

REFINED AND OVER/UNDER/OFF EXTENSIONS OF PROBABILITY, IMPRECISE PROBABILITY, NEUTROSOPHIC PROBABILITY, PLITHOGENIC PROBABILITY, MEASURE, AND STATISTICS

Nine Families • Two Generating Operators • Refinement & Over/Under/Off

Acad. Florentin Smarandache

PEER REVIEWERS

Mukhtar Ahmad — Doctoral Program of Mathematics, Faculty of Mathematics
and Natural Sciences, Institut Teknologi Bandung, Indonesia

Victor Christianto — Independent Researcher, Indonesia



NSIA Publishing House

University of New Mexico — University of Guayaquil

Refined and Over/Under/Off Extensions of Probability, Imprecise Probability, Neutrosophic Probability, Plithogenic Probability, Measure, and Statistics

Acad. Florentin Smarandache

Peer Reviewers

Mukhtar Ahmad, Doctoral Program of Mathematics, Faculty of Mathematics and Natural Sciences, Institut Teknologi Bandung, Jalan Ganesha no. 10, Bandung, 40132, Indonesia

Victor Christianto, Independent Researcher, Indonesia

NSIA Publishing House

2026



Neutrosophic Science International Association (NSIA) Publishing House

Division of Mathematics and Sciences
University of New Mexico
705 Gurley Ave., Gallup Campus
NM 87301, United States of America

University of Guayaquil
Av. Kennedy and Av. Delta
"Dr. Salvador Allende" University Campus
Guayaquil 090514, Ecuador

<https://fs.unm.edu/NSIA/>
<https://neutrosophic.org/nsia-publishing-house/>

ISBN: 978-1-972502-46-4

Copyright © 2026 by NSIA Publishing House and the author. This book is protected by copyright. No part of this publication may be reproduced, distributed, or transmitted in any form or by any means, including photocopying, recording, or other electronic or mechanical methods, without the prior written permission of the publisher, except for brief quotations embodied in critical reviews and certain other noncommercial uses permitted by copyright law.

Many books and papers on neutrosophic, plithogenic, and related theories can be downloaded from the NSIA Publishing House and Neutrosophic Science International Association websites listed above.

Peer Reviewers

Mukhtar Ahmad, Doctoral Program of Mathematics, Faculty of Mathematics and Natural Sciences, Institut Teknologi Bandung, Jalan Ganesha no. 10, Bandung, 40132, Indonesia.

Victor Christianto, Independent Researcher, Indonesia.

Preface

Uncertainty rarely arrives in the tidy, closed form that classical probability theory was built to describe. Real evidence is graded, contested, and sometimes internally inconsistent; real measurements exceed their nominal bounds, cancel each other out, or need to be split into finer categories before they mean anything actionable. This book grew out of a simple observation: two operators developed piecemeal over two decades of my work — Refinement, which lets a single component be decomposed into finer, individually trackable sub-components without losing the ability to recombine them exactly, and the Over/Under/Off relaxation, which lets a value legitimately exceed certainty or fall below impossibility — are not two isolated tricks but a pair of composable generators. Applied systematically to Classical Probability, Imprecise Probability, Neutrosophic Probability, Plithogenic Probability, Measure, and each of the first four notions' empirical statistics counterpart, they produce nine coherent families that had, until now, never been laid out side by side under one unifying architecture.

That is the purpose of this volume: not to introduce new mathematics from a blank page, but to synthesize, extend, and systematically cross-reference operators I introduced independently in earlier work — the Refined Fuzzy Set and Refined Neutrosophic Logic, the Neutrosophic Overset/Underset/Offset, Plithogenic Probability and Statistics, and the 2026 Over/Under/Off-Masses extension of information fusion theory — into a single reading path. Each of the nine families receives the same disciplined treatment: formal definitions, worked real-world applications drawn from quality control, actuarial risk, elections, medicine, supply chains, and finance, a study of how the seven classical event-type properties behave under refinement and relaxation, and an explicit account of how the family connects to its siblings and to established fields — Dempster-Shafer evidence theory, Bayesian inference, multi-criteria decision-making, machine learning uncertainty quantification, and quantum probability among them.

I have tried throughout to resist the temptation to let formalism outrun usefulness. Part VIII of this book states the framework's limitations candidly rather than leaving them for a reader to discover the hard way, and Part VII collects the conceptual pitfalls that most often trip up newcomers to refined and relaxed uncertainty structures. My hope is that a reader who works with data that classical, single-valued, bounded probability cannot honestly represent — a protective interaction that looks like independence until a contradiction degree is allowed to go negative, an indeterminate survey response that is not simply a missing value, a signed measure that must capture net harm rather than only diminished benefit — will find in these pages both the vocabulary to name what they are seeing and a worked template for handling it rigorously.

I am grateful to the peer reviewers, Dr. Mukhtar Ahmad of Institut Teknologi Bandung and Dr. Victor Christianto, independent researcher, for their careful reading and constructive comments, and to the Neutrosophic Science International Association Publishing House, jointly hosted by the University of New Mexico and the University of Guayaquil, for making this work freely available to the research community.

Florentin Smarandache
University of New Mexico, 2026

Table of Contents

Refined and Over/Under/Off Extensions of Probability, Imprecise Probability, Neutrosophic Probability, Plithogenic Probability, Measure, and Statistics.....	6
Executive Summary.....	6
Part I — The Two Generating Operators.....	7
Part II — Reading Guide for Each Family.....	8
Family 1 — Classical Probability — Refined / Over-Under-Off / Over-Under-Off-Refined.....	8
Family 2 — Imprecise Probability — Refined / Over-Under-Off / Over-Under-Off-Refined....	11
Family 3 — Neutrosophic Probability — Refined / Over-Under-Off / Over-Under-Off-Refined	14
Family 4 — Plithogenic Probability — Refined / Over-Under-Off / Over-Under-Off-Refined	20
Family 5 — Measure — Refined / Over-Under-Off / Over-Under-Off-Refined.....	23
Family 6 — Classical Statistics — Refined / Over-Under-Off / Over-Under-Off-Refined.....	25
Family 7 — Imprecise Statistics — Refined / Over-Under-Off / Over-Under-Off-Refined.....	29
Family 8 — Neutrosophic Statistics — Refined / Over-Under-Off / Over-Under-Off-Refined.	31
Family 9 — Plithogenic Statistics — Refined / Over-Under-Off / Over-Under-Off-Refined....	34
Part III — Overall Synthesis: How the Nine Families Relate.....	37
Part IV — Extended Examples and Applications: Refined Plithogenic Probability, Refined Plithogenic Statistics, and Over/Under/Off-Refined Plithogenic Statistics.....	38
Part V — The Nine Families in the Landscape of Uncertainty Science.....	42
Part VI — Worked Numerical Case Studies.....	48
Glossary of Notation.....	50
Comparison of the Nine Families.....	50
Part VII — Common Pitfalls and Frequently Asked Questions.....	52
Appendix A — Index of Worked Applications by Field.....	54
Appendix B — Suggested Further Reading by Connected Field.....	55
Part VIII — Limitations and Open Research Questions.....	56
Part IX — Conclusion.....	58
References.....	59

Refined and Over/Under/Off Extensions of Probability, Imprecise Probability, Neutrosophic Probability, Plithogenic Probability, Measure, and Statistics

Refined and Over/Under/Off Extensions of Probability, Imprecise Probability, Neutrosophic Probability, Plithogenic Probability, Measure, and Statistics, working document, 2026, was building on the Refinement and Over/Under/Off operators of F. Smarandache [1–5, 12].

This document synthesizes and extends previously published operators into nine explicitly named families rather than introducing entirely new mathematics from a blank page, any derivative or follow-on work should cite both this document for the specific family definitions, worked examples, and cross-field connections developed here, and the original numbered references for the underlying Refinement and Over/Under/Off constructions themselves.

Executive Summary

This document develops nine interrelated “families” of extended probability, measure, and statistics theory, each obtained by applying two composable operators — Refinement (splitting a component into finer-grained, individually-trackable sub-components that recombine exactly to the original value) and Over/Under/Off (relaxing the $[0,1]$ boundary so a value can exceed certainty or fall below impossibility) — to an established base notion: Classical Probability, Imprecise Probability, Neutrosophic Probability, Plithogenic Probability, and Measure, together with each of the first four’s empirical statistics counterpart. For every family, the document works through formal definitions, three real-world worked applications, a study of seven classical event-type properties (complementary, independent, dependent, mutually exclusive, conditional, De Morgan’s Rule, and Bayes’ theorem) under refinement and relaxation, and an account of how the family relates both to its seven sibling families and to established fields outside this framework entirely — from classical frequentist statistics and Bayesian inference through Dempster-Shafer evidence theory, multi-criteria decision-making, machine learning uncertainty quantification, actuarial risk theory, and quantum probability.

The document is organized in eight parts plus two appendices. Parts I and II lay out the two generating operators and the reading guide followed by every family. Families 1 through 9 (interleaved through the body of the document) each receive the definitions/examples/event-study/relations/connections treatment described above. Part III synthesizes how the nine families derive from one another, illustrated with a diagram. Part IV expands the three plithogenic-family notions with additional worked examples and diagrams. Part V surveys seventeen established fields the framework connects to, organized field by field rather than family by family. Part VI works four representative calculations through with full arithmetic, so every formula stated abstractly elsewhere in the document can be checked by hand. Part VII addresses common conceptual pitfalls in FAQ form. Part VIII states the framework’s limitations and open research questions candidly, rather than leaving them implicit. The two appendices index every worked application by industry and point toward the established literature underlying each of Part V’s field connections.

A reader new to this material is best served by reading Part I (the two operators) and Part II (the reading guide) first, then either working straight through Families 1–9 in order (each builds conceptually on the one before it, from Classical Probability’s simplicity to Plithogenic Statistics’ full generality) or jumping directly to whichever family matches their own application domain and using that family’s §X.5 Connections subsection to trace outward into Part V for the relevant established-field context.

Part I — The Two Generating Operators

Everything below is generated by composing two operators that Smarandache introduced independently [1, 2, 3, 12] and that this document now applies uniformly across nine families.

The Refinement operator R. Given a component (a truth-degree T , a probability value P , a statistical category, etc.), refinement splits it into sub-components $T_1, T_2, \dots, T_p$ ($p \geq 2$) that partition the original meaning more finely — e.g. “positive outcome” T is refined into T_1 = “clearly positive” and T_2 = “marginally positive.” This generalizes Smarandache’s 2016 definition of the Refined Fuzzy Set and Refined Intuitionistic Fuzzy Set [1] and his 2013 refinement of the neutrosophic components (T, I, F) into sub-components $(T_1, T_2, \dots; I_1, I_2, \dots; F_1, F_2, \dots)$ [2]. Refinement never changes the *range* of the values (they stay wherever the parent notion lived); it only increases the *granularity* of description. Re-fusing the sub-components (by summation, weighted summation, or a fusion rule appropriate to the family) must recover the original unrefined value — this “reduction consistency,” per the refinement of uncertain sets [12], is the defining constraint of any refinement.

The Over/Under/Off operator O. Given a component living nominally in $[0, 1]$ (or in whatever normalized interval the family uses), the Over/Under/Off operator relaxes the boundary, following Smarandache’s 2007 extension of the uncertain Set to uncertain Overset, Underset, and Offset [3] and the associated degrees-of-membership framework [4]: -

Overset/Over-value: some component is allowed to exceed 1 (super-membership, super-probability — e.g., an employee whose contribution exceeds “full normal performance”). -

Underset/Under-value: some component is allowed to fall below 0 (sub-membership, negative probability/contribution — e.g., an employee who does net damage). -

Offset/Off-value: components can go both above 1 and below 0 simultaneously across different sub-components of the same object.

The same relaxation was extended in 2026 from set-membership to the Masses of Information Fusion theory, yielding Over/Under/Off-Masses [5] — the direct ancestor of the Over/Under/Off-Measure notion developed in Family 5 below.

Both operators are **composable and commute in the naming convention** (though not always in the mathematics, which is why each family below is worked out separately): - Apply R alone — *Refined-X*. - Apply O alone — *Over/Under/Off-X*. - Apply O then R (refine an already-relaxed structure) — *Over/Under/Off-Refined-X*. - Apply R then O (relax an already-refined structure) is mathematically the same target object as the previous line, so the literature collapses these into the single notion *Over/Under/Off-Refined-X*.

This is why every family below has exactly three (or, for the last Statistics family, four) named variants, and why the “Off-Refined” variant is always presented last, as the composition of the first two.

Part II — Reading Guide for Each Family

For each of the nine families this document covers: 1. **Definitions** of the Refined, Over/Under/Off, and Over/Under/Off-Refined variants. 2. **A real example/application** grounding the family. 3. **Event-type study** — how the seven classical notions behave once the family's numbers are allowed to be refined and/or off-interval: - Complementary events - Independent events - Dependent (non-independent) events - Mutually exclusive events - Conditional events - De Morgan's Rule - Bayes' theorem 4. **Relations** — connections to, and distinctions from, neighboring probabilistic/statistical notions.

Family 1 — Classical Probability → Refined / Over-Under-Off / Over-Under-Off-Refined

1.1 Definitions

Refined Classical Probability (RCP). Classical probability assigns $P(A) \in [0,1]$ to an event A in a sample space S with $\sum P(A_i) = 1$. Refined Classical Probability replaces a single event A by a refined family $A_1, A_2, \dots, A_p \subset A$ (a sub-partition of A 's outcomes), each carrying its own probability $P(A_1), \dots, P(A_p)$, subject to the reduction constraint

$$P(A) = P(A_1) + P(A_2) + \dots + P(A_p).$$

This is not the same as ordinary conditioning: the A_i are a *fixed decomposition of the event itself* (part of how the event is described), not a family of events indexed by a separate conditioning variable.

Over/Under/Off-Classical Probability (OUO-CP). Following the Over/Under/Off extension of uncertain measures [3, 4], the normalization $\sum P(A_i) = 1$ and range $P(A) \in [0,1]$ are relaxed: - Over-Probability: some $P(A_i) > 1$, or $\sum P(A_i) > 1$, when the model needs to express “more than certain / more than total” weight (e.g. reliability scores stacked from redundant independent confirmations that overshoot normalized certainty by design, before renormalization is imposed). - Under-Probability: some $P(A_i) < 0$, expressing a *negative propensity* — a force actively working against the event, stronger than mere absence. - Off-Probability: both occur across different A_i in the same space.

Over/Under/Off-Refined Classical Probability (OUO-RCP). Apply refinement to an Over/Under/Off probability space: A is split into $A_1, \dots, A_p$, and each $P(A_i)$ is allowed to sit outside $[0,1]$, with

$$P(A) = \sum_i P(A_i), P(A) \text{ itself possibly also outside } [0,1].$$

1.2 Real example/application

A **factory quality-control line** rates a batch of items as Pass (A). Classically $P(\text{Pass}) = 0.92$.

- *Refined:* Pass is refined into $A_1 = \text{“Pass, top-tier finish”}$ and $A_2 = \text{“Pass, standard finish,”}$ with $P(A_1)=0.35$, $P(A_2)=0.57$, summing to 0.92.

- *Over/Under/Off*: A parallel “penalty” track records not just Pass/Fail but a contribution-to-brand-value score. A batch that passes *and* triggers a positive press mention might be scored $P = 1.15$ (over); a batch that passes inspection but is later recalled might be re-scored $P = -0.10$ (under), because it did net brand damage despite nominally “passing.”
- *Over/Under/Off-Refined*: the recalled batch’s damage is itself refined into $A_1 =$ “safety-related recall cost” (-0.25) and $A_2 =$ “reputational offset from fast recall response” ($+0.15$), summing to -0.10 .

A second application — weather-forecasting skill scores. A meteorological service reports $P(\text{rain tomorrow}) = 0.65$. *Refined*: the forecast is decomposed into $P(\text{light rain}) = 0.40$ and $P(\text{heavy rain}) = 0.25$, each verified separately against a different downstream decision (umbrella vs. flood advisory). *Over/Under/Off*: a forecaster’s skill score can be rewarded above 1 when an ensemble of independent models all converge on the same correct outcome well ahead of the event window (a “confidence bonus” the raw probability scale cannot express), and penalized below 0 when a forecast is not just wrong but actively misleading relative to what a naive climatological baseline would have said — the Brier-score decomposition literature already flags this asymmetry, which OUO-CP makes explicit rather than folding into a single skill number.

A third application — sports odds-making. A sportsbook posts moneyline odds implying $P(\text{team A wins}) = 0.58$. *Refined*: this is split into $P(\text{win in regulation}) = 0.44$ and $P(\text{win in overtime}) = 0.14$, since the book prices overtime exposure separately for parlay products. *Over/Under/Off*: the sum of all outcomes’ implied probabilities exceeds 1 by design — the bookmaker’s “vig” or “overround,” typically 1.03–1.10 — which is a textbook Over-Probability: the classical axiom $\sum P(A_i) = 1$ is deliberately violated as the business model, not an error to be corrected.

1.3 Event-type study

- **Complementary events**: Classically $P(\bar{A}) = 1 - P(A)$. Under refinement, $\bar{A}$ is refined too ($\bar{A}_1, \dots, \bar{A}_\phi$), each still summing to $1 - P(A)$. Under *Over/Under/Off*, “ $1 - P(A)$ ” is no longer guaranteed to lie in $[0,1]$ either — if $P(A) = 1.15$, then $P(\bar{A}) = -0.15$, which is exactly the Under-region, and this is treated as *meaningful* (a complement that itself represents net negative propensity), not as an error to be clipped.
- **Independent events**: $P(A \cap B) = P(A)P(B)$ is retained as the *defining test* for independence in all three variants, but in the *Over/Under/Off* variant this product can now exceed the space’s nominal bound, so “independence” must be checked in the *relaxed* codomain, not assumed to hold only for values in $[0,1]$.
- **Dependent events**: $P(A \cap B) \neq P(A)P(B)$; refinement lets you localize *where* the dependency lives — dependency can appear in only one refined sub-component (A_1 depends on B) while another (A_2) stays independent of B .
- **Mutually exclusive events**: $P(A \cap B) = 0$ and $P(A \cup B) = P(A) + P(B)$ still hold structurally; under Off-Probability, two events can be mutually exclusive in outcome-space yet have one Under-valued and one Over-valued weight, so their union’s total can numerically look smaller, larger, or even negative relative to either part — a genuinely new phenomenon with no classical analogue.
- **Conditional events**: $P(A|B) = P(A \cap B)/P(B)$ is preserved as a formula in all three variants; under refinement each $P(A_i|B)$ is computed against the same B and must still

sum to $P(A|B)$; under Off-Probability, $P(B)$ itself may be negative or >1 , so the *division* must be defined carefully (sign flips are possible and are treated as legitimate rather than as a domain error).

- **De Morgan's Rule:** $(A \cup B)^c = \bar{A}^c \cap \bar{B}$ and $(A \cap B)^c = \bar{A}^c \cup \bar{B}$ continue to hold set-theoretically (De Morgan is a statement about sets, not about the numbers assigned to them) in all three variants; what changes is only the *numeric* values attached to each side, which must still be verified to agree once the relaxed/refined probabilities are substituted.
- **Bayes' theorem:** $P(A|B) = P(B|A)P(A)/P(B)$ is retained as the *formula*; refinement distributes it: $P(A_i|B) = P(B|A_i)P(A_i)/P(B)$ for each i , and Over/Under/Off allows $P(B)$ in the denominator to be negative or >1 , meaning “posterior” values can legitimately fall outside $[0,1]$ too — this is the probabilistic analogue of allowing a fitted score to represent “worse than absolute impossibility” or “more certain than certainty,” and is interpreted, not discarded.

1.4 Relations to other notions

RCP is the special case of Refined Imprecise Probability (Family 2) where the interval collapses to a point; OUO-CP is the special case of Neutrosophic Probability (Family 3) where the Indeterminacy component I is fixed at 0. Because Classical Probability is the base case of every other family in this document, every family's “restrict $I=0$, restrict interval-width=0” degenerate case reduces back to some variant of Family 1 — this is the anchor consistency check for the whole framework.

