

**PREDICTION AND OPTIMIZATION APPLICATIONS IN DISASTER
MANAGEMENT**
(AFET YÖNETİMİNDE TAHMİN VE OPTİMİZASYON UYGULAMALARI)

by

Şule Nur SARGIN, B.S.

Thesis

Submitted in Partial Fulfillment
of the Requirements
for the Degree of

MASTER OF SCIENCE

in

INDUSTRIAL ENGINEERING

in the

GRADUATE SCHOOL OF SCIENCE AND ENGINEERING

of

GALATASARAY UNIVERSITY

Oct 2025

ACKNOWLEDGEMENTS

I sincerely appreciate all the invaluable guidance, support and encouragement throughout the entire process of writing this master thesis provided by my esteemed academic advisor, Assoc. Prof. M. Ebru ANGÜN.

I would like to thank my family for their love and support.

Oct 2025

Şule Nur SARGIN

TABLE OF CONTENTS

LIST OF SYMBOLS	vi
LIST OF FIGURES	vii
LIST OF TABLES	viii
ABSTRACT.....	ix
ÖZET	xi
1. INTRODUCTION	1
2. LITERATURE REVIEW	5
2.1 ML and DL Methods in Disaster and Hazard Prediction.....	6
2.2 ML and DL Methods for Risk and Vulnerability Assessment.....	6
2.3 ML and DL Methods for Damage Assessment.....	7
2.4 ML and DL Methods for Post-Disaster Response	8
2.5 ML and DL Methods to Predict the Number of Casualties and Injuries After Earthquakes	9
2.6 Advantages and Disadvantages of Earthquake Casualty Prediction Methods.....	14
3. METHODS	16
3.1 Gaussian Processes.....	16
3.2 The Leave-One-Out Cross-Validation Method for Validation	20
3.3 Chance-constrained Optimization	23
4. NUMERICAL APPLICATIONS	30
4.1 Input and Output Data for Injuries Prediction	30
4.2 Validation for the Prediction of Injuries Through Gaussian Process: Leave-One- Out Cross-Validation	32
4.3 Benchmark Prediction Methods.....	36

5. CONCLUSION	39
REFERENCES.....	40



LIST OF SYMBOLS

AFAD	: Disaster and Emergency Presidency of Turkey
ANN	: Artificial Neural Network
AI	: Artificial Intelligence
BP	: Back Propagation
CNN	: Convolutional Neural Network
DACE	: Design and Analysis of Computer Experiments
DL	: Deep Learning
ELM	: Extreme Learning Machine
FEMA	: Federal Emergency Management Agency
GDP	: Gross Domestic Product
GNIPC	: Gross National Income Per Capita
KNN	: K-Nearest neighborhood
LHD	: Latin Hypercube Designs
LOO-CV	: Leave-One-Out Cross-Validation
LPBoost	: Linear Programming Boost
LR	: Linear Regression
MAD	: Mean Absolute Deviation
MAD	: Mean Absolute Deviation
ML	: Machine Learning
MLEs	: Maximum Likelihood Estimators
NB	: Naïve Bayes
NN	: Neural Network
NDVI	: Normalized Difference Vegetation Index
PAGER	: Prompt Assessment of Global Earthquakes for Response
PGA	: Peak Ground Acceleration
PGV	: Peak Ground Velocity
RF	: Random Forest
QRE	: Quick Rough Estimate
RMSE	: Root Mean Square Error
RNN	: Recurrent Neural Network
SVM	: Support Vector Machine

LIST OF FIGURES

Figure 1.1: The number of the total natural disasters reported from 1900 till 2020..... 3

Figure 1.2: Fault lines that affect Turkey 3



LIST OF TABLES

Table 2.1: Prediction methods for casualties/injuries with given explanatory variables	13
Table 2.2: Advantages and disadvantages of earthquake casualty prediction methods	15
Table 4.1: Input and output data for injuries prediction	31
Table 4.2: Results of LOO-CV with the kernel type isotropic Gaussian	33
Table 4.3: Results of LOO-CV with the kernel type anisotropic Gaussian	33
Table 4.4: Results of LOO-CV with the kernel type anisotropic Matérn-3/2	34
Table 4.5: Results of LOO-CV with the kernel type anisotropic Matérn-5/2	35
Table 4.6: Evaluation of the kernel types	35
Table 4.7: Earthquake damage matrix of the percentage of buildings being in the vulnerability class B	37
Table 4.8: The relationship between the death ratio and the injury ratio with the degrees of damage to the buildings	37
Table 4.9: The number of the residential building and GDPs in 2023	38

ABSTRACT

This study considers the preparedness and response phases of a disaster management problem. One of the biggest problems for disaster management is that accurate predictions for the number of death and injured people are not available. Researchers in disaster management usually use data from the historical disasters as if they were the true numbers of death and injured people. However, in a future disaster, the numbers of injured and death people will be different with probability one. Consequently, recent studies focus on predicting these numbers applying different machine learning techniques such as neural networks, support vector machine, support vector regression, Gaussian process, and multiple linear regression. In this study, we use Gaussian processes to predict the numbers of injured people. This prediction method uses only two types of explanatory variables, namely, population density and the number of undamaged, slightly and moderately damaged buildings. We use the data resulting from the two earthquakes occurred on 6 February 2023 in the southeastern region of Turkey, and we collect the data from different sources. We compare our prediction method with three other benchmark prediction methods. None of the methods provide very accurate predictions. However, Gaussian processes have the advantage of estimating the variance of the predictor in addition to the predictor, and the predictor has normal distribution. We further use these predictors and their variances to solve a chance-constrained optimization problem to find shelter locations. We use randomly generated data for the optimization problem.

Further research has to be done to provide: i- machine learning method with better prediction capacity; ii- solution of the shelter location problem with real data; iii- solution procedure to solve a large-scale optimization problem for shelter location problem.

Keywords : Machine learning, Gaussian process, Chance constraints, Shelter location problem, Disaster management



ÖZET

Bu çalışma ile felaket yönetiminin hazırlık ve cevap aşamasına katkıda bulunulmaya çalışılmaktadır. Felaket yönetiminde yaşanan en büyük sorunlardan biri, felaket yaşanmadan önce, yaralı sayısını az hata ile tahmin etmektir. Felaket yönetimi konusunda çalışan çoğu araştırmacı eski felaket verilerini kesin veriler olarak kullanmaktadır. Oysaki bir sonraki felakette bu verilerin gerçekleşme olasılığı sıfırdır. Bu sebeple yeni çalışmalar makine öğrenmesi metotlarını kullanarak yaralı ve ölü sayılarını tahmin etmeye çalışmaktadır. Bu amaçla kullanılan metotlar yapay sinir ağları, destek vektör makinası (support vector machine), destek vektör regresyonu (support vector regression), Gauss süreci ve çoklu doğrusal regresyondur. Bu çalışmada yaralı sayıları Gauss süreci kullanılarak tahmin edilmektedir. Açıklayıcı değişken olarak nüfus yoğunluğu ve hasarsız, az ve orta hasarlı bina sayıları kullanılmıştır ve yaralı sayıları tahmin edilmeye çalışılmıştır. Veriler, 6 şubat 2023 tarihinde Türkiye'nin güneydoğu bölgesinde gerçekleşen deprem verilerinden oluşmaktadır ve farklı kaynaklardan toplanmıştır. Tahmin metodu, literatürden üç metot ile karşılaştırılmıştır. Metotlardan hiçbirisi az hatalı tahmin üretememiştir. Hataların hesaplanmasında eğitim-test very seti ayırma (leave-one-out cross-validation) metoduyla yapılmıştır. Bu metot her seferinde bir veri hariç diğer verileri eğitim seti olarak kullanıp, parametreleri tahmin etmekte ve daha sonra eğitimde kullanmadığı bir adet test versini kullanarak hatayı hesaplamaktadır. Bu hesaplama very setindeki tüm veriler için ayrı ayrı yapılmaktadır. Bu çalışmada kullanılan Gauss sürecinin diğer metotlara karşı olan en bariz üstünlüğü sadece tahmin ediciyi değil, tahmin edicinin varyansını da bulmaktadır. Ayrıca, tahmin edici normal dağılıma sahiptir. Bu tahminler, barınak seçimi yapabilmek için kullanılan olasılıksal kısıtları olan bir eniyileme probleminde kullanılmıştır. Bu eniyileme problemi rastlantısal olarak üretilmiş bir veri seti için çözülmüştür.

İleride yapılması gereken çalışmalar şu şekilde sıralanabilir: i- Farklı ve daha az hatalı makine öğrenmesi metotlarının dikkate alınması; ii- Olasılıksal kısıtları olan eniyileme probleminin gerçek veri seti ile çözülmesi ve sonuçların tartışılması; iii- Büyük ölçekli ve olasılıksal kısıtları olan problemler için çözüm yöntemlerinin geliştirilmesi.

Anahtar sözcükler : Makine öğrenmesi, Gauss süreci, Olasılıksal kısıtlar, Barınak seçimi problemi, Afet yönetimi



1. INTRODUCTION

A disaster, as defined by the Federal Emergency Management Agency (FEMA), is a “non-routine event that exceeds the capacity of the affected area to respond to it in such a way as to save lives, preserve property, and to maintain the social, ecological, economical, and political stability of the affected region”. Altay and Green (2006) provides one of the earliest and most comprehensive literature surveys on disaster management.

Altay and Green (2006) and FEMA distinguish four phases of the disaster management, namely, mitigation, preparedness, response, and recovery. The mitigation phase is the application of measures that will either prevent the onset of a disaster or reduce the impacts should one occur. Some typical activities of the mitigation phase are zoning and land use controls to prevent occupation of high hazard areas, building codes to improve disaster resistance of structures, and risk analysis to measure the potential for extreme hazards. The preparedness phase prepares the community to respond when a disaster occurs. Some typical activities are recruiting personnel for the emergency services and for community volunteering groups, emergency planning, budgeting for and acquiring vehicles and equipment, maintaining emergency supplies, and development of communication systems. The response phase is the employment of resources and the emergency procedures as guided by plans to preserve life, property, the environment, and the social, economic, and political structure of the community. Typical activities of this phase can be activating the emergency operations plan, activating the emergency operations center, evacuation of threatened population, opening of shelters and provision of mass care, emergency rescue and medical care, and fatality management. Finally, the recovery phase involves the actions taken in the long term after the immediate impact of the disaster has passed to stabilize the community and to restore normalcy. Some activities of this phase are disaster debris clean up, rebuilding of roads,

bridges, and key facilities, reburial of displaced human remains, and full restoration of lifeline services.

