

On Compound Hazards in Cost-Effective Recovery Management

Walaa Bashary, Lukas Halekotte, Andrea Mentges, Daniel Lichte

*German Aerospace Center (DLR), Institute for the Protection of Terrestrial Infrastructures, Germany,
E-mail: walaa.bashary@dlr.de, lukas.halekotte@dlr.de*

The failure of an infrastructure component can compromise the overall performance of a system. Approaches to quantifying component vulnerability typically focus on single-hazard scenarios, neglecting, for example, the effect of consecutive hazards on the system, which could increase the vulnerability of the system over time. This can lead to inaccurate risk or resilience assessments and consequently to inadequate decision-making. In this work, we apply a framework for modeling the evolution of component vulnerabilities in response to compound hazards using state-dependent fragility curves. The approach employs a Markovian state transition model in which a hazard event can alter the state of system components. Using this approach, we illustrate how the interaction between hazards can affect the degradation of components. We show that the compounding effect of two hazardous events depends on both the time lag between the two and the recovery capacity of the affected components. We illustrate this relationship by analyzing hazard time series with different correlation strength, and by defining a minimal recovery rate required to maintain system performance below a prescribed failure threshold. Finally, we incorporate a cost analysis by linking recovery and replacement processes to cost functions, allowing the computation of these costs over time. This allows the identification of appropriate strategies when hazards are correlated versus independent.

Keywords: Natural hazards, compound hazards, fragility curves, Markov chain, infrastructure, recovery function, cost-benefit analysis

1. Introduction

Infrastructures are increasingly exposed to compound events rather than single ones (Moftakhari and AghaKouchak, 2019), a trend that is expected to increase as climate drivers shift (Ward et al., 2020). Compound hazards may arise from the temporal clustering, spatial concurrence or sequential interaction of multiple hazardous processes (Zscheischler et al., 2018; Gill and Malamud, 2014), leading to impacts that differ substantially from those expected under independent events (Kappes et al., 2012).

Natural hazards such as floods, heatwaves, and storms often exhibit temporal overlap, clustering in time due to shared driving processes (Zscheischler et al., 2020). When hazards are temporally correlated, their impacts on system components are not necessarily additive: degradation accelerates, and thus the window for intervention narrows (Ebi, 2024). Despite growing recognition of compound hazards, many approaches continue to assume independent hazard occurrences (Hochrainer-Stigler et al., 2023), potentially un-

derestimating both risk and long-term costs (Stalhandske et al., 2024).

A key limitation of many existing approaches is their reliance on single-hazard fragility curves that do not account for changes in component condition following previous events. While some studies consider multi-hazard fragility or fragility surfaces (Harati and van de Lindt, 2024; Fu et al., 2020), the explicit temporal evolution of component states – particularly when recovery processes are included – has received limited attention. Moreover, the interaction between hazard timing, recovery capacity, and economic cost remains poorly explored.

This paper addresses this gap by proposing a framework that links hazard correlation, component state evolution, and recovery dynamics. Specifically, this work quantifies the impact of hazard timing and temporal correlation on component degradation trajectories, defines the minimum recovery rate required to maintain acceptable performance under repeated hazards, and demonstrates how this recovery rate systematically depends on hazard correlation and timing.

Additionally, the framework links recovery strategies to expected costs, enabling cost–benefit comparisons across different hazard scenarios.

2. Method

2.1. General framework

In this study, we build upon the framework presented in our previous work (Bashary et al., 2025). We assumed that a component's state can be described by a set of mutually exclusive discrete levels, ranging from an undamaged state, S_1 , to a complete failure state, S_n , with each state corresponding to a different degree of vulnerability. Furthermore, the component is potentially subject to m distinct hazards whose intensities are described by a single intensity measure each (all hazard intensity measures then constitute the hazard space: $H_1, \dots, H_m$). The transition from one state to another, as a result of the impact of one particular hazard i , was modeled probabilistically, using a $n \times n$ transition matrix $\mathbf{T}_i$ with $i = 1, \dots, m$ that described the likelihood of moving between states under the influence of each hazard H_i . Each entry of the transition matrix $\mathbf{T}_i$ was defined by the corresponding state-dependent fragility curve, modeled as lognormal cumulative distribution functions.

