



Privacy Risk Assessment Framework for Information System [†]

Raisa Islam ^{*} and Dongwan Shin

Department of Computer Science, New Mexico Institute of Mining and Technology, Socorro, NM 87801, USA

^{*} Correspondence: raisa@cs.nmt.edu

[†] Extended from a non-archival presentation titled: Privacy Risk Assessment Framework for Information System. Presented at the 31st Australasian Conference on Information Security and Privacy, Perth, WA, Australia, 6–9 July 2026.

How to Cite: Islam, R.; Shin, D. Privacy Risk Assessment Framework for Information System. *Pragmatic Cybersecurity* **2026**, *1*(3), 15. <https://doi.org/10.53941/pc.2026.100015>

Received: 21 July 2026

Revised: 6 August 2026

Accepted: 7 August 2026

Published: 1 September 2026

Abstract: Privacy risk assessment is difficult due to uncertainty, incomplete information, and the evolving nature of modern information systems. Existing methods often rely on precise probabilities or stochastic models that do not adequately capture epistemic uncertainty or conflicting evidence. This paper presents a structured Privacy Risk Assessment framework based on the Evidential Reasoning Model and Dempster-Shafer belief theory. The framework decomposes a risk object into hierarchical components, assertions, and evidence, enabling extensible and fine-grained risk computation. Dempster-Shafer theory aggregates heterogeneous evidence while explicitly representing belief, disbelief, and uncertainty, and the model supports incremental updates without full recomputation. A case study using a smart grid environment implementing the Green Button Initiative demonstrates that the approach provides systematic, transparent, and computationally efficient privacy risk evaluation in complex data-sharing systems.

Keywords: privacy; risk assessment; uncertainty; evidential reasoning model; Dempster-Shafer theory; smart grid

1. Introduction

Privacy Risk Assessment (PRA) is a relatively new concept to address privacy threats. Privacy can be defined as the right of data owners to control access to and disclosure of their sensitive information to other parties. Risk can be defined in terms of the likelihood that a threat agent will exploit a vulnerability associated with user data and the resulting direct or indirect harmful impact on the data owner. PRA of any information is a measure to evaluate the possibility of unwanted events occurring where the privacy of any user can be harmed. The risk associated with information has increased in recent years because of the growing number of threats that have emerged daily. Although a series of regulations are mandated by the government, this is not enough to keep up with ever-evolving technologies. PRA can lead to an early identification and prevention of possible privacy breaches [1]. Since an adversary can exploit a user's privacy in any stage of data sharing, PRA requires a systematic analysis of all possible events.

The amount of data collected by contemporary IT systems has increased rapidly, making such systems attractive targets for adversaries [2,3]. Although regulatory frameworks such as the Privacy Act, HIPAA, COPPA, FISMA, and CCPA aim to protect personal data [4–7], regulatory compliance alone is insufficient to address evolving technological threats. Consequently, systematic privacy risk assessment is required to identify potential privacy breaches before data sharing occurs [1]. Since privacy violations may arise at any stage of data collection, processing, or dissemination, PRA must analyze the entire lifecycle of information sharing.

A central challenge in PRA is uncertainty. In risk analysis, uncertainty appears in two forms. Aleatory uncertainty describes randomness in observable events such as system failures or attack frequency. Epistemic uncertainty arises from incomplete knowledge, conflicting evidence, or limited observability of adversary behavior. Privacy risk is dominated by epistemic uncertainty because attacker intent, auxiliary information, and real world data linkage cannot be reliably measured. As various uncertainties and ambiguities exist in the real-world decision-making problems, using precise numbers without any evidences are not competent to describe the assessment value [8,9].



Unlike traditional security risk assessment, privacy disclosure does not depend solely on observable event frequency. It depends on what an adversary can infer using unknown auxiliary information. Consequently, privacy risk is not a probability estimation problem; it is an evidence aggregation problem under adversarial knowledge uncertainty. These limitations motivate the following research questions:

- **RQ1:** Do existing privacy risk assessment methods adequately address uncertainty and complexity in modern interconnected IT systems?
- **RQ2:** How can uncertainty in privacy risk assessment be represented and resolved more effectively?
- **RQ3:** How can a privacy risk assessment framework remain adaptable to evolving technologies and data sharing environments?

To address these questions, a privacy risk assessment framework is developed using the evidential reasoning model (ERM) and Dempster-Shafer (DS) belief theory. ERM supports multi-attribute reasoning by systematically combining heterogeneous evidence and expert judgment, while DS theory provides a formal mechanism for aggregating uncertain and potentially conflicting evidence without requiring precise probability estimates. This makes the approach suitable for privacy settings where attacker behavior, auxiliary information, and data linkage cannot be directly measured.

The framework models privacy exposure by decomposing a risk object into hierarchical components and associated assertions. Evidence collected for each assertion is evaluated and combined using DS belief aggregation to compute the overall risk. The structure enables fine-grained analysis at the information level rather than only at the system level, while allowing new attributes or evidence to be incorporated without redesigning the model.

The primary contributions are:

- We evaluate privacy risk at the level of individual information attributes rather than only at the system level, enabling precise assessment of data-sharing decisions.
- We develop a privacy risk assessment model based on the Evidential Reasoning Model and Dempster-Shafer theory to aggregate uncertain and potentially conflicting evidence.
- We introduce a hierarchical risk decomposition that supports incremental updates and extensibility as system characteristics evolve.
- We demonstrate practicality through a smart grid case study (Green Button Initiative), showing effective privacy risk evaluation in real-world data-sharing environments.

The remainder of this paper is organized as follows. In Section 2, we discussed related work and background. The design of our proposed Privacy Risk Assessment framework is presented in Section 3, followed by complexity analysis in Section 4. In Section 5, we demonstrated the feasibility of proposed framework by extending the framework for smart grid systems. Finally, in Section 6, we summarize our framework and outline future research.

2. Related Work and Background

2.1. Existing Information Risk Assessment Methods

Over the past decades, numerous information risk assessment methods have been proposed from both security and privacy perspectives. These approaches can generally be categorized into three major groups based on how risk is represented and measured: qualitative, quantitative, and semi-quantitative methods.

Qualitative risk assessment relies on descriptive or subjective evaluation of risks. In such approaches, risks are typically ranked based on perceived severity and likelihood using ordinal scales such as high, medium, and low. Techniques such as expert judgment and the Delphi method are commonly employed to capture domain knowledge and collective expertise. Although qualitative methods are relatively simple to implement and require limited quantitative data, their subjective nature may lead to inconsistent evaluations across different assessors.

Quantitative risk assessment, in contrast, attempts to represent risk numerically. These approaches typically compute risk as a function of two primary components: the probability that a harmful event will occur and the magnitude of the potential loss associated with that event. Several quantitative models have been proposed for information security and privacy analysis, including the Factor Analysis of Information Risk (FAIR) model, statistical approaches, and fuzzy mathematics-based models [10–16]. Quantitative approaches provide more precise numerical results and facilitate formal analysis; however, they often require accurate probability estimates and detailed datasets, which are not always available in practice.

Semi-quantitative risk assessment represents a hybrid approach that attempts to bridge the gap between qualitative and quantitative methods [17]. In semi-quantitative frameworks, qualitative assessments are translated into numerical scores or ratings to allow comparison and prioritization of risks. This approach improves consistency relative to purely qualitative evaluations while avoiding the extensive data requirements of fully quantitative models.

Several specific frameworks have been developed within these categories. Levy proposed a data center risk assessment model in which risk is expressed as a time-dependent function of event risk levels and weighted event importance [18]. In this approach, risk is decomposed into probability and impact components, each assigned semi-quantitative values. While the framework provides a structured way to evaluate operational risks, it focuses on event-level analysis and does not provide detailed decomposition of risks associated with specific information attributes.

Freund and Jones introduced the Factor Analysis of Information Risk (FAIR) model in 2014 [10]. FAIR defines an ontology of risk factors and their relationships and assigns numerical values to these factors to estimate overall risk. Cronk later extended FAIR to address privacy risk by refining the decomposition of risk factors and incorporating privacy-related variables [11]. In this extended framework, privacy risk is evaluated based on the potential impact on the data owner and the frequency of successful adversarial actions. Similarly, Sion et al. proposed a data-subject-aware privacy risk assessment model that further extends FAIR by incorporating privacy-specific factors and data subject considerations [19]. Their approach integrates risk inputs into Data Flow Diagram (DFD) models and uses Program Evaluation and Review Technique (PERT) distributions together with Monte Carlo simulations to estimate risk values under uncertainty.

