

SANDIA REPORT

SAND2014-5039

Unlimited Release

Printed July 2014

Process for Estimating Likelihood and Confidence in Post Detonation Nuclear Forensics

John L. Darby and Charles M. Craft

Prepared by
Sandia National Laboratories
Albuquerque, New Mexico 87185 and Livermore, California 94550

Sandia National Laboratories is a multi-program laboratory managed and operated by Sandia Corporation, a wholly owned subsidiary of Lockheed Martin Corporation, for the U.S. Department of Energy's National Nuclear Security Administration under contract DE-AC04-94AL85000.

Approved for public release; further dissemination unlimited.



Sandia National Laboratories

Issued by Sandia National Laboratories, operated for the United States Department of Energy by Sandia Corporation.

NOTICE: This report was prepared as an account of work sponsored by an agency of the United States Government. Neither the United States Government, nor any agency thereof, nor any of their employees, nor any of their contractors, subcontractors, or their employees, make any warranty, express or implied, or assume any legal liability or responsibility for the accuracy, completeness, or usefulness of any information, apparatus, product, or process disclosed, or represent that its use would not infringe privately owned rights. Reference herein to any specific commercial product, process, or service by trade name, trademark, manufacturer, or otherwise, does not necessarily constitute or imply its endorsement, recommendation, or favoring by the United States Government, any agency thereof, or any of their contractors or subcontractors. The views and opinions expressed herein do not necessarily state or reflect those of the United States Government, any agency thereof, or any of their contractors.

Printed in the United States of America. This report has been reproduced directly from the best available copy.

Available to DOE and DOE contractors from

U.S. Department of Energy
Office of Scientific and Technical Information
P.O. Box 62
Oak Ridge, TN 37831

Telephone: (865) 576-8401
Facsimile: (865) 576-5728
E-Mail: reports@adonis.osti.gov
Online ordering: <http://www.osti.gov/bridge>

Available to the public from

U.S. Department of Commerce
National Technical Information Service
5285 Port Royal Rd.
Springfield, VA 22161

Telephone: (800) 553-6847
Facsimile: (703) 605-6900
E-Mail: orders@ntis.fedworld.gov
Online order: <http://www.ntis.gov/help/ordermethods.asp?loc=7-4-0#online>



SAND2014-5039
Unlimited Release
Printed July 2014

Process for Estimating Likelihood and Confidence in Post Detonation Nuclear Forensics

John L. Darby, Charles M. Craft
Sandia National Laboratories
P.O. Box 5800
Albuquerque, New Mexico 87185-MSXXXX

Abstract

Technical nuclear forensics (TNF) must provide answers to questions of concern to the broader community, including an estimate of uncertainty. There is significant uncertainty associated with post-detonation TNF. The uncertainty consists of a great deal of epistemic (state of knowledge) as well as aleatory (random) uncertainty, and many of the variables of interest are linguistic (words) and not numeric. We provide a process by which TNF experts can structure their process for answering questions and provide an estimate of uncertainty. The process uses belief and plausibility, fuzzy sets, and approximate reasoning.

Intentionally Left Blank

Contents

1. Introduction.....	9
2. Classical Probability and Statistics	13
3. Bayesian Approach	15
4. Belief/Plausibility Measure of Uncertainty	23
5. Linguistic Variables	27
6. Evidence and the Calculation of Belief/Plausibility	29
7. Application to Technical Nuclear Forensics.....	35
8. Incorporation of Intelligence Information	41
9. Conclusions and Recommendations	45
Appendix A. Example of LinguisticBelief [®] Software	47
A.1 Variables Used in LinguisticBelief [®]	47
A.2 Creating a LinguisticBelief [®] Model.....	48
Appendix B. Example of Bayesian Update for Unknown Coin	54
Appendix C. Evidence and Calculation of Belief and Plausibility	57
C.1 Difference Between Assigned Evidence, Belief, and Plausibility	57
C.2. Results as a Belief to Plausibility Interval	59
References	65
Distribution	66

Figures

Figure 3-1. Prior for Coin Example: $\beta[3,6]$	16
Figure 3-2. Posterior for Coin Example: $\beta[3,7]$	17
Figure 3-3. Posterior for Coin Example: $\beta[4,16]$	17
Figure 3-4. Posterior for Coin Example: $\beta[12,106]$	18
Figure 3-5. Comparison of Results for Coin Example.	19
Figure 4-1. Any Probability Distribution Is Consistent with Assignment of Evidence: The Figure Illustrates Three Possible Choices for Total Ignorance.	24
Figure 6-1. Example of Assignment of Probabilities for <i>Adversary Level of Technical Expertise</i>	30
Figure 6-2. Example of Assignment of Evidence for <i>Adversary Level of Technical Expertise</i> . ..	30
Figure 7-1. Belief and Plausibility Distribution for “What is the Category of the Device?”	37
Figure 7-2. Forced Probability Distribution for “What is the Category of the Device?”	38
Figure 7-3. Bins for Likelihood and Confidence for “What is the Category of the Device?”	40
Figure 8-1. Focal Elements for Source of Device Before TNF Assessment.	41
Figure 8-2. Belief and Plausibility for Source of Device Before TNF Assessment.	42

Figure 8-3. Belief and Plausibility for Source of Device After TNF Assessment.....	43
Figure A-1. Simple Model in LinguisticBelief®.....	48
Figure A-2. Approximate Reasoning Rule Base for Threat Partially Completed.	49
Figure A-3. Example of Evidence Assigned to Consequence.	49
Figure A-4. Pooled Evidence for Consequence.....	50
Figure A-5. Risk for a Scenario.	51
Figure A-6. Sensitivity Analysis Results for Variable Risk.	52
Figure A-7. Results of Trace for <i>Adversary Perception of Consequence</i>	53
Figure A-8. Actual Model for Terrorist Scenarios.....	53
Figure B-1. Uniform Prior for P(H).	54
Figure B-2. Posterior for P(H) With Two Heads from Two Tosses.	55
Figure B-3. Posterior for P(H) with 10,000 Heads from 10,000 Tosses.	56
Figure C-1. Probability Distribution for No Uncertainty.....	59
Figure C-2. Focal Elements for No Uncertainty.....	60
Figure C-3. Probability Distribution.	60
Figure C-4. Focal Elements for a Probability Distribution.....	60
Figure C-5. Focal Elements Where Information Insufficient to Assign Probability Distribution.	61
Figure C-6. Belief and Plausibility Distribution.	61
Figure C-7. Focal Elements for Unknown Coin: Total Ignorance.....	62
Figure C-8. Belief and Plausibility for Unknown Coin: Total Ignorance.....	62
Figure C-9. Focal Elements for “What is the Category of the Device”?	63
Figure C-10. Belief and Plausibility Distribution for “What is the Category of the Device”?.....	63
Figure C-11. Focal Elements and Event for <i>the Device is Not Cat 1</i>	64

Tables

Table 5-1. Example Rule Base for Approximate Reasoning.....	28
---	----

Nomenclature

GUM	guide to uncertainty in measurement
PDF	probability density function
TNF	technical nuclear forensics

Intentionally Left Blank

1. Introduction

This report proposes an approach for estimating the likelihood and confidence for key questions related to technical nuclear forensics (TNF) that captures the significant uncertainty associated with the estimate. For example, a key question is “What is the category of the device?” Many variables are combined in a complex functional relationship to answer this question, including prompt data, radiochemistry results, and expert opinion. Such questions for TNF have common complexities:

- Significant uncertainty in the body of information results in the need for expert opinion.
- Many of the variables are not numeric.
- The body of information consists of both objective data and the state-of-knowledge experts use to form expert opinion.

Epistemic, or state-of-knowledge, uncertainty is different from aleatory, or stochastic, uncertainty. Techniques that use probability as the measure of uncertainty—either classical or Bayesian—do not capture epistemic uncertainty very well and therefore can significantly underestimate the uncertainty whenever significant epistemic uncertainty is present.^a We propose the mathematics of belief/plausibility to develop the measure of uncertainty for TNF; belief/plausibility is an extension of probability theory that better captures epistemic uncertainty.^{1,12}

Some variables of concern for addressing the key questions are inherently non-numeric in that there is no appropriate numerical range for such variables. For example, the variable *sophistication of material X* is better considered using the choices “crude”, “moderate”, and “exotic” instead of some arbitrary numerical range such as 0 to 100. We should not force an unknown numeric scale when combining many such variables to answer a top level question since the relative numerical ranges for the different variables will affect the answer to the question. Instead, we should model such variables linguistically—that is, using words. We model the words using the mathematics of fuzzy sets.³

To answer questions of concern, we must functionally combine numerous variables. For linguistic variables, we need a process to combine the variables. Some of the variables are numeric, and to consider them in functional combinations with non-numeric variables, we bin the outcomes of the numeric variables into fuzzy numbers and assign the fuzzy number linguistic fuzzy sets. We use a rule base to combine the linguistics of different variables. We use the mathematics of approximate reasoning on fuzzy sets to specify the combinations.^{b,3}

We have developed and applied a process for addressing problems with significant epistemic uncertainty and linguistic variables. The process is based on established mathematical techniques and is implemented in custom software.^{1,3,6,12} We call our implementation of these mathematical techniques the Linguistic Belief technique. We provide an example application of the software in appendix A.

^a In general the value of a random variable is uncertain. Uncertainty can be quantified using the probability measure, axiomatically developed by Kolmogorov, by developing a probability density function (PDF). However, with significant epistemic uncertainty the PDF is not known.

^b For a linguistic variable that is a function of other linguistic variables, a rule base provides a mapping of the linguistic outcomes of the input variables to the linguistic outcomes of the output variable. An example of such a rule base is provided in section 5. The rule base is an implementation of approximate reasoning because the rule base contains fuzzy sets; that is, the functional outcomes are approximate (fuzzy).³

We have applied this process and software tool to numerous problems, such as issues related to terrorist activities and nuclear weapons applications. In a recent project, working with nuclear weapons experts we applied the process to four low-level questions for TNF as examples of its application to TNF.

We are concerned with answering a question such as “What is the category of the device?” We chose to answer the question by binning the answer into categories of devices, 1 through 5, where each category has certain defined characteristics. Since we cannot examine the device, but must infer the category based on the body of information available after an event (real or exercise), we have uncertainty in our answer.

We wish to reflect the uncertainty in our answer using likelihood and confidence. However, likelihood can have many connotations: the frequency of an event, the probability of an event, the expected value of a random variable, and so forth. Similarly, confidence can have many connotations: a confidence interval in classical statistics, the standard deviation of a random variable, the probability a random variable is within a certain interval, and so forth. Therefore, we defer a specific discussion of what we mean by likelihood and confidence until after we have summarized mathematical techniques to address uncertainty. In section 7, we define explicitly what we mean by likelihood and confidence using the technique we recommend—the linguistic belief technique.

We summarize three different techniques for estimating uncertainty and discuss when they are appropriate for use. The three techniques are:

- Classical Probability and Statistics.
- Bayesian.
- Linguistic Belief (our name) using a combination of Belief/Plausibility, Fuzzy Sets, and Approximate Reasoning.

For many variables important to TNF, we argue that the classical technique cannot be used because there are no data. Similarly, we argue that the Bayesian approach should not be used as it significantly underestimates the uncertainty for many variables because the assumed prior probability distribution is unknown and there is limited or no information to generate an accurate posterior probability distribution. Then, we demonstrate how the third approach—linguistic belief—can be applied to capture all the uncertainty inherent in TNF.

We do not provide in this paper detailed mathematical discussions of the techniques, such as discussions of conjugate priors for the Bayesian approach or convolution of focal elements for the belief/plausibility measure of uncertainty. Instead we provide a summary of the techniques and focus on how the fidelity of information available determines which technique should be used. All the techniques are based on proven mathematical concepts discussed at length in the technical literature.^{1,3,12}

Both the classical and Bayesian approaches use the probability measure of uncertainty; the Bayesian approach broadens the classical objective probability to a subjective concept. Linguistic Belief uses the belief/plausibility measure of uncertainty, an extension of the probability measure. The belief/plausibility measure of uncertainty specifically addresses the epistemic uncertainty inherent using expert opinion, whereas the probability measure of uncertainty—using either the classical or Bayesian approach—does not. However, since the belief/plausibility measure of uncertainty reduces to the probability measure of uncertainty when the uncertainty is all aleatory (no epistemic uncertainty), the linguistic belief technique includes the classical and Bayesian techniques as special cases.

The appropriate technique to use depends on the fidelity of the information available. If sufficient data are available, the appropriate technique should be classical or Bayesian. When information is missing or insufficient and state-of-knowledge is used to accommodate the missing information, then the linguistic belief technique is an appropriate approach to estimate uncertainty.

In sections 2 and 3 we discuss classical and Bayesian techniques in some detail to explain their shortcomings for adequately reflecting all the uncertainty inherent in TNF where significant epistemic uncertainty is present. For example, although a Bayesian approach allows subjectivity, it still uses the probability measure of uncertainty; for example, by requiring the assignment of a prior probability distribution.

Intentionally Left Blank

2. Classical Probability and Statistics

For the classical technique, the probability of an event E, $P(E)$, is

$$P(E) = \lim_{\tau \rightarrow \infty} \frac{N(E)}{\tau} \quad (\text{Equation 1})$$

where $N(E)$ is the number of times event E occurs out of τ independent trials.

Note that:

- The classical probability is one value with no uncertainty, and
- The classical probability is never truly known as an infinite number of trials are required to know it precisely.

In basic applications of classical probability it is assumed that $P(E)$ is known. For example, if we assume a coin is “fair” there is an equal likelihood of heads (H) or tails (T) after a toss; that is, $P(H) = 0.5$ and $P(T) = 0.5$.^c

As another example, for a population of motor driven pumps failure to start is described by the binomial distribution, and a parameter of the distribution is the probability of failure of a pump to start on demand, $P(FS)$.^d With a large body of test data, we can estimate $P(FS)$ as the fraction of failures from the sample trials in which a start signal was issued to the pump. For example, if we observe 100 failures to start in 10,000 attempts to start the pump, we estimate $P(FS)$ as 0.01.^e

When there is insufficient information available to assume that $P(E)$ is known, and if we can take a sample, then we can use statistical inference. We take a sample from the population and infer the parameters of an assumed probability distribution. For the pump example, we may take a sample of pumps and evaluate the number of failures to start in a certain number of trials, and from that sample infer $P(FS)$. The larger the number of trials, the more accurately we can estimate $P(FS)$. If we have a large population of pumps, collect and test a sample of 45 pumps, and observe no failures, we are “90 percent confident that $P(FS)$ is no greater than 0.05”; that is, 0.05 is a 90 percent one-sided upper confidence level for $P(FS)$.^f We are 90 percent confident that $P(FS)$ is somewhere in the interval $[0, 0.05]$. The

^c H and T are not independent. H and T are all the mutually exclusive outcomes of the universal set (sample space) for the coin toss, so once we know $P(H)$ we know $P(T) = 1 - P(H)$.

^d Standard probability distributions are defined by parameters. For example, the binomial distribution has two parameters: the number of items and the probability of failure for each item.⁸

^e The number of tests is large; we assume the true $P(FS)$ from equation 1 is the fraction of failures observed from the large number of tests.

^f The sample is for 45 pumps selected at random from a large population of pumps. Sampling can be performed with or without replacement. Sampling without replacement does not return a selected pump back to the population; sampling with replacement does return a selected pump back into the population. Sampling without replacement is described by the hypergeometric probability distribution; sampling with replacement is described by the binomial probability distribution. For a large population, it does not matter if the sampling is performed with or without replacement since the likelihood of selecting the same pump more than one time is small when sampling with replacement from a large population; the hypergeometric distribution approaches the binomial distribution for a large population. For a sample of 45, if the population has more than about 2,100 items the

confidence interval is a measure of our uncertainty. The larger the sample, the smaller the confidence interval, and with a sample of infinite size we know $P(FS)$ with no uncertainty.

Classical probability and statistics use objective measures: they are based solely on data. With sufficient data, or what we call high fidelity information, classical probability is the appropriate technique to use. Unfortunately, in TNF many variables have limited or no data.

estimated confidence is the same whether sampling with or without replacement.¹⁰ Let $UCL_\gamma(x, n)$ denotes the one-sided Upper Confidence Level (UCL) for the probability parameter of the binomial distribution for confidence $(1-\gamma)100$ percent for a sample of size n with x failures (the interval is

$[0, UCL_\gamma(x, n)]$). $UCL_\gamma(x, n) = \frac{(x+1)F_{1-\gamma}(2x+2, 2n-2x)}{(n-x) + (x+1)F_{1-\gamma}(2x+2, 2n-2x)}$, where F is the F-ratio distribution function of classical statistics.^{2,10} For the example here, γ is 0.1, n is 45, and x is 0; $UCL_{0.1}(0, 45)$ is 0.05.¹⁰ The Bayesian approach also defines “confidence” intervals, or more accurately probability intervals.² In contrast to classical probability where parameters are assumed fixed but unknown, a Bayesian approach treats parameters as random variables and therefore can provide intervals with a certain probability that the parameter is in the interval. Our Linguistic Belief approach develops a confidence interval using the belief/plausibility measure of uncertainty; see section 7.

3. Bayesian Approach

A Bayesian approach does not use the classical definition of probability of equation 1. Instead it considers probability as a subjective concept. Objective probability is based solely on data; subjective probability is based on the body of information, including state-of-knowledge as well as data. Since classical probability (objective) and this broader definition (subjective) refer to two completely different concepts, to be precise we should call the former objective probability and the latter subjective probability. We can have a subjective probability distribution for the objective probability of an event, $P(E)$, as defined in equation 1, to reflect our uncertainty in $P(E)$. We treat $P(E)$ as a random variable in the subjective approach, whereas in the classical approach $P(E)$ is a fixed value, perhaps unknown.

Both objective and subjective probability obey the mathematics of a Kolmogorov probability measure.^g

The Bayesian process can be applied to either events or to parameters of a probability distribution.² Here, we focus on a Bayesian consideration of the parameters of a probability distribution. Parameters of a probability distribution are assumed to be random variables for the Bayesian approach; not fixed—but perhaps unknown—values as assumed in the classical technique.^h

As an example of the Bayesian process, assume a pump manufacturer has developed an improved version of an existing pump designed specifically to have a higher reliability for starting on demand. Until the new pump is tested we cannot provide an objective probability for the parameter probability of failure to start, $P(FS)$, because we have no data. However, our subjective opinion is that the new pump is at least as reliable as the old pump since it was designed to be more reliable. Therefore, we assume initially that $P(FS)$ for the new pump is the same as for the old pump, a subjective assignment of probability. As test data for the new model become available, we can use a Bayesian process to update $P(FS)$ for the new pump.

We can always provide a subjective probability of an event based on our body of information. As we shall see, the issue is that with poor information the assumed prior probability distribution may be wrong, and that with insufficient data to update the assumed distribution the updated distribution will also be wrong. In application, the Bayesian approach is as follows:

- Assume a reasonable type of probability distribution for the events of concern. For example, the binomial distribution is assumed for failures of pumps for the pump example.
- Assume probability distributions *for the parameters* of the assumed probability distribution for the event.ⁱ In our pump example, the parameter of interest for the binomial distribution is the probability that a pump fails to start on demand, $P(FS)$. The beta probability distribution is

^g Kolmogorov axiomatically defined a probability measure with three axioms.⁸ The mathematics of probability is based on these three axioms.

