

Reliability to Resiliency

Exploring Interconnected Metrics of System Performance

INCOSE WSRC · 2025

AUTHORS

Vincent P. Paglioni, Aaron Brown, and Steven Conrad

PUBLICATION DATE

July, 2026

SUBMISSION

57

Reliability to Resiliency: Exploring Interconnected Metrics of System Performance



Vincent P. Paglioni
Risk, Reliability & Resiliency
Characterization (R3C) Lab
Department of Systems
Engineering, Colorado State
University
Vincent.paglioni@colostate.edu

Aaron Brown
Department of Systems
Engineering, Colorado State
University
Aaron.brown@colostate.edu

Steven Conrad
Department of Systems
Engineering, Colorado State
University
Steve.conrad@colostate.edu

Copyright © 2025 by the author(s). Permission granted to INCOSE to publish and use.

Abstract

Reliability and resiliency are often treated as distinct system characteristics. However, this approach reinforces unnecessary tradeoffs between the two. A holistic approach to understanding reliability, and resiliency and their relationship to risk ensures safe system operation as increasing complexity threatens our understanding of system safety. This paper reviews the foundations of risk, reliability, and resiliency in the context of engineering systems – particularly socio-technical systems, to approach a unified framework for understanding these critical system characteristics.

Keywords

Resiliency, reliability, risk, system safety.

Introduction

Reliability and resiliency are both fundamental attributes of system performance and safety, and although they are often treated as separate metrics, they are inextricably linked in both concept and practice. These attributes inform our understanding

of how systems behave over time, particularly under stress or failure conditions, yet their analyses are frequently siloed. This separation can obscure important dynamics, particularly where tradeoffs between the two must be managed.

At their core, reliability and resiliency reflect different facets of system robustness:

- Reliability is a measurement of the probability that a system performs its designated function without failure and within specified requirements under expected conditions over a set timeframe.
- Resiliency encompasses a system's ability to absorb shocks, recover critical functionality, and adapt dynamically in response to unanticipated failures or changing conditions.

While a reliable system is designed to avoid failure, a resilient system is designed to survive and recover from it. While these concepts are distinct, they are also interconnected. Although interrelated, often there is conflict between the two practical applications. For example:

- A system designed for maximum reliability may be highly optimized and efficient under nominal

conditions but brittle in the face of unmodeled events. Such a system may lack the capacity to adapt, reroute, or self-heal, thereby possessing low resiliency.

- Conversely, a highly resilient system may tolerate frequent disturbances or component-level failures gracefully, reconfiguring itself or bouncing back with minimal disruption. However, this flexibility can come at the cost of increased complexity or decreased baseline reliability due to a greater number of interacting parts or recovery mechanisms.

This interplay reveals a critical insight: resiliency often becomes visible only when reliability fails. Thus, the two concepts operate on different timelines; reliability dominates during normal operations, while resiliency governs the aftermath of disruption. Both must be considered to fully characterize system behavior over its lifecycle.

Despite their overlap, reliability and resiliency are often evaluated through disjointed methodologies. Traditional reliability engineering emphasizes failure probabilities, mean time between failures (MTBF), and probabilistic risk assessments. These approaches typically focus on pre-failure conditions and may not explicitly account for what happens after a system fails; how it copes, adapts, or recovers.

On the other hand, resiliency analysis often centers on post-failure dynamics: response time, recovery effectiveness, degradation tolerance, and the ability to maintain core functionality during disruption. Some approaches are qualitative or scenario-based and may not quantify failure likelihood at all. This can lead to gaps in understanding where a system is either presumed to be inherently resilient or unrealistically modeled as binary (working or failed) with no middle ground.

Some applications call for a more holistic approach which acknowledges that reliability and resiliency must be co-designed. Systems should not only strive

to minimize the chance of failure but also incorporate pathways for graceful degradation, rapid recovery, and functional reconstitution. Systems engineering, particularly for complex or safety-critical domains (e.g., aerospace, energy, healthcare, humanitarian, sustainability), benefits from frameworks that can simultaneously model both performance under normal conditions and adaptive capacity under abnormal ones.