Despite these improvements, several factors involved in privacy risk estimation remain difficult to quantify, including adversary motivation, threat capability, and probability of adversarial actions. These models therefore rely heavily on probabilistic estimates. While Monte Carlo simulation is effective for modeling stochastic uncertainty, it is less suited for representing epistemic uncertainty arising from incomplete knowledge. Moreover, Monte Carlo approaches are not well equipped to aggregate evidence originating from heterogeneous or potentially conflicting sources, since they assume that all uncertainty can be represented within a single probabilistic framework [20].

Another widely used approach for handling uncertainty in risk assessment is the fuzzy inference system. Unlike classical binary logic, fuzzy logic allows elements to have partial membership in multiple sets simultaneously, enabling reasoning with partial truths and ambiguous information. PRIAM [12] is a fuzzy–probabilistic privacy risk assessment framework that models relationships among harms, feared events, and privacy weaknesses to estimate the likelihood of different risks and determine whether they exceed acceptable thresholds. PRIAM has also been applied in healthcare privacy risk analysis [21]. Additional fuzzy logic-based risk assessment models have been proposed in [13–15]. However, fuzzy inference systems often require large numbers of expert-defined rules, and the number of rules grows exponentially with the number of risk factors. Furthermore, the absence of reliable expert knowledge can significantly limit the effectiveness of these systems [16].

Attack-tree-based approaches have also been used for privacy risk assessment. In [22], the authors proposed an attack-tree-based mechanism for evaluating privacy risk, which was later extended with ontology-based reasoning to support automated risk assessment in IoT ecosystems [23]. In these models, risk is estimated by considering attack attributes, adversarial capabilities, and computational resources available on target devices. Individual attack probabilities are combined using logical AND/OR gate structures. While these approaches effectively model device vulnerabilities and network-layer threats, they do not consider the potential availability of auxiliary public information that could be combined with shared data to compromise user privacy.

Grey system theory has also been applied to risk analysis in environments characterized by incomplete information. Wu and Zhao proposed a model that integrates grey system theory with the Decision Making Trial and Evaluation Laboratory (DEMATEL) method and the Analytic Network Process (ANP) [24]. In this framework, DEMATEL determines causal relationships among risk factors, while ANP calculates the relative importance of evaluation indices. The authors argue that this combination can reduce uncertainty in expert evaluations.

Bayesian network models have likewise been applied to information risk assessment. Wang et al. [25] proposed a Bayesian network-based model that integrates probabilistic reasoning with expert knowledge to estimate value-at-risk. The authors also developed a system architecture incorporating a security knowledge base to support automated risk evaluation. Similarly, the cloud model, a fuzzy–probabilistic semi-quantitative approach, was proposed in [17]. The cloud model represents qualitative concepts using three parameters: expectation, entropy, and excess entropy, which together capture the uncertainty associated with risk factors. Although these approaches provide useful mechanisms for handling uncertainty, they primarily focus on information security risk rather than privacy risk.

Other domain-specific risk assessment methods have also been proposed in [26–28]. However, these approaches are typically designed for specialized environments such as social media platforms and therefore lack general applicability across diverse domains. In addition, a framework combining ERM with DS theory was introduced in [29] for system-level risk analysis. While this approach incorporates evidential reasoning for uncertainty modeling, it focuses primarily on assessing the overall system risk rather than the risk associated with individual information attributes. Consequently, the model lacks the granularity required for fine-grained privacy risk assessment.

To better illustrate the methodological differences among existing privacy risk assessment approaches, Table 1

summarizes representative techniques discussed in the literature. The comparison highlights key characteristics relevant to privacy risk modeling, including the representation of uncertainty, mechanisms for evidence aggregation, the level of granularity at which risk is evaluated, and the ability to handle conflicting information. This comparison reveals several limitations of traditional approaches and motivates the need for a framework capable of systematically integrating heterogeneous evidence while representing uncertainty in a more expressive manner.

2.2. NIST Risk Management Framework (RMF)

The National Institute of Standards and Technology (NIST) RMF was developed to help organizations manage both security and privacy risk by satisfying the requirements derived from laws, regulations, and policies described in the Federal Information Security Modernization Act of 2014 (FISMA), the Privacy Act of 1974, the Office of Management and Budget (OMB) policies, and Federal Information Processing Standards, among others. The RMF provides a dynamic and flexible approach to effectively manage risk for evolving systems with complex and sophisticated threats and vulnerabilities. The objective was to design a technology-neutral methodology that can be applied across different types of information systems. The RMF consists of the following steps:

- Prepare—Initiates the RMF execution within an information system by establishing a context and priorities for managing risk.
- Categorize—Includes identification and classification of the attributes processed, stored, and transmitted in the information system.
- Select—Based on the categorization of attributes, an initial set of control mechanisms is selected to reduce risk to an acceptable level based on an assessment.
- Implement—Apply the selected risk control methods and document how they are deployed within the information system.
- Assess—Evaluate the effectiveness of the implemented mechanism to ensure that it is producing the desired outcomes.
- Authorize—Makes a risk-based decision based on the assessment results regarding whether to authorize the system's operation.
- Monitor—Continually monitor the system for changes and reassess the effectiveness of control mechanisms; modify and update the system when necessary.

While the NIST RMF provides a comprehensive structure for managing risk, it does not mention anything about how to achieve the recommended steps. The IT industry is vast and typically involves numerous interconnected entities. When an entire system is considered as a single entity, the design and implementation of the RMF can become exceedingly challenging. Additionally, it lacks the necessary granularity to address the complexities of the modern IT sector. The rapid evolution of the technology landscape further complicates this issue. The NIST RMF does not readily accommodate the inclusion of new components and technologies, making the framework appear rigid and slow to adapt. This inflexibility can result in significant security and privacy gaps, as the controls specified by the RMF may not adequately address emerging and evolving risks.

2.3. DS Theory of Belief

DS theory, introduced by Dempster [30] and Shafer [31], is an evidence-based reasoning framework for decision making under incomplete or conflicting information. Unlike Bayesian probability, which requires precise likelihood estimates, DS theory represents uncertainty explicitly by allowing belief to be assigned not only to single hypotheses but also to sets of hypotheses. This property makes DS theory suitable for privacy risk assessment where adversary knowledge, auxiliary information, and data linkage cannot be reliably quantified.

Let the frame of discernment be a set of mutually exclusive hypotheses

$$\Theta = \{F_1, F_2, \dots, F_N\}.$$

A basic belief assignment (mass function) distributes belief over subsets of Θ :

$$m : 2^\Theta \rightarrow [0, 1], \quad \sum_{A \subseteq \Theta} m(A) = 1, \quad m(\emptyset) = 0.$$

The value $m(A)$ represents the degree of support provided by the available evidence to proposition A .

Table 1. Comparison of privacy risk assessment approaches.

Approach	Uncertainty Representation	Evidence Aggregation Method	Attribute-Level Granularity	Conflict Handling	Extensibility
Risk Matrix (Likelihood $\times$ Impact)	Deterministic qualitative or ordinal scoring	Simple multiplication or scoring matrix	Low. Typically evaluated at system or scenario level	None	Limited. Hard to extend beyond predefined categories
Bayesian Probabilistic Models	Precise probability distributions representing likelihood of events	Bayesian inference and conditional probability updates	Moderate. Can model variables but requires predefined probabilistic relationships	Limited. Assumes probabilistic consistency across evidence	Moderate. Extension requires redesign of probabilistic structure
Monte Carlo Simulation	Random uncertainty represented through sampling distributions	Large-scale stochastic simulation and statistical estimation	Moderate. Depends on model design and simulation variables	None. Conflicting inputs appear as variance rather than explicit conflict	Moderate. New variables can be added but simulation complexity increases
Attack Graph/Threat Modeling	Scenario-driven risk reasoning based on attack paths	Graph traversal, scoring, or reachability analysis	Low to moderate. Focuses on attack scenarios rather than individual attributes	Limited. Conflicts between evidences are not explicitly modeled	Moderate. Graph structures can be expanded but require redesign
Differential Privacy Risk Analysis	Formal mathematical privacy guarantees through privacy budgets	Analytical privacy bounds and sensitivity analysis	High for statistical query outputs but limited to aggregate data scenarios	Not applicable	Limited to statistical database query environments
ERM with DS Theory (Proposed)	Belief and plausibility intervals representing epistemic uncertainty	Evidence combination using Dempster's rule of combination	High. Supports attribute-level privacy risk evaluation	Explicit modeling of conflicting evidence through belief mass redistribution	High. Additional evidence sources can be incorporated without redesigning the model

From the mass function, two measures are derived. The belief function provides a lower bound of support:

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

while the plausibility function provides an upper bound:

$$Pl(A) = \sum_{B \cap A \neq \emptyset} m(B) = 1 - Bel(\neg A).$$

The interval $[Bel(A), Pl(A)]$ therefore captures the uncertainty associated with a proposition.