^h In classical probability the probability of an event $P(E)$ is assumed to be fixed, but perhaps unknown, as defined in equation 1. A Bayesian approach considers $P(E)$ to be a random variable (not fixed) with uncertainty. Similarly, in classical probability the parameters of a probability distribution are assumed to be fixed, but perhaps unknown, whereas a Bayesian approach treats the parameters as random variables.⁸

ⁱ The uncertainty in the probability distribution for the event is due to the uncertainty in the parameters of the distribution.²

typically assumed for this parameter.^j In our pump example we assume the random variable $P(FS)$ has a beta distribution.

- Gather more information and update the probability distributions for the parameters to reflect the improved body of information.

The initial probability distribution for a parameter is called the prior distribution, or simply the prior. The prior is updated with new information to form the posterior distribution or simply the posterior.

To better illustrate the Bayesian process, consider the following example. We have an unknown coin and wish to estimate the probability it is a fair coin, but we are limited in our ability to examine the coin. The coin is described by two mutually exclusive outcomes: heads (H) and tails (T). The parameter of interest is the probability of heads, $P(H)$.^k Given $P(H)$ we know $P(T)$ since $P(T) = 1 - P(H)$. Using a Bayesian approach we could subjectively assume as a prior for $P(H)$ that the new coin is fair based on our body of information about most other coins with which we have experience.^l For example, we might assume the prior of figure 3-1.^m This is a probability density function for the random variable $P(H)$. This prior was selected to have an expected value of 0.5 to reflect our subjective opinion that the coin is likely fair, and a standard deviation of 0.19 to allow for the coin to be biased to some extent.²

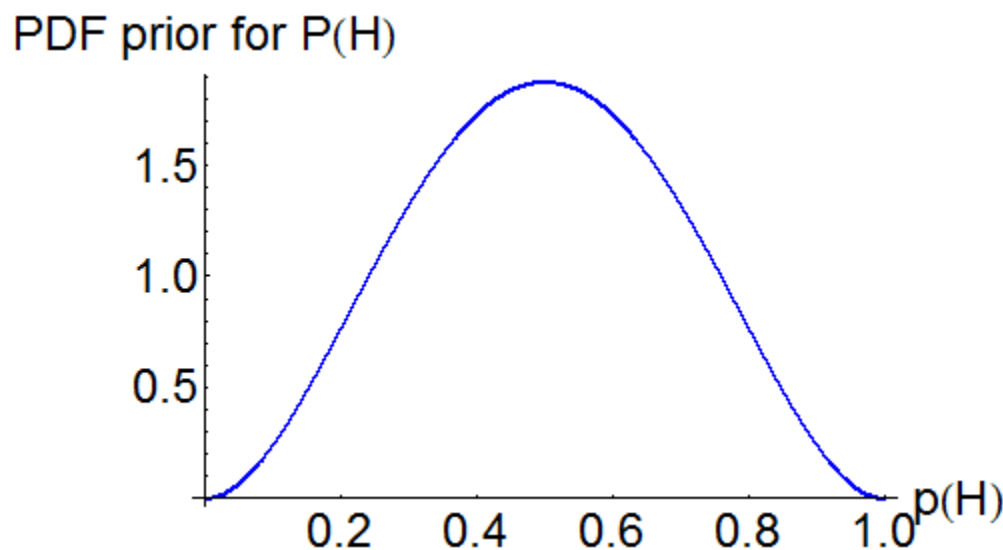


Figure 3-1. Prior for Coin Example: $\beta[3,6]$.

Then, as we test (toss) the new coin we can update the prior with objective information to generate a posterior for $P(H)$. Suppose we toss the coin one time and observe tails. The Bayesian process updates

^j A beta distribution is a very general probability distribution with range $[0, 1]$; the PDF can have many different shapes depending on the values of its two parameters. It is widely used in Bayesian analyses as a conjugate prior for the probability parameter of the binomial distribution.² The parameters of the beta distribution are selected to generate a prior probability distribution; for example, to match the $P(FS)$ for the existing pump.

^k The probability that a certain number of heads results from a specified certain number of tosses is described by the binomial distribution with the parameter $P(H)$.

^l Calculations in this paper for the Bayesian examples were performed with Mathematica.⁷

^m This prior is $\beta[3,6]$ where $\beta[x_0, n_0]$ is the beta distribution with parameters x_0 and n_0 .²

the prior with this objective information to generate a posterior as shown in figure 3-2.ⁿ The posterior has an expected value of 0.43 and a standard deviation of 0.18.

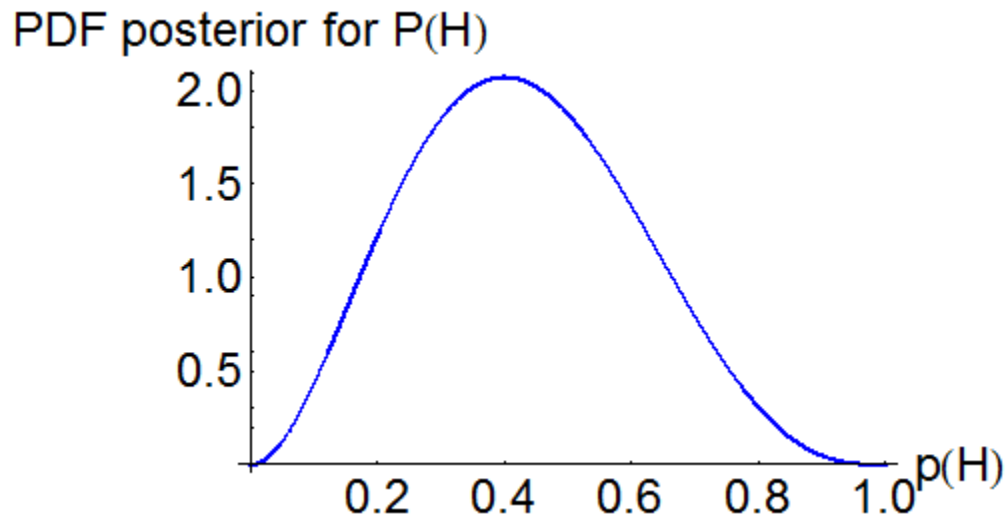


Figure 3-2. Posterior for Coin Example: $\beta[3,7]$.

With more test data, the objective information will further modify the prior. If we toss the coin 10 times and observe one heads, the posterior is as shown in figure 3-3.^o This posterior has expected value 0.25 and standard deviation 0.10. The data indicate the coin is biased toward tails.

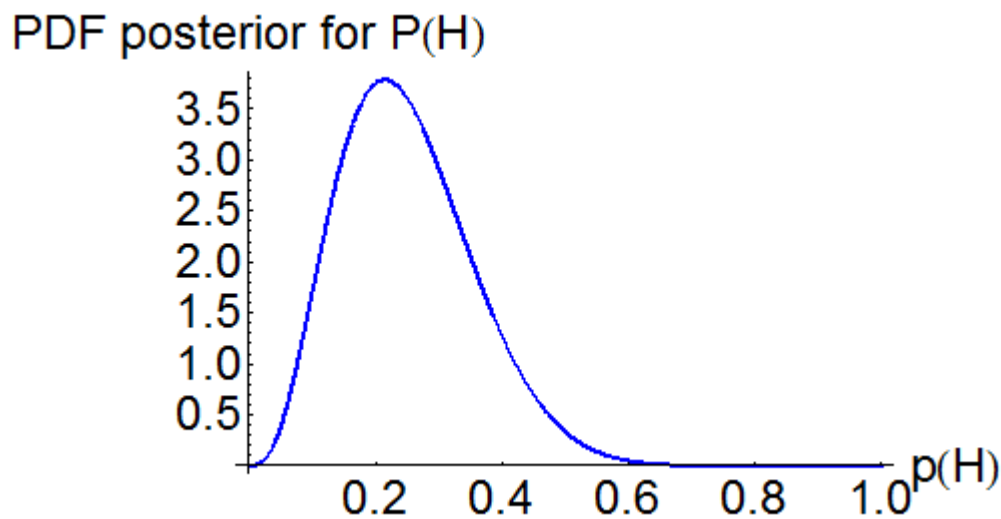


Figure 3-3. Posterior for Coin Example: $\beta[4,16]$.

ⁿ The posterior is $\beta[3 + 0, 6 + 1]$.²

^o The posterior is $\beta[3 + 1, 6 + 10]$.²

Suppose we have 100 tosses with nine heads. The posterior for this situation is shown in figure 3-4.^p The posterior mean is 0.11 and the standard deviation is 0.03.

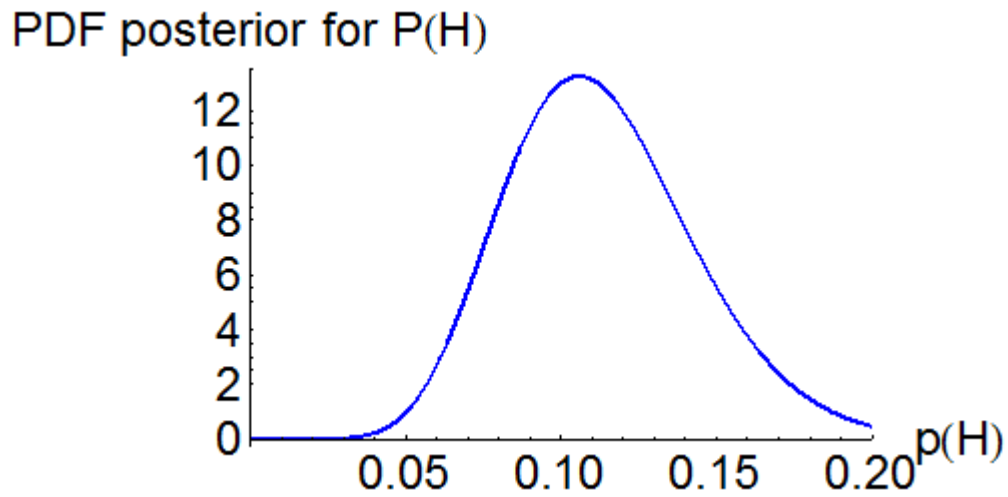


Figure 3-4. Posterior for Coin Example: $\beta[12,106]$.

Suppose the coin is from an ancient civilization where the heads side has raised embossing making the coin biased toward tails, such that the true probability of heads is only 0.1.^q Figure 3-5 compares the probability distributions of figures 3-1 through 3-4, showing how as we obtain more objective data we generate successively more accurate posterior distributions.

^p The posterior is $\beta[3 + 9, 6 + 100]$.²

^q If we could toss the coin many times, heads would turn up 10 percent of the time. See equation 1, where event E is heads.

PDF posterior for $P(H)$

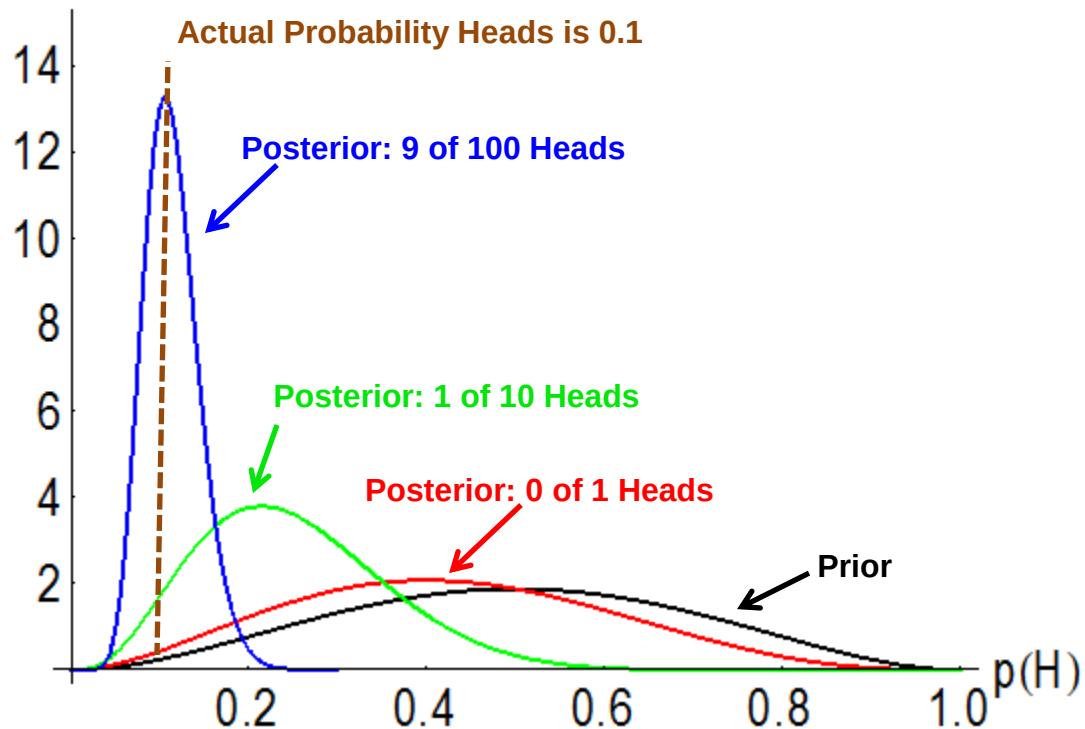


Figure 3-5. Comparison of Results for Coin Example.

With more objective data, the data tends to “dominate” the prior. For example, if we have 10,000 tosses with 1,000 heads, the posterior will have mean 0.1 and standard deviation 0.003, in agreement with the large amount of objective data that the coin is biased 90 percent tails, or 10 percent heads.

With an infinite number of tosses the posterior will have a mean determined by the objective data alone, regardless of the initial prior used, and the mean will be the same as provided by equation 1 using the classical technique. The posterior will have a standard deviation of zero indicating no uncertainty.

However, with little objective information, an update of a poorly guessed prior does not produce a good posterior.

For the coin example, the actual probability of heads is only 0.1, but we used an initial prior shown in figure 3-1, based on our somewhat biased assumption that the coin is most likely fair. Such an assumption is reasonable given our past experience with coins, but in this case it turns out to be grossly wrong. If we have only the results of one toss (with a tails), the posterior of figure 3-2 is highly inaccurate as our posterior best estimate for the probability the coin will come up heads is 0.43. Our “strong” prior overwhelms a “weak” amount of objective data and the first posterior erroneously indicates that the coin is almost fair. Even with more data, say nine tails from 10 tosses, although the posterior of figure 3-3 indicates the coin is biased toward tails, the best estimate probability of heads of 0.25 from the posterior is still not accurate.

Only with sufficient objective data, does the posterior accurately “correct” a poor prior. For the coin example, we need 100 tosses (with nine heads) to produce a posterior with a mean probability of heads (0.11 for our example) that is close to the true probability of heads which is 0.1. And it is not until we have 10,000 tosses that the posterior accurately reflects the correct distribution: probability of heads 0.1 with essentially no uncertainty.

The bottom line is this: without sufficient objective data, the posterior distribution produced by updating an inaccurate prior distribution remains inaccurate. (Conversely, appendix B shows that even with an ill-informed prior, with sufficient objective data the posterior is accurate.)

The Bayesian technique is appropriate when we have insufficient objective data for using the classical technique but when our body of information supports the assignment of a reasonably accurate prior distribution, or where we can gather significant new information to update the prior distribution. We call this case medium fidelity information.

We now consider situations in which it is not clear what prior distribution to use and sufficient information to update the prior cannot be obtained. This situation occurs often in TNF.

Given our example with an unfair coin, we are now a bit more wary of what initial distribution to assume. Perhaps a better way to evaluate the unknown coin is to not assume the coin is fair, but to assume the coin has equal probability for heads, $P(H)$, being any value in the interval $[0, 1]$. In this case, our prior distribution for $P(H)$ is a uniform probability distribution over $[0, 1]$. If we could toss the coin, using the Bayesian process we could update $P(H)$. However, if we cannot toss the coin, we cannot update this prior and our probability for $P(H)$ remains uniformly distributed over $[0, 1]$.^f With this assumption, the best estimate for $P(H)$ (the mean of the probability distribution) is 0.5 and the standard deviation for this uniform distribution is 0.29, so we could say $P(H) = 0.5 \pm 0.29$, with the understanding that the uncertainty here is one standard deviation.⁵ This estimate of $P(H)$ does include uncertainty. However, we assumed a uniform probability distribution for $P(H)$ and with total ignorance we have no justification for that assumption. *With total ignorance any assumed prior probability distribution is equally valid.* For example, with total ignorance any of the following assumed prior distributions would be equally valid: $P(H)$ uniform over $[0, 1]$, $P(H)$ with expected value near zero (biased tails), $P(H)$ with expected value of 0.5 (fair), $P(H)$ with expected value near 1.0 (biased heads), $P(H)$ a triangular probability distribution over $[0, 1]$ with expected value 0.7 (skewed toward 1), or $P(H)$ a lognormal distribution truncated at 0 and 1.

We face the following conundrum. Using the probability measure of uncertainty, even subjectively in a Bayesian sense, we are forced to assume a probability distribution; but with little information—no information for the case of total ignorance—we have no basis for preferentially selecting one distribution over another. However by a priori selecting a distribution we discard most of the epistemic uncertainty associated with our state-of-knowledge. Because data are limited, without sufficient information to update a poorly guessed prior distribution, the posterior distribution is also inaccurate, both for the expected value and for the uncertainty (that is, the standard deviation).

To accommodate these situations where we cannot make a wise choice for the prior distribution and have little data to update the distribution, we need an extension of the probability measure of

^f As shown in appendix B, even with a poor prior, given sufficient objective data the updated posterior will be accurate. Our problem for this example (and for many of the variables important for TNF) is that not only do we not know the prior, we have essentially no information upon which to update the prior.

⁵ For a uniform probability distribution over $[a, b]$, the mean is $(a + b)/2$, and the variance (square of the standard deviation) is $(b - a)^2/12$.⁸

uncertainty to better capture the epistemic uncertainty associated with the body of information available to us. The belief/plausibility measure of uncertainty discussed in the next section is such an extension.

Intentionally Left Blank

4. Belief/Plausibility Measure of Uncertainty

The belief/plausibility measure of uncertainty was developed by mathematicians and logicians to extend the probability measure of uncertainty to better consider epistemic uncertainty.^{1,3,12} Conceptually, belief/plausibility allows us to provide an estimate of the result being somewhere within a collection of outcomes, without forcing us to assign a probability of each outcome.[†] When we do not know the probability distribution but we force the use of a probability, we throw away uncertainty, specifically epistemic uncertainty. With belief and plausibility we do not throw away uncertainty.

Belief and plausibility are lower and upper bounds, respectively, on probability; for the situation where there is no epistemic uncertainty, only aleatory uncertainty, both belief and plausibility are the same, that is, the probability. Instead of requiring a probability distribution to specify the likelihood of each possible outcome as classical and Bayesian approaches do, belief/plausibility allows the assignment of evidence to collections of outcomes. Appendix C provides a simple explanation of the concept of assigning evidence.

