

Quantifying and Reducing HRA Uncertainty in Advanced Reactors: Bayesian Inference Approach

Krzysztof Radziszewski^{a*}, Tatsuya Sakurahara^a

^a Risk Analysis and Reliability Engineering (RARE) Laboratory, Department of Mechanical Engineering and Materials Science, Swanson School of Engineering, University of Pittsburgh, Pittsburgh, PA, USA

* Corresponding author, Email: krr97@pitt.edu

Abstract:

Advanced nuclear reactors, such as small modular reactors and microreactors, introduce new challenges for Human Reliability Analysis (HRA) due to differences in design characteristics, operational context, information availability, and design maturity compared to conventional light-water reactors (LWRs). One distinct characteristic of these advanced reactors is increased digitization and automation, which can affect operators' performance. A literature review of HRA applications for microreactors shows that existing studies largely rely on qualitative categorization, lump-sum frequencies from generic databases, or simplified nominal human error probability (HEP) values.

The U.S. Nuclear Regulatory Commission (NRC) has developed the Integrated Human Event Analysis System for Event and Condition Assessment (IDHEAS-ECA) as the state-of-the-art HRA method to support risk-informed regulation. The associated database, documented in IDHEAS-DATA (RIL 2025-01), is built primarily from simulator data and other sources specific to conventional LWRs. In general, increased digitization and automation in advanced reactors can shift cognitive failure modes and performance-influencing factors toward contexts where existing human performance data are relatively sparse.

This paper addresses two research questions: (i) how much additional uncertainty can limited data introduce into HRA for an advanced reactor context, and (ii) if the uncertainty is unacceptably large, how much additional data is needed to reduce it? A Bayesian inference approach is proposed to construct probability distributions representing uncertainties in model parameters of IDHEAS-ECA, including base HEPs and performance-influencing factor (PIF) weights. Probabilistic programming is implemented using the PyMC Python package. The proposed approach is demonstrated using synthetic data designed to reflect the structure and quantity of the IDHEAS-DATA dataset. While application to the actual IDHEAS-DATA dataset is ongoing, the results from the synthetic data presented here provide an order-of-magnitude insight into the HEP uncertainty associated with IDHEAS-ECA.

1. INTRODUCTION

Advanced nuclear reactors, such as microreactors and Small Modular Reactors (SMRs), introduce new challenges for Human Reliability Analysis (HRA) due to operational conditions and contexts that distinguish them from existing plants, induced by new operational paradigms, automation and risk profile shift, operating experience data and information availability, and design maturity [1-3].

The authors have conducted a literature review to examine HRA for microreactors published from 2020 to 2025. The results have been very limited: Eight references that quantitatively analyzed human errors have been identified. Two of them categorized frequencies of human failure events by assigning specific qualitative descriptions to the likelihood (e.g., Anticipated, Unlikely, Extremely Unlikely, and Beyond Extremely Unlikely) rather than assessing numerical probabilities [4, 5]. Another group of studies [6-

8] relied on historical databases to derive event frequencies during transportation. Accident databases from transportation (FARS, CAROL, and U.S. DOT) served as the basis for estimating frequencies of transportation accident scenarios, while these studies focused on estimating aggregated probabilities at the scenario level without detailed analysis of human actions. Human error remained implicit in the composite event frequencies used as event-tree branching probabilities. Only one study applied a formal HRA method but implemented a simplified form of SPAR-H by using only the base HEPs [9]. Detailed task analysis, evaluation of performance-shaping factors (PSFs), or dependency evaluation were not implemented.

To enhance the level of detail in advanced reactor HRA, one natural pathway is to apply existing HRA methods; however, we need to carefully evaluate the applicability of the selected HRA method to advanced reactor contexts and conditions, such as digitization and higher levels of automation. Given the distinct characteristics of advanced reactors and their limited operating experience, existing HRA methods could lack supporting human performance data on human failure modes and PSFs unique to advanced reactors. Limited data can introduce large uncertainty in HEP quantification. This raises two research questions: (i) how much additional uncertainty can the limited data introduce into HRA for an advanced reactor context, and (ii) if the uncertainty is unacceptably large, how much additional data is needed to reduce it? The objective of this work is to provide insights into these research questions using Bayesian analysis.

