



Sponsor: DoD
Dept. No.: J83G
Contract No.: W15P7T-13-A802/D070
Project No.: 0714D070-VA

The views, opinions, and/or findings contained in this report are those of The MITRE Corporation and should not be construed as official government position, policy, or decision unless so designated by other documentation.

Distribution Statement A:
Approved for Public Release; Distribution Unlimited. 14-0500

©2014 The MITRE Corporation.
All rights reserved.

McLean, VA

A Measurable Definition of Resiliency Using “Mission Risk” as a Metric

Scott Musman
Seli Agbolosu-Amison

February 2014

Report Documentation Page				Form Approved OMB No. 0704-0188	
Public reporting burden for the collection of information is estimated to average 1 hour per response, including the time for reviewing instructions, searching existing data sources, gathering and maintaining the data needed, and completing and reviewing the collection of information. Send comments regarding this burden estimate or any other aspect of this collection of information, including suggestions for reducing this burden, to Washington Headquarters Services, Directorate for Information Operations and Reports, 1215 Jefferson Davis Highway, Suite 1204, Arlington VA 22202-4302. Respondents should be aware that notwithstanding any other provision of law, no person shall be subject to a penalty for failing to comply with a collection of information if it does not display a currently valid OMB control number.					
1. REPORT DATE FEB 2014		2. REPORT TYPE		3. DATES COVERED 00-00-2014 to 00-00-2014	
4. TITLE AND SUBTITLE A Measurable Definition of Resiliency Using 'Mission Risk' as a Metric				5a. CONTRACT NUMBER	
				5b. GRANT NUMBER	
				5c. PROGRAM ELEMENT NUMBER	
6. AUTHOR(S)				5d. PROJECT NUMBER	
				5e. TASK NUMBER	
				5f. WORK UNIT NUMBER	
7. PERFORMING ORGANIZATION NAME(S) AND ADDRESS(ES) MITRE Corporation,7525 Colshire Drive,McLean,VA,22102-7539				8. PERFORMING ORGANIZATION REPORT NUMBER	
9. SPONSORING/MONITORING AGENCY NAME(S) AND ADDRESS(ES)				10. SPONSOR/MONITOR'S ACRONYM(S)	
				11. SPONSOR/MONITOR'S REPORT NUMBER(S)	
12. DISTRIBUTION/AVAILABILITY STATEMENT Approved for public release; distribution unlimited					
13. SUPPLEMENTARY NOTES					
14. ABSTRACT In the cyber world, there has been shift in mindset from trying to prevent attacks from occurring and succeeding to developing tools and techniques that can make systems resilient in the face of incidents. Unfortunately, progress in this area has been hampered by the fact that we lack concrete methods that allow us to evaluate when, and by how much, modifications to a system contribute to making it more resilient. Part of the problem is that the term 'resilience' itself lacks a clear definition that supports measurable metrics that would allow two like systems to be compared against each other, or would enable the measurement of how different resiliency techniques can improve a system's resiliency when they are applied. In this paper we will review and discuss the terminology and definitions that have been proposed and used for describing the terms 'resilience' and 'resiliency' with respect to cyber and other systems. Ultimately, we address the deficiencies of these previous definitions by choosing a definition for resilience that equates to the inverse of 'mission risk' that is adequately qualified by the context in which it applies. In selecting a measurement (or estimated measurement) based on risk as our resilience metric, we have chosen a resilience definition that is clearly defined measurable, and has a sound theoretical grounding. Our computable metric makes it possible to perform like-to-like systems comparisons that allow us to measure the resiliency of a system, and to use this measurement to evaluate how resiliency methods are able to improve the resiliency of a system.					
15. SUBJECT TERMS					
16. SECURITY CLASSIFICATION OF:			17. LIMITATION OF ABSTRACT Same as Report (SAR)	18. NUMBER OF PAGES 20	19a. NAME OF RESPONSIBLE PERSON
a. REPORT unclassified	b. ABSTRACT unclassified	c. THIS PAGE unclassified			

Abstract

In the cyber world, there has been shift in mindset from trying to prevent attacks from occurring and succeeding to developing tools and techniques that can make systems resilient in the face of incidents. Unfortunately, progress in this area has been hampered by the fact that we lack concrete methods that allow us to evaluate when, and by how much, modifications to a system contribute to making it more resilient. Part of the problem is that the term “resilience” itself lacks a clear definition that supports measurable metrics that would allow two like systems to be compared against each other, or would enable the measurement of how different resiliency techniques can improve a system’s resiliency when they are applied.

In this paper we will review and discuss the terminology and definitions that have been proposed and used for describing the terms “resilience” and “resiliency” with respect to cyber and other systems. Ultimately, we address the deficiencies of these previous definitions by choosing a definition for resilience that equates to the inverse of “mission risk” that is adequately qualified by the context in which it applies. In selecting a measurement (or estimated measurement) based on risk as our resilience metric, we have chosen a resilience definition that is clearly defined, measurable, and has a sound theoretical grounding. Our computable metric makes it possible to perform like-to-like systems comparisons that allow us to measure the resiliency of a system, and to use this measurement to evaluate how resiliency methods are able to improve the resiliency of a system.

Table of Contents

1	Introduction.....	1
2	Considerations for Measuring Resilience	2
3	A Definition of Resilience	9
4	Example of Estimating Resiliency.....	10
5	Comparison with Previous Work and Definitions	12
6	Discussion.....	14
7	Summary	16
8	References.....	17

List of Figures

Figure 1. Performance Curve Illustrating Aspects of Resilience.....	3
Figure 2. Performance Curves of Two Systems with a Minimum Acceptable Performance Level.....	4
Figure 3. Two Systems Where the Integral of the Curves Represents the Overall Performance ...	4
Figure 4. The Effects of Choosing a Different Time Window on Calculating Resiliency	5
Figure 5. Two Systems that Perform the Same in the Face of an Incident.....	6
Figure 6. Performance of the System With and Without an Incident	7
Figure 7. The Performance of Two Systems Given Two Different Incidents	8

1 Introduction

A number of different definitions of the term “resilience” exist. In our work trying to evaluate and then determine how to combine different resiliency techniques together, we have been frustrated by the qualitative nature of the existing resiliency definitions [Haimes, 2009; Wood, 2005]. While it is understood that developing and deploying systems that are resilient in the face of incidents and attacks is highly desirable, there is a need to know which techniques are the best ones to choose given the incident environment. In addition, there is a need to be able to compare the resiliency of different systems or system versions within the same incident environment. In this paper, we propose a definition of resiliency that supports its measurement and enables comparison of like systems (e.g., two different designs for an airplane).