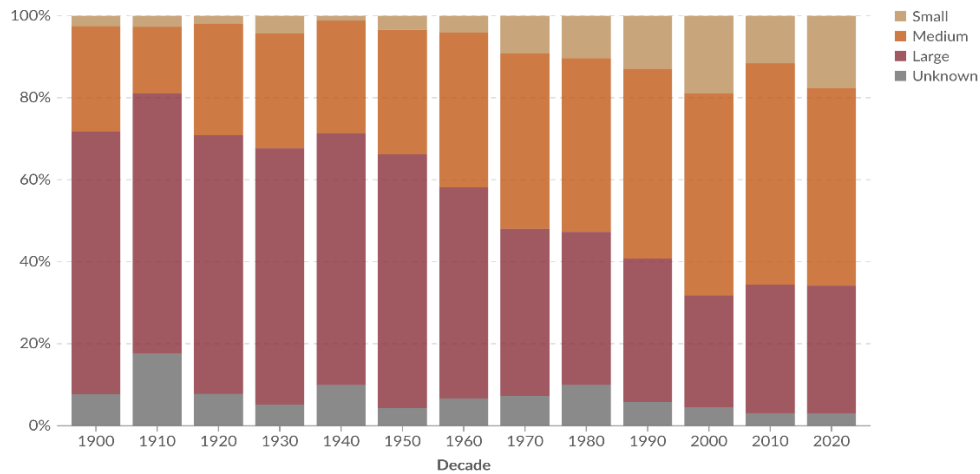
The following figure shows how the total numbers of the reported natural disasters change since 1900; see The International Disaster Database EM-DAT, CRED / UC Louvain (2025). These natural disasters include extreme weather, flood, extreme temperature, wet mass movement, wildfire, drought, earthquake, volcanic activity, dry mass movement, and glacial lake outburst flood. The reported numbers of disasters in year 2024 for volcanic activity, dry mass movement, and glacial lake outburst flood are seven, one, and one, respectively; so, the majority of the reports belong to the other disaster types. In particular, for 2024, the reported number of earthquakes is 14. EM-DAT classifies an event as “small” if the number of deaths does not exceed five, the number of affected people is at most 1500, and the event causes economic damages not exceeding 13 million US dollars. The economic threshold for previous years is adjusted for inflation accordingly. A “large-impact” event is the one that causes the death of at least 50 people, affects more than 150,000 people, and causes economic damages more than 320 million US dollars. An event which does not fall under small or large-impact is classified as “medium”. Also, the following limitation exists to compare the number of disasters over the years: “Time biases result from unequal reporting quality and coverage over time. Technologies and initiatives can be considered as responsible for the dominant trend observed. Therefore, it is challenging to infer insight into the actual drivers of disasters such as climate change, population growth, or disaster risk management”.

Turkey is located in a seismically active area within the complex zone of collision between the Eurasian plate and both the African and Arabian plates. Much of the country is affected by two fault zones, namely, the North Anatolian fault and the East Anatolian Fault; see Figure 1.2 from https://en.wikipedia.org/wiki/List_of_earthquakes_in_Turkey. Furthermore, the western part of the country is also affected by the Hellenic Arc. Finally, the easternmost part of Turkey lies on the western end of the Zagros fold and thrust belt. This location implies that earthquake is the most important disaster type in Turkey that causes the deaths of many people and devastates cities.

Global number of reported disasters by size per decade

Our World
in Data

'Small' events have 5 or fewer deaths, less than 1,500 people affected, or reported economic damages below US\$13 million. 'Large' events have at least 50 deaths, more than 150,000 people affected, or reported economic damages of at least US\$320 million. At least part of the increase in small and medium reported events is due to improved reporting over time.



Data source: EM-DAT, CRED / UCLouvain (2025)

OurWorldinData.org/natural-disasters | CC BY

Figure 1.1 The number of the total natural disasters reported from 1900 till 2020



Figure 1.2 Fault lines that affect Turkey

This study aims to determine the shelter locations in the aftermath of an earthquake. The number of injuries is predicted using a well-known machine learning method, namely,

Gaussian process. We compare our prediction method with two simple methods from the literature and with the multiple linear regression, which is one of the most used method for prediction of casualties. Our prediction method is simple and fast, and does not require the collection of many inputs. Furthermore, the Gaussian process prediction method provides not only the predictor but also the variance of the predictor, which is not available in neural networks. Moreover, the predictor has a normal distribution. We compute the root mean square error (RMSE), mean absolute deviation (MAD), R-square values for all prediction methods; we also compute the t -statistic to test the hypothesis that the error in prediction is zero. This validation is achieved through the leave-one-out cross-validation method that considers all but one input as the training set and the left-out input as the test set, and repeats the same procedure until all inputs are considered once in the training set and once in the test set. Our findings about the Gaussian process method can be improved; yet, they are in general better than the benchmark methods.

We then use our prediction in the shelter location selection problem, which is modeled as a chance-constrained optimization problem. We generate 25 hypothetical data and solve the problem over all samples to find the shelter locations. Many extensions can be considered to improve the current study including better machine learning techniques to predict injuries as benchmark methods, applying the shelter location problems using realistic data, and developing a solution method when the problem becomes large-scale.

2. LITERATURE REVIEW

This section reviews several machine learning (ML) techniques applied in different phases of the disaster management. There is an increasing trend in using ML to analyze and process big data coming from various data sources for informed decision making in disaster management. Recent reviews are Linardos et al. (2022) that focuses on publications from 2017 till 2021, Sun et al. (2020) that focuses on the use of artificial intelligence (AI) in disaster management, and Yu et al. (2018) that focuses on the use and potential of big data in natural disaster management. In particular, Sun et al. (2020) provides a figure which maps several AI techniques to their different application areas such as hazard risk assessment, vulnerability assessment, early warning systems, disaster detection, event mapping, damage assessment, disaster rescue, and relief and resource allocation.

In the rest of this section, we present four subsections with a few recent references that apply ML techniques applied to disaster and hazard prediction, risk and vulnerability assessment, damage assessment, and post-disaster response. These subsections are not directly related to the thesis; yet, they provide a general perspective to the readers and show the wide range of areas where ML techniques are applicable. The fifth subsection, however, provides a wide range of ML techniques that are used to predict the number of fatalities and/or injured disaster victims. Also, this subsection focuses on a single type of disaster, namely, earthquakes. The final subsection presents the advantages and disadvantages of several prediction methods.

2.1 ML and DL Methods in Disaster and Hazard Prediction

In these studies, data are analyzed in order to predict upcoming event or escalation of an event to minimize the unpredictability of disasters. Such studies refer to the mitigation and preparedness phases of disaster management.

Yuan and Moayedi (2020) proposes an optimization of the multilayer perceptron classification technique for landslide prediction where they apply evolutionary methods such as ant colony optimization, particle swarm optimization, genetic algorithm, and evolutionary strategy.

Sankaranarayanan et al. (2019) proposes a flood prediction method based on deep NN, which uses temperature and rainfall intensity. Their method outperforms traditional ML methods such as support vector machine (SVM), K-nearest neighborhood (KNN), and naïve Bayes (NB) methods.

Huang et al. (2018) proposes a forecasting method that applies fuzzy neural network (NN) combined with the locally linear embedding algorithm to predict the daily precipitation during typhoons. Their method outperforms the interpolation method of the European Centre for Medium-Range Weather Forecasts and stepwise regression approach.

Asim et al. (2017) applies different ML techniques to predict earthquake magnitude in the Hindukush region, Pakistan. These techniques were pattern recognition NN, recurrent NN (RNN), random forest (RF), and linear programming boost (LPBoost) ensemble classifier, where the LPBoost ensemble outperforms the other methods with respect to prediction accuracy.

2.2 ML and DL Methods for Risk and Vulnerability Assessment

Amin and Ahn (2021) proposes an awareness based educational system that applies a deep learning-based (DL-based) detection algorithm to identify the indoor objects that can become harmful during earthquakes.

Prasad et al. (2022) implements different ML techniques such as KNN, logit boost, boosted regression tree, and rotation forest to map the flood vulnerability. They apply these methods individually as well as with adabag as an ensemble. They find that the ensemble technique of adabag-rotation forest achieves the highest area under the curve value.

Nsengiyumva and Valentino (2020) applies NB tree, logistic model tree, and RF to predict the vulnerability of a specific area to landslides. They find that NB tree achieves the best results with respect to accuracy and precision.

Shirzadi et al. (2017) develops a method based on NB tree and random subspace ensemble to study landslide susceptibility mapping of a specific region in Iran.

Sriram et al. (2019) proposes a deep NN-based causal approach to perform vulnerability assessment of electric outages and roadway closures after extreme weather events.

2.3. ML and DL Methods for Damage Assessment

Presa-Reyes and Chen (2020) proposes a convolutional NN (CNN) architecture that is two-streamed network. The proposed architecture is trained by the pre- and post-hurricane aerial images, and can overcome the shortcomings of the previous CNN applications by effectively differentiate up to four damage levels.

Akshya and Priyadarsini (2019) applies K-means clustering and SVM to classify several areas according to their vulnerability to flood. They use aerial images of several areas captured by a drone as input to their classification.

Dotel et al. (2020) proposes a CNN-based approach to assess the impact of flood in urban and rural regions. They obtain pre- and post-flood satellite images of several regions and their proposed CNN identifies areas with maximal change.

Li et al. (2018) proposes a method based on CNNs and class activation maps to identify and assess damage in disaster-affected area. The method consists of a CNN that

classifies images into the classes damage and no damage, a class activation mapping, and a damage severity score.

2.4 ML and DL Methods for Post-Disaster Response

In the aftermath of a disaster, supply of relief aid is one of the main activities conducted for the disaster response phase. Lin et al. (2020) proposes a dynamic model to estimate demand for relief supplies using big data; this paper uses various sources of data including crowdsourcing, web mapping, and social media data, and applies multi-layer perceptron NN for estimation.

Li et al. (2020) carries out a qualitative research to construct a disaster medical rescue decision table based on the rough set theory. This paper applies a genetic algorithm to classify various disasters based on common medical features. As a result, the paper proposes recommendations for medical emergency rescue management.

Ehara et al. (2020) proposes a system using drones to recognize disaster victims and their status to recommend them to rescue teams for immediate help. Their system uses supervised ML techniques to classify the status of the victim; i.e., whether that victim is standing, sitting, or lying down on the ground.

Reynard and Shirgaokar (2019) aims to identify whether Twitter data can be used in planning response phase during or after disaster. This paper applies ML techniques to identify positive, negative or neutral sentiment from the tweets containing damage or transportation related information.

Chaudhuri and Bose (2020) collects geo-tagged images from earthquake-hit regions and applies a DL method for classification of these images to identify survivors in debris. This paper also find that DL methods classify images more accurately than traditional ML methods and use less computational resources.

Basu et al. (2019) investigates various supervised classification methods, unsupervised pattern matching and unsupervised information retrieval methods to propose a post-

disaster response management for resource needs and availabilities based on information obtained from tweets. This paper finds that if good quality training data are available from prior events, then classification approaches perform better. If, on the other hand, such data are not available, then unsupervised retrieval methods outperform supervised classification approaches.

