

1-2021

Establishing a frame of reference for measuring disaster resilience

Christopher W. Zobel
Virginia Tech

Cameron A. MacKenzie
Iowa State University, camacken@iastate.edu

Milad Baghersad
Florida Atlantic University

Yuhong Li
Old Dominion University

Follow this and additional works at: https://lib.dr.iastate.edu/imse_pubs



Part of the [Operational Research Commons](#), and the [Systems Engineering Commons](#)

The complete bibliographic information for this item can be found at https://lib.dr.iastate.edu/imse_pubs/271. For information on how to cite this item, please visit <http://lib.dr.iastate.edu/howtocite.html>.

This Article is brought to you for free and open access by the Industrial and Manufacturing Systems Engineering at Iowa State University Digital Repository. It has been accepted for inclusion in Industrial and Manufacturing Systems Engineering Publications by an authorized administrator of Iowa State University Digital Repository. For more information, please contact digirep@iastate.edu.

Establishing a frame of reference for measuring disaster resilience

Abstract

Due to the increasing occurrence of disruptions across our global society, it has become critically important to understand the resilience of different socio-economic systems, i.e., to what extent those systems exhibit the ability both to resist a disruption and to recover from one once it occurs. In order to characterize this ability, however, one must be able to quantitatively measure the relative level of resilience that a given system displays in response to a disruptive event. Such a measurement should be easily understandable and straightforward to implement, but it should also utilize a consistent frame of reference so that one can properly compare the relative performance of different systems and assess the relative effectiveness of different resilience investments. With this in mind, this paper presents an improved approach for measuring system resilience that supports better decision making by providing both consistency and flexibility across different contexts. The theoretical basis for the approach is developed first, and then its advantages and limitations are illustrated in the context of several different practical examples.

Keywords

Predicted resilience, Disaster operations management, Quantitative modeling, Decision support

Disciplines

Operational Research | Systems Engineering

Comments

This is a manuscript of an article published as Zobel, Christopher W., Cameron A. MacKenzie, Milad Baghersad, and Yuhong Li. "Establishing a frame of reference for measuring disaster resilience." *Decision Support Systems* 140 (2021): 113406. DOI: [10.1016/j.dss.2020.113406](https://doi.org/10.1016/j.dss.2020.113406). Posted with permission.

Establishing a frame of reference for measuring disaster resilience

Christopher W. Zobel

Department of Business Information Technology, Virginia Tech, Blacksburg, VA

Cameron MacKenzie

Department of Industrial and Manufacturing Systems Engineering, Iowa State University, Ames, IA

Milad Baghersad

Department of Information Technology & Operations Management, Florida Atlantic University, Boca Raton, FL

Yuhong Li

Department of Information Technology and Decision Sciences, Old Dominion University, Norfolk, VA

Abstract

Due to the increasing occurrence of disruptions across our global society, it has become critically important to understand the resilience of different socio-economic systems, i.e., to what extent those systems exhibit the ability both to resist a disruption and to recover from one once it occurs. In order to characterize this ability, however, one must be able to quantitatively measure the relative level of resilience that a given system displays in response to a disruptive event. Such a measurement should be easily understandable and straightforward to implement, but it should also utilize a consistent frame of reference so that one can properly compare the relative performance of different systems and assess the relative effectiveness of different resilience investments. With this in mind, this paper presents an improved approach for measuring system resilience that supports better decision making by providing both consistency and flexibility across different contexts. The theoretical basis for the approach is developed first, and then its advantages and limitations are illustrated in the context of several different practical examples.

Keywords:

Predicted resilience; Disaster operations management; Quantitative modeling; Decision support

1. Introduction

The National Academies' Committee on Science, Engineering, and Public Policy defines disaster resilience to be the ability to prepare and plan for, absorb, recover from, or more successfully adapt to actual or potential adverse events [1]. This definition encapsulates both proactive and reactive capabilities and activities, and it applies to individuals and businesses and communities both on a physical level and on a socio-economic level. The inherent richness and complexity of the concept of resilience provides the opportunity for researchers in many different disciplines to characterize and quantify different aspects of resilient behavior from their own perspective. This makes the concept both useful and accessible across a wide variety of different discipline-specific contexts.

The last few years have seen a rise in the relative impact of disasters on our global society [2], and it thus is becoming increasingly important to understand not only what makes a system resilient but also how one might act to improve its capacity for resilience. In order to develop such an understanding, however, we must first be able to determine a baseline level of resilience and then assess the extent to which an investment in time and resources has improved that baseline resilience. This, in turn, implies that we must define a quantitative measure of resilience that can serve as the basis for such a comparison.

From an engineering perspective, a common approach for assessing disaster resilience is to look at the physical characteristics of a system, and to consider how system loss evolves over time by studying the extent to which the system is initially damaged and the amount of time needed for it to regain its normal functionality [3–7]. On the socio-economic side of things, there are also many studies that adopt a social sciences perspective and focus instead on less dynamic indicators of a community's overall capacity for resilient behavior, such as a home ownership, crime rate, medical capacity, and employment [8–10]. The resilience framework discussed below is derived from the first of these two perspectives, and it thus focuses on characterizing the dynamic loss and recovery experienced by a system during a disaster event. As part of the discussion, however, we will show that the approach can

easily be adapted to help analyze socio-economic indicators of resilient behavior, and thus that it is applicable across a broad variety of contexts.

Whether a decision maker is analyzing different systems after a disaster in order to compare the extent to which they were affected, or simulating different disaster-related scenarios in order to assess the predicted effectiveness of different resilience investments, having a shared frame of reference will ensure more valid and actionable results. The approach below thus seeks to clarify and strengthen the theoretical and practical basis for quantitatively measuring resilience by explicitly providing for a specific frame of reference to support the proper comparison of different resilience outcomes [11]. It does so by building on the "predicted resilience" measure of Zobel [11,12], which was developed to compare the resilience of different potential (i.e., predicted) scenarios in a simulation environment.

The original predicted resilience measure allows decision makers to proactively compare the expected effectiveness of different (future) resilience investments in advance of an actual disruption, but it also can be used descriptively to characterize observed resilience behaviors both during and after a disaster [7,13]. The main contribution of this paper is to examine and extend several of the predicted resilience measure's fundamental assumptions about measuring loss and recovery time, and then to rigorously derive specific parameter thresholds that allow for extending the general approach to a wider variety of situations. The overall focus of these extensions is on improving the general technique so that it can be used more broadly and effectively to support decision makers in their efforts to quantify and compare the relative resilience of different systems or scenarios.

We begin our discussion with a brief overview of the literature on quantifying disaster resilience and a look at the background and motivation for the original "predicted resilience" measure. We then provide an in-depth discussion of the threshold parameters that extend the theoretical and practical relevance of the original measure by providing a valid and consistent frame of reference for resilience

calculations. In order to explore the implications of adopting these parameters, and to demonstrate the usefulness of the new framework, we then use the new resilience measure to illustrate three different instances of resilient behavior associated with the impacts of Superstorm Sandy on New York City. Finally, we conclude with a discussion of the managerial implications of the approach and provide several possible further extensions.

2. Background and Motivation

2.1. Measuring Resilience

A variety of approaches have been developed for characterizing and measuring disaster resilience from different perspectives. Within the social sciences, for example, a significant amount of research has been focused within the area of community resilience [8]. Such research often involves identifying measurable indicators of a community's capacity for resilient behavior, such as the percentage of the population who are non-native speakers and the ratio of large businesses to small businesses within the community [10,14,15]. These individual indicators are then normalized and assigned weights in order to create a single aggregated measure of community resilience [8]. In part because such indicator variables do not tend to change significantly over time, the results are mixed in terms of how well approaches such as this can describe a community's response to, and recovery from, an actual disaster event [15–18].

The study of supply chain resilience tends to focus on the ability of firms and supply chains to return to normal operations following a disruptive event [19–22], and it also often involves defining indicators, in order to either qualitatively or quantitatively characterize a firm's capacity for such resilience [23–26]. In particular, metrics that account for the network of suppliers, firms, and customers can be particularly important for assessing supply chain and business resilience [27–29]. Other approaches, in turn, have

focused less on measuring supply chain resilience directly and more on assessing the factors that contribute to supply chain resilience [30,31].

From an engineering perspective, quantifying resilience typically involves measuring the performance of an engineered system over time with a particular focus on the system performance during and after the disruptive event [3,5,32]. Several different approaches have been taken to quantify resilient behavior in this context, including assessing resilience as a time-dependent measure of the ratio of restored performance to total lost performance [33,34], and assessing uncertainty by modeling multiple possible performance trajectories [35–38]. Network resilience has also been explored in the context of understanding interdependencies among infrastructure components, in order to measure loss of functionality following a disruption and/or the time until the entire network recovers [34,39].

We focus in this paper on a particular type of engineering-based approach for measuring resilience that is based on the theoretical concept of the disaster resilience triangle [3–5,40] (see Figure 1). Given a time series response curve that represents some important measure of system performance, Figure 1 provides an illustration of both the immediate effect of a sudden impact disaster and the system's subsequent response behavior. The area above that response curve then can be used as a quantitative measure of the *loss of resilience* in the system due to the occurrence of a specific disaster event [3].

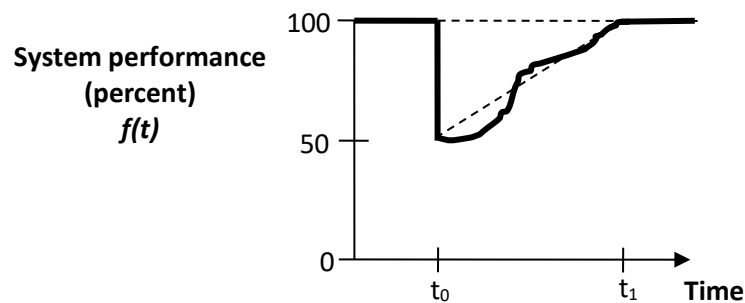


Fig 1. The original resilience triangle (adapted from [3])

Bruneau et al. [3] introduced the notion that such a curve may be used to represent either a physical characteristic of the system (such as the relative percent of households in a community that have electric power during the week after a major storm) or a social or organizational or economic characteristic (such as the relative level of employment during the year after a major hurricane). This allows the concept to be applied in a variety of different contexts in order to help capture some of the complexity and multi-dimensional nature of overall system resilience.