Consider the earlier example of the unknown coin. With total ignorance, all we know is that $P(H)$ is somewhere in the interval $[0, 1]$, so *evidence* is assigned to the entire interval $[0,1]$. This assignment of evidence is *not* a probability distribution; we make no assumption as to how $P(H)$ is distributed over $[0, 1]$; all we know is $P(H)$ is somewhere within the interval $[0, 1]$. This assignment of evidence is different than assigning a probability distribution; for example, a uniform *probability* distribution for $P(H)$ mentioned before, forces equal probability that $P(H)$ is any value within $[0, 1]$.

Unlike either classical or Bayesian approaches that use a probability distribution, the belief and plausibility approach makes no assumption of a probability distribution over the values in the interval $[0,1]$. For example, the belief that that coin is two-headed is zero and the plausibility that the coin is two-headed is one; the probability is somewhere in the belief/plausibility interval $[0,1]$.[‡] This result has uncertainty that reflects our body of information: total ignorance in this case. Any assumed probability distribution is consistent with the belief/plausibility interval. For example:

- If we assume the coin is fair, $P(H)$ is 0.5, which is within the interval $[0, 1]$.
- If we assume the coin is two-headed, $P(H)$ is 1.0, also within the interval $[0,1]$.
- If we assume the coin is two-tailed, $P(H)$ is 0.0, also within the interval $[0, 1]$.
- If we assume $P(H)$ is uniform over the closed interval $[0, 1]$, any specific value of $P(H)$ is within $[0, 1]$.

Any assumed probability distribution for $P(H)$ is within the belief to plausibility interval generated based on the assignment of evidence. Figure 4-1 notionally illustrates this for the case of total ignorance.

[†] For a continuous variable over the real numbers, collections of outcomes are intervals. For a discrete variable, either numeric or linguistic, collections of outcomes are subsets of the set of all outcomes, or as discussed later, a collection of outcomes is an element of the power set of the set of outcomes.

[‡] $[$ and $]$ denote inclusion and $($ and $)$ denote bounds. For example, $[0,1]$ is the interval containing all real numbers between 0 and 1 including 0 and 1; $(0,1)$ denotes all real numbers within the interval where the greatest lower bound is 0 and the least upper bound is 1.

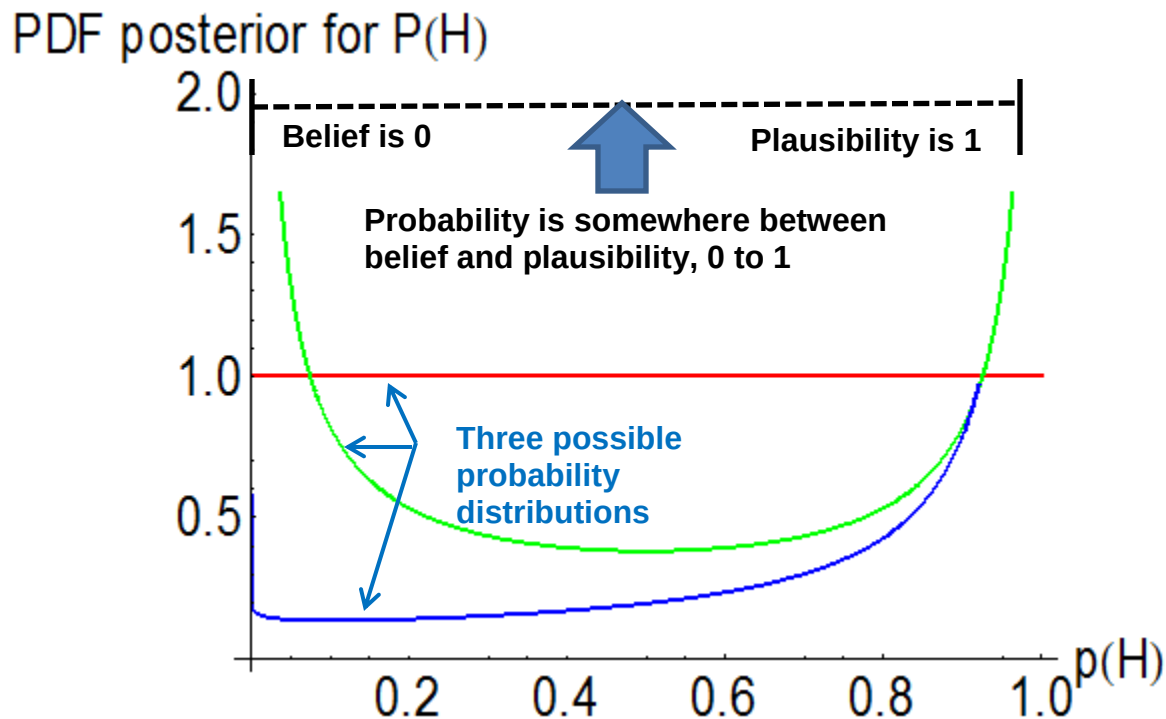


Figure 4-1. Any Probability Distribution Is Consistent with Assignment of Evidence: The Figure Illustrates Three Possible Choices for Total Ignorance.

For TNF, we have many important variables for which we have no objective data, so we cannot use classical probability and statistics. Also, we have many variables for which our body of information is insufficient to assign an accurate prior probability distribution and for which no or limited information is available to perform a Bayesian update. Therefore, we need to use the belief/plausibility measure of uncertainty to not throw away the significant amount of epistemic uncertainty for our problem, thereby leading decision makers to think that we have significantly more certainty (that is, less uncertainty) in our answers to key questions than is actually present.

In TNF exercises we have had cases where the error bars from different laboratory's estimates of the same parameters do not overlap, leading to great consternation in making a final joint estimate of the parameter and interpretation of the meaning of that parameter. This situation is the result of discarding significant amounts of epistemic uncertainty in estimating the error bounds on parameters of interest, which in turn leads to the situation where final judgments are improperly caveated and may give a false sense of our state-of-knowledge.

Although assignment of a probability distribution cannot be supported for many of the important variables, a probability distribution can be assigned to some variables for which sufficient information is available, using either classical or Bayesian techniques. Belief/plausibility seamlessly incorporates probability as a special case; therefore some—but not all—of the variables for questions pertinent to TNF can be assigned probability distributions.

The experts involved in the TNF evaluation do have a significant body of information. Therefore, we are not faced with the extreme case of total ignorance discussed above in the unknown coin example.

However, for many variables the body of information is not sufficient to assign a probability distribution. The body of information is used to assign evidence to collections of outcomes from which belief and plausibility can be calculated; if the evidence is assigned only to individual outcomes, both belief and plausibility are the same, the probability.

The belief/plausibility technique is appropriate when our variables are dominated by epistemic uncertainty. We call this the low fidelity information case.

Intentionally Left Blank

5. Linguistic Variables

Often variables are not numeric. For example, consider the variable *adversary level of technical expertise*. If we try to model this variable numerically, the question arises as to what is the appropriate scale? Is the range 1-to-100, 0-to-1,000, or some other range? One approach is to normalize the forced numeric range, for example 0 to 1. However, we combine many different variables in complex functional relationships to answer a question of interest, and some variables may be more important than others. A forced numeric normalization that has the same range, say 0 to 1, for all variables loses the relative importance among the variables. Therefore, different normalized ranges are necessary, say 0 to 1 for one variable and 0 the 1,000 for another. However, for the large number of variables of concern, an accurate, relative forced numeric normalization of all variables each with respect to one another that is applicable for each functional relationship is impossible to achieve.

We should not force a numeric scale as it will artificially drive the results of our evaluation. Instead, we should model the variable linguistically, that is, using words. For example, we could model *adversary level of technical expertise* as {minimal, moderate, advanced} with clear definitions for each linguistic outcome, such as “advanced” meaning “possessing an advanced degree or equivalent training or experience”.

The TNF questions have both numeric and non-numeric variables. However, when some variables are naturally linguistic, then the reasoning must be all linguistic; that is, it is not possible to combine numeric and linguistic reasoning. In mixed variable type problems such as TNF, the numeric variables are assigned linguistic variables that are fuzzy numbers so they can be combined with non-numeric variables using approximate reasoning.

Approximate reasoning for our use is simply a rule base for combining linguistic variables.^v For example, assume we wish to combine *adversary level of technical expertise* with *resources available* to form a new variable *likelihood material XYZ is in device*. Suppose we model *adversary level of technical expertise* as above, that is {minimal, moderate, advanced}, *resources available* as {few, some, significant}, and *likelihood material XYZ is in device* as {low, medium, high}. With these variables we form the rule base shown in table 5-1. These variables, their possible linguistic outcomes, and the rule base for the functional combination are derived from the state-of-knowledge of the respective subject matter experts.

The linguistic outcomes are inherently fuzzy sets in that they are not precisely defined. For example we do not provide a precise definition of the transition from few to some for adversary level of technical expertise. The fuzziness is consistent with the precision the experts need to reason on the question of interest by defining the approximate reasoning rules.

^v Since we are dealing with subjective words, we need to specify the “mathematics” for how words combine, and we do that with a rule base. Since the words in the rule base are fuzzy, the rule base is an approximate reasoning rule base.

Table 5-1. Example Rule Base for Approximate Reasoning

Adversary Level of Technical Expertise	Resources Available	Likelihood Material XYZ is in Device
Minimal	Few	Low
Minimal	Some	Low
Minimal	Significant	Low
Moderate	Few	Low
Moderate	Some	Medium
Moderate	Significant	Medium
Advanced	Few	Low
Advanced	Some	Medium
Advanced	Significant	High

For this example, *likelihood material XYZ is in device* is determined from the rule that the *likelihood material XYZ is in device* is the minimum (lowest capability) of either *adversary level of technical expertise* or *resources available*. For example, even if the *adversary level of technical expertise* is advanced, if the *resources available* is few, the *likelihood material XYZ is in device* is low, as indicated in the rule base.

The rule base captures how outcomes of linguistic variables combine functionally. The rule base does not capture our uncertainty in the outcome of the variables. For example, the rule base of table 5-1 states that if *adversary level of technical expertise* is advanced and if *resources available* is significant, then *likelihood material XYZ is in device* is high; or, if *adversary level of technical expertise* is moderate and if *resources available* is few, then *likelihood material XYZ is in device* is low. For a given TNF event (real world or exercise), the analysts have uncertainty for the variables, reflected by the assignment of evidence to basic variables specific to that event as discussed in the next section.

A complete model consists of both basic and rule-based variables. Basic variables are those assigned evidence by subject matter experts. A rule-based variable is a variable that is a function of other variables (basic and or rule based) as defined by an approximating reasoning rule base.

Once evidence is assigned to basic variables, uncertainty (belief and plausibility) is calculated for all the variables in the model. For the rule-based variables, uncertainty is calculated using the mathematics of convolution for the belief to plausibility measure using the functional relationships specified in the rule base. The end result is an estimate of uncertainty for every variable in the model.^w

^w If we were using probability distributions, the uncertainty for all the variables would be estimated by convolution of the probability distributions for all the basic variables per the functional relationships specified in the rule base. The belief/plausibility measure of uncertainty has a similar convolution process discussed in the literature.^{1,3,12} A simple explanation of the convolution process for belief/plausibility is provided in a report.¹¹

6. Evidence and the Calculation of Belief/Plausibility

Belief and plausibility are calculated based on evidence, which is the body of information available. Mechanistically, evidence is proportioned to elements of the power set (set of all subsets) for a given variable. The subsets of the power set to which evidence is proportioned are called focal elements. The proportioning of evidence is sometimes called assigning evidence and focal elements are sometimes called degrees of evidence.^x The focal elements are in general not probabilities, but represent proportioned evidence from which lower and upper limits on probability (belief and plausibility, respectively) can be calculated.^y Appendix C discusses these concepts in more detail and provides a simple discussion showing how focal elements are produced by weighting the body of information available (proportioning evidence).

Consider the variable *adversary level of technical expertise* from the previous section. The set of all outcomes for this variable is {minimal, moderate, advanced}. A probabilistic assignment of evidence is a special case where each collection of outcomes assigned evidence is a unique outcome. Each focal element contains only one outcome, and for this special case the value of the focal element is the probability of the outcome. The value of a focal element is designated by m .

For example, we could assign 20 percent chance (0.2 probability) to minimal, 10 percent chance to moderate, and 70 percent chance to advanced. This probabilistic assignment is illustrated in figure 6-1. The probabilities of the outcomes are: 0.2 for minimal, 0.1 for moderate, and 0.7 for advanced.

^x More details for the following summary are in the references.^{1,3,12} Evidence is the body of information available. We proportion (assign) evidence to elements of the power set of the outcomes for a given discrete variable. In the linguistic belief process the outcomes are fuzzy sets in the approximate reasoning rules, but are treated as crisp (classical) sets for the assignment of evidence. The power set is the set of all subsets. For example, if the variable has outcomes {a, b, c}, the power set is {null, {a}, {b}, {c}, {a, b}, {a, c}, {b, c}, {a, b, c}}. A set with n elements has 2^n elements in the power set. Each element of the power set that is assigned evidence is called a focal element. (The null set is never assigned evidence since the original set is required to contain all mutually exclusive outcomes.) If there is certainty in the outcome, say a, there is one focal element with value 1.0 assigned to {a}. If all the uncertainty is aleatory, the focal elements are the singletons of the power set, {a}, {b}, and {c}, and each focal element is a probability. With epistemic uncertainty, we have insufficient information to assign a probability to each singleton, and in general the focal elements are not singletons, for example {a, b}. The focal element value $m(A)$ assigned to the element A of the power set expresses the proportion to which all available and relevant evidence supports the claim that an outcome of the variable will be exactly in A (in A alone, and not in any subset of A). The sum of all focal elements is 1.0. From the focal elements, belief and plausibility can be calculated. For element A of the power set, the belief is the degree of evidence that the outcome will be in A or any subset of A. The plausibility is the degree of evidence that the outcome will be in A, any subset of A, or any set that overlaps A. The literature sometimes calls a focal element a basic probability assignment; that nomenclature is not used here, since a focal element is not a probability (unless all the focal elements are singletons) in the sense of probability theory but is used to calculate bounds on probability: belief and plausibility.

^y If each focal element is assigned to only one individual outcome, we have the special case where belief and plausibility are equal, the probability. For this special case, the focal elements for each outcome are the probabilities for each outcome. But in general, a focal element is not a probability.

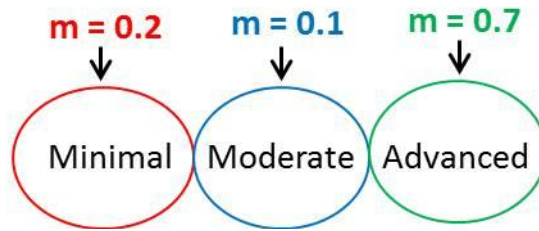


Figure 6-1. Example of Assignment of Probabilities for Adversary Level of Technical Expertise.

However, with the epistemic uncertainty inherent in the body of information available, we may not be able to assign probabilities.

In general, an assignment of evidence allows consideration of collections of outcomes, and the set of all collections of outcomes is the power set for the set of outcomes. The power set consists of all combinations of any number of outcomes, plus the null set (no outcomes); for example, null, {minimal}, {minimal, advanced}, and {minimal, moderate, advanced} are elements of the power set for the set of outcomes for *adversary level of technical expertise*.² If a set of outcomes has “n” elements, there are “n” individual outcomes, and the power set has 2^n elements. For *adversary level of technical expertise*, there are $2^3 = 8$ elements in the power set.

Just as outcomes are assigned a probability if we use the probability measure of uncertainty, the appropriate elements of the power set are assigned evidence if we use the belief/plausibility measure of uncertainty. Focal elements are in general not probabilities, but the bounds on probability (belief and plausibility) are calculated using the focal elements. Figure 6-2 is an example assignment of evidence to three specific collections of outcomes; that is, to three focal elements. Focal elements are represented graphically by circles that contain collections of outcomes assigned evidence. Subject matter experts weight the focal elements (assign evidence) by considering the credibility of the evidence. The value (weight assigned) of a focal element is designated by m; for example, in figure 6-2, the focal element containing {minimal, moderate} has weight $m = 0.2$.

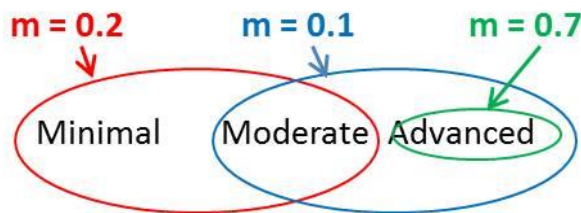


Figure 6-2. Example of Assignment of Evidence for Adversary Level of Technical Expertise.

² Order does not matter for a combination; order does matter for a permutation. For example, {Minimal, Advanced} and {Advanced, Minimal} are the same combination but are different permutations.

The probability of each outcome using the evidence in figure 6-1 is as follows:

- Somewhere in the belief to plausibility interval [0, 0.2] for minimal.
- Somewhere in the belief to plausibility interval [0, 0.3] for moderate.
- Somewhere in the belief to plausibility interval [0.7, 0.8] for advanced.

Appendix C discusses how belief and plausibility are calculated from the focal elements. For a specific outcome, the belief is the value of a focal element that contains only the outcome of concern, and the plausibility is the sum of all of all focal elements that contain the outcome of concern (either alone or with other outcomes) . For example, for moderate no focal element contains only moderate so the belief for moderate is zero. Two focal elements, with values 0.1 and 0.2, contain moderate (in addition to other outcomes) so the plausibility for moderate is 0.3. Similarly, for advanced one focal element with value 0.7 contains only advanced so the belief for advanced is 0.7; two focal elements with values 0.7 and 0.1 contain advanced so the plausibility for advanced is 0.8.

The belief and plausibility are easy to visualize using the graphical representation of the focal elements as in figure 6-2. An outcome *will* happen to the extent the evidence is specific to the outcome, and this is the belief or the lower bound on the probability for the outcome. An outcome *may* happen to the extent that the evidence overlaps the outcome to any degree, and this is the plausibility or upper bound on the probability for the outcome. (For the special probabilistic case of figure 6-1, each assignment of evidence is to individual outcomes only so belief and plausibility are the same.)

The sum of all the assigned evidence is normalized to 1.0; therefore, each focal element in each of figures 6-1 and 6-2 have an assigned value between zero and one, and the sum of the values of all the focal elements is 1.0.

Evidence is the body of information available, including both objective data and the state-of-knowledge. If the body of information is all objective, the proportioning of evidence as focal elements specific to individual outcomes as in figure 6-1 is a probability assignment, and the focal elements are the respective probabilities of the possible outcomes. However, if the body of information is dominated by state-of-knowledge with significant epistemic uncertainty, there is insufficient information to assign a probability to each individual outcome, and the focal elements are not probabilities, but represent proportioned evidence from which lower and upper limits on probability (belief and plausibility, respectively) can be calculated. Here, focal elements can include one or more outcomes as in figure 6-2. The credibility of the evidence as considered by subject matter experts is the basis for the assignment of evidence to focal elements. Consider the following example.

Assume we have a variable which we term *presence of isotope ABC*, and suppose this variable has the possible values {yes, no}.^{aa} This example is trivial, but it illustrates the concept of assignment of evidence to focal elements in the belief and plausibility context. If the radiochemistry unambiguously indicates the presence of isotope ABC and there is no uncertainty in this result, then the analyst would assign 1.0 to the focal element containing only yes and 0 to the focal element containing only no. In making this assignment the analyst is absolutely certain, based on the radiochemical and subsequent sample measurements, that isotope ABC is present in the sample. In this special case of absolute certainty, we actually have a probability distribution because each possible outcome has a unique value, namely 1 for yes and 0 for no; but there is no uncertainty.