2.5 ML and DL Methods to Predict the Number of Casualties and Injuries After Earthquakes

Casualty/injury prediction methods can be classified into two major types: Empirical approaches and analytical approaches. Empirical approaches base fatality predictions on seismic data such as magnitude, intensity, population density, and the numbers of casualties/injuries from historical earthquakes data. Analytical approaches consider soil conditions, building inventories, demographic data of the disaster region in addition to seismic data. As a result, analytical approaches provide predictions with smaller variances; yet, they require the collection of more data for longer periods. Below, we summarize different prediction models from both empirical and analytical approaches.

Bastami and Soghrat (2017) predicts the number of earthquake fatalities using linear regression, where they use distance to the epicenter or peak ground acceleration (PGA) as explanatory variables; because the PGA's are not available for villages, they also estimate PGA using linear regression, and this time, distance to the epicenter and focal depth are the explanatory variables.

Firuzi et al. (2022) models uncertainties in the number of fatalities of the past earthquakes and the ground motion parameters such as PGA, peak ground velocity, and modified Mercalli intensity by fuzzy regression and Monte Carlo simulation. To predict fatalities, this paper provides a simple model that multiplies the fatality rate with the exposed population.

Tang et al. (2019) proposes a two-stage approach to rapidly predict the number of casualties in the aftermath of an earthquake; they show that the cumulative number of casualties follows a logarithmic curve in time with a sharp increase in usually the first

72 hours. This paper first estimates fatality rates at different earthquake intensity levels, applying empirical regression method in Prompt Assessment of Global Earthquakes for Response (PAGER) system. Then, this paper estimates the expected number of fatalities by multiplying the fatality rates at each intensity levels with the number of fatalities. This paper requires that the number of fatalities accumulated over time from past earthquake data are available.

Liu et al. (2025) proposes a Quick Rough Estimate (QRE) method to predict the final total number of fatalities. This method directly depends on the cumulative fatality numbers available so far and on the rescue efforts, and it indirectly depends on the earthquake magnitude. The QRE estimates the death toll by quasi-linear and adaptive estimation method. Furthermore, as more data become available, the method predicts with more accuracy. This paper provides an empirical estimation method using the fatality data, which accumulates over time, provided by Disaster and Emergency Presidency of Turkey (AFAD) for M_S 8.0-7.9 earthquake doublet occurred in 2023 in the south east region of Turkey.

Li et al. (2021) proposes a fatality estimation method, multiplying the fatality rate with the population density exposed to the earthquake. Their proposed fatality rate uses fatality records normalized by the Human Development Index. To estimate the fatality rate, they consider three different models, namely, the logarithmic-linear model, the logarithmic-exponential model, and the lognormal cumulative model.

Shi et al. (2024) proposes a casualty prediction model for Taiwan using fatality rate, population exposure, and gross national income per capita (GNIPC) as inputs. Their method assumes that this fatality ratio follows a lognormal cumulative distribution function. They further show that the prediction models employed for different geographical areas differ significantly.

Gul and Guneri (2016) estimates the number of injuries applying artificial neural network (ANN), where they use earthquake occurrence time, earthquake magnitude, and population density as predictors. They use data from 21 earthquakes occurred in Turkey for network training.

Zhou et al. (2023) proposes a spatial evaluation model of earthquake-affected population based on the correlation characteristics of the influencing factors and the back propagation (BP) neural network. These influencing factors are divided into environmental factors, which consist of elevation, slope angle, population density, per capita Gross Domestic Product (GDP), distance to fault, distance to river, and normalized difference vegetation index (NDVI), and into seismic factors, which consist of PGA, peak ground velocity (PGV), and distance to epicenter. Of these ten influencing factors, they find that per capita Gross Domestic Product (GDP) and PGA had a stronger correlation with the earthquake-affected population, where the correlation coefficient is given by Spearman rank correlation coefficient.

Xin et al. (2020) establishes an extreme learning machine (ELM) network to predict earthquake fatalities based on data from 84 groups of victims in China. This paper also finds that ELM has better robustness and generalization capability than BP neural network and SVM. Xin et al. (2020) initially considers nine inputs, namely, earthquake occurrence time, magnitude, epicenter intensity, population density, epicenter distance, fortification level, building damage area, earthquake disaster geographical environment, and whether there are significant precursors. They further apply screening to these nine inputs through Principal Component Analysis, and find that epicenter intensity, population density, earthquake occurrence time, and damage area of buildings are the important inputs to be further used in their analysis.

Fang et al. (2020) proposes a casualty prediction and rescue model using a network model of the disaster area. The casualty prediction uses the damage matrix of a specific type of building with respect to the earthquake intensities. They further provide a matrix with the number of casualties with respect to the damage levels of the buildings. These matrices are obtained through sampling of the old data and simulation. Their approach requires the knowledge of all building types in the disaster area, which may not be available instantly.

Furthermore, Xu et al. (2025) proposes a spatiotemporal casualty assessment method of earthquake-induced falling debris, accounting for human emergency behavior through

agent-based simulation. They first obtain a high-resolution population distribution using Kriging interpolation, then they apply agent-based simulation for the emergency behaviors of the individuals during earthquake, and they use physics engine to simulate the distribution of falling debris in urban building cluster; so, they obtain a dynamic and spatiotemporal assessment method.

Ren et al. (2024) proposes an earthquake casualty assessment model for residential areas, which consider residential population, time dependent residential occupancy rate, proportion of age groups of the residents, probability of being in some level of damage for a building, where the damage levels are intact, slightly damaged, moderately damaged, severely damaged, and collapsed, and casualty rate of residents within an age group when the building is in a certain level of damage as inputs. Furthermore, this paper assumes that the seismic vulnerability function of a residential building follows a log-normal distribution, the earthquake occurrence time follows a uniform distribution, and the number of occupants of a building follows a binomial distribution. Under a given seismic intensity, the procedure generates a big sample of occurrence times, building damage states, and occupancy numbers to estimate casualty rates. These casualty rates are then used to obtain their distribution.

Huang and Jin (2018) first considers nine inputs, but applies PCA to select the most important ones among them and finally, it considers only the following five inputs, namely, epicenter intensity, magnitude, building damage area, earthquake occurrence time, and population density. The evolution of the number of casualties and injuries follow a two-stage rule during an earthquake; i.e., these numbers increase very fast in the early hours of an earthquake, then stabilize and decrease until the rescue ends. Because this pattern is similar to a Gaussian curve, Huang and Jin (2018) uses a corrected partial Gaussian curve as a function epicenter intensity to estimate casualties, where magnitude, building damage area, earthquake occurrence time, and population density are used for correction.

Table 2.1 summarizes the methods applied to predict the number of casualties and injuries after an earthquake.

Table 2.1 Prediction methods for casualties/injuries with given explanatory variables

Publication	Explanatory (input) variable	Prediction method
Bastami and Soghrat (2017)	PGA, distance to the epicenter, focal depth	Linear regression
Firuzi et al. (2022)	PGA, peak ground velocity, modified Mercalli intensity	Fuzzy regression and Monte Carlo simulation
Tang et al. (2019)	Accumulated number of fatalities at different intensity levels of earthquake	Empirical regression in PAGER
Liu et al. (2025)	Accumulated number of fatalities over time, rescue efforts, earthquake magnitude (indirectly)	QRE, which uses quasi-linear and adaptive estimation method
Li et al. (2021)	Estimated fatality rate, population density exposed to earthquake	Empirical estimate which multiplies fatality rate with exposed population density
Shi et al. (2024)	Fatality rate, population exposure, GNIPC	Empirical estimate which multiplies fatality rate with exposed population density and normalized GNIPC
Gul and Guneri (2016)	Earthquake occurrence time, earthquake magnitude, population density	ANN
Zhou et al. (2023)	Elevation, slope angle, population density, per capita GDP, distance to fault, distance to river, NDVI, PGA, PGV, distance to epicenter	Spearman rank correlation coefficient, BP neural network
Xin et al. (2020)	Epicenter intensity, population density, earthquake occurrence time, damage area of buildings	ELM
Fang et al. (2020)	Detailed knowledge about buildings such as reinforced concrete or steel structures (A-level), brick wall bearing structures (B-level), rural houses without formal design (C- and D-level)	Build the damage matrix using old data, simulation, and information about buildings. Then using damage matrix, find the death rate given in Table 5 of Fang et al. (2020).

		Number of fatalities is given by death rate multiplied by the population.
Ren et al. (2024)	Residential population, time-dependent residential occupancy rate, proportion of age groups of the residents, probability of being in some level of damage for a building, casualty rate of residents	Monte Carlo simulation, analytical model for earthquake casualty assessment
Huang and Jin (2018)	Epicenter intensity, magnitude, building damage area, earthquake occurrence time, population density	An analytical corrected partial Gaussian curve

2.6 Advantages and Disadvantages of Earthquake Casualty Prediction Methods

This section presents the advantages and disadvantages of the methods applied for casualty prediction. In particular, the advantages and the disadvantages of the two most applied methods, namely, NN and decision trees are given in detail below. Those of the other methods are summarized in Table 2.2, which is based on Table 2.3 of Huang and Jin (2018).

Most models for earthquake-casualty prediction are built using neural networks (NN) or linear regression (LR) methods. However, the NN methods have the following disadvantages: i- The model convergence is slow because most of the NN models are trained using BP algorithm; ii- The prediction accuracy of the model is affected by random initialization of the connection weight; iii- Interpretation of the model is poor; iv- The network structure has significant uncertainty because it can only be determined by human experience and several experiments; v- The generalization ability of the model is poor due to its propensity for overfitting. These NN methods have the advantages of nonlinear fitting ability and high prediction accuracy.

The decision tree is another model type, which is usually preferred in ML studies. The advantages of the decision tree models over NN methods are: i- They have faster training speed under the same data scale; ii- They have less difficulty in parameter adjustment; iii- They are able to model nonlinear relationships between inputs and outputs without requiring input scaling; iv- The decision tree, which is a set of if-then rules, is easier to interpret and use for decision-making by the end users; v- The decision tree model with pruning strategy has better generalization ability. Furthermore, if the number of inputs is big and there is less data, tree-based ensemble models outperform NN models; however, as the number of data increase, the advantages of using NN models become apparent.

Table 2.2 Advantages and disadvantages of earthquake casualty prediction methods

Methods	Advantages	Disadvantages
Logarithmic function	Applicable to specific areas	Poor in stability and accuracy
Linear regression	Applicable to specific areas	Low generalization and poor stability
High-order nonlinear fitting	Suitable for specific region and large-scale earthquakes	Limited to the study of low order parameters and poor stability
ANN	Strong self-adaptability and high precision	Difficult data collection and unstable prediction
Probability estimation	General applicability	Poor prediction
Casualty index	Theoretical significance	Difficult to acquire data and to popularize

For prediction, we will use Gaussian process (also known as Kriging) that we will introduce in Chapter 3. Gaussian process can predict accurately when there is a small amount of data. Furthermore, a property of Gaussian process is that it provides not only a point predictor but also its variance; the methods above cannot provide the predictor variance. We will also compare our prediction methods with some of the methods we present in this section in Chapter 4.