Cimellaro, Reinhorn, and Bruneau [41] and Zobel [11] extended the original idea of Bruneau et al. [3] by defining a direct measure of resilience in terms of the area *beneath* the time series curve. A number of other research efforts also expanded on Bruneau's initial work by looking at considerations such as uncertainty [42], multi-dimensionality [4], slow-onset disasters [43], multi-event disasters [44], and non-linear disaster recovery [7].

The discussion below builds on the work of Zobel [11,12], which focuses on characterizing the tradeoffs between the loss suffered by the system due to a disaster event and the subsequent system recovery time, by using the relative area beneath the curve. Given a parameter, T^* , which is typically chosen to represent the maximum allowable recovery time for the process being modeled, Zobel defines *predicted disaster resilience* to be the relative amount of functionality retained by the system over time. This is equal to the ratio of the area beneath the curve for the disrupted system to the area beneath the curve if no disruption has taken place (See Figure 2), and it is calculated as follows:

$$R(X, T) = \frac{T^* - \frac{XT}{2}}{T^*} = 1 - \frac{XT}{2T^*} \quad X \in [0,1], T \in [0, T^*]. \quad (1)$$

where X is the initial loss in system performance as a fraction and T is the time until recovery. It is important to note that equation (1), in keeping with Bruneau et al.'s [3] resilience triangle, assumes instantaneous loss and a constant rate of recovery. This implies that, in addition to T^* , equation (1)

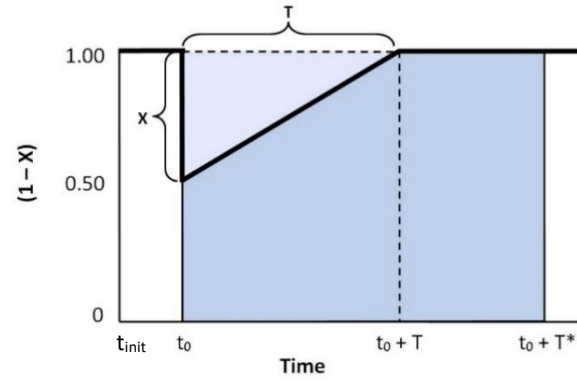


Fig 2. Predicted resilience as a proportion of T^*

depends only on the product of the initial loss (X), as a percentage of total available functionality, and the length of time to recovery (T). Consequently, given the same maximum recovery time, two different systems (or two characteristics of the same system) may have exactly the same calculated resilience value even though they have very different levels of initial resistance and very different recovery times.

To explicitly capture these tradeoffs between loss and recovery, and to characterize their relationship to system resilience, Zobel [11] identified a simple transformation of equation (1) into a series of parallel hyperbolic curves. This new representation, in which each curve symbolizes the set of (X , T) values that corresponds to a single level of resilience, provides a visual indication of the relative contribution of both loss and recovery time to the single, calculated resilience value.

In order to extend this technique and enable it to be applied to multiple consecutive events, with non-linear recovery behaviors, Zobel and Khansa [44] generalized equation (1) by adopting a piecewise linear response curve to capture the timing of the different events. This maintained the ability to use the resilience tradeoff curves to represent the interplay between loss and recovery time, but involved switching to average loss $\bar{X}$ over the duration of the disruption instead of initial loss as the variable of interest (See equation (2)). This simple change, however, allows for a more generalized equation for resilience that no longer assumes either instantaneous loss or linear recovery behavior. Furthermore,

because average loss is independent of the trajectory of the recovery process, it also allows the concept to be applied in the context of more general, non-piecewise linear response curves:

$$R(\bar{X}, T) = \frac{T^* - \bar{X}T}{T^*} = 1 - \frac{\bar{X}T}{T^*} \quad \bar{X} \in [0,1], T \in [0, T^*]. \quad (2)$$

This more generalized measure of resilience will be the basis for the discussion that follows.

From a theoretical standpoint, the resilience curves illustrated in Figure 3 provide a simple yet descriptive approach for quantifying and visualizing resilience, and thus they enable different systems to be compared with respect to how well they react to a disaster. Each resilience curve (contour line) demonstrates the tradeoff between the average loss $\bar{X}$ and the time until recovery T with greater resilience occupying the lower left-hand corner of the figure. This provides support for a decision maker

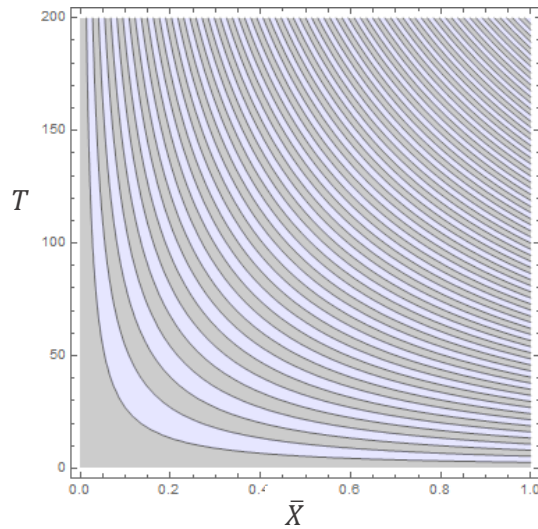


Fig 3. Resilience tradeoff curves

to analytically compare the results of investing in different resilience-building strategies, so that it is possible to make a more informed decision about the best approach for strengthening their system.

MacKenzie and Zobel [45] extend the model in equation (2) to determine how best to allocate resources in order to optimize system resilience.

2.2. Adjusting Resilience Values

Regardless of the functional form of any given quantitative resilience measure, it is extremely important to recognize that such a measure will only be useful if it produces results that a decision maker can accept as reasonable. For example, consider the following:

Suppose that 30% of the houses in a community lose power (and thus heat) after a major snowstorm, but recovery begins immediately and all power is restored after three days. If the T^* value, as given in (2), is set at seven days, then the calculated resilience for the community in this situation (assuming a constant rate of recovery) would be around 94%. However, given that there was a relatively large number of households without power, and given the significant impact that power loss during a snowstorm has on families and on the community, the town's emergency manager might be more likely to feel that the actual exhibited amount of resilience was closer to a lower value, such as 75%, thus leaving more perceived room for improvement (and more leverage for convincing the town council to press for infrastructure improvements before the next storm). Such a scenario suggests that a theoretically calculated resilience value might need to be recalibrated, i.e. adjusted upward or downward, based on a decision maker's perceptions about how much resilience the system actually exhibited in response to the disruption. This can be done by adjusting the frame of reference within which that resilience behavior is being measured. Once a proper frame of reference is determined for that particular resilience context, it can serve as the basis for all future analysis of behavior within that context by that decision maker. The ability to implement such adjustments can make the resilience measure a much more useful, and thus more relevant, tool in the long run.

In order to address this opportunity, Zobel [12] originally introduced an approach for modeling differing perceptions about resilience, such as these, which used the *perceived* level of resilience to adjust a set of calculated values in order to better match expectations. This approach allows for adjusting the reported value of the overall level of resilience, as well as the values associated with the sub-measures of loss and recovery time, by incorporating an appropriate set of parameters into equation (1). In the following discussion, we further explore this notion of adjusting resilience, but we do so from a slightly different perspective that focuses specifically on the value of the T^* parameter. In particular, by relaxing the typical assumption that T^* represents the maximum possible (or acceptable) recovery time, we can provide a decision maker with more flexibility in determining the relative resilience of different systems. This, in turn, leads to the opportunity to significantly strengthen the applicability of the resilience quantification approach to a number of additional situations.

We begin our discussion by defining the broader context within which an appropriate value for T^* can be established, and we show that within an appropriate range the actual value of this parameter can vary without affecting the relative ranking of any resilience values that are calculated from it. We follow this by showing that an equivalent threshold parameter, X^* , can be defined for the value of X , and then demonstrate that this leads to opportunities for applying the concept of quantitative resilience measurement in a much more generalized context.

3. Exploring the T^* parameter

Before we consider the theoretical basis for how one might adjust the value of the T^* parameter in our generalized resilience equation (equation (2)), it is important to recognize that there may be situations in which the decision maker cannot adjust the parameter's value because it is assigned a fixed value, either at the organizational level or at the industry level. For example, as a result of performing a Business Impact Analysis an organization will typically define a threshold parameter, called the

maximum tolerable period of disruption (MTPD), for key products and their critical activities [46–48]. The MTPD is analogous to the T^* parameter in that it represents the “time frame within which the impacts of not resuming activities would become unacceptable to the organization” [49]. Because it is defined at the organizational level, however, an individual decision maker who adopts the MTPD for use in calculating resilience, as the value for T^* , will not typically have the flexibility to change it.

The discussion below focuses on the T^* parameter as defined by Zobel [11,12]. However, a number of other research efforts have also proposed similar concepts for quantitatively measuring resilience. For example, Ouyang et al. [50,51] explicitly use a fixed large value of T in order to compare the resilience of multiple events. Adjetey-Bahun et al. [52] also use a fixed parameter T that represents the length of the simulation study over which they are collecting multiple observations of resilient systems. Zhao et al. [53] use the term *maximum tolerable recovery time* and Li et al. [54] use the term *maximum allowable recovery time* to refer to threshold parameters that are similar to the MTPD. Zhang et al. [55] specifically define a T^* parameter to represent the “largest rapidity” value to be considered in their calculations, and Kong et al. [56] provide a t^* value that they use as “the specific time set for resilience assessment.” Finally, both Gong et al. [57] and Zhao et al. [58] use a parameter t_n that defines the time at which performance assessment ends.