As a representative example of the state-of-the-art HRA methods, this study uses the Integrated Human Event Analysis System for Event and Condition Assessment (IDHEAS-ECA) approach [10], developed by the U.S. Nuclear Regulatory Commission (NRC). Quantification of HRA parameters in IDHEAS-ECA, such as base HEPs and Performance Influencing Factor (PIF) weights, has been supported by a collection of human performance databases available in IDHEAS-DATA [11]. The IDHEAS-DATA package presents 27 tables with multiple types of quantitative and qualitative data for various states of Cognitive Failure Modes (CFMs), base HEPs, and PIF weights. These probability and weight values, however, can involve parameter uncertainty, as they were estimated from a limited number of data points for each combination of their states. Step 8 of IDHEAS-ECA [10] requires the user to analyze uncertainties and perform sensitivity studies, but the guidance largely focuses on the time-based probability P_t , leaving the cognitive probability P_c without thorough uncertainty quantification. This paper focuses on uncertainty quantification for the cognitive portion of IDHEAS-ECA.

The remainder of this paper is organized as follows. Section 2 describes an illustrative model and problem setup currently used in the study. Section 3 presents the mathematical formulation and computational implementation of Bayesian inference. Section 4 reports the results from the preliminary implementation. Section 5 concludes the paper and suggests directions for future work.

2. ILLUSTRATIVE MODEL AND PROBLEM SETUP

As an illustrative model for demonstrating the proposed approach, this paper utilizes a simplified model, given in Equation (1), that has a structure similar to the IDHEAS-ECA model [10]. The purpose of introducing the simplified model, rather than using the actual IDHEAS-ECA model, is to facilitate iterative trial-and-error in computational execution and code debugging and enhance the interpretability of the results while maintaining relevance to IDHEAS-ECA. Ongoing research by the authors applies the proposed approach to the IDHEAS-ECA model.

The illustrative model used in this paper is defined by Eq. (1).

$$Y = \left[1 - \prod_{i=1}^2 (1 - \theta_i) \right] \left[1 + \sum_{j=1}^2 (w_j - 1) \right] \quad (1)$$

where

- Y: Model output corresponding to an HEP in IDHEAS-ECA,

- θ_1 and θ_2 : Influencing factors with 3 discrete states, corresponding to the Base Human Error Probability (BHEP) states in IDHEAS-ECA,
- w_1 and w_2 : Weight factors with five discrete states, corresponding to PIF impact weights for the given attributes in IDHEAS-ECA.

The first term of Eq. (1) follows the series system logic, where the combined basic HEP is determined by the union of θ_1 and θ_2 . The second term introduces a multiplicative weight interaction that scales the HEP based on the joint contribution of w_1 and w_2 using linear multipliers, which are consistent with the PIF treatment in IDHEAS-ECA.

For the base HEP factors (θ_1 and θ_2), ordering constraints are imposed to represent the monotonically increasing severity states across their states:

$$\theta_{i,1} < \theta_{i,2} < \theta_{i,3}, \quad i = 1, 2 \quad (2)$$

The objective of the analysis is to estimate the model parameters in Eq. (1), including $\theta_{i,s}$, $i \in \{1, 2\}$, $s \in \{1, 2, 3\}$, and $w_{j,t}$, $j \in \{1, 2\}$, $t \in \{1, 2, 3, 4, 5\}$, based on a given dataset.

This study generates a synthetic dataset with a structure and quantity similar to those of the IDHEAS-DATA. The data is stored in a CSV format, and for each row, 5 columns are included: observed human performance as a binomial pair (x, n), where x is the number of errors, and n is the number of trials; the state of each base HEP factor (θ_1 and θ_2); and the state of each PIF weight (w_1 and w_2). The synthetic dataset used in this paper contains 21 rows.