When multiple independent evidence sources exist, DS theory combines them using Dempster's rule of combination. For two belief assignments m_1 and m_2 , the combined belief $m = m_1 \oplus m_2$ is

$$m(A) = \frac{\sum_{B_1 \cap B_2 = A} m_1(B_1)m_2(B_2)}{1 - K},$$

where the conflict coefficient

$$K = \sum_{B_1 \cap B_2 = \emptyset} m_1(B_1)m_2(B_2)$$

measures disagreement between the evidence sources. The normalization factor $(1 - K)^{-1}$ redistributes conflicting belief to consistent hypotheses.

In this work, DS theory is used to aggregate heterogeneous evidence associated with privacy assertions. Its ability to represent belief, disbelief, and residual uncertainty allows privacy risk to be evaluated without assuming precise probabilities.

3. PRA Framework

Digital services and online platforms routinely collect large volumes of personal information, including contact data, financial records, usage patterns, and behavioral traces [3]. The protection of such data requires a systematic method for evaluating privacy exposure before information is shared or processed. We therefore propose a PRA framework that quantifies both the likelihood and impact of privacy disclosure. The information under evaluation is modeled as a *risk object*, which is decomposed into multiple components to enable granular analysis of individual risk contributors and their influence on the overall risk value.

The framework applies ERM to represent assessments derived from heterogeneous evidence sources. Each component is associated with an evidentiary assessment describing whether privacy requirements are satisfied. DS theory provides the mathematical mechanism for aggregating evidence and resolving conflicting observations through belief normalization. Unlike probability-only approaches, DS theory explicitly represents belief, disbelief, and uncertainty, making it suitable for privacy scenarios characterized by incomplete or ambiguous information.

The DS rule of combination further supports hierarchical risk evaluation. The operator $\oplus$ is commutative and associative, allowing evidence to be aggregated across different components and levels of the structure without dependence on ordering. Idempotence prevents duplicated evidence from artificially increasing belief, and continuity ensures that small changes in evidence produce proportionally small changes in belief assignment [32]. These properties enable recursive aggregation and incremental updates within the proposed framework.

The PRA process consists of three phases: (1) specification of the evidential structure; (2) assignment of evidence strength; and (3) computation of overall risk.

During the specification phase, the framework defines the components and their relationships. A component represents a distinct functional part of the risk object and may contain subcomponents or assertions. Each component is assigned a relative weight based on expert judgment. Assertions describe potential privacy violations associated with the component, and each assertion is evaluated using relevant data attributes. Because different attributes expose privacy at different levels, data properties are identified to support precise risk estimation.

The following data properties are considered in the assessment:

- Uniqueness: data that alone can identify an individual (e.g., SSN).
- Quasi-identifier: attributes that are non-identifying individually but may identify a person when combined (e.g., birthdate, ZIP code, gender).
- Granularity: the level of detail in the data; fine-grained measurements (per-second usage) reveal more information than aggregated values (monthly usage).
- Severity: the potential harm caused by disclosure or misuse of the data (e.g., medical records imply high impact).

- Volume: the amount of available data; larger datasets increase inference capability.
- Purity/format: data accuracy and consistency; higher purity indicates less noise, while perturbation reduces reliability.

The next phase involves gathering belief values that support each piece of evidence to evaluate the strength and reliability of the evidence. Domain experts play a central role in the belief assignment process. Their contributions include evaluating the relevance of the evidence, establishing its mapping to assertions, and assigning belief, disbelief, and uncertainty values based on domain knowledge. Their input serves as the primary source of high-confidence belief assignment, particularly in scenarios where domain-specific risk semantics and contextual nuances cannot be adequately captured by automated methods. However, the framework is not restricted to human-driven inputs. Mass values can also be derived from empirical or machine-generated evidence when available. For example, if a predictive system demonstrates an observed accuracy of 80%, this performance metric can be incorporated into the belief assignment as a measure of evidential support. Similarly, historical data or statistical analysis can be used to derive belief distributions.

This step is crucial to ensure that the risk assessment is founded on robust, credible evidence, which forms a solid foundation for accurate risk determination. Each piece of evidence may contribute to multiple assertions related to the same or different components within the system. Additionally, it is possible to assign a weight to each piece of evidence, reflecting its relative importance and reliability in the overall assessment.

The final phase focuses on synthesizing the belief values to determine the overall level of privacy risk. This involves aggregating the evidence, taking into account the belief values and assigned weights. The culmination of this phase is the development of a detailed risk profile, which provides decision-makers with the necessary insights. This comprehensive risk profile can be used for informed and effective decision-making.

3.1. Structural Overview

Figure 1 presents the multi-layered structure of our proposed PRA framework where the risk of each component is assessed in a granular manner, resulting in a detailed understanding of potential outcomes. This detailed assessment is achieved by decomposing the overall risk into smaller, manageable elements. Each of these elements can be analyzed separately, providing a comprehensive evaluation of the entire risk object. The "Risk Object", $\mathcal{R}$, signifies the primary entity or process being evaluated for risk. A risk object comprises at least one or more components, denoted by $\mathcal{C}$. Formally, we can represent the risk object as $\mathcal{R} = \{C_1, C_2, \dots, C_i, \dots, C_n\}$, where i ranges from 1 to n .

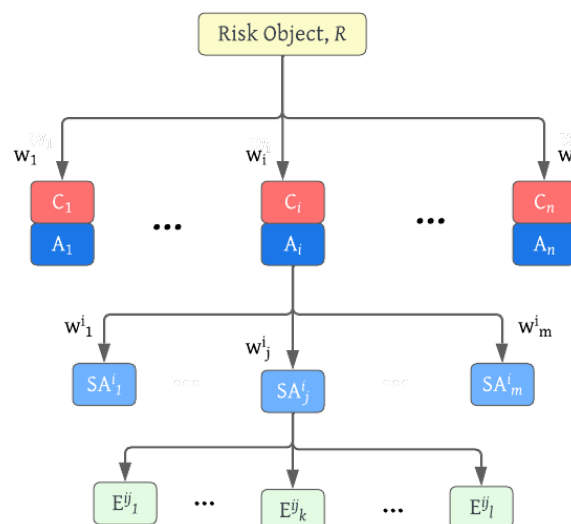


Figure 1. Structural overview of the proposed risk assessment diagram.

Each component C_i is associated with a main assertion, denoted by A_i , establishing a one-to-one relationship between components and main assertions. Mathematically, we can express this as $A_i = f(C_i)$, where f is a mapping function. Further, each assertion A_i may have multiple sub-assertions SA_j^i , where j ranges from 0 to m , indicating the number of sub-assertions for the i -th assertion. For simplicity, Figure 1 shows only a single level of sub-assertions.

The leaf nodes of each component represent evidence supporting the assertion, denoted by $\mathcal{E}$. It is possible for one evidence to support multiple assertions. We can denote the evidence as $\mathcal{E}_k^{ij}$, where k ranges from 1 to l , indicating the number of evidences for the j -th sub-assertion of the i -th assertion.

The hierarchical structure of the risk object can be expressed as:

$$\begin{aligned}\mathcal{R} &= \{\mathcal{C}_1, \mathcal{C}_2, \dots, \mathcal{C}_i, \dots, \mathcal{C}_n\} \\ \mathcal{C}_i &\rightarrow \mathcal{A}_i = \{\mathcal{SA}_1^i, \mathcal{SA}_2^i, \dots, \mathcal{SA}_j^i, \dots, \mathcal{SA}_m^i\} \\ \mathcal{SA}_j^i &= \{\mathcal{E}_1^{ij}, \mathcal{E}_2^{ij}, \dots, \mathcal{E}_k^{ij}, \dots, \mathcal{E}_l^{ij}\}\end{aligned}$$

The framework employs weights associated with each component which allows for a thorough evaluation by considering the relative importance of each component to the overall risk profile. Mathematically, the risk contribution of each component $\mathcal{C}_i$ to the risk object $\mathcal{R}$ can be represented as $w_i \cdot R(\mathcal{C}_i)$, where $R(\mathcal{C}_i)$ denotes the risk associated with component $\mathcal{C}_i$. Therefore, the total risk of the risk object $\mathcal{R}$ can be expressed as a weighted sum of the risks of its components:

$$R(\mathcal{R}) = \sum_{i=1}^n w_i \cdot R(\mathcal{C}_i)$$

3.2. Process Description

Figure 2 presents the process flow for information risk management in the PRA framework. Two entities are involved: the *data owner* and the *system*. The data owner is the individual or organization that legally controls the data and determines how it is collected, processed, and shared. The system collects, stores, and uses the information to provide services through applications such as web portals or mobile interfaces. Data owners supply information to obtain services.

Since privacy preferences vary across individuals and contexts, the data owner specifies an acceptable risk threshold through the system interface, which can be modified at any time. The system collects relevant evidence and expert assessments (trust values) and stores them in the database.