$$\mathbf{T}_i = \begin{bmatrix} (S_1 \rightarrow S_1) & \cdots & (S_1 \rightarrow S_n) \\ \vdots & \ddots & \vdots \\ 0 & \cdots & (S_n \rightarrow S_n) \end{bmatrix} \quad (1)$$

In general, real-world systems undergo external interventions that potentially counteract some degree of degradation at certain time steps. This recovery capacity was modeled, similarly, by a recovery matrix, $\mathbf{R}$ that described the rate at which the component's state returned to an intact state.

$$\mathbf{R} = \begin{bmatrix} (S_1 \rightarrow S_1) & \cdots & 0 \\ \vdots & \ddots & \vdots \\ (S_n \rightarrow S_1) & \cdots & (S_n \rightarrow S_n) \end{bmatrix} \quad (2)$$

To model the evolution of the component's state (i.e., probability state vector P_t) over time, we adopted a discrete Markov chain, where the com-

ponent's probability state vector P_t at time t is given by:

$$P_t = P_{t-1} \left[\prod_{i=1}^m \mathbf{T}_i \right] \mathbf{R}, \quad (3)$$

where P_{t-1} is the state probability vector at time step $t - 1$. $\mathbf{T}_i$ are the transition matrices for the hazards, assuming that there is a predefined order with which hazards affect the component. These transition matrices are time-dependent, as they are parametrized by hazard intensities that vary over time and are modeled as random variables drawn from their probability distribution. At each time step, realizations of the hazard intensities are sampled and then used in the corresponding transition matrix. Lastly, $\mathbf{R}$ is the recovery matrix, which takes place only after the effect of all hazards, within the time step.

2.2. Demonstrator

To illustrate the use case, we adopt a system demonstrator made of a total of three states (intact S_1 , damaged S_2 and failed S_3) under the influence of two hazards. We assume that H_1 is only capable of pushing the component to S_3 , therefore the transitions are depicted by two fragility curves: f_{13} if the component was in S_1 prior to the event, and f_{23} if the component was instead in S_2 . For H_2 , we assume that the hazard is only capable of pushing the component from S_1 to S_2 (fragility curve g_{12}). This reflects a dynamic where one hazard is capable of bringing the component to failure, regardless of the current state, while the other hazard only stresses the component (but does not bring it to failure). Consequently, the component's evolution over time is given by:

$$P_t = P_{t-1} \mathbf{T}_2 \mathbf{T}_1 \mathbf{R}, \quad (4)$$

where P_t and P_{t-1} are the state probability vector at time step t and $t - 1$, respectively. $\mathbf{T}_1$ and $\mathbf{T}_2$ are the transition matrices for the two hazards, assuming that the hazards affect the component in a predefined order. Lastly, $\mathbf{R}$ is the recovery matrix, which in this demonstrator contains only one element (r_{21} , recovery from S_2 to S_1), and it takes

place only after the effect of both hazards, within the time step. The demonstrator is implemented through Monte Carlo simulations based on 1000 realizations of randomly drawn time series of hazard intensities, for 300 time steps, and presented in terms of the median outcomes of S_3 . A summary of the simulation parameters used in this study is presented in Table 1.

Table 1. Simulation Parameters

Parameter	Value
Log-normal	
f_{13}	(3, 0.4)
f_{23}	(2, 0.3)
g_{12}	(1, 0.6)
BVN – independent	
Mean vector	[2.0, 0.5]
Covariance matrix	$\begin{bmatrix} 3 & 0 \\ 0 & 1 \end{bmatrix}$
BVN – low correlation	
Mean vector	[2, 0.5]
Covariance matrix	$\begin{bmatrix} 3 & 0.5\sqrt{3} \\ 0.5\sqrt{3} & 1 \end{bmatrix}$
BVN – high correlation	
Mean vector	[2, 0.5]
Covariance matrix	$\begin{bmatrix} 3 & \sqrt{3} \\ \sqrt{3} & 1 \end{bmatrix}$

Adopting the model presented above, we investigate how the timing between correlated hazards affects component degradation, with a particular focus on the time available for recovery between successive events. Specifically, we examine how closely spaced hazards constrain recovery actions and limit the component’s ability to return to an undamaged state, and therefore accelerate degradation (Section 3.1). Building on this, we analyze the recovery effort required to counteract these effects (Section 3.2) and assess the associated costs for different recovery strategies (Section 3.3).