^{aa} We could consider an additional outcome “maybe” to be used in the case the isotope of concern was identified by the radchemistry analysis, but could have come from contamination, not from the device. For simplicity of illustration, this additional complication is not included in the example.

Next consider the case where the radiochemistry and subsequent sample measurements are not so clear. Suppose the measurements were at or slightly above the limit of detection, so the actual presence of the isotope is ambiguous; however, due to some physics reason the analyst believes that isotope ABC should be present. In this case, the analyst has insufficient information (evidence) to assign a probability distribution (focal elements containing only individual outcomes); for example, a probability (focal element) 0.7 to yes and a probability (focal element) of 0.3 to no. Such a probabilistic assignment would state that a large number of additional samples obtained and analyzed in an identical manner would contain the isotope 70 percent of the time. However, such an assignment of probabilities would not represent the true situation, and there is no way the analyst could sufficiently and identically replicate the collection, preparation, and measurement of the sample to make such an assignment of probabilities. Rather, the analyst has this one measurement and his subject matter expertise to inform his decision.

Based on this body of information which includes this ambiguous measurement and the analysts' state-of-knowledge about the physics of what might be going on, the analyst feels that the focal elements for the variable *presence of isotope ABC* should be 0.7 for yes and 0.3 for both yes and no. With this assignment of evidence, the analyst believes (it is very credible) that yes is most likely the correct answer, but is unwilling to say with certainty that yes is *the* answer. The analyst has hedged the bet about the presence of isotope ABC by proportioning 30 percent of the evidence to both yes and no. This is not a probability assignment, because the analyst could not proportion the evidence across each outcome uniquely; that is, the analyst could not assign a probability of 0.3 to no based on the body of information.^{bb} The probability of no is somewhere in the belief/plausibility interval [0, 0.3], and the probability of yes is somewhere in the belief to plausibility interval [0.7, 1.0].^{cc}

References discuss the inclusion of epistemic uncertainty for radiochemistry results.^{4,9} In the Guide to Uncertainty in Measurement (GUM) approach epistemic uncertainty is referred to as Type B uncertainty; aleatory uncertainty is Type A uncertainty.

A reviewer of this paper provided the following insights relative to GUM.⁵

1. *Radiochemists base their Classical statistics uncertainty calculations on GUM (the ISO publication "Guide to the Expression of Uncertainty in Measurement"), in which epistemic uncertainty is referred to as "Type B" uncertainty (Gaussian uncertainty is "Type A") and is usually described as a rectangular probability distribution between, say, a and b (such as 0,1 in your presentation). The crucial questions are, how reliable is the estimate of Type B uncertainty, and how are the two types of uncertainty combined?*
2. *The above leads to the question of the meaning of the radiochemical analysis uncertainties reported by each laboratory for each radionuclide. Are they based on (1) counting statistics only, (2) counting statistics combined 'correctly' with 'correct' epistemic analytical uncertainties, or (3)*

^{bb} If the analyst tested a large number of samples and the results of each test were conclusive as to the presence of the isotope, then a probability assignment could be made. For example, if 70 percent of the samples contain the isotope and 30 percent do not, the probability of yes is 0.7 and the probability of no is 0.3.

^{cc} Appendix C discusses how belief and plausibility are calculated from the focal elements. For example, the focal element value 0.3 for {yes, no} is the degree of evidence assigned by the analyst to the credibility of the outcome being somewhere within {yes, no} but not further localized. "Somewhere within but not further localized" is called "being *exactly* within" in the literature.³ Here, exactly within {yes, no} means within {yes, no} but not further localized as to yes or no. The focal element 0.7 for {yes} is the degree of evidence assigned by the analyst to the outcome being yes. The belief for yes is the focal element for yes (0.7) and the plausibility for yes is the sum of all focal elements containing yes ($0.7 + 0.3 = 1.0$). The belief for no is zero since no focal element is assigned to no; the plausibility for no is 0.3 since a focal element with value 0.3 overlaps no.

counting statistics combined with epistemic analytical uncertainties where the combination or the epistemic value itself is 'incorrect'. Epistemic sources of error (e.g., non-random errors of calibration measurements, unreliability of individual analysts) of course exist in the laboratory as in the nuclear forensics world. The resulting problems with reported uncertainties of measurement when situations (1) or (3) occur can only be brought to light by laboratory inter-comparison programs to examine laboratory values that deviate from the known and the mean, have different uncertainty ranges than similar laboratories, or have ranges that do not overlap with those from other laboratories.

Our Linguistic Belief technique allows an estimate of the Type B uncertainty consistent with the state-of-knowledge, and has the appropriate mathematical approach built-in for considering both aleatory and epistemic uncertainty. The existing GUM approach is a special case of Linguistic Belief where a probability distribution is used to reflect the epistemic uncertainty.

For problems with significant epistemic uncertainty, focal elements based on the evidence are determined using expert opinion. In practice, we find that experts grasp the concept of assignment of evidence with a few hours of training, and readily use it whenever the body of information does not allow the assignment of a probability distribution. Of course, probability assignments are used for variables where sufficient information is available, but for many of the dominant variables the amount of information is not sufficient to support assignment of a probability distribution to the various outcomes of a variable.

Intentionally Left Blank

7. Application to Technical Nuclear Forensics

This section presents a relatively abstract discussion of the application of our process to post detonation TNF.

Consider the question “What is the category of the device?” We need to develop a model that provides the logical equation that can be used to evaluate any specific device. The model need only be developed once; it can then be later used to evaluate each specific event (real or exercise). A very simple analogy follows to emphasize the difference between the model (the equation) and the evaluation of an event (evaluating the equation given values for the variables). Consider the equation

$$Z = A + (B * C).$$

If we evaluate this equation for a specific case such as A of 3, B of 4 and C of 10 we have

$$Z = 3 + (4 * 10) = 43.$$

For another specific case where A is -6, B is 3, and C is 4,

$$Z = (-6) + (3 * 4) = 6.$$

Here the model is the equation and the specific values are the two instances of the evaluation of the model. This example should be rudimentary to all.

The body of information and the model for TNF are complex. For example, objective data with aleatory uncertainty include:

- Statistics of radioactive decay
- Measurement uncertainty in reference standards
- Variability in calibrations measurements
- Variability in measured constants (for example, decay constants)
- Measured variability in filter blank corrections

Examples of areas with significant epistemic uncertainty for TNF include:

- Incomplete understanding of physics
- Approximations in the physics models
- Restricted parameter space
- Imprecise knowledge of modeled parameters
- Applying models tuned for a selected design space to other design spaces
- Temporal variations in laboratory procedures
- Corrections for environmental contributions

The complex body of information can be organized into the following classes of information to structure the evaluation of questions of interest for TNF:

- Measured values
- Evaluated parameters
- Material modeling
- Device modeling



These classes are listed in increasing amount of epistemic uncertainty (decreasing amount of aleatory uncertainty); for example, measured values have more aleatory and less epistemic uncertainty than device modeling. However, there is epistemic uncertainty in all categories. For example, measured values have epistemic uncertainty associated with: different results from different laboratories for the mean $\pm$ standard deviation of the amount of an isotope, and corrections for environmental contributions. For the latter level categories, most of the uncertainty is epistemic.

Each class of information is a functional combination of constituent variables and in turn the different classes are combined with a functional relationship to produce a top level variable for “What is the category of the device?” These complicated relationships are created using expert opinion to generate the equation that evaluates to the answer to question “What is the category of the device?” The equation is a linguistic equation with the functional relationships described by approximate reasoning; it intrinsically allows a subsequent evaluation to capture uncertainty by the considering all (linguistic) outcomes for each variable.

A given event is evaluated by evaluating the equation— specifically, proportioning the evidence as focal elements (assigning evidence) across the family of the fuzzy sets for each basic variable in the model based on the body of information available.^{dd} Then, the evidence is convolved per the functional relationships using the mathematics of belief/plausibility up to the top level variable in the model. The result is a belief to plausibility interval for each outcome for every variable in the model, including the top level variable that represents the question of interest.

Capturing the functional structure of a model and evaluating it has been automated in a custom tool called LinguisticBelief[®]. This tool includes a sensitivity analysis capability to determine the basic variables that most affect the result. Appendix A discusses LinguisticBelief[®].

A model for determining the category of device has not been developed yet and will be classified once developed.^{ee} Here, we illustrate results of a notional model using the linguistic belief technique. Assume the results of the model are as shown in figure 7-1. The device cannot be assigned one category with certainty; furthermore, due to considerable epistemic uncertainty the probability that the device is a specific category is not precisely known.^{ff} This situation is reflected in general by the probability for a

^{dd} The model consists of basic variables that are assigned evidence, and rule-based variables that are functional combinations of basic and other rule-based variables. Appendix A provides a simple model as implemented in the LinguisticBelief[®] software.

^{ee} As part of an earlier project, classified models were developed and evaluated for four lower level questions related to TNF. These models are available in classified documents.

^{ff} The focal elements for the basic variables in the model are not singletons due to epistemic uncertainty, and as focal elements are convoluted up to the top level variable (the category of the device) per the functional relationships of the model, the probabilities of outcomes 1, 2, and 3 are no better known than being within a belief to plausibility interval. The focal elements for the basic variables do rule out the device being either category 4 or 5, so these categories each have zero probability.

category being within a belief to plausibility interval, belief being a lower bound and plausibility an upper bound on the probability. When belief and plausibility collapse to a single value, that value is the probability; here, the probability for categories 4 and 5 is zero. In figure 7-1, the category is 1, 2, or 3, since the probability for either category 4 or 5 is zero. The probability that the device is category 1, 2 or 3 is indicated by the belief to plausibility intervals in the figure. For example, the probability that the device is category 2 is somewhere within $[0, 0.5]$; the belief the device is category 2 is 0 and the plausibility the device is category 2 is 0.5. Similarly, the probability the device is category 1 is somewhere within $[0, 0.1]$, and the probability the device is category 3 is somewhere within $[0.5, 1.0]$.

Note the considerable epistemic uncertainty represented by the large belief to plausibility intervals for the probabilities for categories 2 and 3.

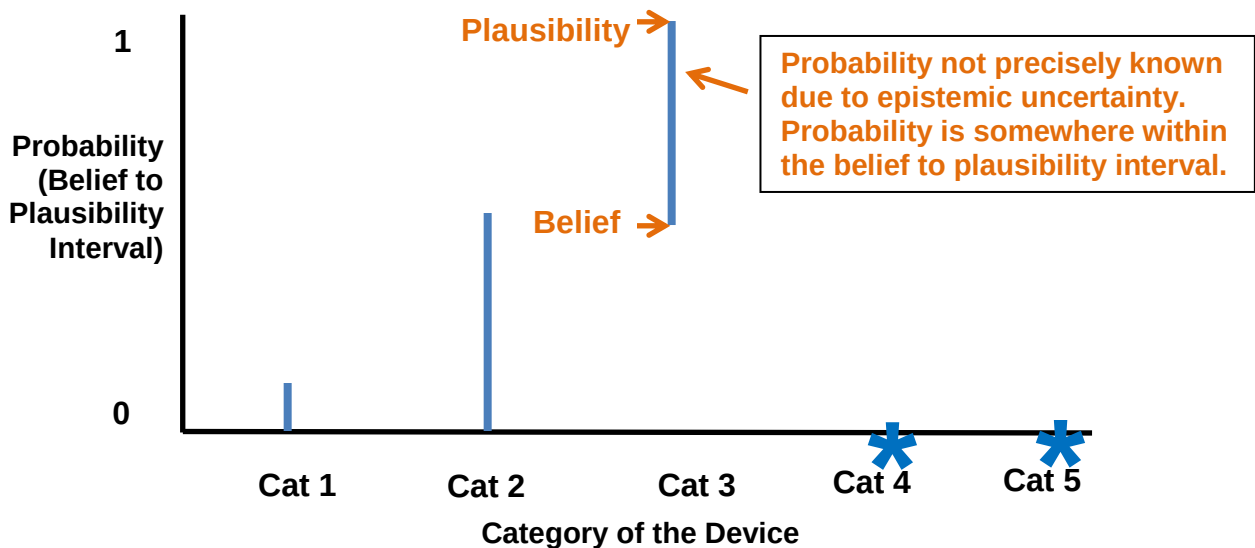


Figure 7-1. Belief and Plausibility Distribution for “What is the Category of the Device?”

Without using the belief/plausibility measure of uncertainty, and forcing the use of a probability measure over each outcome, we disregard epistemic uncertainty. If we force the use of a probability measure, the evidence must be proportioned as focal elements for each individual outcome.⁸⁸

If we had forced the assignment of a probability distribution across the fuzzy sets for every basic variable in the model for the category of a device, our results would be as indicated in figure 7-2, where the probability of each category is a point value. For example, we could arbitrarily force the probabilities of

⁸⁸ For example, suppose a basic variable is the *presence of isotope ABC* with outcomes {yes, no} as developed previously in section 6. The credibility of the evidence supports the following focal elements assigned by subject matter experts: 0.7 to yes and 0.3 to {yes, no}, resulting in the probability of No being somewhere in the belief/plausibility interval $[0, 0.3]$, and the probability of Yes being somewhere in the belief to plausibility interval $[0.7, 1.0]$. If we had forced the use of a probability measure, we could proportion the 0.3 assigned to {yes, no} as 0.15 equally to yes and no, resulting in a probability of 0.85 for yes and 0.15 for no. However there is no basis for how we break up the 0.3 for {yes, no} to a probability for both yes and no; we are equally justified assigning 0.25 to yes and 0.05 to no, resulting in a probability of 0.95 for yes and 0.05 for no. By forcing the use of a probability measure we are forcing the exclusion of epistemic uncertainty.

category 1, 2, and 3 to be 0.05, 0.15, and 0.80, respectively, as indicated in figure 7-2.^{hh} If we do this, our result has thrown away uncertainty. Alternately, we could have forced the probabilities for categories 1, 2, and 3 to be 0.1, 0.3, and 0.6, respectively. In both cases, we would have forced the probability of each outcome to be a precise value, indicating that our body of information was sufficient to know these probabilities, *whereas in reality our body of information was insufficient to know the probability of each outcome precisely*. All we can say is the probability of each outcome is bounded within an interval, that interval referred to as a lower bound of belief and an upper bound of plausibility. A more precise estimate of the probability would require that we reduce our epistemic uncertainty by improving our state of knowledge for all the basic variables- for example with replicated trials- to support the assignment of probability distributions to each basic variable.

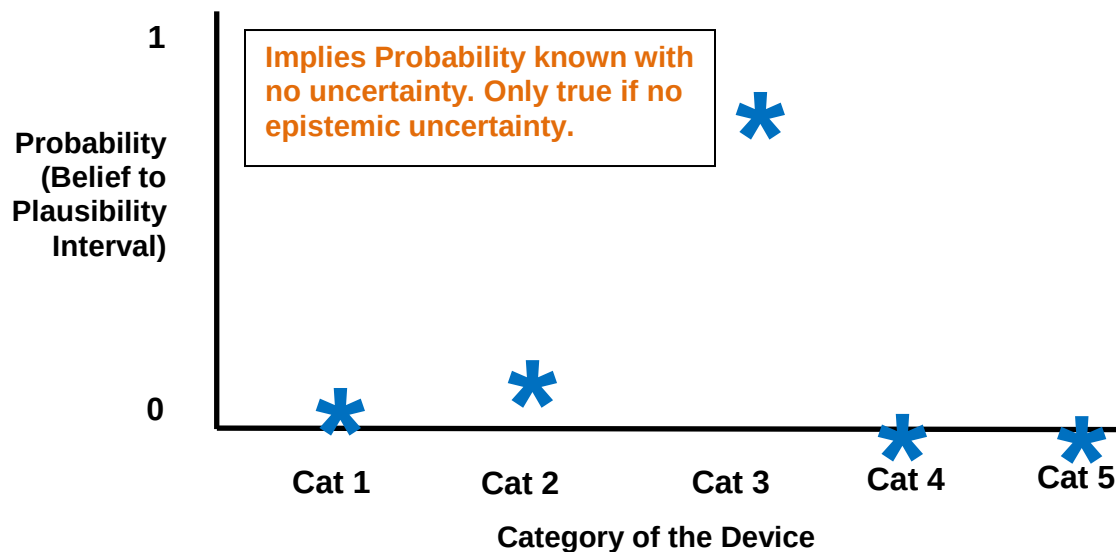


Figure 7-2. Forced Probability Distribution for “What is the Category of the Device?”

The more realistic result is shown in figure 7-1 where the probability for category 2 and 3 is only known within a rather large interval. By forcing the use of a probability assignment we have the results of figure 7-2 where the probability for category 1, 2, and 3 is estimated as a single value. However, the more precise estimates of probability in figure 7-2 are not accurate; they only result by forcing the use of a probability measure, and for cases where epistemic uncertainty does not support the assignment of a probability distribution we have artificially thrown away a great deal of the true uncertainty by selecting unsupportable probability distributions for all of the basic variables in the model.

Comparing figures 7-1 and 7-2, it is obvious that the forced use of probability distributions not supported by the evidence results in significant under-estimate of the uncertainty.

^{hh} Since the evidence did not provide precise probabilities, an assignment of a precise probability is just a guess. Only if the evidence resulted in each focal element having only one outcome do we know the probabilities precisely, and in this case both belief and plausibility are the same: the probability. Actually, evidence is assigned to the basic variables in the model and uncertainty in each rule-based variable is calculated by convolution, so forcing probability is accomplished by forcing all the focal elements of each basic variable to have only one outcome. Typically, to force a probability the value of a focal element is uniformly distributed as a probability across each element in the focal element. For example, for the focal element is {a, b, c} with value 0.3, we would assign probability 0.1 to each outcome a, b, and c.

Probability does capture aleatory (stochastic or random) uncertainty. For example, in figure 7-2 the device category is uncertain as reflected by no one category having a probability of one. However, with epistemic uncertainty, as emphasized previously, the probability of an outcome is not known resulting in the additional uncertainty for the probability of each outcome expressed as a belief to plausibility interval as in figure 7-1. The aleatory uncertainty is captured by a probability for an outcome, and the epistemic uncertainty is captured by uncertainty in the probability for an outcome.

We can now address likelihood and confidence. The results of the belief and plausibility distribution can be used to summarize the likelihood and confidence for a question of interest for TNF. One such way to define likelihood and confidence is as follows. For each outcome (answer to the question) likelihood is taken to be a point estimate of the probability, hopefully a best estimate. The belief to plausibility interval covers any assumed probability distribution consistent with the evidence, and any assumed probability distribution will provide a probability somewhere in the interval. For a point estimate probability, we assume all evidence containing more than one outcome is evenly distributed specifically to each outcome (a uniform probability distribution); this will give our point estimate probability such that the point estimate probabilities over all outcomes always sum to 1.0.ⁱⁱ As subsequently discussed we will only use the point estimate likelihoods to determine a linguistic bin for the results; that is the numeric point estimate values in themselves are not of major importance for our conclusions. Confidence is taken to be the interval containing the probability, specifically, the belief to plausibility interval. Our definition of confidence is similar in concept to that of a confidence interval in classical statistics (previously discussed) in that the probability is somewhere in an interval, here a belief to plausibility interval. These are numeric definitions of likelihood and confidence. Note that the likelihood and confidence easily fall-out from the results using the Linguistic Belief technique.