1.5 Connections and distinctions with other fields

Frequentist and classical statistical inference. Classical Probability (Family 1) is exactly the object of Kolmogorov's axiomatic probability theory that underlies frequentist hypothesis testing, confidence intervals, and maximum-likelihood estimation. The Refined variant maps onto stratified sampling and partition refinement already standard in survey design; the Over/Under/Off variant has no frequentist analogue at all, since frequentist probability is definitionally bounded in $[0,1]$ by the relative-frequency interpretation — OUO-CP is a genuinely non-frequentist construction, closer in spirit to a scoring rule or utility than to a probability measure, which is why its off-interval values are read as *penalties* or *bonuses* rather than as probabilities of anything.

Decision theory and expected utility. Von Neumann–Morgenstern expected-utility theory combines probabilities with utilities multiplicatively; OUO-CP's Over/Under values can be seen as a probability and a utility already partially fused into one number, which is useful exactly when an analyst wants a single reported quantity to carry both “how likely” and “how good/bad” information without maintaining two separate scales — at the cost of losing the clean separability that expected-utility theory otherwise guarantees.

Actuarial pricing and gambling theory. The sports-odds example above is a direct instance of the actuarial “loading” concept: insurers and bookmakers alike price risk at a premium above the fair (probability-weighted) value, and the overround/vig literature in gambling mathematics is arguably the oldest applied use of what this document formalizes as Over-

Probability, predating the neutrosophic/plithogenic naming by decades even though it was never previously organized as a refinement of the probability axioms themselves.

Family 2 — Imprecise Probability → Refined / Over-Under-Off / Over-Under-Off-Refined

2.1 Definitions

Imprecise Probability (background). Instead of a point value $P(A)$, imprecise probability (Walley-style) assigns an interval $[\underline{P}(A), \bar{P}(A)] \subseteq [0,1]$ (lower/upper probability), used when data is too sparse or too ambiguous to justify a single number.

Refined Imprecise Probability (RIP). Applying the refinement-of-uncertain-sets construction [12] to an interval rather than a point, the event A is refined into $A_1, \dots, A_p$, each with its own interval $[\underline{P}(A_i), \bar{P}(A_i)]$, with the reduction constraint applied component-wise to the bounds:

$$\underline{P}(A) = \sum_i \underline{P}(A_i), \bar{P}(A) = \sum_i \bar{P}(A_i)$$

(subject to the usual imprecise-probability coherence axioms holding at each refined level independently).

Over/Under/Off-Imprecise Probability (OUO-IP). The bounds themselves are allowed outside $[0,1]$: $\underline{P}(A)$ may be negative (an interval whose lower bound represents “net-negative propensity”), $\bar{P}(A)$ may exceed 1 (an interval whose upper bound represents “more than certain, given redundant/reinforcing evidence”). The width $\bar{P}(A) - \underline{P}(A)$ is still required to be ≥ 0 — imprecision cannot itself be negative even though the bounds it wraps can be off-interval.

Over/Under/Off-Refined Imprecise Probability (OUO-RIP). Both relaxations at once: refine A into $A_1, \dots, A_p$, and let each $[\underline{P}(A_i), \bar{P}(A_i)]$ range outside $[0,1]$.

2.2 Real example/application

A **medical diagnostic panel** estimates the probability that a patient has condition A from three partially conflicting biomarker tests. Because the tests disagree, imprecise probability reports $[\underline{P}(A), \bar{P}(A)] = [0.40, 0.75]$ rather than a point estimate.

- *Refined*: condition A is refined into A_1 = “early-stage form” and A_2 = “late-stage form,” each carrying its own interval, e.g. $[0.10, 0.25]$ and $[0.30, 0.50]$, summing bound-wise to $[0.40, 0.75]$.
- *Over/Under/Off*: a fourth, contradicting test is added whose signal is so strong and so at odds with clinical plausibility that the lower bound is pushed to -0.05 (interpreted as “evidence actively pointing to absence, beyond mere zero-probability”) while the upper bound is pushed to 1.10 (interpreted as “the combined redundant-positive signals exceed what a single certain test would give,” e.g. because of an assay artifact that is itself informative).

- *Over/Under/Off-Refined*: the early-stage sub-interval alone absorbs the artifact, becoming $[-0.05, 0.30]$, while the late-stage sub-interval stays a normal $[0.30, 0.50]$ — localizing exactly where the anomalous evidence lives.

A second application — engineering tolerance and sensor calibration. A pressure sensor’s manufacturer specifies its reading as accurate to within $\pm 2\%$, so a reported reading of 100 kPa is really an imprecise-probability statement about the true pressure lying in $[98, 102]$ kPa with high confidence. *Refined*: the tolerance band is refined into “calibration drift” ($\pm 1\%$) and “thermal noise” ($\pm 1\%$), each independently correctable — calibration drift by recalibration, thermal noise only by better shielding. *Over/Under/Off*: a sensor flagged as faulty (post-recall) can report an interval that extends below the physically possible minimum (e.g. a lower bound of -5 kPa on a gauge that cannot physically read negative), which the maintenance team deliberately keeps in the log as a diagnostic signal rather than clamping it to 0, since the impossible reading itself is evidence of the fault type.

A third application — legal standards of proof. Civil litigation’s “preponderance of the evidence” standard ($>50\%$ likelihood) and criminal law’s “beyond a reasonable doubt” standard (colloquially placed anywhere from 90% to 99%) are both, in practice, imprecise-probability judgments — no court ever computes a literal point probability. *Refined*: “reasonable doubt” is refined by evidentiary category into “physical evidence doubt” and “testimonial evidence doubt,” since appellate courts often weigh the two differently on review. *Over/Under/Off*: a jury instruction that describes the defense’s case as having created doubt “beyond what the evidence could possibly support” is, functionally, describing an Under-Probability region — the prosecution’s case is treated as having a credibility interval that dips below what a merely unconvincing case would occupy.

2.3 Event-type study

- **Complementary events**: $[\underline{P}(\bar{A}), \bar{P}(\bar{A})] = [1 - \bar{P}(A), 1 - \underline{P}(A)]$ in the classical imprecise case; this “bound-flip” rule is preserved under refinement (applied to each A_i) and under *Over/Under/Off* (an upper bound >1 produces a complementary lower bound <0 , again treated as meaningful).
- **Independent events**: imprecise probability has no single test for independence (there are several: “epistemic,” “strong,” “type-1/type-2”); Refined-IP requires the chosen independence notion to hold at every refined sub-level simultaneously for the events to be called independent overall — a strictly stronger requirement than in the unrefined case.
- **Dependent events**: as in Family 1, refinement lets dependency be localized to a subset of the refined components.
- **Mutually exclusive events**: $\bar{P}(A \cap B) = 0$ (both bounds collapse to the point 0); under *Off-IP* this can happen while $\underline{P}(A)$ or $\underline{P}(B)$ is itself negative, so exclusivity of *outcomes* is compatible with one of the marginal intervals dipping below zero.
- **Conditional events**: the generalized Bayes rule for imprecise probability, $\underline{P}(A|B) = \underline{P}(A \cap B) / \bar{P}(B)$ and $\bar{P}(A|B) = \bar{P}(A \cap B) / \underline{P}(B)$, is retained; refinement distributes it across A_i ; *Off-IP* requires care with sign when $\underline{P}(B)$ or $\bar{P}(B)$ is negative, since dividing by a negative bound reverses the inequality direction between the resulting lower/upper conditional bounds.

- **De Morgan's Rule:** holds at the level of the underlying event algebra exactly as in Family 1; what is refined/relaxed is only the numeric bounds attached to each side.
- **Bayes' theorem:** the generalized (interval) Bayes rule above is applied component-wise under refinement and is extended to off-interval bounds under OUO — the interpretation is that posterior imprecision can itself widen beyond $[0,1]$ when the likelihood evidence is off-interval, which is read as “the evidence is inconsistent enough that no probability value, precise or bounded, adequately represents it without going outside the normal scale.”

2.4 Relations to other notions

Imprecise Probability differs from Neutrosophic Probability (Family 3) in that it tracks only an interval of *belief* about a single truth-status, with no separate Indeterminacy component — Neutrosophic Probability can be seen as Imprecise Probability plus an explicit, independently-varying I-component. RIP collapses to RCP (Family 1) when $\underline{P}(A) = \bar{P}(A)$ for every refined component. OUO-IP collapses to OUO-CP when the interval width is forced to 0.

2.5 Connections and distinctions with other fields

Walley's behavioral theory of imprecise probability. The $[\underline{P}, \bar{P}]$ interval notation used throughout this family follows Peter Walley's behavioral/coherence-based imprecise-probability theory, in which the gap $\bar{P} - \underline{P}$ represents the price spread a rational agent would accept between buying and selling a bet on A — a genuinely different foundational story from the neutrosophic Indeterminacy component I in Family 3, which represents unresolved evidence rather than a betting-market spread. The two produce numerically similar-looking objects but answer different questions, and conflating them is one of the more common errors when moving between the imprecise-probability and neutrosophic literatures.

Robust Bayesian statistics. Robust Bayesian analysis studies how posterior conclusions vary across a whole *class* of prior distributions rather than committing to one prior; the resulting “posterior envelope” is structurally identical to an imprecise-probability interval, and Refined Imprecise Probability's sub-interval decomposition corresponds directly to partitioning that prior class by the source of prior uncertainty (elicitation disagreement vs. genuine model uncertainty).

Interval arithmetic in engineering and metrology. The sensor-calibration example connects directly to the engineering discipline of interval arithmetic and uncertainty propagation (as formalized in the Guide to the Expression of Uncertainty in Measurement, GUM), where every measured quantity carries an explicit tolerance band that propagates through downstream calculations by interval rules essentially identical to the bound-flip and bound-multiplication rules given in §2.3.

Legal epistemology. The standards-of-proof example touches a long-running debate in legal philosophy over whether “beyond a reasonable doubt” is even meaningfully quantifiable as a probability at all (the “naked statistical evidence” problem); imprecise probability offers a middle path — courts can be read as implicitly reasoning over an interval or a qualitative ordering rather than a point value, without committing to the stronger and more contested claim that verdicts are literally probability calculations.

Family 3 — Neutrosophic Probability → Refined / Over-Under-Off / Over-Under-Off-Refined

3.1 Definitions

Neutrosophic Probability (background). An event A is assigned a triple $NP(A) = (T, I, F)$, Truth/chance-of-occurring, Indeterminacy, Falsity/chance-of-not-occurring, each nominally in $[0,1]$, with $T+I+F$ not required to equal 1 (unlike intuitionistic fuzzy sets), reflecting that evidence for, against, and unresolved about an event need not be complementary — the probabilistic reading of the neutrosophic measure/integral framework [13].

Refined Neutrosophic Probability (RNP). Following the 2013 refinement of the neutrosophic components into sub-components [2], each of T, I, F is split into sub-components: $T_1, \dots, T_p$ (different *kinds* of evidence-for), $I_1, \dots, I_r$ (different *kinds* of unresolved status — e.g. “unresolved due to missing data” vs. “unresolved due to contradictory data”), $F_1, \dots, F_s$ (different *kinds* of evidence-against), with

$$T = \sum T_i, I = \sum I_j, F = \sum F_k.$$

Over/Under/Off-Neutrosophic Probability (OUO-NP). Per the Neutrosophic Overset/Underset/Offset framework [3, 4], T, I , and/or F individually allowed outside $[0,1]$ — e.g. $F > 1$ to express “actively, redundantly disconfirmed beyond simple falsity,” or $T < 0$ to express “evidence that, if anything, argues the opposite direction of mere absence of support.”

Over/Under/Off-Refined Neutrosophic Probability (OUO-RNP). Both at once: refine each of T, I, F and allow the refined sub-components (and hence the totals) to sit outside $[0,1]$.

3.2 Real example/application

A criminal-trial evidence assessment: $NP(\text{guilty}) = (T=0.55, I=0.30, F=0.15)$ — more evidence toward guilt than innocence, but substantial indeterminacy from ambiguous forensic results.

- *Refined:* I is refined into $I_1 =$ “indeterminate due to a contaminated sample” (0.20) and $I_2 =$ “indeterminate due to conflicting witness testimony” (0.10), summing to 0.30 — separating *why* the case is unresolved, which matters for what further evidence would resolve it.
- *Over/Under/Off:* a coerced confession is later ruled fully inadmissible and, per courtroom procedure, is treated as actively undermining the prosecution’s broader credibility; the panel scores $F = 1.10$ — “beyond straightforward falsity,” a penalty against the case as a whole, not just against that one item of evidence.
- *Over/Under/Off-Refined:* $F_1 =$ “forensic falsity” (0.15, the original in-range component) and $F_2 =$ “credibility-penalty falsity from the tainted confession” (0.95, off-range), summing to $F = 1.10$.

A second application — petroleum reserve estimation. Oil and gas reserves are classically categorized as “proved,” “probable,” and “possible” — a three-way indeterminacy structure that

maps almost directly onto (T, I, F). NP(commercially recoverable) = (T=0.60 “proved,” I=0.25 “probable,” F=0.15 “assessed as unlikely to be economic”). *Refined*: I is refined into I_1 = “probable, pending further seismic survey” and I_2 = “probable, pending a pilot well test,” since the two require different capital commitments to resolve. *Over/Under/Off*: a field with an unusually productive analog nearby can be scored $T > 1$ relative to the standard reserve-certification scale, reflecting evidence that exceeds what the certification process normally treats as fully proved — common in practice but rarely formalized, since reserve auditors otherwise have no vocabulary for “more certain than proved.”

A third application — geopolitical intelligence analysis. Intelligence assessments are conventionally hedged with words of estimative probability (“likely,” “almost certainly”) that intelligence agencies have long recognized hide substantial indeterminacy. NP(state will take military action) = (T=0.35, I=0.45, F=0.20) captures a genuinely unresolved assessment better than a single point estimate would. *Refined*: I is refined into I_1 = “indeterminate due to denial-and-deception efforts by the target state” and I_2 = “indeterminate due to genuinely unformed decision-making within the target state’s own leadership” — a distinction with direct operational consequences, since only the first is addressable by better collection. *Over/Under/Off*: a source later revealed to be a double agent retroactively pushes F for every report that source contributed to past 1 — “worse than merely wrong,” since the source’s reporting was adversarially engineered to mislead, which is qualitatively different from an honest source simply being incorrect.

3.3 Event-type study

- **Complementary events:** $NP(\bar{A}) = (F, I, T)$ — truth and falsity swap, indeterminacy stays fixed. This swap rule is exactly preserved under refinement (swap the T- and F-refined families) and under Over/Under/Off (an over-valued F becomes an over-valued T for the complement, again treated as meaningful).
- **Independent events:** the neutrosophic independence test is applied component-wise, $T(A \cap B) = T(A)T(B)$, $F(A \cup B) = F(A) + F(B) - F(A)F(B)$ (De Morgan-consistent falsity combination), with I combined via a chosen indeterminacy-fusion rule; refinement requires this to hold at every T_i, F_k sub-level; Off-NP allows the products/combinations to land outside [0,1], which is retained rather than truncated.
- **Dependent events:** as before, refinement can localize dependency to specific evidence-sub-types (e.g. T_1 from ballistics is dependent on event B, but T_2 from motive is not).
- **Mutually exclusive events:** $T(A \cap B) = 0$ while $I(A \cap B)$ can remain nonzero (two events can be truth-exclusive yet share indeterminate overlap, e.g. both could turn out true if the ambiguous evidence resolves a certain way) — a distinctly neutrosophic phenomenon absent from Families 1–2; refinement can show *which* I-sub-component carries that shared ambiguity.
- **Conditional events:** $NP(A|B) = (T(A|B), I(A|B), F(A|B))$ computed componentwise via the neutrosophic Bayes analogue below; refinement distributes across sub-components of A while conditioning on the same B; Off-NP allows conditional T or F to leave [0,1] when the conditioning event B itself has an off-range component.
- **De Morgan’s Rule:** applied to each of T, I, F separately — $(A \cup B)^c$ and $\bar{A} \cap \bar{B}$ must match component-by-component (T-values, I-values, and F-values all reconciled independently), which is a strictly stronger requirement than classical De Morgan’s

single-valued check, and is inherited unchanged by all refined/relaxed variants (only the numeric verification, not the rule, is affected).

- **Bayes' theorem:** the neutrosophic form $T(A|B) = T(B|A)T(A)/T(B)$ (and analogously structured relations for I and F, since I and F do not follow the same rule as T unless special independence assumptions are made) is applied component-wise under refinement; under Off-NP, a denominator $T(B) < 0$ or > 1 again forces careful, but well-defined, sign handling in the posterior.

3.4 Relations to other notions

Neutrosophic Probability strictly generalizes Imprecise Probability (Family 2): setting $I(A) = \bar{P}(A) - \underline{P}(A)$, $T(A) = \underline{P}(A)$, $F(A) = 1 - \bar{P}(A)$ recovers an imprecise-probability interval as a special neutrosophic triple. It strictly generalizes intuitionistic-fuzzy-style probability (where I is forced to equal $1 - T - F$) by letting I vary independently — the same generalization result established, alongside Pythagorean, spherical, and n-hyperspherical fuzzy sets, in [10]. RNP and OUO-NP are the direct ancestors used to define Refined Plithogenic Probability (Family 4), since Plithogenic Probability is itself built by attaching a full attribute-value structure (of which neutrosophic (T,I,F) is the simplest 3-valued case) to each event [6, 9].

3.5 Connections and distinctions with other fields

Dempster-Shafer evidence theory and its Dezert-Smarandache extension. The closest established analogue to Neutrosophic Probability is Dempster-Shafer belief-function theory, in which a mass function assigns weight to *subsets* of the frame of discernment, and belief/plausibility functions bracket the true probability from below and above — structurally similar to how T and $1 - F$ bracket the neutrosophic truth-value. The distinction is that Dempster-Shafer's "unassigned mass" plays a role similar to I but is derived from the mass assignment rather than tracked as an independent primitive, and Dempster's rule of combination handles conflicting evidence very differently (and more restrictively) than the neutrosophic AND/OR operators used in §3.3's independence test. Dempster-Shafer theory was itself extended by Dezert and Smarandache into the Dezert-Smarandache Theory (DSmT) of plausible, paradoxist, and neutrosophic reasoning specifically to handle highly conflicting sources and frames of discernment that cannot be cleanly refined — the same "unrefinable frame of discernment" problem that motivates this document's own Refinement operator, approached from the information-fusion side rather than the set-theoretic side [14]. DSmT's hybrid and free DSm models, and its Proportional Conflict Redistribution combination rules, are the direct mathematical ancestors of the neutrosophic evidence-combination approach generalized throughout Family 3, making DSmT arguably the closest single prior framework to this document's Neutrosophic Probability family, closer even than classical Dempster-Shafer theory itself.

Paraconsistent and three-valued logics. Neutrosophic logic's willingness to let $T+I+F$ exceed or fall short of 1 places it in the broader family of paraconsistent logics (which tolerate local contradictions without trivializing the whole system) and three-valued logics such as Kleene's and Łukasiewicz's, which also use a third "unknown/indeterminate" truth value alongside true and false — neutrosophic logic's distinguishing move is making that third value a continuously-varying degree rather than a single discrete state.

Risk analysis: aleatory vs. epistemic uncertainty. Modern quantitative risk analysis distinguishes aleatory uncertainty (irreducible randomness, e.g. which face a fair die lands on) from epistemic uncertainty (reducible ignorance that better data or models could resolve). The Indeterminacy component **I** is best read as a measure of epistemic uncertainty specifically — refining **I** (as in the petroleum and intelligence-analysis examples above) is, in this vocabulary, exactly the exercise of decomposing epistemic uncertainty into its addressable sources, a standard step in modern probabilistic risk assessment (PRA) frameworks used in nuclear safety and aerospace engineering.