To align our study of these interconnections, we adopt the following operational definitions drawn from the literature:

- Reliability is “the probability that a system successfully performs its mission under stated conditions for a given period of time” (Modarres & Groth, 2023).

Resiliency has been variously described as:

- “The ability to prevent something bad from happening, prevent something bad from becoming worse, or recover from something bad once it has happened” (Boring, 2010).
- “The ability to respond to an unanticipated disturbance [...] and then to resume normal operations quickly and with minimum decrement in performance” (Fairbanks et al., 2014);
- “The capability of recovering safely and efficiently from abnormal events” (Patriarca et al., 2018); and
- “The capability of a system to anticipate, monitor, plan for, respond to, learn from and adapt in the face of unexpected events and maintain system safety” (Paglioni, 2024).

These perspectives underscore the multidimensional nature of resiliency, incorporating anticipation, absorption, recovery, and adaptation. The relationship between risk, reliability, and resiliency is context-dependent but essential for systems engineering. Risk, as an explicit function of both probability (of failure) and consequences, often serves as the bridge between reliability and resiliency. As such, understanding these three dimensions together allows for better-informed tradeoffs and more robust system architectures. This paper discusses risk, reliability, and resiliency first as distinct systems characteristics, and then presents an

approach to viewing them as a holistic construct by casting resiliency as part-and-parcel with risk and reliability. To aid in conceptualizing these concepts – both distinctly and as a whole – this paper incorporates a running “toy example” of a nuclear power plant.

Background

Risk, reliability, and resiliency may have slightly different definitions across different contexts in engineering. This section provides the approaches to these three system characteristics as we use them in this study of their interconnections.

Risk

Risk can be defined at different levels. Broadly, a *risk* is the possibility of unwanted consequences. That is to say, risk is *inherently* uncertain and negative. More specifically, risk is typically defined for a system using the “risk triplet,” three questions that characterize the magnitude of the uncertainty and negative consequences (Kaplan & Garrick, 1981):

1. *Scenario*: What can go wrong? What accident sequences are possible in the system?
2. *Consequence*: What happens if an accident sequence occurs?
3. *Probability*: How likely is this to occur?

Functionally, risk is often characterized as the product of probability and some quantified consequence, as in Equation (1), as the scenario is used to structure the risk analysis. That is, the event trees or fault trees used to assess risks assume a top event or event sequence (scenario s_i). The risk of a scenario $F(s_i)$ is the *expected* consequence of that scenario; the product of consequences $C(s_i)$ and scenario probability $\Pr(s_i)$.

$$F(s_i) = C(s_i) \times \Pr(s_i) \quad (1)$$

The risk of an entire system can be found as the sum of risks across all considered scenarios as in Equation (2). Total system risk is the expected consequence of any envisioned scenario.

$$F(\text{sys}) = \sum_{i=1}^n F(s_i) = \sum_{i=1}^n C(s_i) \times \Pr(s_i) \quad (2)$$

Risk thus captures the probability and consequences of adverse events occurring at *some* point in a system’s life. However, the definition of risk offered in the risk triplet and Equations (1) and (2) does not explicitly address the temporal aspect of risk. The probability associated with a risk changes over time as failure rates increase due to aging (or decrease due to maintenance and replacement), and the probabilities of external events change.

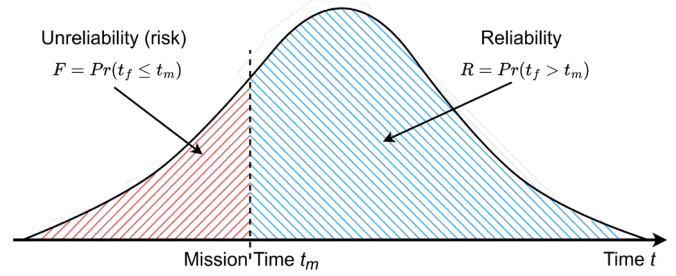


Figure 1. Risk and Reliability as probabilistic metrics in a time-to-failure distribution.