Table 1 presents representative data entries. The data structure is intended to resemble the IDHEAS-DATA structure. Entries of types #1 and #2 in Table 1 capture the state of θ_1 and θ_2 , respectively; the remaining columns are set to n/a to isolate the contribution of each factor individually. The entries of types #3 and #4 capture the joint influence of θ_1 and θ_2 and one active PIF weight at a time, while the other PIF column is set to “n/a.”

Table 2 summarizes the number of data entries per state for each model parameter in the synthetic dataset. As a general tendency, the harsher the state is, the fewer data entries are available, consistent with the data-entry distribution observed in the IDHEAS-DATA.

Table 1: Sample Data Entries from Synthetic Dataset

Data Entry	Observed Performance	θ_1 State	θ_2 State	w_1 State	w_2 State
1	(2, 100)	1	n/a	n/a	n/a
2	(10, 56)	n/a	3	n/a	n/a
3	(6, 24)	2	2	3	n/a
4	(10, 94)	1	1	n/a	5
...	...	...	...	...	...

Table 2: Number of Data Entries per State for θ and w in the Synthetic Dataset

Variable	State 1	State 2	State 3	State 4	State 5
θ_1	8	7	2	—	—
θ_2	7	8	1	—	—
w_1	8	2	2	0	0
w_2	10	0	0	1	1

3. METHODOLOGY

3.1. Mathematical Formulation

To quantify the uncertainty in the estimates of the model parameters, i.e., $\theta_{i,s}$ and $w_{j,t}$ in Eq. (1), this research utilizes Bayesian inference, which is selected for three reasons: (i) fusion of multiple data sources facilitates parameter estimation under limited empirical data; (ii) the posterior distribution can be updated sequentially as new evidence becomes available, enabling continuous refinement of HRA inputs; and (iii) it can compute the full posterior distribution, rather than point estimates or summary statistics such as confidence intervals.

Let $\mathbf{\Omega} = \{\theta_{\{i,s\}}, w_{\{j,t\}}\}$ denote the collection of model parameters, where $i \in \{1, 2\}$, $s \in \{1, 2, 3\}$, $j \in \{1, 2\}$, $t \in \{1, 2, 3, 4, 5\}$. Using Bayes' Theorem, the posterior distribution of $\mathbf{\Omega}$, given evidence E , is expressed as:

$$\pi(\mathbf{\Omega}|E) = \frac{L(E|\mathbf{\Omega})\pi_0(\mathbf{\Omega})}{\int \pi_0(\mathbf{\Omega})L(E|\mathbf{\Omega})d\mathbf{\Omega}}, \quad (3)$$

where $\pi(\mathbf{\Omega}|E)$ is the posterior distribution, $L(E|\mathbf{\Omega})$ is the likelihood function, $\pi_0(\mathbf{\Omega})$ is the prior distribution, and E is the evidence, namely, the synthetic dataset (Table 1).

For the problem described in Section 2, prior distributions for the θ and w parameters are specified as uniform distributions:

$$\pi_0(\theta_{i,1}) = U(0,1), \quad i = 1, 2 \quad (4)$$

$$\pi_0(\theta_{i,j}) = U(\theta_{i,j-1}, 1), \quad i = 1, 2; j = 2, 3 \quad (5)$$

Since $\theta_{i,j}$ represent base HEPs, a uniform distribution over $[0, 1]$ is a natural choice of non-informative priors. Equation 5 uses $\theta_{i,j-1}$ as the lower bound to enforce the ordering constraint in Equation (2).

Based on the IDHEAS-DATA [11], weight multipliers are assumed to lie within the range $[1, 5]$, with the exception of $w_{i,1}$, which is fixed at 1 as it represents the nominal state. Uniform priors are specified as given in Equations (6) and (7).

$$w_{i,1} = 1, \quad (6)$$

$$\pi_0(w_{i,l}) = U(1, 5), \quad (7)$$

where $i = 1, 2$ and $l = 2, 3, 4, 5$.