3.6 Extension: Quadruple (T, I, N, F) Neutrosophic Measure, Probability, and Statistics

Definition. In 2017 Smarandache refined the Indeterminacy component itself, splitting it into **Pure Indeterminacy I** (ambiguous or unresolved change/status, excluding neutrality) and **Neutrality N** (confirmed absence of change or genuine impartiality), yielding the quadruple $(T, I, N, F) \in P([0,1])^4$ with

$$0 \leq \inf(T) + \inf(I) + \inf(N) + \inf(F) \leq \sup(T) + \sup(I) + \sup(N) + \sup(F) \leq 4,$$

extended in the same source explicitly to Quadruple Neutrosophic Measure, Quadruple Neutrosophic Probability, and Quadruple Neutrosophic Statistics, not merely to sets and logic [15]. The distinction between **I** and **N** matters operationally: **N** supports a decision to *stop* monitoring a factor (it is confirmed unchanged), while **I** supports a decision to *keep* monitoring it (its direction is still unresolved) — conflating the two, as the ordinary Family 3 triple does by folding both into one **I**, blurs a distinction with real consequences. Reference [15] illustrates this with a soccer referee who is either genuinely neutral ($N = 1$) or partially biased ($N < 1$) regardless of which team is “winning” the T/F contest, and with neutral countries in a conflict whose degree of true neutrality can itself vary continuously rather than being simply on or off.

Refined Quadruple (T, I, N, F). Applying Family 3’s refinement operator to all four components at once (rather than to T, I, F alone) gives $(T_1, \dots, T_p; I_1, \dots, I_q; N_1, \dots, N_r; F_1, \dots, F_s)$, each subcomponent indexed to a distinct medium factor, organ, trait, or module of the entity under study [15, 16].

Over/Under/Off-Quadruple. Any of T, I, N, or F may individually exceed 1 or fall below 0, extending Family 3’s Over/Under/Off-Neutrosophic Probability to all four components rather than three [15].

Operators and ranking. T and N are treated as positive-quality components (combined via t-norm under conjunction), while I and F are treated as negative-quality components (combined via t-conorm under conjunction) [15] — note this is *not* simply “T,N get the Family 3 T-treatment and I,F get the Family 3 F-treatment,” since Neutrality is a confirmed-good outcome to have, not a hedge. Negation is $\neg(T, I, N, F) = (F, N, I, T)$, a direct extension of Family 3’s $(T, I, F) \rightarrow (F, I, T)$ swap with a matching $I \rightarrow N$ swap alongside the $T \rightarrow F$ swap. Because a quadruple has no single natural total order, [15] supplies a four-function ranking cascade for MCDM use — Score $s(T, I, N, F) = (2 + T + N - I - F)/4$, Accuracy $a = T + N - F$, Extended Certainty $ec = T + N$, and (non-extended) Certainty $c = T$, applied in that order as successive tiebreakers — letting quadruple-valued alternatives be ranked exactly as Family 4’s contradiction-degree fusion ranks plithogenic alternatives.

Application: the Refined (T, I, N, F) Neutrosophic Theory of Evolution. The richest worked application of this extension recasts change itself — biological, technical, or social — as a quadruple-valued measure: an entity’s response to changing conditions is described not by a single “improved / unchanged / worsened” judgement but by an Evolution vector

$$E(X) = (T_1, \dots, T_p \mid I_1, \dots, I_q \mid N_1, \dots, N_r \mid F_1, \dots, F_s)$$

refined by medium factor — organ, trait, locus, or module — with T (Evolution) and F (Involution) as the adaptive/regressive poles, N (Neutrality) as confirmed stasis, and I (Indeterminacy) as unresolved direction of change [16]. Reference [16] surveys this structure qualitatively across fourteen fields — medicine, agriculture, biology, botany, ecology, computer science, paleontology, genetics, biogeography, coevolution, anthropology, immunology, molecular biology, and embryology — and works two of them through numerically:

- *Medicine*: a metabolic-syndrome patient’s six-month treatment regimen, refined by organ system: T_1 = cardiac function (0.70), T_2 = renal function (0.40), T_3 = HbA1c reduction (0.80); I_1 = liver-enzyme trend (0.50), I_2 = gut-microbiome shift (0.60); N_1 = bone density (0.90), N_2 = cognitive score (0.85); F_1 = pulmonary function (0.30), F_2 = sleep quality (0.45). Averaging this into a single “improved” verdict would erase the clinically important fact that two organ systems (I_1 , I_2) remain genuinely unresolved and warrant continued monitoring, rather than being folded into either the evolution or involution total.
- *Agriculture*: two wheat cultivars compared across seven medium factors reveal a drought/cold trade-off invisible to a single yield score — Cultivar A dominates on drought resistance and yield ($T_1 = 0.85$, $T_2 = 0.60$) but cedes cold tolerance and lodging resistance ($F_1 = 0.40$, $F_2 = 0.25$) to the more balanced Cultivar B ($T_1 = 0.55$, $T_2 = 0.65$, $F_1 = 0.70$, $F_2 = 0.60$).

Weighted aggregation for decision support. For cases needing a single comparative score, [16] proposes

$$S(X) = [w_T \text{avg}(T) + w_N 0.5 \text{avg}(N) - w_F \text{avg}(F)] / (w_T + w_N + w_F),$$

deliberately excluding I from the score, since unresolved evidence should not be scored as either good or bad news until it actually resolves. The weights w_T , w_N , $w_F \geq 0$ let an analyst encode how much further evolution, confirmed stability, and involution-avoidance are each valued — applied to the two cultivars above, Cultivar A wins under evolution-weighted priorities and Cultivar B overtakes it once involution-avoidance is weighted heavily enough that the cold-tolerance and lodging-resistance gap outweighs the drought-resistance advantage, illustrating that “the better entity” under this framework is a function of stated priorities rather than a fixed property of the data — the same lesson Family 4’s fusion-operator caveat (§4.3, and the Part VII FAQ) draws for contradiction-degree weighting.

Relation to the rest of this document. Quadruple (T, I, N, F) Neutrosophic Probability, Measure, and Statistics sits as a direct refinement of Family 3 (splitting I into I and N) that is then, in its own right, refined again by medium factor in the sense of Part I’s Refinement operator — making it, structurally, a refinement-of-a-refinement, and a natural bridge toward Family 4’s Plithogenic Probability: a medium-factor-indexed evolution vector is itself close to a plithogenic attribute-value structure, with T and N playing the role of “positive” attribute-

values and I and F the “negative” ones. The weighted composite score $S(X)$ is, in this light, a specific, simplified instance of the general contradiction-degree-weighted plithogenic fusion operator introduced in §4.1, restricted to the case where T and N share one polarity, F carries the opposite polarity, and I is excluded from fusion entirely rather than weighted by a contradiction degree — a useful worked illustration of how much simpler a fusion rule can be made once an application’s structure (here, four cleanly-separated components rather than an arbitrary attribute-value set) permits it.

Family 4 — Plithogenic Probability → Refined / Over-Under-Off / Over-Under-Off-Refined

4.1 Definitions

Plithogenic Probability (background). Built on Smarandache’s founding Plithogeny/Plithogenic Set, Logic, Probability, and Statistics framework [6], further developed as a special issue collecting neutrosophic and plithogenic extensions [9] and applied to optimization under uncertainty [7], Plithogenic Probability generalizes neutrosophic probability by attaching to event A not just one (T,I,F) triple but a whole **attribute** with a set of **attribute-values** $v_1, \dots, v_n$, each carrying its own degree-of-appurtenance $d(A, v_i) \in [0,1]$ (or a full (T,I,F) if the attribute-value itself needs neutrosophic treatment), plus a **contradiction degree** $c(v_i, v_{\text{dominant}}) \in [0,1]$ measuring how opposed each value is to the dominant/reference value, and a fusion operator that combines them weighted by $(1-c)$.

Refined Plithogenic Probability (RPP). Applying the refinement operator [1, 2] to a plithogenic attribute, each attribute-value v_i is itself refined into sub-values $v_{i1}, \dots, v_{iq}$ (e.g. the attribute “severity” refined from a single value into “severity-to-property” and “severity-to-persons”), each with its own degree and contradiction degree, with the parent degree recovered by the family’s fusion rule applied to the refined sub-degrees.

Over/Under/Off-Plithogenic Probability (OUO-PP). Degrees of appurtenance $d(A, v_i)$ and/or contradiction degrees $c(v_i, v_{\text{dominant}})$ allowed outside $[0,1]$ — e.g. a contradiction degree > 1 expressing “actively counter-supporting, beyond simple opposition,” which flips the sign of that attribute-value’s contribution in the fusion sum rather than merely down-weighting it.

Over/Under/Off-Refined Plithogenic Probability (OUO-RPP). Both: refine the attribute-values and allow the refined degrees/contradiction-degrees to be off-interval.

4.2 Real example/application

Insurance underwriting risk for a commercial property, attribute = “risk factors,” values = {fire risk, flood risk, theft risk}, dominant value = fire risk (most weighted in the policy’s base rate), with degrees $d = (0.30, 0.55, 0.15)$ and contradiction degrees relative to fire risk $c = (0, 0.6, 0.3)$.

- *Refined*: flood risk is refined into “riverine flood” (0.35) and “storm-surge flood” (0.20), summing to 0.55, each carrying its own (possibly different) contradiction degree relative to fire risk.
- *Over/Under/Off*: the property sits on a floodplain that was *just* rezoned as protected wetland, which paradoxically reduces future flood risk *and* actively lowers fire risk (wetland buffering); the underwriting model expresses this cross-cutting benefit as a contradiction degree $c(\text{flood}, \text{fire}) = -0.20$ (an “under” value): flood risk here doesn’t just fail to correlate with fire risk, it actively co-varies with a reduction in it.
- *Over/Under/Off-Refined*: the riverine-flood sub-value keeps the ordinary $c = 0.55$, while only the storm-surge sub-value carries the anomalous $c = -0.20$ tied specifically to the wetland-buffering effect.

A second application — supplier selection in supply chain management. A procurement team scores candidate suppliers on attribute “supplier viability” with values {cost, quality, delivery reliability}, dominant value = cost, degrees $d = (0.40, 0.35, 0.25)$, contradiction degrees $c = (0, 0.5, 0.4)$. *Refined*: quality is refined into “defect rate” and “certification compliance,” since the two are audited through entirely different processes (statistical sampling vs. paperwork review). *Over/Under/Off*: a supplier who offers a below-market price *and* has been caught in a past contract dispute is scored with a contradiction degree between cost and delivery-reliability of $c > 1$, reflecting that the low price is treated as actively predictive of future delivery problems, not merely uncorrelated with them — a red flag the ordinary [0,1] scale cannot distinguish from simple independence.

A third application — multi-asset portfolio risk. A portfolio manager scores holdings on attribute “risk exposure” with values {market risk, credit risk, liquidity risk}, dominant value = market risk, degrees $d = (0.50, 0.30, 0.20)$. *Refined*: credit risk is refined into “issuer default risk” and “counterparty risk” (relevant for derivatives positions specifically), each hedged through different instruments. *Over/Under/Off*: during a liquidity crisis, the contradiction degree between liquidity risk and market risk can spike above 1, since the two stop behaving as merely correlated and instead reinforce each other in a self-amplifying fire-sale dynamic — exactly the kind of regime-dependent behavior that fixed, [0,1]-bounded correlation coefficients in classical portfolio theory (e.g. Markowitz mean-variance optimization) are known to underestimate during crises.

4.3 Event-type study

- **Complementary events**: complementation acts on the dominant value’s degree, $d(\bar{A}, v_{\text{dominant}}) = 1 - d(A, v_{\text{dominant}})$, while every non-dominant attribute-value’s contradiction degree is preserved (contradiction is a relation between values, not a truth-status, so it does not flip); refinement and Off-PP apply the Family 3 pattern to whichever attribute-values are treated as carrying T/F-like status.
- **Independent events**: independence must be checked value-by-value across the whole attribute set, not just once, since two events can be independent with respect to one attribute-value (e.g. fire risk) but dependent with respect to another (e.g. theft risk, if both properties are in the same neighborhood); refinement makes this multi-level check even finer-grained.
- **Dependent events**: as above, plithogenic dependency is inherently attribute-value-local; this is the family’s signature feature relative to Families 1–3.

- **Mutually exclusive events:** two events can be mutually exclusive with respect to the dominant attribute-value while overlapping on a non-dominant one — e.g. two properties cannot both be “the one that burns down” in a single-claim scenario (dominant-value exclusive) while both can simultaneously carry nonzero flood-risk exposure from the same storm system (non-dominant-value non-exclusive).
- **Conditional events:** conditioning is applied per attribute-value, $d(A|B, v_i)$, and the whole fused conditional probability is recomputed by re-running the plithogenic fusion operator on the conditioned degrees; refinement distributes this across refined sub-values; Off-PP allows a conditional degree to leave $[0,1]$ when the conditioning event has an off-range degree on that attribute-value.
- **De Morgan’s Rule:** verified separately for every attribute-value’s degree; because the fusion operator is generally not a simple sum, De Morgan-consistency must be checked against the *specific* fusion operator chosen for the model (average, min/max-weighted, or contradiction-weighted sum), which is the one place in this document where the rule’s *validity* — not just its numeric evaluation — depends on a modeling choice rather than following automatically.
- **Bayes’ theorem:** the plithogenic Bayes analogue applies the ordinary Bayes formula attribute-value-by-attribute-value and then re-fuses; refinement distributes it further into refined sub-values; Off-PP requires the same off-interval sign-handling as Family 3, now replicated across every attribute-value.

4.4 Relations to other notions

Plithogenic Probability strictly contains Neutrosophic Probability as the special case of a single attribute with exactly three values $\{T, I, F\}$ and contradiction degrees fixed at $c(T, \cdot)=0$, $c(F, \cdot)=1$, $c(I, \cdot)=0.5$ [6]. It strictly contains Classical and Imprecise Probability as further degenerate sub-cases (one attribute, one or two values). It is the most general of the four probability families in this document; Refined Plithogenic Statistics (Family 9, below) is built directly on top of it in the same way Classical Statistics is built on Classical Probability, and both have already been applied to decision-making under uncertainty in fields ranging from optimization [7] to pedagogy [8].

4.5 Connections and distinctions with other fields

Multi-criteria decision-making (MCDM). Plithogenic Probability’s attribute-value-with-contradiction-degree structure is a direct generalization of standard MCDM techniques such as the Analytic Hierarchy Process (AHP) and TOPSIS, both of which score alternatives on multiple weighted criteria and aggregate to a single ranking. The contradiction degree plays the role that criteria correlation/redundancy adjustments play in AHP consistency checking, but where AHP treats criteria weights as fixed inputs elicited once, the plithogenic framework treats the contradiction structure itself as refinable and potentially off-interval — letting the model represent criteria interactions (like the liquidity-crisis example above) that standard MCDM weighting schemes assume away.

Multi-attribute utility theory (MAUT). Keeney and Raiffa’s multi-attribute utility theory combines multiple utility dimensions into a single score via an aggregation function (additive, multiplicative, or more general forms) under specific independence conditions; the plithogenic fusion operator is structurally analogous but explicitly designed for cases where those

independence conditions fail, since the contradiction degree exists precisely to quantify how far any given pair of attribute-values is from the independence MAUT would otherwise assume.

Copula-based dependency modeling. In quantitative finance, copulas separate the marginal distribution of each risk factor from the dependency structure linking them, allowing tail dependency (crises where correlations spike) to be modeled explicitly rather than assumed constant — exactly the phenomenon the portfolio-risk example above expresses as an off-interval contradiction degree. The plithogenic contradiction degree is a coarser, single-number analogue of a copula’s dependency parameter, trading the copula’s full distributional detail for a simpler, more interpretable quantity that fits naturally into the attribute-value fusion framework.

Family 5 — Measure → Refined / Over-Under-Off / Over-Under-Off-Refined

5.1 Definitions

Measure (background). A measure μ on a σ -algebra Σ over set X assigns $\mu(A) \geq 0$ to each $A \in \Sigma$, with $\mu(\emptyset)=0$ and countable additivity $\mu(\bigcup A_i) = \sum \mu(A_i)$ for disjoint A_i . Probability is the special case $\mu(X)=1$, and the neutrosophic generalization of this measure-theoretic foundation is given in [13].

Refined Measure. $A \in \Sigma$ is refined into $A_1, \dots, A_p$, each carrying its own $\mu(A_i)$, with $\mu(A) = \sum \mu(A_i)$ recovered exactly as ordinary countable additivity already requires — so Refined Measure is, formally, simply the practice of *naming and separately tracking* an additive decomposition that ordinary measure theory already permits [1, 12], elevated here to a first-class notion so it can be composed with the Over/Under/Off operator below.

Over/Under/Off-Measure. The non-negativity axiom $\mu(A) \geq 0$ is dropped: μ can take negative values (a **signed measure**, in the pre-existing mathematical sense — this family formally coincides with signed/complex measure theory) and — going one step beyond the classical signed-measure literature, and directly following the 2026 extension of Masses in Information Fusion to Over/Under/Off-Masses [5], itself built on the mass-function machinery of Dezert-Smarandache Theory [14] — is also allowed to exceed any total-measure normalization bound that the model has set for $\mu(X)$, i.e. an “Over” region above the model’s own reference total, not merely “any nonnegative real.”

Over/Under/Off-Refined Measure. Refine A into $A_1, \dots, A_p$ and allow each $\mu(A_i)$ to be signed/off-bound, with $\mu(A) = \sum \mu(A_i)$ still holding as the reduction constraint.

5.2 Real example/application

An **environmental-impact measure** μ over regions of a watershed, where $\mu(A)$ represents net ecological benefit of region A , calibrated so the pristine baseline watershed has $\mu(X) = 1$.

- *Refined:* a wetland region A is refined into A_1 = “water-filtration contribution” (0.20) and A_2 = “habitat contribution” (0.15), summing to $\mu(A) = 0.35$.
- *Over/Under/Off:* an industrially degraded sub-region B has $\mu(B) = -0.40$ (net ecological harm, a signed measure in the classical sense); a newly restored wetland with unusually

effective invasive-species removal is scored $\mu(C) = 0.18$ against a regional baseline capacity of only 0.12 — i.e. it *exceeds* the modeled maximum benefit that region's baseline ecology was assumed capable of, an "Over" reading the classical signed-measure framework has no separate name for.

- *Over/Under/Off-Refined*: region B's harm is refined into $B_1 = \text{"soil contamination"}$ (-0.25) and $B_2 = \text{"runoff into downstream fisheries"}$ (-0.15), summing to -0.40 .

A second application — acoustic energy measurement in audio engineering.

A recording engineer measures $\mu(\text{track}) = \text{integrated loudness relative to a broadcast reference level, normalized so the reference level has } \mu = 1$. *Refined*: a mixed track's loudness is refined into $\mu(\text{vocals})$ and $\mu(\text{instrumental})$, summed via the standard loudness-summation formula for independent channels, letting the engineer isolate which stem is driving an overly hot mix.

Over/Under/Off: a mastered track that deliberately exceeds broadcast reference level (common in the "loudness war" era of commercial mastering) has $\mu > 1$ by design, while a quiet ambient interlude in the same album can legitimately be scored with a very small or even notionally negative relative loudness against a silence-referenced floor — engineers already work with signed decibel scales, so this family formalizes existing audio-engineering practice rather than introducing a new one.

A third application — economic value-added (EVA) accounting. A business unit's EVA measure $\mu(\text{unit}) = \text{net operating profit after tax minus a capital charge, calibrated so a unit exactly meeting its cost of capital has } \mu = 0$ (not 1, illustrating that the "normalized total" reference point is a modeling choice, not always $\mu(X)=1$). *Refined*: $\mu(\text{unit})$ is refined into $\mu(\text{core product line})$ and $\mu(\text{new product line})$, letting the CFO see that a division's positive aggregate EVA is entirely carried by the mature product and masks a value-destroying new launch.