Previous research has recommended choosing a large enough value for T^* to effectively exceed the maximum observed recovery time for any system being analyzed [7,12,53–55]. If this is so, then the process of dividing by T^* in equation (2) will always normalize the resilience measure to a value between zero and one. Just given this heuristic recommendation, however, the actual choice of T^* could be somewhat arbitrary, since many different values could satisfy the recommendation that T^* exceed the maximum observed recovery time. This implies that, given the same initial loss and the same recovery trajectory, two different decision makers could choose different values of T^* in equation (2), and receive different calculated resilience values for exactly the same system. In order to compare the performance

of two different systems within the same context, therefore, it is critical that those systems are compared on the basis of a shared T^* value.

3.1. Justifying a shared T^* value

We easily can show that it is necessary to include a shared T^* in our equation for resilience, in some form, by directly considering the result of not doing so. Suppose, for example, that we define a related alternate measure that includes only the area under the response curve between time 0 and time T , without mention of T^* :

$$R'(\bar{X}, T) = T - \bar{X}T = T(1 - \bar{X}) \quad \bar{X} \in [0,1], T \geq 0. \quad (3)$$

This equation represents the difference between complete functionality for a length of time T and the area above the disrupted curve for which resilience is to be measured $\bar{X}T$. This provides a very straightforward way of calculating the total amount of functionality retained over time, while allowing one to compare the resilience of different systems subject to the same disruptive event. At the same time, however, such a measure also could lead to inconsistent results. Suppose, for example, that one system has an initial loss of functionality of 50% with a constant rate of recovery that takes 6 weeks, and a second identical system has an initial loss of 40% with constant recovery for 3 weeks. The total loss suffered by the first system (i.e., the area of the resilience triangle) is then 1.5, which is greater than the total loss suffered by the second system (0.6). However, since each calculation uses its own value of T , the calculated resilience value of the first system (4.5) is actually greater than that of the second (2.4). Given that resilience represents the ability to resist against and then recover from a disruption, one would expect that a comparison of any two similar systems (over the same time period and subject to the same disruption) should result in the system with less actual total loss being assigned a higher value for its measured resilience. This equation implies that a decision maker should actually lengthen the time until recovery T in order to improve resilience, which is nonsensical.

This issue is caused by the use of two different time frames to compare the two systems, and it easily can be resolved by specifying a shared time frame with respect to which the loss is assessed. We may therefore simply subtract the loss, in each case, from the total area representing complete functionality for a defined common period T^* (i.e., where the total area is $(1-X_{t_{init}})T^*$, given $X_{t_{init}} = 0$):

$$R''(\bar{X}, T) = T^* - \bar{X}T \quad \bar{X} \in [0,1], T \geq 0 \quad (4)$$

Given this adjustment, the system with more loss will always be the one with a lower level of resilience since $\bar{X}^{(1)}T^{(1)} > \bar{X}^{(2)}T^{(2)} \Rightarrow T^* - \bar{X}^{(1)}T^{(1)} < T^* - \bar{X}^{(2)}T^{(2)}$ for any value of T^* .

However, because resilience, as we have defined it, represents the ability to resist and recover from a disaster event, it does not make sense for a resilience measure to take on negative values. This means that the value of T^* must always be greater than $\bar{X}T$ for any choice of $\bar{X}$ and T . Choosing T^* to be equal to the maximum (observed or actual) value of T guarantees that it will always be large enough, but smaller values also may satisfy this constraint because of the moderating effect of the average loss. Our generalized measure of resilience that was presented in equation (2) is then simply a normalization of equation (4), achieved by dividing R'' by T^* . This not only restricts the measure to the $[0,1]$ interval but also serves to make resilience unitless, as is customary in practice.

In a broader sense, although this discussion demonstrates that a shared T^* value is important for consistency across a single decision-making context, different T^* values may be appropriate in different contexts, and the actual value of T^* can be chosen to meet a decision maker's needs, as long as it satisfies the constraint discussed above. As T^* is increased, the associated calculated resilience values will also increase, although their relative ranking will stay the same. Having the flexibility to change the value of T^* gives an opportunity to fine-tune the actual resilience values to better match the decision maker's perceptions about what they should be, and the result is similar to that of adjusting the α parameter discussed in [12]. Having this flexibility also allows for analyzing hypothetical situations, such

as a “black swan” event that results in a substantially longer recovery time than was previously observed for a given system. Although increasing T^* to enable calculating the system’s resilience to such an event will also increase the relative resilience of that system to all previously analyzed events, the rescaled resilience values will still be comparable, even in the new context, because their relative rank ordering will remain the same.

3.2. T^* for different types of systems

It is important to recognize that even given the ability to adjust T^* to match a decision maker's perceptions about system behavior, the calculated resilience value from equation (2) still might not appropriately reflect the actual performance of a given system if the context is too broad. For example, if the same T^* is used to compare systems that typically take months to recover, such as physical infrastructure elements, against systems that typically take minutes or hours to recover, such as communications networks or power grids, then the systems with the much shorter recovery times will tend to have significantly higher resilience values, simply because they suffer far less relative loss. It is reasonable to expect, however, that one would want to be able to talk about the relative resilience of buildings with respect to each other and the relative resilience of computer networks with respect to each other, and to have a system that measures 95% resilience, for example, in a consistent manner in each case. This implies that we must allow for assigning a *different* T^* value to each type of system.

In order to equate the resilience level of multiple types of systems in the context of the same disruptive event, we therefore must slightly adjust the formulation for resilience that was introduced in equation (2). We accomplish this by simply letting T_c^* be a fixed choice of T^* for a given class, c , of systems that each have an equivalent range of recovery times. We may then rewrite the resilience formula as follows, for any given class c :

$$R_c(\bar{X}, T) = \frac{T_c^* - \bar{X}T}{T_c^*} = 1 - \frac{\bar{X}T}{T_c^*} \quad \bar{X} \in [0,1], T \in [0, T_{max}], T_c^* \geq \bar{X}T_{max} \quad (5)$$

where T_{max} is the largest value of T that is of interest. This allows resilience to be compared consistently both within a given class of systems, as well as between different classes of systems, where each class has its own corresponding T_c^* value. The resulting resilience values will always be unitless and measured on the interval $[0, 1]$, regardless of whether they refer to a physical infrastructure component, or a communication network, or a socioeconomic system measure.

Different individual decision makers or organizations may each choose different T_c^* values for the same class of systems, in order to reflect their varying perceptions about what, for example, 95% resilience actually means. The key, however, is to ensure that the perceived resilience assessment is internally consistent for all relevant decision makers, across all systems that are assigned to that particular class.

4. Defining X^*

Our discussion so far has been representing loss and recovery time using two different types of measure. The average loss, $\bar{X}$, measures the unitless *percentage* of system functionality lost, whereas the recovery time, T , measures the number of time units elapsed until recovery. If we select $T^* \geq T_{max}$, however, then the ratio T/T^* represents the corresponding percentage of total recovery time taken, relative to T^* . Equation (2) (and, by extension, equation (5)) thus allows us to instead graph $\bar{X}$ against the ratio, T/T^* , each defined on the interval $[0,1]$. This, in turn, makes it straightforward to compare on the same plot the relative resilience of systems with different T^* values. What this also does, however, is suggest the possibility of further incorporating a more general analogous upper bound for the loss, X^* , into the resilience equation.

With this in mind, let X^* represent a fixed amount of loss against which the average loss of different systems can be compared, such that $X^* \geq \bar{X}_{max}$. Equation (2) then becomes:

$$R(\bar{X}, T) = \frac{X^*T^* - \bar{X}T}{X^*T^*} = 1 - \frac{\bar{X}T}{X^*T^*} \quad \bar{X} \in [0, X^*], T \in [0, T^*] \quad (6)$$

If $X^* = 1$, representing 100% loss of system functionality as a worst-case scenario, then equation (6) reduces to the original equation (2) and equation (5). Assuming here that both X^* and T^* are upper bounds on their respective variables simplifies both the implementation and interpretation of the

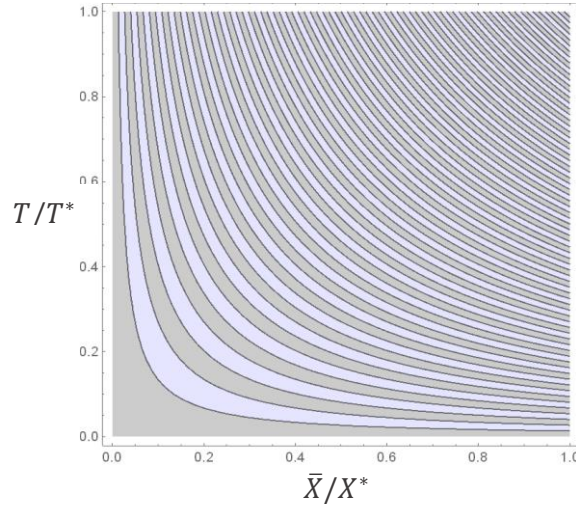


Fig 4. Generalized predicted resilience curves

resilience measure. In general, however, this assumption could be relaxed as long as the product of X^* and T^* is greater than the largest possible value of $\bar{X}T$. We may now adjust the resilience curves from Figure 3 to incorporate the addition of X^* into the resilience equation, along with the ratio between T and T^* , giving a more generalized version of the original resilience curves (See Figure 4).

We can also further generalize both equation (6) and Figure 4 by incorporating T_c^* and then introducing a corresponding X_c^* , in order to reflect that a given class of systems may share the same X^* value. This provides the following:

$$R_c(\bar{X}, T) = 1 - \frac{\bar{X}T}{X_c^*T_c^*} \quad \bar{X} \in [0, X_c^*], T \in [0, T_c^*] \quad (7)$$

for each class of systems, $c \in C$. Due to its generality, this can also be applied to calculating the resilience of different dimensions of the same complex system.

In our discussion above, we referred to the maximum tolerable period of disruption (MTPD) from the business continuity management (BCM) literature as being similar to the T^* parameter. The BCM literature also discusses the concept of the minimum business continuity objective (MBCO), which is more closely related to the X^* parameter but plays a slightly different role. Whereas X^* would typically represent an upper bound on the total amount of loss possible, so that no individual loss values are expected to exceed it, the MBCO is explicitly defined to be a lower bound on the range of “acceptable” operating conditions for a given organization [46]. The MTPD is then an upper bound on the length of time that the system’s operations remain at an unacceptable level below the MBCO [46] (See Figure 5).