The likelihood function $L(E|\mathbf{\Omega})$ is defined as a binomial distribution:

$$L(E|\mathbf{\Omega}) = \text{Bin}(n, p), \quad (8)$$

where n is the number of trials included in the dataset, and p is the human error probability per trial. p is obtained from the model defined in Equation (1). The binomial distribution is selected due to the nature of the data structure, where each entry records x failures out of n trials.

3.2. Computational Implementation

Bayesian inference is implemented using PyMC (version 5.28.5), a probabilistic programming library in Python [12]. PyMC performs posterior estimation via Markov Chain Monte Carlo (MCMC) sampling

algorithms, which scale effectively to high-dimensional models. PyMC leverages PyTensor for automatic differentiation and uses the No-U-Turn Sampler (NUTS) for efficient posterior exploration.

In PyMC, the sampling process consists of two phases. In the first phase, a user-specified number of initial iterations are discarded as burn-in. These samples are excluded from the subsequent posterior inferences. In the second phase, the sampler draws the retained posterior samples used for inference. The number of burn-in iterations and retained draws are set to 2,000 and 2,000 per chain, respectively, based on the following convergence criteria:

- Trace plots are visually inspected to assess chain behavior, which informs the selection of the number of burn-in iterations. Each trace plot displays the posterior samples of each unknown of interest as a function of the number of MCMC iterations for each chain. Well converged chains are characterized by stable, dense scatters oscillating around a constant median. If any chain presents a visually identifiable drift, slow mixing, or artifacts across the MCMC iterations, convergence is not achieved, and the sample size needs to be increased.
- Values of $\hat{R}$ close to 1.0 are targeted to select the adequate burn-in size. $\hat{R}$ is the ratio of between-chain variance to within-chain variance across multiple Markov chains. Values close to 1.0 indicate that chains have converged, while values significantly above 1.0 suggest insufficient convergence.
- Values of the effective sampling size (ESS) are targeted at 1,000 per chain to select the number of retained draws. Because consecutive MCMC samples are autocorrelated, they provide less information per sample than independent samples. The ESS accounts for this autocorrelation by estimating the equivalent number of independent samples. Two types of ESS are evaluated: ESS_{bulk} , which reflects the convergence of estimates near the center of the posterior distribution; and ESS_{tail} , which represents the convergence in tails of the posterior distributions. Values below a few hundred per chain indicate unreliable posterior estimates.