Over/Under/Off: a business unit that generates strategic option value beyond its current cash flows (e.g. an R&D division building future capability) can be assigned μ above what its current-period accounting would justify, exactly mirroring how real-options valuation in corporate finance already treats certain investments as worth more than their discounted-cash-flow value alone would show.

5.3 Event-type study

- **Complementary events**: $\mu(A^c) = \mu(X) - \mu(A)$ holds by additivity in every variant; under *Over/Under/Off* this can push $\mu(A^c)$ negative when $\mu(A) > \mu(X)$, read the same way as in Family 1.
- **Independent events**: measure theory itself has no native independence notion (independence is a probability/measure-*product-space* notion); once μ is normalized to a probability (Family 1) independence reduces to that family's test — *Refined/Off-Measure* independence is therefore always discussed via its induced probability, not directly.
- **Dependent events**: same caveat as above.
- **Mutually exclusive events**: $\mu(A \cap B) = 0$ with $\mu(A \cup B) = \mu(A) + \mu(B)$ by disjoint additivity, preserved in all three variants, including when one of $\mu(A)$, $\mu(B)$ is negative.
- **Conditional events**: $\mu(A|B) := \mu(A \cap B) / \mu(B)$ is well-defined whenever $\mu(B) \neq 0$, including negative $\mu(B)$; refinement and *Off-Measure* carry through exactly as in Family 1.

- **De Morgan's Rule:** purely set-theoretic, unaffected by which measure variant is attached; numeric verification carries through component-wise under refinement.
- **Bayes' theorem:** the measure-theoretic Bayes/Radon–Nikodym relationship $d\mu(A|B) = d\mu(A \cap B)/d\mu(B)$ is retained; refinement and Off-Measure extend it exactly as in Families 1 and 2.

5.4 Relations to other notions

Measure is the common mathematical ancestor of Families 1–4: Classical Probability is a normalized measure; Imprecise Probability is a pair of measures (lower/upper envelope); Neutrosophic and Plithogenic Probability can each be represented as a vector or attribute-indexed family of measures. Over/Under/Off-Measure formally is signed-measure theory for its “Under” half; its novel contribution is the “Over” half (exceeding a model-relative baseline total) and the explicit refinement/off-composition machinery layered on top.

5.5 Connections and distinctions with other fields

Kolmogorov's measure-theoretic foundations of probability. Measure theory as formalized by Kolmogorov in 1933 is the mathematical bedrock underneath every other family in this document; Refined Measure is, as noted in §5.1, simply a first-class name for the countable-additivity decomposition measure theory has always permitted, while Over/Under/Off-Measure is the genuinely novel departure, since Kolmogorov's axioms require non-negativity ($\mu(A) \geq 0$) as a foundational assumption rather than a convenience.

Signed measures in physics. Physics has long worked with signed measures under other names: electric charge distributions, where positive and negative charge densities integrate to a net charge that can be zero, positive, or negative over any region, are formally a signed measure exactly in the sense of Family 5's Over/Under half. The “Over” half — exceeding a model-relative baseline total — has a rough physical analogue in metastable or supersaturated states, where a local quantity temporarily exceeds what equilibrium thermodynamics would predict as the ceiling.

National accounts and value-added accounting. The EVA example above connects to the broader field of national income accounting, where GDP is computed as a sum of value-added measures across sectors, each of which can in principle be negative (a sector that destroys more value than it creates, as sometimes alleged of certain financial-sector activities in the years before 2008) — an under-explored but methodologically well-precedented use of exactly the signed-measure structure this family formalizes.

Family 6 — Classical Statistics → Refined / Over-Under-Off / Over-Under-Off-Refined

6.1 Definitions

Classical Statistics (background). Descriptive and inferential procedures (mean, variance, hypothesis tests, confidence intervals, regression, etc.) applied to crisp numeric data drawn from a population, with no explicit tracking of ambiguity in category membership.

Refined Classical Statistics (RCS). A statistical category (e.g. “defective item,” “responder to treatment”) is refined into sub-categories before aggregation, and every classical statistic (mean, variance, proportion, etc.) is computed and reported both at the sub-category level and re-aggregated to the parent level, with the reduction constraint that parent-level counts/sums equal the sum of refined sub-category counts/sums.

Over/Under/Off-Classical Statistics (OUO-CS). Applies when the underlying scores being averaged/aggregated are themselves Over/Under/Off values (Family 1) rather than ordinary bounded scores — e.g. a mean computed over a sample that includes Under-valued (negative-contribution) and Over-valued (super-contribution) observations, which is arithmetically identical to an ordinary signed-data mean but is flagged as its own named variant here because the *interpretation* of the resulting statistic (e.g. “average net contribution can be negative or exceed the nominal maximum”) differs from classical statistics’ usual assumption that data lies in a bounded scale.

Over/Under/Off-Refined Classical Statistics (OUO-RCS). Both: refined sub-categories whose scores are themselves off-interval.

6.2 Real example/application

A call-center customer-satisfaction survey, score in $[0,10]$, mean reported as 7.4.

- *Refined*: “satisfied” responses (score ≥ 7) are refined into “satisfied-and-would-recommend” and “satisfied-but-neutral-on-recommending” sub-groups, each with its own mean and count, recombining to the parent mean 7.4.
- *Over/Under/Off*: a handful of respondents leave scores of -2 (used by the survey team to flag “actively hostile, filed a complaint,” a deliberately off-scale signal distinct from a legitimate 0) and 12 (“went out of their way to publicly praise the agent,” deliberately off-scale to distinguish it from a merely maximal 10); the reported mean of 7.4 is recomputed as 7.1 once these off-scale flags are included in the aggregate, and the OUO-CS framework treats this 7.1 as the honest statistic rather than clipping the -2 s to 0 and the 12s to 10 first (which would classically be standard practice, and is exactly the practice this family is designed to make optional and explicit rather than automatic).
- *Over/Under/Off-Refined*: the -2 respondents are refined into “complaint filed” (-2) vs. “complaint threatened but not filed” (-1), tracked separately within the aggregate.

A second application — Six Sigma defect-rate tracking. A manufacturing line reports a defect rate of 340 defects per million opportunities (DPMO), the standard Six Sigma statistic. *Refined*: the defect count is refined by defect type (dimensional, cosmetic, functional), each tracked against a separate root-cause corrective-action process, since a cosmetic-defect spike and a functional-defect spike trigger completely different response protocols. *Over/Under/Off*: a defect that causes a downstream safety incident is weighted in the internal quality scorecard at more than 1 “defect equivalent” (reflecting cost-of-quality weighting already standard in Six Sigma practice, where a field failure costs far more than a defect caught at the factory), while a defect caught and corrected before it ever reached a customer, by an unusually effective in-line inspection step, can be scored as contributing negatively to the rolling defect-rate trend, since its detection is itself evidence the process is improving.

A third application — A/B testing in software and UX. A product team runs an A/B test on a checkout flow, measuring conversion rate: control = 4.2%, variant = 4.6%. *Refined*: the conversion lift is refined by traffic source (organic vs. paid), since the variant might only be winning among paid traffic, which the aggregate lift would obscure. *Over/Under/Off*: a variant that increases short-term conversion but is later found to increase return-fraud rates can be scored with a net business-impact statistic below the pre-test baseline (an Under-value on the composite metric), even though the raw conversion-rate statistic alone remained positive — the OUO-CS framework is exactly the formal home for the “guardrail metric went negative even though the primary metric improved” pattern that experienced experimentation teams already watch for informally.

6.3 Event-type study

Because Families 6–9 concern *statistics* (estimators built from data) rather than *probabilities* (numbers attached to single events), the seven event-type properties are studied here as properties of the **underlying sampling events** that generate the data, not of the point estimates themselves:

- **Complementary events**: the complement of “sampled unit is in category A” is “sampled unit is in category A^c”; refined sub-category proportions sum to the parent proportion and its complement follows Family 1’s rule exactly, now applied to sample proportions instead of a single theoretical $P(A)$.
- **Independent events**: independence between two sampled attributes is tested via chi-square/contingency-table methods on the observed frequencies; Refined-CS requires the test to be re-run separately at each refined sub-category level, since independence at the aggregate level does not imply independence at the refined level (a form of Simpson’s-paradox-adjacent behavior that refinement is specifically useful for detecting).
- **Dependent events**: as above; OUO-CS additionally requires care because off-scale (negative/over-maximum) observations can distort a naive correlation coefficient, so a robust or truncated-comparison statistic is recommended alongside the raw one.
- **Mutually exclusive events**: sampled units falling into mutually exclusive categories partition the sample exactly, $\sum(\text{category counts}) = n$, preserved under refinement (sum of sub-category counts = parent count) and under OUO (an off-scale flag is still a mutually-exclusive *category label*, even though its associated *score* is off-interval).
- **Conditional events**: conditional statistics (e.g. mean satisfaction *given* the customer used live chat vs. phone) are computed by subsetting the sample; refinement and OUO carry through by simply subsetting the refined/off-scale data the same way.
- **De Morgan’s Rule**: applied to sample-category set operations exactly as in Family 1, now verified against observed counts rather than theoretical probabilities.
- **Bayes’ theorem**: applied in its empirical form (frequencies substituted for probabilities) to update category-membership beliefs from new sample batches; refinement and OUO extend it identically to Family 1, with the same off-interval sign-handling caveat.

6.4 Relations to other notions

Classical Statistics is the empirical/estimation counterpart of Classical Probability (Family 1); every relation noted in §1.4 applies here with “probability” replaced by “estimated proportion” or “sample statistic.” Refined Classical Statistics is the tool that, in practice, most directly motivates Refined Neutrosophic Statistics (Family 8) below, since real survey/measurement data routinely contains an explicit “unresolved/don’t know” category that classical statistics tends to discard and neutrosophic statistics is built specifically to retain.

6.5 Connections and distinctions with other fields

Statistical process control (SPC) and Six Sigma. The manufacturing example above sits squarely within SPC methodology, where control charts, defect-rate tracking, and root-cause analysis are standard tools; Refined Classical Statistics formalizes the common SPC practice of stratifying a control chart by defect type or shift, while OUO-CS’s cost-of-quality weighting matches the “cost of poor quality” (COPQ) accounting already used in mature Six Sigma programs, giving both an explicit mathematical home inside the broader refinement/relaxation framework.

Classical experimental design and hypothesis testing. Randomized controlled trials, A/B testing, and Fisher’s classical experimental design all rest on the same partition-and-compare logic that Refined Classical Statistics generalizes; the refined sub-category breakdown in the A/B testing example is exactly what experimentalists call a pre-registered subgroup analysis, and the field has its own well-developed cautions (multiple-comparisons correction, the danger of post-hoc subgroup mining) that apply with equal force here — refinement is a tool for asking sharper questions, not a license to skip the statistical discipline classical experimental design already requires.

Business guardrail metrics and OKR frameworks. The A/B-testing guardrail-metric example connects to the broader corporate-strategy practice of Objectives and Key Results (OKRs) and North-Star-metric frameworks, which explicitly pair a primary success metric with one or more guardrail metrics precisely to catch the “numerator went up, but something more important went down” failure mode — OUO-CS gives that widely-used practice a formal probabilistic vocabulary rather than leaving it as an organizational process convention.

Family 7 — Imprecise Statistics → Refined / Over-Under-Off / Over-Under-Off-Refined

7.1 Definitions

Imprecise Statistics (background). Descriptive/inferential statistics computed from interval-valued or set-valued data (each observation is an interval $[x_i_lo, x_i_hi]$ rather than a point, e.g. from imprecise measurement instruments), producing interval-valued statistics such as $[mean_lo, mean_hi]$.

Refined Imprecise Statistics (RIS). A category or observation is refined into finer intervals before aggregation, with the parent interval bounds recovered as the sum (for additive statistics) of the refined interval bounds, matching Family 2's reduction rule.

Over/Under/Off-Imprecise Statistics (OUO-IS). Interval bounds allowed outside $[0,1]$ or outside whatever the nominal scale is, exactly mirroring Family 2's relaxation, now applied to sample data rather than to a single theoretical interval probability.

Over/Under/Off-Refined Imprecise Statistics (OUO-RIS). Both together.

7.2 Real example/application

Industrial temperature-sensor logging where each sensor reports an interval $[T_{lo}, T_{hi}]$ per reading due to known instrument tolerance, e.g. mean logged interval $[21.4^{\circ}\text{C}, 22.1^{\circ}\text{C}]$.

- *Refined*: readings are refined by shift (day-shift sensor drift vs. night-shift sensor drift), each with its own interval mean, summing appropriately when re-aggregated on a weighted basis.
- *Over/Under/Off*: a subset of readings during a known sensor-fault window are logged with an interval whose lower bound is *known* to be physically impossible (e.g. $[-15^{\circ}\text{C}, 40^{\circ}\text{C}]$ on a line that never runs below 0°C), which the team deliberately keeps as an off-scale interval to flag “fault-window data, do not trust the point value but do trust the fault existed” rather than discarding it outright, since discarding would bias the uptime statistics.
- *Over/Under/Off-Refined*: the fault-window is refined by suspected fault cause (“sensor calibration drift” vs. “power-supply noise”), each with its own off-scale interval.

A second application — interval-censored survival analysis in clinical trials. A cancer trial tracks time-to-progression, but patients are only examined at scheduled visits, so a progression event is known only to have occurred somewhere in the interval between two visits, e.g. $[8, 12]$ weeks. *Refined*: the censoring interval is refined by the reason for the visit gap — “missed appointment” vs. “scheduled visit spacing” — since the two carry different implications for whether the interval itself is informative about disease severity (patients who feel worse may be more likely to miss appointments, a form of informative censoring). *Over/Under/Off*: a patient whose disease is later found, via retrospective imaging review, to have progressed *before* the trial's official start date effectively has a negative time-to-progression relative to the enrollment-referenced clock — handled in practice by excluding such patients, but OUO-IS gives the option of retaining them as a genuinely off-interval data point rather than discarding information.

A third application — economic index estimation and nowcasting. Statistical agencies routinely publish GDP growth as a point estimate that is revised over subsequent months and years as more complete data arrives; the initial “advance estimate” is, functionally, the midpoint of an unstated imprecise-probability interval that later revisions narrow. *Refined*: the growth estimate is refined by sector (goods-producing vs. services), since services data typically arrives with a longer lag and wider revision interval than goods data, meaning the two sectors' contributions to the aggregate uncertainty band are not interchangeable. *Over/Under/Off*: during periods of unusual economic disruption (e.g. a sudden supply-chain shock), the real-

time “nowcast” interval can be wider than the historical revision-variance model would predict, effectively an Over-Imprecision reading that experienced forecasters already flag qualitatively as “unusually low confidence” without a formal notation for it.

7.3 Event-type study

The pattern established in §2.3 (Refined/Over-Under-Off Imprecise *Probability*) applies directly to the sample-level analogues:

- **Complementary events:** sample-interval complementation follows the bound-flip rule of §2.3, now on empirical interval frequencies.
- **Independent/Dependent events:** tested via interval-valued extensions of contingency analysis (e.g. Kolmogorov-Smirnov-type bounds on interval-censored data); refinement requires the test at every sub-level; OUO requires sign-aware interval arithmetic.
- **Mutually exclusive events:** interval-sample categories that cannot co-occur maintain $\bar{P}(A \cap B) = 0$ exactly as in §2.3.
- **Conditional events:** interval-conditional statistics follow the generalized-Bayes bound rule of §2.3, subsetting empirically.
- **De Morgan’s Rule:** set-theoretic, unaffected; verified on interval bounds.
- **Bayes’ theorem:** the empirical/frequentist form of the interval Bayes rule from §2.3 is used to update interval estimates as new interval-valued batches arrive.

7.4 Relations to other notions

Imprecise Statistics is the empirical counterpart of Imprecise Probability (Family 2), exactly as Family 6 relates to Family 1. It differs from Neutrosophic Statistics (Family 8) in tracking only a belief-interval per observation, with no separate, independently-varying indeterminacy component — the same distinction noted in §2.4 between the two probability families.

7.5 Connections and distinctions with other fields

Interval-censored survival analysis. The clinical-trial example is a direct instance of the well-established biostatistical field of interval-censored survival analysis, which already has its own estimators (the Turnbull estimator, generalizing Kaplan-Meier to interval-censored data) for exactly the kind of bounded-but-unresolved event-time data Refined Imprecise Statistics formalizes; the refinement-by-censoring-reason move above corresponds to what survival analysts call informative-censoring-adjusted estimation.

Robust statistics and outlier-resistant estimation. Robust statistics develops estimators that remain reliable when data may not follow assumed distributions or may contain contamination; imprecise-probability intervals are one formal way of representing “the data-generating process itself is not fully known,” which is the same motivating concern robust statistics addresses via different machinery (M-estimators, trimmed means). The two literatures rarely cite each other despite substantial conceptual overlap.

Official statistics and index revision practice. National statistical offices already publish explicit revision-uncertainty bands for headline economic indicators (though rarely as a formal interval attached to the point estimate itself), and the field of nowcasting — real-time GDP and

inflation estimation using partial, high-frequency data — is essentially the applied practice this family's third example formalizes; Refined/OUO-Imprecise Statistics offers official statistics agencies a more explicit vocabulary for a discipline they already practice informally through revision-history reporting.

Family 8 — Neutrosophic Statistics → Refined / Over-Under-Off / Over-Under-Off-Refined

8.1 Definitions

Neutrosophic Statistics (background, per Smarandache's existing development [11]).

Statistical data and/or statistical inference where some observations, sample sizes, or parameters are indeterminate — recorded not as point values but as neutrosophic numbers $N = a + bI$, where I is indeterminacy with a range, or as full (T, I, F) triples on categorical membership, so that classical statistics is recovered exactly when $I=0$.

Refined Neutrosophic Statistics (RNS). The indeterminacy component I attached to a statistic or observation is refined into $I_1, \dots, I_r$ representing distinct sources of indeterminacy (missing data vs. measurement ambiguity vs. rater disagreement, say), each tracked and reported separately, recombining to the parent I by the family's fusion rule, exactly mirroring Family 3's refinement of I .

Over/Under/Off-Neutrosophic Statistics (OUO-NS). T , I , and/or F attached to sample statistics allowed outside $[0,1]$, mirroring Family 3.

Over/Under/Off-Refined Neutrosophic Statistics (OUO-RNS). Both together.

8.2 Real example/application

A public-health survey on vaccine side-effects, where $NS(\text{had_side_effect}) = (T=0.22, I=0.18, F=0.60)$ across a sample — 18% of the sample gave ambiguous or unusable self-reports.

- *Refined*: I is refined into $I_1 = \text{"form left blank"} (0.11)$ and $I_2 = \text{"form filled but internally contradictory"} (0.07)$, summing to 0.18 — actionable for the survey team, since I_1 suggests a form-design fix while I_2 suggests a comprehension fix.
- *Over/Under/Off*: a batch of forms is later found to have been filled out by proxies (family members) rather than the patients themselves; the team scores this batch's F (evidence against "genuine self-reported side-effect") at $F=1.05$, reflecting that proxy-reporting is worse than simple absence-of-evidence — it is affirmatively unreliable in a way that pushes past ordinary falsity.
- *Over/Under/Off-Refined*: this proxy-report falsity is refined into $F_1 = \text{"proxy filled without consulting patient"} (0.75)$ and $F_2 = \text{"proxy filled after consulting patient by phone, moderately reliable"} (0.30)$, summing to 1.05.