3. METHODS

This section introduces the prediction method, namely, Gaussian process (also known as Kriging), which is one of the ML methods. To validate the predictive model, we will use leave-one-out cross-validation, which is again one of the most used methods for validation in ML. Finally, we will present the chance-constrained optimization problem that we will apply to find shelter locations.

3.1 Gaussian Processes

Our presentation of the Gaussian process is based on Rasmussen and Williams (2006) and Chapter 5 of Kleijnen (2015).

Gaussian process models the output of an experiment as follows:

$$y(\mathbf{x}) = \phi(\mathbf{x})^T \boldsymbol{\beta} + \mathbf{M} \quad (3.1)$$

where $\mathbf{x}$ is a k -dimensional input vector (i.e., $\mathbf{x} \in R^k$), y is a scalar output (i.e., $y \in R$), $\phi(\mathbf{x})$ is a set of fixed basis functions (e.g., $\phi(\mathbf{x})$ consists of $[1, \mathbf{x}, \mathbf{x}^2]$ if $E(y)$ is given by a second-order polynomial of $\mathbf{x}$, $E[y(\mathbf{x})] = \beta_0 + \beta_1 x + \beta_2 x^2$), $\boldsymbol{\beta}$ is a vector of unknown coefficients, $\mathbf{M}$ follows a Gaussian process with zero mean vector and a covariance matrix $\boldsymbol{\Sigma}_M$, and the superscript T denotes the transpose of a vector or a matrix. This $\boldsymbol{\Sigma}_M$ is also known as kernel in ML. For simplicity, we assume that $\phi(\mathbf{x})$ is 1, so (3.1) becomes an ordinary Gaussian process (or Kriging):

$$y(\mathbf{x}) = \beta_0 + \mathbf{M}. \quad (3.2)$$

From (3.2), the expected value of y is $(y) = \beta_0$. Furthermore, we assume constant variance for y , so $Var(y) = \tau^2$.

The kernel Σ_M presents how the correlation between two outputs $y(\mathbf{x})$ and $y(\mathbf{x}')$ decreases, and it depends on the distance between $\mathbf{x}$ and $\mathbf{x}'$; if the distance becomes bigger, this correlation goes to zero (i.e., $y(\mathbf{x})$ and $y(\mathbf{x}')$ become uncorrelated), and if the distance becomes smaller, then this correlation gets closer to 1. Let h denote the Euclidean distance between $\mathbf{x}$ and $\mathbf{x}'$:

$$h = \sqrt{\sum_{j=1}^k (x_j - x'_j)^2}.$$

Then, an *isotropic correlation function* is given by

$$\rho(\mathbf{x}, \mathbf{x}') = \frac{\sigma(h)}{\tau^2}$$

where $\sigma(h)$ is the covariance between two vectors $\mathbf{x}$ and $\mathbf{x}'$. We further present the most popular *anisotropic correlation functions* such as Gaussian

$$\rho(\boldsymbol{\theta}, \mathbf{x}, \mathbf{x}') = \prod_{j=1}^k \exp[-\theta_j (x_j - x'_j)^2] = \exp\left[\sum_{j=1}^k -\theta_j (x_j - x'_j)^2\right], \quad (3.3)$$

Matérn with parameters $\nu = 3/2$ and $\nu = 5/2$

$$r = r(\boldsymbol{\theta}, \mathbf{x}, \mathbf{x}') = \sqrt{\frac{\sum_{j=1}^k (x_j - x'_j)^2}{\theta_j^2}}$$

$$\rho(\nu = 3/2, r) = (1 + \sqrt{3}r)\exp(-\sqrt{3}r)$$

$$\rho(v = 5/2, r) = (1 + \sqrt{5}r + \frac{5r^2}{3})\exp(-\sqrt{5}r)$$

where $\boldsymbol{\theta} = (\theta_1, \theta_2, \dots, \theta_k)$ is the vector of length scale parameters with $\theta_j \geq 0$.

The unknown parameters of the model in (3.2) are $\boldsymbol{\psi} = (\beta_0, \tau^2, \theta_1, \theta_2, \dots, \theta_k)$, and these parameters are usually estimated by the maximum likelihood estimation method. Let $\mathbf{X}$ be an $(n \times k)$ matrix of input vectors and $\mathbf{y}(\mathbf{X})$ be an $(n \times 1)$ vector of its corresponding output. This $\mathbf{y}(\mathbf{X})$ has a multivariate normal distribution. Therefore, the log-likelihood function is

$$\begin{aligned} l(\boldsymbol{\theta}, \beta_0, \tau^2) = & -\ln \left[(2\pi)^{\frac{n}{2}} \right] - 0.5 \ln[\det(\tau^2 \mathbf{R}_M(\boldsymbol{\theta}))] \\ & - 0.5(\mathbf{w} - \beta_0 \mathbf{1})^T (\tau^2 \mathbf{R}_M(\boldsymbol{\theta}))^{-1} (\mathbf{w} - \beta_0 \mathbf{1}), \end{aligned} \quad (3.4)$$

where $\ln(\cdot)$ denotes the natural logarithm of a positive scalar, $\det(\cdot)$ denotes the determinant of a matrix, $\mathbf{R}_M(\boldsymbol{\theta})$ is the $(n \times n)$ correlation matrix (i.e., $\boldsymbol{\Sigma}_M(\boldsymbol{\theta}) = \tau^2 \mathbf{R}_M(\boldsymbol{\theta})$), n is the number of input vectors, $\mathbf{w}$ is the $(n \times 1)$ vector of output realizations, and $\mathbf{1}$ is the vector with all components equal to one of appropriate dimensions. The maximum likelihood estimators (MLEs) $\hat{\boldsymbol{\psi}} = (\hat{\beta}_0, \hat{\tau}^2, \hat{\theta}_1, \hat{\theta}_2, \dots, \hat{\theta}_k)$ are obtained by minimizing the function that follows from (3.4)

$$\min_{\boldsymbol{\theta} \geq 0} \ln[\det(\tau^2 \mathbf{R}_M(\boldsymbol{\theta}))] + (\mathbf{w} - \beta_0 \mathbf{1})^T (\tau^2 \mathbf{R}_M(\boldsymbol{\theta}))^{-1} (\mathbf{w} - \beta_0 \mathbf{1}). \quad (3.5)$$

The minimization in (3.5) is accomplished through the following iterative algorithm.

Algorithm:

Initialization: Give an initial value for the vector $\hat{\boldsymbol{\theta}}$, so that $\hat{\mathbf{R}}_M = \mathbf{R}_M(\hat{\boldsymbol{\theta}})$.

While MLEs do not converge

Compute the generalized least squares estimator of β_0 :

$$\hat{\beta}_0 = (\mathbf{1}^T \hat{\mathbf{R}}_M^{-1} \mathbf{1})^{-1} \mathbf{1}^T \hat{\mathbf{R}}_M^{-1} \mathbf{w}.$$

Compute the MLE of τ^2 :

$$\hat{\tau}^2 = \frac{(\mathbf{w} - \hat{\beta}_0 \mathbf{1})^T \hat{\mathbf{R}}_M^{-1} (\mathbf{w} - \hat{\beta}_0 \mathbf{1})}{n}.$$

Solve the following problem to update $\hat{\boldsymbol{\theta}}$:

$$\min_{\hat{\boldsymbol{\theta}} \geq \mathbf{0}} \hat{\tau}^2 \det(\hat{\mathbf{R}}_M)^{-n}. \quad (3.6)$$

Update $\hat{\mathbf{R}}_M$ with the new found $\hat{\boldsymbol{\theta}}$.

End while

The output of the Algorithm is vector of the MLEs $\hat{\boldsymbol{\psi}} = (\hat{\beta}_0, \hat{\tau}^2, \hat{\theta}_1, \hat{\theta}_2, \dots, \hat{\theta}_k)$. The minimization problem in (3.6) is difficult because the objective function is multimodal. This implies that different MLEs may result from different software packages or from initializing the same package from different starting values. Besides MLEs, cross-validation can also be used to estimate $\boldsymbol{\theta}$; see Rasmussen and Williams (2006, pp. 116-124).

Now, given $\hat{\beta}_0, \hat{\boldsymbol{\Sigma}}_M$, and $\mathbf{X}$, which is the $(n \times k)$ matrix with input vectors as its rows, the predictor $\hat{y}(\mathbf{x}_* | \hat{\beta}_0, \hat{\boldsymbol{\Sigma}}_M, \mathbf{X})$ at a new point $\mathbf{x}_*$ is given by

$$\hat{y}(\mathbf{x}_* | \hat{\beta}_0, \hat{\boldsymbol{\Sigma}}_M, \mathbf{X}) = \hat{\beta}_0 + \hat{\boldsymbol{\sigma}}_M^T(\mathbf{x}_*) \hat{\boldsymbol{\Sigma}}_M^{-1} (\mathbf{w} - \hat{\beta}_0 \mathbf{1}) \quad (3.7)$$

and the variance of the predictor is given by

$$s[\hat{y}(\mathbf{x}_*|\hat{\beta}_0, \hat{\Sigma}_M, \mathbf{X})] = \hat{\tau}^2 - \hat{\sigma}_M^T(\mathbf{x}_*)\hat{\Sigma}_M^{-1}\hat{\sigma}_M(\mathbf{x}_*) + \frac{[1 - \mathbf{1}^T\hat{\Sigma}_M^{-1}\hat{\sigma}_M(\mathbf{x}_*)]^2}{\mathbf{1}^T\hat{\Sigma}_M^{-1}\mathbf{1}}, \quad (3.8)$$

where $\hat{\sigma}_M(\mathbf{x}_*)$ is the $(n \times 1)$ vector of covariances of the n old input vectors with the new point $\mathbf{x}_*$. For notational simplicity, we denote (3.7) and (3.8) by $\hat{y}(\mathbf{x}_*)$ and $s[\hat{y}(\mathbf{x}_*)]$, respectively.

To estimate $\hat{\psi}$, we use a free-of-charge MATLAB toolbox called Design and Analysis of Computer Experiments (DACE), which is well-documented in Lophaven et al. (2002) and can be obtained from <http://www2.imm.dtu.dk/pubdb/pubs/1460-full.html>. DACE allows its users to choose among three mean functions and seven kernels. The three mean functions mean that $\phi(\mathbf{x})$ can be 1, $[1, \mathbf{x}]$, or $[1, \mathbf{x}, \mathbf{x}^2]$. These seven kernels, however, exclude Matérn kernels, which are available in the MathWorks software on <https://ch.mathworks.com/help/stats/regressiongp.html>. All software packages for Gaussian processes provide $\hat{y}(\mathbf{x}_*)$ and $s[\hat{y}(\mathbf{x}_*)]$ after giving $\mathbf{X}$ and $\mathbf{w}$. Some packages including DACE also provide $\nabla\hat{y}(\mathbf{x}_*)$.

Finally, to obtain more accurate estimates of $\hat{\psi}$, it is preferable to obtain the input matrix $\mathbf{X}$ according to some space filling design; the most popular design for Gaussian process is the Latin Hypercube Designs (LHD). Other designs are orthogonal array, uniform, maximum entropy, minimax, maximin, integrated mean square prediction error, and “optimal designs”.