Consider the results of figure 7-1. For each category (1 through 5) the likelihood can be taken as a point estimate, obtained as previously described; the confidence can be taken as the size of the belief to plausibility interval. With this approach we have the following numeric likelihood and confidence results, assuming the point estimates as indicated result from the process just described:

- Cat 1 has likelihood of 0.05 and confidence of [0, 0.1]
- Cat 2 has likelihood of 0.25 and confidence of [0, 0.5]
- Cat 3 has likelihood of 0.70 and confidence of [0.5, 1.0]
- Cat 4 has likelihood of 0 with perfect confidence
- Cat 5 has likelihood of 0 with perfect confidence

We can also summarize likelihood and confidence linguistically.^{jj} If the experts can define linguistic bins for probability, such as [0, 0.2] for low, (0.2, 0.6] for medium, and (0.6, 1.0] for high, we can summarize the likelihood linguistically.^{kk} These numeric probability bins are for illustrative purposes only; in application they would be assigned by experts based on their connotations of low, medium, and high.^{ll}

ⁱⁱ Typically, the point estimate probabilities obtained in this manner for any outcome, are close to the values mid-way between the belief and the plausibility for that outcome.

^{jj} Other ways to summarize likelihood and confidence can be developed if desired by the TNF community. A linguistic definition is consistent with the “estimative language” approach considered by the intelligence community.

^{kk} (indicates a bound and [indicates includes. For example, (0.2, 0.6] is all real numbers greater than 0.2 up to and including 0.6.

^{ll} The TNF community may wish to define the bin ranges differently, and/or define more bins. If necessary, the bin endpoints could be fuzzy, but we do not discuss fuzzy bins in this paper.

Using these bins:

- Cat 1 has low likelihood
- Cat 2 has medium likelihood
- Cat 3 has high likelihood
- Cat 4 has low likelihood
- Cat 5 has low likelihood

We can consider the confidence linguistically by determining the range of linguistic bins covered by the belief to plausibility interval. If we consider both belief and plausibility together and both are within the same bin, then we have high confidence that the outcome is within the linguistic bin determined by the point estimate likelihood. That is, the belief to plausibility interval is contained within one bin and our uncertainty does not change our linguistic conclusion based on the point estimate likelihood. If belief and plausibility are not both within the same bin, we do not have high confidence that the outcome is in the bin determined by the point estimate likelihood, as the outcome could be in another bin. If the belief to plausibility interval spans two bins we assign moderate confidence; if the belief to plausibility interval spans three bins we assign low confidence.

This is best illustrated by an example. For example, figure 7-3 graphically summarizes the results of figure 7-1 indicating the linguistic bins.

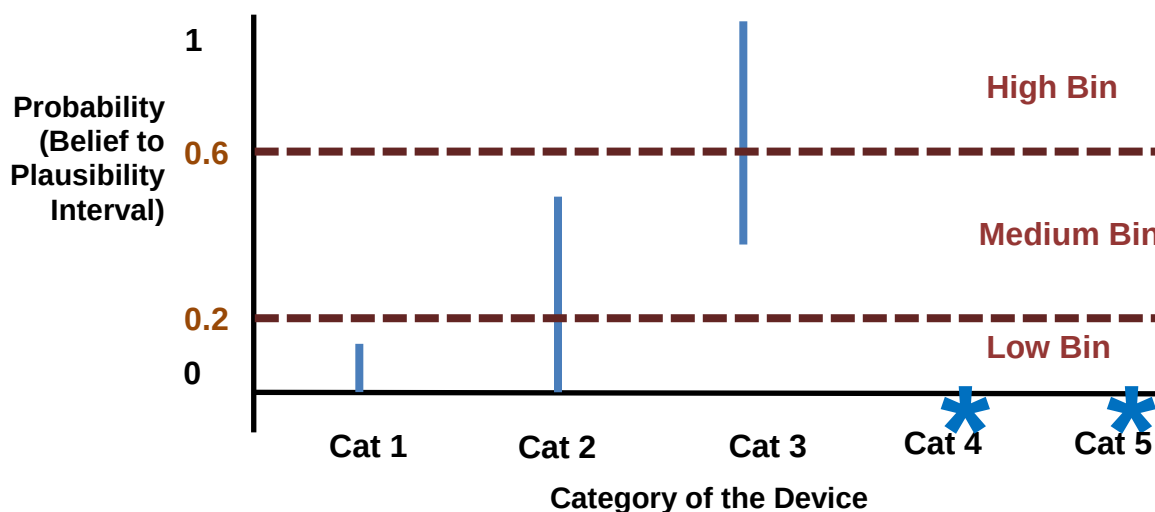


Figure 7-3. Bins for Likelihood and Confidence for “What is the Category of the Device?”

The results can be summarized linguistically as follows:

- Cat 1 has low likelihood with high confidence
- Cat 2 has medium likelihood with moderate confidence
- Cat 3 has high likelihood with moderate confidence
- Cat 4 has low likelihood with high confidence
- Cat 5 has low likelihood with high confidence

8. Incorporation of Intelligence Information

The TNF information can be combined with intelligence information to provide an estimate of the source of the device. Since intelligence information is mostly linguistic and has considerable epistemic uncertainty, our approach can seamlessly incorporate intelligence information.

For example, suppose the fuzzy sets for *source of device* are designated as {known, few, many, unknown}.^{mm} Few and many are fuzzy, and are based on the analysts' connotation for these linguistics; for example few could be on the order of 10 locations.

Suppose that after the event, but before the TNF assessment is completed, the credibility of the body of knowledge (purely intelligence information) leads the analysts to assign the evidence as depicted in figure 8.1.ⁿⁿ The analysts believe the source is not known and not few, but could be many or unknown. The analysts believe the source is likely unknown but hedge the assignment of evidence to include many resulting in the focal elements in the figure. The assignment of evidence for the source being unknown is 0.8, and unknown and many cover the entire possible outcomes based on the available intelligence information.

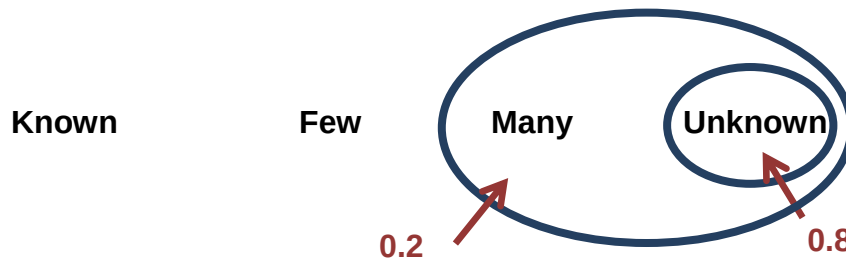


Figure 8-1. Focal Elements for Source of Device Before TNF Assessment.

^{mm} Few means known no more precisely than a few source locations, and many means known no more precisely than many source locations.

ⁿⁿ More than one expert (analyst) can assign evidence, resulting in pooled focal elements based on weighting each focal element from each analyst (typically equal weighting is used). The LinguisticBelief[®] code has a utility that performs this pooling.

The result for the source before the TNF assessment is completed is shown in figure 8-2.

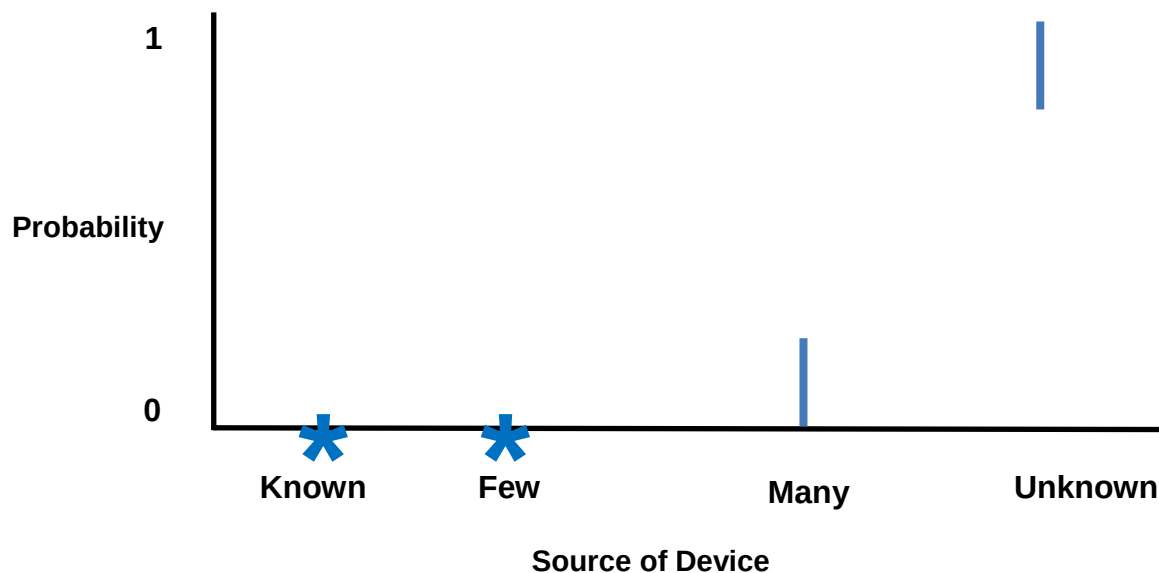


Figure 8-2. Belief and Plausibility for Source of Device Before TNF Assessment.

Based solely on the intelligence information, the source is many with a probability somewhere in the belief to plausibility interval $[0, 0.2]$. The source is unknown with a probability somewhere in the belief to plausibility interval $[0.8, 1.0]$. Using the likelihood/confidence bins discussed for figure 7-3, many has low likelihood with high confidence, and unknown has high likelihood with high confidence. Without more information, the source is essentially unknown.

The TNF assessment may provide definitive information, which if combined with the intelligence information would reduce the uncertainty in the source. For example, suppose the TNF assessment indicates the presence of a particular material or radionuclide with high probability, and the intelligence community believes a device with that material/radionuclide can only come from certain sources. A notional result for the source after the TNF assessment is combined with the intelligence information is shown in figure 8-3; it indicates the improved state of knowledge about the source of the device with the TNF assessment added to the body of knowledge.^{oo}

^{oo} In actual application, the TNF information and the intelligence information would be combined functionally using an approximate reasoning rule base.

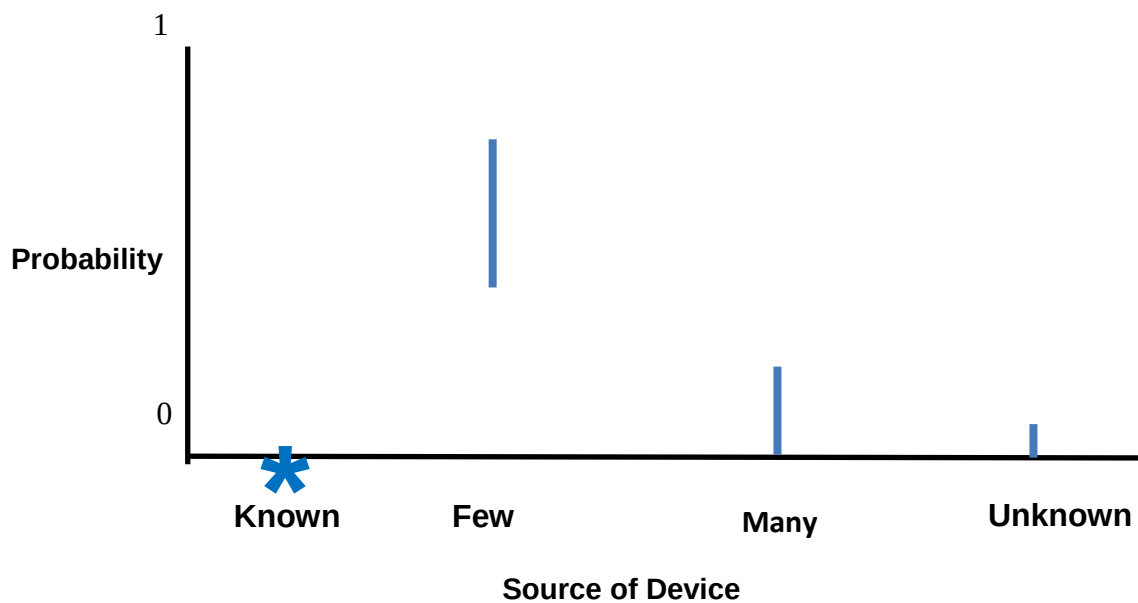


Figure 8-3. Belief and Plausibility for Source of Device After TNF Assessment.

The source of the device is more precisely known by including the TNF assessment in the body of information, compared to the source as evaluated using the intelligence information alone. With the TNF assessment the probability the device source is few is in the belief to plausibility interval $[0.45, 0.8]$, the probability the source is many is in the interval $[0, 0.2]$, and the probability the device source is unknown is in the interval $[0, 0.08]$. Using the likelihood/confidence bins discussed for figure 7-3, few has high likelihood with moderate confidence, many has low likelihood with high confidence, and unknown has low likelihood with high confidence.

Intentionally Left Blank

9. Conclusions and Recommendations

TNF evaluation of pertinent questions after an event requires consideration of complex functional relationships among many variables. Most of these variables have uncertainty, and although for some variables the uncertainty is aleatory and can be considered using the probability measure of uncertainty, many of the most important variables have significant epistemic uncertainty, which the probability measure of uncertainty is insufficient to capture.

Unless the epistemic uncertainty is accurately captured, the uncertainty in the answer to pertinent questions will be highly underestimated, thereby leading decision makers to think that we have significantly less uncertainty (more certainty) in our answers than is actually present.

For variables with significant epistemic uncertainty, classical probability cannot be used as there are no objective data. A subjective probability can be assigned in a Bayesian sense, but the body of information available is too poor to assign an appropriate prior, and little information can be gathered to perform a Bayesian update to generate a more accurate posterior. Since the Bayesian approach requires specification of a probability distribution, and for many variables there is little basis for selecting one distribution over another, selection of any probability distribution discards most of the epistemic uncertainty present.

With significant epistemic uncertainty, the probability measure cannot capture all the uncertainty, even if probability is considered subjectively as in a Bayesian approach. An extension of probability, the belief plausibility measure of uncertainty, can capture the epistemic uncertainty. Since belief/plausibility reduces to probability as a special case, it allows probability distributions to be assigned for variables where appropriate, but uses the broader assignment of evidence in general.

Many of the variables of concern for TNF are not numeric. Forced assignment of a numerical scale to such variables will bias the results based on the unknown scale used. Such variables should be considered linguistically, with outcomes specified as words (fuzzy sets) instead of numbers.

To combine linguistic variables to answer a question of concern, approximate reasoning can be used.

A process is needed that does the following:

- Capture the complex functional relationship among all variables pertinent to a question of concern.
- Consider both aleatory and epistemic uncertainty for each variable.
- Propagate uncertainty through the functional relationships to provide uncertainty for answers to the question of concern.
- Handle linguistic variables.

We have applied the mathematics of belief/plausibility, fuzzy sets, and approximate reasoning as previously developed by mathematicians and logicians, into a process called linguistic belief. We have implemented this process in custom software called LinguisticBelief[®] and have applied the process using the software to numerous complex problems, including four low-level questions for TNF.

The model for the answer to a question is developed once using a team of experts. This process provides the equation for evaluating the question. Given an event, the model is evaluated (that is, the equation is

evaluated) based on using the specific body of information available for that event to assign focal elements for the basic variables in the model.

We recommend that our technique be applied to develop a model for questions pertinent for TNF.

Appendix A. Example of LinguisticBelief[®] Software

The linguistic belief process described in this paper has been automated in the LinguisticBelief[®] custom Java code.⁶ The code implements the complex mathematics of the process and is used with subject matter experts to generate a model and evaluate a specific event. We have applied the software to numerous issues over the last seven years, including nuclear weapons applications, terrorist radiological sabotage scenarios, and selected issues for TNF. In this unclassified report we use a simple example that illustrates application of the software. The example is evaluation of the risk of a terrorist attack.

We use expert judgment to create the risk model, specify approximate reasoning rules, and assign evidence to variables for specific scenarios.

We define the risk of a terrorist scenario as:

$$Risk = Threat \times Vulnerability \times Consequence \quad (\text{Equation A1})$$

where “ $\times$ ” denotes convolution per an approximate reasoning rule base, not algebraic multiplication.

A physical security scenario includes adversary resources, the attack plan, and the target. Threat is the likelihood of the scenario. Vulnerability is the likelihood that the threat is successful in causing consequence. Consequence is the result of a successful scenario.

We evaluate threat from the perspective of the adversary (the terrorists) and vulnerability and consequence from the perspective of the defender (us). The adversary and defender each have different areas of uncertainty. For example, the adversary has more uncertainty than the defender about vulnerability, since the adversary has less knowledge of the possible security measures in place. The defender has significant epistemic uncertainty about the threat, but the adversary has no uncertainty about the threat as the adversary *is* the threat.

Since the adversaries have a choice of scenarios, they select a scenario based on their perception of the combination of vulnerability and consequence.

$$Threat = Adversary\ Perception\ of\ Vulnerability \times Adversary\ Perception\ of\ Consequence \quad (\text{Equation A2})$$

A.1 Variables Used in LinguisticBelief[®]

For ease of illustration, we will limit our example to the variables in equations A1 and A2. That is, we treat Vulnerability, Consequence, Adversary Perception of Vulnerability, and Adversary Perception of Consequence as basic variables to which evidence is assigned. In practice, these variables are further developed into numerous constituent variables. For example, consequence can be further developed as a combination of fatalities, injuries, economic loss, damage to national morale, fear in the populace, and so forth.

Note that fatalities, injuries, and economic loss are “hard” consequences, meaning they can be defined numerically. However, damage to national morale and fear in the populace are “soft” consequences in that we do not know the appropriate numeric scale to use; for these variables a purely linguistic

description is better than the forced use of an arbitrary numeric scale. Since we will be combining variables, many of which cannot be appropriately described numerically, we will treat all variables linguistically.

Each variable is either a basic or a rule-based variable. Basic variables are not developed further, and rule-based variables are formed by combinations of other variables, either rule-based or basic. For our simple example the basic variables are: vulnerability, consequence, adversary perception of vulnerability, and adversary perception of consequence. The rule-based variables are: threat and risk.

For each variable, we define linguistic fuzzy sets. For example, for Threat we define the fuzzy sets as {unlikely, credible, and likely}. For vulnerability we define the fuzzy sets {low, marginal, and high} and for consequence we define the fuzzy sets {low, moderate, major, and catastrophic}. The rest of the variables are similarly described with fuzzy sets.

A.2 Creating a LinguisticBelief[®] Model

To create the model for risk, the variables and their fuzzy sets are entered into LinguisticBelief[®], with the result indicated in figure A-1.