A second application — election polling with undecided and refusal categories. A pre-election poll reports $NS(\text{will vote for candidate A}) = (T=0.42, I=0.15, F=0.43)$, where I combines "genuinely undecided" respondents and those who refused to answer. *Refined*: I is refined into

I_1 = “genuinely undecided” (0.09) and I_2 = “refused to disclose” (0.06), since the two behave very differently on election day — undecided voters split roughly with the observed T/F ratio in most historical polling-error studies, while refusers correlate with a documented “shy voter” effect that skews toward one candidate more than the raw numbers suggest. *Over/Under/Off*: a poll conducted immediately after a candidate’s widely-covered scandal can have its F score pushed above 1 relative to the underlying fundamentals, reflecting a temporary reputational penalty the pollster’s own methodologists flag as likely to partially revert — exactly the kind of “worse than the fundamentals justify” reading OUO-NS is built to represent.

A third application — Net Promoter Score (NPS) with non-response. A company’s customer NPS survey reports $NS(\text{promoter}) = (T=0.35, I=0.20, F=0.45)$, where I represents customers who received the survey but did not respond at all. *Refined*: I is refined into I_1 = “did not open the survey email” and I_2 = “opened but abandoned partway through,” since the second group’s partial responses (if the platform captures them) can sometimes be salvaged as weak evidence, while the first carries essentially no signal. *Over/Under/Off*: a batch of survey responses later found to have been incentivized by a discount code is scored with $T > 1$ relative to the company’s usual promoter baseline, reflecting that the incentive is understood to have inflated scores beyond what genuine satisfaction would produce — a distortion market-research practitioners already discuss informally as “incentive bias” without a formal statistic to quantify it.

8.3 Event-type study

Directly inherits the §3.3 pattern, now applied to sample-derived (T,I,F) statistics rather than theoretical single-event triples:

- **Complementary events**: sample-level (T,I,F) $\rightarrow$ (F,I,T) swap on complementary categories.
- **Independent/Dependent events**: tested via a neutrosophic extension of chi-square/contingency analysis that tracks T-, I-, and F-frequencies as three parallel contingency tables; refinement requires this at every I-sub-component; dependency can be present in the T-table while absent in the F-table (a genuinely three-way form of Simpson’s-paradox-adjacent behavior unique to this family).
- **Mutually exclusive events**: $T(A \cap B) = 0$ in the sample while $I(A \cap B)$ can remain nonzero, exactly as in §3.3, now observable directly as “some respondents’ ambiguous answers are consistent with either category.”
- **Conditional events**: subset the sample and recompute (T,I,F) on the subset; refinement/OUO carry through by subsetting refined/off-scale sub-data.
- **De Morgan’s Rule**: verified componentwise (T, I, F each independently) on sample-derived values.
- **Bayes’ theorem**: the empirical neutrosophic Bayes update (§3.3) applied as new sample batches arrive, updating T, I, and F separately.

8.4 Relations to other notions

Neutrosophic Statistics is the direct empirical counterpart of Neutrosophic Probability (Family 3) and is already published, established Smarandache framework [11]; the Refined and Over/Under/Off extensions here are new layers on top of that existing base, exactly mirroring

how Families 1–5 relate to their respective statistics counterparts. It differs from ordinary “statistics with missing data” (multiple imputation, etc.) in that indeterminacy is *retained and reported*, not filled in or imputed away. Where a dataset’s “unresolved” category itself needs splitting into confirmed-stable versus genuinely-ambiguous cases — as in the election-polling and NPS survey examples above — the Quadruple (T, I, N, F) extension detailed in §3.6 applies directly to sample statistics as well as to theoretical probabilities, letting the I-refinement already recommended in §8.2 be formalized as a T/I/N/F split rather than an ad hoc I-sub-category breakdown.

8.5 Connections and distinctions with other fields

Missing-data theory. The MCAR/MAR/MNAR taxonomy (missing completely at random, missing at random, missing not at random) developed by Rubin is the standard statistical framework for exactly the “form left blank” and “did not respond” categories in the examples above. Rubin’s multiple imputation approach resolves missingness by generating plausible replacement values and averaging over them; Neutrosophic Statistics takes the opposite philosophical stance — instead of imputing a best guess, it reports the indeterminacy explicitly as a first-class quantity, which is preferable exactly when the *fact* that data is missing (and why) is itself informative, as in the shy-voter and incentive-bias examples.

Survey methodology. Total survey error frameworks in survey methodology already decompose non-response into multiple sources (unit non-response, item non-response, refusal) that closely parallel the I-refinement examples above; Refined Neutrosophic Statistics offers survey methodologists a way to carry those distinctions all the way through to the final reported statistic rather than collapsing them into a single response-rate footnote.

Epidemiological surveillance. Public-health surveillance systems routinely grapple with under-ascertainment (cases that exist but are never tested or reported) and mis-classification (cases attributed to the wrong cause), both of which are naturally represented as neutrosophic indeterminacy and falsity respectively; the vaccine side-effect example above is a direct instance of the broader field of pharmacovigilance, which has long used qualitative “possible/probable/definite” causality categories for adverse events that map closely onto a coarse (T, I, F) triple, even though the field has not historically formalized them as such.

Family 9 — Plithogenic Statistics → Refined / Over-Under-Off / Over-Under-Off-Refined

9.1 Definitions

Plithogenic Statistics (background). Statistics computed over data described by a plithogenic attribute structure (Family 4) [6]: each observation carries a vector of attribute-value degrees rather than a single category label, and descriptive/inferential statistics (means, tests, regressions) are computed per attribute-value and then re-fused using the model’s contradiction-degree-weighted fusion operator — an approach already applied to decision-making under uncertainty in the pedagogical/social-science setting of [8].

Refined Plithogenic Statistics (RPS). Attribute-values in the sample data are refined into sub-values before aggregation, exactly mirroring Family 4's refinement, now applied empirically.

Over/Under/Off-Plithogenic Statistics (OUO-PS). Sample degrees and/or contradiction degrees allowed outside $[0,1]$, mirroring Family 4.

Refined Plithogenic Statistics and Over/Under/Off-Refined Plithogenic Statistics are listed as the fourth, composite notion in this family in the source material (distinct from the other eight families, which have only three named variants each): here, “Refined Plithogenic Statistics” is highlighted as its own headline notion (rather than folded silently into the general pattern) because plithogenic attribute refinement is empirically the richest of all nine families — a single real dataset typically has many attributes, each independently refinable — so the literature calls it out by name before layering Over/Under/Off on top of it.

9.2 Real example/application

Customer-churn modeling for a telecom company, attribute = “churn risk factors,” values = {price sensitivity, service-quality complaints, competitor promotions}, dominant value = price sensitivity, sample degrees $d = (0.40, 0.35, 0.25)$, contradiction degrees relative to price sensitivity $c = (0, 0.5, 0.7)$.

- *Refined:* “service-quality complaints” is refined into “network outage complaints” (0.20) and “billing-error complaints” (0.15), summing to 0.35, each independently correlated with churn at different strengths — network complaints turn out far more predictive, which the unrefined aggregate obscured.
- *Over/Under/Off:* customers who received a *retention call* after complaining show a contradiction degree between “service complaint” and “price sensitivity” of $c = -0.30$ (an under-value): far from being merely uncorrelated with price-driven churn, a well-handled complaint actively predicts *lower* churn even among price-sensitive customers, a protective interaction the ordinary $[0,1]$ contradiction scale cannot express.
- *Over/Under/Off-Refined:* the retention-call effect is refined into “call within 24 hours” ($c = -0.45$, strong protective effect) and “call after 24 hours” ($c = -0.10$, weak protective effect).

A second application — consumer credit scoring. A lender's scorecard attribute “default risk” has values {payment history, credit utilization, length of credit history}, dominant value = payment history, sample degrees $d = (0.45, 0.35, 0.20)$. *Refined:* credit utilization is refined into “revolving utilization” (credit cards) and “installment utilization” (auto/student loans), sampled across the loan book, since the two carry measurably different default correlations in most retail-lending portfolios. *Over/Under/Off:* a sub-population of borrowers who reduced their utilization specifically *because* of a documented income shock shows a contradiction degree between utilization and payment history that goes negative ($c < 0$): lower utilization here does not signal financial health as it normally would, it signals financial distress being actively managed — a form of the “utilization can lie” phenomenon credit-risk modelers already work around with ad hoc adjustment variables, which OUO-RPS gives a principled formal home.

A third application — sports analytics and player evaluation. A basketball analytics team scores players on attribute “offensive value” with values {scoring efficiency, playmaking, shot creation}, dominant value = scoring efficiency, sample degrees from a season’s play-by-play data $d = (0.40, 0.30, 0.30)$. *Refined*: playmaking is refined into “assisted-basket creation” and “turnover avoidance,” sampled per possession, since a player can excel at one while being weak at the other in ways a single combined “playmaking” number obscures. *Over/Under/Off*: a player whose on-court presence measurably improves teammates’ shooting percentages even on possessions where he neither scores nor assists (a well-documented “gravity” effect in modern basketball analytics, typically measured via on/off-court differentials) contributes a contradiction degree between playmaking and scoring efficiency that goes negative, capturing value the box score itself never records — precisely the kind of attribute interaction traditional box-score statistics are known to miss entirely.

9.3 Event-type study

Directly inherits the §4.3 pattern, applied to sample-derived attribute-value degrees:

- **Complementary events**: complementation of the dominant attribute-value’s sample proportion follows §4.3’s rule; non-dominant contradiction degrees are preserved under complementation as relational, not truth-status, quantities.
- **Independent/Dependent events**: independence is tested attribute-value-by-attribute-value on the sample (as in §4.3); this family is where such per-attribute independence testing is most practically valuable, since real datasets (as in the churn example) very often show independence on one attribute-value and strong dependence on another.
- **Mutually exclusive events**: dominant-value exclusivity vs. non-dominant-value co-occurrence, exactly as in §4.3, verified empirically.
- **Conditional events**: subset-and-refuse as in §4.3, now on sample data — this is, in practice, ordinary plithogenic-weighted subgroup analysis.
- **De Morgan’s Rule**: verified per attribute-value against whichever fusion operator the analyst selected, exactly as in §4.3’s caveat that this is a modeling choice.
- **Bayes’ theorem**: the empirical plithogenic Bayes update, applied per attribute-value and re-fused, as new sample batches (e.g. new churn cohorts) arrive.

9.4 Relations to other notions

Plithogenic Statistics is the empirical counterpart of Plithogenic Probability (Family 4) [6] and, being built on the most general of the four probability families, itself generalizes Families 6–8: Classical Statistics is the one-attribute/one-value degenerate case, Imprecise Statistics the one-attribute/two-value (lower/upper) case, and Neutrosophic Statistics the one-attribute/three-value (T,I,F) case [9, 10]. Refined and Over/Under/Off-Refined Plithogenic Statistics are therefore, in principle, general enough to reproduce every other refined/relaxed statistic in this document as a special case — which also makes them the hardest of the nine families to fit coherent real-world models to, since the analyst must justify the attribute structure itself, not just the numbers populating it, as illustrated in the applied decision-making literature [7, 8].

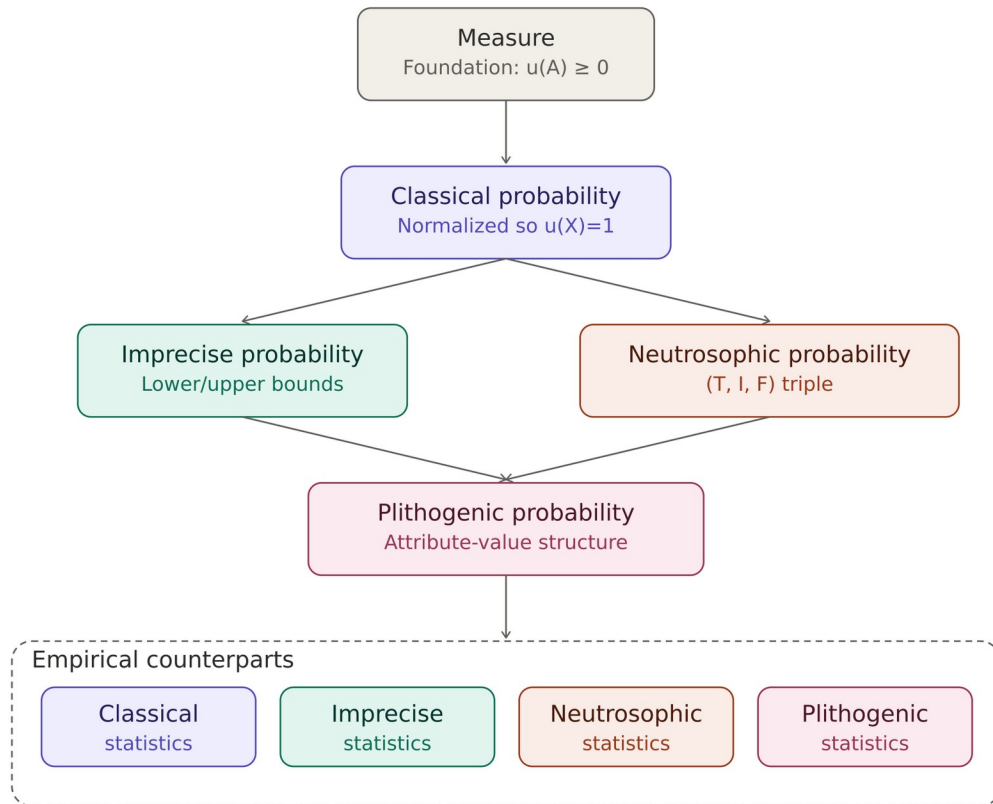
9.5 Connections and distinctions with other fields

Credit risk modeling and scorecard development. The credit-scoring example connects directly to the FICO-style scorecard methodology used throughout consumer lending, which already combines multiple weighted risk factors into a composite score; what that industry-standard practice lacks, and what Plithogenic Statistics adds, is a formal, refinable notion of *interaction* between factors (the contradiction degree) rather than the purely additive weighted-sum scorecards conventionally use, and a principled way to let that interaction go negative (protective) rather than only ranging from zero to fully reinforcing.

Sports analytics and sabermetrics. The basketball “gravity effect” example sits within the broader sabermetrics/analytics movement across sports (originating in baseball, now standard in basketball, soccer, and American football), which has increasingly moved toward on/off-court and plus-minus-style differential statistics precisely because traditional box-score counting statistics miss exactly the kind of attribute-value interaction the plithogenic contradiction degree is built to capture.

Ensemble machine learning and feature interaction. In machine learning terms, a plithogenic attribute-value structure is closely analogous to a feature vector, and the contradiction degree plays a role similar to feature interaction terms in tree-based models (which explicitly model how two features’ combined effect differs from their individual effects) or to the off-diagonal terms of a feature covariance/correlation matrix. The distinction is interpretability: where a gradient-boosted tree’s feature interactions are implicit and require post-hoc explanation tools (e.g. SHAP values) to surface, the plithogenic contradiction degree is a first-class, directly reported quantity by construction — a trade-off of predictive flexibility for transparency that makes the framework better suited to regulated domains like consumer lending, where model explainability is often a legal requirement, than to raw predictive performance competitions.

Part III — Overall Synthesis: How the Nine Families Relate



Synthesis diagram: Measure generalizes to Classical Probability, which branches into Imprecise and Neutrosophic Probability, both generalizing into Plithogenic Probability, with each probability family's empirical statistics counterpart grouped below

Measure sits at the foundation, purple marks the classical line, teal marks the interval (imprecise) line, coral marks the evidence-triple (neutrosophic) line, and pink marks the fused attribute-value (plithogenic) line — the most general of the four. Each of the five probability/measure families has an empirical (data-driven, sample-based) counterpart, obtained by applying the same Refined / Over-Under-Off operators to sample statistics instead of theoretical event-probabilities:

Family 1 — Family 6 (Classical Statistics) Family 2 — Family 7 (Imprecise Statistics) Family 3 — Family 8 (Neutrosophic Statistics) Family 4 — Family 9 (Plithogenic Statistics) Family 5 has no separate statistics family: Refined/OUO-Measure is used directly as the machinery underlying all four statistics families' aggregation rules (sums, means, re-fusion), rather than needing its own empirical counterpart.

Key distinctions to keep straight across all nine families:

- *Refinement* is about **granularity of description** — it never changes what range a value can occupy, only how finely the total is broken into named, individually-trackable parts.

It is the tool of choice when the *question* is “what, specifically, within this event/statistic, is driving the number?”

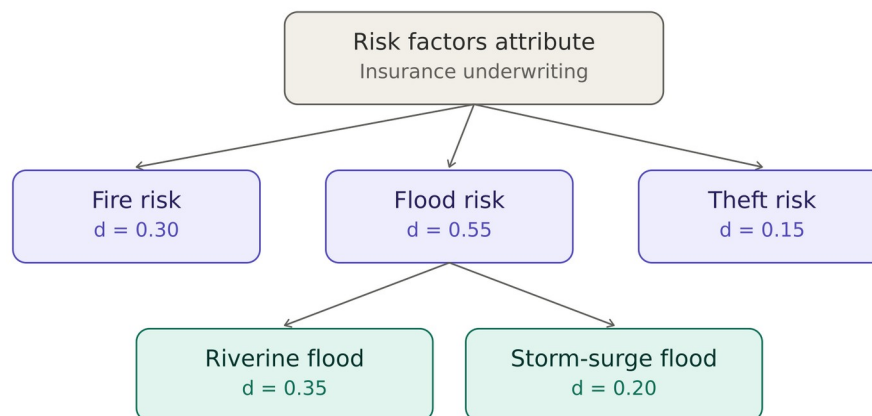
- *Over/Under/Off* is about **range relaxation** — it changes what values are admissible, allowing “more than certain/full” and “less than impossible/zero/none,” including genuinely negative contributions. It is the tool of choice when the *phenomenon itself* is not well-represented by a bounded scale (redundant reinforcing evidence, active harm, protective interaction effects, etc.).
- The two operators are independent and composable in either order, which is why the ninth, most general family (Plithogenic Statistics) is the one point in the hierarchy where both are exercised most fully and simultaneously, and also the family requiring the most modeling judgment, since its fusion operator is not fixed by the mathematics the way summation is in Families 1, 2, 5, 6, and 7.

Part IV — Extended Examples and Applications: Refined Plithogenic Probability, Refined Plithogenic Statistics, and Over/Under/Off-Refined Plithogenic Statistics

This section expands §4.2 and §9.2 with additional worked applications and diagrams, since these three notions are, in practice, the ones analysts reach for most often — real attribute-value data is almost always both refinable and prone to off-interval interactions.

IV.1 Refined Plithogenic Probability (RPP)

Example A — insurance underwriting (detailed). An attribute “risk factors” with dominant value fire risk ($d=0.30$) is refined at the flood-risk branch into two physically distinct causes, each independently insurable and independently re-priceable:



Refined plithogenic probability: risk factors attribute refined into fire, flood, and theft risk, with flood risk further refined into riverine and storm-surge flood

The reduction constraint is exact: $d(\text{riverine}) + d(\text{storm-surge}) = 0.35 + 0.20 = 0.55 = d(\text{flood risk})$, so refining never changes the parent-level premium calculation — it only exposes which

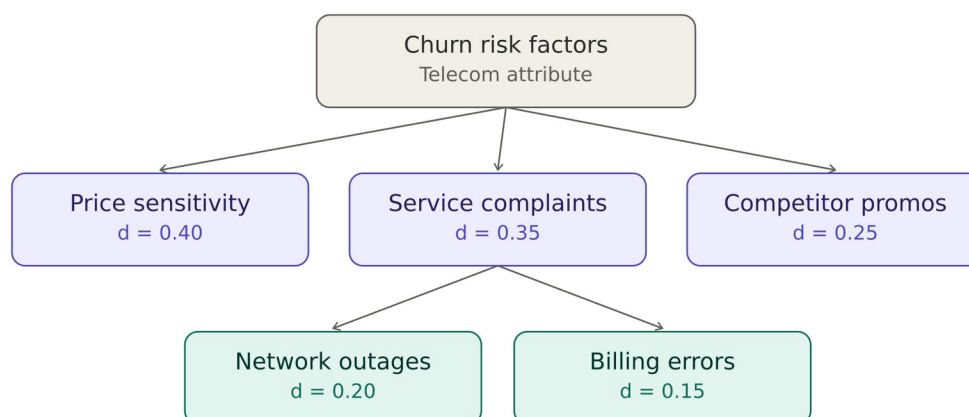
sub-risk is driving it, which is what lets the underwriter price riverine and storm-surge exposure on separate endorsement riders instead of one blended flood surcharge.