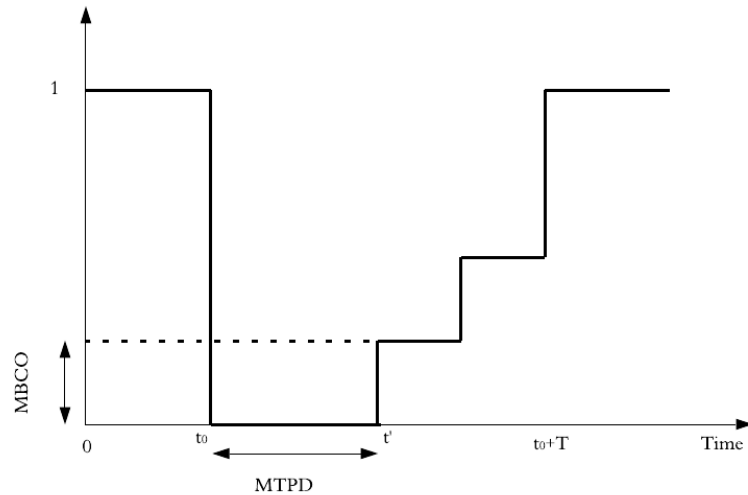


Fig 5. Graphical representation of MTPD and MBCO (adapted from [46])

4.1. Justifying X^*

The importance of X^* , from a practical perspective, can be illustrated by recognizing that a system can effectively fail well before it loses all functionality. For example, if a communications network has guaranteed its customers a certain minimum amount of throughput, then this minimum value may become the threshold against which acceptable levels of functionality are compared, rather than a throughput of zero. Similarly, if an automobile is damaged in an accident, the level of damage above which the vehicle would be replaced instead of repaired is almost certainly less than 100%. In each of these cases, defining X^* to be a value less than 100% allows for calculating the relative resilience of different systems with respect to their actual intended behavior and not some arbitrary notion of "total" loss.

Perhaps more importantly, however, the inclusion of X^* as an upper bound on $\bar{X}$ more effectively supports the application of the resilience curves to other contexts, particularly situations in which a disruption in the original system is exhibited as a percentage *gain* in "functionality" of $\bar{X}$, relative to X^* , rather than a percentage loss of $\bar{X}$. For example, a surge in emergency room visits would require a hospital to exhibit resilient behavior in order to respond to the resulting disruption in normal operations. In this case, the ratio of the two parameters allows for them both to take on negative values without changing the range of $R(\bar{X}, T) \in [0, 1]$ or necessitating a new approach to graphing the tradeoffs between loss and recovery time.

Practically, one can elicit X^* and T^* from a decision maker in a manner similar to how one assesses a value or utility function for a single attribute in multicriteria decision making [59]. For a given attribute, the decision maker is asked to select a level that represents the truly unsatisfactory point at which that attribute provides no value. That unsatisfactory point is then assigned a value or utility of 0, and any level of the attribute worse than that point also receives a value or utility of 0 [60]. Similarly, in the context of assessing resilience, a decision maker can set X^* to be the level of *performance* that is truly unsatisfactory or unacceptable and for which they effectively receive no value. A decision maker can

then also identify the length of recovery time, T^* , which is truly unsatisfactory and beyond which any other recovery time is also unsatisfactory. Repeating this exercise for each class of systems allows us to compare resilience among those different systems because the product X^*T^* reflects an equally unsatisfactory situation in each case, within the mind of the decision maker.

5. Application / Case Study

To demonstrate the generality of the proposed method for calculating the resilience of a system, we now provide three different examples that each have unique characteristics. These three examples show that the proposed approach can be applied to the calculation of resilience in a variety of settings.

Among all of the natural hazards that have significantly affected New York City over the last 20 years, Superstorm Sandy, which made landfall on October 29, 2012, had one of the most significant impacts. In order to illustrate the benefit of using X^* and T^* to adopt specific frames of reference and define the relative resilience of different systems, we will look at and compare several different aspects of the impacts of Sandy on New York City. In the examples below, we choose values for T^* and X^* that are based on historical data, either directly or indirectly, in order to establish a solid initial foundation for comparison. As previously discussed, and as revisited below, these values could be adjusted, if necessary, in order to better fit a decision maker's preferences about the most appropriate frame of reference to be used.

5.1. Electrical power outages

As a result of Superstorm Sandy, the Con Edison electrical power infrastructure network was disrupted for nearly two weeks across different parts of the city [7]. As the primary electrical service provider in New York City, Con Edison served a total of approximately 3.3 million customers, and at one point during the storm more than 750,000 of these customers were without electric power. Table 1 gives the

reported outages for each day, beginning on the day before the storm made landfall and continuing until full recovery was achieved, based on press releases issued on those days.

According to Kenward et al. [61], a power outage can be formally defined as a large electrical disturbance during which at least 50,000 customers lose power for at least an hour. For the sake of our analysis, we consider smaller disturbances to fall within the range of “normal” operating conditions (i.e., satisfying intended system performance). Setting a lower threshold for loss such as this is analogous to defining a minimum business continuity objective (MBCO) value, as discussed above, and it allows us to more clearly identify and characterize the extreme conditions due to the storm. The resulting net loss data is given in the last column of Table 1.

Table 1

Con Edison customers without power.

<i>Date</i>	<i>Customers without power</i>	<i>Amount above 50,000</i>
10/28/2012	<i>No outages reported</i>	0
10/29/2012	68,700	18,700
10/30/2012	779,000	729,000
10/31/2012	720,000	670,000
11/1/2012	674,000	624,000
11/2/2012	570,000	520,000
11/3/2012	280,800	230,800
11/4/2012	177,000	127,000
11/5/2012	156,800	106,800
11/6/2012	117,900	67,900
11/7/2012	70,050	20,050
11/8/2012	65,000	15,000
11/9/2012	28,300	0
11/10/2012	19,637	0
11/11/2012	<i>Specific numbers not available</i>	0
11/12/2012	<i>No outages reported</i>	0

In 1977, the entire electrical grid for Con Edison in New York City had lost power for up to 25 hours in different parts of the city [62]. In response, many improvements were made to the system, including efforts to improve system readiness for emergencies, efforts to improve command and control

capabilities, and revisions of the overall system design [62]. Although the subsequent disruption due to Superstorm Sandy ultimately affected fewer customers overall, the outages due to Sandy lasted for a much longer period of time. Since the outage in 1977 was effectively resolved within a single day, the total system-wide loss that it represented does not necessarily provide a good threshold against which to compare the Sandy loss. Instead, because the loss due to Sandy occurred after changes were made post-1977, we take the Sandy results as the basis against which to compare any future losses.

We may therefore set X^* to be equal to the maximum daily loss observed during that storm: 729,000 (net beyond the 50,000 threshold). Since the last day before the first instances of power loss

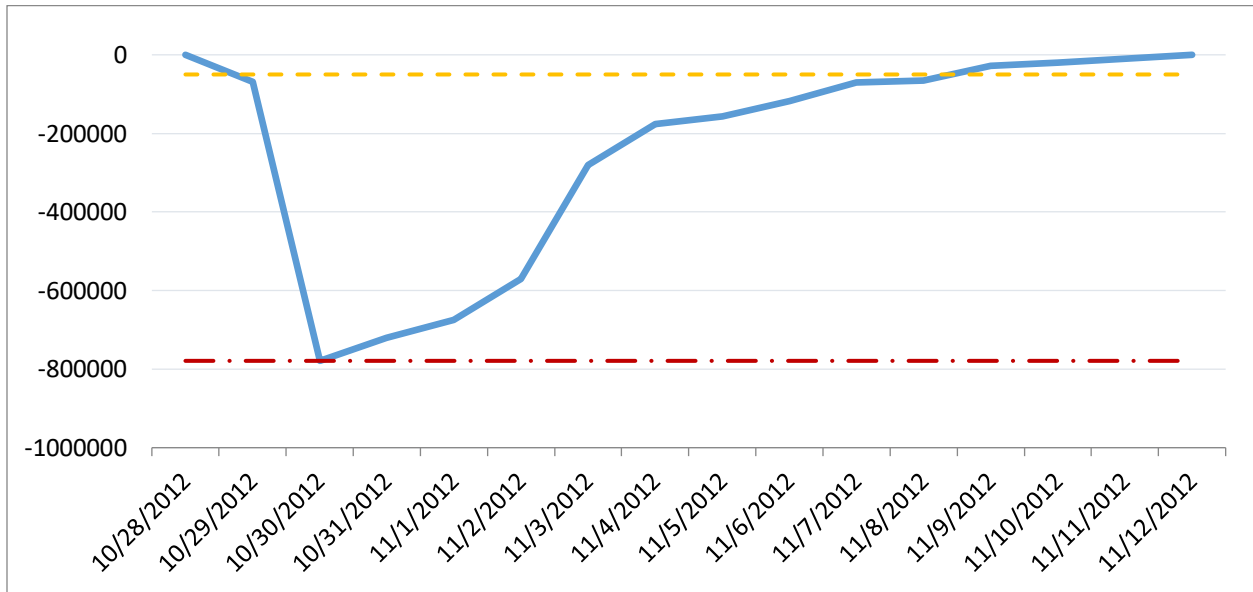


Fig 5. Electrical power functionality curve for Superstorm Sandy

were reported was October 28th and it was reported that there were no more instances of loss on November 12th, we set $T^* = 16$. Applying these thresholds to the data in Table 1 and taking the average loss from October 29th to November 8th generates a resilience value of $R = 0.732$, where $\bar{X} = 284,477$ and $T = 11$. The corresponding response curve representing the loss and recovery of electrical power is given in Figure 5.

Had X^* been set equal to the loss of Con Edison's entire customer base in New York City (3.3M), the measured resilience would have been 0.941. As discussed in [12], however, even though the given resilience measure is a relative and not an absolute value, it is important to have it provide a reasonable indicator of the amount of loss suffered over time, relative to the worst case scenario. Using this larger resilience value could be considered more beneficial to the power company because it implies a robust response, but it is somewhat misleading since it does not reflect the relative significance of the impact of the event. Setting X^* to the maximum historical loss value better captures the impact of the storm relative to past behavior and it more clearly shows the opportunity to improve the future response.