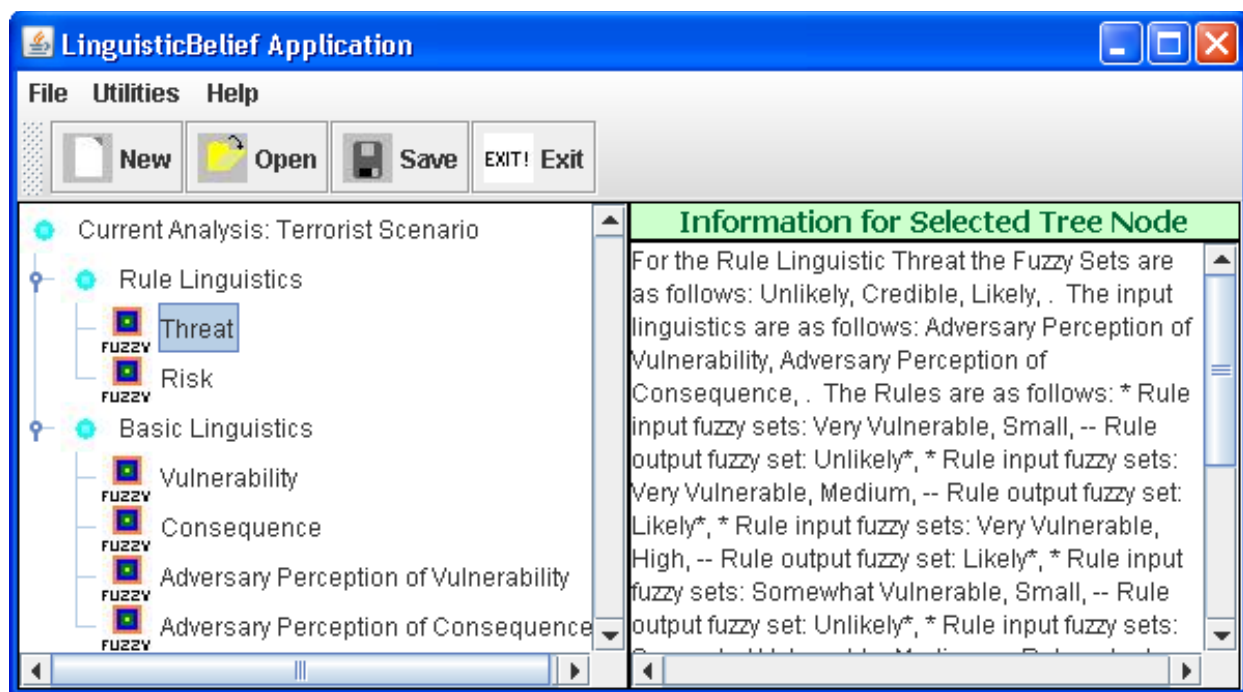


Figure A-1. Simple Model in LinguisticBelief[®].

The left panel in figure A-1 shows the variables in a tree structure. The right panel displays the current state of a selected variable, a node selected in the tree in the left panel. In figure A-1, the current state of threat is displayed. The state of the selected variable consists of the following: its fuzzy sets, its rule base if a rule-based variable, its focal elements if assigned (basic variable) or if calculated (rule-based variable), and its belief and plausibility if calculated.

For each rule-based variable, the variable is defined in terms of its constituent variables using an approximate reasoning rule base. Figure A-2 shows the approximate reasoning rule base for threat partially completed. The rules are developed using expert judgment.

Rules for RuleLinguistic: Threat		
Fuzzy Set for Input Linguistic: Adversary Perception of Vulnerability	Fuzzy Set for Input Linguistic: Adversary Perception of Consequence	Output Fuzzy Set
Very Vulnerable	Small	
Very Vulnerable	Medium	
Very Vulnerable	High	Likely
Somewhat Vulnerable	Small	Unlikely
Somewhat Vulnerable	Medium	Credible
Somewhat Vulnerable	High	
Not Vulnerable	Small	Unlikely
Not Vulnerable	Medium	
Not Vulnerable	High	

Specify Output Fuzzy Set for Selected Rule
Choices Are:

Accept Rules as Shown
Cancel
Unlikely
Credible

Figure A-2. Approximate Reasoning Rule Base for Threat Partially Completed.

Once all the rules have been created, the model is complete. The risk of a specific terrorist scenario is evaluated by assigning evidence (focal elements) to each basic variable. The evidence is assigned using expert judgment. Figure A-3 is an example of evidence assigned to consequence by one expert.

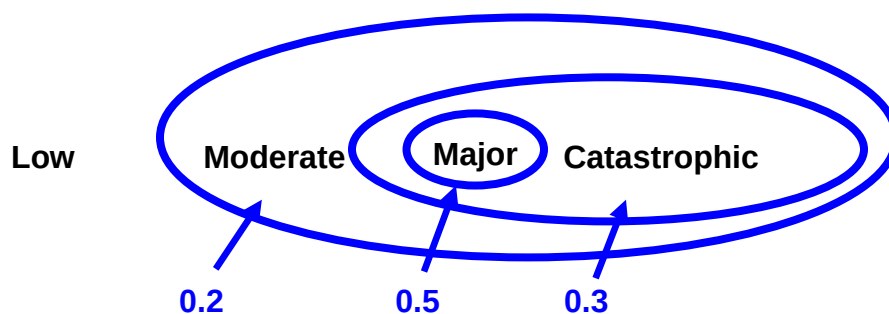


Figure A-3. Example of Evidence Assigned to Consequence.

The experts may not assign the same evidence, and the PoolEvidence code (a utility for LinguisticBelief[®]) is used to pool the focal elements from all the experts into one set of focal elements. The pooling weights each expert equally. Figure A-4 is an example of pooled evidence from two experts for consequence.

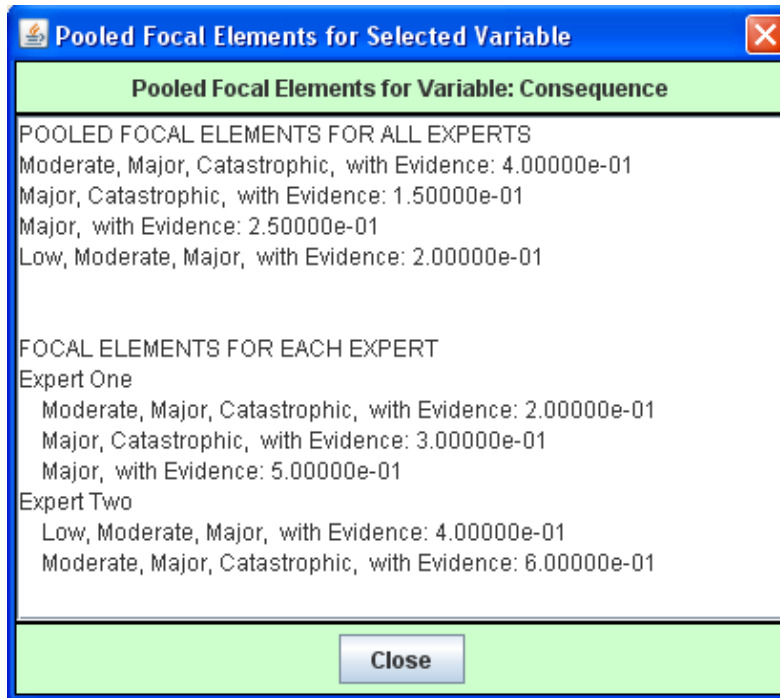


Figure A-4. Pooled Evidence for Consequence.

After all the pooled evidence is entered into LinguisticBelief[®], the scenario may be evaluated. Using the mathematics of belief/plausibility, LinguisticBelief[®] convolutes the evidence for the basic variables to produce focal elements for each rule-based variable. The variables are assumed to be non-interacting (independent). The belief and plausibility of any variable (basic or rule-based) can then be evaluated. Figure A-5 provides example results for risk for a scenario using dummy data. Two graphs are provided in figure A-5. The fuzzy sets ordered from “best” to “worst” in the view of the defender in the graphs. The top graph is the likelihood of a fuzzy set provided as a belief to plausibility interval. The bottom graph is the likelihood of exceedance of a fuzzy set.

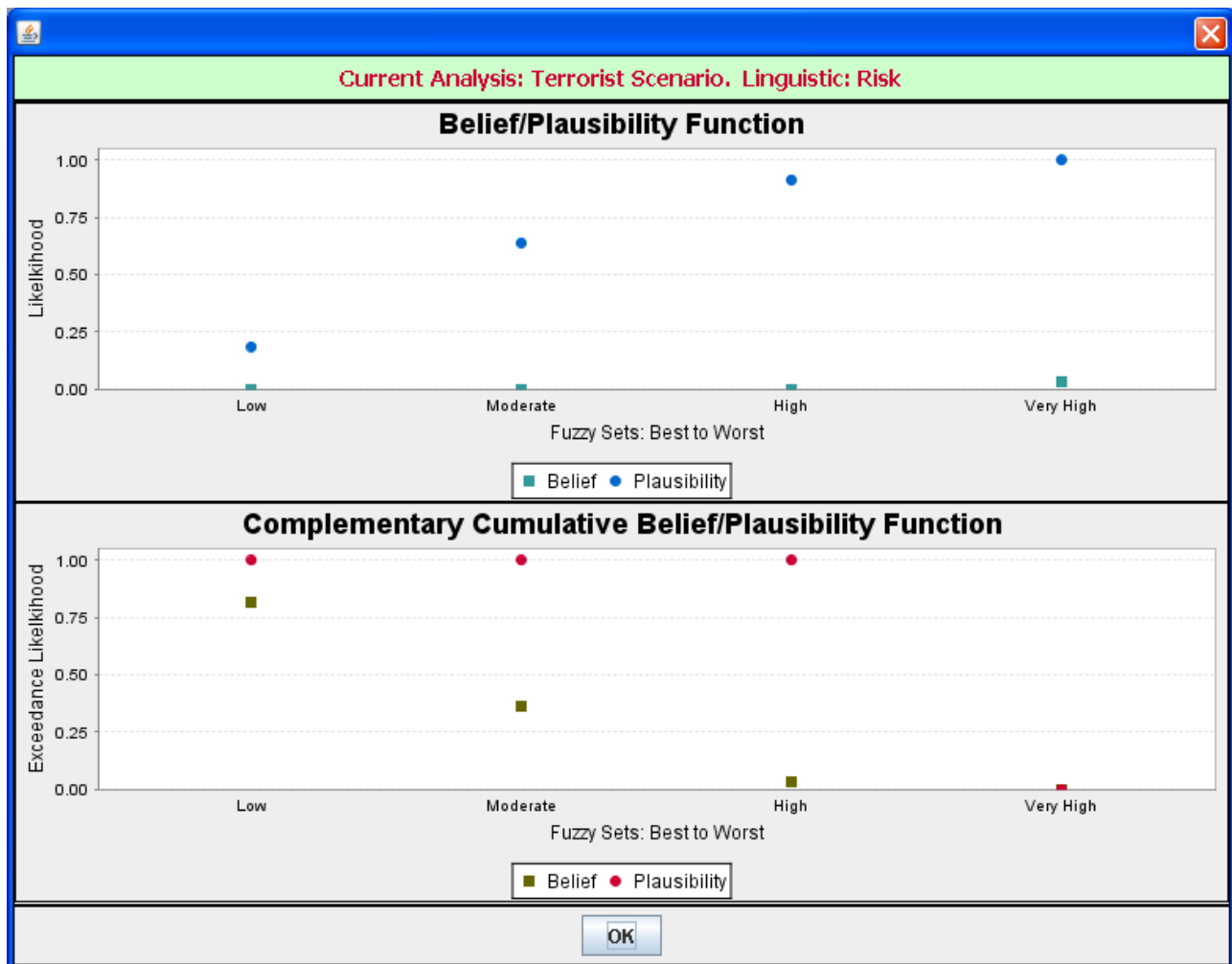


Figure A-5. Risk for a Scenario.

Scenarios are ranked from highest to lowest concern from the view of the defender as follows. Scenarios are ranked by non-zero plausibility of exceeding the “worst” fuzzy set (decreasing). Scenarios with equal ranking by plausibility are sub-ranked by belief of exceeding the “worst” fuzzy set (decreasing). The scenario in figure A-5 has a ranking Exceeds High with Plausibility 1.0 and Belief 0.03. If another scenario had plausibility/belief of exceeding High of 1.0/0.4 it would be ranked higher. If another scenario had zero plausibility of exceeding high, but plausibility/belief of exceeding moderate of 1.0/0.9 it would be ranked lower.

LinguisticBelief® includes a sensitivity analysis. Figure A-6 shows the results of the sensitivity analysis for the top level variable risk. Risk is most sensitive to consequence as indicated in the figure.

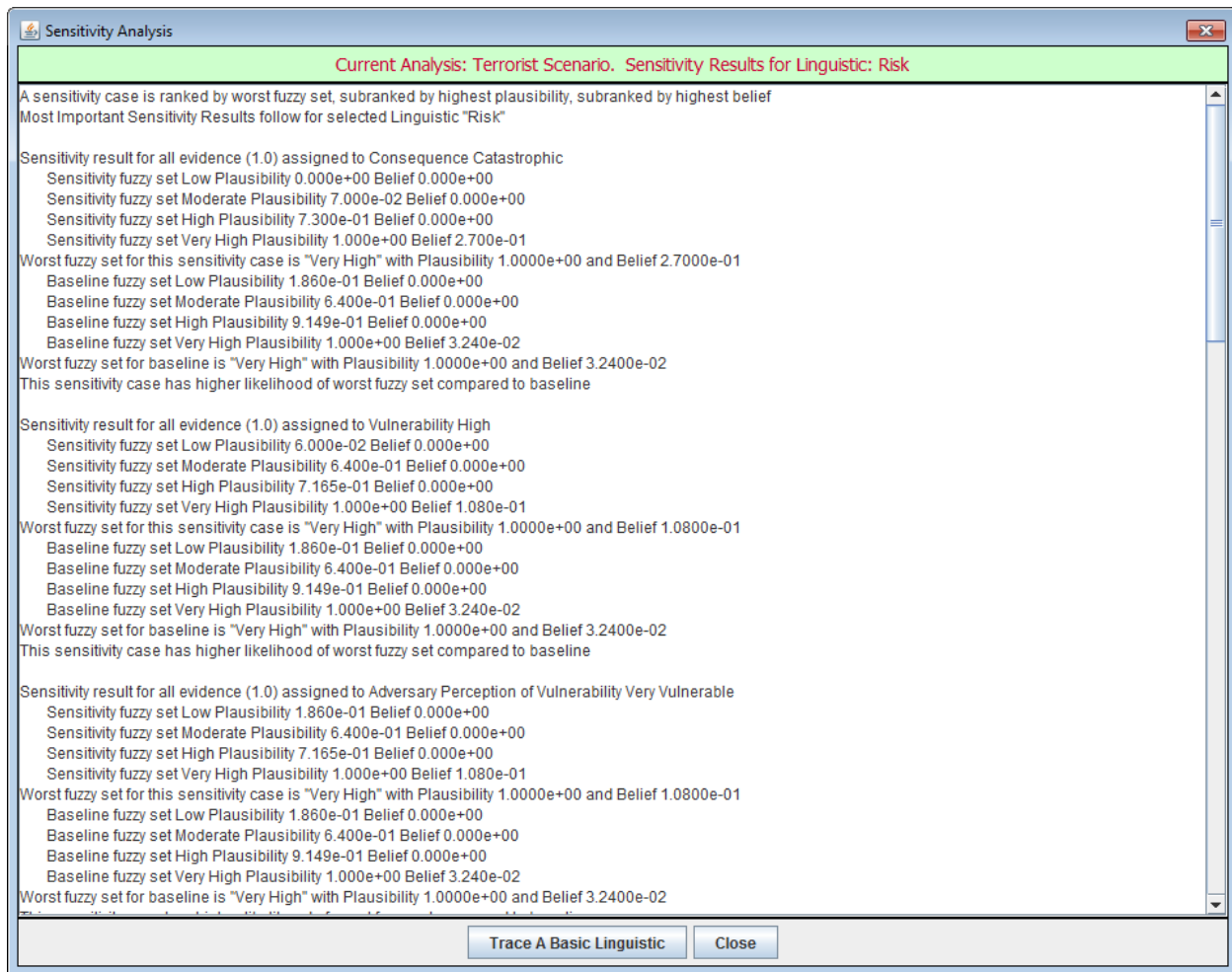


Figure A-6. Sensitivity Analysis Results for Variable Risk.

LinguisticBelief® has the capability to "trace" a basic variable, showing all the rule based variables to which it contributes, either directly or indirectly. Figure A-7 shows the results of a trace for *adversary perception of consequence*. *Adversary perception of consequence* is a direct input for threat, and in turn threat is an input for risk (so *adversary perception of consequence* is an indirect input for risk).

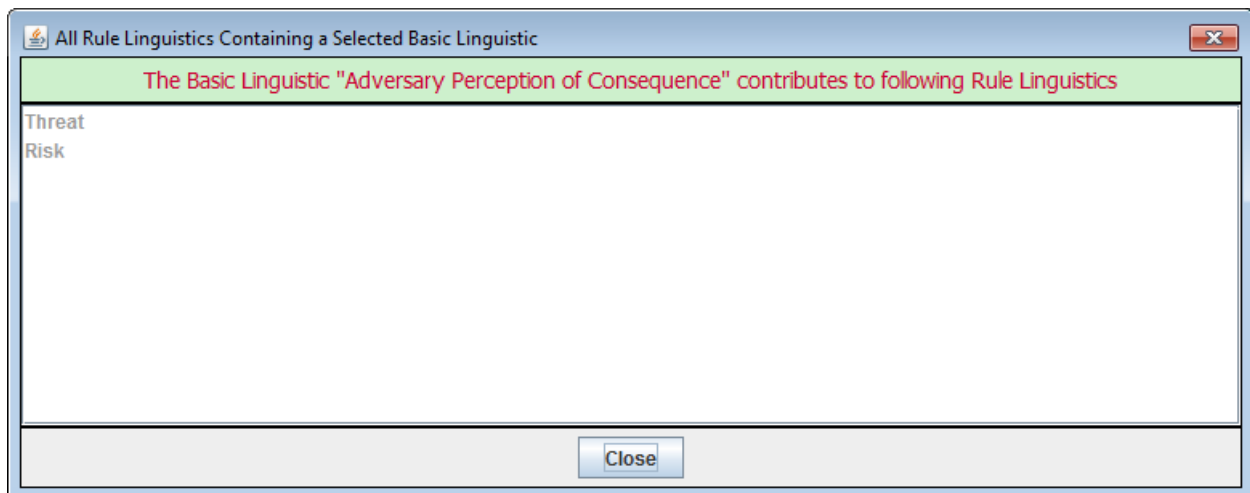


Figure A-7. Results of Trace for *Adversary Perception of Consequence*.

An actual application consists of many variables. Figure A-8 is an example of a more detailed model of terrorist scenarios. Here, the variables on the right hand sides of equation A-1 and A-2 are further developed.