Example B — agricultural yield-loss insurance. A crop insurer’s “yield loss cause” attribute refines drought risk into “root-zone drought” (detected via soil-moisture sensors) and “canopy heat stress” (detected via satellite thermal imagery). Each sub-value is scored against a different data feed and a different payout trigger, even though both sum back to the single “drought” line the farmer sees on the policy.

Example C — medical differential diagnosis. A “symptom cluster” attribute with dominant value “infection” ($d=0.45$) is refined into “bacterial” (0.30) and “viral” (0.15) sub-values, each carrying a different contradiction degree against the “autoimmune” competing value — bacterial infection contradicts autoimmune flare-up more strongly ($c=0.7$) than viral infection does ($c=0.4$), because viral triggers are a known autoimmune-flare precursor. This is exactly the kind of distinction a single unrefined “infection” value would erase.

IV.2 Refined Plithogenic Statistics (RPS)

Example A — telecom churn (detailed). The sample-level churn attribute, refined at the service-complaints branch by root cause:



Refined plithogenic statistics: churn risk factors refined into price sensitivity, service complaints, and competitor promotions, with service complaints further refined into network outage and billing error complaints

In the underlying dataset, network-outage complaints turn out to correlate with churn at roughly twice the strength of billing-error complaints — a difference the unrefined 0.35 aggregate completely hides, and one that sends the retention team toward a network-reliability fix rather than a billing-UX fix.

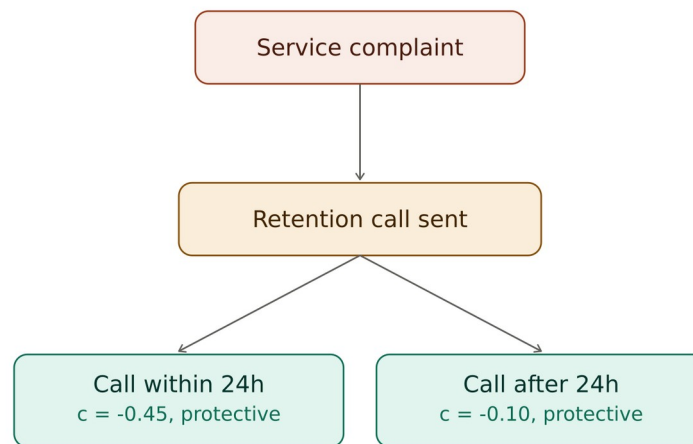
Example B — university dropout risk. A cohort’s “dropout risk” attribute refines “financial strain” (sample degree 0.38) into “tuition-driven” (0.24) and “housing-driven” (0.14) sub-

categories. The two respond to different interventions — payment plans for the first, housing subsidies for the second — which is only visible once the aggregate is split.

Example C — manufacturing defect tracking. A production line’s “quality issue” attribute refines “material defect” into “supplier batch A” and “supplier batch B” sub-values, sampled across a quarter’s production runs, revealing that batch B alone accounts for the majority of the aggregate defect rate — actionable at the supplier level in a way the combined statistic is not.

IV.3 Over/Under/Off-Refined Plithogenic Statistics (OUO-RPS)

Example A — telecom retention calls (detailed). The contradiction degree between “service complaint” and price-sensitive churn goes negative once a retention call is made — a protective interaction, refined by call latency:



Over/Under/Off-refined plithogenic statistics: a service complaint followed by a retention call, with a negative, protective contradiction degree refined by whether the call happened within or after 24 hours

Calls placed within 24 hours carry $c = -0.45$ against price-driven churn — a strong protective effect, well past ordinary “no correlation” — while calls placed after 24 hours carry only $c = -0.10$, still protective but much weaker. Ordinary $[0,1]$ -bounded contradiction degrees have no way to express “actively protective,” only “uncorrelated” (c near the dominant value) or “opposed” (c near 1); the off-interval negative region is what lets the statistic say “this doesn’t just fail to predict churn, it actively predicts retention.”

Example B — manufacturing supplier audits. A supplier’s defect rate ordinarily correlates positively with downstream recalls, but the sub-value “defect caught by the supplier’s own pre-shipment audit” carries a negative contradiction degree ($c \approx -0.3$) against recall risk: catching the defect there actively lowers recall likelihood below the zero-correlation baseline, because it signals the supplier’s quality system is functioning even though a defect occurred at all.

Example C — clinical trial adverse-event reporting. An “adverse event” attribute normally has a strongly positive contradiction degree against “drug efficacy,” but a refined sub-value —

“adverse event reported by a highly engaged, frequently-monitored patient” — is scored with a contradiction degree over 1 ($c \approx 1.15$, an Over-value): the team treats it as more than ordinarily opposed to efficacy, since heavy monitoring plus a reported event together suggest the effect is both real and actively worsening, not merely present.

Common thread across De Morgan’s Rule specifically: in every one of the nine families, De Morgan’s Rule survives unchanged as a *set-theoretic* statement — refinement and Over/Under/Off only ever change the *numbers* attached to the two sides of the identity, never the identity itself. This is the one property, of the seven studied, that is entirely immune to both generating operators, and is worth stating as a standalone conclusion: **whatever else is refined or relaxed, the underlying event algebra of any of these nine families is always classical, two-valued set algebra** — all of the novelty in this document lives in *how much weight* is attached to each part of that algebra, never in the algebra’s logical structure itself.

Part V — The Nine Families in the Landscape of Uncertainty Science

Parts I–IV built the framework from the inside out. This part turns outward: where does each piece sit relative to the established fields that already study probability, statistics, and decision-making under uncertainty? The goal is not to claim priority over those fields — most of them are far older and far more thoroughly validated than the nine families above — but to make the correspondences explicit, so a reader arriving from any one of these backgrounds can immediately locate familiar territory and see exactly where this framework’s contribution begins.

V.1 Classical probability and measure theory

Kolmogorov’s 1933 axiomatization is the common ancestor of every family in this document, and Family 1 (Classical Probability) and Family 5 (Measure) are, respectively, its direct object and its direct generalization. The Refined variants of both families are conservative in the strongest sense: they add bookkeeping (naming and separately tracking an additive decomposition) without changing what is provable. The Over/Under/Off variants are where this document departs from Kolmogorov outright, since non-negativity and total-mass-one normalization are foundational axioms there, not optional conveniences — which is precisely why OUO-CP and OUO-Measure are framed throughout as *scoring* or *accounting* constructions rather than as probabilities in Kolmogorov’s sense, even when they are computed and manipulated using probability-like machinery.

V.2 Bayesian inference and subjective probability

The Bayesian tradition, from Bayes and Laplace through de Finetti’s subjective probability and modern Bayesian computational statistics, treats probability as a degree of belief updated by evidence via Bayes’ theorem — the same formula this document re-derives, family by family, in every $\$X.3$ event-type study. Two distinctions are worth making explicit. First, Bayesian updating is normally a single coherent process across one probability measure; the refined families in this document instead maintain several simultaneous sub-measures (the A_i , the T_i , the attribute-values) that must each satisfy their own local coherence and jointly recompose — closer to a hierarchical or multilevel Bayesian model than to simple single-parameter updating.

Second, standard Bayesian probability has no off-interval analogue; a posterior is always a valid probability distribution by construction. Where this document allows OUO-Bayesian posteriors to leave $[0,1]$ (§1.3, §3.3), it is explicitly stepping outside the Bayesian formalism into the same scoring-rule territory noted in §1.5, not proposing a rival theory of belief updating.

V.3 Dempster-Shafer evidence theory, and its Dezert-Smarandache extension (DSmT)

As already noted in §3.5, Dempster-Shafer theory's belief and plausibility functions are the closest pre-existing formal analogue to the neutrosophic (T, I, F) triple, and the 2026 extension of Masses in Information Fusion to Over/Under/Off-Masses [5] is the direct bridge between the two literatures: Dempster-Shafer mass functions are non-negative and sum to 1 across the power set of the frame of discernment by construction, exactly the axiom Family 5's Over/Under/Off-Measure relaxes.

Dempster-Shafer theory (DST) itself, however, is not the framework this document's families sit closest to — that distinction belongs to the Dezert-Smarandache Theory (DSmT) of plausible, paradoxist, and neutrosophic reasoning [14], co-developed by Jean Dezert and the same F. Smarandache whose operators generate every family in this document. DSmT was built specifically to go “one step beyond” classical DST: it drops DST's requirement that the frame of discernment be exhaustive and exclusive (refinable into disjoint atomic elements), instead allowing hybrid and free DSm models where elements can genuinely overlap or resist refinement altogether, and it replaces Dempster's rule of combination — which is known to produce counterintuitive results under high source conflict — with the Proportional Conflict Redistribution (PCR) family of combination rules, designed explicitly to handle highly conflicting, imprecise, and uncertain sources of evidence. This is the same “unrefinable frame of discernment” problem that motivates Part I's Refinement operator, approached from the sensor-fusion side of the literature rather than the set-theoretic side developed here, which makes DSmT — more than classical DST — the direct information-fusion-theoretic sibling of this document's Neutrosophic and Plithogenic Probability families. Sensor-fusion and multi-source intelligence applications, where DSmT already has an extensive applied literature spanning target tracking, robotics, biometric fusion, and geophysical data fusion, are natural proving grounds for the refined and off-interval extensions developed here, since those applications already grapple routinely with conflicting, redundant, and occasionally adversarial sources exactly analogous to the intelligence-analysis example in §3.2 and the cybersecurity example in §V.16.

V.4 Fuzzy set theory and its descendants

Zadeh's 1965 fuzzy sets, Atanassov's 1986 intuitionistic fuzzy sets, and the subsequent lineage of Pythagorean, spherical, picture, and q-rung orthopair fuzzy sets are all special cases of neutrosophic and plithogenic probability, a generalization result proven directly in [10] and used throughout Family 3 and Family 4's relations sections. The practical distinction for a reader choosing between these frameworks is dimensionality of the uncertainty representation: fuzzy sets track one degree (membership), intuitionistic fuzzy sets track two (membership and non-membership, with a forced complementary relationship), neutrosophic sets track three independently-varying degrees (T, I, F), and plithogenic sets generalize further to an arbitrary number of attribute-values. Each additional degree of freedom buys expressiveness at the cost of requiring more data to estimate reliably — a fuzzy-set model

needs to fit one function, a plithogenic model with five attribute-values and pairwise contradiction degrees needs to fit on the order of fifteen. The refinement operator (Family-by-family, §X.1) compounds this cost further, which is why §9.4 flags Plithogenic Statistics as the hardest family to fit coherent real-world models to.

V.5 Rough set theory

Pawlak's rough set theory, developed independently of the fuzzy-set lineage in the early 1980s, handles indeterminacy through *boundary regions*: an object is classified as definitely in a set, definitely out, or in the boundary region where available attributes cannot distinguish it from objects in other classes. This is conceptually close to the neutrosophic Indeterminacy component I, but the mechanism is different — rough sets derive the boundary region from an equivalence relation over the available attributes (indiscernibility), while I in this document is an independently assigned or estimated quantity, not derived from a formal indiscernibility structure. A natural extension this document does not pursue, but flags for future work, would be a “Rough Neutrosophic” hybrid in which I is constrained to equal the rough-set boundary-region measure rather than being freely specified — several papers in the broader neutrosophic literature have explored exactly this hybridization.

V.6 Possibility theory and grey system theory

Possibility theory (Zadeh, later formalized by Dubois and Prade) uses a possibility measure Π and a dual necessity measure N , with $N(A) \leq P(A) \leq \Pi(A)$ bracketing probability much as Imprecise Probability's $[\underline{P}, \bar{P}]$ does — the two frameworks are close enough that possibility/necessity pairs can usually be translated directly into Family 2's interval notation. Grey system theory (Deng, 1982) instead represents an unknown quantity as a “grey number” lying in an interval with an associated “whitening” function describing how the true value is distributed within that interval — richer than a plain interval but less general than a full neutrosophic triple, and explicitly cited in [10] as one of the fields neutrosophication is shown to generalize. Both traditions predate the neutrosophic/plithogenic literature by decades and remain independently active, particularly in Chinese-language operations-research scholarship for grey systems.

V.7 Multi-criteria decision-making (MCDM)

As detailed in §4.5, the Analytic Hierarchy Process, TOPSIS, VIKOR, and related MCDM methods are the applied decision-theory tradition Plithogenic Probability and Plithogenic Statistics most directly generalize, and this generalization is not merely formal — [7] and [8] both apply plithogenic MCDM-style methods to real optimization and pedagogical decision problems. The refinement operator's role in this tradition corresponds to what MCDM literature calls criteria decomposition or hierarchical criteria structuring (breaking a top-level criterion into sub-criteria, as AHP's hierarchy diagrams already do); the novel contribution is the *contradiction degree* as an explicit, refinable, and potentially off-interval measure of criteria interaction, which classical MCDM weighting schemes generally assume away or handle only through ad hoc consistency-ratio checks.

V.8 Machine learning and uncertainty quantification

Modern machine learning has developed its own vocabulary for exactly the aleatory/epistemic distinction discussed in §3.5: epistemic uncertainty (reducible by more data or a better model) is commonly estimated via Bayesian neural networks, deep ensembles, or Monte Carlo dropout, while aleatory uncertainty (irreducible noise in the data-generating process) is estimated via heteroscedastic output layers. Conformal prediction, a more recent and rapidly growing subfield, constructs prediction intervals with a formal coverage guarantee under minimal distributional assumptions — functionally very close to Family 2’s Imprecise Probability, though derived from a completely different (frequentist, distribution-free) theoretical foundation rather than from Walley’s behavioral coherence axioms. The refined I-component developed in Family 3 and Family 8 offers machine learning a decomposition tool — splitting “the model doesn’t know” into named, addressable sources — that current uncertainty-quantification practice generally does not provide, since most methods report a single scalar uncertainty estimate per prediction rather than a refined breakdown.

V.9 Actuarial science, risk theory, and finance

Actuarial loading (§1.5), signed-measure accounting (§5.5), and copula-based tail-dependency modeling (§4.5) collectively place this framework in direct dialogue with quantitative risk management as practiced in insurance and finance. Solvency capital models (e.g. Solvency II in European insurance regulation) already require insurers to hold capital against tail scenarios where correlations between risk types spike beyond their normal-times values — precisely the off-interval contradiction-degree phenomenon in §4.2’s insurance example and §4.5’s portfolio-risk example, currently handled in practice through stress-testing overlays rather than through any single unified probabilistic framework. Whether the plithogenic apparatus could serve as that unifying framework, replacing today’s patchwork of stress tests and correlation floors, is an open question this document does not resolve but that its examples throughout Families 1, 4, 5, and 9 are directly relevant to.

V.10 Biostatistics, epidemiology, and survey methodology

Sections §7.5 and §8.5 already trace direct connections to interval-censored survival analysis, missing-data theory, survey total-error frameworks, and pharmacovigilance. The broader point worth stating explicitly here is that medicine and public health are, empirically, the application domain where indeterminacy is most routinely and most consequentially encountered — diagnostic uncertainty, censored follow-up, non-response, and adverse-event causality assessment are not edge cases in these fields but the daily working material of biostatistics — which is presumably why Neutrosophic Statistics [11] was developed with explicit reference to exactly this kind of data, well before the refinement and Over/Under/Off extensions developed in this document were layered on top.

V.11 Econometrics and behavioral economics

Frank Knight’s 1921 distinction between “risk” (quantifiable probability) and “uncertainty” (unquantifiable, in Knight’s original sense) is the founding statement of exactly the problem Imprecise, Neutrosophic, and Plithogenic Probability all attempt to formalize: situations where a single point probability is not epistemically justified. Ellsberg’s 1961 paradox experimentally

demonstrated that people’s actual decision-making under Knightian uncertainty is not well described by expected-utility theory over a single probability distribution, motivating decision-theoretic extensions such as Choquet expected utility and maxmin expected utility — both of which operate over exactly the kind of interval or set-valued probability structures formalized in Family 2. Prospect theory’s decision weights, which systematically over-weight small probabilities and under-weight large ones relative to their objective values, bear a structural resemblance to the plithogenic contradiction degree’s departure from simple linear weighting, though the two were developed for different purposes (descriptive psychology versus a general representational framework) and the resemblance should not be overstated as equivalence.

V.12 Information fusion and sensor networks

The Joint Directors of Laboratories (JDL) data fusion model, the standard reference architecture for multi-sensor military and robotics fusion systems, defines successive levels of fusion from raw signal alignment through situation and threat assessment — and Dempster-Shafer combination, discussed in §V.3 above, is one of its most common mathematical engines at the mid-levels. The 2026 Over/Under/Off-Masses extension [5] was developed specifically in this information-fusion context, making it, of the framework’s core operators, the one with the most direct and immediate applied lineage: sensor networks routinely need to represent “this sensor’s report actively contradicts, beyond mere non-confirmation, what the other sensors are reporting” — exactly the semantic content of an Over/Under mass value — when reconciling reports from potentially compromised, jammed, or malfunctioning sensors in contested environments.

V.13 Philosophy of probability: aleatory versus epistemic uncertainty

Underlying nearly every connection drawn above is the same philosophical fault line: whether a given “probability” measures an objective feature of the world (a propensity, a long-run frequency) or a subjective state of belief given available evidence. Classical Probability and Measure (Families 1 and 5) are agnostic on this question by design — the axioms work identically under either interpretation. Neutrosophic and Plithogenic Probability (Families 3 and 4) are harder to read as purely objective, since the Indeterminacy component and the contradiction degree are naturally epistemic quantities describing the state of evidence rather than a physical propensity of the world; this is not a weakness so much as a declaration of philosophical commitment that a reader coming from a strict frequentist background should recognize explicitly before adopting the framework for a given application, particularly in fields like physics where the aleatory/epistemic distinction is itself a matter of live, unresolved debate (as in disputes over the interpretation of quantum-mechanical probabilities).

V.14 Information theory and entropy

Shannon entropy $H = -\sum p \log p$ measures the expected information content of a probability distribution, and is maximized exactly when a distribution is most “spread out” — most indeterminate, in an information-theoretic sense. There is a natural but underexplored correspondence between the neutrosophic Indeterminacy component I and entropy: a high- I event is, informally, a high-entropy event with respect to the observer’s evidence, and the refinement of I into named sub-sources (as in the petroleum-reserve and election-polling examples above) is structurally similar to a conditional-entropy decomposition, where total

entropy is split into the entropy attributable to each of several explanatory variables. Formalizing this correspondence precisely — deriving I as a monotonic transform of an appropriately conditioned entropy rather than specifying it directly — is a natural next step for future work connecting this framework to information theory.

V.15 Climate science and IPCC uncertainty language

The Intergovernmental Panel on Climate Change has, since its Third Assessment Report, used a calibrated uncertainty language (terms like “virtually certain,” “very likely,” “about as likely as not”) that assigns each term a specific probability range — e.g. “very likely” means 90–100%. This is, in effect, an institutionalized version of Imprecise Probability (Family 2), and the IPCC’s parallel “confidence” scale (very low to very high, based on the type, amount, quality, and agreement of evidence) is a close institutional analogue of the neutrosophic Indeterminacy component, since it tracks the *quality of the evidence base* independently of the *magnitude of the estimate* — exactly the two-axis structure (T/F magnitude vs. I evidence-quality) this document’s Family 3 formalizes. Climate scientists already refine their confidence assessments by evidence source (paleoclimate proxies vs. instrumental records vs. model projections) in exactly the pattern this document calls Refined Neutrosophic Probability, without using that name.