3.2. The Leave-One-Out Cross-Validation Method for Validation

Given the input matrix $\mathbf{X}$ and the vector of realizations of output $\mathbf{w}$, this validation method removes one input vector, say, input vector $\mathbf{x}_i$, from $\mathbf{X}$ and its corresponding output w_i from $\mathbf{w}$, and estimates $\hat{\psi}_{-i}$ without this combination. Then, it obtains the predictor $\hat{y}(\mathbf{x}_i)$ and its variance $s[\hat{y}(\mathbf{x}_i)]$ from (3.7) and (3.8). Now, the error in prediction can be computed by $\hat{e}(\mathbf{x}_i) = w_i - \hat{y}(\mathbf{x}_i)$. Therefore, the leave-one-out cross-validation (LOO-CV) method always considers a training set of size $(n - 1)$ and test set of size 1; yet, it repeats this experience for all n input combinations, resulting in n error

values. Below, we give an algorithm for LOO-CV. This section is mainly based on Kleijnen (2015).

Algorithm:

Input: Input matrix $\mathbf{X}$ and vector of output realizations $\mathbf{w}$

For $i = 1 : n$

Delete input $\mathbf{x}_i$ from $\mathbf{X}$ and output w_i from $\mathbf{w}$

Estimate $\hat{\boldsymbol{\psi}}_{-i}$ through MLEs

Compute $\hat{y}(\mathbf{x}_i) = \hat{\beta}_{0;-i} + \hat{\boldsymbol{\sigma}}_{M;-i}^T(\mathbf{x}_i)\hat{\boldsymbol{\Sigma}}_{M;-i}^{-1}(\mathbf{w}_{-i} - \hat{\beta}_{0;-i}\mathbf{1})$ and

$$s[\hat{y}(\mathbf{x}_i)] = \hat{t}_{-i}^2 - \hat{\boldsymbol{\sigma}}_{M;-i}^T(\mathbf{x}_i)\hat{\boldsymbol{\Sigma}}_{M;-i}^{-1}\hat{\boldsymbol{\sigma}}_{M;-i}(\mathbf{x}_i) + \frac{[1 - \mathbf{1}^T\hat{\boldsymbol{\Sigma}}_{M;-i}^{-1}\hat{\boldsymbol{\sigma}}_{M;-i}(\mathbf{x}_i)]^2}{\mathbf{1}^T\hat{\boldsymbol{\Sigma}}_{M;-i}^{-1}\mathbf{1}}$$

using $\hat{\boldsymbol{\psi}}_{-i}$

Compute $\hat{e}(\mathbf{x}_i) = w_i - \hat{y}(\mathbf{x}_i)$

End for

We use these n prediction errors $\hat{e}(\mathbf{x}_i)$ to compute root mean square error (RMSE), mean absolute deviation (MAD), and the determination coefficient (R-square) as criteria for model accuracy. We also perform a hypothesis testing which uses a t -statistic.

RMSE is sensitive to extreme values (small or large) and MAD is more robust to extreme values. The smaller the RMSE or MAD value is, the better the prediction is. Moreover, R-square describes how well the inputs explain the output, and its values ranges from zero to one, where a value close to one is better. RMSE is defined as

$$RMSE = \sqrt{\frac{\sum_{i=1}^n \hat{e}^2(\mathbf{x}_i)}{n}}, \quad (3.9)$$

MAD is defined as

$$MAD = \frac{\sum_{i=1}^n |\hat{e}(\mathbf{x}_i)|}{n}, \quad (3.10)$$

where $|\cdot|$ is the absolute value of a scalar, R-square is defined as

$$R - square = 1 - \frac{\sum_{i=1}^n (w_i - \bar{w})^2}{\sum_{i=1}^n \hat{e}^2(\mathbf{x}_i)} \quad (3.11)$$

where $\bar{w}$ is the sample average of the output realizations; i.e., $\bar{w} = \frac{\sum_{i=1}^n w_i}{n}$.

Furthermore, we perform the following hypothesis testing, where

$$H_0: e = 0$$

$$H_1: e \neq 0.$$

We use the following test statistic, which has a t -distribution with $(n - 1)$ degrees of freedom:

$$t = \frac{\bar{e}}{s(\bar{e})}, \quad (3.12)$$

where $\bar{e}$ is the sample average of the prediction errors and $s(\bar{e})$ is the sample standard deviation of $\bar{e}$

$$\bar{e} = \frac{\sum_{i=1}^n \hat{e}(\mathbf{x}_i)}{n}, s(\bar{e}) = \sqrt{\frac{\sum_{i=1}^n (\hat{e}(\mathbf{x}_i) - \bar{e})^2}{n(n-1)}}.$$

After computing the test statistic in (12), we compute the following p -value:

$$p - value = 2 * Prob\{T_{n-1} > t\} \quad (3.13)$$

where T_{n-1} follows a t -distribution with $(n - 1)$ degrees of freedom. Then, we compare the resulting p -value with $\alpha = 10\%$, $\alpha = 5\%$, and $\alpha = 1\%$, and comment on the results.

3.3 Chance-constrained Optimization

Chance-constrained optimization is one of the optimization model types to handle problems under uncertainty. It is first introduced in Charnes and Cooper (1959) and further developed in Miller and Wagner (1965). A chance constraint is also known as a probabilistic constraint. It can consist of a single constraint to be satisfied with a prespecified probability, which is then called an individual chance constraint. A chance constraint which includes multiple constraints to be satisfied simultaneously with a prespecified probability is called a joint chance constraint. For more details, we refer to Birge and Louveaux (1997).

The determination of the shelter locations before an earthquake occurs is an important part of earthquake preparedness. In Turkey, the candidate shelter locations are selected by the Turkish Red Crescent if the area is vulnerable to earthquakes, like İstanbul. The organization considers several criteria such as distance to healthcare institutions, electrical infrastructure, and sanitary system, and rank the candidate locations. Each candidate location is given either a zero or one with respect to every criterion, and the final weight of each location is computed as the convex combination of all its scores. Then, the candidate locations can be sorted in a descending order. After an earthquake occurs, the Turkish Red Crescent considers the sorted list of candidate locations to build shelters sequentially until the number of shelters is enough to accommodate all affected people.

Kılıcı et al. (2015) improves the selection criteria by adding the following criteria: i-) utilization rates of shelters; ii-) distance between the potentially affected people and the shelter locations; and iii-) pairwise utilization difference of the open shelter areas.

In terms of capacity, it is assumed that at least 3.5 square meters are allocated to each person in a shelter area. Further assumptions of the problem are given below.

- The set of candidate locations is known in advance.
- There is a maximum capacity (measured in squared meters) for each shelter location.
- Each shelter location has a weight computed a priori.
- The utilization rate of each shelter must be above a value specified in advance.
- Each district must be assigned to its closest shelter.
- The population of each district is assumed to be concentrated in the district centroid.
- The problem is to be formulated as a max-min problem, where the minimum weight of the open shelters is maximized.

We first introduce our notation to be used hereafter. Then, we present a deterministic version of the shelter location problem, assuming that we know for sure the total number of affected people at each district. Finally, we consider the randomness in the number of affected people, predict these numbers through Gaussian processes, and introduce two sets of chance constraints to handle the randomness.

Sets:

- J : Set of candidate shelter locations, indexed by j
- M : Set of districts, indexed by m

Parameters:

- u_j : Weight of candidate shelter location j
- $l_{j;m}$: Distance between candidate shelter location j and district m
- q_j : Capacity of candidate shelter location j in square meters
- d_m : Deterministic number of affected people at district m
- β : Threshold for the minimum utilization rate of a shelter
- γ_j : Probability that shelter j does not have enough capacity to handle all affected people assigned to it

δ_j : Probability that shelter j is used below the minimum utilization rate

Variables:

y_j : 0-1 decision variable, takes value 1 if shelter in location j is open, and 0 otherwise

$z_{j;m}$: 0-1 decision variable, takes value 1 if shelter in location j is allocated to district m , and 0 otherwise

u_{min} : Minimum weight among the open shelters

The distances $l_{j;m}$ are sorted in ascending order; we denote by $j_{m(1)}$ the closest shelter to district m , so $j_{m(r)}$ denotes the r th closest shelter to that district, where $r = 1, \dots, |J|$.

Now, the deterministic problem is given by

$$\max \quad u_{min} \quad (3.14)$$

$$\text{subject to } u_{min} \leq u_j y_j + (1 - y_j), \quad j \in J \quad (3.15)$$

$$\sum_{j \in J} z_{j;m} = 1, \quad m \in M \quad (3.16)$$

$$z_{j_{m(1)};m} \geq y_{j_{m(1)}}, \quad m \in M \quad (3.17)$$

$$z_{j_{m(r)};m} \geq y_{j_{m(r)}} - \sum_{s=1}^{r-1} y_{j_{m(s)}}, \quad m \in M, r = 2, \dots, |J| \quad (3.18)$$

$$\sum_{m \in M} d_m z_{j;m} \leq q_j y_j, \quad j \in J \quad (3.19)$$

$$\sum_{m \in M} d_m z_{j;m} \geq \beta q_j y_j, \quad j \in J \quad (3.20)$$

$$y_j \in \{0, 1\}, j \in J \quad (3.21)$$

$$z_{j,m} \in \{0, 1\}, j \in J, m \in M. \quad (3.22)$$

We can explain the objective function and the constraints as follows. The objective function in (3.14) maximizes the minimum weight among all open shelters. The set of constraints in (3.15) makes sure that u_{min} is less than or equal to the weight of any open shelter; so, together with the maximization in the objective, u_{min} is equal to the smallest weight of open shelters. The set of constraints in (3.16) implies that every district is assigned to one shelter. The sets of constraints in (3.17) and (3.18) require detailed explanation. The set of constraints in (3.17) assigns each district to its first closest shelter if that shelter is open. If the first closest shelter to that district is not open, then the set of constraints in (3.18) assigns that district to the second closest shelter if that one is open; if not, it assigns to the third closest shelter, etc. We further give a simple example with four candidate shelters, namely, S1, S2, S3, and S4. Suppose that these candidate shelters are sorted in an ascending order in terms of distances to a certain district, say m' , as S3 (the first closest), S4 (the second closest), S1 (the third closest), and S2 (the fourth closest). Then, the constraint (3.17) for district m' becomes

$$z_{S3,m'} \geq y_{S3}.$$

Furthermore, the set of constraints in (3.18) becomes

$$z_{S4,m'} \geq y_{S4} - y_{S3}$$

$$z_{S1,m'} \geq y_{S1} - y_{S3} - y_{S4}$$

$$z_{S2,m'} \geq y_{S2} - y_{S3} - y_{S4} - y_{S1}.$$

The set of constraints in (3.19) implies that the total number of affected people assigned to a certain candidate shelter cannot exceed its capacity. The set of constraints in (3.20) makes sure that the capacity usage of an open candidate shelter exceeds a minimum utilization rate. Finally, the sets of constraints in (3.21) and (3.22) define the domains of the decision variables.

