health Workers in Southern Kayonza, Rwanda. *J Community Health* (2015) 40:625-632
- [267] Mitsunaga, T., Hedt-Gauthier, B., Ngizwenayo, E., Farmer, D. B., Karamaga, A., Drobac, P., Basinga, P., Hirschhorn, L., Ngabo, F., Mugeni, C. (2013). Utilizing community health worker data for program management and evaluation: Systems for data quality assessments and baseline results from Rwanda. *Social Science and Medicine*, 85, 87-92.
- [268] S. Martello and P. Toth. *Knapsack Problems*. Wiley, New York, 1990
- [269] Morton, D. P., R. Wood. 1998. On a stochastic knapsack problem and generalizations. D. Woodruff, ed. *Advances in Computational and Stochastic Optimization, Logic Programming and Heuristic Search*. Kluwer, Boston, 149-168.
- [270] Kleywegt, A. J., J. D. Papastavrou. 2001. The dynamic and stochastic knapsack problem with random sized items. *Oper. Res.* 49(1) 26-41.
- [271] Qin, Y., Wang, R., Vakharia, A. J., Chen, Y., and Seref, M. M. H. The newsvendor problem: Review and directions for future research. *European Journal of Operational Research*, vol. 213, no. 1, pp. 361-374, 2011.
- [272] Kogan, K. Unbounded knapsack problem with controllable rates: the case of a random demand for items. *J Oper Res Soc* (2003) 54: 594. doi:10.1057/palgrave.jors.2601554
- [273] Turken N, Tan Y, Vakharia A, Wang L, Wang R, Yenipazarli A (2012) The multi-product newsvendor problem: Review, extensions, and directions for future research. Choi T-M, ed. *Handbook of Newsvendor Problems*, Chap. 1 (Springer, New York).

- [274] Hadley, G., & Whitin, T. (1963). Analysis of inventory systems. Englewood Cliffs, NJ: Prentice-Hall.
- [275] Nahmias, S., & Schmidt, C. P. (1984). An efficient heuristic for the multi-item newsboy problem with a single constraint. *Naval Research Logistics Quarterly*, 31(3), 463-474.
- [276] Zhang, B., Xu, X., & Hua, Z. (2009). A binary solution method for the multi-product newsboy problem with budget constraint. *International Journal of Production Economics*, 117, 136-141.
- [277] Yadav, P., O. Stapleton, and L. Van Wassenhove. 2013. Learning from Coca-Cola. *Stanford Social Innovation Review* 11(1):51-55.
- [278] Shields L, Twycross A (2003) The difference between incidence and prevalence. *Paediatric Nursing*. 15, 7, 50.
- [279] Indrayan A. Medical Biostatistics. 2nd ed. Boca Raton, FL: Chapman & Hall/CRC; 2008
- [280] Mason CA, Kirby RS, Sever LE, Langlois PH. 2005. Prevalence is the preferred measure of frequency of birth defects. *Birth Defects Res A Clin Mol Teratol* 73:690-692. Mason CA, Kirby RS, Sever LE, Langlois PH. 2005. Prevalence is the preferred measure of frequency of birth defects. *Birth Defects Res A Clin Mol Teratol* 73:690-692.
- [281] Prüss-Üstün, A., Mathers, C., Corvalán, C., and Woodward, A. Introduction and methods: assessing the environmental burden of disease at national and local levels. Geneva, World Health Organization, 2003. (WHO Environmental Burden of Disease Series, No. 1).
- [282] Anand S., Hanson K. DALYs: Efficiency Versus Equity. *World Development*. 1998;26(2):307-10.
- [283] Williams A. Calculating the Global Burden of Disease: Time for a Strategic Appraisal? *Health Economics*. 1999;8(1):1-8.
- [284] Salomon J. A., Murray C. J. L. The Epidemiologic Transition Revisited: Compositional Models for Causes of Death by Age and Sex. *Population and Development Review*. 2002;28(2):205-28
- [285] Murray, C. J. L., J. A. Salomon, C. D. Mathers, and A. D. Lopez. 2002. Summary Measures of Population Health: Concepts, Ethics, Measurement, and Applications. Geneva: World Health Organization.
- [286] Murray, C. J. L., A. Tandon, J. A. Salomon, C. D. Mathers, and R. Sadana. 2002. "New Approaches to Enhance Cross-Population Comparability of Survey Results". In Summary Measures of Population Health: Concepts, Ethics, Measurement, and Applications, ed. C. J. L. Murray, J. A. Salomon, C. D. Mathers, and A. D. Lopez, 421-32. Geneva: World Health Organization.
- [287] Ministry of Health, Liberia. Revised National Community Health Services Policy 2016-2021.
- [288] World Health Organization The Community Health Worker: Working Guide, Guidelines for Training, Guidelines for Adaptation WHO, Geneva (1987)
- [289] Liberia Institute of Statistics and Geo-Information Services (LISGIS), Ministry of Health and Social Welfare [Liberia], National AIDS Control Program [Liberia], and ICF International. 2014. Liberia Demographic and Health Survey 2013. Monrovia, Liberia: Liberia Institute of Statistics and Geo-Information Services (LISGIS) and ICF International.

