

Uncertainty Quantification in Reactor Safety: Time Dependency, Available Information, and Measurability

Martin Wortman and Ernest Kee

HazTechRisk.Org

The Organization for Public Awareness of Hazardous Technology Risks

2221 Market St. Suit 100

Galveston, TX 77550

wortman@haztechrisk.org and kee@haztechrisk.org

21/02/2022

Abstract

Uncertainty Quantification (*UQ*) in the analysis of reactor safety is challenging. The power of *UQ* derives from its grounding in probability theory.¹ But, certain important safety related events are not probability-measurable which is problematic for risk-analytic methodologies that rely on *UQ* computations. In this note we identify why the dynamics of available engineering information plays an essential role in governing the fidelity of uncertainty quantification, and why uncertainty quantification in reactor safety is limited by the un-measurability of certain critical events. We provide an historical example providing a practical context for our observations. Finally we discuss the implications of measurability for regulatory decision-making governed by recent nuclear industry legislation that advocates increased use of risk informed, performance-based regulation for advanced reactor licensing.

Keywords— Uncertainty Quantification, FMEA, Failure Modes, Stopping Times, Filtration

¹We hold that probability theory is the only viable theory of uncertainty.

*But there are also unknown unknowns —
the ones we don't know we don't know.*

— Donald Rumsfeld

1 Introduction

Engineered designs are created with an intent to avoid failures – or at least understand the uncertainty surrounding when and what failures may occur. This is especially true when protections are designed for hazardous systems (*e.g.*, nuclear reactors). Engineers are well aware that any device or system can break down unexpectedly due to previously undiscovered failure modes. Failure modes that are first discovered under circumstances that if not immediately addressed might exceed to catastrophe are especially troubling. Extensive records are kept on such “near miss incidents”^{123450[-]0}, for example the Nuclear Regulatory Commission (*NRC*) Licensee Event Report (*LER*) database which documents many such failures found in many different designs. Following near miss incidents, root cause analysis is used to isolate the immediate cause for the breakdown. Engineers proficient in root cause analysis can determine if a breakdown arose from something that can be assigned to physics, improper operation, or improper maintenance. Knowing that unanticipated failure modes will be experienced, safety margin and defense in depth are primary protections against dangers that may be present. Engineers use logic structures, such as Failure Mode and Effects Analysis (*FMEA*), HazOp, fault trees, and event trees, to help study scenarios that have little, or no back up and may require additional consideration of safety margin or defense in depth. If a single point of failure cannot be avoided, they are reluctant to accept the design unless the device has been tested thoroughly and includes substantial safety margin. Regular inspections can help ensure the design retains safety margin over the device service life.

We use the terminology unanticipated failure modes to indicate failure modes can only be found at a future time. The potential for their appearance in service is well understood by engineers. Prior to release of potentially hazardous processes or potentially hazardous products, good engineering practice dictates a cautious approach with a gradual ramp up to full production.

Effective maintenance organizations address unexpected failures in critical equipment by aggressively pursuing root cause determination followed up with prompt corrective action. Although quantitative risk assessments work with probabilities of functional failures, the functional failures they work with are collections of one or more failure modes including any unexpected ones. The importance of unanticipated failure modes is that they cannot be included in quantitative assessments such as Probabilistic Risk Assessment (*PRA*).

The 2019 Nuclear Energy Innovation and Modernization Act (*NEIMA*) asks for, in part, “risk-informed” regulation by the *NRC*. While the *NRC* currently uses risk-informed decision-making to guide certain policy decisions as well as prioritizing inspection and enforcement, they define risk-informed to include prudent engineering practices that include defense in depth and safety margin. Reliance on such engineering practices are captured in the *NRC* definition of risk-informed decision-making (*NRC*, 2013):

“A ‘risk-informed’ approach to regulatory decision-making represents a philosophy whereby risk insights are considered together with other factors to establish requirements that

better focus licensee and regulatory attention on design and operational issues commensurate with their importance to health and safety. A ‘risk-informed’ approach enhances the traditional approach by: (a) allowing explicit consideration of a broader set of potential challenges to safety, (b) providing a logical means for prioritizing these challenges based on risk significance, operating experience, and/or engineering judgment, (c) facilitating consideration of a broader set of resources to defend against these challenges, (d) explicitly identifying and quantifying sources of uncertainty in the analysis, and (e) leading to better decision-making by providing a means to test the sensitivity of the results to key assumptions. Where appropriate, a risk-informed regulatory approach can also be used to reduce unnecessary conservatism in deterministic approaches, or can be used to identify areas with insufficient conservatism and provide the bases for additional requirements or regulatory actions.”