In equations (3.19) and (3.20), we assume that the number of affected people d_m per district is known. However, we can only predict that number using a ML method, e.g., Gaussian process. Therefore, in these constraints, we replace the deterministic d_m by its Gaussian process prediction $\hat{y}_m(\mathbf{x})$, where $\mathbf{x}$ is the input data used to predict the number of affected people at district m . Hence, the constraints (3.19) and (3.20) are replaced by the following chance constraints

$$Prob \left\{ \sum_{m \in M} 3.5 \hat{y}_m(\mathbf{x}) z_{j;m} \leq q_j y_j \right\} \geq 1 - \gamma_j, \quad j \in J \quad (23)$$

$$Prob \left\{ \sum_{m \in M} 3.5 \hat{y}_m(\mathbf{x}) z_{j;m} \leq \beta q_j y_j \right\} \leq \delta_j, \quad j \in J \quad (24)$$

where the set of constraints in (3.23) implies that the assignment of a candidate shelter to districts is accomplished in such a way that the capacity constraint for that shelter holds with a high probability (i.e., the prespecified γ_j value is small), and the set of constraints in (3.24) implies that this assignment also limits the under utilization of that shelter with a high probability (i.e., the prespecified δ_j value is also small). Then, we solve the chance-constrained problem which consists of the objective function in (3.14), and constraint functions in (3.15) through (3.18), (3.23) and (3.24), and (3.21) and (3.22).

Note that due to the Gaussian process prediction, $\hat{y}_m(\mathbf{x})$ is normally distributed with mean in (3.7) and variance in (3.8); furthermore, $3.5 \hat{y}_m(\mathbf{x})$ is also normally distributed with mean multiplied by 3.5 of (3.7) and variance multiplied by 3.5^2 of (3.8). Moreover, the sum of normally distributed random variables (whether they are dependent or independent) is also normally distributed. Notice that if we use a multivariate Gaussian process approach to predict the numbers of affected people for all districts simultaneously, we can also predict correlated $\hat{y}_m(\mathbf{x})$'s for all districts. However, there are works provided by Kleijnen and Mehdad (2014) and Lim and Wu (2022) which show that multivariate Gaussian process can be inferior than univariate Gaussian process. Therefore, we also ignore possible correlations between the numbers of affected people at different districts

Now, we define the random total number of affected people assigned to shelter j by D_j

$$D_j = \sum_{m \in M} 3.5 \hat{y}_m(\mathbf{x}) z_{j;m}, \quad j \in J$$

which is normally distributed with mean

$$E[D_j] = \sum_{m \in M} 3.5 E[\hat{y}_m(\mathbf{x})] z_{j;m}, \quad j \in J$$

and variance

$$Var[D_j] = \sum_{m \in M} 3.5^2 Var[\hat{y}_m(\mathbf{x})] z_{j;m}^2, \quad j \in J$$

where an estimate of $E[\hat{y}_m(\mathbf{x})]$ is given by (7) and an estimate of $Var[\hat{y}_m(\mathbf{x})]$ is given by (8). Hence, the constraints (3.23) and (3.24) become

$$Prob\{D_j \leq q_j y_j\} \geq 1 - \gamma_j, \quad j \in J$$

$$Prob\{D_j \leq \beta q_j y_j\} \leq \delta_j, \quad j \in J$$

so that

$$Prob\left\{Z_j \leq \frac{q_j y_j - E[D_j]}{\sqrt{Var[D_j]}}\right\} \geq 1 - \gamma_j, \quad j \in J$$

$$Prob\left\{Z_j \leq \frac{\beta q_j y_j - E[D_j]}{\sqrt{Var[D_j]}}\right\} \leq \delta_j, \quad j \in J$$

where Z_j has a standard normal distribution. Now, the above inequalities can be rewritten as

$$\Phi\left(\frac{q_j y_j - E[D_j]}{\sqrt{Var[D_j]}}\right) \geq 1 - \gamma_j, \quad j \in J$$

$$\Phi\left(\frac{\beta q_j y_j - E[D_j]}{\sqrt{\text{Var}[D_j]}}\right) \leq \delta_j, \quad j \in J$$

where $\Phi(\cdot)$ is the cumulative distribution function of the standard normal distribution. Now, the deterministic equivalent of the chance constraints (3.23) and (3.24) become

$$q_j y_j - \sum_{m \in M} 3.5 E[\hat{y}_m(\mathbf{x})] z_{j,m} - z_{1-\gamma_j} \sqrt{\sum_{m \in M} 3.5^2 \text{Var}[\hat{y}_m(\mathbf{x})] z_{j,m}^2} \geq 0, \quad j \in J \quad (3.25)$$

$$\beta q_j y_j - \sum_{m \in M} 3.5 E[\hat{y}_m(\mathbf{x})] z_{j,m} - z_{\delta_j} \sqrt{\sum_{m \in M} 3.5^2 \text{Var}[\hat{y}_m(\mathbf{x})] z_{j,m}^2} \leq 0, \quad j \in J \quad (3.26)$$

where $z_{1-\gamma_j}$ and z_{δ_j} are $(1 - \gamma_j)$ -quantile and δ_j -quantile of the standard normal distribution, respectively. The constraints in (3.25) and (3.26) are nonlinear, and Kınay et al. (2018) obtains their linear approximations that enable them to solve large-scale problems. Yet, we will consider these nonlinear constraints and solve only small to medium scale problems using an off-the-shelf software in the numerical examples.

4. NUMERICAL APPLICATIONS

This section consists of six subsections. In the first subsection, we present input and output data we use to predict the number of injuries (i.e., the number of affected people) after an earthquake occurs. In the second subsection, we validate our prediction using the leave-one-out cross validation method and various statistical criteria that we introduce in Chapter 3. Then, in subsection four, we apply our prediction to different data sets that we find in the literature. In that subsection, we also apply different prediction methods to our input and output data. Therefore, this subsection explains why Gaussian process is important in prediction. In subsection five, we solve the shelter location problem using the numbers of injuries predicted through Gaussian processes. Finally, in subsection six, we solve the deterministic version of the shelter location problem that we introduce in Chapter 3 as a benchmark of the chance-constrained problem, and we discuss the solutions that we obtain from our problem and the benchmark problem.

4.1. Input and Output Data for Injuries Prediction

Based on our literature survey in Chapter 2, we conclude that the inputs to predict injuries, which are often used by many other researchers, are population density (i.e., population divided by the area of the province), and the total number of undamaged, slightly damaged, and moderately damaged buildings. Other common inputs are the longitude and the latitude of the province and the time of the earthquake. Initially, we included these inputs as well; yet, after our validation phase, we decided to focus on only the first two inputs, namely, the population density and the number of buildings, because increasing the number of inputs did not improve the prediction accuracy. We use the input and output data of the two earthquakes of magnitudes M_S 8 and M_S 7.9 which hit the southeastern part of Turkey on 6 February 2023. Because we consider the

cumulative data of these two earthquakes, we don't consider the magnitude as an input variable.

We collect the input data from different sources. We obtain the population data from the Report of 2023 Kahramanmaraş and Hatay Earthquakes provided by the Presidency of Strategy and Budget of the Republic of Turkey (Table 1: Population of Provinces based on 2022 data, page 9). We collect the areas of provinces (in square kilometers) and the number of injuries at every province (i.e., the outputs) from Wikipedia at https://tr.wikipedia.org/wiki/2023_Kahramanmara%C5%9F_depremleri. Finally, the number of undamaged, slightly damaged, and moderately damaged buildings are from the Preliminary Investigation Report of Istanbul Technical University for 6 February 2023 Earthquakes (Table 3.4., page 87). We present these inputs and outputs below in Table 4.1.

Table 4.1 Input and output data for injuries prediction

Province	Population	Area	Population density x_1	Number of buildings x_2	Number of injuries w
Adana	2,271,106	13,844	164.26	7,305	7,450
Adıyaman	635,169	7,337	86.57	23,617	17,499
Diyarbakır	1,804,880	15,168	118.99	25,482	902
Elazığ	591,497	9,383	63.03	2,321	379
Gaziantep	2,154,051	6,803	316.63	122,924	25,276
Hatay	1,686,043	5,600	301.07	49,227	30,762
Kahramanmaraş	1,177,436	14,525	81.06	47,034	9,243
Kilis	147,919	1,521	97.25	5,194	754
Malatya	812,580	12,313	65.99	17,368	9,108
Osmaniye	559,405	3,767	148.50	30,341	2,606
Urfa	2,170,110	19,242	112.77	33,642	8,919

We first compute the Pearson correlation coefficients (see, e.g., Law (2024)) between the population density and the number of injuries, and the number of undamaged, slightly and moderately damaged buildings and the number of injuries through

$$\hat{\rho}_{x_j,y} = \frac{\sum_{i=1}^n (x_{j,i} - \bar{x}_j)(w_i - \bar{w})}{\sqrt{\sum_{i=1}^n (x_{j,i} - \bar{x}_j)^2} \sqrt{\sum_{i=1}^n (w_i - \bar{w})^2}}$$

where $n = 11$ is the number of input/output data, $x_{j,i}$ is the i th input for $j = 1, 2$, $\bar{x}_j$ is the sample average of the 11 inputs, w_i is the i th output, and $\bar{w}$ is the sample average of the 11 outputs. The Pearson correlation between the density and the number of injuries is estimated as $\hat{\rho}_{x_1,y} = 0.75$ and the one between the number of buildings and the number of injuries as $\hat{\rho}_{x_2,y} = 0.68$. The two positive estimated correlations imply that as these inputs increase the number of injuries also increases. Furthermore, we conclude that these two inputs will provide accurate predictions of the number of injuries. Further validation studies will be done in the next subsection. As a further input related to the economic development level of each province, we consider the percentage of the GDP per district and estimate the Pearson correlation. This correlation is 0.48, which is smaller than the correlations of the other two inputs; so, we don't consider the percentage of the GDP per district as an input.

4.2. Validation for the Prediction of Injuries Through Gaussian Process: Leave-One-Out Cross-Validation

We use the input/output data in Table 4.1 to estimate the parameters of the Gaussian process. These parameters are sometimes called hyperparameters. This estimation is done by DACE, a free-of-charge MATLAB toolbox; see Lophaven et al. (2002). DACE estimates the parameters of the Gaussian process through the MLE method described in Chapter 4. To use DACE, we need to select a kernel type; so, we consider the most common kernel types, namely, Gaussian, and Matérn with parameters $\nu = 3/2$ and $\nu = 5/2$. DACE does not consider Matérn kernels; so, for these kernel types, we use Mathworks software on <https://ch.mathworks.com/help/stats/regressiongp.html>. Both softwares estimate the parameters of the Gaussian process, and the predictor $\hat{y}(\mathbf{x})$ and the prediction error $s[\hat{y}(\mathbf{x})]$. We can further select an anisotropic versus an isotropic kernel for the Gaussian kernel type. We consider four possible combinations, namely,

Gaussian isotropic and anisotropic, Matérn-3/2 anisotropic, and Matérn-5/2 anisotropic. Then, we select the combination which gives the most accurate results with respect to the criteria introduced in Chapter 4. In general, if one has some knowledge about the order of differentiability of the predictor, she can use this knowledge to select the kernel type; see Rasmussen and Williams (2006).