Risk is evaluated whenever new information is shared or a third party requests data. A risk object contains multiple components. If a component has not been evaluated, its risk is computed using Algorithm 1; otherwise, the stored value is reused and combined with other components according to their weights.

Algorithm 1 PRA Framework

```

class Component
  children: Array[] Component           ▷ List of sub-components
  count: int                           ▷ Number of children
  m: Array[] float                     ▷ m-values for  $p, q$  &  $\{p, q\}$ 

  procedure RISK_ESTIMATION
    if count = 0 then                               ▷ Evidence node
      return m
    end if

    for each child in children do
      child.m ← child.Risk_Estimation()
      m = m ⊕ child.m                               ▷ Apply DS rule of combination here
    end for

    return m                                         ▷ return overall risk value

  end procedure

end class

```

Algorithm 1 performs recursive risk estimation over a hierarchical structure of components. Leaf nodes (evidence) return their m-values, while higher-level nodes combine child results using the Dempster-Shafer rule of combination to obtain an aggregated risk value. This decomposition enables complex systems to be analyzed by aggregating simpler evidential elements.

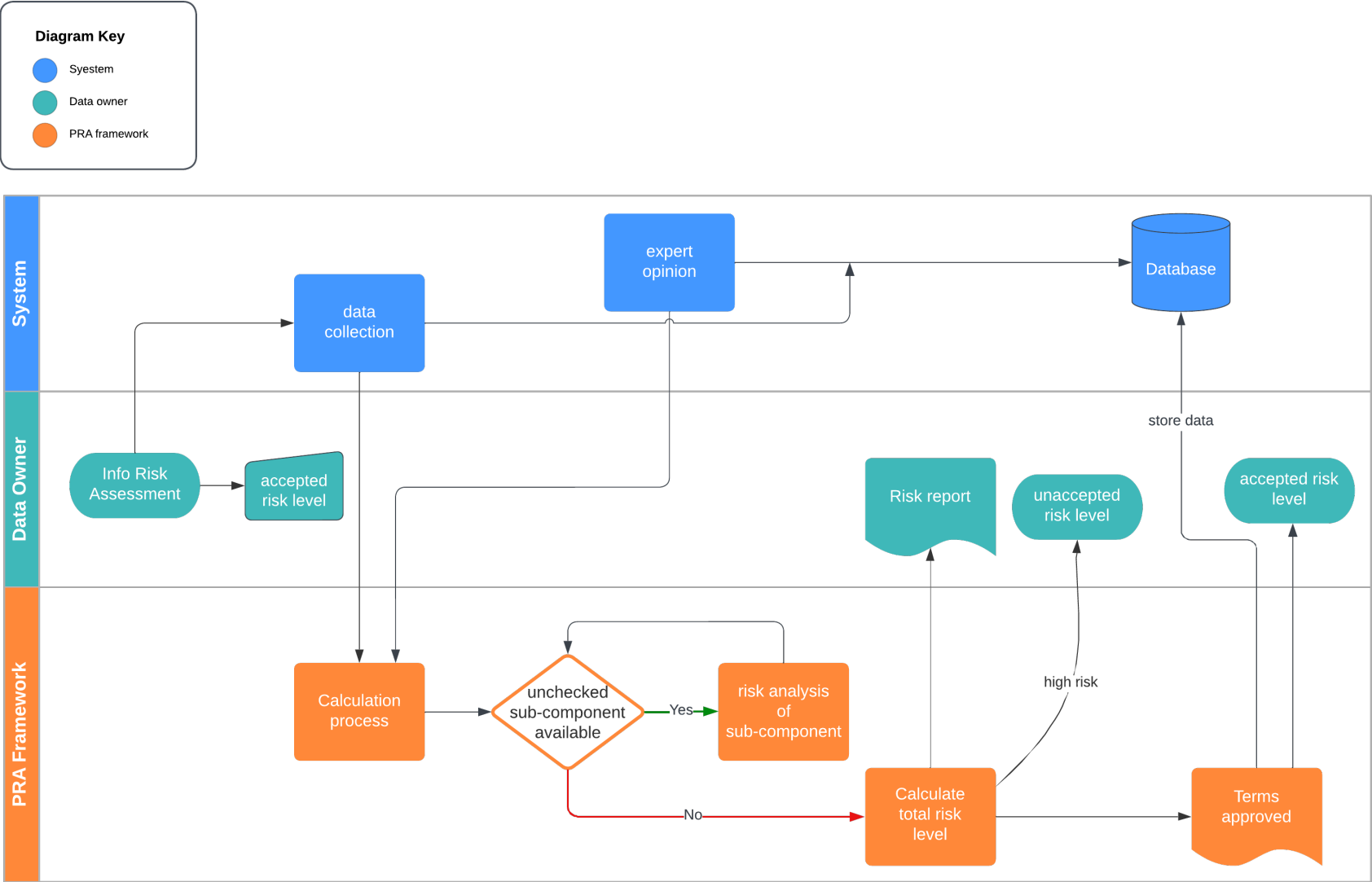


Figure 2. Process flow diagram.

3.3. Interpretation Mapping

The DS aggregation produces belief, disbelief, and uncertainty values rather than a single probability. Therefore a decision interpretation is required to translate evidential outcomes into operational privacy actions.

Let p denote the proposition that sharing a requested attribute is privacy-safe. After evidence aggregation, each request yields a belief interval $[Bel(p), Pl(p)]$, where $Bel(p)$ represents supported safety and $Pl(p)$ represents the maximum possible safety under remaining uncertainty.

The decision policy is defined as:

- If $Bel(p) \geq \tau_{allow}$, the request is approved.
- If $Pl(p) \leq \tau_{deny}$, the request is rejected.
- Otherwise, the request is uncertain and requires user notification or additional evidence.

Here τ_{allow} and τ_{deny} are user-defined risk tolerance thresholds. The uncertainty mass $m(\{p, q\})$ directly indicates lack of knowledge rather than risk, allowing the system to distinguish insufficient evidence from actual privacy danger.

This mapping converts evidential reasoning outputs into actionable privacy decisions while preserving uncertainty information, enabling risk-aware data sharing rather than probability-based approval.

4. Complexity Analysis

4.1. DS Rule of Combination

Let's assume, for a subset $A \subseteq \Theta$ with two independent items of evidence and respective belief assignment functions are represented by $m_1(A)$ and $m_2(A)$. Then the combined m-value using Dempster's rule of combination is computed from m_1 and m_2 by adding all products of the form $m_1(B)m_2(C)$ where B and C are selected from the subsets of Θ in all possible ways such that their intersection is A . Thus,

$$m(A) = \frac{\sum_{B \cap C = A} m_1(B)m_2(C)}{K}$$

$$K = 1 - \sum_{B \cap C = \phi} m_1(B)m_2(C)$$

Calculating these values involves iterating over all pairs of subsets. The associative and commutative properties ensure that the rule generates the same result regardless of the grouping and order of the evidence. If we have x elements in Θ , we need to evaluate intersections for 2^x possible combinations (size of the power set).

The computational complexity of this process depends on the number of pieces of evidence and the number of possible values that each piece of evidence can take. In our PRA framework, the frame of discernment Θ has two elements, first one supports the assertion, p , and the other that does not support the assertion, q . Θ with elements $\{p, q\}$ has the power set $\{\phi, \{p\}, \{q\}, \{p, q\}\}$

From this, we can calculate the normalization constant and combined m-values,

$$K = 1 - \{m_1(p)m_2(q) + m_1(q)m_2(p)\} \quad (1)$$

$$m(p) = \frac{m_1(p)m_2(p) + m_1(p)m_2(\{p, q\}) + m_1(\{p, q\})m_2(p)}{K} \quad (2)$$

$$m(q) = \frac{m_1(q)m_2(q) + m_1(q)m_2(\{p, q\}) + m_1(\{p, q\})m_2(q)}{K} \quad (3)$$

$$m(\{p, q\}) = \frac{m_1(\{p, q\})m_2(\{p, q\})}{K} \quad (4)$$

Since the size of power set is constant, calculating normalization constant and combining two evidences takes constant time i.e., $O(1)$. Combining the first two mass functions gives an intermediate mass function, which can be combined with the third mass function, and so on. In this case, we will have $n - 1$ combination operations for n pieces of evidence, thus cost of combining is $(n - 1) \cdot O(1) = O(n)$.

4.2. PRA Framework

Let's assume the tree representation of the PRA framework has a height h , with the root node positioned at level 0 and all evidence represented by leaf nodes at level h . Each node represents either a component or an assertion, while leaf nodes represent evidences. From the root node, all nodes up to and including level $h - 2$ have,

on average, k nodes each. At level $h - 1$, each node typically has an average of n evidences. In real-world scenarios, the number of evidences n is significantly greater than both k and h , and neither h nor k are large numbers.