The un-measurability arrival times of undiscovered failure modes reinforces the need for defense in depth and safety margin to be included to help ensure the unbounded level of risk not included in uncertainty quantification is taken into account in protective system regulation.²

We find a relatively large literature that would inform investigators on the nature of predictive modeling under uncertainty induced by random variables and probabilities that appear as random variables. A few of the authoritative studies and texts that we have reviewed over time are .

In the following, we briefly review a relevant, non-consequential, example of an unexpected failure mode from an operating nuclear power plant and then provide the analytical formalism clarifying why catastrophes arising from unanticipated failure modes defy uncertainty quantification. We conclude with a brief discussion asserting that over-reliance on *UQ* should be avoided in risk analyses of safety-critical protections.

2 An example

On December 18, 1995, the South Texas Project (STP) Unit 1 commercial nuclear reactor experienced an unexpected reactor trip due to a main and auxiliary transformer lockout while operating at 100% power as described in Head (1996). The transformer lockouts triggered a series of protective system actuations:

1. The generator output breaker tripped open,
2. The main turbine tripped, causing a reactor trip, and
3. All four 13.8 kV auxiliary bus breakers opened resulting in a loss of offsite power to the A Train Engineered Safeguards Features Bus.

The operators entered the emergency operating procedure for reactor/turbine trip which required verification of subcritical reactor and control rods fully inserted. However, three control rods, designated F10, C9, and N7, representing 3 out of a total of 57 control rods required to fully insert, indicated they had instead stopped at about 6 steps above the bottom of the core as required.³ All the control rods that had stopped prior to complete insertion were operating in a new fuel design.

²For example, Title 10 of the Code of Federal Regulations Part 50, Appendix A “General Design Criteria” requires such prudent engineering practices.

³One step corresponds to an increment of less than 1/2 of an inch of vertical motion.

Proper operation of the reactor control rods is verified prior to full power operation after the reactor core is replaced during a refueling outage. In protective operation, the rods fall due to the force of gravity in two basic stages, a free fall stage followed by a braking stage. They are tested to verify they fall within a prescribed length of time, from time of release to the time they start braking. This test was performed successfully as normally done for the new core installed after the refueling outage prior to the event on December 18.⁴

2.1 Root cause

Other similar events started occurring at other reactor plants after the *STP* October 18 event (Crutchfield, 1996). Root cause was investigated at the *STP* over the next few days and it was concluded that the fuel assemblies, which had been exposed to radiation in at least one fuel cycle were becoming susceptible to buckling as partially described by Kee (1996). The root cause was found by developing a physical process theory model that was fit to data. Investigators developed new understanding of certain aspects of the control rod fluid shear and a fuel assembly annealing process similar to a spring and viscous damper with hysteresis in series, to explain the physics Kee et al. (2005), and Kee and Björnkvist (1997). Various unrelated control rod insertion events have been observed, root cause isolated, and the knowledge base for these kinds of events continues to grow larger for example, (Lagiewsk, 1982; Jordan, 1986).

2.2 Corrective and compensatory actions

Because the root cause investigation revealed that the control rods would insert to a point where the incremental reactivity control was very small, the consequence analysis showed that, due to safety margin included in their design, the control rods would continue to meet their reactivity control design requirements with continued margin to safety. To validate the root cause and consequence analysis, the *STP* proposed a compensatory testing plan to the *NRC* which was accepted and implemented. To fully address the root cause and consequence analysis, the *STP* suggested a modification to the fuel assembly designs being purchased in the future that was implemented by the vendor.

2.3 Observations

As described by Head (1996), the reactor trip revealed the existence of the new failure mode in the fuel assemblies that was itself triggered by a previously unknown failure mode due to a pinched pilot wire induced by maintenance activity. The timing of the pilot wire failure mode revealed the fuel assembly failure mode at an unexpected time during full power operation. The annealing of in-vessel components is a well-known process and aspects were always included in, for example, fuel assembly and Boiling Water Reactor (*BWR*) channel box designs. However, the process of annealing on the approach to Euler buckling discovered by Ringhals plant was unknown to be a problem in Pressurized Water Reactor (*PWR*) fuel assemblies. The fluid shear process in the braking section of the fuel assembly had never been investigated, at least in the academic literature, until *STP* made empirical observations and Kee et al. derived the necessary equations now published in (Kee et al., 2005, Section 4.2.1).