In Table 4.2 through 4.5, we show the results of the leave-one-out cross validation (LOO-CV) using each time a different combination.

Table 4.2 Results of LOO-CV with the kernel type isotropic Gaussian

Province	Number of injuries w_i (output)	Number of predicted injuries $\hat{y}_i$	Error in prediction $e_i = w_i - \hat{y}_i$
Adana	7,450	11,449	-3,999
Adıyaman	17,499	8,263	9,236
Diyarbakır	902	10,441	-9,539
Elazığ	379	6,731	-6,352
Gaziantep	25,276	9,776	15,500
Hatay	30,762	9,043	21,719
Kahramanmaraş	9,243	16,067	-6,824
Kilis	754	7,506	-6,752
Malatya	9,108	15,867	-6,759
Osmaniye	2,606	4,941	-2,335
Urfa	8,919	5,415	3,504

Table 4.3 Results of LOO-CV with the kernel type anisotropic Gaussian

Province	Number of injuries w_i (output)	Number of predicted injuries $\hat{y}_i$	Error in prediction $e_i = w_i - \hat{y}_i$
Adana	7,450	12,088	-4,638
Adıyaman	17,499	4,667	12,832

Diyarbakır	902	12,408	-11,506
Elazığ	379	11,293	-10,914
Gaziantep	25,276	8,763	16,513
Hatay	30,762	8,215	22,547
Kahramanmaraş	9,243	12,420	-3,177
Kilis	754	11,258	-10,504
Malatya	9,108	19,463	-10,355
Osmaniye	2,606	11,082	-8,476
Urfa	8,919	5,854	3,065

Table 4.4 Results of LOO-CV with the kernel type anisotropic Matérn-3/2

Province	Number of injuries w_i (output)	Number of predicted injuries $\hat{y}_i$	Error in prediction $e_i = w_i - \hat{y}_i$
Adana	7,450	6,408	1,042
Adıyaman	17,499	5,949	11,550
Diyarbakır	902	11,840	-10,938
Elazığ	379	11,252	-10,873
Gaziantep	25,276	8,847	16,429
Hatay	30,762	10,069	20,693
Kahramanmaraş	9,243	24,206	-14,963
Kilis	754	11,214	-10,460
Malatya	9,108	7,491	1,617
Osmaniye	2,606	11,101	-8,495
Urfa	8,919	10,989	-2,070

Table 4.5 Results of LOO-CV with the kernel type anisotropic Matérn-5/2

Province	Number of injuries w_i (output)	Number of predicted injuries $\hat{y}_i$	Error in prediction $e_i = w_i - \hat{y}_i$
Adana	7,450	6,106	1,344
Adıyaman	17,499	6,039	11,460
Diyarbakır	902	10,082	-9,180
Elazığ	379	11,252	-10,873
Gaziantep	25,276	8,865	16,411
Hatay	30,762	10,249	20,513
Kahramanmaraş	9,243	24,269	-15,226
Kilis	754	11,214	-10,460
Malatya	9,108	7,220	1,888
Osmaniye	2,606	11,155	-8,549
Urfa	8,919	11,144	-2,225

Now, we evaluate each kernel type with respect to RMSE, MAD, $R - square$, and the hypothesis testing that we introduce in Chapter 3, equations (3.9), (3.10), (3.11), (3.12), and (3.13). Smaller values of RMSE and MAD are preferable, and values closer to 1 of $R - square$ is preferable. Furthermore, p -values are evaluated considering the most common type I error rates, namely, $\alpha = 10\%$, 5% , and 1% .

Table 4.6 Evaluation of the kernel types

Kernel type	RMSE	MAD	$R - square$	t statistic	p -value
Gaussian isotropic	10,001	8,410.8	0.0507	0.2132	0.8354
Gaussian anisotropic	11,784	10,411	0.3163	-0.1126	0.9126
Matérn-3/2	8,970.9	7,640.3	0.2957	0.4498	0.6624
Matérn-5/2	9,034.4	7,752.2	0.2790	0.4700	0.6484

Based on Table 4.6, Matérn-3/2 performs best with respect to RMSE and MAD; yet, its R_square value is relatively low. Gaussian anisotropic performs better in terms of R_square value. Nevertheless, we use Gaussian process with the Matérn-3/2 kernel type to predict injuries. Finally, all computed p -values are bigger than 10%, 5%, and 1%. So, we cannot reject the null hypothesis $H_0: e = 0$ for all kernel types.

4.3. Benchmark Prediction Methods

We compare our prediction applying Gaussian process with benchmark prediction methods such as multiple linear regression, the method proposed by Shi et al. (2024), and the method proposed by Zhiming et al. (2020).

For the multiple linear regression model, we consider the following full-quadratic model

$$y = \beta_0 + \beta_1 x_1 + \beta_2 x_2 + \beta_{12} x_1 x_2 + \beta_{11} x_1^2 + \beta_{22} x_2^2 + \varepsilon$$

where $\beta = (\beta_0, \beta_1, \dots, \beta_{22})$ is the vector of unknown coefficients and ε is normally distributed noise with mean zero and unknown variance σ^2 . In this model, x_1 is the population density, x_2 is the number of undamaged, slightly and moderately damaged buildings, and y is the number of injuries. We will estimate β by the classic least-squares method.

Shi et al. (2024) estimates the number of fatalities with a simple expression. We change this expression to estimate the number of injuries as follows:

$$\text{Estimation} = R * \frac{GNIP_p}{GNIP_T} * \text{Population}$$

where R is the ratio of injury, $GNIP_p$ is the gross national income per capita for that province, $GNIP_T$ is the total gross national income per capita of the region hit by the same earthquake. This ratio of the gross national income and the sum of the gross national incomes over all provinces reflects the difference between provinces due to

their different economic growth levels; such a socio-economic factor is also used in Li et al. (2021).

Zhiming et al. (2020) uses the following two tables, Tables 4.7 and 4.8, that are obtained from Yin and Yang (2004). Table 4.7 shows the probability of a B vulnerability class building (i.e., brick wall bearing structure) be damaged under various earthquakes of different Mercalli intensities. Furthermore, Table 4.8 shows the probability of an injury or death for a given damage level. This simple method finds the damage level of a specific building for a given earthquake intensity. Then, it multiplies with the number of people living in that building with the corresponding death or injury ratio.

Table 4.7 Earthquake damage matrix of the percentage of buildings being in the vulnerability class B

Earthquake intensity	Undamaged	Slightly damaged	Moderately damaged	Severely damaged	Destroyed
VI	66.97	26.08	5.42	1.36	0.17
VII	63.37	23.11	8.97	3.60	0.96
VIII	48.25	24.12	16.27	8.25	3.10
IX	28.67	22.97	23.39	16.08	8.90
X	8.42	14.04	24.43	28.05	25.06

Table 4.8 The relationship between the death ratio and the injury ratio with the degrees of damage to the buildings

Damage state	Death ratio	Injury ratio
Undamaged	0	0
Slightly damaged	0	1/10,000
Moderately damaged	1/100,000	1/1,000
Severely damaged	1/1,000	1/200
Destroyed	1/30	1/8

The first earthquake on 6 February 2023 occurred at 1.17 am, and the second occurred at 10.24 am. Therefore, in the application of Zhiming et al. (2020), we only use the number of the residential buildings. These numbers are obtained from the Report of 2023 Kahramanmaraş and Hatay Earthquakes provided by the Presidency of Strategy and Budget of the Republic of Turkey (Table 11: The total number of buildings in provinces affected by the earthquake, page 26). Furthermore, the GDP in 2023 of each province is obtained from the Statistical Institute of Turkey (<https://www.tuik.gov.tr/>). The Mercalli intensity of both earthquakes is XII. Below in Table 4.9, we provide the number of residential buildings and the GDP in 2023 for every province. The sum of all GDPs (i.e., GDP_T) is $2,651,158,937 \times 10^3$.

Table 4.9 The number of the residential building and GDPs in 2023

Province	Number of residences	GDP in 2023 ($\times 10^3$)
Adana	404,502	555,079,118
Adıyaman	107,242	97,271,966
Diyarbakır	199,138	250,767,582
Elazığ	106,569	121,471,844
Gaziantep	269,212	493,893,693
Hatay	357,467	343,609,143
Kahramanmaraş	219,351	233,007,142
Kilis	33,399	33,169,591
Malatya	159,896	148,679,409
Osmaniye	128,163	104,148,760
Urfa	347,902	270,060,689

5. CONCLUSION

This study considers the preparedness and response phases of a disaster management problem. It is well-known that accurate predictions of casualties can help rescue people immensely. The disaster management literature usually uses the data collected from the historical earthquakes to predict the number of casualties. Recently, with the development of machine learning tools, there are better prediction methods to estimate the number of casualties such as various types of neural networks, support vector machines, support vector regressions, Gaussian processes, and multiple linear regression. Some of these methods, including Gaussian processes, can provide accurate estimates even with a small amount of data. Therefore, this study uses Gaussian process to predict the number of injuries. A further advantage of Gaussian processes compared to neural networks is that they provide both the predictor and the predictor variance, and the predictor is normally distributed. We compare our prediction with three benchmark methods: None of the methods provide perfect results. Yet, the normality of the prediction by Gaussian process is further used in our optimization problem. After predicting the number of injuries, we solve the shelter location problem, which is formulated as a chance-constrained optimization problem.

There are many extensions of the current study such as considering better machine learning methods to provide more accurate predictions, considering real data set to solve the optimization problem, and develop a solution procedure for a large-scale optimization problem with chance-constraints. These further studies are left for future research.

REFERENCES

- Akshya, J., & Priyadarsini, P. L. K. (2019, February). A hybrid machine learning approach for classifying aerial images of flood-hit areas. In 2019 International conference on computational intelligence in data science (ICCIDS) (pp. 1-5). IEEE.
- Altay, N., & Green III, W. G. (2006). OR/MS research in disaster operations management. *European journal of operational research*, 175(1), 475-493.
- Amin, M. S., & Ahn, H. (2021). Earthquake disaster avoidance learning system using deep learning. *Cognitive Systems Research*, 66, 221-235.
- Asim, K. M., Martínez-Álvarez, F., Basit, A., & Iqbal, T. (2017). Earthquake magnitude prediction in Hindukush region using machine learning techniques. *Natural Hazards*, 85(1), 471-486.
- Bastami, M., & Soghrat, M. R. (2017). An empirical method to estimate fatalities caused by earthquakes: the case of the Ahar–Varzaghan earthquakes (Iran). *Natural hazards*, 86(1), 125-149.
- Basu, M., Shandilya, A., Khosla, P., Ghosh, K., & Ghosh, S. (2019). Extracting resource needs and availabilities from microblogs for aiding post-disaster relief operations. *IEEE Transactions on Computational Social Systems*, 6(3), 604-618.
- Birge, J. R., & Louveaux, F. (1997). *Introduction to stochastic programming*. New York, NY: Springer New York.
- Charnes, A., & Cooper, W. W. (1959). Chance-constrained programming. *Management science*, 6(1), 73-79.
- Chaudhuri, N., & Bose, I. (2020). Exploring the role of deep neural networks for post-disaster decision support. *Decision Support Systems*, 130, 113234.
- Cui, S., Yin, Y., Wang, D., Li, Z., & Wang, Y. (2021). A stacking-based ensemble learning method for earthquake casualty prediction. *Applied Soft Computing*, 101, 107038.