The difference between these two scenarios clearly illustrates the importance of explaining and justifying the choice of X^* and T^* , particularly if there is no standardized approach for determining their values in a given situation. In order for a resilience measurement to be validated by someone other than the original decision maker, the context within which it was originally calculated must be made clear. This information can help to resolve any differences in the interpretation of the results.

5.2. Socio-economic dimensions

In the case of Con Edison's electrical network, 100% functionality was defined to be any service level where fewer than 50,000 customers in the network were without power. Resilience was then measured with respect to the average percent of functionality lost over time, relative to the difference between this level and the X^* threshold. The following two examples provide a very different, yet still relevant, context within which to analyze resilience, and help to illustrate that the approach is broadly applicable. Both of the examples discussed below are based on using the number of non-emergency 311 service requests over time as an indicator of different aspects of the community's resilience to the storm [13,63]. All 311 requests received in New York City since 2004 are available through its Open Data initiative (<https://nycopendata.socrata.com/>), and the data consists of highly structured multi-

dimensional records that represent more than 100 different types of complaints assigned to specific agencies for processing (See [63] for more details).

We focus here on two specific types of complaints that can be used to represent aspects of the city's resilience from a socio-economic perspective: *Street Light Condition*, which serves as an indicator of the mobility of the population, and *Damaged Tree*, which provide an indication of the population's concerns about both safety and quality of life. For both types of complaint, the individual request data from the three years preceding Superstorm Sandy (2009-2011) was used to generate a 95% prediction interval for the number of requests the City would have received if the storm had not actually occurred. See [63] for more details on calculating the 95% prediction interval. In each case, the deviation from "100%" performance is then measured with respect to the extent to which the number of complaints either exceeded the upper bound of that interval or fell short of its lower bound. The same time frame was chosen for analysis as was used for the power loss data: October 28th to November 12th, 2012.

5.2.1. *Street Light Condition*

With respect to the *Street Light Condition* complaint type, the number of service requests fell within the predicted range except for a few days during the actual storm when it dropped below the lower bound of the prediction interval (See Figure 6). In this case, falling short of the lower bound of the prediction interval doesn't imply loss of functionality in the same way as the loss of electrical power above. Instead, it likely reflects a change in the population's mobility as a result of the disruption. With more people staying home as a result of the storm and its after effects, one would expect fewer complaints about issues such as street lights being out. This decrease in the number of complaints serves as an indicator that the storm had an impact not just on the physical infrastructure of the City but also on the City's population, through its interaction with that infrastructure.

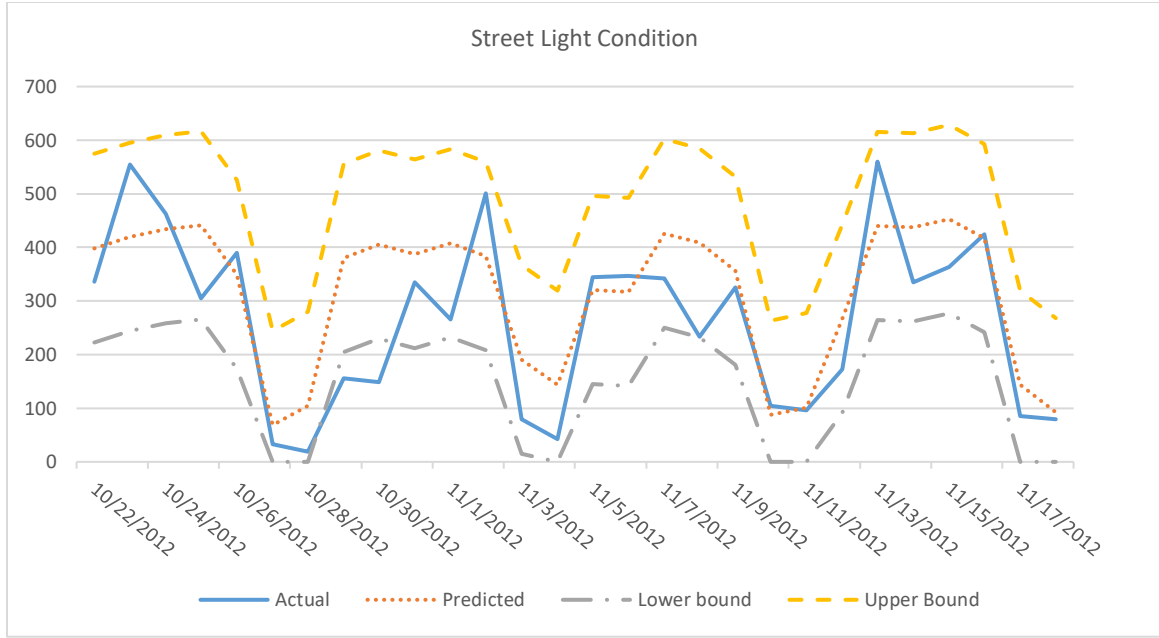


Fig 6. Prediction interval and number of actual requests for *Street Light Condition* indicator

Because any observations that fall within the prediction interval can be considered to correspond to our notion of 100% functionality (i.e., desired performance), we set the lower bound of the prediction interval to represent 100% functionality. In this case, the X^* threshold is set equal to the lower bound for any given day. The *actual* value of X^* thus varies with t , and we may represent it as X_t^* . The value of X_t is defined as the amount by which the lower bound exceeds the actual number of requests received and $X_t = 0$ when the number of requests is larger than the lower bound. The parameter T is the total number of time periods when the number of requests received is less than the lower bound or when $X_t > 0$.

We may address the issue of having such a variable X^* value by recognizing that the ratio in the resilience equation (6) is equivalent to dividing the total area above the performance curve – the actual loss ($\bar{X}T$) – by the total area below the performance curve if no disruption had occurred – the total possible loss (X^*T^*). In our current example, however, since X^* varies over time, the area associated

with the total possible loss may instead be written as $\bar{X}^*T^*$, where $\bar{X}^*$ is the *average* value of the X_t^* threshold from time 0 to time T^* . This reduces to X^*T^* when X^* is constant. Similarly, $\bar{X}$ is the average value of X_t when $X_t > 0$ or $\bar{X} = (\sum_{\forall t, X_t > 0} X_t)/T$. Thus we may generalize equation (6) substituting in the mean value for X^* :

$$R(\bar{X}, T) = \frac{\bar{X}^*T^* - \bar{X}T}{\bar{X}^*T^*} = 1 - \frac{\bar{X}T}{\bar{X}^*T^*} \quad \bar{X} \in [0, \bar{X}^*], T \in [0, T^*] \quad (8)$$

Equation (8) provides a method to calculate resilience when X^* changes over time that is consistent with the original concept that resilience is measured as the ratio of area above the actual performance curve to the area of the total possible loss. The resilience value for *Street Light Condition* is thus $R = 0.940$, with $\bar{X} = 64.31$ and $\bar{X}^* = 133.92$, and with $T = 2$ and $T^* = 16$.

5.2.2. Damaged Tree

The first two examples were similar in that they both represented *loss* of "functionality", which is the typical measure used to analyze the resilience of both infrastructure and socio-economic systems. In this next example, however, there is a positive deviation from expected or desired performance instead. As illustrated in Figure 7, the number of complaints about damaged trees during and after Superstorm Sandy greatly exceeded the bounds of the prediction interval for a number of days during and after the storm. Being able to specify a value for X^* (or $\bar{X}^*$) becomes critically important for calculating resilience in this case because otherwise there is no upper bound for this type of deviation.

Although X^* could simply be chosen to be related to the maximum possible number of requests that could be received, this is difficult to estimate and would most likely lead to an artificially high resilience value, even under extremely severe conditions. Instead, therefore, we initially adopt the approach taken in the Con Edison example and define X^* based on the maximum number of requests made on any given day within the entire 3-year data range (4251 requests). As in the Con Edison example, this

maximum number of requests was actually achieved during Superstorm Sandy. The corresponding X^* value for any given day is then equal to the difference between this constant value (4251 requests) and the upper bound of the prediction interval on that day, and X is the difference between the actual

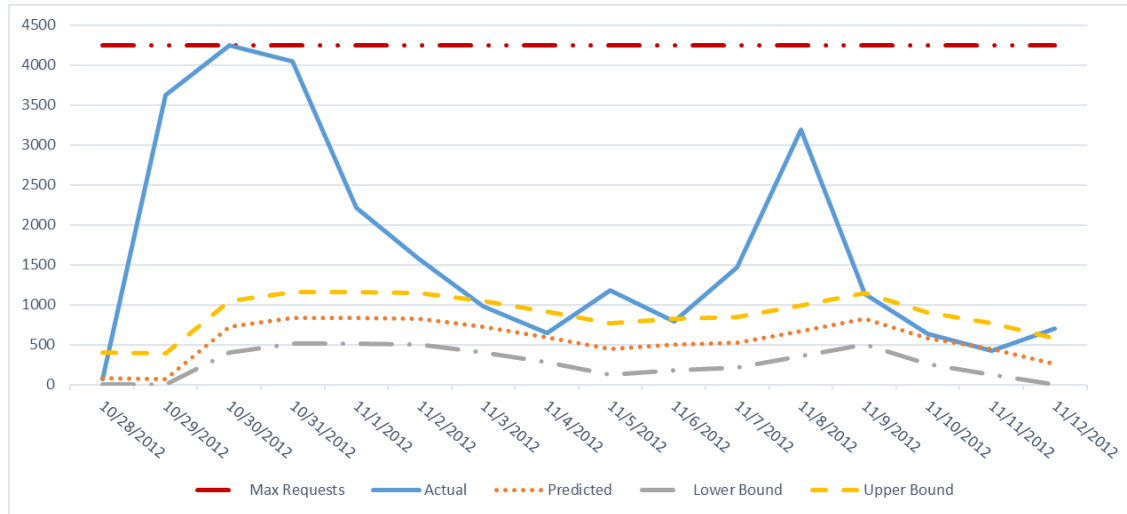


Fig 7. Prediction interval and number of actual requests for *Damaged Tree* indicator – variable X^*

number of requests received and this same upper bound. If the number of requests is less than the upper bound, then $X = 0$ for that day. Analyzing the *Damaged Tree* data using this approach gives $\bar{X} = 1573.87$ and $\bar{X}^* = 3369.5$, with $T = 9$ and $T^* = 16$, resulting in a resilience value of $R = 0.737$.