⁴“Refueling outage” refers to a regularly scheduled period of time when approximately 1/3 of the reactor core is replaced with new fuel.

It can be said that the safety margin included in the as-designed reactivity requirements of the control rods prevented a return to critical with the control rods not fully inserted. This observation helps support the need for such prescriptive regulations as those included in Title 10 CFR Part 50 that require safety margins and defense in depth for safety-critical functions. Root cause analysis wisely acknowledges the possible existence of unanticipated failure modes with plans to learn from their discovery.

3 Protection Availability

Our observations about uncertainty quantification for protections are made with respect to a filtered probability space $(\Omega, \mathcal{F}, \{\mathcal{F}_t\}_{t \geq 0}, P)$, where Ω is the set of possible outcomes for a protective system over its lifecycle, $\mathcal{F}$ is a σ -algebra on Ω capturing all events of predictive modeling interest, and P is a probability of the measurable space $(\Omega, \mathcal{F})$. The filtration $\{\mathcal{F}_t\}_{t \geq 0}$ is the analytical construct used to capture the flow of engineering information over time. $\{\mathcal{F}_t\}_{t \geq 0}$ contains all P -measurable collections of outcomes in Ω known to modelers by time $t \geq 0$. This filtration has the following properties for all $s, t \geq 0$,

1. $\mathcal{F}_0$ contains all P -null sets (completeness),
2. $\lim_{s \downarrow 0} \mathcal{F}_{t+s} \triangleq \bigcap_{s > 0} \mathcal{F}_s = \mathcal{F}_t$ (right-continuity),
3. $\mathcal{F}_t \subseteq \mathcal{F}_{t+s}$ (monotonicity),
4. $\lim_{t \rightarrow \infty} \mathcal{F}_t \triangleq \bigvee_{t \geq 0} \mathcal{F}_t = \mathcal{F}$ (convergence).

These properties are normative and intuitively understandable. Property 2 indicates that engineering information is not necessarily pre-visible (*e.g.*, it is not possible to know that a failure will occur the instant before it actually occurs). Property 3 asserts that information gained through discovery is not lost over time. Properties 3 and 4 acknowledge that one cannot be convinced that all useful modeling information will be revealed in finite time.⁵

Consider, now, the random variable $X_t : \Omega \rightarrow \{0, 1\}$ indicating the state of system protections where

$$X_t = \begin{cases} 1, & \text{protections are available at time } t \\ 0, & \text{otherwise.} \end{cases}$$

We call $(X_t)_{t \geq 0}$ the *protection availability process*, and we normatively take its trajectories to be right-continuous.⁶ We also take $(X_t)_{t \geq 0}$ to be adapted to the filtration $\{\mathcal{F}_t\}_{t \geq 0}$. This is to say, at any time $t \geq 0$, X_t is $\mathcal{F}_t$ -measurable, and we have that, $\sigma((X_s)_{0 \leq s \leq t}) \subseteq \mathcal{F}_t$.⁷ Of course, at any time $t \geq 0$, there is far more engineering information in $\mathcal{F}_t$ than simply the partial protection availability trajectories that comprise $\sigma(X_t)$. Information related to protection design, maintenance records, historical environmental conditions, *etc* are also events that can appear in the filtration $\{\mathcal{F}_t\}_{t \geq 0}$.

⁵Even though it is reasonable to assert by definition that any consequential failure mode will be discovered in finite time, there is way to recognize that *all* consequential failure modes will have been discovered by any historically observable time.

⁶Because $(X_t)_{t \geq 0}$ is right-continuous with left-limits, it belongs to the class of càdlàg ensuring that it has a left-continuous version with respect to the probability measure P .

⁷In other words, the “natural history” of the protection availability process is contained within the filtration capturing the flow of engineering information.