- Dotel, , S., Shrestha, A., Bhusal, A., Pathak, R., Shakya, A., & Panday, S. P. (2020, March). Disaster assessment from satellite imagery by analysing topographical features using deep learning. In Proceedings of the 2020 2nd international conference on image, video and signal processing (pp. 86-92).
- Ehara, K., Aljehani, M., Yokemura, T., & Inoue, M. (2020, January). Individual status recognition system assisted by UAV in post-disaster. In 2020 IEEE International Conference on Consumer Electronics (ICCE) (pp. 1-3). IEEE.
- EM-DAT, CRED / UCLouvain (2025) – with major processing by Our World in Data. “Disasters – EM-DAT” [dataset]. EM-DAT, CRED / UCLouvain, “Natural disasters” [original data].
- Fang, Z., Huang, J., Huang, Z., Chen, L., Cong, B., & Yu, L. (2020). An earthquake casualty prediction method considering burial and rescue. *Safety science*, 126, 104670.
- Federal Emergency Management Agency (FEMA). Multi-hazard loss estimation methodology, earthquake model, Hazus-MH 2.1 Technical Manual, 2012. https://www.researchgate.net/publication/328175418_HAZUS_MH_21_Earthquake_Model_Technical_Manual.
- Firuzi, E., Hosseini, K. A., Ansari, A., & Tabasian, S. (2022). Developing a new fatality model for Iran's earthquakes using fuzzy regression analysis. *International Journal of Disaster Risk Reduction*, 80, 103231.
- Gul, M., & Guneri, A. F. (2016). An artificial neural network-based earthquake casualty estimation model for Istanbul city. *Natural hazards*, 84(3), 2163-2178.
- Huang, X., & Jin, H. (2018). An earthquake casualty prediction model based on modified partial Gaussian curve. *Natural Hazards*, 94(3), 999-1021.
- Huang, Y., Jin, L., Zhao, H. S., & Huang, X. Y. (2018). Fuzzy neural network and LLE algorithm for forecasting precipitation in tropical cyclones: Comparisons with interpolation method by ECMWF and stepwise regression method. *Natural Hazards*, 91(1), 201-220.
- Istanbul Technical University, The Preliminary Investigation Report of for 6 February 2023 Earthquakes. Retrieved on 05/01/2026, from https://haberler.itu.edu.tr/docs/default-source/default-documentlibrary/2023_itu_deprem_on_raporu.pdf, last

- Kılıcı, F., Kara, B. Y., & Bozkaya, B. (2015). Locating temporary shelter areas after an earthquake: A case for Turkey. *European journal of operational research*, 243(1), 323-332.
- Kınay, Ö. B., Kara, B. Y., Saldanha-da-Gama, F., & Correia, I. (2018). Modeling the shelter site location problem using chance constraints: A case study for Istanbul. *European Journal of Operational Research*, 270(1), 132-145.
- Kleijnen, J. P. C. (2015) *Design and Analysis of Simulation Experiments*, Springer.
- Kleijnen, J. P., & Mehdad, E. (2014). Multivariate versus univariate Kriging metamodels for multi-response simulation models. *European Journal of Operational Research*, 236(2), 573-582.
- Law, A. (2024). *Simulation Modeling and Analysis*, sixth edition. McGraw-Hill.
- Li, T., Li, Z., Zhao, W., Li, X., Zhu, X., Pan, S., ... & Li, J. (2020). Analysis of medical rescue strategies based on a rough set and genetic algorithm: A disaster classification perspective. *International Journal of Disaster Risk Reduction*, 42, 101325.
- Li, X., Caragea, D., Zhang, H., & Imran, M. (2018, August). Localizing and quantifying damage in social media images. In *2018 IEEE/ACM International Conference on Advances in Social Networks Analysis and Mining (ASONAM)* (pp. 194-201). IEEE.
- Li, Y., Xin, D., & Zhang, Z. (2021). A rapid-response earthquake fatality estimation model for mainland China. *International Journal of Disaster Risk Reduction*, 66, 102618.
- Lim, C. Y., & Wu, W. Y. (2022). Conditions on which cokriging does not do better than kriging. *Journal of Multivariate Analysis*, 192, 105084.
- Lin, A., Wu, H., Liang, G., Cardenas-Tristan, A., Wu, X., Zhao, C., & Li, D. (2020). A big data-driven dynamic estimation model of relief supplies demand in urban flood disaster. *International Journal of Disaster Risk Reduction*, 49, 101682.
- Linardos, V., Drakaki, M., Tzionas, P., & Karnavas, Y. L. (2022). Machine learning in disaster management: recent developments in methods and applications. *Machine Learning and Knowledge Extraction*, 4(2).
- Liu, Y., Wang, Z., & Zhang, X. (2025). Estimating the final fatalities using early reported death count from the 2023 Kahramanmaraş, Türkiye, MS 8.0–7.9 earthquake doublet and revising the estimates over time. *Earthquake Research Advances*, 5(2), 100331.

- Lophaven, S. N., Nielsen, H. B., Sondergaard, J. (2002) DACE: A Matlab Kriging toolbox, version 2.0., IMM Technical University of Denmark, Kongens, Lyngby.
- Miller, B. L., & Wagner, H. M. (1965). Chance constrained programming with joint constraints. *Operations research*, 13(6), 930-945.
- Nie, W., Fan, X., Wang, J., Wang, L., Qi, Y., & Liu, M. (2024). Fine-scale spatiotemporal earthquake casualty risk assessment considering building function types. *International Journal of Disaster Risk Reduction*, 112, 104806.
- Nsengiyumva, J. B., & Valentino, R. (2020). Predicting landslide susceptibility and risks using GIS-based machine learning simulations, case of upper Nyabarongo catchment. *Geomatics, Natural Hazards and Risk*, 11(1), 1250-1277.
- Prasad, P., Loveson, V. J., Das, B., & Kotha, M. (2022). Novel ensemble machine learning models in flood susceptibility mapping. *Geocarto International*, 37(16), 4571-4593.
- Presa-Reyes, M., & Chen, S. C. (2020, August). Assessing building damage by learning the deep feature correspondence of before and after aerial images. In 2020 IEEE conference on multimedia information processing and retrieval (MIPR) (pp. 43-48). IEEE.
- Ren, Y., Ma, D., Wang, W., Wang, Z., Zhao, X., & Zhu, M. (2024). Probabilistic assessment of earthquake casualties in residential areas. *International Journal of Disaster Risk Reduction*, 113, 104902.
- Reynard, D., & Shirgaokar, M. (2019). Harnessing the power of machine learning: Can Twitter data be useful in guiding resource allocation decisions during a natural disaster? *Transportation research part D: Transport and environment*, 77, 449-463.
- Sankaranarayanan, S., Prabhakar, M., Satish, S., Jain, P., Ramprasad, A., & Krishnan, A. (2020). Flood prediction based on weather parameters using deep learning. *Journal of Water and Climate Change*, 11(4), 1766-1783.
- Shapira, S., Novack, L., Bar-Dayan, Y., & Aharonson-Daniel, L. (2016). An integrated and interdisciplinary model for predicting the risk of injury and death in future earthquakes. *PLoS One*, 11(3), e0151111.
- Shi, Y., Li, Y., & Zhang, Z. (2024). Reevaluating earthquake fatalities in the Taiwan Region: Toward more accurate assessments. *Seismological Research Letters*, 95(3), 1939-1948.

- Shirzadi, A., Bui, D. T., Pham, B. T., Solaimani, K., Chapi, K., Kavian, A., Shahabi, H. & Revhaug, I. (2017). Shallow landslide susceptibility assessment using a novel hybrid intelligence approach. *Environmental Earth Sciences*, 76(2), 60.
- Sriram, L. M. K., Ulak, M. B., Ozguven, E. E., & Arghandeh, R. (2019). Multi-network vulnerability causal model for infrastructure co-resilience. *IEEE Access*, 7, 35344-35358.
- Sun, W., Bocchini, P., & Davison, B. D. (2020). Applications of artificial intelligence for disaster management. *Natural Hazards*, 103(3), 2631-2689.
- Tang, B., Chen, Q., Liu, X., Liu, Z., Liu, Y., Dong, J., & Zhang, L. (2019). Rapid estimation of earthquake fatalities in China using an empirical regression method. *International Journal of Disaster Risk Reduction*, 41, 101306.
- Presidency of Strategy and Budget, The Report of 2023 Kahramanmaraş and Hatay Earthquakes. Retrieved on 05/01/2026, from <https://www.sbb.gov.tr/wp-content/uploads/2023/03/2023-Kahramanmaras-ve-Hatay-Depremleri-Raporu.pdf>.
- Williams, C. K., & Rasmussen, C. E. (2006). *Gaussian Processes for Machine Learning*. Cambridge, MA: MIT press.
- Xing, H., Junyi, S., & Jin, H. (2020). The casualty prediction of earthquake disaster based on Extreme Learning Machine method. *Natural Hazards*, 102(3), 873-886.
- Xu, Z., Zhu, Y., Fan, J., Zhou, Q., Gu, D., & Tian, Y. (2025). A spatiotemporal casualty assessment method caused by earthquake falling debris of building clusters considering human emergency behaviors. *International Journal of Disaster Risk Reduction*, 117, 105206.
- Yin, Z. Q., & Yang, S. W. (2004). Earthquake loss analysis and fortification standards. *Geosciences* 2018, 8, 165
- Yu, M.; Yang, C.; Li, Y. Big data in natural disaster management: A review. *Geosciences* 2018, 8, 165
- Yuan, C., & Moayedi, H. (2020). Evaluation and comparison of the advanced metaheuristic and conventional machine learning methods for the prediction of landslide occurrence. *Engineering with Computers*, 36(4), 1801-1811.
- Zhao, M., Jiang, W., Yan, G., Zhang, X., & Ma, R. (2021). Instant prediction of earthquake casualties for early rescue planning: a joint Poisson mixed modeling approach. *International Journal of Disaster Risk Reduction*, 58, 102178.

- Zhou, H., Che, A., Shuai, X., & Zhang, Y. (2023). A spatial evaluation method for earthquake disaster using optimized BP neural network model. *Geomatics, Natural Hazards and Risk*, 14(1), 1-26.
- Zhu, X., & Sun, B. (2017). Study on earthquake risk reduction from the perspectives of the elderly. *Safety science*, 91, 326-33.