In this paper, we will review and discuss the terminology and definitions that have been proposed and used for describing the term “resiliency”. In many of these cases, the focus has been on qualitative descriptions or on comparing resiliency to other terms such as robustness or survivability. Often, when there has been an attempt to compute quantitative metrics, there has been a tendency to overly simplify resilience and to focus on trying to compute one or more ordinal, scalar numbers as metrics (typically performance metrics) [Haimes, 2009; Vugrin et al., 2010]. This approach often ignores many of the considerations that must be taken into account to adequately describe the context where a resiliency metric might be valid. Ultimately, we address the deficiencies of these previous definitions by choosing a definition of resilience that equates to the inverse of “mission risk” [Haimes, 1991; Haimes, 2009; Wood, 2005; Ayubb, 2013]. By selecting a measurement (or estimated measurement) based on risk as our resilience metric, we have chosen a resilience definition that is clearly defined, measurable, and has a sound theoretical grounding. Since risk relies on both the likelihood of events occurring as well as changes in utility (value) when these events occur, we are provided with a computable metric. This metric will be fully qualified by the contextual characteristics or attributes that define the set of events, utility scores, system description (e.g., the system boundaries), and even the timeframes over which the metric can be considered valid. These attributes provide the considerations that should be used to clearly define resilience as a computable metric.

2 Considerations for Measuring Resilience

While we look for quantitative methods that allow us to evaluate the resiliency of a system (e.g., cyber systems and biological systems), we must first describe all of the different factors that make its measurement so complex. Probably the most important issue with formulating a workable definition is that there is often a tendency to oversimplify what is being defined. There have been many proposals for resiliency metrics, and many of them focus on performance. More importantly, many of them focus on implicitly useful metrics [Bodeau et al., 2012] that largely imply resilience (in the face of a number of implicit assumptions) but that may not always hold in general. As an example, consider the use of network performance metrics (throughput, latency, etc.) in a system that relies heavily on the exchange of information between networked components. Although these network performance metrics and the set of incidents over which they apply can be determined, it is easy to consider that in one mission context what might be important would be to ensure that some mission-critical data can be exchanged between networked components within a mission-dependent timeframe. In another mission context (e.g., video conferencing), it may be that maintaining the sustained throughput of the network is more important. So to consider system resiliency, we must be explicit about all of the factors, mission and otherwise, that contribute to the resiliency definition. We are calling this “mission risk” since it allows us to characterize how well the intended function of the system will be achieved in the face of actual or potential incidents. The rest of this section explores the different aspects of this problem.

System Boundaries Matter

An important factor in determining resilience is to consider the system boundary of any system for which you are trying to estimate the resilience. Consider estimating the resilience of a power plant that produces electricity. Given the power plant itself, it is possible to consider the various types of attacks and incidents to which it might be subject, and estimate the circumstances under which it can produce power, and/or how long it might take to recover from any incidents that might cause it to stop producing power for any period of time. Given this context, it is possible to compare whether one power plant design might be more resilient than another. If, however, we expand the boundary of what we consider to be the “system” to include the delivery of power to customers, not only would the resiliency of an individual power plant represent only the resiliency of part of the system, the overall system is different, and the events that it might be expected to withstand are likely to be also different (e.g., it is unlikely that high winds would have any impacts on a power plant, but they are known to have impacts on the lines that transmit the power to consumers).

Performance and Resilience

People often think that resilience is mainly about performance. Though this is often true, as we will show, this is not always so; often performance is not the only thing to consider when thinking about a system’s resilience. But since performance is often an important aspect of resilience, we should consider how it contributes. Consider the mission performance curve illustrated in Figure 1. The figure shows the mission-oriented performance of a system in response to some incident. The performance of the system can be something as simple as the monetary profit the system produces. However, for most of the systems we are likely to try to analyze, the metrics are complex ones that would require a multi-attribute utility function. This utility function would be defined by reaching agreement across multiple stakeholders as to what constitutes an acceptable way to measure the performance of the system. Irrespective of which

(mission-specific) multi-attribute performance metric is selected, when considering the resilience of a system, the choice of performance metric needs to be described and documented explicitly, so that it is clear what system properties are being optimized for the purpose of resiliency decision making.

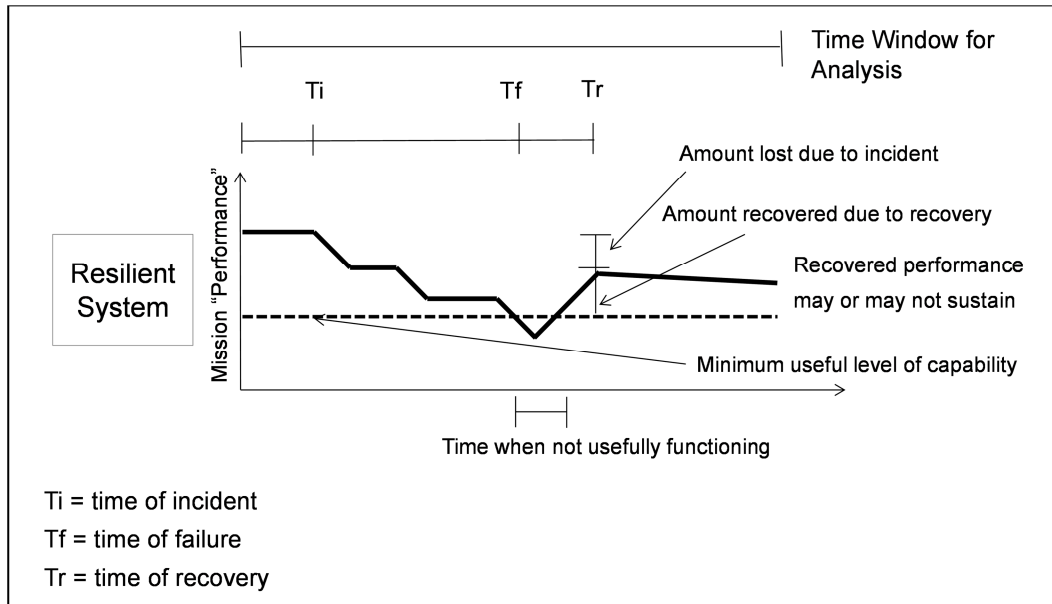


Figure 1. Performance Curve Illustrating Aspects of Resilience

The curve above illustrates many of the aspects of performance that people mention when considering resiliency. The curve shows how the mission-oriented performance of the system has been reduced after the incident occurs, and it shows the time-to-failure. The time-to-failure illustrates the amount of time it takes after the incident occurs before the system performance drops below an acceptable level. Time-to-recovery represents the amount of time it then takes, either because of an automatic response or as a result of a manual intervention, to recover the performance of the system. The recovered performance level may either return to the original pre-incident performance levels, or may only represent a partial recovery. After that, the recovered performance can either sustain at the recovered level or not. People sometimes refer to "time-to-failure," "time-to-recovery," and "recovered performance" as examples of resilience metrics [Sheard and Mostashari, 2008], since there is often an implicit relationship between them and the resilience of a system. Longer "time-to-failure" is usually better. Shorter "time-to-recovery" is usually better, and a recovery to pre-incident performance levels is typically what is wanted. But as we will demonstrate, such metrics are not always appropriate.