3. Results

3.1. Modeling temporally correlated hazards

Real-world settings show how hazards often have common drivers, which results in hazard inten-

sities being temporally correlated. To take into account the effect of hazard correlation, we model hazard occurrences using a bivariate normal distribution (BVN), in three correlation settings, representing no correlation (i.e., independent hazards), low correlation, and high correlation (Figure 1).

The timing of subsequent hazards can be a critical factor in determining their compound impact. When two hazards occur in quick succession, the system has less time to recover from the impact of the first before being subjected to the second hazard, which potentially accelerates the system’s degradation. To investigate this, we examine how varying the time lag between two correlated hazards affects the degradation of the component state. To this end, we first create correlated time series of intensities for each hazard and then shift one of the two time series by a number of time steps (representing the time lag between the two hazards), so that high intensities of H_2 are to be expected a corresponding number of time steps later.

Our results (Figure 2) show that without recovery (Fig. 2(a) and Fig. 2(c)), the influence of the lag is absent, as more time between hazard events does not improve the component’s state (it cannot return to its unimpaired state S_1). However, when a recovery mechanism is present (Fig. 2(b) and Fig. 2(d)), the effect of the lag is considerable. In fact, increasing the time lag between hazard events reduces their compound impact as the component has more time to return to its intact state S_1 . Ultimately, once a critical time lag is exceeded, two subsequent hazards can no longer be considered to present a compounding effect. This phenomenon is evident in the curves of Figure 2, where for larger lags the trajectories for correlated hazards increasingly resemble those for independent hazards.

3.2. Minimal recovery rate and logistic function fit

The previous results have shown how the interplay between the timing of hazard events and the recovery dynamics of a component influence the progression of its degradation. We now focus on quantifying the recovery effort required to coun-

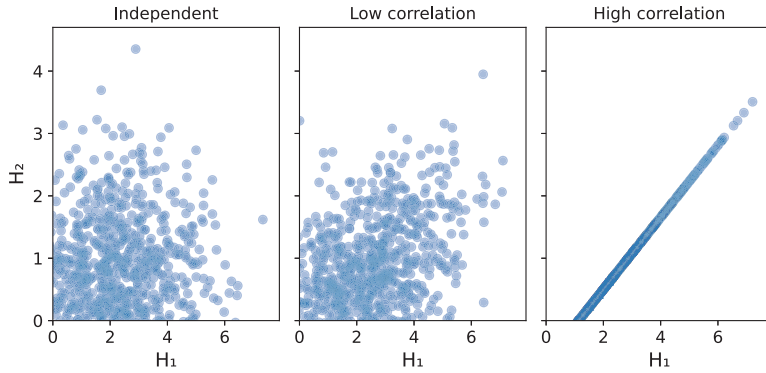


Fig. 1. Scatter plots of bivariate normal samples illustrating independent, low-correlation, and high-correlation cases.