Risk is deeply connected to both reliability (e.g., Figure 1) and resiliency. Where risk identifies the probability and consequences of system failure, reliability investigates the probability of system success (i.e., the complement of risk), and resiliency is concerned with the *recovery* from failure. So, where risk characterizes the failure of a system, reliability addresses the success criteria for a system.

Risk can be described as a function of the system as a whole, the constituent subsystems, or even at the component level. The basis for risk analysis, particularly in approaches such as failure modes and effects analysis (FMEA) and probabilistic risk assessment (PRA), is that small (i.e., component-level) failures can propagate to produce system-level risk. For a nuclear power plant, total system risk is the sum of (probability × consequences) over all identified accident sequences (event trees in the PRA). Thus, a statement to the effect of “the risk of core damage is low” expresses that, over all identified sequences, the probability of damage to the reactor core is sufficiently small (e.g., $< 1 \times 10^{-7}$) to outweigh the severe consequences.

Reliability

Reliability, as shown in Figure 1, is the probability that a system (or subsystem, component, etc.) will perform its mission under stated conditions for the desired, or mission, lifetime (Modarres & Groth, 2023). That is, reliability of a system $R(sys)$ is the probability of *nonfailure* during the mission life. Reliability is therefore proportional to the complement of risk as shown in Equation (3). They are not strictly equal because risk, in addition to probability of failure, includes measures of consequence severity.

$$R(sys) \propto 1 - \sum_{i=1}^n C(s_i) \times \Pr(s_i) \quad (3)$$

Like risk, reliability is an uncertain system characteristic and quantified with a probability distribution. Unlike risk, however, reliability is cast explicitly as a temporal characteristic that expresses, e.g., the useful life of a system and/or the mean time to (or between) failures (MTTF or MTBF, respectively). The reliability of a system at time t is the probability that a failure will occur *after* that time, as per Equation (4).

$$R(t) = \Pr(t_f > t) = 1 - F(t) \quad (4)$$

Expressing reliability as in Equation (4) immediately identifies it as the complement of system risk. The connection between reliability and resiliency, however, is less clear-cut. Reliability, as the absence of system failure, does not typically consider post-failure actions, except in the course of repairable system modeling. Although only a partial treatment of resilience (i.e., only modeling *reactive* resilience), repairable system modeling offers insights into how reliability and resiliency might be unified through the lens of *availability*, or the fraction of time a system spends operating.

Reliability, as the probability of nonfailure, is often characterized at the component level, potentially to the detriment of system-level insights (Aalund & Paglioni, 2025b). While system-level approaches to reliability are being developed for complex systems (Aalund & Paglioni, 2025a), the understanding of

system reliability is still best understood as the probability that all critical components function as designed for the duration of the system mission. In the case of a nuclear reactor, for example, reliability is the probability that its components (e.g., pumps, valves, steam generators, etc.) function for the design life. That is, reliability does not consider consequences (i.e., it is not strictly the probability of avoiding core damage).

Resiliency

As described, resilience is widely understood as the ability of systems to withstand or cope with risks, shocks or stressors. Within the understanding of resilience is also the ability to continue to maintain certain key functions or structures. In this regard, resilience relates to reliability in that a state of system performance is desired and measurable. However, it is in the manner in which this performance is maintained that we see differences between resilience and reliability.

Holling's seminal work in 1973 established the core elements of resilience around evolving stable domains, and the ability to absorb changes to state variables, driving variables, and parameters, and still persist (Holling, 1973). These elements laid the foundation for anticipating system surprises and transformation in the face of novel, and unexpected circumstances. Stability was defined as performance in the system near an equilibrium (Pimm, 1984), whereas Holling introduced resilience to describe behaviors away from these stable points.