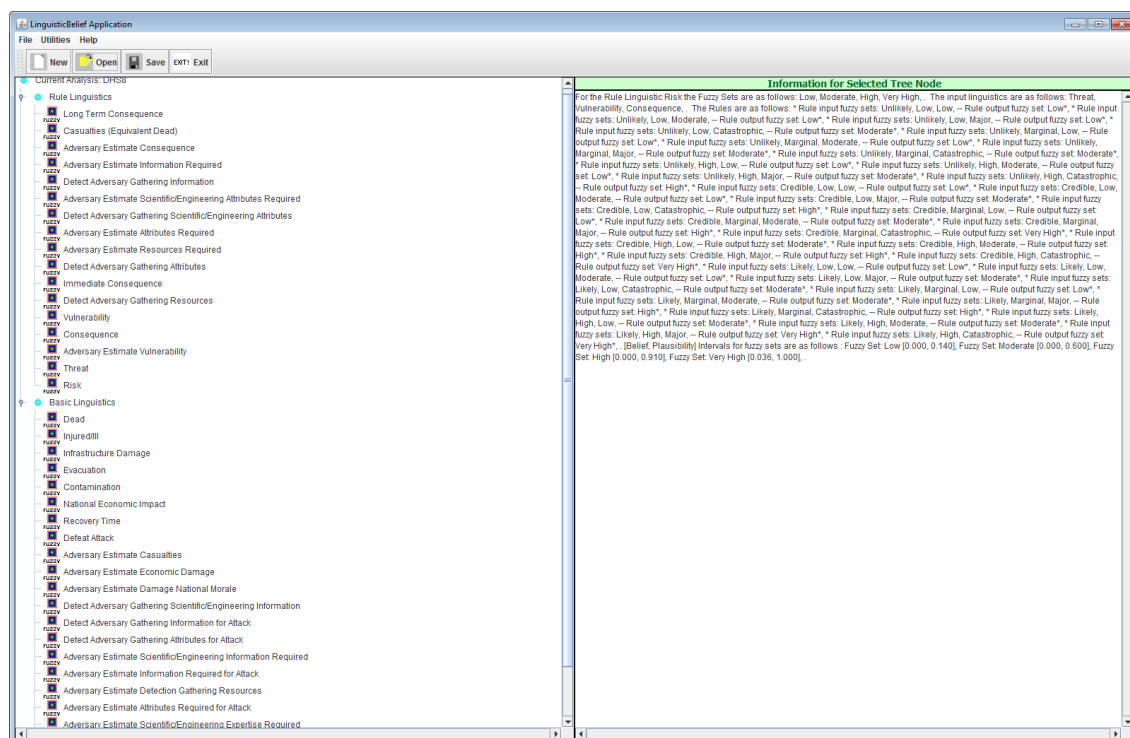


Figure A-8. A Model for Scenarios.

Appendix B. Example of Bayesian Update for Unknown Coin

Section 3 discussed the problem of not being able to adequately specify a prior probability distribution for the probability that an unknown coin will come up heads, $P(H)$, when tossed. With total ignorance our body of information is null, so we have no basis for selecting one prior over another.

Even with a poor prior, given sufficient objective information (many tosses of the coin) the updated posterior will be accurate. So, if we could gather sufficient information we could use the Bayesian approach to generate an accurate posterior probability distribution. The problem is that for the unknown coin (and many variables important for TNF) we cannot gather that information and we are forced to use a poor prior with limited information with which to update to an accurate posterior.

This appendix shows that with sufficient information (tosses) for the coin we can generate an accurate posterior even starting with an ill-informed prior. Here we use the beta distribution, β , for the distribution of the parameter $P(H)$; $p(H)$ is a specific value of the random variable $P(H)$.^{pp}

Assume as a prior $P(H)$ is a uniform probability distribution over $[0, 1]$ ($\beta[1,2]$). This prior is shown in figure B-1. The mean of the prior is 0.5 and the standard deviation is 0.29.²

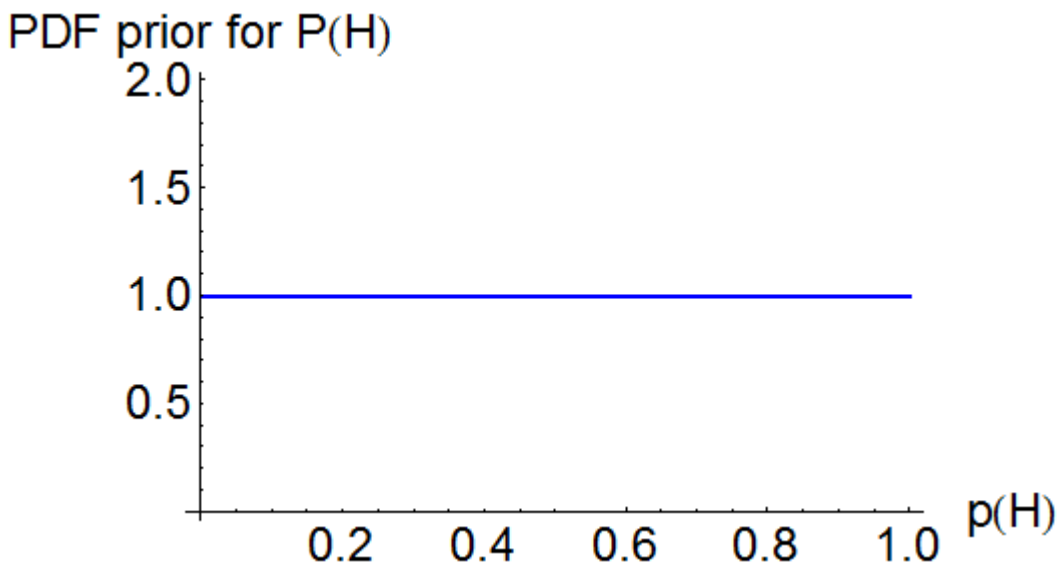


Figure B-1. Uniform Prior for $P(H)$.

If we toss the coin twice and the result is heads both times, our posterior is as shown in figure B-2.^{qq}

^{pp} As mentioned in section 3, a beta distribution is a very general probability distribution with range $[0,1]$; the PDF can have many different shapes depending on the values of its two parameters. It is widely used in Bayesian analyses as a conjugate prior for the probability parameter of the binomial distribution.² The PDF for the beta distribution is provided in textbooks that discuss probability distributions.²

The mean of this posterior is 0.75 and the standard deviation is 0.19.²

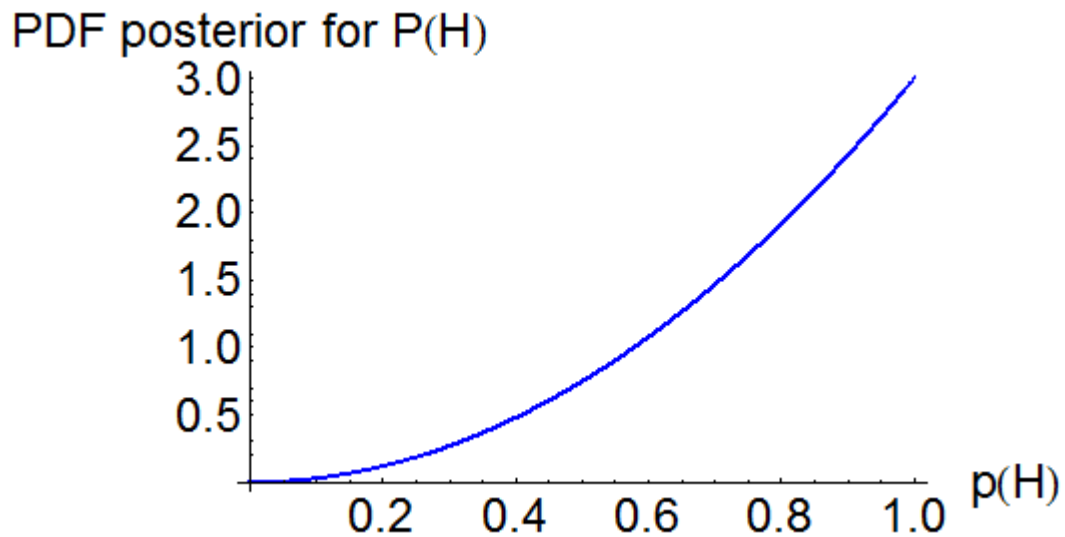


Figure B-2. Posterior for P(H) With Two Heads from Two Tosses.

As indicated in Figure B-2, the posterior of P(H) is biased towards heads, based on the objective data that both of two tosses turned up heads.

If we toss the coin 10,000 times and the result is heads for each toss, our posterior is as shown in figure B-3.^{rr} The mean of this posterior is 1.00 and the standard deviation is 0.00010.²

^{qq} The equations for generating the posterior of figure B-2 can be found in a reference.² Here, the posterior is $\beta[1 + 2, 2 + 2]$.

^{rr} Here the posterior is $\beta[1 + 10,000, 2 + 10,000]$.

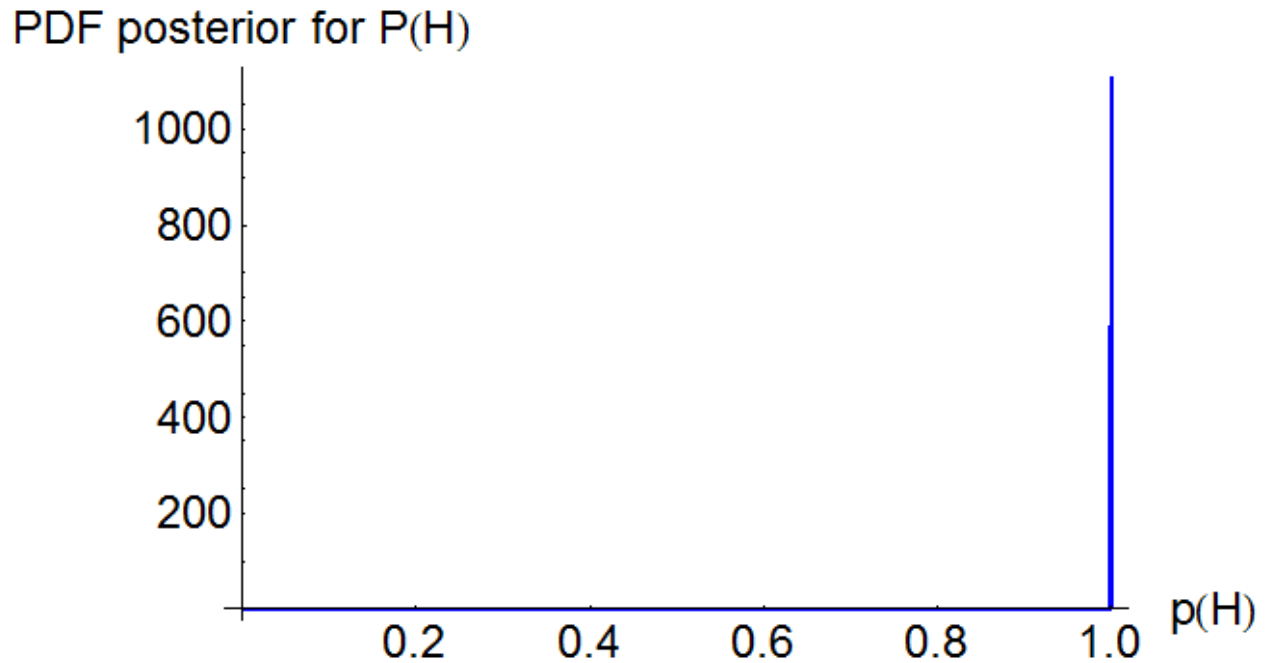


Figure B-3. Posterior for P(H) with 10,000 Heads from 10,000 Tosses.

As indicated in figure B-3, the posterior of P(H) indicates that the coin is two-headed, essentially with no uncertainty, based on the objective data of a great number of tosses, all turning up heads. With an infinite number of tosses, P(H) will be a delta function at the specific value p(H) of 1.0.

Therefore, even a poorly informed prior will generate an accurate posterior given sufficient objective data with which the prior can be updated.

Unfortunately, in TNF we have no ability to gather sufficient data to allow one to generate an accurate posterior distribution function.

Appendix C. Evidence and Calculation of Belief and Plausibility

By using examples, this appendix provides a brief tutorial on the assignment of evidence and how belief and plausibility are calculated from the focal elements that are assigned evidence.

C.1 Difference Between Assigned Evidence, Belief, and Plausibility

Evidence is the body of information available. Assigned evidence is evidence proportioned as focal elements. The value of a focal element is not a probability (except for the special case where all the focal elements are assigned to individual outcomes); limits for probability (belief, the lower limit, and plausibility, the upper limit) are calculated based on the focal elements.

Assigned evidence can be viewed as information weighted based on expert opinion, by the expert(s) consideration of the credibility of the information. A simple example illustrates this interpretation. This example uses a continuous variable, but the concept applies equally to a discrete variable.

Suppose on 31 December 2005, I am visiting Grandma and Grandpa, and after dinner we are discussing the age of a distant relative, Jack, who is not present.

Grandpa says “I think Jack was born sometime in 1980 or later.”

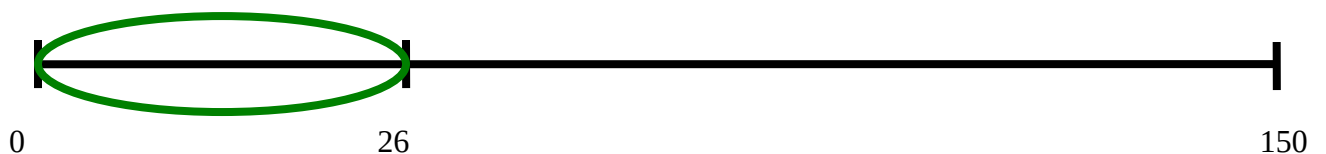
Grandma says “I think Jack is a teenager.”

Jack has definite age, but there is uncertainty in our body of information about his age.

The sample space is the set of all unique ages: $S = [0, 150]$ years.^{ss}

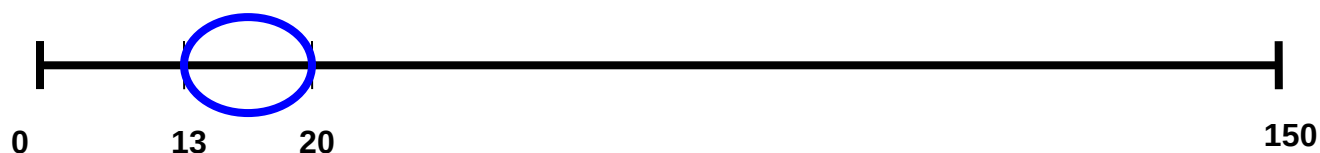
Our *evidence* has two parts: what Grandpa says and what Grandma says.

Evidence part 1: What Grandpa says. This is evidence for Jack’s age being somewhere *exactly* in $(0, 26)$.



^{ss} $[a, b]$ contains all values between a and b including a and b. $[a, b)$ contains all values between a and b, including a but not including b.

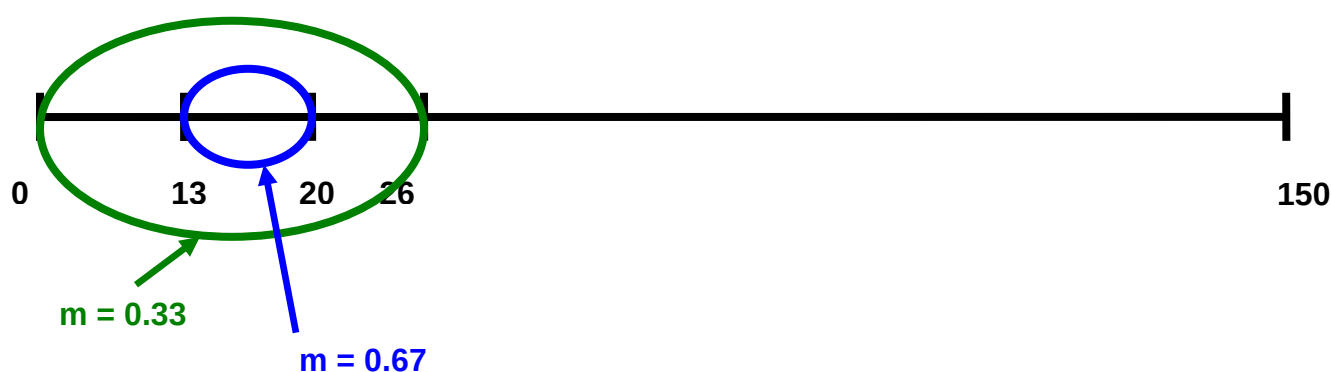
Evidence part 2: What Grandma says. This is evidence for Jack's age being somewhere *exactly* in $[13, 20]$.



The body of information is solely what Grandma and Grandpa say about Jack's age. Here, I am acting as an expert assessing the credibility of the body of information to subjectively assign evidence.

I know that Grandma has a better memory than Grandpa about relatives, so I decide to weight Grandma's evidence twice as much as Grandpa's evidence: two-thirds for Grandma and one-third for Grandpa.

Our focal elements are as follows: m is the value of a focal element.



Note that even though $(0, 26)$ contains $[13, 20]$ the value of the focal element 0.33 for $(0, 26)$ is less than the value of the focal element 0.67 for $[13, 20]$, because the value for a focal element assigned to an interval is the evidence of being *exactly* in that interval and *not* localized within any subinterval.

Since $(0, 26)$ contains $[13, 20]$, the belief for $(0, 26)$ will be greater than or equal to than the belief for $[13, 20]$. The belief for an interval is the sum of all focal elements *in that interval or any other interval within that interval*.

Since $(0, 26)$ contains $[13, 20]$, the plausibility for $(0, 26)$ will be greater than or equal to than the plausibility for $[13, 20]$. The plausibility for an interval is the sum of all focal elements *in any interval that overlaps that interval (any interval not disjoint with that interval)*. The plausibility is always greater than or equal to the belief, since it includes all the assigned evidence for belief plus the additional assigned evidence from the overlapping intervals.

For Jack, the belief for $[13, 20]$ is 0.67 and the plausibility for $[13, 20]$ is 1.0. The belief for $(0, 26)$ is 1.0 and the plausibility for $(0, 26)$ is 1.0. Based on the evidence, we conclude that Jack is not 26 years old or older, and the probability that Jack is a teenager is somewhere in the belief/plausibility interval 0.67 to 1.0.

A given expert can provide more than one piece of evidence. That is, Grandma could have provided **both** pieces of evidence. For example, Grandma could say my best recollection is that Jack is a teenager, but I think I remember Aunt Maude telling me that he is not yet 26. So I will assign evidence of two-thirds to teenager and one-third to not being yet 26.

C.2. Results as a Belief to Plausibility Interval

The belief and plausibility for an event of interest are calculated from the focal elements as follows. Belief is the sum of all evidence contained within the event of interest, and plausibility is the sum of all evidence that overlaps to any extent the event of interest. To be more precise, the belief and plausibility for any subset A is calculated as follows, where $m(B)$ is a focal element for subset B:

$$Bel(A) = \sum_{B|B \subseteq A} m(B)$$

$$Pl(A) = \sum_{B|A \cap B \neq \emptyset} m(B)$$

That is, all focal elements contained within A contribute to the belief for A, and all focal elements that overlap A (to any extent) contribute to the plausibility of A. (Focal elements contained with A also overall A.)