Hence, at level i (for $1 \leq i < h - 1$), the number of nodes at that level is k^{i-1} (Figure 3). At level $h - 1$, each of the k^{h-2} nodes has n children, totaling to $k^{h-2} \cdot n$ nodes. Thus total number of nodes,

$$\begin{aligned}
 T &= 1 + k + k^2 + \dots + k^i + \dots + k^{h-2} + k^{h-2} \cdot n \\
 &= \sum_{i=0}^{h-2} k^i + k^{h-2} \cdot n \quad \triangleright \text{the first summation term is a geometric series} \\
 &= \frac{k^{h-1} - 1}{k - 1} + k^{h-2} \cdot n \\
 &\approx k^{h-1} + k^{h-2} \cdot n \\
 &= k \cdot k^{h-2} + k^{h-2} \cdot n \\
 &= k^{h-2} \cdot (n + k)
 \end{aligned}$$

Section 4.1 already showed that the cost of combining n pieces of evidence is $O(n)$, hence cost of combining T nodes is, $O(T) = O(k^{h-2} \cdot (n + k))$.

To integrate a new sub-component C_i into the framework at level i , we must first determine the computational cost associated with this sub-component. If we treat the new component C_i as a smaller risk object, its cost analysis mirrors that of the entire risk object $\mathcal{R}$. Let's assume the sub-component C_i contains T_i nodes. Given that the DS rule of combination is associative, combining the belief value of the new component with the pre-calculated value will yield the updated value efficiently. Therefore, the integration cost for C_i is $O(T_i) + O(1) = O(T_i)$.

Now, let's consider the scenario where we integrate an evidence node instead of a sub-component. Since an evidence node represents a single entity, the term $O(T_i)$, which accounts for the number of nodes in a sub-component, simplifies to $O(1)$. This simplification occurs because the evidence node does not require any further decomposition or combination with other nodes. Therefore, the integration cost of an evidence node is $O(1)$.

The computational cost of the proposed framework remains tractable. First, DS aggregation operates over a fixed frame $\Theta = \{p, q\}$, resulting in constant-time combination and linear complexity $O(n)$ with respect to the number of evidence items. Second, the hierarchical structure enables modular computation, where evidence updates (e.g., addition or removal) are localized to affected branches, avoiding global recomputation. Most component-level risk values are precomputed and stored. Thus, during a risk evaluation request, only a limited subset of affected sub-components is processed, followed by bottom-up aggregation along the corresponding branches of the hierarchy. This design significantly reduces runtime overhead and ensures efficient operation in practice.

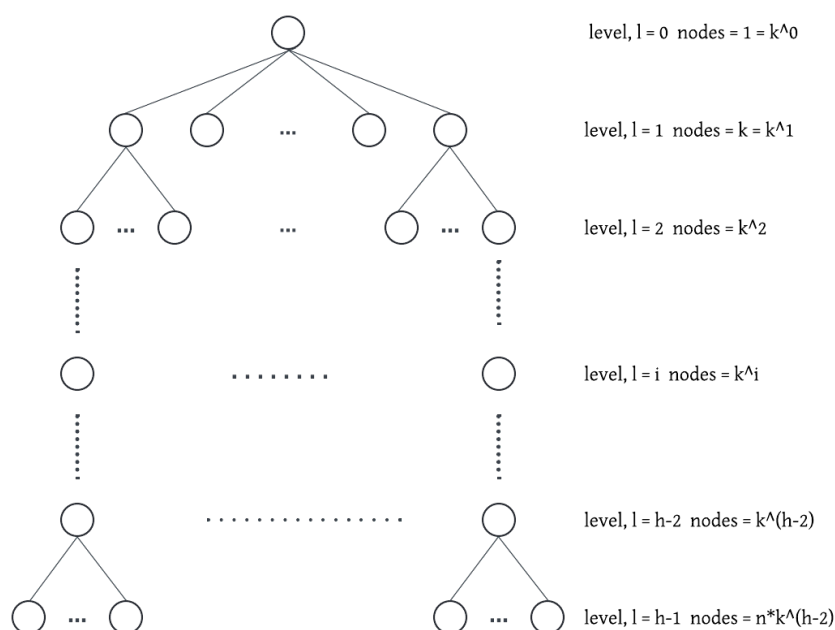


Figure 3. Node calculation of PRA framework.

5. Use Case: Smart Grid System

In this section, we demonstrate the feasibility of the PRA framework by applying it to a smart grid (SG) system, specifically the Green Button Initiative (GBI) (<https://www.energy.gov/data/green-button> (accessed on 17 February 2026)). SG operates as an information system that collects, processes, stores, and exchanges consumer energy usage data to support monitoring and service delivery. GBI is a SG application that allows users to download and share their energy usage data with third parties. Assume *SGcom* is a fictitious SG company that implements GBI with the following policies.

- (1) Data Accessibility:
 - (a) consumers have easy access to their energy usage data through a user-friendly web portal interface.
 - (b) only registered consumers can access their energy usage data using valid login credentials.
- (2) Data Standardization: to ensure interoperability and compatibility with all parties involved in the SG system, we adopted the common ontology proposed in [33].
- (3) Privacy and Security Policy:
 - (a) Access Controls: access to consumer energy usage data is strictly controlled and limited to authorized personnel who require it to perform their duties, with authorization enforced using OAuth 2.0.
 - (b) Data Encryption: all energy usage data transmitted between the consumer's devices, company servers, and any third-party applications is encrypted using industry-standard encryption protocol, i.e., SSL encryption to ensure that data remains confidential and secure during transmission.
 - (c) Incident Response Plan: to minimize the impact on consumers and mitigate any potential damage in the event of a data breach or security incident, prompt notification is sent to the affected individuals, authorities, and stakeholders, as well as steps for containing and remedying the breach.
 - (d) Regular Audits and Compliance Checks: conducted to ensure that privacy and security measures remain effective and aligned with industry best practices and regulatory requirements. Any identified vulnerabilities or non-compliance issues are promptly addressed and remediated.
 - (e) Data Retention Policies: ensure that consumer data are not retained longer than necessary.
- (4) Customer Data Privacy Protection:
 - (a) Adhere to the California Consumer Privacy Act (CCPA) regulations to safeguard customer data privacy rights.
 - (b) Consent Mechanisms: prior to accessing or sharing energy usage data, consumers will be required to provide explicit consent.
 - (c) Compliance Monitoring and Reporting: appoint a designated compliance officer responsible for monitoring adherence to CCPA regulations.

5.1. Phase 1: Model Specification

The GBI implementation consists of three components. *OpenESPI-DataCustodian-Java* allows consumers to download energy usage data and approve third-party authorization requests. *OpenESPI-ThirdParty-Java* enables third parties to request authorization and obtain authenticated usage data. *OpenESPI-Common-Java* provides shared schema, data structures, and persistence services for both applications. Each component is associated with a primary assertion that may include multiple sub-assertions.

An SG ontology enumerates system entities and their associated data attributes [33], enabling privacy risk to be evaluated at the attribute level. Quasi-identifiers are particularly important because combinations of non-unique attributes may permit indirect identification or sensitive inference.

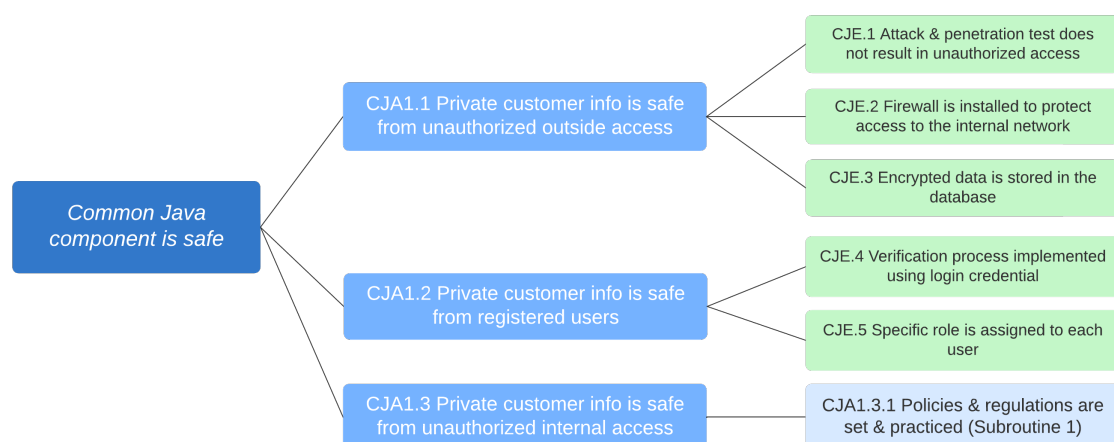
The attributes correspond to several privacy risk categories. Direct identifiers (e.g., SSN, state ID) introduce explicit identification risk. Energy usage logs, appliance characteristics, and consumption patterns reveal behavioral information and support burglary and presence inference. Consumer profile and debt information enable discrimination, targeted advertising, and financial inference. Vehicle type, charging behavior, and smart-meter location allow mobility and location tracking, while infrastructure schedules and detailed usage timing expose occupancy patterns. These mappings provide the evidential basis for constructing assertions and evaluating privacy exposure.