Concepts of resilience have evolved to include discussions of single vs multiple stable points as illustrated in Figure 2. In the classical sense, authors describe resilience as the time to return to steady state conditions following a disturbance around a single equilibrium condition (Neubert & Caswell, 1997). The characteristics of resilience described by the shape of the bowl, equilibrium as a ball, and the time need to return the ball to its stable point after disturbance (Scheffer et al., 1993). Holling described this as engineering resilience (Holling, 1996) as these views closely align with engineering's goal to find and design an optimal structure and maintain

that structure. Conversely, a second view of resilience is one that observes multiple stable states. In these cases, resilience is measured by the magnitude of disturbance that can be absorbed before the system state changes or the structure of the system is altered. These views on multiple states have been considered ecological resilience to mirror these states transitions found in nature (Holling, 1996).

More recently, resilience has been extended to include aspects of adapting in the face of change driving the system to more desirable states (Folke, 2016). Resilience engineering formed as a response to overcoming limitations in risk assessment by addressing the dynamic nature of systems (Woods & Wreathall, 2003). In this domain reliability is improved by building in capacity to manage unpredictable disturbances (Hollnagel et al., 2006; Yodo & Wang, 2016). However, while accounting for dynamic and adapting system pressures, resilience engineering still focuses on preserving a single state in the system while improving the recovery time.

In the context of the nuclear power plant example, resilience may be understood as the resumption of operation following a perturbation. Note that this is distinct from both risk (as the confluence of probability and consequence of, e.g., operation failure) and reliability (as the probability of non-failure). Resiliency thus includes aspects of risk and reliability, while adding the probability (and result) of recovery.

Complex Systems Reliability and Resiliency

Tradeoffs between reliability and resiliency

Figure 2. Single versus multiple stable points as views on resilience.

A key distinction between these two views is the existence and transformation between multiple stable states. If we consider only a single state, we would then define stability as forces that keep a system near this single state. This view leads to measuring stability in terms of return time and behaviours away from and returning to a single equilibrium point (Bodin & Wiman, 2004). If instead, we view a system as having multiple stable points we would measure stability along the transition from existing states of equilibrium through growth, conservation, collapse, and reorganization. This process of change was described by Gunderson and Holling as the panarchy cycle (Gunderson & Holling, 2003).

As discussed, reliability and resiliency, although interdependent, entail tradeoffs when considering system design and operation. High-reliability systems may be fragile or brittle when faced with unexpected events. For example, the Space Shuttle was designed to be reliable, but that pursuit resulted in over 3,000 single points of failure or “Criticality 1” items (Marais et al., 2004), and all but eliminated its potential for resiliency. In the face of unexpected events (e.g., the emergence of the “O-Ring Problem”), the Shuttle, and indeed the *enterprise* proved unable to adapt. In addition to high reliability potentially leading to reduced resiliency, the pursuit of reliability, as the avoidance of system failure, may mask the importance of the adaptation and recovery mechanisms that enforce resiliency. That is,

resiliency may be inadvertently neglected during the design phase of systems.

On the other hand, some strategies for pursuing resiliency can decrease reliability. For example, redundancy can ensure the continuity of operation in the event of failure, but increases the complexity of the system overall. The increased complexity and addition of more components and interactions can complicate reliability analysis and decrease system reliability (Marais et al., 2004). That complexity and tight coupling (between components) contribute to an inability to absorb failures is the foundation of Perrow’s “Normal Accident Theory” (Perrow, 1999). Tightly coupled systems can respond to perturbations quickly, but this resiliency jeopardizes system reliability by increasing complexity and raising the probability of quick system collapse under unexpected/unexamined conditions.