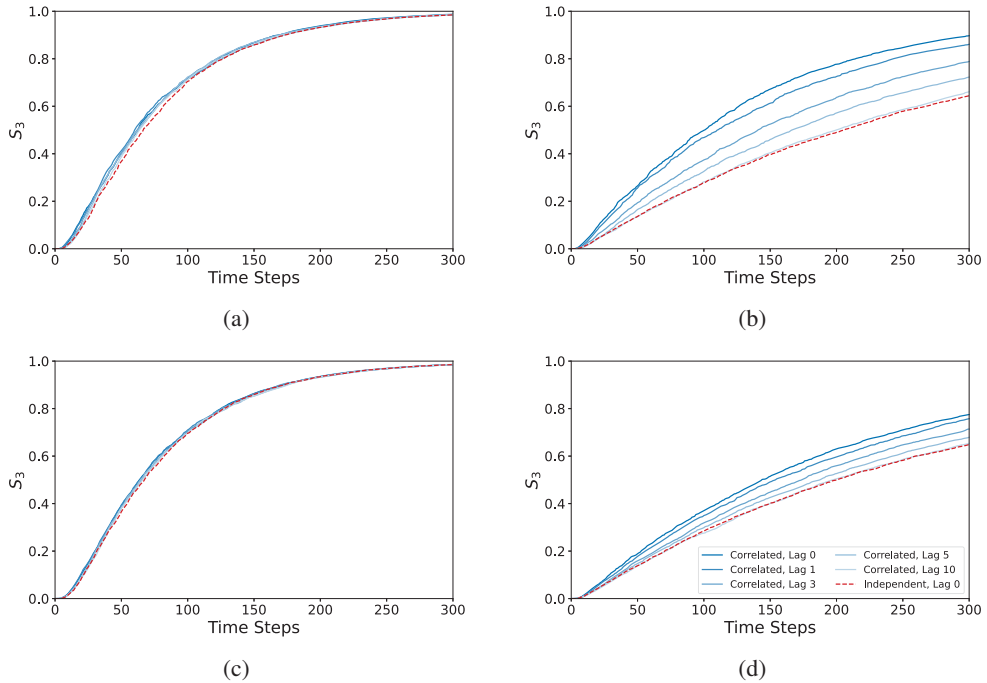


Fig. 2. Temporal evolution of S_3 probability under different recovery rates and correlation strength. Within each subplot, trajectories represent median outcomes corresponding to different lag values. (a) High correlation and no recovery ($r_{21} = 0$). (b) High correlation and recovery ($r_{21} = 0.2$). (c) Low correlation and no recovery. (d) Low correlation and recovery ($r_{21} = 0.2$).

teract the effect of compound hazards. In particular, we examine how high the recovery rate must be to keep the probability of being in the failed

state S_3 below a specified critical level. This provides a way to compare how different hazard interaction structures (correlated versus uncorrelated)

translate into different recovery requirements for achieving the same level of system performance.

To quantify this recovery requirement, we define $r_{21,min}$ as the minimum recovery rate from S_2 to S_1 required to keep the probability of $S_3(t)$ below a predefined threshold ($S_{3,th}$), for the given time horizon. In a real system, this formulation corresponds to a decision-making setting in which the characteristics of the hazard process are not directly controllable, whereas the recovery process can be influenced through the allocation of resources (i.e., recovery effort). The threshold $S_{3,th}$ represents, for example, a reliability or performance requirement or constraints imposed by design. Under this interpretation, $r_{21,min}$ quantifies the minimum recovery capacity that must be ensured to meet a prescribed performance target over the considered time horizon.

To determine the minimum recovery rate required for a given failure probability threshold $S_{3,th}$, we proceed as follows: for a fixed hazard interaction setup, which is defined by a specific correlation strength and lag, we simulate the component state's evolution over a range of recovery rates r_{21} . For each value of r_{21} , we compute the probability of the component being in S_3 at the selected time horizon (300 time steps). The results exhibit a nonlinear relationship between r_{21} and S_3 which depends on the temporal structure of the hazard interaction. In particular, in the absence of a lag, the impact of the recovery is limited, as the temporal clustering of the hazards does not allow sufficient time to transition back to S_1 . Consequently, increasing r_{21} leads to only minor reductions in S_3 resulting in an almost vertical trend. As the lag increases, the influence of the hazard correlation weakens. For example, for lag = 10 the curve is approaching the independent curve, which effectively decouples the hazards. Even a small lag (lag = 1) is sufficient to reduce temporal clustering and allow the recovery process to act more effectively. This difference is most evident as r_{21} approaches 1: without a lag, a direct transition pathway through state S_2 remains active, whereas for any lag scenario this pathway is effectively removed, and all curves converge to the same limit. In contrast, for r_{21} approaching 0,

all cases converge since recovery has no influence.

These dynamics are approximated by a four-parameter logistic function. However, due to the boundary behavior described above, the fit is less accurate for small lags (e.g., lag = 1), where the curve exhibits sharper transitions that are not fully captured by the logistic form.

$$f(x) = \frac{L}{1 + e^{-k(x-x_0)}} + d, \quad (5)$$

where L denotes the upper asymptote, k represents the steepness of the curve, x_0 is the inflection point at which the function transitions most rapidly, and d accounts for the shift in r_{21} . The parameters were estimated using a nonlinear least-squares approach. Figure 3 illustrates the minimum recovery rate required to maintain the probability of S_3 below $S_{3,th}$ for a given number of time steps and for various temporal lags for the low correlation scenarios.