For the purposes of predictive modeling, we must be careful to stipulate the status of engineering information. Note that the likelihood that system protections are failed at time t is given by

$$P(X_t = 0) = E[(1 - X_t)] \triangleq E[(1 - X_t)|\mathcal{F}] = E[E[(1 - X_t)|\mathcal{F}_t]]. \quad (1)$$

$\mathcal{F}_0$ is assumed to be complete and contains all information characterizing the protection design space with $t = 0$ taken as the time of deployment. So, all events associated with failure modes identified in an original design *FMEA* are included in $\mathcal{F}_0$.

Equation (1) reveals that quantifying the uncertainty about the operational integrity of protections at time t requires information that is not necessarily available. To see this, examine the set difference $A_t \triangleq \mathcal{F} \setminus \mathcal{F}_t$. Intuitively, A_t contains all as yet undiscovered engineering information at time t . It follows from eq. (1) that $P(X_t = 0)$ can be quantified with respect to probability measure P if and only if $P(A_t) = 0$. That is, we must be sure that there is no remaining undiscovered information that would be valuable in predicting the unavailability of protections. $P(A_t) = 0$ is guaranteed only in the limit as stipulated in Property 4, where a numerical value for $P(A_t) = 0$ is available only to a clairvoyant. Hence, engineering predictive modeling must be contented with quantifying uncertainty only up to the currently available engineering information, or

$$P(X_t = 0|\mathcal{F}_t) = E[(1 - X_t)|\mathcal{F}_t]. \quad (2)$$

It is important to appreciate that system protections are useful only if they are available exactly when needed. Suppose that T is the time of arrival of some initiating event that can possibly lead to catastrophe. Here, $T : (\Omega, \mathcal{F}) \rightarrow (\mathbb{R}^+, \mathcal{B}(\mathbb{R}^+))$. Catastrophe can occur only if $X_T = 0$.

Remark 1 *Simply characterizing the state of protections with respect to the flow of engineering information over time is generally inadequate to inform the likelihood of successful protections; protections must hold when at the random times when initiating events occur. In practical circumstances, $P(X_T = 0) \neq P(X_t = 0)$, when T is the arrival time of an initiating event.*⁸

3.1 Initiating event arrivals that are stopping times.

When the events $\{T \leq t\}$ are included in $\mathcal{F}_t$ for all $t \geq 0$, T is said to be an $\{\mathcal{F}_t\}_{t \geq 0}$ stopping time. It is well known that any $\{\mathcal{F}_t\}_{t \geq 0}$ stopping time T must also be $\mathcal{F}_T$ -measurable, where by definition

$$\mathcal{F}_T \triangleq \{A \in \mathcal{F} : A \cap \{T \leq t\} \in \mathcal{F}_t\}. \quad (3)$$

It follows directly that $\mathcal{F}_T \subseteq \mathcal{F}_t$. Thus, $P(X_T = 0|\mathcal{F}_t) = E[(1 - X_T)|\mathcal{F}_t]$ is well defined.

Clearly, stopping times play an important role in *UQ* supporting risk analysis and regulatory oversight. It should be appreciated that generally $\mathcal{F}_T$ contains engineering information that is *not* necessarily included in the protection hardware design/maintenance space (*e.g.*, information derived from prior understandings of weather patterns, political unrest, economic activity, *etc.*). Under the circumstance that the arrival time of an initiating event is an $\mathcal{F}_t$ -stopping time for all $t \geq 0$, we know that $P(X_T = 0|\mathcal{F}_t)$ is well defined. However, being well defined *does not* imply that numerical values for $P(X_T = 0|\mathcal{F}_t)$ are attainable (Wortman et al., 2021).

⁸This fact bears directly upon the validity of *UQ* methods when applied to risk informed, performance-based regulatory oversight.

3.2 Initiating event arrivals that are *not* stopping times.

There are a host of practical reasons why the arrival time of an initiating event might not be measurable with respect to the filtration $\{\mathcal{F}_t\}_{t \geq 0}$. In particular, if it happens that the initiating event reveals a previously undiscovered system protection failure mode, a catastrophe can ensue. Since the failure mode is heretofore unknown, its arrival time T is *not* $\mathcal{F}_t$ -measurable for any $t \geq 0$. This is to say $E[(1 - X_T)|\mathcal{F}_t]$ is not well defined which implies that neither $P(X_T = 0|\mathcal{F}_T)$ nor $P(X_T = 0|\mathcal{F}_t)$ can be quantified. Thus, undiscovered protection failure modes are problematic for any regulatory oversight strategy that relies completely on UQ .