Depending on the system, and the mission it performs, there may be some acceptable level of system performance that must be maintained (over some time period). In this context, the resiliency of the system may be binary: as long as the performance stays above the minimum level, it's fine, but below the minimum level it might as well be zero. In Figure 2 the performance of two systems is depicted. Because the performance of System 1 drops below the minimum acceptable level, it is reasonable to say that System 2 is more resilient than System 1. Again, depending on the mission context, it may not even matter whether the performance is below the minimum level for a short period of time or for a long period of time. Also, some system designers/operators are more risk averse than others, and for the risk averse it is often the worst possible performance outcome that matters, rather than the mean or integral value of the performance curve over some time interval.

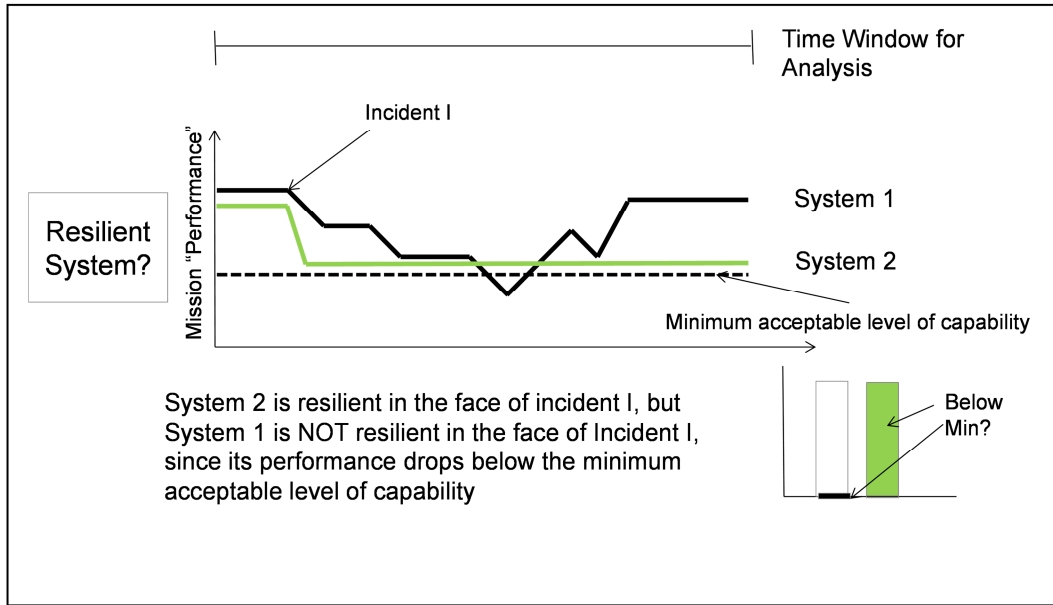


Figure 2. Performance Curves of Two Systems with a Minimum Acceptable Performance Level

An incident-perturbed performance curve as shown in Figures 1 and 2 contains all the performance information that is necessary for computing resiliency. Depending on the system and mission context, however, different people may choose to compute a different metric given the curve. Popular examples of a metric are binary (given some threshold), the minimum, the average, or the integral under the curve. As an example, in Figure 3, if the performance curve represents profit per unit of time, calculating the integral under the curve would indicate that System 1 is more resilient than System 2 in the face of this incident. No one metric choice is necessarily the correct one to use for all circumstances, but for anyone claiming that a system is being made more resilient, whichever metric is chosen should be explicitly documented as being part of how the system's resiliency is being estimated.

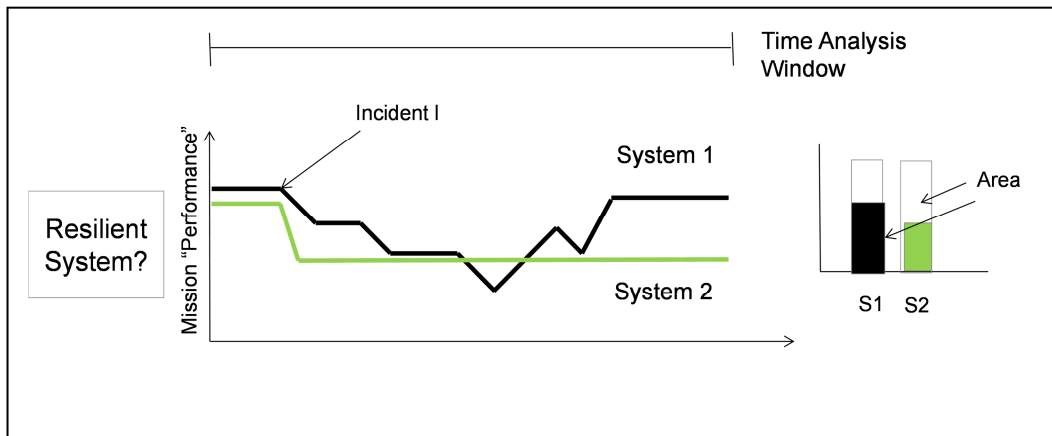


Figure 3. Two Systems Where the Integral of the Curves Represents the Overall Performance

Resilience Depends on the Analysis Timeframe

Above we introduced some common aspects of a system's resilience as being the time-to-failure and time-to-recovery, given the system's mission-oriented performance. In most circumstances,

as in the example in Figure 3, a system that can recover from an incident more quickly will be considered more resilient than a system that recovers more slowly. There is almost always a relationship between time and resiliency, and in exactly the same way that it is necessary to explicitly specify the system performance metric being used to measure resiliency, and to define which function to apply to the performance curve, it is necessary to define the time interval over which a resiliency estimate is valid.

Figure 4 shows how two different systems perform after an incident occurs. If we look at the performance using time window 1 (TW1) and curve metrics such as “min,” “mean,” or the “integral,” it is apparent that we will consider System 1 more resilient than System 2. If, however, we continued to monitor the performance of the systems over a longer time interval, the performance of System 1 might continue to decline, while the performance of System 2 might stabilize. Using the same “min,” “mean,” or “integral” metrics for time window 2 (TW2), it becomes clear that System 2 is more resilient than System 1. Thus, the resiliency estimate is valid only for the time interval selected and so its chosen value must also be explicitly documented as part of any resilience calculation.

Usually there are operational motivations for selecting a time window over which to compute. Sometimes it’s because having a longer time to failure might provide additional time to respond to the incident (e.g., to evacuate a burning building), which might save lives. There are a number of examples in which people have focused on using incident recovery metrics as their measure of resiliency [Vugrin et al., 2010], and this remains entirely consistent with our view of estimating resiliency since it merely involves selecting a specific timeframe associated with the system performance after the occurrence of an incident. Sometimes the timeframe is dependent on the complete lifecycle of the incidents of concern. Particularly in the physical world, where you might be considering how society continues to function in the face of damage to buildings and infrastructure from events such as hurricanes, floods, or earthquakes, the incident lifecycle can be quite long, since it must consider the manual actions of rebuilding and adapting. Choosing a timeframe to use for computing resilience is always going to be system- and mission-dependent.