The logistic fit provides a useful approximation of the simulation results, allowing interpolation between thresholds without the need to rerun the simulations. This is useful for exploring how the minimum recovery rate would change under different system configurations, e.g., hazard correlation and temporal lag.

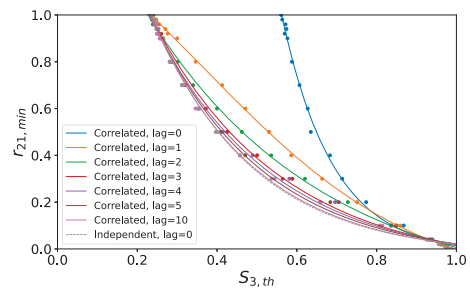


Fig. 3. Logistic-function fits of the minimum recovery rate required to maintain the probability of the failed state below a given threshold for the low correlation scenario. Curves show correlated hazards with different lags, while the dashed line represents the independent hazard case. Dots are the simulation results.

3.3. Cost–benefit analysis

The objective of this section is to illustrate how the economic feasibility of implementing recovery capacity, for example by adopting the calculated $r_{21,min}$ from Section 3.2), can be examined within the proposed framework. To evaluate these economic implications, we associate a cost function to the recovery process that goes from the intermediate damage state S_2 to the intact state S_1 , governed by the recovery rate r_{21} . Since the framework is probabilistic, the recovery process is modeled in terms of expected interventions – which constitute the preventive interventions. At each time step, the expected number of interventions from S_2 to S_1 is given by the product of the probability of being in state S_2 and the recovery rate r_{21} . Multiplying this quantity by the unit intervention cost ($c_{21} = 10$) and accumulating over time returns the cumulative expected cost:

$$C_{21}(T) = \sum_{\tau=1}^T S_2(\tau) r_{21} c_{21} \quad (6)$$

Additionally, we compute the expected cost at the end of the simulation horizon to replace all components that are in S_3 . These interventions are replacement interventions ($c_{31} = 250$) and their expected costs are computed as:

$$C_{31}(T) = S_3(T) c_{31} \quad (7)$$

Figure 4 shows the cumulative cost at time T as a function of the recovery rate r_{21} for various lags within the low correlation scenario. We see that with increasing r_{21} (i.e., higher effort), the associated costs also increase. Additionally, the figure shows the cost associated with the replacement intervention at time T . In this case, costs diminish nonlinearly with the increase of the recovery rate r_{21} (as the probability of the component being in S_3 decreases). Considering the costs for intervention and replacement for different time lags allows us to highlight the economic implications that short time spans between correlated hazards might have. In particular, we see how for higher lags, costs for preventive interventions increase due to the increase of the time window between

events where interventions are possible. However, the associated cost of replacement intervention decreases as eventually the components are less probable to be in the failed state.

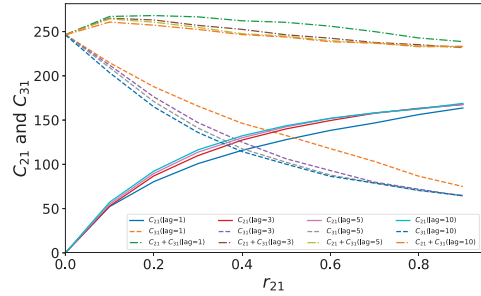


Fig. 4. Cumulative expected recovery cost (over 300 time steps) as a function of the recovery rate r_{21} for various lags and low correlation scenario. Solid lines represent r_{21} cost; dashed lines indicate r_{31} cost; Dash-dot lines are the total cost.