The ERM diagram defines the assessment structure. Relevant assessments are identified for each component, and organizational policies serve as the primary source for specifying the supporting evidence.

Figures 4–8 present the ERM diagram of each component of the GBI and all the evidences are described in Table 2. Evidence related to Common-Java components is prefixed with CJE for simplicity, while evidence associated with Data-Custodian components and Third-party components is prefixed with DCE and TPE, respectively. The Third-party ERM includes a subroutine, and evidence related to this subroutine begins with the prefix 'sub'.

Table 2. Descriptions of evidence.

Evidence	Description
CJE.1	An attack and penetration test is a simulated attack to evaluate and identify vulnerabilities that could be exploited by real attackers. The test report provides the details of the vulnerabilities discovered, the methods used to exploit them, and recommendations for remediation.
CJE.2	Firewall is a barrier between the internal network and external networks; it helps to protect data from outside attacks by monitoring and controlling network traffic.
CJE.3	Data is encrypted within the database to ensure that attackers cannot decipher its contents.
CJE.4	Each user must log in using their credentials; common-java checks if the credentials are valid or not.
CJE.5	Each registered user is assigned a certain role and they can not access any data outside their role.
DCE.1	All customers are required to register through the utility's web portal and create a username and password. These credentials are used for subsequent logins.
DCE.2	No additional data is transmitted other than the requested data.
DCE.3	No additional process is running other than data transfer.
DCE.4	Customer information is encrypted using SSL during transmission to prevent unauthorized parties from accessing or using the data, even if the transmitted data are intercepted.
TPE.1	All third-parties are provided unique login credentials which are used for future login.
TPE.2	OAuth2 is used to enable the third-party to access customers' energy data securely and with their consent.
TPE.3	See DCE.4.
TPE.4	Third-party applications must be registered with the utility company before they can access customer data through the Green Button Initiative.
TPE.5	The Third-Party must agree to utility company terms as part of the onboarding process.
TPE.6	Utility companies may conduct regular audits to verify whether third parties use customer data in accordance with the stated purposes; however, such audits are generally not performed, which is considered negative evidence in the risk assessment.
sub.2.1	Unique attributes may reveal too much about the customer.
sub.2.2	Quasi-identifiers can help infer information about the customer, so it is important not to reveal such information.
sub.2.3	Availability of public records related to requested attributes might assist in inferring additional information about the customer.
sub.2.4	Background knowledge about the customer can lead to inferences about private information.
sub.2.5	The privacy risks associated with the requested attributes must be carefully evaluated to minimize their potential exploitation for fraudulent activities.
sub.2.6	Requested attributes must not reveal the customer's monetary condition.
sub.2.7	Time-series data can reveal patterns and routines that may disclose sensitive information about the customer.
sub.2.8	Sharing fine-grained data can reveal highly detailed and specific information about individuals.

**Figure 4.** ERM for OpenESPI-Common-Java component.

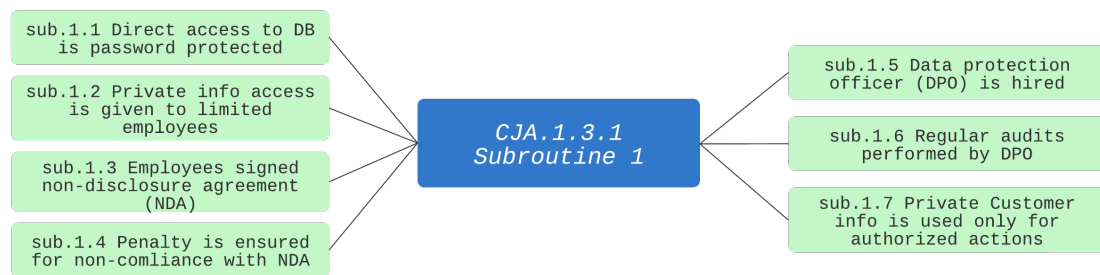


Figure 5. ERM for OpenESPI-Common-java subroutine 1.

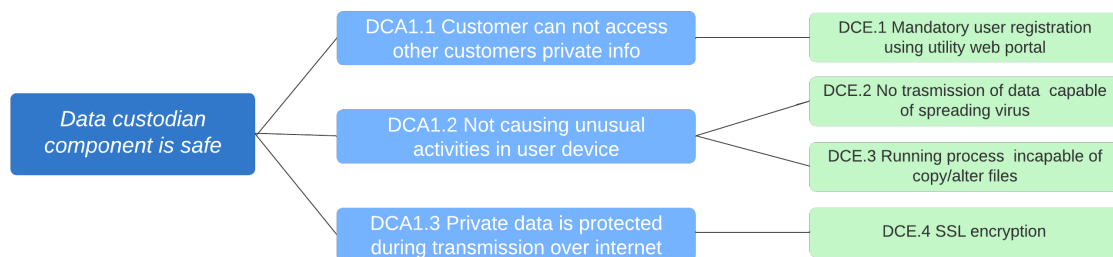


Figure 6. ERM for OpenESPI-DataCustodian-Java component.

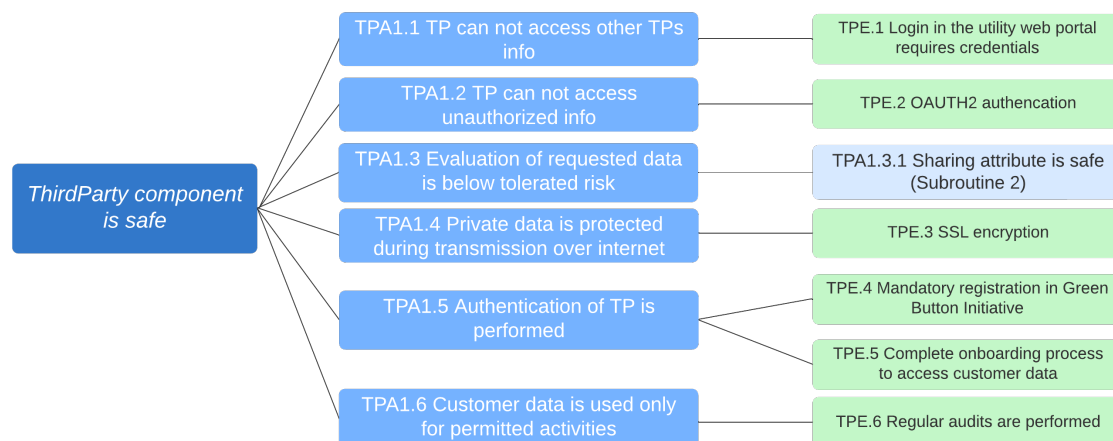


Figure 7. ERM for OpenESPI-ThirdParty-Java component.

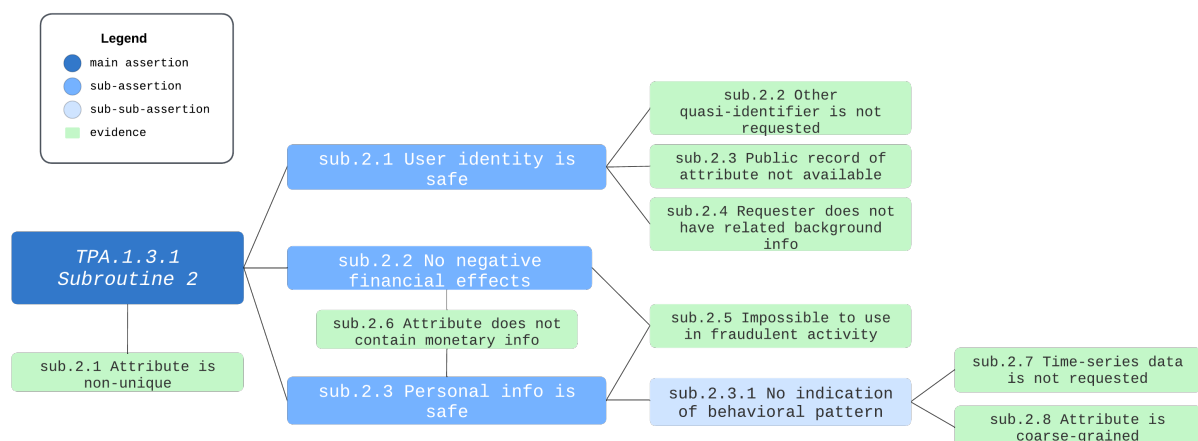


Figure 8. ERM for an attribute.