Synergies connecting reliability and resiliency

While reliability and resiliency can be at odds, there are certain synergies obtained by considering both when designing and operating systems. Dual designing for reliability and resiliency requires pursuing reliability by means *other* than “simply” introducing redundancy. There are some components in complex systems that can produce the coupling required for high reliability while also supporting flexible operations to cope with unexpected events. The tradeoffs between reliability and resiliency explored above, are extrinsic tradeoffs, reinforced by the separation of the characteristics, rather than inherent contradictions. As discussed in (Marais et al., 2004), pursuing redundancy is an inefficient (i.e., ineffective and costly) method for achieving resiliency, in addition to its tendency to reduce reliability.

High reliability organizations (HROs) pursue reliability via two major pathways: proactive (preventative) measures and reactive (recovery) actions (Sutcliffe, 2011). That is, reliability-seeking socio-technical systems prize resiliency as part-and-parcel with reliability, precisely because of its emphasis on failure recovery. It is interesting that this bridge is identified through the lens of high reliability

organizations, which feature significant human involvement. Humans, though often blamed for industrial accidents, are the most flexible “components” in any system, and help enforce resiliency (Boring, 2010; Paglioni, 2024).

Bridging Reliability and Resiliency

As previously discussed, there are significant overlaps between reliability and resiliency. Assuming a resiliency perspective provides a new approach to understanding reliability – as not just the occurrence of failure, but of an *unrecovered* failure. Reliability offers an opportunity to more robustly quantify system resiliency – particularly reactive resiliency. To reiterate, the *reactive* portion of resiliency characterizes a system’s response to a perturbation, whereas *proactive* resiliency identifies appropriate planning for perturbations.

Resiliency Degree	Reliability Definition	Risk Definition	Figure 3
Better than new	System is repaired to a state better than when system was new.	System risk lower than when system was new.	Recover to Line (1)
Good-as-new	System is resilient such that reliability following repair $R=1$.	System risk equal to when system was new.	Recover to Line (2)
Same-as-old	System reliability following repair equivalent to its reliability immediately prior to failure.	System risk is equal to immediately prior to failure.	Recover to Line (3)
Worse-than-old	System reliability following repair lower than immediately prior to failure.	System risk is higher than immediately prior to failure.	Recover to Line (4)

Table 1. Degrees of system resiliency.

Reactive resiliency can be visualized by Figure 3. The recovery time of a system is the interval

between the occurrence of a failure/perturbation, and the resumption of operation. Figure 3 further identifies that, as for repair, resiliency does not *necessarily* imply a resumption of operation at the same “level.” That is, a system may be resilient in degrees, as shown in Table 1 and Figure 3.

The resiliency degree, similar to the rejuvenation or repair effectiveness parameter in repairable system modeling (Kaminskiy & Krivtsov, 2015), characterizes the restoration capacity of the resilient system. This approach to resiliency also explicitly links resiliency to reliability by casting the resiliency degree in terms of the recovered reliability. Furthermore, this approach provides a pathway to a first-approximation quantification of reactive resiliency, considering both the time required to effect the resiliency actions (t_{res}) and degree of resiliency Δ_{res} . Equation (5) outlines a first approximation of quantified reactive resiliency ρ_r .

$$\begin{aligned}\rho_r &= (1 - \Delta_{res}) \cdot \frac{t_{mission} - t_{fail} - t_{res}}{t_{mission} - t_{fail}} \\ &= (1 - \Delta_{res}) \cdot \left(1 - \frac{t_{res}}{t_{mission} - t_{fail}}\right) \quad (5)\end{aligned}$$

In Equation (5), Δ_{res} is the “gap” between perfect and effective resiliency; That is, Δ_{res} is the difference between as-new reliability (i.e., $R = 1$) and post-resiliency reliability. The times $t_{mission}$, t_{fail} , and t_{res} , are the required lifetime, failure time, and time of operation resumption, respectively. The two key components of Equation (5), recovery time and recovery degree, ensure that the quantification of resiliency is responsive to both characteristics of system resiliency. The time required to effect resiliency actions, t_{res} , may take any value (i.e., $t_{res} \in [0, \infty)$). As $t_{res} \rightarrow 0$, resiliency approaches the resiliency degree $1 - \Delta_{res}$. If the time required to effect resiliency actions exceeds the mission time, i.e., the system is not recovered prior to the anticipated end of life, the system resiliency is zero, and the system is “perfectly non-resilient.”