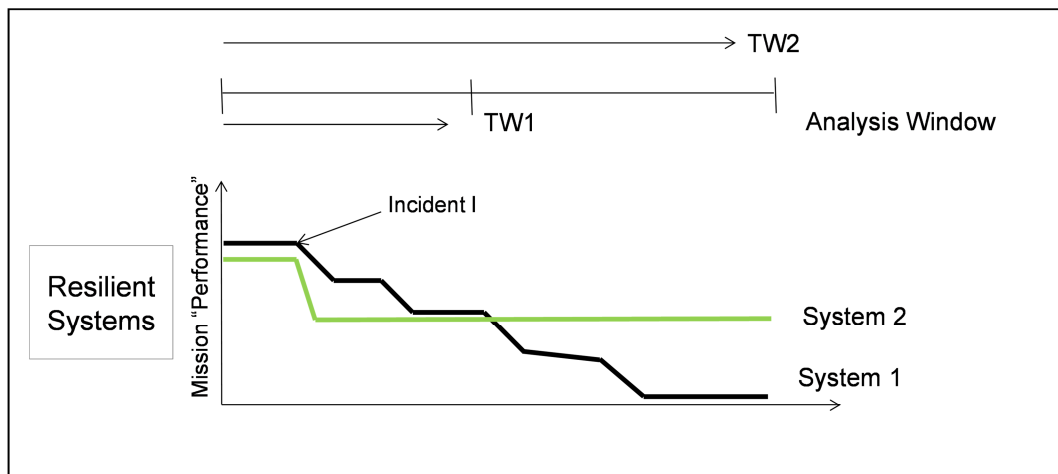


Figure 4. The Effects of Choosing a Different Time Window on Calculating Resiliency

Resilience Is Also About the Capacity to Withstand Incidents

In the previous examples we presented resiliency examples that are based on comparing how a system’s mission performance is affected in the face of specific incidents. While this is an

important aspect of resilience, it is not the only aspect to consider. Consider the performance curves for two different systems shown in Figure 5. Although it looks only like a single performance curve is present in the figure, there are actually two coincident curves that are indistinguishable from each other. What is depicted are two systems that successfully withstand an incident without having any performance impact. An example of how this might be possible might be to have fault-tolerant systems that contain hot spares for a component that fails or is compromised. But what if System 1 has two hot spares while System 2 has only one? In the face of a single compromise neither system is impacted; however, most people will intuitively consider that System 1 is more resilient than System 2, because it has the capacity to withstand more compromising events. What Figure 5 illustrates is that focusing only on performance in the face of incidents is sometimes an inadequate way to represent resiliency. While in some circumstances it is perfectly acceptable to present an assessment of resiliency that is based purely on performance, there are also many situations where resiliency is about the possibility that some number of compromising events might occur. Consider the following definition for estimating risk:

$$\text{Mission Risk} = P(\text{bad-event-occurs}) \times \text{the-value-of-loss-given-event}$$

Since we associate resiliency with the inverse of risk, this calculation adequately captures and represents how System 1 is more resilient than System 2, since the likelihood that two component failures occur is almost certainly less¹ than the likelihood that a single component fails. Unlike the general case of calculating risk, this formulation is acceptable because we're estimating risk in only a specific context, considering how "like" systems respond to the same set of incidents over the same evaluation timeframe, under the same assumptions.

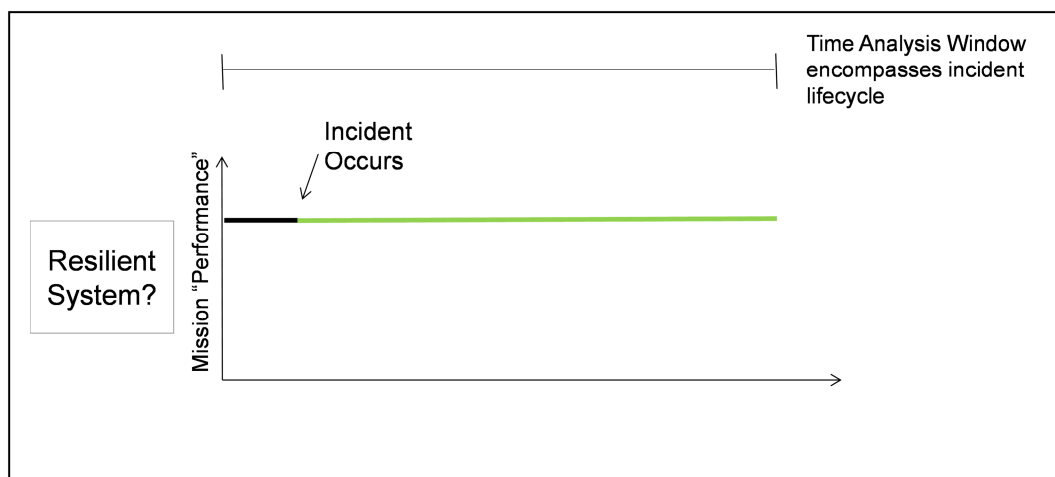


Figure 5. Two Systems that Perform the Same in the Face of an Incident

Cost/benefit tradeoffs are usually involved in seeking to achieve additional resiliency, but we do not cover that issue in this paper.

The Occurrence of Incidents Is Uncertain

While it is often reasonable to talk about a system's resilience as the capacity to withstand specific incidents, in general systems are designed for and must operate in an environment in which one or more incident types "might" occur. Some incidents are more likely to occur than

¹ When dealing with deliberate attacks against a system rather than random failures, the ability of an adversary to subvert a component usually implies that other instances of the same type of component in the system are also susceptible to the same act.

others. Some incidents have more impacts than others, and from a systems engineering perspective, resiliency is about getting the best overall system performance given the uncertainties associated with these incidents.

In the previous section, we discussed system resiliency in the context of having the capacity to withstand incidents. What we were actually describing were situations where incidents that could cause some level of performance degradation were less likely to occur for one system than the other. So, in this current context a comparative measure of the resilience of two or more systems can be expressed as the probability (or likelihood) that some incidents may bring the system to a failure state (or some other pre-determined definition of bad performance).