Considering this result, however, it is important to recognize that choosing to define X^* relative to a fixed maximum value effectively implies that the maximum possible number of "additional" requests

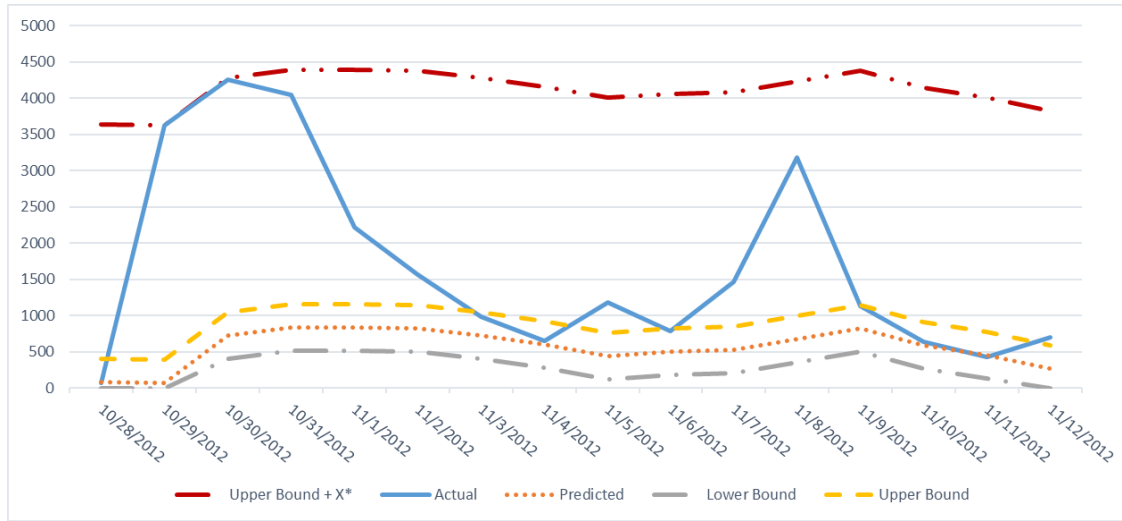


Fig 8. Prediction interval and number of actual requests for *Damaged Tree* indicator – constant X^*

above the upper limit of the prediction interval will be less when the predicted number of requests is greater, and more when the predicted number calls is smaller. This does not necessarily align with behavior that one might actually expect to see. Instead, one could make the case that if the additional requests are associated specifically with the event causing the disruption, then they should be relatively independent of the number of requests that one would normally expect to see. This implies that it would be reasonable to set the maximum number of additional requests so that it is constant over time.

With this in mind, therefore, we recalculate the resilience by instead setting $X^* = 3236.71$ across all time intervals, which provides a constant X^* but a variable actual upper bound (See Figure 8). This value is simply the difference between the number of requests (4251) and the upper bound of the prediction interval (1047.54) on the day that had the most requests (10/30/2012). This gives $\bar{X}^* = X^* = 3236.71$, with $\bar{X}$, T , and T^* as before, and a new resilience value of $R = 0.726$. This second approach produces a resilience value that is similar to that of the initial approach, but it is easier to justify by putting X and X^* on the same scale.

6. Discussion

Focusing the calculation of the resilience measure on the ratios of $\bar{X}/X^*$ and T/T^* , rather than just on $\bar{X}$ and T , provides a significant advantage because it allows for simultaneously comparing different resilience outcomes, or even different systems, on the same set of resilience curves (See Figure 4), regardless of whether or not they share the same X^* and/or T^* values. In order to compare such resilience values across different systems, the decision maker must simply ensure that each system has a single, appropriate frame of reference defined by a specific X^* and T^* . Different systems may share exactly the same threshold parameter values, as in the case of T^* above, or the decision maker may adopt different parameter values to define each frame of reference, as in the case of X^* above.

With this in mind, Table 2 summarizes the resilience values and parameters from the three examples above, each of which can be thought of as representing a different dimension of New York City's overall resilience behavior (out of many possible dimensions) in response to Hurricane Sandy. Figure 9 then illustrates the relative tradeoffs between loss and recovery time in each case.

For each dimension presented in Table 2, the T^* value represents the maximum recovery time and the X^* value represents the maximum loss. This consistent interpretation allows us to compare the individual dimensions in terms of the extent to which they were each impacted by the storm. The results thus show that all three dimensions are somewhat similar in their (normalized) average loss values, over the course of the storm, but very different in the (normalized) length of time that it took them to recover.

Table 2 shows that the Power Network dimension suffered more relative loss than the Damaged Tree dimension but also that it recovered more quickly. This is likely related to the relatively greater importance of the network and the need to restore power as quickly as possible. In contrast, the relative increase in the number of complaints about damaged trees was less extreme. These complaints persisted for a longer period of time, however, perhaps because resources were diverted to help with

recovery operations in more critical dimensions. In the end, the extended recovery time in the Damaged Tree dimension was significant enough that the total amount of disruption (relative to its specific frame of reference) was equivalent to that experienced by the power network.

Table 2

Resilience values for multiple dimensions in NYC, for Superstorm Sandy.

<i>Dimension</i>	$\bar{X}$	X^*	T	T^*	$\bar{X}/X^*$	T/T^*	R
Electrical Power Network	-284477	-729000	11	16	0.390	0.688	0.732
Street Light Condition	-64.3	-133.9 ⁽¹⁾	2	16	0.480	0.125	0.940
Damaged Tree	944.3	3236.7	15	16	0.292	0.938	0.726

⁽¹⁾ Average X^* over T^* interval

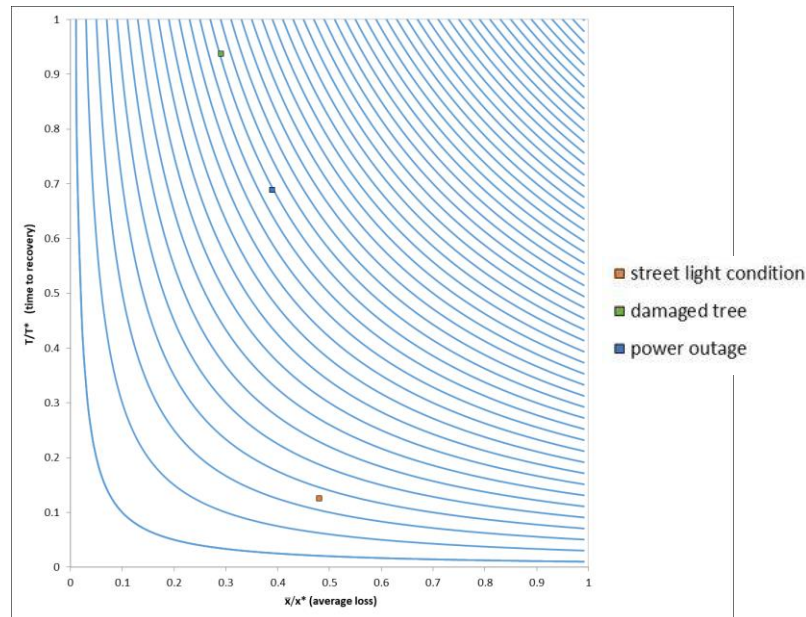


Fig 9. Relative resilience values

Table 2 also indicates that the Street Light Condition dimension has a much higher resilience value than the other two dimensions because of its similar average loss value and relatively short recovery time. Even though its (normalized) average loss is more than 20% higher than that of both of the other dimensions, the number of street light complaints exceeded the prediction interval for only 2 out of the 16 days. Thus, its T/T^* ratio is much smaller than the Power Network and Damaged Tree. Similar to the

Power Loss dimension, Street Lights are a critical functionality, particularly in a very urban area such as New York City. A decrease in the number of complaints could indicate a drop in the overall mobility of the population due to the effects of the storm. It could also signal a shift in the population's priorities towards other issues, such as dealing with the loss of power at home or cleaning up damaged property. That the number of calls increased back into the normal range after such a short period of time indicates that the daily activities of the population were able to resume relatively quickly.

Being able to directly compare these different dimensions of New York City's response to Hurricane Sandy allows for decision makers to develop a better understanding of the complexity of the disaster's overall impacts. By clarifying the relationship between loss and recovery time within each of the different dimensions, the City can also identify the need to invest relatively more resources into either recovery activities or loss prevention activities, in each case. Furthermore, the results could provide support for comparing the relative effectiveness of such investments *across* those dimensions, and potentially help the decision makers to identify the need to shift resources from one aspect of the City's overall disaster response efforts to another. A more complete analysis of New York City's overall resilience to a disaster such as Hurricane Sandy would, of course, require more than three data points – this set of examples is merely intended to illustrate the possibilities for supporting improved decision making. For this reason, it is important to recognize that this particular approach to measuring resilience can also be used to compare the relative resilience of not only other systems or dimensions but also multiple simulated outcomes from the same system. Given an appropriate simulation model, it thus becomes possible to analyze the distribution of a large number of potential outcomes associated with the probabilistic impacts of a disaster event across many different contexts.

From a practical perspective, although one could calculate each of the parameter values ($\bar{X}$, X^* , T , T^*) separately, if one has discrete empirical data with which one wishes to calculate resilience, such as in

our examples above, the easiest approach to calculating resilience is to sum both the individual actual deviations and the total possible deviations, and then subtract the ratio of those two values from 1:

$$R(\bar{X}, T) = 1 - \frac{\bar{X}T}{\bar{X}^*T^*} = 1 - \frac{(\text{total actual deviation})}{(\text{total possible deviation})} = 1 - \frac{\sum_{t=1}^{T^*} X_t}{\sum_{t=1}^{T^*} X_t^*} \quad (9)$$

This makes it easier, in particular, to deal with situations where the actual value of X^* varies over time, and it corresponds exactly to measuring the area beneath the (generalized) resilience triangle.