The degree of resiliency is $1 - \Delta_{res}$, where Δ_{res} is the difference between as-new reliability (i.e., $R = 1$) and post-resiliency actions reliability R' . Thus, there are three “domains” for Δ_{res} , namely ($\Delta_{res} >$

0), ($\Delta_{res} = 0$), and ($\Delta_{res} < 0$), corresponding to recovery to worse-than-new, same-as-new, and better-than-new, respectively. A system that is instantaneously resilient would thus have a resiliency ($\rho_s < 1$), ($\rho_s = 1$), and ($\rho_s > 1$) for worse-than-new, same-as-new, and better-than-new resiliency effectiveness. Thus, while resiliency has a lower bound of zero (perfectly non-resilient), there is technically no upper bound, as some resiliency actions could improve a system beyond its original state.

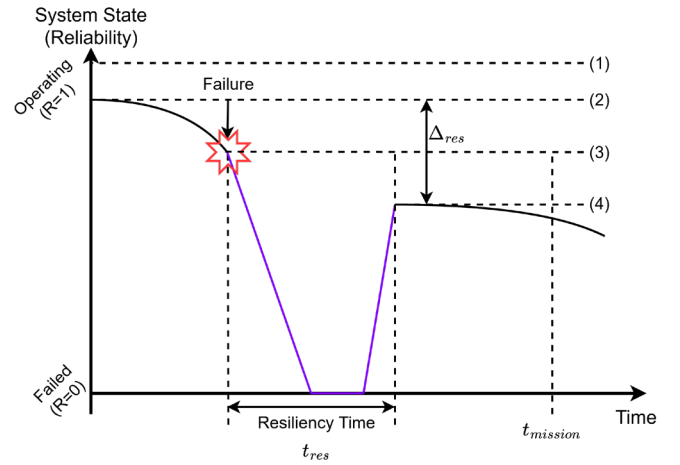


Figure 3. System resiliency visualization.

This quantification of resiliency ensures that the metric encapsulates both the time required and the recovery capacity, which are the main components of reactive resiliency. Casting reactive resiliency in terms of reliability reinforces the relationship between these two deeply-connected system characteristics.

In the nuclear power plant example, the resiliency degrees (Table 1) relate to the end-state achieved by resiliency actions on components, subsystems, or the system as a whole. Perhaps best conceptualized through the lens of a repairable component, such as a reactor coolant pump (RCP), (reactive) resiliency is the ability, and effectuation, of repairing the pump and restore the nuclear power plant to operation. “Better than new” resiliency would entail, e.g., upgrading the pump to a newer model; “Good-as-new” could entail replacing the pump with a new version of the same model; “Same-as-old” resiliency could mean repairing the immediate

damage/fault/failure to bring the pump back into operating condition; “Worse-than-old” resiliency might require a “functional repair” where the pump is restored to operation, but in a degraded state.

These component-level resiliency actions contribute to system resiliency. The nuclear power plant can shut down if an active coolant pump fails; at the very least, the system resiliency is degraded on loss of a reserve/redundant coolant pump. Thus, the component resiliency metric translates to system resiliency directly: a “better than new” pump would result in a “better than new” system.

Conclusions

Risk, reliability, and resiliency are clearly intertwined system characteristics, and the optimization of any may lead to degrading the others. For example, optimizing a system to minimize risk may in fact jeopardize resiliency, while overemphasizing resiliency (particularly reactive resiliency) may decrease reliability. However, viewing them as interconnected metrics helps to avoid these issues and allows for multi-objective approaches to system improvement. Improving resiliency through means other than instilling redundancy can simultaneously lower system risk and improve reliability.