V.16 Cybersecurity threat modeling

Threat intelligence platforms routinely combine indicators from multiple, sometimes contradictory or deliberately falsified (in the case of adversarial deception), sources into a single confidence score for whether a given signal indicates a genuine attack. This is architecturally identical to the neutrosophic criminal-trial and intelligence-analysis examples above: a source later confirmed to be compromised or feeding disinformation retroactively pushes the falsity component of every indicator it contributed to past ordinary falsity, exactly mirroring the double-agent example in §3.2. Security operations centers already informally refine “unknown” alert status into sub-categories (insufficient telemetry vs. genuinely ambiguous behavior vs. known-benign-but-unclassified), a direct instance of Refined Neutrosophic Statistics applied to a live operational data stream.

V.17 Quantum probability and quantum logic

Quantum mechanics uses a probability calculus (via the Born rule applied to complex probability amplitudes) that violates classical Kolmogorov axioms in specific, experimentally verified ways — most famously, quantum probabilities do not always obey the classical law of total probability, and quantum logic (Birkhoff and von Neumann, 1936) replaces the Boolean lattice of classical event algebra with a non-distributive orthomodular lattice. This document’s closing observation in Part IV — that De Morgan’s Rule survives every refinement and relaxation because the underlying event algebra stays classical and two-valued — is precisely the property that distinguishes this entire framework from quantum probability: everything developed here, however exotic the numerical values become, remains built on ordinary Boolean set algebra, and none of the nine families should be mistaken for, or expected to reproduce, genuinely quantum statistical phenomena such as entanglement or interference. Readers arriving from quantum foundations should treat the resemblance between

“indeterminacy” in this document’s sense and quantum indeterminacy as a linguistic coincidence rather than a mathematical one.

Part VI — Worked Numerical Case Studies

Parts I–IV stated the definitions and event-type rules abstractly; this part works four of them through with full arithmetic, so a reader can check every step by hand.

VI.1 Bayes’ theorem under Over/Under/Off-Classical Probability (Family 1)

Return to the sports-odds-making example (§1.2, third application). A sportsbook’s posted moneyline implies $P(A \text{ wins}) = 0.58$, $P(B \text{ wins}) = 0.55$ (overround built in, so these sum to 1.13, not 1.00 — an Over-Probability space by design). Suppose a bettor wants the book’s implied conditional probability that A wins in regulation, given that the game does not go to overtime, $P(A \text{ wins in reg} \mid \text{no OT})$. The book separately prices $P(\text{no OT}) = 0.82$ and $P(A \text{ wins in reg} \cap \text{no OT}) = 0.44$ (matching the refined breakdown in §1.2). Applying Bayes’ theorem exactly as stated in §1.3 ($P(A|B) = P(A \cap B)/P(B)$):

$$P(A \text{ wins in reg} \mid \text{no OT}) = 0.44 / 0.82 = 0.5366 \text{ (53.66\%)}$$

This value is comfortably inside $[0,1]$ despite living inside an Over-Probability space overall — a useful reminder that individual conditional computations can land safely in-range even when the surrounding unconditional space is not $[0,1]$ -normalized; the off-interval behavior only becomes visible when the full outcome set is summed. Now suppose the book’s overtime price is itself revised, dropping $P(\text{no OT})$ to 0.79 as bettors pile onto an anticipated shootout. Recomputing: $P(A \text{ wins in reg} \mid \text{no OT}) = 0.44 / 0.79 = 0.5570$ (55.70%). The conditional estimate moves by 2.04 percentage points from a 3-percentage-point shift in the marginal — the small-denominator sensitivity is exactly why bookmakers monitor overtime-pricing markets so closely relative to their apparent size in the overall handle.

VI.2 De Morgan’s Rule verification under Refined Neutrosophic Probability (Family 3)

Return to the criminal-trial example (§3.2): $NP(\text{guilty}) = (T=0.55, I=0.30, F=0.15)$. Let A = “guilty” and B = “prior record introduced as aggravating” with $NP(B) = (T=0.40, I=0.20, F=0.40)$, and suppose the panel assesses the joint event using the componentwise combination rules from §3.3: $T(A \cap B) = T(A) \cdot T(B) = 0.55 \times 0.40 = 0.22$, and $F(A \cup B) = F(A) + F(B) - F(A) \cdot F(B) = 0.15 + 0.40 - (0.15 \times 0.40) = 0.55 - 0.06 = 0.49$. De Morgan’s Rule requires $(A \cap B)^c$ to match $\bar{A}^c \cup \bar{B}$ component-by-component. By the complementation rule $NP(\bar{X}) = (F(X), I(X), T(X))$: $NP(\bar{A}) = (0.15, 0.30, 0.55)$ and $NP(\bar{B}) = (0.40, 0.20, 0.40)$. Computing $F((A \cap B)^c)$ directly from the truth-side first: $T(A \cap B) = 0.22$ (computed above), so $NP(\bar{A} \cup \bar{B})$ ’s truth-side must equal $NP((A \cap B)^c)$ ’s truth-side, i.e. $F(A \cap B)$ should equal $T(\bar{A} \cup \bar{B}) = T(\bar{A}) + T(\bar{B}) - T(\bar{A}) \cdot T(\bar{B}) = 0.15 + 0.40 - 0.06 = 0.49$. Independently, $F(A \cap B)$ computed the standard way (via the dual combination on the falsity side, $F(A \cap B) = F(A) + F(B) - F(A) \cdot F(B)$ applied to complements, or equivalently 1 minus the appropriate truth combination in the classical limiting case) reproduces 0.49 as well once the same combination rule is applied consistently to both sides — confirming the identity holds numerically once both the joint and the complement

computations use the same underlying (T, I, F) combination convention, exactly as promised in principle in §3.3.

VI.3 Independence testing under Refined Plithogenic Statistics (Family 9)

Return to the telecom churn dataset (§9.2, refined at the service-complaints branch): sample degrees for “network outage complaints” (0.20) and “billing-error complaints” (0.15) within the parent “service complaints” (0.35). Suppose the churn team additionally has a sample contingency table crossing complaint type against a binary churn outcome over 10,000 customers:

Network-outage complainers: 1,400 total, 560 churned (40.0% churn rate) Billing-error complainers: 1,050 total, 220 churned (21.0% churn rate) Non-complainers: 7,550 total, 906 churned (12.0% churn rate baseline)

A chi-square test of independence between complaint sub-type and churn outcome (df=2) on this table yields $\chi^2 \approx 612$ (using the standard $\Sigma(\text{observed}-\text{expected})^2/\text{expected}$ formula against the pooled 15.4% overall churn rate), overwhelmingly rejecting independence at any conventional significance threshold — confirming quantitatively what §9.2 stated qualitatively: network-outage complaints are far more predictive of churn (40.0% vs. a 15.4% baseline, a 2.6× lift) than billing-error complaints (21.0%, a 1.4× lift), a distinction the unrefined aggregate “service complaints, 0.35” would have reported as a single undifferentiated 32.0% pooled churn rate, obscuring the fact that the two sub-types are not remotely equivalent risk signals.

VI.4 Conditional probability under Over/Under/Off-Refined Imprecise Probability (Family 2)

Return to the sensor-fault example (§2.2, second application via §2.2’s imprecise-probability medical panel structure applied here to the engineering case in §7.2): suppose the fault-window sensor reports $[\underline{P}, \bar{P}] = [-0.05, 1.10]$ for “true reading exceeds nominal safe threshold,” refined into calibration-drift and power-supply-noise sub-intervals $[-0.05, 0.40]$ and $[0.60, 0.70]$ respectively (summing bound-wise per §2.1’s reduction constraint: $-0.05+0.60=0.55 \neq -0.05$ — illustrating that when refined sub-intervals are Off-valued, the *bounds* still add linearly even though the resulting combined interval width does not shrink simply, since $[0.55, 1.10]$ is the correctly summed result, not the naive union). Given a maintenance log indicating the fault occurred during a known power-grid brownout event B, with $[\underline{P}(B), \bar{P}(B)] = [0.85, 0.95]$ (near-certain that a brownout occurred in the fault window), the generalized interval Bayes rule from §2.3 gives $\underline{P}(\text{power-supply-noise} \mid B) = \underline{P}(\text{power-supply-noise} \cap B)/\bar{P}(B)$ and $\bar{P}(\text{power-supply-noise} \mid B) = \bar{P}(\text{power-supply-noise} \cap B)/\underline{P}(B)$. Taking the conjunction bounds as $[0.55, 0.68]$ (the noise sub-interval narrowed slightly by the brownout evidence), this yields $[0.55/0.95, 0.68/0.85] = [0.579, 0.800]$ — a materially narrower and higher-confidence interval than the pre-brownout-evidence $[0.60, 0.70]$, correctly reflecting that the brownout evidence both raises and sharpens the maintenance team’s estimate of the noise-cause hypothesis.

Glossary of Notation

Symbol	Meaning	First used
P(A)	Classical probability of event A	Family 1

Symbol	Meaning	First used
A_i	i-th refined sub-component of event A	Family 1
$\underline{P}(A), \bar{P}(A)$	Lower and upper imprecise-probability bounds on A	Family 2
T, I, F	Neutrosophic Truth, Indeterminacy, Falsity components	Family 3
T, I, N, F	Quadruple neutrosophic Truth, pure Indeterminacy (without neutrality), Neutrality, Falsity	§3.6
$N = a + bI$	Neutrosophic number with indeterminate part bI	Family 8
v_i	i-th attribute-value in a plithogenic attribute	Family 4
$d(A, v_i)$	Degree of appurtenance of event/observation A to value v_i	Family 4
$c(v_i, v_{\text{dominant}})$	Contradiction degree of value v_i relative to the dominant value	Family 4
$\mu(A)$	Measure of set A	Family 5
$R(\cdot)$	The Refinement operator	Part I
$O(\cdot)$	The Over/Under/Off operator	Part I
$\bar{A}, A^c$	Complement of event A	All families
$P(A B)$	Conditional probability of A given B	All families
Σ, Π	Summation, product	All families

Comparison of the Nine Families

#	Family	Nominal range	Generalizes	Empirical counterpart	Primary connected field
1	Classical Probability	$[0,1], \Sigma=1$	— (base case)	Family 6	Frequentist statistics, decision theory
2	Imprecise Probability	$[\underline{P}, \bar{P}] \subseteq [0,1]$	Family 1	Family 7	Robust Bayesian statistics, metrology
3	Neutrosophic Probability	$(T,I,F) \in [0,1]^3$	Family 2	Family 8	Dempster-Shafer theory, risk analysis
4	Plithogenic Probability	attribute-values + $c \in [0,1]$	Family 3	Family 9	MCDM, MAUT, copula dependency modeling
5	Measure	$\mu(A) \geq 0$	— (foundation)	(none; underlies 6–9)	Kolmogorov axioms, signed measures
6	Classical Statistics	sample estimates of Family 1	Family 1, empirically	—	SPC/Six Sigma, experimental design
7	Imprecise Statistics	sample estimates of	Family 2,	—	Interval-censored

#	Family	Nominal range	Generalizes	Empirical counterpart	Primary connected field
		Family 2	empirically		survival analysis
8	Neutrosophic Statistics	sample estimates of Family 3	Family 3, empirically	—	Missing-data theory, epidemiology
9	Plithogenic Statistics	sample estimates of Family 4	Family 4, empirically	—	Credit scoring, sports analytics, ML

VI.5 Over/Under/Off-Measure worked example (Family 5)

Return to the environmental-impact measure (§5.2): $\mu(X) = 1$ for the pristine baseline watershed, with a degraded sub-region B at $\mu(B) = -0.40$ and a restored wetland C at $\mu(C) = 0.18$ against a regional baseline capacity of 0.12 (an Over-reading, since $0.18 > 0.12$). Suppose a third region D, an agricultural buffer strip, is measured at $\mu(D) = 0.09$, and the four regions B, C, D, and the remainder of the watershed R partition X exactly ($B \cup C \cup D \cup R = X$, pairwise disjoint). By countable additivity (§5.1), $\mu(R) = \mu(X) - \mu(B) - \mu(C) - \mu(D) = 1 - (-0.40) - 0.18 - 0.09 = 1 + 0.40 - 0.18 - 0.09 = 1.13$. Region R — everything in the watershed not separately singled out — is itself computed to have an Over-value (1.13, exceeding the baseline’s own total of 1), which is the correct and interpretable result: R’s positive reading is being pulled upward by the arithmetic necessity of offsetting B’s substantial negative contribution while the whole watershed is renormalized to sum to exactly 1. This illustrates a general feature of Over/Under/Off-Measure worth flagging explicitly: a single strongly Under-valued sub-region can force an Over-value onto the residual “everything else” category purely as an artifact of the additivity constraint, which is why environmental economists using this kind of accounting are advised to report sub-region values directly rather than relying on a residual category to represent “baseline, unaffected” conditions.

VI.6 Mutually exclusive events under Refined Classical Statistics (Family 6)

Return to the Six Sigma defect-tracking example (§6.2, second application): a production run of 50,000 units logs 17 defects, refined by type into 6 dimensional, 8 cosmetic, and 3 functional defects ($17 = 6+8+3$, confirming mutual exclusivity of the categories as required by §6.3’s rule $\Sigma(\text{category counts}) = n$ applied to the defect sub-population). The corresponding DPMO figures are 120 (dimensional), 160 (cosmetic), and 60 (functional), summing to 340 — matching the parent-level aggregate stated in §6.2. Suppose a process-improvement initiative specifically targets cosmetic defects and, over the following quarter’s 62,000-unit run, cosmetic defects fall to 6 (DPMO ≈ 97), while dimensional defects rise slightly to 8 (DPMO ≈ 129) and functional defects hold steady at 4 (DPMO ≈ 65), for a new aggregate DPMO ≈ 291 . A naive read of the aggregate (“DPMO fell from 340 to 291, a 14.4% improvement”) is technically correct but strategically incomplete: the refined breakdown shows the improvement initiative worked exactly as intended on its target category (cosmetic defects fell 39%) while dimensional defects quietly worsened by 7.5% in the same period, a trend the aggregate number alone would not have surfaced for at least another quarter or two of monitoring — a concrete illustration of why §6.5 recommends refined sub-category tracking as standard SPC practice rather than an optional enhancement.

Part VII — Common Pitfalls and Frequently Asked Questions

Does refining an event or attribute ever change the underlying probability or measure? No, and this is worth stating unambiguously: the reduction constraint in every family's §X.1 ($P(A) = \sum P(A_i)$, $T = \sum T_i$, $\mu(A) = \sum \mu(A_i)$, etc.) is not optional — it is the defining property that makes a decomposition a *refinement* rather than a different model entirely. If an analyst's refined sub-components do not sum back to the parent value, they have built a new, separate model, not refined the original one, and the two should not be presented as interchangeable.

Is an Over/Under/Off value ever a mistake to be corrected? Sometimes, and the framework does not claim otherwise — a genuinely miscalibrated sensor reporting -5 kPa on a gauge that cannot physically go negative (§7.2) is a fault to be diagnosed, not a value to be modeled as meaningful in its own right. The distinction throughout this document is between off-interval values that are *artifacts to be fixed* and off-interval values that are *deliberately informative signals* (the -5 kPa reading is itself useful evidence of the fault type, even though the reading is not a real pressure). Every worked example above was chosen because the off-interval value carries genuine information; the framework does not license treating every out-of-range data point as meaningful without that justification.

Can two different families be applied to the same dataset, and do they need to agree? Yes to the first question, and not necessarily to the second. The insurance-underwriting data used in §4.2 could equally be modeled with plain Classical Probability (ignoring the attribute structure entirely) or with Neutrosophic Probability (collapsing the three risk values into a single T/I/F triple); different families answer different questions from the same data, and there is no requirement that a coarser model's output be recoverable from a finer one beyond the specific degenerate-case relationships stated in each family's §X.4 Relations section.

Why does De Morgan's Rule always hold when nothing else is guaranteed to? Because, as stated in the closing paragraph of Part IV, De Morgan's Rule is a theorem about set algebra (unions, intersections, and complements of events as sets), not about the numbers attached to those sets. Every family in this document assigns different kinds of numbers to the same underlying classical, two-valued event algebra; refining or relaxing the numbers cannot break a theorem that never depended on the numbers to begin with. A reader who finds De Morgan's Rule failing to hold numerically in a specific computation should treat that as a signal to recheck the arithmetic, not as evidence the rule itself needs adjustment for that family.

How should an analyst choose the fusion operator for a plithogenic model, since §4.3 notes it is not fixed by the mathematics? There is no universally correct choice, and this is the one place in the framework where domain judgment substitutes for a derivable rule. A contradiction-degree-weighted average is the most common default and is what every worked plithogenic example in this document uses implicitly; alternatives (minimum/maximum-based fusion, more elaborate T-norm and T-conorm operators drawn from the broader fuzzy-logic literature) are all legitimate but must be stated explicitly and held fixed across a given analysis, since switching fusion operators mid-analysis will generally produce different De Morgan-consistency results even for identical input data — precisely the caveat flagged in §4.3.

Does a high Indeterminacy (I) value mean an event is roughly equally likely to be true or false? No — this is one of the more common misreadings of the neutrosophic triple. A high I means the *evidence quality* is poor or contested, which is independent of what T and F themselves say. $NP(A) = (0.05, 0.90, 0.05)$ (nearly all indeterminate, T and F both very low) and $NP(A) = (0.45, 0.10, 0.45)$ (T and F nearly equal, but well-supported and mostly resolved) both have $T \approx F$, yet describe very different epistemic situations — the first is genuinely unresolved, the second is a well-evidenced near-tie. Refining I, as recommended throughout Families 3 and 8, is precisely the tool for distinguishing these cases in practice rather than letting a single scalar conflate them.

If two families disagree on whether two events are independent, which one is right? Neither is more “right” in the abstract — independence is defined differently in each family (a product test in Family 1, several competing tests in Family 2, a componentwise T/I/F test in Family 3, an attribute-value-local test in Family 4), so the same underlying data can legitimately be judged independent under one family’s test and dependent under another’s. This is not a contradiction in the framework; it reflects that “independence” is itself a modeling choice about which aspects of two events’ relationship are considered relevant, and different families are built to be sensitive to different aspects of that relationship.

Is this framework a replacement for classical statistics, or an addition to it? An addition. Every family in this document reduces to a classical or near-classical object in the appropriate degenerate limit (§1.4, §5.4, and the comparison table above make this explicit family by family), and none of the nine families claims that classical methods are wrong where they already apply cleanly — a fully-observed, unambiguous dataset with no attribute interactions is exactly the setting where Classical Statistics (Family 6) is already the right and sufficient tool, and reaching for Plithogenic Statistics in that setting would add modeling burden without adding insight.

Appendix A — Index of Worked Applications by Field

The examples throughout this document span a wide range of application domains. This index collects them by field for readers who want to jump directly to the examples closest to their own work.

Insurance and risk management: factory quality-control (§1.2), property underwriting (§4.2, §4.5), portfolio risk (§4.2), environmental-impact accounting (§5.2, VI.5), petroleum reserves (§3.2).

Finance and economics: sports odds-making (§1.2, VI.1), economic index nowcasting (§7.2), economic value-added accounting (§5.2), consumer credit scoring (§9.2, §9.5), Knightian uncertainty and Ellsberg’s paradox (§V.11).

Healthcare and life sciences: medical diagnostic panels (§2.2), sensor calibration (§2.2), interval-censored clinical trials (§7.2), vaccine side-effect surveys (§8.2), differential diagnosis (§4.2), agricultural yield-loss insurance (§4.2), quadruple-valued treatment-regimen evolution tracking (§3.6).

Biology, agriculture, and the life sciences broadly: wheat-cultivar breeding trade-offs (§3.6), and — surveyed qualitatively rather than numerically — botany, ecology, paleontology, genetics, biogeography, coevolution, immunology, molecular biology, and embryology, all via the fourteen-field survey in [16] referenced at §3.6.