4. RESULTS

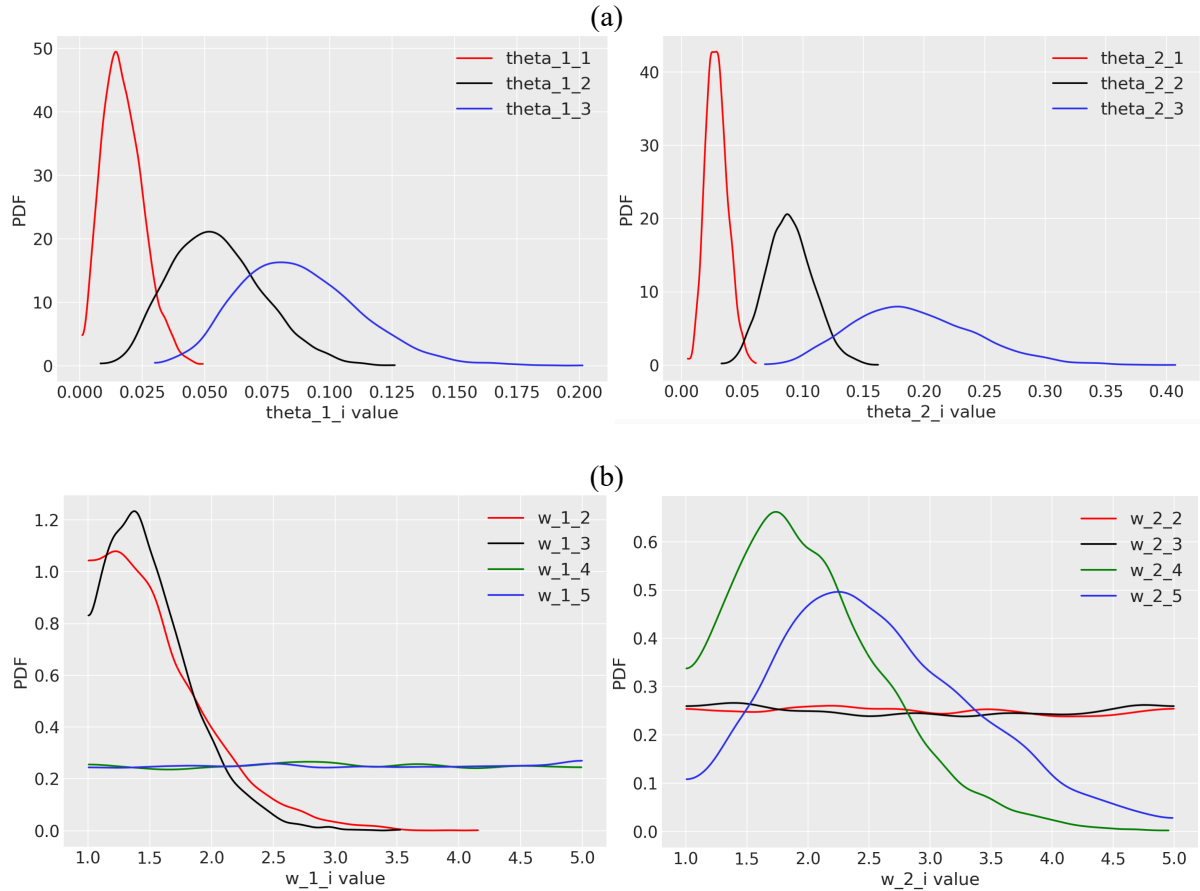
Section 4.1 presents the results of the Bayesian inference applied to the synthetic dataset described in Section 2. Section 4.2 demonstrates the propagation of the parameter uncertainty through the model to the HEP outputs and an implementation of global sensitivity analysis to identify significant sources of uncertainty that dominate the HEP uncertainty.

4.1. Characterization of Uncertainty for HRA Model Parameters

Figure 1 presents the marginal posterior distributions of each model parameter. The parameters $\theta_{1,1}$, $\theta_{1,2}$, $\theta_{2,1}$, and $\theta_{2,2}$ show relatively narrow posteriors, indicating that the dataset was highly informative for these parameters. In contrast, posteriors of $\theta_{1,3}$ and $\theta_{2,3}$ are relatively flat, with broad probability density, indicating that the observed data provided limited information for these parameters. These observations are supported by the number of data entries in Table 2: seven to eight data entries are available for States 1 and 2 of these parameters, while only one or two data entries are provided for State 3.

The posterior distributions of $w_{1,2}$ and $w_{1,3}$ are noticeably narrower than their priors, namely, a uniform distribution in $[1, 5]$, suggesting that the dataset provides relatively strong evidence that indicates smaller values in the specified range. The weight factor parameters $w_{1,4}$, $w_{1,5}$, $w_{2,2}$, and $w_{2,3}$ exhibit notably wide posteriors that mirror their priors, suggesting that the dataset is uninformative for these parameters. This observation is consistent with the counts of data entries shown in Table 2, where zero data entries are available for $w_{1,4}$, $w_{1,5}$, $w_{2,2}$, and $w_{2,3}$.

Figure 1: Posterior Distributions for θ_{ij} and w_{ij} .