This paper provides an overview of the current understanding of risk, reliability, and resiliency. Further, we demonstrate one approach to unifying the qualification and quantification of these metrics, by casting resiliency in terms of risk and reliability. However, there remains much work to be done to implement a unified approach to risk, reliability and resiliency. Principally, the terminological basis of the three metrics needs to be revised and updated, particularly to settle on a coherent definition of resiliency. Secondly, the mathematical approaches and modeling frameworks applicable to risk, reliability, and resiliency unification need to be developed. This work serves as the first piece in a broader approach to improving our understanding of risk, reliability, and resiliency.

Future work in this area will focus on the areas of need identified above, and principally on the

mathematics and modeling architectures that can bridge reliability and resiliency assessment. Despite the body of research in resiliency and resiliency engineering, it lacks the well-defined models and mathematics that characterize reliability engineering. On the other hand, resiliency engineering focuses on aspects of system behavior that may be more relevant to the operation of modern complex systems (i.e., the ability to recover from, rather than simply avoid, failures). Thus, both fields stand to gain from each other. This work developed a first attempt at mathematically characterizing *reactive* resiliency, drawing on similarities to repairable system modeling. Our future work will continue exploring the application of repairable system modeling techniques to the characterization of resiliency, and then developing a suitable modeling architecture.

References

- Aalund, R., & Paglioni, V. (2025a, June). Systems Engineering Approach to Design for Reliability (DfR). *Proceedings of ESREL SRA-E 2025*. European Safety and Reliability Conference (ESREL) and Safety and Society for Risk Analysis Europe (SRA-E), Stavanger, Norway.
- Aalund, R., & Paglioni, V. P. (2025b). Enhancing Reliability in Embedded Systems Hardware: A Literature Survey. *IEEE Access*, 13, 17285–17302. IEEE Access. <https://doi.org/10.1109/ACCESS.2025.3534138>
- Bodin, P., & Wiman, B. L. B. (2004). Resilience and other stability concepts in ecology: Notes on their origin, validity and usefulness. *The ESS Bulletin*, 2(2).
- Boring, R. L. (2010). Bridging resilience engineering and human reliability analysis. *10th International Conference on Probabilistic Safety Assessment and Management 2010, PSAM 2010*, 2, 1353–1360.
- Folke, C. (2016). Resilience (Republished). *Ecology and Society*, 21(4). <https://doi.org/10.5751/ES-09088-210444>