It is reasonable, at this juncture, to consider the extent to which protection failure modes not identified in $\mathcal{F}_0$ are problematic. We begin with some observations that are normatively justified in the engineering pedagogy:

- The discovery time T of a specific heretofore undiscovered failure mode is random and measurable with respect to probability space $(\Omega, \mathcal{F}, P)$.
- T is not an $\{\mathcal{F}_t\}_{t \geq 0}$ stopping time.
- It is impossible to quantify the probability of events that depend on T .
- There exists some finite time $\tau < \infty$ such that $P(T \leq \tau) = 1$. That is, any consequential failure mode will almost surely be found over that lifecycle of protections, else cannot be consequential.
- It is possible that a failure mode is first discovered through catastrophe postmortem analysis (*i.e.*, the failure mode caused catastrophe upon its first appearance).
- For any time $t \geq 0$ a non-clairvoyant cannot rule out the possibility of remaining undiscovered protection failure modes.

If undiscovered consequential failure modes are in play, the filtration $\{\mathcal{F}_t\}_{t \geq 0}$ cannot yet have converged. From the perspective of predictive modeling, it is understood that A_t must contain the undiscovered failure mode information, and thus $P(A_t|\mathcal{F}) > 0$. But, clearly, $A_t \notin \mathcal{F}_t$ and thus $P(A_t|\mathcal{F}_t)$ is not well defined ... that is, a non-clairvoyant is unable to quantify the inaccessible event probabilities related to undiscovered failure modes.

However, if a modeler is sufficiently confident in a non-clairvoyant belief that that all consequential failure modes have been discovered, then all information impacting protection design will be found in the tail σ -algebra $\mathcal{T}$, where

$$\mathcal{T} \triangleq \bigcap_{t \geq 0} \sigma((X_s)_{s > t}).$$

By definition, $\mathcal{T} \subset \mathcal{F}$ and characterizes design information the *remote future* of $(X_t)_{t \geq 0}$. In the remote future of protections, there can be no undiscovered consequential failure modes.

Typically, UQ methodologies (*e.g.*, *PRA*, Quantitative Risk Analysis (*QRA*), Probabilistic Safety Analysis (*PSA*)) implicitly rely on the assumption that events associated with protections are $\mathcal{T}$ -measurable, because this assumption guarantees the existence of

$$X = \lim_{t \rightarrow \infty} X_t \text{ and } \bar{X} = \lim_{t \rightarrow \infty} \frac{1}{t} \int_0^t X_s ds.$$

Expected values of these X and $\bar{X}$ and their corresponding statistical estimators play essential roles in UQ . Informally, when time $t = 0$ is taken to be in the remote future of the protection availability

process $(X_t)_{t \geq 0}$, then $\mathcal{T} \subset \mathcal{F}_t$ and X and $\bar{X}$ each become $\mathcal{F}_t$ -measurable, and numerical estimates of their respective expected values can (in principle) be computed using historical data collected up through time t .⁹

Importantly, modelers should appreciate that enabling computation of UQ statistics nearly always requires assuming that $(X_t)_{t \geq 0}$ be $\mathcal{T}$ -measurable. And, assuming that $(X_t)_{t \geq 0}$ is $\mathcal{T}$ -measurable implicitly ignores the possibility of undiscovered consequential failure modes that might lead to catastrophe. This leaves open the engineering question, “*How much time must elapse before one might reasonably trust that all consequential protection failure modes have been discovered?*”. Mathematics dictates that only a clairvoyant can answer this question with complete confidence. Of course, no one is clairvoyant and undiscovered protection failure modes present especially difficult problems for engineers. Hence, because protection failure mode discovery times defy UQ , care must be exercised when applying UQ methodologies to studies of protection efficacy.