4.2. Uncertainty Propagation and Sensitivity Analysis

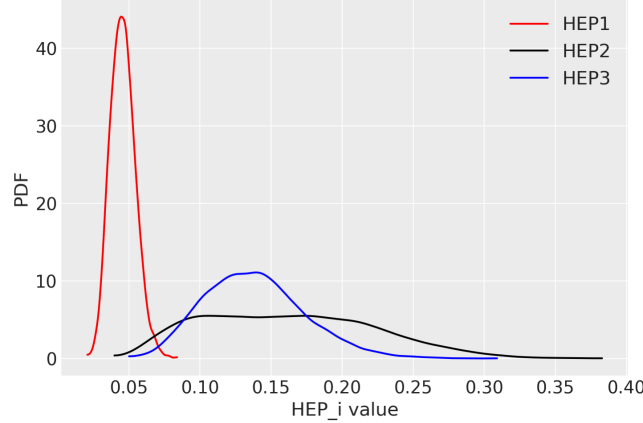
The posterior distributions for the model parameters obtained in Section 4.1 are propagated through the model to the output Y in Equation (1), representing an HEP estimate. Here, Y is calculated for three illustrative cases, as defined in Table 3, with uncertainty propagation using Monte Carlo simulation.

Table 3: Three illustrative cases used for demonstrating uncertainty propagation.

Case	θ_1 State	θ_2 State	w_1 State	w_2 State
1	1	1	1	1
2	1	1	2	2
3	1	1	3	5

Figure 2 shows the kernel density estimates of Y for three illustrative cases (HEP1, HEP2, and HEP3). These distributions reflect the total uncertainty in the HEP outputs propagated from the posterior distributions of the model parameters. HEP2 presents the largest uncertainty due to the inclusion of $w_{2,2}$, for which no data point is available in the given synthetic dataset, and, hence, the posterior distribution is identical to the non-informative prior distribution, as shown in Figure 1(b).

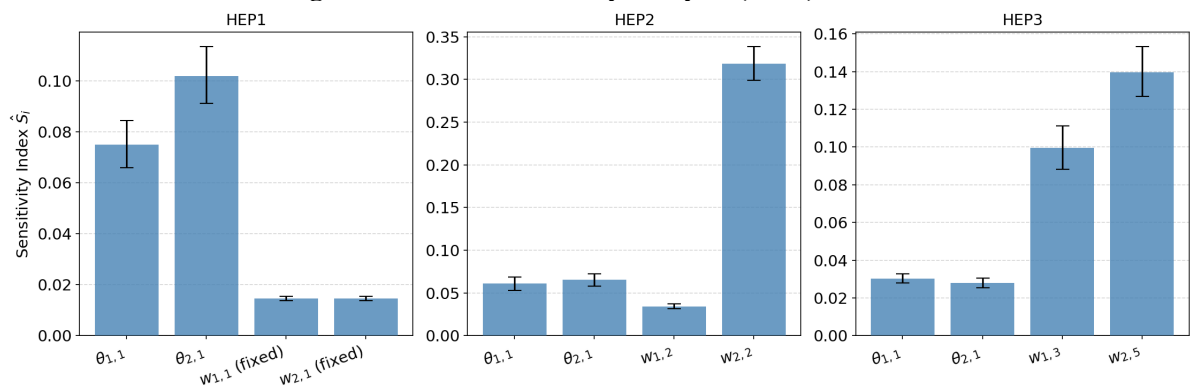
Figure 2: Uncertainty distributions for the HEP outputs.



One of the benefits of constructing explicit probability distributions for the model parameters is that quantitative sensitivity analysis can be conducted to identify the most significant sources of uncertainty. To quantify the relative contribution of each model parameter to the uncertainty in the HEP outputs, a global sensitivity analysis (GSA) is performed using a cumulative distribution function (CDF)-based sensitivity index proposed by Liu and Homma [13]. This method is advantageous when the sensitivity evaluation should be based on the entire output distribution rather than specific summary statistics or the center of the posterior distributions [14]. For each model parameter X_i , the unconditional CDF of the output Y is compared with conditional CDFs obtained by fixing X_i at a certain sampled value. The area between the unconditional and conditional CDFs is used to quantify the sensitivity of the output distribution to each input. The CDF-based sensitivity indices are computed using the single-loop approach developed by Plischke et al. [15] to estimate GSA indices from a given input-output dataset, i.e., the MCMC draws from Bayesian inference after discarding the burn-in samples.