5.2. Phase 2: Belief Assignment to Evidence

Belief assignment distributes probability mass to subsets of the frame of discernment, representing the degree of belief supported by the evidence. This step quantifies confidence in competing hypotheses. In this framework, the frame Θ has two elements: evidence supports the assertion or does not support the assertion. The mass function m is defined over the power set of Θ , producing three values: belief, disbelief, and uncertainty (excluding the null set). Each evidence must assign values to all three.

In this use case, a hybrid, human-AI collaborative approach has been adopted. LLMs are used only to generate preliminary, illustrative mass assignments for hypothetical evaluation, leveraging their trained knowledge base to produce context-aware initial estimates. Evidence descriptions, relevant attributes, and the power set of Θ were provided to the models to generate initial m -values; however, these values are not treated as authoritative. Additionally, domain experts were consulted to assign m -values by incorporating contextual understanding and domain-specific knowledge. The resulting belief assignments therefore reflect both data-driven consistency and expert-informed judgment, aligning with the evidential reasoning paradigm and ensuring robust representation of epistemic uncertainty.

5.3. Phase 3: Risk Calculation

Risk calculation is the final step in our PRA framework. After collecting evidence and assigning belief values, we employed Algorithm 1 to evaluate the risks associated with information sharing. The DS combination operation utilized in this process is detailed in the Equations (1)–(4) presented in Section 4.1. The implementation and detailed usage instructions are publicly available on GitHub (<https://github.com/raisa-nmt/PRA-Framework> (accessed on 27 April 2026)).

Let's take sub2.1 into consideration. The initial mass assignments are:

- GPT-4o: {0.8, 0.2, 0.0}
- GPT-4: {0.8, 0.1, 0.1}
- Expert-1: {0.7, 0.1, 0.2}
- Expert-2: {0.85, 0.05, 0.1}

Let's combine the two LLM-generated assignments:

$$\begin{aligned}
 K_{LLM} &= 1 - \{(0.8 \times 0.1) + (0.2 \times 0.8)\} = 0.76 \\
 m_{LLM}(p) &= \frac{(0.8 \times 0.8) + (0.8 \times 0.1) + (0.0 \times 0.8)}{0.76} = 0.947 \\
 m_{LLM}(q) &= \frac{(0.2 \times 0.1) + (0.2 \times 0.1) + (0.0 \times 0.1)}{0.76} = 0.053 \\
 m_{LLM}(\{p, q\}) &= \frac{0.0 \times 0.1}{0.76} = 0.000
 \end{aligned}$$

Thus,

$$m_{LLM} = \{0.947, 0.053, 0.000\}.$$

Next, the two expert assignments are combined:

$$\begin{aligned}
 K_{EXP} &= 1 - \{(0.7 \times 0.05) + (0.1 \times 0.85)\} = 0.88 \\
 m_{EXP}(p) &= \frac{(0.7 \times 0.85) + (0.7 \times 0.1) + (0.2 \times 0.85)}{0.88} = 0.949 \\
 m_{EXP}(q) &= \frac{(0.1 \times 0.05) + (0.1 \times 0.1) + (0.2 \times 0.05)}{0.88} = 0.028 \\
 m_{EXP}(\{p, q\}) &= \frac{0.2 \times 0.1}{0.88} = 0.023
 \end{aligned}$$

Thus,

$$m_{EXP} = \{0.949, 0.028, 0.023\}.$$

Finally, the aggregated LLM and expert belief assignments are combined:

$$K_{final} = 1 - \{(0.947 \times 0.028) + (0.053 \times 0.949)\} = 0.923$$

$$m_{final}(p) = \frac{(0.947 \times 0.949) + (0.947 \times 0.023) + (0.000 \times 0.949)}{0.923} = 0.997$$

$$m_{final}(q) = \frac{(0.053 \times 0.028) + (0.053 \times 0.023) + (0.000 \times 0.028)}{0.923} = 0.003$$

$$m_{final}(\{p, q\}) = \frac{0.000 \times 0.023}{0.923} = 0.000$$

Due to the associative property of DS combination, the order of aggregation does not affect the final result. The final combined mass for sub2.1 converges to:

$$\{0.997, 0.003, 0.0\}$$

which matches the value reported in Table 3. This process is repeated for all other sub-assertions.

For the purpose of this study, we assume that a third party has requested summary usage data for research purposes which includes current billing period, consumption and daily usage (<https://s3-us-west-2.amazonaws.com/technical.greenbuttonalliance.org/library/sample-data/TestGBDataOneYearDailyBinnedMonthly.xml> (accessed on 14 July 2025).) The requested information includes attributes such as the time period of energy consumption, usage in kilowatt-hours, and the cost of usage. Table 3 presents the attribute-level m-values for the daily energy consumption attribute, which were collected in Phase 2.

The results indicate that most sub-assertions strongly support the proposition that sharing the attribute is privacy-safe, as reflected by high belief values and minimal uncertainty. In contrast, evidence sub2.7 exhibits a significantly higher disbelief component, indicating potential privacy risks associated with time-series data. This differentiation demonstrates the ability of the framework to identify specific risk factors rather than producing a uniform risk score. Based on the aggregated belief values, the resulting belief interval satisfies the acceptance threshold defined in Section 3.3, indicating that the requested data can be shared under the current risk tolerance. This demonstrates how the framework translates evidential aggregation into actionable privacy decisions.

Table 3. Attribute-level evidence combination.

	sub2.1	sub2.2	sub2.3	sub2.4	sub2.5	sub2.6	sub2.7	sub2.8
GPT-4o	{0.8, 0.2, 0.0}	{0.9, 0.1, 0.0}	{0.7, 0.2, 0.1}	{0.8, 0.15, 0.05}	{0.85, 0.1, 0.05}	{0.75, 0.2, 0.05}	{0.6, 0.3, 0.1}	{0.7, 0.2, 0.1}
GPT-4	{0.8, 0.1, 0.1}	{0.9, 0.05, 0.05}	{0.7, 0.2, 0.1}	{0.6, 0.2, 0.2}	{0.9, 0.05, 0.05}	{0.85, 0.1, 0.05}	{0.3, 0.6, 0.1}	{0.65, 0.25, 0.1}
Expert-1	{0.7, 0.1, 0.2}	{0.8, 0.1, 0.1}	{0.8, 0.1, 0.1}	{0.6, 0.25, 0.15}	{0.85, 0.05, 0.1}	{0.6, 0.3, 0.1}	{0.2, 0.7, 0.1}	{0.55, 0.3, 0.15}
Expert-2	{0.85, 0.05, 0.1}	{0.8, 0.0, 0.2}	{0.7, 0.1, 0.2}	{0.55, 0.15, 0.3}	{0.8, 0.05, 0.15}	{0.6, 0.2, 0.2}	{0.25, 0.65, 0.1}	{0.4, 0.5, 0.1}
	{0.997, 0.003, 0}	{0.999, 0.001, 0}	{0.99, 0.01, 0}	{0.968, 0.031, 0.001}	{0.999, 0.001, 0}	{0.985, 0.015, 0}	{0.149, 0.851, 0.001}	{0.881, 0.118, 0.001}
Combined	{1.0 (approx), 0, 0}							

6. Conclusions

This paper presented a privacy risk assessment framework based on the ERM and DS theory. The framework evaluates privacy risk at the level of specific shared information by decomposing a risk object into components, assertions, and evidences. Through DS evidence combination, heterogeneous and potentially conflicting observations are aggregated while explicitly representing belief, disbelief, and uncertainty, which is essential in privacy scenarios where attacker knowledge and auxiliary information are not directly observable.

The analysis shows that evidence aggregation is linear in the number of evidences, while the initial construction cost depends on the depth of the evidential hierarchy. Once the model is established, new evidences or sharing requests require only localized recomputation, enabling scalable and continuous privacy decision-making in dynamic systems.

A smart grid case study based on the Green Button Initiative demonstrated the practical feasibility of the framework. The results indicate that the proposed approach can systematically quantify privacy exposure and support informed data-sharing decisions in real-world environments.

Our immediate future work will extend the framework to additional domains and explore ML/AI-based automated evidence collection, weight assignment, and structural optimization to further enhance the proposed framework in terms of scalability, efficiency, and adaptability.

Author Contributions

R.I.: Conceptualization, Methodology, Data Curation, Formal Analysis, Writing—Original Draft; D.S.: Supervision, Validation, Writing—Review & Editing. All authors have read and agreed to the published version of the manuscript.

Funding

This work was partially supported at the Secure Computing Laboratory at New Mexico Tech by the grant from the National Science Foundation (NSF-DGE-1946650).