While Figure 4 shows how costs increase with recovery effort, it does not indicate whether the recovery rate required to achieve a desired level of performance is economically feasible. In particular, if the minimum recovery rate needed to substantially reduce the failure probability entails prohibitively high costs, it becomes relevant to assess whether lower recovery rates are still acceptable. To quantify the benefit of increasing the recovery rate r_{21} , we compute the reduction in the probability of S_3 relative to the no-recovery baseline scenario (Figure 5). Also in this case the figure displays a nonlinear trend with higher lags resulting in a higher benefit (i.e., more time to intervene successfully) compared to lower lags (where the time between events is scarce). Notably, very low recovery rates (e.g., 0.1) lead to the highest total cost, indicating the existence of an economically suboptimal region. Although, from a cost perspective, no recovery can appear equivalent to high recovery investment, this equivalence may not hold in terms of system performance. Evaluating component-level performance is therefore important, but beyond the scope of this work.

To assess cost–benefit trade-off associated with the recovery intervention, economic and perfor-

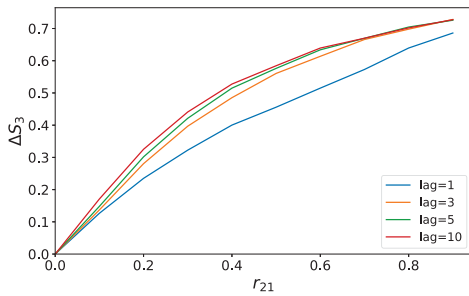


Fig. 5. Benefit of preventive intervention as a function of the recovery rate r_{21} (at 300 time steps) for various lags and low correlation scenario

formance outcomes are linked. Figure 6 brings together the information conveyed separately in Figures 4 and 5, by plotting the expected preventive cost against the reduction in the final failure probability for different values of r_{21} (as each point represents a distinct recovery strategy) and for various lags in the low correlation scenario.

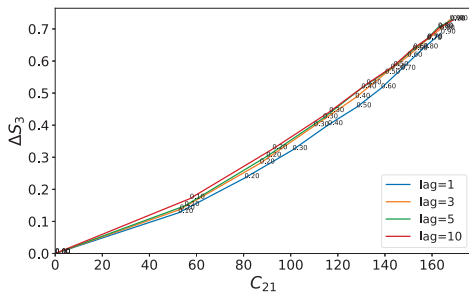


Fig. 6. Cost-benefit analysis for different lags and recovery rates (numbers on the curves). Benefit is represented as reduction in the final probability of the failed state as a function of the preventive intervention cost.

This plot allows us to identify recovery strategies that balance affordability and effectiveness. For instance, for a desired benefit, the associated cost could be potentially unfeasible, therefore an increase of the budget might be needed or alternatively, a reduction in the expected benefit is considered acceptable. Moreover, the presence of multiple lag scenarios can be interpreted as a source of uncertainty in the timing of future events: when the lag is not known a priori, re-

covery strategies must be evaluated not only in terms of expected cost and benefit, but also in terms of their robustness across a range of possible temporal configurations.

4. Discussion and conclusion

This work investigated how temporally correlated hazards influence the degradation, recovery, and recovery cost of infrastructure components, with a particular focus on the interaction between hazard timing and correlation strength. By adopting state-dependent fragility curves and a Markovian state-transition framework, we explicitly represented the evolution of component condition under compound hazards and incorporated recovery interventions as state-dependent processes. This allowed us to move beyond static, single-hazard fragility assessments and to explore how component vulnerability evolves dynamically over time.

The results highlight that the effects of compound hazards cannot be understood solely in terms of hazard intensity, frequency, and fragility. When hazards are closely spaced in time, recovery opportunities are limited, leading to accelerated degradation and a higher probability of reaching failure states. Conversely, increased time lags between hazards decouple their impacts, provided that recovery mechanisms are sufficiently effective.

We proposed a methodology to address the identification of recovery intervention by calculating a minimum recovery rate required to maintain acceptable performance levels under compound hazards. By defining a failure probability threshold and varying recovery rates, time lags, and correlation strength, we showed that recovery requirements increase nonlinearly as hazards become more temporally clustered. The observed relationship between recovery demand and the failure probability threshold is well captured by logistic functions, providing a quick and compact representation of recovery needs under different configurations. By linking recovery actions to explicit cost functions, we demonstrated how economic considerations could be accounted for. This relationship is ultimately expressed by a cost-benefit analysis.