Consider the example in section 6 for the variable *presence of isotope ABC* with fuzzy sets {yes, no}. If the radiochemistry is such that we know with certainty the answer is yes, we assign probability of 1.0 to yes, and zero probability to no, resulting in the probability distribution of figure C-1.

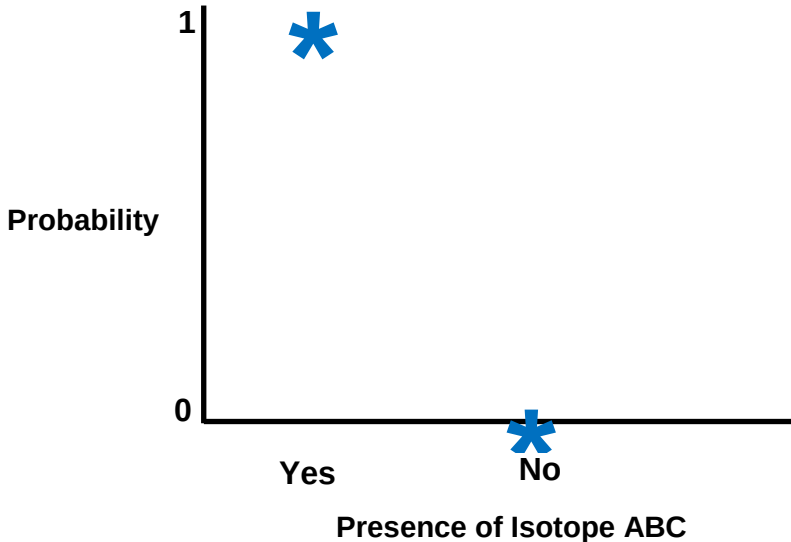


Figure C-1. Probability Distribution for No Uncertainty.

For this case with no uncertainty, this results from all evidence assigned to one focal element containing one fuzzy set, yes, as indicated in figure C-2.



Figure C-2. Focal Elements for No Uncertainty.

If the radiochemistry is dominated by aleatory uncertainty, we may assign the following probabilities: 0.3 to no, and 0.7 to yes, resulting in the probability distribution of figure C-3.

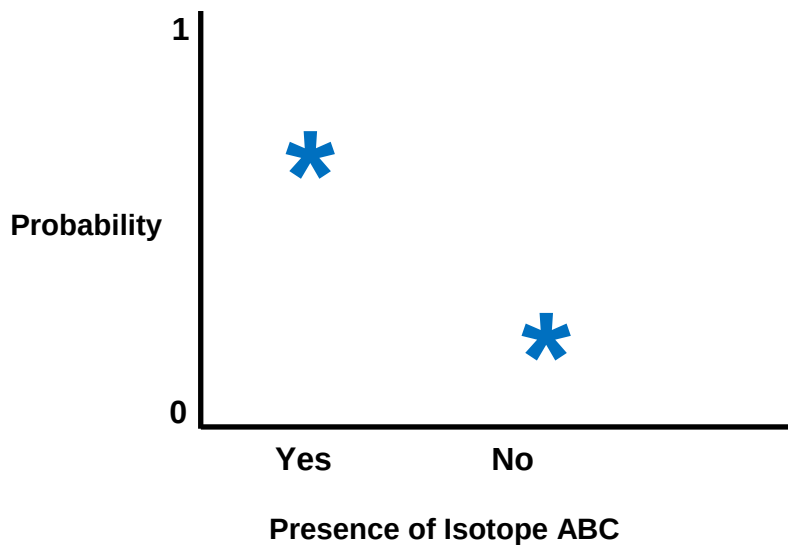


Figure C-3. Probability Distribution.

For this case using a probability distribution, this results from evidence assigned to focal elements each containing only one fuzzy set (singletons of the power set of the set of outcomes), as indicated on figure C-4. We have sufficient information to assign a probability to each outcome; that is, our evidence is sufficient to know the relative likelihood of each separate outcome. Each focal element contains only one outcome and both belief and plausibility are the same value for the focal element (outcome)—the probability. Probability is a degenerate case of belief/plausibility.



Figure C-4. Focal Elements for a Probability Distribution.

However, suppose the radiochemistry has significant epistemic uncertainty, as previously discussed in section 6. In this case we may assign the following evidence: 0.7 to yes, and 0.3 over {yes, no}. This is a hedge assignment; we think the answer is yes but hedge it to be within {yes, no}. This is *not* a probability assignment, because we cannot assign evidence to each specific outcome. The focal elements for this situation are shown in figure C-5.

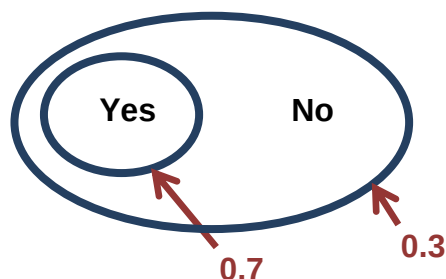


Figure C-5. Focal Elements Where Information Insufficient to Assign Probability Distribution.

The probability of no is somewhere in the belief/plausibility interval $[0, 0.3]$, and the probability of yes is somewhere in the belief to plausibility interval $[0.7, 1.0]$. The result for this case is shown in figure C-6 where probability is somewhere within a belief to plausibility interval.

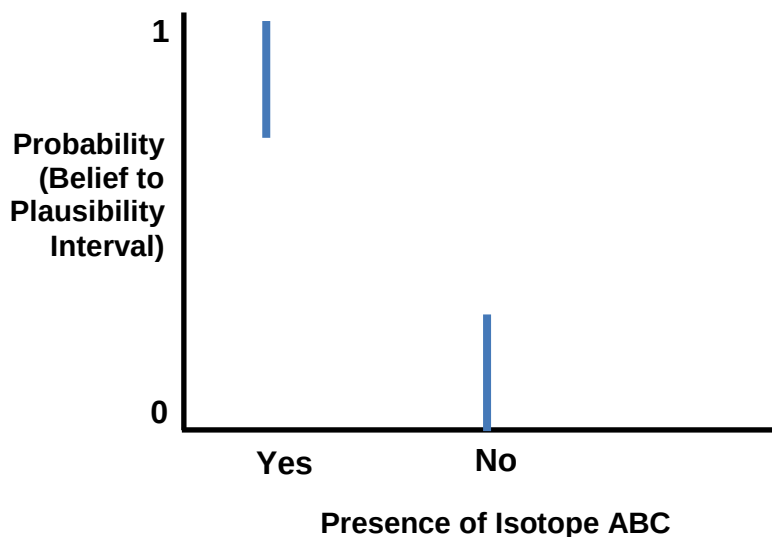


Figure C-6. Belief and Plausibility Distribution.

For any event the belief is no greater than the plausibility. The probability of the event is somewhere in the belief to plausibility interval. If the body of knowledge is sufficient to assign a probability distribution, the belief and plausibility are the same: the probability. Belief is a lower bound on probability and Plausibility is an upper bound on probability.

The one focal element for a case of total ignorance—the unknown coin discussed in section 3—is shown in figure C-7; the results are shown in figure C-8.

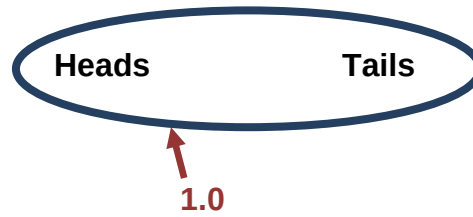


Figure C-7. Focal Elements for Unknown Coin: Total Ignorance

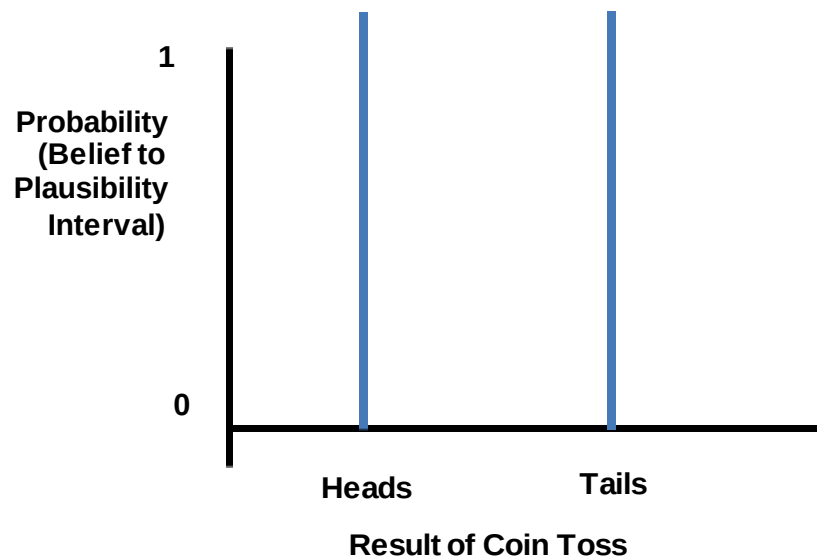


Figure C-8. Belief and Plausibility for Unknown Coin: Total Ignorance

Consider the focal elements of figure C-9 for the question “What is the category of the device?” that result from an evaluation of a complex model with many constituent variables with evidence for each basic variable provided by more than one expert.

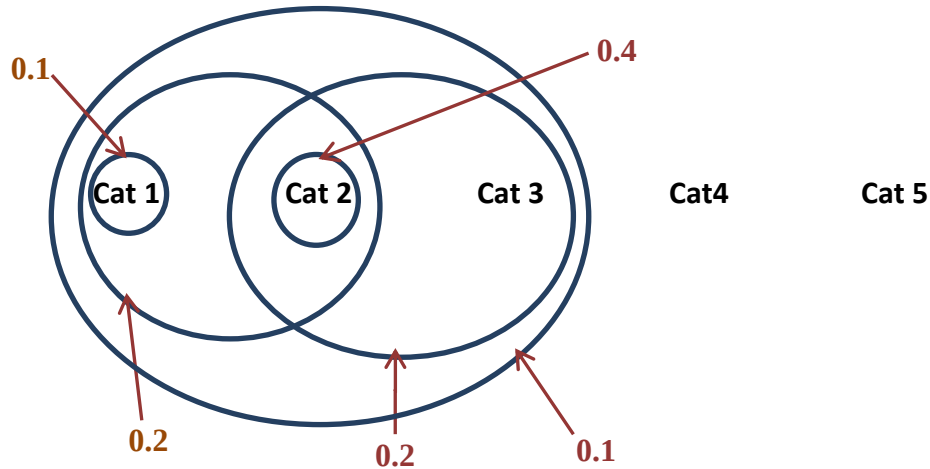


Figure C-9. Focal Elements for “What is the Category of the Device”?

The result for the focal elements of figure C-9 is given in figure C-10. The belief to plausibility interval for the probability for Cat 1 is [0.1, 0.4]. The belief to plausibility interval for the probability for Cat 2 is [0.4, 0.9]. The belief to plausibility interval for the probability for Cat 3 is [0, 0.3]. The belief to plausibility interval for the probability for Cat 4 is [0, 0]. The belief to plausibility interval for the probability for Cat 5 is [0, 0].

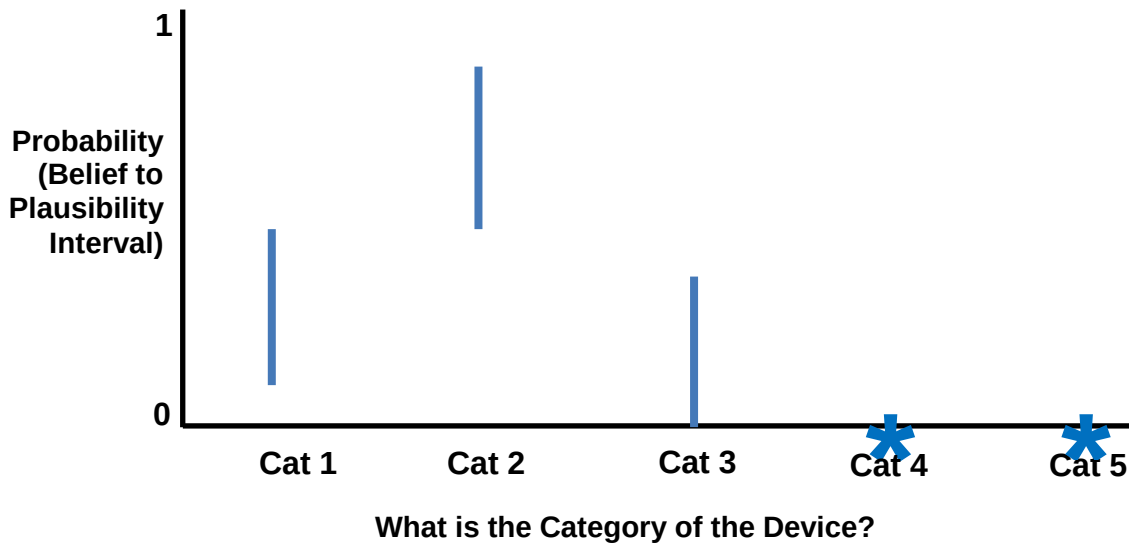


Figure C-10. Belief and Plausibility Distribution for “What is the Category of the Device”?

As a final example of the calculation of belief and plausibility from focal elements, assume the event of concern is *the device is not Cat 1*. The focal elements and this event of concern are shown in figure C-11.

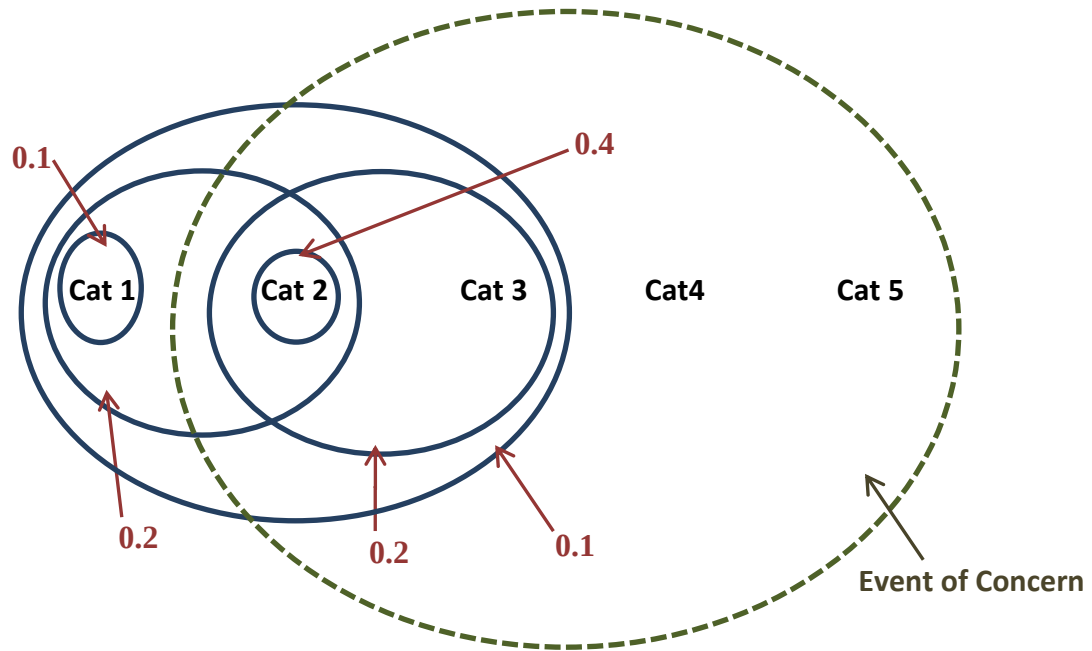


Figure C-11. Focal Elements and Event for *the Device is Not Cat 1*

The belief for *the device is not Cat 1* is 0.6 and the plausibility is 0.9. (Note: the plausibility that it could be Cat 1 is 0.4 so the belief that the event is NOT Cat 1 is $1 - \text{plausibility of Cat 1} = 1 - 0.4 = 0.6$.^{tt})

^{tt} For any event A, let A^c represent not A. $\text{Belief}(A) = 1 - \text{Plausibility}(A^c)$, and $\text{Plausibility}(A) = 1 - \text{Belief}(A^c)$.³

References

1. Helton, J.C., Johnson, J.D., and Oberkampf, W. L., "An Exploration of alternative approaches to the representation of uncertainty in model predictions", Reliability Engineering and System Safety 85 (2004) 39-71.
2. Bayesian Reliability Analysis, H.F. Martz and R. A. Waller, Wiley, 1982.
3. Fuzzy Sets and Fuzzy Logic, Theory and Applications, G.R. Klir and B. Yuan, Prentice Hall PTR, 1995.
4. "Guide To The Expression Of Uncertainty In Measurement", ISO, Geneva (1993), (ISBN 92-67-10188-9).
5. Comments from Bernd Kahn, Georgia Tech, e-mail to Charles Craft, Sandia, September 3, 2013.
6. "LinguisticBelief[®]: A Java Application for Linguistic Evaluation Using Belief, Fuzzy Sets, and Approximate Reasoning", J. L. Darby, SAND2007-1299, March, 2007.
7. Mathematica software version 9.0, Wolfram Research, Inc.
8. Probability and Statistics for the Engineering, Computing, and Physical Sciences, E.R. Dougherty, Prentice-Hall, 1990.
9. "Quantifying Uncertainty in Analytical Measurement", EURACHEM / CITAC Guide, Second Edition, QUAM:2000.P1.
10. "Sample Sizes for Confidence Limits For Reliability", J. L. Darby, SAND2110-0550, Feb, 2010.
11. "Evaluation of Risk from Acts of Terrorism: The Adversary/Defender Model using Belief and Fuzzy Sets", J.L. Darby, SAND2006-5777, September, 2006.
12. A Mathematical Theory of Evidence, G. Shafer, Princeton University Press, 1976.

Distribution

Doris Bruner, AFTAC/DO
Julie Gostic, AFTAC/TMA

Bill Ulicny, DHS/DNDO/NTNFC
Bill Daitch, DHS/DNDO/NTNFC
Frank Wong, DHS/DNDO/NTNFC

Jill Franks, DTRA/J9TNF
Kevin Mueller, DTRA/J9TNF

Jim Blankenship, FBI

Phil Cole, NSTech/NCNS

Bernd Kahn, Georgia Tech

Kevin Carney, INL

Lawrence Livermore National Laboratory

Allen Ross
Bill Dunlop
Roger White
Chad Giacamozi
Nathan Wimer
Yves Dardenne

Los Alamos National Laboratory

Jim Cooley
Chuck Wilkerson
Mike MacInnes
Bob Rundberg
Hugh Selby
James Gattiker
Scott Vander Weil

Vince Jodoin, ORNL

John Wacker, PNNL

Sandia National Laboratories

MS 0405	Jared McLaughlin, 0433
MS 0405	Todd Hinnerichs, 0435
MS 0405	John Darby, 0435
MS 0405	David Lartonoix, 0241
MS 0421	Carla Ulibarri, 0249
MS 0492	Jeffrey Brewer, 0411
MS 0748	Jon Helton, 1341
MS 0757	Gregory Wyss, 6612
MS 0829	Edward Thomas, 0415
MS 1218	Betty Biringer, 5944
MS 1219	Charles Craft, 5943
MS 1219	Sharif Heger , 5944
MS 1219	Sarah Soisson, 5943
MS 0899	Technical Library, 9536 (electronic copy)