Consider the systems shown in Figure 6 Similar to Figure 5, this figure depicts the performance of two systems that perform identically in the face of some incident, making it appear that there is only one performance curve. Suppose, for example, System 1 is made of plastic parts and System 2 is made of titanium parts, and the probability that plastic parts will fail is higher than the probability that titanium parts will fail. So $P(\text{fail} | \text{incident}) = 0.6$ for System 1 and $P(\text{fail} | \text{incident}) = 0.1$ for System 2. Since the performance of both systems is the same given the incident, a performance-based description of resilience is not informative, and we must consider an alternative measurement. In this case, expressing resiliency as a probability (or likelihood) that the performance-impacting event will occur is more useful, and intuitively most people will agree that the system with titanium parts is more resilient. We are using the term “more resilient” in this context because we are also considering the system response in the face of deliberate actions of misuse rather than just random failure.

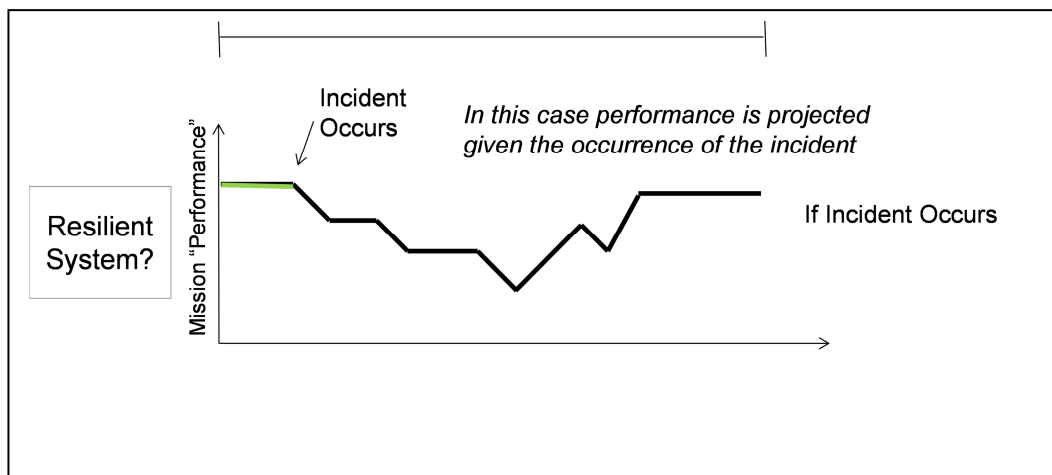


Figure 6. Performance of the System With and Without an Incident

Usually, however, a system must withstand multiple incident types, and system performance in the face of different incidents can be expressed as multiple performance curve metrics as described above. A system may do well in the face of some incident types (magnitudes, durations, etc.) but less well for others. Estimating the comparative resiliency of the system, then, depends on estimating a composite, overall performance metric given the relative likelihood that each type of incident might occur. Typically, for convenience, people will consider that the different types of incidents are mutually exclusive of each other, but in reality they may not be, and that can make the estimation of resiliency much more complex. That is because it relies on estimating a joint probability distribution that scales exponentially with the number of dependent variables, and may also involve estimating the system performance given the occurrence of multiple simultaneous incidents.

Figure 7 shows another example of two systems with different performance metrics for two different types of incidents. If we consider that the incidents that that we are considering are mutually exclusive of each other (that is, we are concerned that Incident 1 or Incident 2 might occur but are not concerned that both might occur simultaneously), and if $P(\text{Incident1})=0.95$ is much more likely than $P(\text{incident2})=0.05$, then, based on the performance curves, System 2 would be more resilient than System 1 (assuming you are measuring performance as integral under the curve). If the circumstances were reversed and $P(\text{incident2})=0.95$ was much more likely than $P(\text{incident1})=0.05$, then System 1 would be more resilient. One can think of this example as measuring the overall system output performance given a number of Bernoulli trials, where the number of trials for each incident type is in the ratio of the incident probabilities.

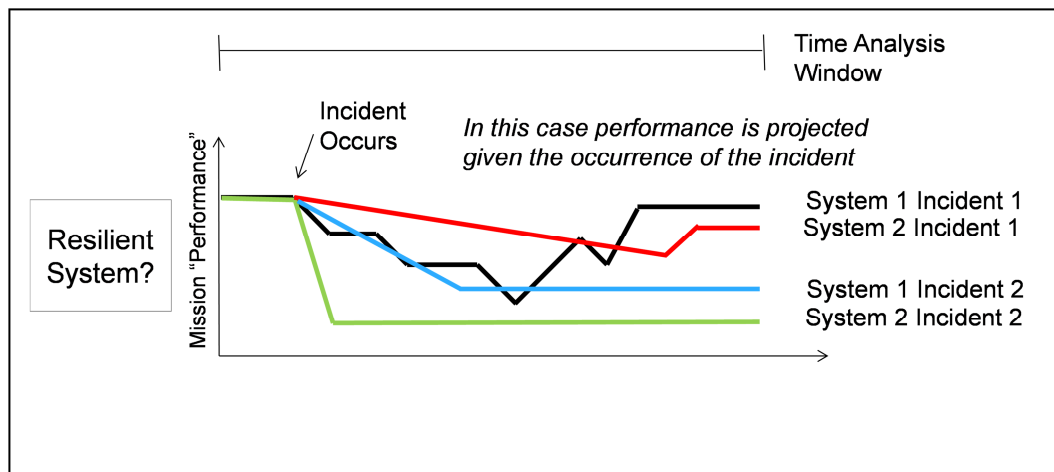


Figure 7. The Performance of Two Systems Given Two Different Incidents

The Scale of Deployment May Affect Resiliency for Those Who Are Risk Averse

Another consideration relating to estimating resiliency is to understand the differences associated with the scale of deployment for resilience mechanisms. So far we have discussed the performance of systems as being more or less resilient. In reality, different mechanisms that can make a system more resilient may work differently from each other. One mechanism may work with certainty on some incidents. Another may, in addition, also work on other incidents, but with some variance. As a result of the variance, the number of times a mechanism is deployed can matter for risk-averse decision makers.

For example, consider two Host-based Intrusion Prevention Systems (HIPS). One (HIPS 1) reduces the number of successful Denial of Service (DoS) attacks from 10% to 5% with certainty. The other (HIPS 2) reduces the number of attacks from 10% to 0% with certainty $P(0.8)$. If there is only a single deployment of this mechanism, then HIPS 2 is the best choice. But as the number of mechanisms deployed increases, the variance of HIPS 2 means that there will be intervals of time where the second mechanism may perform worse than HIPS 1, which has a guaranteed performance of reducing the number of attacks to 5%. For a risk averse decision maker (someone who cares about the worst possible outcome), HIPS 1 is a better choice.

3 A Definition of Resilience

Based on the discussion above, we have come up with the following general definition useful for estimating the resilience of a system.

Resilience is:

The persistence under uncertainty of a system's mission-oriented performance in the face of some set of disturbances that are likely to occur given some specified timeframe.

This definition is very similar to the definition proposed by Ayyub [Ayyub, 2013], but it is more explicit in that it specifically calls out the fact that an estimate of resilience depends on the set of disturbances over which it is valid, as well as the fact that it exists within an operational timeframe.