Several limitations remain and should be addressed for future research. First, the analysis focuses on a single component represented by a small number of discrete damage states. While this abstraction facilitates the model construction, real infrastructure systems consist of multiple interacting components, potentially with different characteristics. Second, recovery processes are represented using constant rates, whereas in practice recovery capacity may vary over time due to resource constraints, prioritization, or disruptions. Third, the hazard generation model, while flexible, remains simplified and does not capture all possible features of real hazard processes, such as long-term shifts or the use of multiple intensity measures to describe the hazard. In summary, future research will aim to adopt the framework in a real environmental setting and demonstrate its applicability using evidence-based data.

Overall, this study demonstrates that accounting for hazard dynamics, recovery interventions, and cost within a unified modeling framework provides insights that are missing in traditional single-hazard approaches.

References

- Bashary, W., L. Halekotte, A. Mentges, and D. Lichte (2025, November). Addressing compound hazards by using state-dependent fragility functions. In *2025 9th International Conference on System Reliability and Safety (ICSRS)*, pp. 11–16. IEEE.
- Ebi, K. L. (2024, December). Understanding the risks of compound climate events and cascading risks. *Dialogues on Climate Change* 2(1), 33–37.
- Fu, X., H.-N. Li, G. Li, and Z.-Q. Dong (2020, April). Fragility analysis of a transmission tower under combined wind and rain loads. *Journal of Wind Engineering and Industrial Aerodynamics* 199, 104098.
- Gill, J. C. and B. D. Malamud (2014, October). Reviewing and visualizing the interactions of natural hazards: Interactions of natural hazards. *Reviews of Geophysics* 52(4), 680–722.
- Harati, M. and J. W. van de Lindt (2024, October). Data-driven machine learning for multi-hazard fragility surfaces in seismic resilience analysis. *Computer-Aided Civil and Infrastructure Engineering* 40(6), 698–720.
- Hochrainer-Stigler, S., R. Šakić Trogrlić, K. Reiter, P. J. Ward, M. C. de Ruiter, M. J. Duncan, S. Torresan, R. Ciurean, J. Mysiak, D. Stuparu, and S. Gottardo (2023, May). Toward a framework for systemic multi-hazard and multi-risk assessment and management. *iScience* 26(5), 106736.
- Kappes, M. S., M. Keiler, K. von Elverfeldt, and T. Glade (2012, July). Challenges of analyzing multi-hazard risk: a review. *Natural Hazards* 64(2), 1925–1958.
- Moftakhari, H. and A. AghaKouchak (2019, October). Increasing exposure of energy infrastructure to compound hazards: cascading wildfires and extreme rainfall. *Environmental Research Letters* 14(10), 104018.
- Stalhandske, Z., C. B. Steinmann, S. Meiler, I. J. Sauer, T. Vogt, D. N. Bresch, and C. M. Kropf (2024, March). Global multi-hazard risk assessment in a changing climate. *Scientific Reports* 14(1).
- Ward, P. J., V. Blauhut, N. Bloemendaal, J. E. Daniell, M. C. de Ruiter, M. J. Duncan, R. Emberson, S. F. Jenkins, D. Kirschbaum, M. Kunz, S. Mohr, S. Muis, G. A. Riddell, A. Schäfer, T. Stanley, T. I. E. Veldkamp, and H. C. Winsemius (2020, April). Review article: Natural hazard risk assessments at the global scale. *Natural Hazards and Earth System Sciences* 20(4), 1069–1096.
- Zscheischler, J., O. Martius, S. Westra, E. Bevacqua, C. Raymond, R. M. Horton, B. van den Hurk, A. AghaKouchak, A. Jézéquel, M. D. Mahecha, D. Maraun, A. M. Ramos, N. N. Ridder, W. Thiery, and E. Vignotto (2020, June). A typology of compound weather and climate events. *Nature Reviews Earth & Environment* 1(7), 333–347.
- Zscheischler, J., S. Westra, B. J. J. M. van den Hurk, S. I. Seneviratne, P. J. Ward, A. Pitman, A. AghaKouchak, D. N. Bresch, M. Leonard, T. Wahl, and X. Zhang (2018, May). Future climate risk from compound events. *Nature Climate Change* 8(6), 469–477.