- [290] Jarrah, Z., K. Wright, C Suraratdecha, and D. Collins. 2013. Costing of Integrated Community Case Management in Senegal. Submitted to USAID by the TRAction Project: Management Sciences for Health. Available from http://www.msh.org/sites/msh.org/files/msh_costing_of_integrated_community_case_management_analysis_report_sene.pdf
- [291] Nefdt R, Ribaira E, Diallo K. Costing commodity and human resource needs for integrated community case management in thie differing community health strategies of Ethiopia, Kenya and Zambia. *Ethiop Med J*. 2014 Oct;52 Suppl 3:137-49.
- [292] Ministry of Health. Government of Ghana: National Community Health Worker (CHW) Program 2014 Available from http://1millionhealthworkers.org/files/2014/04/GH1mCHW_Roadmap_2014-04-20_Final_compressed.pdf
- [293] USAID | DELIVER PROJECT, Task Order 4. 2014. Quantification of Health Commodities: A Guide to Forecasting and Supply Planning for Procurement. Arlington, Va.: USAID | DELIVER PROJECT, Task Order 4. Available from http://deliver.jsi.com/dlvr_content/resources/allpubs/guidelines/QuantHealthComm.pdf
- [294] USAID | DELIVER PROJECT, Task Order 3. 2011. Guidelines for Managing the Malaria Supply Chain: A Companion to the Logistics Handbook. Arlington, Va.: USAID | DELIVER PROJECT, Task Order 3.
- [295] Watson, Noel, Loren Bausell, Andrew Ingles, and Naomi Printz. 2014. Malaria Seasonality and Calculating Resupply: Applications of the Look-Ahead Seasonality Indices in Zambia, Burkina Faso, and Zimbabwe. Arlington, Va.: USAID | DELIVER PROJECT, Task Order 7.
- [296] Briet O, Vounatsou P, Gunawardene DM, Galppaththy GNL, Amerasinghe PH: Models for short-term malaria prediction in Sri Lanka. *Malar J* 2008, 7:76.
- [297] Abellana R, Ascaso C, Aponete J, Saute F, Nhalungo D, Nhacolo A, Alonso P: Spatio-seasonal modeling of the incidence rate of malaria in Mozambique. *Malar J* 2008, 7:228.
- [298] Cancre N, Tall A, Rogier C, Faye J, Sarr O, Trape JF, Spiegel A, Bois F: Bayesian analysis of an epidemiological model of *Plasmodium falciparum* malarial infection in Ndiop, Senegal. *Am J Epidemiol* 2000, 152:760-770.
- [299] Chatterjee C, Sarkar RR: Multi-step polynomial regression method to model and forecast malaria incidence. *PLoS ONE* 2009, 4:e4726.
- [300] Strengthening Pharmaceutical Systems. 2011. Manual for Quantification of Malaria Commodities: Rapid Diagnostic Tests and Artemisinin-Based Combination Therapy for First-Line Treatment of *Plasmodium Falciparum* Malaria. Submitted to the US Agency for International Development by the Strengthening Pharmaceutical Systems Program. Arlington, VA: Management Sciences for Health.
- [301] Netessine, S. and Rudi, N. (2003) Centralized and Competitive Inventory Models with Demand Substitution. *Operations Research* 51(2):329-335.
- [302] Gudex, C., Dolan, P.H., Kind, P., Thomas, R., and Williams, H.A. *European Journal of Public Health*. 1997(7);4.

- [303] Tan-Torres Edejer T, Baltussen R, Adam T, Hutubessy R, Acharya A, Evans DB, et al., editors. Making choices in health: WHO guide to cost-effectiveness analysis. Geneva: World Health Organization; 2003.
- [304] Mathers CD, Vos T, Lopez AD, Salomon J, Ezzati M, editors. National burden of disease studies: a practical guide. Edition 2.0. Geneva: World Health Organization; 2001.
- [305] Adleman, Dan, Barnes-Schuster, Dawn, and Eisenstein, Don; Operations Quadrangle: Business Process Fundamentals, The University of Chicago Graduate School of Business, 1999, p. 5.
- [306] Institute of Medicine (US) Committee on the Economics of Antimalarial Drugs; Arrow KJ, Panosian C, Gelband H, editors. Saving Lives, Buying Time: Economics of Malaria Drugs in an Age of Resistance. Washington (DC): National Academies Press (US); 2004. 2, The Cost and Cost-Effectiveness of Antimalarial Drugs. Available from: <https://www.ncbi.nlm.nih.gov/books/NBK215621/>
- [307] WHO. 1996. Investing in Health Research and Development: Report of the Ad Hoc Committee on Health Research Relating to Future Intervention Options. TDR/Gen/96.1 . Geneva: World Health Organization.
- [308] WHO. Inventory Management. (MDS-3: Managing Access to Medicines and Health Technologies, Chapter 23) (2012)
- [309] Leung N-HZ, Chen A, Yadav P, Gallien J (2016) The Impact of Inventory Management on Stock-Outs of Essential Drugs in Sub-Saharan Africa: Secondary Analysis of a Field Experiment in Zambia. PLoS ONE 11(5): e0156026
- [310] Robberstad B, Strand T, Black RE, Sommerfelt H. Cost-effectiveness of zinc as adjunct therapy for acute childhood diarrhoea in developing countries. Bull World Health Organ. 2004 Jul;82(7):523-31.
- [311] Kale P. L., Hinde J. P. & Nobre F. F. Modeling diarrhea disease in children less than 5 years old. Ann Epidemiol 14, 371-377 (2004). <http://www.sciencedirect.com/science/article/pii/S1047279703002771>
- [312] Xu, Z., Hu, W., Zhang, Y., Wang, X., Zhou, M., Su, H., Huang, C., Tong, S., and Guo, Q. (2015). Exploration of diarrhoea seasonality and its drivers in China. Scientific Reports, 5, 8241. <http://doi.org/10.1038/srep08241>
- [313] Hadley, G., & Whitin, T. (1963). Analysis of inventory systems. New jersey: Prentice-Hall.