The following definitions apply to the definition above:

- **System** defines the “scope” (or boundaries) over which the analysis applies.
- **Performance** is a function of mission requirements, outcomes, or objectives, and is measured as a form of output, throughput, etc. metrics that indicate how well the system contributes to achieving those objectives.
- **Uncertainty** relates to the probability or likelihood of the events and disturbances that the system may experience given the analysis timeframe.
- **Persistence** is a matter of the system enduring the events and/or recovering and continuing the performance of its operation.
- **Disturbance** is synonymous with incidents, an event, or attack that would be likely to have an impact on the system's performance.
- **Timeframe** is the time interval over which the performance, uncertainty, and persistence measures apply.

This definition of resilience is, in effect, a risk metric. Risk is conventionally measured as a functional combination of the likelihood that an undesirable event will occur and the severity of the event's consequences, where it is understood that these factors are evaluated with respect to some specified timeframe and with some degree of uncertainty. In this definition of resilience, the undesirable event is “some number of disturbances” and the consequences are measured as the persistence of system performance.

“Mission risk” is the risk that a mission will fail to meet its objectives. If a mission depends on a system, mission objectives are translated into Measures of Performance (MoP) and Measures of Effectiveness (MoE) for mission tasks and/or Key Performance Parameters (KPP) for the system. In our definition of resilience, persistence of performance can be expressed using one or a combination of these metrics computed over the chosen timeframe.

4 Example of Estimating Resiliency

For an example of how mission risk can reflect resiliency, and how resiliency can be calculated and applied to decision-making, consider the following:

An e-commerce company sells products via a web service that allows their customers to view their products and purchase them. Having grown from a small startup, the company relies on a single web server that runs ruby scripts. The company is now worried that unavailability of the web server (whether from a system or software failure or from a malicious attack that causes the server to crash) will affect their profits. Their IT guru has suggested two possible solutions.

1. To replicate the server, so that if one becomes unavailable the other will continue to serve customers.
2. To implement a fast recovery solution, so that the time required to reconstitute an unavailable server so that it is up and running again is significantly reduced so that fewer customer orders will be lost.

To analyze these options, the company needs to consider the frequency of customer orders, the average value of those orders, and how long it currently takes to recover from an unavailable server. In a typical 12-hour operating day, the company receives 1000 orders, each averaging \$200 profit. When the server becomes unavailable, the time required getting it back up and running is usually three hours. So the typical impact of a server-unavailable incident is about \$50,000. These server-unavailable incidents occur on average about once every 3 months (90 days), so the probability that the server will be unavailable on any given day is 1.11 percent.

For proposed solution #1, replicating the server doesn't reduce the impact of an incident should both servers fail, but it does reduce the likelihood that they will both be down at the same time. Conservatively, since a server crash is only one of the reasons why the service might be unavailable to customers, IT has estimated that the replication solution reduces the likelihood that both servers will be unavailable by at least 75%.

For proposed solution #2, using a fast recovery server doesn't reduce the chances that the server will be unavailable, but it does reduce the impact of each unavailability event since the time to recover to operating condition has been reduced to ~35 minutes. Hence, using solution #2, instead of having an impact of \$50k, the income loss (impact) is now less than \$10k.

Using our mission risk formulation, for the initial state of the system without using any possible mitigation, the mission risk is as follows:

Since mission risk = $P(\text{bad-event-occurs}) \times \text{the-value-of-loss-given-event}$
the nominal operational risk of the system is

$$0.01111 \times \$50,000 = 555.5$$

If solution #1 is used, then the mission risk becomes:

$$(0.01111 \times 0.25) \times \$50,000 = 138.9$$

If solution #2 is used, then the mission risk becomes

$$0.01111 \times \$10,000 = 111.1$$

Referring back to our definition of resilience, we must also describe the resilience context. In this example the mission risk estimate assumes that the system definition encompasses only the operation of the web server being able to take customer orders. Other business functions, such as payroll, inventory, etc. are outside the scope of this analysis. Mission performance is based on an expected rate of orders of an average value that would be received during the period during which the server is unavailable. Uncertainty in this example is based on the fact that the original likelihood of an incident is 1.11%, and that one of the resiliency techniques reduces this likelihood. Persistence considers the typical duration of these incidents, and how one of the resilience techniques reduces that interruption. Disruption in this case only considers unavailability events, and does not consider theft of customer information, or other disruptions such as the modification of information in the system. Finally, the timeframe associated with this metric is a 12-hour business day, since it is expected to take less than a business day to resolve the incident.

Based on the risk metric, both proposed solutions are shown to be effective in terms of making the e-commerce site more resilient in the face of server-unavailability events. Assuming we are comfortable with the performance estimates associated with applying the different solutions, solution #2 is shown to be a more effective approach than solution #1.

Although this example is greatly over-simplified, and as depicted does not include an adequate description of the nuances of DoS events against the e-commerce site, and how the different solutions address them, it is only intended to illustrate the application of the mission risk metric as guidance to a decision maker. Obviously, the hard part of performing this analysis is to come up with accurate estimates for mission impacts, and to determine how the resiliency techniques can reduce the likelihood of the bad events occurring. Although these are non-trivial problems, one must consider why anyone would ever choose to implement a resilience technique if it were not possible to make any assessments about its expected performance (be they measured or subjective assessments).

For the purpose of comparison of different resilience solutions, this approach works because each version of the system is performing the same mission; the mission performance is based on the same metric in each case; the timeframe over which we are performing the resiliency analysis is the same; and in this case the resilience comparison is valid only in the context of server-unavailability events. If we were to expand the resiliency context to include other types of incidents, such as unavailability of the internet connection, modification of data, and/or the interception of confidential information (e.g., customer information or their credit card details) on the server, then the resilience estimate would be different.

5 Comparison with Previous Work and Definitions

The term “resiliency” has multiple connotations and is multi-dimensional depending on the context in which it is used [Haimes, 1991; Gates. 2011]. The concept of resilience and related research on its definition and metrics have appeared in multiple domains, including cyber systems and biological systems [Sheard and Mostashari, 2008; Vugrin et al., 2010]. Most of the definitions, however, provide insufficient context, and hence do not support clearly measureable metrics [Wood, 2005; Gilbert, 2010]. Sheard and Mostashari, who provide a survey of other peoples’ works, conclude that a measureable definition of resiliency that would support metric development should include such multiple attributes as time frame (short-term recovery, long-term recovery, etc.), events (disturbance, perturbations, attacks, etc.), system definition (cyber systems, ecological systems, etc.), required actions (failover, recover, etc.), and preserved qualities (system function, structure etc.). We have found that the work by Sheard and Mostashari provides the most comprehensive description of the factors that relate to resiliency, and hence it offers the strongest basis for developing a measurable definition of resiliency. Despite this, they do not actually develop any workable metrics themselves.