Given our emphasis on the flexibility that a decision maker has in choosing X^* and T^* , it is important to recognize that different decision makers may choose different values for these threshold parameters in a given context. Because of this, it is very important for each decision maker not only to be able to explicitly specify the frame of reference (X^* , T^*) within which a resilience value has been calculated, but also to be able to justify why those particular parameter values were chosen. If a standardized approach to specifying those values is adopted by all parties, such as adopting a consistent definition of “maximum recovery time” or setting X^* equal to 100% functionality, then comparing the resilience assessments of different decision makers can be relatively straightforward. If there is a difference in the approach taken to define the different X^* and T^* values, however, then each decision maker may have to explicitly justify the different frames of reference that he or she adopted. It then may be necessary, potentially after further discussion, to adjust some of the parameter thresholds in order for the group as a whole to agree on the validity of the generated resilience values.

7. Conclusions

This paper has extended previous work by demonstrating the importance and benefit of using explicit threshold values for recovery time (T^*) and functional loss (X^*) to define an appropriate frame of reference for calculating resilience. Although these parameters can theoretically take on a range of values, setting them at or above the expected maximum values for loss and recovery time can

significantly help with the interpretability of the results. There may be situations in which the specific choices for X^* and T^* can be determined objectively, however a decision maker can also choose subjective values for them in order to better reflect his or her interpretation of the resilience outcomes. Even if the threshold parameter values are chosen objectively (e.g., setting T^* equal to the maximum historical recovery time), those values may need to be adjusted as conditions change and as more information becomes available about the range of outcomes that could occur. Any such adjustment will define a new frame of reference upon which to compare any new outcomes against the existing ones, but the relative ranking of those existing values will be maintained.

The addition of the X^* parameter, in particular, allows the original concept of predicted resilience to be further generalized so that very different resilience indicators may be compared on a consistent basis, even if they operate on very different loss scales. The inclusion of the T^* parameter provides the same benefit for systems with varying recovery time scales. Taken together, they provide a significant step towards more fully representing the complex, multi-dimensional nature of a real-world system that has been disrupted by a disaster event. This allows decision makers such as community leaders, emergency managers, and infrastructure owners and operators, to directly compare resilience behaviors across different dimensions, so that they can describe complex system behaviors more effectively. From a decision-making standpoint, this supports not just descriptive analyses of past events that can be used to inform future infrastructure management decisions but also prescriptive analyses, based on simulating future disruptions, that involve comparing the relative effectiveness of different resilience investments.

Particularly because there can be flexibility in a decision maker's choice of X^* and T^* , it is important to recognize that assigning both of them very large values will result in any corresponding calculated resilience values being relatively close to 1. This subsequently implies consistently resilient behavior. On the other hand, if the same T^* and X^* are assigned very small values then the corresponding

resilience values may instead be very close to 0, even though their actual measured loss and recovery time values haven't changed. This illustrates the importance of understanding the context behind the choice of values for these threshold parameters. Even if it is possible to choose a wide range of values for the parameters it is critical to be able to justify the resulting frame of reference as not only valid but meaningful, otherwise any associated resilience measurements will have questionable utility and may even be misleading to someone else. A clear explanation and justification of the context of any calculated resilience value is thus extremely important, particularly if it is going to be used in comparing resilience behaviors across different systems or dimensions.

The approach presented in this paper can be extended in several ways. First of all, in this study we assumed that the disrupted system recovers to its original level after a period of time. Some disruptions, however, may have a permanent impact on a system and full recovery may not be possible. In this case, the system will recover to a level lower than its original level. The area under the curve can still be measured in this situation, however, even if the response curve never returns to its original value, because T^* is not restricted to representing a maximum possible recovery time. Future research should therefore extend the resilience measure presented in this paper to account for such permanent loss.

It is also important to recognize that complex systems, such as organizations or communities, have several different performance dimensions. The resilience measure discussed above can be applied to evaluate the performance of each dimension separately. However, there still remains opportunity to look at different ways to combine these measures, as in [13], and thus to provide an improved overall assessment of resilience performance. Since the resilience calculation within a single dimension can reflect the subjective perception of that particular aspect of the overall system response to a disruption, a weighted average that appropriately combined the results of multiple dimensions into a single value could also reflect the subjective resilience perception of an entire system or organization.

Acknowledgements

This work was supported by the National Science Foundation (NSF-CRISP #1541155 and NSF-NRT #1735139).

References

- [1] National Research Council, Disaster resilience: a national imperative, Washington, DC: The National Academies Press, 2012. <https://doi.org/10.17226/13457>.
- [2] UNISDR, Sendai Framework for Disaster Risk Reduction: 2015-2030, 2015. <https://www.undrr.org/publication/sendai-framework-disaster-risk-reduction-2015-2030>.
- [3] M. Bruneau, S. Chang, R. Eguchi, G.C. Lee, T.D. O'Rourke, A.M. Reinhorn, M. Shinozuka, K. Tierney, W.A. Wallace, D. von Winterfeldt, A framework to quantitatively assess and enhance the seismic resilience of communities, *Earthq. Spectra*. 19 (2003) 733–752. <http://earthquakespectra.org/doi/abs/10.1193/1.1623497> (accessed January 8, 2015).
- [4] M. Bruneau, A. Reinhorn, Exploring the concept of seismic resilience for acute care facilities, *Earthq Spectra*. 23 (2007). <https://doi.org/10.1193/1.2431396>.
- [5] G.P. Cimellaro, A.M. Reinhorn, M. Bruneau, Seismic resilience of a hospital system, *Struct Infrastruct Eng*. 6 (2010). <https://doi.org/10.1080/15732470802663847>.
- [6] D.M. Simpson, C.B. Lasley, T.D. Rockaway, T.A. Weigel, Understanding critical infrastructure failure: examining the experience of Biloxi and Gulfport, Mississippi after Hurricane Katrina, *Int. J. Crit. Infrastructures*. 6 (2010) 246–276. <https://doi.org/10.1504/IJCIS.2010.033339>.
- [7] C.W. Zobel, Quantitatively representing nonlinear disaster recovery, *Decis. Sci*. 45 (2014) 1053–1082. <https://doi.org/10.1111/dec.12103>.

- [8] S.L. Cutter, L. Barnes, M. Berry, C. Burton, E. Evans, E. Tate, J. Webb, A place-based model for understanding community resilience to natural disasters, *Glob. Environ. Chang.* 18 (2008) 598–606. <https://doi.org/10.1016/j.gloenvcha.2008.07.013>.
- [9] F.H. Norris, S.P. Stevens, B. Pfefferbaum, K.F. Wyche, R.L. Pfefferbaum, Community Resilience as a Metaphor, Theory, Set of Capacities, and Strategy for Disaster Readiness, *Am. J. Community Psychol.* 41 (2008) 127–150. <https://doi.org/10.1007/s10464-007-9156-6>.
- [10] S.L. Cutter, C.G. Burton, C.T. Emrich, Disaster Resilience Indicators for Benchmarking Baseline Conditions, *J. Homel. Secur. Emerg. Manag.* 7 (2010). <https://doi.org/10.2202/1547-7355.1732>.
- [11] C. Zobel, Comparative visualization of predicted disaster resilience, *Proc. 7th Int. ISCRAM Conf.* (2010) 1–6. <http://www.iscram.org/ISCRAM2010/Papers/191-Zobel.pdf> (accessed January 8, 2015).
- [12] C.W. Zobel, Representing perceived tradeoffs in defining disaster resilience, *Decis. Support Syst.* 50 (2011) 394–403. <https://doi.org/10.1016/j.dss.2010.10.001>.
- [13] C.W. Zobel, M. Baghersad, Analytically comparing disaster resilience across multiple dimensions, *Socioecon. Plann. Sci.* 69 (2020) 1–14. <https://doi.org/10.1016/j.seps.2018.12.005>.
- [14] S.L. Cutter, The landscape of disaster resilience indicators in the USA, *Nat. Hazards.* 80 (2016) 741–758. <https://doi.org/10.1007/s11069-015-1993-2>.
- [15] W.G. Peacock, S.D. Brody, W.A. Seitz, W.J. Merrell, A. Vedlitz, S. Zahran, R.C. Harriss, R. Stickney, Advancing resilience of coastal localities: Developing, implementing, and sustaining the use of coastal resilience indicators: A final report, 2010.
- [16] C.G. Burton, A Validation of Metrics for Community Resilience to Natural Hazards and Disasters Using the Recovery from Hurricane Katrina as a Case Study, *Ann. Assoc. Am. Geogr.* 105 (2015)

- 67–86. <https://doi.org/10.1080/00045608.2014.960039>.
- [17] L.A. Bakkensen, C. Fox-Lent, L.K. Read, I. Linkov, Validating Resilience and Vulnerability Indices in the Context of Natural Disasters, *Risk Anal.* 37 (2017) 982–1004.
<https://doi.org/10.1111/risa.12677>.
- [18] I. Noy, The macroeconomic consequences of disasters, *J. Dev. Econ.* 88 (2009) 221–231.
<https://doi.org/https://doi.org/10.1016/j.jdeveco.2008.02.005>.
- [19] M. Christopher, H. Peck, Building the resilient supply chain, *Int. J. Logist. Manag.* 15 (2004) 1–14.
<https://doi.org/10.1108/09574090410700275>.
- [20] E. Brandon-Jones, B. Squire, C.W. Autry, K.J. Petersen, A contingent resource-based perspective of supply chain resilience and robustness, *J. Supply Chain Manag.* 50 (2014) 55–73.
<https://doi.org/10.1111/jscm.12050>.
- [21] Y. Sheffi, *The Resilient Enterprise: Overcoming Vulnerability for Competitive Advantage*, 2005.
<https://ideas.repec.org/b/mtp/titles/0262693496.html>.
- [22] R. Dubey, A. Gunasekaran, S.J. Childe, T. Papadopoulos, C. Blome, Z. Luo, Antecedents of Resilient Supply Chains: An Empirical Study, *IEEE Trans. Eng. Manag.* 66 (2019) 8–19.
<https://doi.org/10.1109/TEM.2017.2723042>.
- [23] S. Hosseini, D. Ivanov, A new resilience measure for supply networks with the ripple effect considerations: a Bayesian network approach, *Ann. Oper. Res.* (2019).
<https://doi.org/10.1007/s10479-019-03350-8>.
- [24] T.J. Pettit, J. Fiksel, K.L. Croxton, Ensuring supply chain resilience: development of a conceptual framework, *J. Bus. Logist.* 31 (2010) 1–21. <https://doi.org/10.1002/j.2158-1592.2010.tb00125.x>.