Figure 3 shows the GSA results for the three cases (HEP1, HEP2, and HEP3). In the HEP1 case, the dominant uncertainty sources are θ_1 and θ_2 , which is consistent with the HEP1 setup, where both w_1 and w_2 are fixed at their base state with a multiplier of unity. In the HEP2 case, the dominant uncertainty stems from w_2 . This can be explained by the fact that, as shown in Table 2, the synthetic data contains no data entries with the second state of w_2 . In the HEP3 case, w_1 and w_2 are relatively large uncertainty sources, which is consistent with the observation that the synthetic dataset includes only two data entries for the third state of w_1 and one data entry for the fifth state of w_2 . These insights from the GSA results can be used to prioritize further data collection efforts with the aim of reducing HEP uncertainty.

Figure 3: Global sensitivity analysis (GSA) results.



The Bayesian analysis in this paper uses the simplified model in Equation (1) and the synthetic dataset. However, due to the structural similarity of the model to IDHEAS-ECA and of the synthetic dataset to IDHEAS-DATA, the preliminary results in this paper can provide order-of-magnitude insights into the potential uncertainty in current IDHEAS-ECA calculations.

- The 95th/5th percentile ratio for HEP uncertainty ranges from 1.9 to 3.6 among the three cases in Figure 2. This range is comparable to, or at least not significantly larger than, the typical error factors of other PRA inputs, such as initiating event frequencies and common cause failure model parameters.
- PIF states supported by fewer than three quantitative data entries are the dominant sources of HEP uncertainty. Such data sparsity can produce substantially broader HEP uncertainty, which can lead to 95th/5th percentile ratios of 5 to 10.

It should be noted that these insights are preliminary, derived from the illustrative analysis presented in this paper. These insights are being verified in ongoing analysis by the authors using the full IDHEAS-ECA model and the IDHEAS-DATA database.

5. CONCLUSION

This paper proposes a Bayesian inference method to quantify parameter uncertainty in IDHEAS-ECA, motivated by the challenge of applying existing HRA methods in advanced-reactor contexts where relevant human performance data can be sparse. The proposed method constructs posterior probability distributions for base HEP parameters and PIF weight multipliers directly from observed human performance data through probabilistic programming using the PyMC package.

In this paper, the proposed method is demonstrated on a synthetic dataset designed to reflect the structure and data entry characteristics of IDHEAS-DATA. For the three illustrative HEP cases studied, the 95th/5th percentile ratio of the resulting HEP uncertainty ranges from 1.9 to 3.6, while it can reach 5 to 10 when key PIF states lack empirical data support (e.g., < 3 data entries). The preliminary findings suggest that uncertainty at the level of typical PRA input error factors can be achieved under reasonable human performance data coverage, but sparse PIF state data can add substantial uncertainty. Global sensitivity analysis can be performed using the posterior samples to identify which specific parameters dominate HEP uncertainty, providing a systematic basis for prioritizing future data collection efforts.

Ongoing efforts focus on applying the proposed method to the full IDHEAS-ECA model and the actual IDHEAS-DATA database. Global sensitivity analysis will be extended across the full parameter space to identify specific base HEP and PIF combinations for which additional data collection would provide the greatest reduction in HEP uncertainty in the advanced reactor contexts.

Acknowledgements

This research has been supported by the Department of Mechanical Engineering and Materials Science at the University of Pittsburgh. The authors also thank members of the Risk Analysis and Reliability Engineering (RARE) Laboratory for their feedback and collegial support throughout this work.

The authors used Grammarly (<https://app.grammarly.com/>) to improve the grammar and readability of this paper. The manuscript was carefully reviewed and edited by the authors, who take full responsibility for the accuracy of the manuscript content.