Vugrin defined resilience in this manner: *“Given the occurrence of a particular disruptive event (or set of events), the resilience of a system to that event (or events) is the ability to reduce efficiently both the magnitude and duration of the deviation from targeted system performance levels”* [Vugrin et al., 2010]. This definition points to a computable resilience metric that is based on system performance, but does not take into consideration any uncertainty associated with the disruptive events that might occur, hence giving them all equal weighting. Inherently, this becomes a limitation to the resiliency definition in Vugrin’s paper, since, as we discussed in the previous section, there are a number of circumstances where uncertainties should be considered in resiliency definitions. Primarily, this is because of the different degrees of impacts given different incident types, and there are many examples where a system can be made much more resilient overall by eliminating low-impact but highly likely events, rather than by eliminating impactful but unlikely events.

Wood defined resilience in systems as *“a system’s ability to adapt or absorb disturbance, disruption and change”* [Wood, 2005]. Including attributes such as system and disturbance in Wood’s definition helps facilitate the development of a computable metric. Gilbert and his associates [Gilbert et al., 2010] defined cyber resiliency as *“the ability to provide and maintain an acceptable level of service, in the face of faults and challenges to normal operations.”* Using this definition, the level of service provides an adequate surrogate for measuring system performance, accepting that “service” might be a compound multi-attribute metric obtained by combining several system performance metrics together. However, this definition doesn’t clearly differentiate whether resilience is being defined in the context of all possible events (the faults and challenges to normal operation) or perhaps just some. In both of these definitions, other attributes (timeframe, uncertainty, etc.) that could support such a measureable definition of resilience were not clearly discussed or enumerated.

Haimes and others [Haimes, 2009; Bishop et al., 2011; Ford et al., 2012] all defined resiliency metrics while discussing them in terms of qualitative measures. Haimes defined system resiliency as *“the ability of the system to withstand a major disruption within acceptable degradation parameters and to recover within an acceptable time and composite cost and risks.”* Bishop et al. defined resilient systems as *“systems expected to maintain their operations under attack or failure,”* but stated that *“they are also expected to remain mission-capable, that is, to*

reconfigure or recover in order to restore their original state.” Bishop et al. compared resiliency to survivability, robustness, reliability, etc. in an attempt to develop a measureable definition of resiliency. Ford et al. suggested defining the term resiliency by considering specific systems, tasks, outputs, and other variable conditions. The definitions of Haimes, Bishop et al., and Ford et al. all discuss and consider the various factors that would enable the development of measureable resiliency metrics. They all suggest that resiliency metrics are highly contextual, that a holistic approach to each specific system should be considered, and that it is impractical to generalize a set of quantitative metrics for different systems. For example, Ford et al. stated that recovery time resiliency metrics used for ecological systems might not be appropriate for computing systems because of their different missions, recovery options, etc. Ford et al. and Haimes also identified the relationship between resilience and risk by suggesting how improving a system’s resilience offers significant advantage in managing its risk, and they discussed examples of how the resulting risks can be measured in terms of recovery time and/or composite costs (that can be calculated in a variety of ways). These costs can include the expected value of risks or the conditional value of the extremes, given that the inputs and outputs are probabilistic. Among the three authors, Haimes’s discussion is the one that is supportive of the fact that the resilience metric of a system can be measured as a probability or as the inverse of a risk estimate. This last argument is reinforced by the fact that the probabilistic nature of inputs and outputs warrants the consideration of a probabilistic measurement of mission risk [Haimes, 2009]. These statements, though limiting to the definitions of Bishop et al. and Ford et al., all support the discussions in our previous sections on resiliency considerations. Nonetheless, even Haimes ultimately failed to transform his qualitative discussion into quantitative computable metrics.

Holling defined ecological resilience as *“the persistence of relationships within a system and is a measure of the ability of these systems to absorb changes of state variables, driving variables, and parameters, and still persist”* [Holling, 1993]. Ayyub defined resilience in multi-hazard environments as *“the persistence under uncertainty of a system’s performance in the face of disturbances”* [Ayyub, 2013]. Together, Holling and Ayyub indicate the importance of needing multiple contextual attributes to support the development of resilience metrics. Despite their identification of several of the necessary contextual attributes in each of these definitions, neither of them considers the timeframe over which such a computable metric would apply. As discussed in the previous section, time and resilience are interrelated. For instance, for a given system, one resiliency technique might be best suited for one mission timeframe (e.g., short-term performance) while another could be best for a different timeframe (e.g., long-term performance).

Based on the literature, for example Sheard and Mostashari [2008], and from our own discussions on considerations for resiliency, we conclude that when developing a measureable definition of resiliency, the following must be taken into consideration: it is system specific (e.g., it is impractical to compare a computer system with an ecological system); a holistic approach should be taken given the system under consideration; sufficient attributes of the contextual system characteristics and its threat environment need to be specified; and the resiliency metric can be a performance, a probability, or a risk measure.

6 Discussion

The contribution of this paper is to present a quantifiable definition of resiliency that would allow systems (cyber systems, biological systems, etc.) and engineers to evaluate how changes to a system can make them more resilient, or to allow the resiliency of different systems (or design options) to be compared against each other. Since we have defined resiliency in terms of the persistence of performance, and hence in mission terms, we do not define an absolute scale for resiliency – it is only reasonable to support a relative comparison of resiliency. In the same way as it is not possible to compare on an absolute scale how an electrical power plant will withstand an earthquake with how a city will withstand a hurricane, our definition of resiliency is only valid for comparison given “like systems” (i.e., given the context in which it is defined).

In computing a measurable metric for resilience, the contextual factors associated with that measurement are vitally important. Any proposed resiliency metric that does not adequately specify the context in which it is valid leaves no room for its proper evaluation. Based on our study, the factors that must be specified are:

$S = \{S1, S2, \dots S_n\}$, the definition of the system being evaluated

$I = \{I1, I2, \dots I_m\}$, the set of incidents the resiliency metric covers

$T = T_w$, the time window over which the metric is calculated

$P = \{P_{f1}, P_{f2}, \dots P_{fk}\}$, the performance metric used to measure system performance

$M = \{\min, \max, \text{mean}, \text{etc.}\}$, the metric used to evaluate the performance graph

$U = \{P(I1), P(I2), \dots P(I_m)\}$, the uncertainty that each incident in the set will occur

S defines the scope of the system in the form of the set of functional capabilities it provides. I is the set of incidents that are considered in estimating the resiliency of the system. T is the overall timeframe for which the resiliency estimate is valid. P is a metric that represents how the performance of the system is characterized in the face of incidents; P can be the combination of multiple metrics P_f . M is the choice of metric that is used to evaluate the system performance P given the anticipated system performance in the face of incidents given the time window T , and where min, max, mean, integral, etc. represent popular examples. U represents the set of probabilities that each incident in the set I will occur.