Data Availability Statement

Implementation of Algorithm-1 and high-resolution images used in this paper are available on GitHub at: <https://github.com/raisa-nmt/PRA-Framework>.

Conflicts of Interest

The authors declare no conflict of interest.

Use of AI and AI-Assisted Technologies

During the preparation of this work, the authors used GPT-4 and GPT-4o to generate preliminary, illustrative belief assignments for selected evidence in the case study. These AI-generated values were used as initial, context-aware estimates. After using these tools, the authors reviewed and edited the generated content as needed and take full responsibility for the content and conclusions presented in the published article.

References

1. Cuijpers, C.; Koops, B.-J. Smart Metering and Privacy in Europe: Lessons from the Dutch Case. In *European Data Protection*; Gutwirth, S., Leenes, R., de Hert, P., et al., Eds.; Springer: Dordrecht, The Netherlands, 2012; pp. 269–293.
2. Javaid, M.; Haleem, A.; Singh, R.P.; et al. Towards Insighting Cybersecurity for Healthcare Domains: A Comprehensive Review of Recent Practices and Trends. *Cyber Secur. Appl.* **2023**, *1*, 100016.
3. Quinn, K. Why We Share: A Uses and Gratifications Approach to Privacy Regulation in Social Media Use. *J. Broadcast. Electron. Media* **2016**, *60*, 61–86.
4. California Consumer Privacy Act of 2018. Available online: https://leginfo.ca.gov/faces/codes_displayText.xhtml?division=3.&part=4.&lawCode=CIV&title=1.81.5 (accessed on 5 April 2026).
5. Children’s Online Privacy Protection Rule (“COPPA”). Available online: <https://www.ftc.gov/legal-library/browse/rules/childrens-online-privacy-protection-rule-coppa> (accessed on 5 April 2026).
6. Federal Information Security Modernization Act of 2014. Available online: <https://www.cisa.gov/topics/cyber-threats-and-advisories/federal-information-security-modernization-act> (accessed on 5 April 2026).
7. Health Insurance Portability and Accountability Act of 1996 (HIPAA). Available online: <https://www.cdc.gov/phlp/php/resources/health-insurance-portability-and-accountability-act-of-1996-hipaa.html> (accessed on 5 April 2026).
8. Zhou, M.; Liu, X.B.; Yang, J.B.; et al. A Revised Evidential Reasoning Model Based on Interval Information. In Proceedings of the 2009 Sixth International Conference on Fuzzy Systems and Knowledge Discovery, Tianjin, China, 14–16 August 2009; Volume 4, pp. 239–243.
9. Van der Bles, A.M.; van der Linden, S.; Freeman, A.L.J.; et al. Communicating Uncertainty About Facts, Numbers and Science. *R. Soc. Open Sci.* **2019**, *6*, 181870.
10. Freund, J.; Jones, J. *Measuring and Managing Information Risk: A FAIR Approach*; Butterworth-Heinemann: Woburn, MA, USA, 2014.
11. Cronk, R.J.; Shapiro, S.S. Quantitative Privacy Risk Analysis. In Proceedings of the 2021 IEEE European Symposium on Security and Privacy Workshops (EuroS&PW), Vienna, Austria, 6–10 September 2021; pp. 340–350.
12. De, S.; Métayer, D. PRIAM: A Privacy Risk Analysis Methodology. In *Data Privacy Management and Security Assurance*; Springer International Publishing: Cham, Switzerland, 2016; pp. 221–229.
13. Yucel, G.; Cebi, S.; Hoege, B.; et al. A Fuzzy Risk Assessment Model for Hospital Information System Implementation. *Expert Syst. Appl.* **2012**, *39*, 1211–1218.
14. Abdymanapov, S.A.; Muratbekov, M.; Altynbek, S.; et al. Fuzzy Expert System of Information Security Risk Assessment on the Example of Analysis Learning Management Systems. *IEEE Access* **2021**, *9*, 156556–156565.
15. Pokorádi, L. Fuzzy Logic-Based Risk Assessment. *Acad. Appl. Res. Mil. Sci.* **2002**, *1*, 63–73.
16. Al Sharif, R.; Pokharel, S. Risk Analysis with the Dempster-Shafer Theory for Smart City Planning: The Case of Qatar. *Electronics* **2021**, *10*, 3080.
17. Tang, Y.; Wang, L.; Yang, L.; et al. Information Security Risk Assessment Method Based on Cloud Model. In Proceedings of the 25th IET Irish Signals & Systems Conference 2014 and 2014 China-Ireland International Conference on Information and Communications Technologies (ISSC 2014/CICT 2014), Limerick, Ireland, 26–27 June 2014; pp. 258–262.
18. Levy, M. A Novel Framework for Data Center Risk Assessment. In Proceedings of the 2020 11th IEEE Annual Ubiquitous Computing, Electronics & Mobile Communication Conference (UEMCON), New York, NY, USA, 28–31 October 2020; pp. 0148–0154.

19. Sion, L.; Van Landuyt, D.; Wuyts, K.; et al. Privacy Risk Assessment for Data Subject-Aware Threat Modeling. In Proceedings of the 2019 IEEE Security and Privacy Workshops (SPW), San Francisco, CA, USA, 19–23 May 2019; pp. 64–71.
20. Haas, P.J. Monte Carlo Methods for Uncertain Data. In *Encyclopedia of Database Systems*; Liu, L., Özsu, M.T., Eds.; Springer: New York, NY, USA, 2018; pp. 2299–2306.
21. Wairimu, S.; Fritsch, L. Modelling Privacy Harms of Compromised Personal Medical Data—Beyond Data Breach. In Proceedings of the 17th International Conference on Availability, Reliability and Security, Vienna, Austria, 23–26 August 2022; pp. 1–9.
22. Asif, W.; Ray, I.G.; Rajarajan, M. An Attack Tree Based Risk Evaluation Approach for the Internet of Things. In Proceedings of the 8th International Conference on the Internet of Things, Santa Barbara, CA, USA, 15–18 October 2018; pp. 1–8.
23. Markovic, M.; Asif, W.; Corsar, D.; et al. Towards Automated Privacy Risk Assessments in IoT Systems. In Proceedings of the 5th Workshop on Middleware and Applications for the Internet of Things, Rennes, France, 10–11 December 2018; pp. 15–18.
24. Wu, T.; Zhao, G. A New Security and Privacy Risk Assessment Model for Information System Considering Influence Relation of Risk Elements. In Proceedings of the 2014 Ninth International Conference on Broadband and Wireless Computing, Communication and Applications, Guangdong, China, 8–10 November 2014; pp. 233–238.
25. Wang, L.; Wang, B.; Peng, Y. Research the Information Security Risk Assessment Technique Based on Bayesian Network. In Proceedings of the 2010 3rd International Conference on Advanced Computer Theory and Engineering (ICACTE), Chengdu, China, 20–22 August 2010; Volume 3, pp. V3-600–V3-604.
26. Ke, C.; Wu, J.; Xiao, F.; et al. A Privacy Risk Assessment Scheme for Fog Nodes in Access Control System. *IEEE Trans. Reliab.* **2022**, *71*, 1513–1526.
27. Pellungrini, R.; Pappalardo, L.; Simini, F.; et al. Modeling Adversarial Behavior Against Mobility Data Privacy. *IEEE Trans. Intell. Transp. Syst.* **2022**, *23*, 1145–1158.
28. Wang, Y.; Aghasaryan, A.; Shrihari, A.; et al. Intelligent Reactive Access Control for Moving User Data. In Proceedings of the 2011 IEEE Third International Conference on Privacy, Security, Risk and Trust and 2011 IEEE Third International Conference on Social Computing, Boston, MA, USA, 9–11 October 2011; pp. 942–950.
29. Sun, L.; Srivastava, R.; Mock, T. An Information Systems Security Risk Assessment Model Under Dempster-Shafer Theory of Belief Functions. *J. Manage. Inf. Syst.* **2006**, *22*, 109–142.
30. Dempster, A.P. A Generalization of Bayesian Inference. *J. R. Stat. Soc. Ser. B Methodol.* **1968**, *30*, 205–232.
31. Shafer, G. *A Mathematical Theory of Evidence*; Princeton University Press: Princeton, NJ, USA, 1976; Volume 42.
32. Shenoy, P.P. On Distinct Belief Functions in the Dempster-Shafer Theory. In Proceedings of the Thirteenth International Symposium on Imprecise Probability: Theories and Applications, Oviedo, Spain, 11–14 July 2023; Volume 215, pp. 426–437.
33. Islam, R.; Cerny, T.; Shin, D. Ontology-Based User Privacy Management in Smart Grid. In Proceedings of the 37th ACM/SIGAPP Symposium on Applied Computing, Online, 25–29 April 2022; pp. 174–182.