- Gunderson, L., & Holling, C. (2003). Panarchy: Understanding Transformations In Human And Natural Systems. *Bibliovault OAI Repository, the University of Chicago Press*, 114. [https://doi.org/10.1016/S0006-3207\(03\)00041-7](https://doi.org/10.1016/S0006-3207(03)00041-7)
- Holling, C. S. (1973). Resilience and Stability of Ecological Systems. *Annual Review of Ecology, Evolution, and Systematics*, 4(Volume 4, 1973), 1–23. <https://doi.org/10.1146/annurev.es.04.110173.000245>
- Holling, C. S. (1996). Engineering Resilience versus Ecological Resilience. In P. C. Schulze (Ed.), *Engineering Within Ecological Constraints* (1st ed., pp. 31–45). National Academy Press.
- Hollnagel, E., Woods, D. D., & Leveson, N. (Eds.). (2006). *Resilience Engineering: Concepts and Precepts*. CRC Press Taylor & Francis Group.
- Kaminskiy, M., & Krivtsov, V. (2015). Geometric G1-Renewal process as repairable system model. *2015 Annual Reliability and Maintainability Symposium (RAMS)*, 1–6. <https://doi.org/10.1109/RAMS.2015.7105174>
- Kaplan, S., & Garrick, B. J. (1981). On The Quantitative Definition of Risk. *Risk Analysis*, 1(1), 11–27. <https://doi.org/10.1111/j.1539-6924.1981.tb01350.x>
- Marais, K., Dulac, N., & Leveson, N. (2004, March 29). *Beyond Normal Accidents and High Reliability Organizations: The Need for an Alternative Approach to Safety in Complex Systems*. MIT Engineering Systems Division Symposium, Cambridge, MA. <http://sunnyday.mit.edu/papers/hro.pdf>
- Modarres, M., & Groth, K. M. (2023). *Reliability and Risk Analysis* (2nd ed.). CRC Press, Taylor & Francis Group.
- Neubert, M. G., & Caswell, H. (1997). Alternatives to Resilience for Measuring the Responses of Ecological Systems to Perturbations. *Ecology*, 78(3), 653–665. [https://doi.org/10.1890/0012-9658\(1997\)078%5B0653:ATRFMT%5D2.0.CO;2](https://doi.org/10.1890/0012-9658(1997)078%5B0653:ATRFMT%5D2.0.CO;2)
- Paglioni, V. (2024, October). *A System Objectives-based Proposal for Modeling Resilience in Nuclear Power Plant Operations*. Pacific Basin Nuclear Conference 2024, Idaho Falls.
- Perrow, C. (1999). *Normal Accidents: Living with High Risk Technologies—Updated Edition*. Princeton University Press; eBook Collection (EBSCOhost). <https://search.ebscohost.com/login.aspx?direct=true&AuthType=cookie,ip,url,cpid&custid=s4640792&db=nlebk&AN=430832&sitelink=bsi-live>
- Pimm, S. L. (1984). The complexity and stability of ecosystems. *Nature*, 307(5949), 321–326. <https://doi.org/10.1038/307321a0>
- Scheffer, M., Hosper, S. H., Meijer, M.-L., Moss, B., & Jeppesen, E. (1993). Alternative equilibria in shallow lakes. *Trends in Ecology & Evolution*, 8(8), 275–279. [https://doi.org/10.1016/0169-5347\(93\)90254-M](https://doi.org/10.1016/0169-5347(93)90254-M)
- Sutcliffe, K. M. (2011). High reliability organizations (HROs). *Best Practice & Research Clinical Anaesthesiology*, 25(2), 133–144. <https://doi.org/10.1016/j.bpa.2011.03.001>
- Woods, D., & Wreathall, J. (2003). *Managing Risk Proactively: The Emergence of Resilience Engineering*.
- Yodo, N., & Wang, P. (2016). Engineering Resilience Quantification and System Design Implications: A Literature Survey. *Journal of Mechanical Design*, 138(111408). <https://doi.org/10.1115/1.4034223>

Biography



Vincent Paglioni

Dr. Vinnie Paglioni is an Assistant Professor in the Department of Systems Engineering at Colorado State University, where his research and teaching center around risk, reliability, and resiliency analysis for complex engineering systems. He holds a B.S. in Nuclear & Radiological Engineering from Georgia Tech (2017) and M.S. (2022) and Ph.D. (2023) in Reliability Engineering from the University of Maryland, College Park.



Aaron Brown

Dr. Aaron Brown, faculty in Colorado State University's Systems Engineering Department, has over two decades of experience in academia and industry. His research focuses on humanitarian engineering and equitable energy access. He has led the creation of innovative engineering programs, conducted research in renewable energy and social impact systems, and contributed to aerospace projects including the Mars Curiosity landing mechanism and Hubble robotics mission. His work has earned honors such as the Martin Luther King Jr. Peace Award and the U.S. State Department's "Visionary Partner" recognition.



Steven Conrad

Steve Conrad is an Associate Professor in the Department of Systems Engineering at Colorado State University. He holds B.S. degrees in Optical Engineering and Psychology from the University of Arizona, an M.S. in Environmental Technology Management from Arizona State University, and a Ph.D. in Resource and Environmental Management from Simon Fraser University, Canada. His research focuses on decision-making and resilient system modeling, principally for the water-energy nexus. Dr. Conrad previously chaired the Future Waters Research Excellence Cluster at the University of British Columbia and served as Associate Director of Simon Fraser University's Pacific Water Research Centre.