6. REFERENCES

- [1] C. Blackett, "Adapting Human Reliability Analysis For Small Modular Reactors," in *34th European Safety and Reliability Conference (ESREL)*, Krakow, Poland, 2024.
- [2] A. Bye, J. A. Julius, and R. Boring, "Challenges for Human Reliability Analysis in New Nuclear Power Plant Designs," in *16th International Conference on Probabilistic Safety Assessment and Management (PSAM 2022)*, Honolulu, HI, 2022.

- [3] C. Blackett, E. Olson, and B. L. Tottie, "A Taxonomy of Human Tasks for Human Reliability Analysis of Small Modular Reactors," in *the 35th European Safety and Reliability and the 33th Society for Risk Analysis Europe Conference*, Stavanger, Norway, 2025.
- [4] D. Gerstner and J. Andrus, "Application of a Qualitative Risk-Informed, Performance-Based Approach for the MARVEL Microreactor at the Idaho National Laboratory," Idaho National Laboratory, 2024.
- [5] M. W. Patterson, "MARVEL Hazard Evaluation ECAR-6440," Department of Energy, 2023.
- [6] T. M. O'Connell, "Risk-Informed Airfield Selection and Transportation of Department of Defense Nuclear Microreactors," 2024.
- [7] Y. A. Alzahrani and M. Prins, "Integrated probabilistic risk assessment framework for transporting microreactors with a case study for road, rail, and maritime transportation modes," *npj Clean Energy*, 2025.
- [8] M. Prins, "A Probabilistic Risk Assessment for the Transportation of Department of Defense Micro Reactors with the Effects of an Accident on the Surrounding Population," 2023: ASME.
- [9] G. A. Coles and T. A. Ikenberry, "Development and Demonstration of a Risk Assessment Approach for Approval of a Transportation Package of a Transportable Nuclear Power Plant for Domestic Highway Shipment," U.S. Department of Energy, 2024.
- [10] J. Xing, Y. J. Chang, and J. DeJesus Segarra, "Integrated Human Event Analysis System for Event and Condition Assessment (IDHEAS-ECA)," in "NUREG-2256," U.S. Nuclear Regulatory Commission, Office of Nuclear Regulatory Research, Washington, DC, 2022/10 2022. [Online]. Available: <https://www.nrc.gov/reading-rm/doc-collections/nuregs/staff/sr2256/index>
- [11] J. Xing and J. DeJesus Segarra, "The Integrated Human Event Analysis System - Human Reliability Database (IDHEAS-DATA)," in "Research Information Letter," U.S. Nuclear Regulatory Commission, Office of Nuclear Regulatory Research, Washington, DC, 2025/01 2025, vol. 2025-01. Accessed: 2026/04/01. [Online]. Available: <https://www.nrc.gov/docs/ML2435/ML24358A050.pdf>
- [12] O. Abril-Pla *et al.*, "PyMC: A Modern and Comprehensive Probabilistic Programming Framework in Python," *PeerJ Computer Science*, vol. 9, p. e1516, 2023, doi: 10.7717/peerj-cs.1516.
- [13] Q. Liu and T. Homma, "A new importance measure for sensitivity analysis," *Journal of nuclear science and technology*, vol. 47, no. 1, pp. 53–61, 2010.
- [14] T. Sakurahara, S. Reihani, E. Kee, and Z. Mohaghegh, "Global Importance Measure Methodology for Integrated Probabilistic Risk Assessment," *Proceedings of the Institution of Mechanical Engineers, Part O: Journal of Risk and Reliability*, vol. 234, no. 2, pp. 377–396, 2020.
- [15] E. Plischke, E. Borgonovo, and C. L. Smith, "Global sensitivity measures from given data," *European Journal of Operational Research*, vol. 226, no. 3, pp. 536–550, 2013/05/01/ 2013, doi: <https://doi.org/10.1016/j.ejor.2012.11.047>.