- [25] J.B. Rice, F. Caniato, Building a secure and resilient supply network, *Supply Chain Manag. Rev.* 7 (2003) 22–30.
- [26] M. Baghersad, C.W. Zobel, Assessing the extended impacts of supply chain disruptions on firms: An empirical study, *Int. J. Prod. Econ.* 231 (2021) 107862.
<https://doi.org/https://doi.org/10.1016/j.ijpe.2020.107862>.
- [27] S. Elluru, H. Gupta, H. Kaur, S.P. Singh, Proactive and reactive models for disaster resilient supply chain, *Ann. Oper. Res.* 283 (2019) 199–224. <https://doi.org/10.1007/s10479-017-2681-2>.
- [28] S. Hosseini, D. Ivanov, A. Dolgui, Review of quantitative methods for supply chain resilience analysis, *Transp. Res. Part E Logist. Transp. Rev.* 125 (2019) 285–307.
<https://doi.org/https://doi.org/10.1016/j.tre.2019.03.001>.
- [29] Y. Li, C.W. Zobel, Exploring supply chain network resilience in the presence of the ripple effect, *Int. J. Prod. Econ.* 228 (2020) 107693.
<https://doi.org/https://doi.org/10.1016/j.ijpe.2020.107693>.
- [30] R. Mogre, S. Talluri, F. D’Amico, A Decision Framework to Mitigate Supply Chain Risks: An Application in the Offshore-Wind Industry, *IEEE Trans. Eng. Manag.* 63 (2016) 316–325.
<https://doi.org/10.1109/TEM.2016.2567539>.
- [31] A. Jabbarzadeh, B. Fahimnia, F. Sabouhi, Resilient and sustainable supply chain design: sustainability analysis under disruption risks, *Int. J. Prod. Res.* 56 (2018) 5945–5968.
<https://doi.org/10.1080/00207543.2018.1461950>.
- [32] S. Hosseini, K. Barker, J.E. Ramirez-Marquez, A review of definitions and measures of system resilience, *Reliab. Eng. Syst. Saf.* 145 (2016) 47–61. <https://doi.org/10.1016/j.res.2015.08.006>.
- [33] D. Henry, J.E. Ramirez-Marquez, Generic metrics and quantitative approaches for system

resilience as a function of time, *Reliab. Eng. Syst. Saf.* 99 (2012) 114–122.

<https://doi.org/10.1016/j.ress.2011.09.002>.

- [34] K. Barker, J.E. Ramirez-Marquez, C.M. Rocco, Resilience-based network component importance measures, *Reliab. Eng. Syst. Saf.* 117 (2013) 89–97.
<https://doi.org/https://doi.org/10.1016/j.ress.2013.03.012>.
- [35] C.A. MacKenzie, C. Hu, Decision making under uncertainty for design of resilient engineered systems, *Reliab. Eng. Syst. Saf.* 192 (2019) 106171.
<https://doi.org/https://doi.org/10.1016/j.ress.2018.05.020>.
- [36] R. Gahi, C.A. MacKenzie, C. Hu, Design optimization for resilience for risk-averse firms, *Comput. Ind. Eng.* 139 (2020) 106122. <https://doi.org/https://doi.org/10.1016/j.cie.2019.106122>.
- [37] B.M. Ayyub, Systems Resilience for Multihazard Environments: Definition, Metrics, and Valuation for Decision Making, *Risk Anal.* 34 (2014) 340–355. <https://doi.org/10.1111/risa.12093>.
- [38] R. Pant, K. Barker, J.E. Ramirez-Marquez, C.M. Rocco, Stochastic measures of resilience and their application to container terminals, *Comput. Ind. Eng.* 70 (2014) 183–194.
<https://doi.org/https://doi.org/10.1016/j.cie.2014.01.017>.
- [39] D.L. Alderson, G.G. Brown, W.M. Carlyle, Operational Models of Infrastructure Resilience, *Risk Anal.* 35 (2015) 562–586. <https://doi.org/10.1111/risa.12333>.
- [40] K. Tierney, M. Bruneau, Conceptualizing and Measuring Resilience: A Key to Disaster Loss Reduction, *TR News.* (2007) 14–18.
- [41] G.P. Cimellaro, A.M. Reinhorn, M. Bruneau, Framework for analytical quantification of disaster resilience, *Eng. Struct.* 32 (2010) 3639–3649. <https://doi.org/10.1016/j.engstruct.2010.08.008>.

- [42] S.E. Chang, M. Shinozuka, Measuring improvements in the disaster resilience of communities, *Earthq. Spectra*. 20 (2004) 739–755. <https://doi.org/10.1193/1.1775796>.
- [43] C.W. Zobel, L. Khansa, Quantifying Cyberinfrastructure Resilience against Multi-Event Attacks, *Decis. Sci.* 43 (2012) 687–710. <https://doi.org/10.1111/j.1540-5915.2012.00364.x>.
- [44] C.W. Zobel, L. Khansa, Characterizing multi-event disaster resilience, *Comput. Oper. Res.* 42 (2014) 83–94. <https://doi.org/10.1016/j.cor.2011.09.024>.
- [45] C.A. Mackenzie, C.W. Zobel, Allocating resources to enhance resilience, with application to superstorm Sandy and an electric utility, *Risk Anal.* 36 (2016) 847–862.
<https://doi.org/10.1111/risa.12479>.
- [46] S.A. Torabi, H. Rezaei Soufi, N. Sahebjamnia, A new framework for business impact analysis in business continuity management (with a case study), *Saf. Sci.* 68 (2014) 309–323.
<https://doi.org/https://doi.org/10.1016/j.ssci.2014.04.017>.
- [47] H. Hassel, A. Cedergren, Exploring the Conceptual Foundation of Continuity Management in the Context of Societal Safety, *Risk Anal.* 39 (2019) 1503–1519. <https://doi.org/10.1111/risa.13263>.
- [48] N. Sahebjamnia, S.A. Torabi, S.A. Mansouri, Integrated business continuity and disaster recovery planning: Towards organizational resilience, *Eur. J. Oper. Res.* 242 (2015) 261–273.
<https://doi.org/10.1016/j.ejor.2014.09.055>.
- [49] ISO, Security and resilience — Business continuity management systems — Requirements (ISO 22301:2019), 2019. <https://www.iso.org/standard/75106.html>.
- [50] M. Ouyang, L. Dueñas-Osorio, Multi-dimensional hurricane resilience assessment of electric power systems, *Struct. Saf.* 48 (2014) 15–24. <https://doi.org/10.1016/j.strusafe.2014.01.001>.

- [51] Y. Ouyang, Z. Wang, H. Yang, Facility location design under continuous traffic equilibrium, *Transp. Res. Part B Methodol.* 81 (2015) 18–33.
<https://doi.org/https://doi.org/10.1016/j.trb.2015.05.018>.
- [52] K. Adjetey-Bahun, B. Birregah, E. Châtelet, J.L. Planchet, A model to quantify the resilience of mass railway transportation systems, *Reliab. Eng. Syst. Saf.* 153 (2016) 1–14.
<https://doi.org/10.1016/j.ress.2016.03.015>.
- [53] S. Zhao, X. Liu, Y. Zhuo, Hybrid Hidden Markov Models for resilience metrics in a dynamic infrastructure system, *Reliab. Eng. Syst. Saf.* 164 (2017) 84–97.
<https://doi.org/https://doi.org/10.1016/j.ress.2017.02.009>.
- [54] R. Li, Q. Dong, C. Jin, R. Kang, A new resilience measure for supply chain networks, *Sustainability.* 9 (2017) 144.
- [55] C. Zhang, J. Kong, S.P. Simonovic, Restoration resource allocation model for enhancing resilience of interdependent infrastructure systems, *Saf. Sci.* 102 (2018) 169–177.
<https://doi.org/https://doi.org/10.1016/j.ssci.2017.10.014>.
- [56] J. Kong, S.P. Simonovic, C. Zhang, Sequential Hazards Resilience of Interdependent Infrastructure System: A Case Study of Greater Toronto Area Energy Infrastructure System, *Risk Anal.* 39 (2019) 1141–1168. <https://doi.org/10.1111/risa.13222>.
- [57] J. Gong, F. You, Resilient design and operations of process systems: Nonlinear adaptive robust optimization model and algorithm for resilience analysis and enhancement, *Comput. Chem. Eng.* 116 (2018) 231–252. <https://doi.org/https://doi.org/10.1016/j.compchemeng.2017.11.002>.
- [58] S. Zhao, F. You, Resilient supply chain design and operations with decision-dependent uncertainty using a data-driven robust optimization approach, *AIChE J.* 65 (2019) 1006–1021.

<https://doi.org/10.1002/aic.16513>.

- [59] R.L. Keeney, Value-focused thinking, Harvard University Press, 1996.
- [60] C.W. Kirkwood, Strategic decision making, Duxbury Press. 149 (1997).
- [61] A. Kenward, U. Raja, Blackout: Extreme Weather , Climate Change and Power Outages, Clim. Cent. (2014) 23.
- [62] Federal Energy Regulatory Commission, The Con Edison power failure of July 13 and 14, 1977, US Department of Energy Washington, DC, 1978. <https://doi.org/10.2172/6673953>.
- [63] C.W. Zobel, M. Baghersad, Y. Zhang, An Approach for Quantifying the Multidimensional Nature of Disaster Resilience in the Context of Municipal Service Provision, in: Urban Disaster Resil. Secur., 2018: pp. 239–259. https://doi.org/10.1007/978-3-319-68606-6_15.