Given that these attributes define the context in which the resiliency metric is valid, there are still a number of different ways the resiliency of a system can be reported.

1. It can be reported as a performance score $\{P, T, M\}$ for the system $\{S\}$ mission-oriented performance in the face of incidents $\{I\}$
2. It can be reported as a probability $\{U\}$ for the system $\{S\}$ that might have some stated mission-oriented performance $\{P, M, T\}$ impact, given incidents $\{I\}$
3. It can be reported as a risk score $[R = \{U\} \times \{P, T, M\}]$ considering the performance impacts given the probabilities $\{U\}$ of incidents $\{I\}$

Each of these statements about a system’s resilience is perfectly valid, given the context for which it is defined. A statement of type 1 assumes that there is no uncertainty in the occurrence of the incidents. A statement of type 2 normalizes out the performance aspect of resiliency and reports on the probability that the incidents would cause an undesirable level of performance (e.g., how likely is it that the incidents would be able to reduce the mission performance to an

unacceptable level). A single system can also be described by multiple resiliency statements, each with a different context, and where each statement might be a different one of the above forms. A system may be very resilient in the face of one type of incident, but may be less so for other incident types. It may also have yet another set of resilience characteristics for the combination of possible incidents, given their relative likelihoods.

Our definition of resilience adequately encompasses metrics for resilience that have been reported by others. The main difference is that we require that the context for the resilience metric to be explicitly defined, rather than leaving it implicit, as is often the case. For example, the metric “time to recovery” (which is commonly used) is encompassed by our definition. This metric is incorporated by being specific about which incident(s) are being recovered from, and by choosing a time window that starts at the point of worst performance, given the incident(s).

In our analysis, we have found no reason to differentiate between resiliency for structural systems, cyber systems, cyber physical systems, biological systems, or even cultural systems. “Resiliency” is a term that depends on its context, and so the resiliency of each system must be considered in the context of the environment in which it operates, the incidents it might face, and any operational (mission-oriented) needs that the system must accomplish. By defining the complete resiliency context (as described above), our definition of resilience should apply for any given system description, making it clear why one system design might be better than another, and how a system’s modification can lead to a more resilient outcome.

7 Summary

This paper was motivated by our desire to be able to compute resilience metrics for cyber systems and improve upon them by applying different resilience techniques [Bodeau and Graubart, 2011; Goldman, 2010]. Since each resilience technique is only likely to be effective against a specific incident or subset of incidents, the goal is to be able to determine which combination of resilience techniques is best suited to making the system more resilient. A good way to achieve this is to optimize the resiliency decisions using a computable metric. Unfortunately, our review of the literature on resilience found mostly qualitative definitions and failed to identify a sufficiently generic yet computable definition for resilience that we could use as our metric.

In this paper we propose a computable definition for resilience, based on mission risk. As has been described by others, resilience is a complex concept that requires the consideration of many factors to define, and is one that must further be defined by the context in which it is considered.

In our definition, resilience can be computed as being either:

- A utility-based performance metric that indicates how well the system responds in the face of one or more incidents (where the incidents are assumed to have occurred)
- A probability that some events might occur to bring the system to some specified unacceptable level of performance
- A risk estimate that combines the probability of incidents with the system utility-based measure of performance changes that result when the incidents occur

A single system can be characterized by more than one resilience statement. A system may be resilient in the face of one type of incident, but not another. The resiliency of a system can be evaluated in the context of only a subset of the incidents that are possible, because other incidents might be out of scope. Clearly, any metric computed for resiliency must be qualified by its context.

The context needed to define a system's resilience depends on specifying the *system* and its boundaries, the set of *incidents* being considered, the *timeframe* over which the system performance is being analyzed, how *system utility* is being estimated, the *metric used to evaluate performance* (e.g., min, max, average) over the time interval, and the *uncertainty* associated with the incidents that might occur.

Although there is no absolute scale for estimating resiliency, selecting a context in which to make a resiliency computation makes it possible to compare “like” systems and to evaluate whether modifications to a system make it more or less resilient.

8 References

1. Haimes Y.Y. Total risk management. *Risk Analysis*, 1991;11(2):169–171.
2. Gates, R.T. Science and Technology (S&T) Priorities for Fiscal Years 2013-17 Planning, Memorandum from the Secretary of Defense, Washington, D.C., April 19, 2011.
3. Woods D.D. Creating foresight: Lessons for resilience from Columbia. Pp. 289–308 in Farjoun M, StarbuckWH (eds). *Organization at the Limit: NASA and the Columbia Disaster*. Malden, MA: Wiley-Blackwell, 2005.
4. Gilbert, S. Disaster Resilience: A Guide to the Literature. NIST Special Publication 1117, Office of Applied Economics, Engineering Laboratory, National Institute of Standards and Technology, Gaithersburg, MD. September 2010.
5. Haimes, Y.Y. (2009). On the definition of resilience systems. *Risk Analysis*, 29(4), 498 – 501.
6. Bishop, M., Carvalho, M., Ford, R., and Mayon, L. Resilience is more than availability. In *Proceedings of the New Security Paradigms Workshop (NSPW)*, 2011.
7. Ford, R., Carvalho, M., Mayron, L., and Bishop, M. Towards Metrics for Cyber Resilience. *21st EICAR Annual Conference Proceedings* pp. 151–159. May 2012.
8. Holling C.S. Resilience and stability of ecological systems. *Annual Review of Ecology and Systematics*, 4(1):1–23; 1993.
9. Ayyub, B.M. Systems Resilience for Multihazard Environments: Definition, Metrics, and Valuation for Decision Making. *Risk Analysis*. 2013.
10. Sheard, S. and Mostashari, A. “A Framework for System Resilience Discussions”, 18th Annual International Symposium of INCOSE, Utrecht, Netherlands, June 15-19, 2008.
11. Bodeau, D., and Graubart, R. *Cyber Resiliency Engineering Framework* (MITRE Technical Report MTR1-10237). Bedford, MA: MITRE Corporation., 2011
12. Bodeau, D., Graubart, R., LaPadula, L., Kertzner, P., Rosenthal, A., and Brennan, J. *Cyber Resiliency Metrics, Version 1.0, Rev. 1* (MITRE Technical Report MP12-0053). Bedford, MA: MITRE Corporation., 2012
13. Goldman, H. *Building Secure, Resilient Architectures for Cyber Mission Assurance* (MITRE Technical Report 10-3301). Bedford, MA: MITRE Corporation., 2010
14. Vugrin, E.D., Warren, D.E., Ehlen, M.A., and Camphouse, R. C. (2010a). A framework for assessing the resilience of infrastructure and economic systems. In K. Gopalakrishnan and S. Peeta, eds., *Sustainable and resilient critical infrastructure systems: simulation, modeling, and intelligent engineering* (pp. 77-116). Berlin: Springer-Verlag, Inc., 2010