4 Discussion

Unknown-unknowns, including undiscovered protection failure modes, are nothing new to engineering design. Engineering practices such as stress testing, accelerated life testing, burn-in, and advanced physics-based simulations have been developed and refined with the specific objective of discovering failure mechanism. Additionally, engineers routinely adopt non-probabilistic strategies including defense in depth and layers of protection for mitigating the consequences of yet undiscovered failure modes. Quantitative risk methodologies have been increasing in popularity in many areas of cost-benefit decision-making and regulatory oversight. Further, elaborate data bases are constructed and maintained so that discoveries can be shared across various designs. In particular, the *NRC* has developed and implemented data recording and sharing protocols specifically intended to capture and alert the entire Nonetheless, quantitative risk analysis unconditionally produces optimistic estimates of predicted protection performance that cannot be overcome. An assessment that estimates forward cost as the product of initiating event frequency, protection failure probability, and consequence, will, of course, underestimate the cost. Engineers tasked with developing processes and designs that may pose a hazard as a result of a failure in use routinely compensate for unmeasurable uncertainty by adding defense in depth and safety margin in areas where either experience informs them as to the level of uncertainty or they lack knowledge.

Although the example we show is from nuclear power, our observations apply equally to other hazardous industrial sectors such as chemical processing, oil and gas exploration and production, public and freight transportation, and others. They also apply to repairable as well as unrepairable systems although understanding the behaviors is nuanced.

References

Apostolakis, G. and S. Kaplan (1981). Pitfalls in risk calculations. *Reliability Engineering* 2(2), 135–145.

⁹Assuming that $(X_t)_{t \geq 0}$ is $\mathcal{T}$ -measurable, in no way implies that an $\mathcal{F}_t$ -measurable stopping time T will be either $\mathcal{T}$ -measurable or independent of $(X_t)_{t \geq 0}$. For example an initiating event arriving at time T could possibly damage system protections triggering a protection failure. Thus, even when assuming that the protection availability process is $\mathcal{T}$ -measurable, estimating $E[(1 - X_T)|\mathcal{F}_t]$ can be perilous.

- Crutchfield, D. M. (1996, 03). Bulletin 96-01: Control Rod Insertion Problems. OMB No. 3150-0012, NRCB 95-02.
- Doorn, N. and S. Hansson (2011). Should Probabilistic Design Replace Safety Factors? *Philosophy & Technology* 24(2), 151 – 168.
- Hansson, S. O. (2009). From the casino to the jungle: Dealing with uncertainty in technological risk management. *Synthese* 168(3), 423 – 432.
- Head, S. M. (1996, 10). Turbine Trip and Reactor Trip Due to Main Transformer Lockout. South Texas, Unit 1, LICENSEE EVENT REPORT (LER), ACCESSION #: 9610150257, to Nuclear Regulatory Commission.
- Jordan, E. L. (1986, 08). Information Notice No. 86-68: Stuck Control Rod. SSINS No.: 6835, IN 86-68.
- Kaplan, S. and B. J. Garrick (1981). On the quantitative definition of risk. *Risk analysis* 1(1), 11–27.
- Kee, E. (1996, 10). Experience with Incomplete Control Rod Insertion in Fuel with Burnup Exceeding Approximately 40 GWD/MTU. In B. N. Laboratory (Ed.), *Twenty-Fourth Water Reactor Safety Information Meeting*, Volume 1 of NUREG/CP-0157, Accession No. ML042230069, Bethesda, MD, pp. 185–200. Houston Lighting & Power Co., South Texas Project: Nuclear Regulatory Commission.
- Kee, E. J. and L. Björnkvist (1997, 03). Application of a semi-empirical rod drop model for studying rod insertion anomalies at south texas project and ringhals unit 4. In *Proceedings of the 1997 ANS International High Burnup Fuel Performance Meeting*.
- Kee, R. J., M. E. Coltrin, and P. Glarborg (2005). *Chemically reacting flow: theory and practice*. John Wiley & Sons.
- Lagiewsk, J. S. (1982, 01). Control Rod Assembly Inoperable. Calvert Cliffs, LICENSEE EVENT REPORT (LER), ACCESSION #: 3391984002.
- NRC (2013, 11). Glossary of Risk-Related Terms in Support of Risk- Informed Decisionmaking. NUREG 2122, ML13311A353, Nuclear Regulatory Commission, Washington, D.C.
- Prigogine, I. and I. Stengers (1997). *The end of certainty*. Simon and Schuster.
- Rasmussen, N. C. (1975). Reactor Safety Study: An Assessment of Risk in U.S. Commercial Nuclear Power Plants, WASH-1400. Technical report, Nuclear Regulatory Commission.
- Wortman, M. A., E. J. L. Kee, and P. Kannan (2021). Is core damage frequency an informative risk metric? arXiv preprint no. 2105.09798.