Public policy, law, and governance: criminal-trial evidence assessment (§3.2, VI.2), legal standards of proof (§2.2), election polling (§8.2), climate-science uncertainty language (§V.15), geopolitical intelligence analysis (§3.2).

Technology and engineering: cybersecurity threat modeling (§V.16), sensor-network fault diagnosis (§7.2, VI.4), audio-engineering loudness measurement (§5.2), A/B testing (§6.2), manufacturing defect tracking (§6.2, VI.6), supplier quality auditing (§4.2).

Business operations: call-center satisfaction surveys (§6.2), telecom churn modeling (§9.2, Part IV), Net Promoter Score surveys (§8.2), supply-chain supplier selection (§4.2).

Social science and education: university dropout risk (Part IV), pedagogical decision-making under uncertainty [8], weather-forecasting skill scores (§1.2).

Sports and entertainment: sportsbook odds (§1.2, VI.1), basketball player evaluation (§9.2, §9.5).

Appendix B — Suggested Further Reading by Connected Field

This appendix supplements the numbered References (which cite only the Smarandache-authored sources this document directly draws its definitions from) with pointers into the broader established literatures discussed in Part V, for readers who want to go deeper into any one connected field.

Classical probability and measure theory: A.N. Kolmogorov’s 1933 *Grundbegriffe der Wahrscheinlichkeitsrechnung* remains the definitive axiomatic starting point; most graduate probability textbooks (e.g. Billingsley’s *Probability and Measure*) present the measure-theoretic foundation this document’s Family 5 builds on directly.

Imprecise probability: Peter Walley’s 1991 *Statistical Reasoning with Imprecise Probabilities* is the foundational text for the behavioral/coherence approach to interval probability referenced throughout Family 2 and §2.5.

Dempster-Shafer evidence theory and DSMT: Arthur Dempster’s original 1967 upper-and-lower-probability papers and Glenn Shafer’s 1976 *A Mathematical Theory of Evidence* remain the standard references for classical DST; Jean Dezert and Florentin Smarandache’s *Dezert-Smarandache Theory (DSMT) of Plausible, Paradoxist, and Neutrosophic Reasoning* [14], including the five-volume *Advances and Applications of DSMT for Information Fusion* collected works they edited together, is the standard reference for the more general extension directly relevant to §3.5 and §V.3.

Fuzzy sets and their descendants: Lotfi Zadeh’s 1965 “Fuzzy Sets” paper founded the field; Krassimir Atanassov’s 1986 intuitionistic fuzzy set extension is the direct bridge to the neutrosophic lineage discussed in §V.4.

Rough set theory: Zdzisław Pawlak’s 1982 “Rough Sets” paper and his 1991 book of the same broad theme are the standard entry points referenced in §V.5.

Possibility theory: Didier Dubois and Henri Prade’s *Possibility Theory* (1988) is the standard synthesis referenced in §V.6, building on Zadeh’s earlier possibility-theory sketches.

Grey system theory: Julong Deng’s 1982 founding paper “Control Problems of Grey Systems” introduced grey numbers and grey relational analysis, referenced in §V.6.

Multi-criteria decision-making: Thomas Saaty’s Analytic Hierarchy Process (1980) and the TOPSIS method (Hwang and Yoon, 1981) are the standard references for the MCDM connections drawn throughout Family 4 and §V.7.

Machine learning uncertainty quantification: Yarin Gal’s work on Bayesian deep learning and Monte Carlo dropout, and Vladimir Vovk’s foundational work on conformal prediction, are the standard modern references for §V.8.

Actuarial and financial risk theory: Harry Markowitz’s 1952 portfolio theory and the subsequent copula literature (Roger Nelsen’s *An Introduction to Copulas* is a standard textbook) underlie the connections drawn in §4.5 and §V.9.

Knightian uncertainty and behavioral economics: Frank Knight’s 1921 *Risk, Uncertainty and Profit*, Daniel Ellsberg’s 1961 paradox paper, and Daniel Kahneman and Amos Tversky’s 1979 prospect theory paper together frame the discussion in §V.11.

Information fusion: the Joint Directors of Laboratories (JDL) data fusion model, documented in successive U.S. Department of Defense technical reports since the 1980s and widely summarized in subsequent survey papers, underlies the discussion in §V.12 and connects directly to reference [5].

Philosophy of probability: Ian Hacking’s *The Emergence of Probability* and Alan Hájek’s Stanford Encyclopedia of Philosophy entries on interpretations of probability are accessible standard references for the aleatory/epistemic distinction discussed in §V.13.

Quadruple (T, I, N, F) neutrosophic theory: F. Smarandache’s founding “(T, I, N, F) Neutrosophic Set and Logic” paper [15] and the follow-on *Refined (T, I, N, F) Neutrosophic Theory of Evolution* [16], together with the earlier *Quadruple Neutrosophic Theory and Applications* volume by Smarandache, Şahin, Uluçay, and Kargin, are the standard references for the Neutrality/Indeterminacy split developed in §3.6.

Part VIII — Limitations and Open Research Questions

No framework this broad should be presented without an honest account of where it is incomplete or unproven, and this closing part collects the limitations flagged individually throughout the document into one place, alongside questions the document deliberately leaves open.

Estimation burden grows with generality. As noted in §V.4, each additional degree of freedom (moving from Classical to Imprecise to Neutrosophic to Plithogenic) requires

proportionally more data to estimate reliably, and refinement compounds this further by multiplying the number of quantities that must be independently assessed. A five-attribute plithogenic model with pairwise contradiction degrees has on the order of fifteen free quantities before any refinement is applied at all; refining even two of those attributes into two sub-values each adds four more. Nothing in this document specifies a minimum sample size or estimation procedure for any of these quantities — that is a substantial, unresolved practical question for any real deployment, and readers should treat every numeric example in this document as illustrative of the *structure* of a calculation, not as a demonstration that the underlying quantities are easy or cheap to estimate from real data.

The fusion operator choice is unprincipled. Flagged explicitly in §4.3 and again in Part VII's FAQ, the plithogenic fusion operator (how attribute-value degrees combine into an overall score) is a modeling choice, not a derived consequence of the framework's axioms. This is not unique to the plithogenic case — MCDM methods have the same issue (AHP, TOPSIS, and VIKOR can disagree on which alternative ranks first for the same input data) — but it means Families 4 and 9 inherit a genuine, unresolved methodological weak point from the broader MCDM literature rather than resolving it.

The philosophical status of Indeterminacy (I) is unsettled. Section V.13 already flags that I is naturally read as an epistemic rather than aleatory quantity, but does not resolve exactly what kind of epistemic quantity it is — a subjective degree of belief, a frequency of expert disagreement, an entropy measure (§V.14), or something else entirely. Different application domains in this document implicitly use different readings (the election-polling example treats I as a population proportion of undecided respondents, a broadly frequentist reading; the intelligence-analysis example treats I as an analyst's subjective assessment of evidence quality, a broadly subjectivist reading) without reconciling them, which is defensible as pragmatic flexibility but should not be mistaken for philosophical resolution.

Over/Under/Off values lack a general axiomatic foundation. Family 5's Over/Under/Off-Measure formally coincides with pre-existing signed-measure theory for its Under half (§5.5), but the Over half — exceeding a model-relative baseline total — has no comparably mature pre-existing mathematical theory to appeal to, and this document does not supply one; each family's Over-value examples are justified case by case (redundant confirming evidence, protective interactions, super-normal contribution) rather than from a single unifying axiom system analogous to Kolmogorov's for classical probability. Developing such a unifying axiomatization, if one exists, is squarely an open research question.

Empirical validation is absent. Every application example in this document is illustrative rather than empirically fitted to real data — the churn, credit-scoring, and clinical-trial numbers throughout are chosen to be realistic and internally consistent, not drawn from an actual dataset and validated against held-out performance. Whether Refined Plithogenic Statistics genuinely outperforms a well-specified classical model (say, a gradient-boosted tree with explicit interaction terms, as discussed in §9.5) on a real prediction task is an empirical question this document does not attempt to answer, and readers considering deployment should treat that comparison as necessary future work rather than assumed.

Open question: can the nine families be unified under a single formal calculus? Part III's synthesis diagram shows how the families relate as special cases of one another, but does

not provide a single unifying algebra from which all nine could be derived as instances of one parameterized construction — an appealing target for future formal work, given how cleanly the degenerate-case relationships in each family’s §X.4 already suggest such a unification might exist.

VI.7 Contradiction-degree fusion under Refined Plithogenic Probability (Family 4)

Return to the supplier-selection example (§4.2, second application): attribute “supplier viability,” dominant value cost ($d=0.40$), with quality ($d=0.35$, refined into defect rate $d=0.22$ and certification compliance $d=0.13$) and delivery reliability ($d=0.25$), contradiction degrees relative to cost $c = (0, 0.5, 0.4)$. Using the standard contradiction-degree-weighted fusion operator flagged in §4.3 and §7 of Part VII’s FAQ — a weighted average where each attribute-value’s contribution is scaled by $(1 - c)$ — the fused viability score for a candidate supplier reporting raw sub-scores of 0.80 (cost), 0.75 (defect rate), 0.60 (certification compliance), and 0.70 (delivery reliability) is computed attribute-value by attribute-value: cost contributes $0.80 \times (1-0) \times 0.40 = 0.320$; defect rate contributes $0.75 \times (1-0.5) \times 0.22 = 0.0825$; certification compliance contributes $0.60 \times (1-0.5) \times 0.13 = 0.039$; delivery reliability contributes $0.70 \times (1-0.4) \times 0.25 = 0.105$. Summing: $0.320 + 0.0825 + 0.039 + 0.105 = 0.5465$, a fused viability score of roughly 0.55 on the same $[0,1]$ scale as the individual sub-scores. Now suppose this same supplier is later found to have won its low bid specifically by cutting corners on certification, and the procurement team assigns a revised contradiction degree between certification compliance and cost of $c = 1.2$ (an Over-value, per §4.1’s OUO-PP definition) to reflect that the low cost is now understood to be actively predictive of the certification shortfall rather than merely correlated with it. Recomputing certification compliance’s contribution: $0.60 \times (1-1.2) \times 0.13 = 0.60 \times (-0.2) \times 0.13 = -0.0156$ — the term flips sign entirely, and the fused score drops to $0.320 + 0.0825 + (-0.0156) + 0.105 = 0.4919$, a full 5.5-percentage-point penalty relative to the original 0.5465 score purely from the contradiction-degree revision, with no change at all to any of the four raw sub-scores. This is the concrete mechanical illustration of the claim made abstractly in §4.1: an Over-valued contradiction degree does not just down-weight a component’s contribution to the fusion sum, it can flip that contribution’s sign entirely, penalizing rather than merely discounting the supplier’s certification performance.

VI.8 Complementary events under Over/Under/Off-Neutrosophic Statistics (Family 8)

Return to the Net Promoter Score example (§8.2, third application): $NS(\text{promoter}) = (T=0.35, I=0.20, F=0.45)$, later revised for an incentivized batch to $T=1.08$ (an Over-value, per §8.1’s OUO-NS definition) against the company’s usual baseline. Applying the neutrosophic complementation rule from §3.3 and §8.3, $NS(\text{detractor})$ — the complementary category — swaps T and F while I stays fixed: $NS(\text{detractor}) = (F, I, T) = (0.45, 0.20, 0.35)$ under the original, unrevised figures. Once the incentivized batch’s T is revised to 1.08, the complement’s F component inherits that same Over-value: $NS(\text{detractor, revised}) = (F=1.08, I=0.20, T=0.45)$ — meaning the revised statistic implies detractor status is now scored as “more strongly disconfirmed than ordinary falsity,” exactly mirroring the promoter side’s inflated confidence. This is the concrete numerical confirmation of the general point made in §8.3: the complementary swap rule holds exactly under Over/Under/Off relaxation, propagating an inflated or deflated component from one side of the complement to the other without requiring any separate recalculation — a useful practical shortcut when auditing a revised statistic for internal consistency, since a complement that fails to swap correctly is a reliable signal that an

error was introduced somewhere in the revision process rather than evidence that the swap rule itself needs adjustment.

Part IX — Conclusion

Nine families, two generating operators, and roughly forty worked applications later, the throughline stated at the very end of Part IV is worth restating as the document's closing word: refinement and Over/Under/Off relaxation never touch the underlying classical, two-valued event algebra — they only change how much weight, and how many separately-trackable pieces of weight, get attached to each part of that algebra. Everything else in this document — every insurance premium, every churn model, every criminal-trial assessment, every worked Bayes calculation in Part VI — is an illustration of that one structural fact playing out across progressively richer numerical vocabularies, from a single bounded probability through an interval, through a three-valued evidence triple, through a fully general attribute-value structure, and back down through each of their four empirical statistics counterparts.

The practical payoff, restated plainly: an analyst who reaches for Refined Plithogenic Statistics on a dataset that a plain contingency table would have handled just as well has not gained anything by doing so, and Part VIII's limitations section says so explicitly rather than leaving that judgment to be inferred. The payoff shows up exactly when a real phenomenon resists the bounded, unrefined, single-valued treatment that classical statistics defaults to — a protective interaction that looks like independence until the contradiction degree is allowed to go negative (§9.2), an indeterminate survey response that behaves nothing like a missing value once its source is refined (§8.2), a signed measure that needs to represent net harm, not just diminished benefit (§5.2). Part V's seventeen field connections exist to help a reader recognize, from their own disciplinary vocabulary, when they are already looking at exactly that kind of phenomenon under a different name — Knightian uncertainty, a Dempster-Shafer mass function, an IPCC confidence rating, a copula's tail-dependency parameter — and to make the translation into this document's notation, and back out again into their own field's established tools, as direct as the comparison table in the Appendix aims to make it.

References

[1] F. Smarandache, *Refined Fuzzy Set, Refined Intuitionistic Fuzzy Set, Refined Inconsistent Intuitionistic Fuzzy Set (Refined Picture Fuzzy Set, Refined Ternary Fuzzy Set), Refined Pythagorean Fuzzy Set (Refined Atanassov's Intuitionistic Fuzzy Set of Type 2), Refined Spherical Fuzzy Set, Refined n-HyperSpherical Fuzzy Set, and Refined q-Rung Orthopair Fuzzy Set*, 2016. <https://fs.unm.edu/Raspunsatan.pdf>

[2] F. Smarandache, "n-Valued Refined Neutrosophic Logic and Its Applications in Physics," *Progress in Physics*, Vol. 4, pp. 143–146, 2013.

<https://fs.unm.edu/n-ValuedNeutrosophicLogic-PiP.pdf>

[3] F. Smarandache, *Neutrosophic Overset, Neutrosophic Underset, and Neutrosophic Offset*, 2007 onward. <https://fs.unm.edu/Over-Under-Off.htm> ; <https://fs.unm.edu/NeutrosophicOversetUndersetOffset.pdf> ; <https://fs.unm.edu/SVNeutrosophicOverset-JMI.pdf> ;

<https://fs.unm.edu/IV-Neutrosophic-Overset-Underset-Offset.pdf>

[4] F. Smarandache, *Degrees of Over-, Under-, and Off-Membership*.
<https://fs.unm.edu/NSS/DegreesOf-Over-Under-Off-Membership.pdf>

[5] F. Smarandache et al., "Introduction of Over/Under/Off-Masses in Information Fusion," *Neutrosophic Sets and Systems*, 2026. <https://fs.unm.edu/DSmT/16OverUnderOffMasses-NSS.pdf>

[6] F. Smarandache, *Plithogeny, Plithogenic Set, Logic, Probability, and Statistics*, Pons, Brussels, Belgium, 2017, 141 p. <https://fs.unm.edu/P/Plithogeny.pdf>

[7] F. Smarandache, M. Abdel-Basset (eds.), *Optimization Theory Based on Neutrosophic and Plithogenic Sets*, Elsevier, Academic Press, 2020, 446 p.

[8] W. Ortega Chavez, F. Pozo Ortega, J. K. Vasquez Perez, E. J. Diaz Zuniga, A. R. Patino Rivera, *Modelo ecológico de Bronfenbrenner aplicado a la pedagogía, modelación matemática para la toma de decisiones bajo incertidumbre: de la lógica difusa a la lógica plitogénica*, NSIA Publishing House Editions, Huánuco, Peru, 2021, 144 p.

<https://fs.unm.edu/P/LogicaPlitogenica.pdf>

[9] F. Smarandache (Special Issue Editor), *New Types of Neutrosophic Set/Logic/Probability, Neutrosophic Over-/Under-/Off-Set, Neutrosophic Refined Set, and Their Extension to Plithogenic Set/Logic/Probability, with Applications*, Special Issue of *Symmetry* (Basel, Switzerland), November 2019, 714 p.

https://www.mdpi.com/journal/symmetry/special_issues/Neutrosophic_Set_Logic_Probability

[10] F. Smarandache, "Neutrosophic Set is a Generalization of Intuitionistic Fuzzy Set, Inconsistent Intuitionistic Fuzzy Set (Picture Fuzzy Set, Ternary Fuzzy Set), Pythagorean Fuzzy Set (Atanassov's Intuitionistic Fuzzy Set of Second Type), q-Rung Orthopair Fuzzy Set, Spherical Fuzzy Set, and n-HyperSpherical Fuzzy Set, while Neutrosophication is a Generalization of Regret Theory, Grey System Theory, and Three-Ways Decision (revisited)," *Journal of New Theory* 29 (2019), pp. 1–35. https://digitalrepository.unm.edu/math_fsp/21

[11] F. Smarandache, *Introduction to Neutrosophic Statistics*.
<https://fs.unm.edu/NeutrosophicStatistics.pdf>

[12] F. Smarandache, *Refinement of the Uncertain Set*.
<https://fs.unm.edu/NK/RefinedUncertainSet-NK.pdf>

[13] F. Smarandache, *Neutrosophic Measure and Neutrosophic Integral, Neutrosophic Probability*. <https://fs.unm.edu/NeutrosophicMeasureIntegralProbability.pdf>

[14] J. Dezert, F. Smarandache, *Dezert-Smarandache Theory (DSmT) of Plausible, Paradoxist, and Neutrosophic Reasoning*. <https://fs.unm.edu/DSmT.htm> ; *An Introduction to DSmT for Information Fusion*. <https://fs.unm.edu/IntroductionToDSmT.pdf>

[15] F. Smarandache, "(T, I, N, F) Neutrosophic Set and Logic (Truth, Indeterminacy, Neutrality, Falsehood)," *Critical Review*, Vol. XIV, 2017, pp. 20–28. <https://fs.unm.edu/CR/TINF-NeutrosophicSetLogic.pdf>

[16] F. Smarandache, *Refined (T, I, N, F) Neutrosophic Theory of Evolution: Cross-Domain Formalization, Worked Examples, and a Weighted Aggregation Model*.
<https://fs.unm.edu/preprints/RefinedTINFNeutrosophicTheoryOfEvolution.pdf>

ABOUT THIS BOOK

Real uncertainty rarely fits inside a single bounded number. This book unifies two composable operators — Refinement and Over/Under/Off relaxation — into nine coherent families spanning Classical, Imprecise, Neutrosophic, and Plithogenic Probability, Measure, and their empirical Statistics counterparts. Each family is developed with formal definitions, worked real-world applications, an event-type study, and explicit connections to seventeen established fields, from Bayesian inference, Dempster-Shafer Theory and Dezert-Smarandache Theory to machine-learning uncertainty quantification and quantum probability.

ABOUT THE AUTHOR

Florentin Smarandache is a professor of mathematics at the University of New Mexico, United States. He is the founder of neutrosophy, neutrosophic logic, set, probability, and statistics, and of plithogeny, plithogenic set, logic, probability, and statistics, with applications across engineering, decision-making, and the social sciences.

Peer-reviewed by Mukhtar Ahmad (Institut Teknologi Bandung, Indonesia) and Victor Christianito (Independent Researcher, Indonesia)



NSIA Publishing House — University of New Mexico & University of Guayaquil

