



EXTENDING (T, I, N, F)

NEUTROSOPHIC LOGIC

to Measure, Probability, and Statistics

A Unified Framework, Worked Examples, and a Research Program

Truth · Pure Indeterminacy · Neutrality · Falsehood

Second, Expanded Edition

Acad. Florentin Smarandache

University of New Mexico, Mathematics and Science Department

Gallup, NM 87301, USA · smarand@unm.edu

Member, Academy of Science Peloritana dei Pericolanti, University of Messina, Italy

Mathematical development drafted with AI-assisted research synthesis (ChatGPT, Gemini, and Claude)

August 2026

NSIA PUBLISHING HOUSE

NSIA

Publishing House

***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 978-1-972502-44-0

Peer-Reviewers:

Maikel Leyva-Vázquez
Universidad de Guayaquil, Guayas, ECUADOR
maikel.leyvav@ug.edu.ec

Victor Christianto
Malang Institute of Agriculture (IPM), Malang, INDONESIA
victorchristianto@gmail.com

© 2026 Florentin Smarandache. Published by NSIA Publishing House.
This is a research synthesis and extension exercise, reviewed by the peer reviewers listed above and offered as a starting point for further, more rigorous work.

Extending (T, I, N, F) Neutrosophic Logic to Measure, Probability, and Statistics

Building on: Florentin Smarandache, "(T, I, N, F) Neutrosophic Set and Logic", Critical Review, Vol. XIV, 2017, pp. 20–28. <https://fs.unm.edu/CR/TINF-NeutrosophicSetLogic.pdf>

Abstract

The 2017 paper cited above refines the classical neutrosophic triple (Truth, Indeterminacy, Falsehood) into a quadruple (T, I, N, F) by splitting Indeterminacy into Pure Indeterminacy (I — genuine unknown-ness, absence of information) and Neutrality (N — a real third stance that is neither true nor false by nature, not by ignorance). It works this refinement out fully at the level of Sets and Logic — giving the operators, negation, and a total-order ranking system — and names Measure, Probability, and Statistics as the natural next layer without developing them.

This booklet carries that refinement upward through the rest of the classical mathematical stack: measure theory, integration, probability (in both a non-normalized, evidence-theoretic reading and a normalized, partition reading), random variables and distributions, descriptive and inferential statistics, and a set of reduction theorems showing how classical probability, fuzzy sets, and ordinary three-component neutrosophic theory all fall out as special cases. This expanded second edition adds a full chapter placing TINF measure, probability, and statistics alongside twelve other major uncertainty frameworks — Dempster-Shafer evidence theory, fuzzy measures, possibility theory, imprecise probability, rough sets, several orthopair fuzzy-set families, Belnap-Dunn four-valued logic, Bayesian and frequentist statistics, robust and interval statistics, random sets, quantum probability, and plithogenic statistics — and more than doubles the worked-application chapter to twenty-seven fields, spanning medicine, engineering, defense, finance, law, epidemiology, political science, biology, sport, insurance, robotics, logistics, education, criminal justice, astronomy, agriculture, energy, misinformation research, decentralized finance, and human-resources analytics.

This is a synthesis and extension exercise, reviewed by the peer reviewers listed on the imprint page, and it should be read as a research proposal and worked-example collection rather than a finished, exhaustively refereed theory. Wherever the underlying mathematics needed a choice (e.g., normalized vs. non-normalized probability, or which t-norm to use), that choice is flagged explicitly rather than presented as settled.

Preface: Scope, Method, and a Terminology Note

This booklet extends the author's own 2017 paper on (T, I, N, F) Neutrosophic Set and Logic into the Measure/Probability/Statistics layers that paper named but did not build. The mathematical development was prepared with substantial AI-assisted drafting and research synthesis — several independent passes with different AI assistants (ChatGPT, Gemini, and Claude) were generated, compared, and merged under the author's direction and review, and are now extended a second time with a dedicated survey of how the (T, I, N, F) framework sits alongside the other major mathematical theories of uncertainty. Rather than picking one draft, this version merges them: the more rigorous definition/axiom/proof scaffolding from one line of drafting, the worked numerical examples and t-norm/t-conorm mechanics from another, the additional application domains suggested across all of them, and — new in this edition — an explicit, chapter-length map of how TINF measure, probability, and statistics relate to Dempster-Shafer evidence theory, fuzzy and possibility theory, imprecise probability, rough sets, Belnap-Dunn four-valued logic, Bayesian and frequentist statistics, and several other established frameworks.

What is genuinely new here, and what is not. To be precise about credit: the source paper explicitly introduces the term "(T, I, N, F) Neutrosophic Set/Logic/Measure/Probability/Statistics" and states the refinement of I into I and N. What it does not do is work out the Measure/Probability/Statistics layers — no axioms, no random variables, no estimators, no worked numerical examples are given for those layers in the source, and it does not situate the quadruple relative to the wider uncertainty-theory literature. That is the gap this booklet fills. The claim made here is therefore narrower and more defensible than "this paper invented (T,I,N,F) probability"; it is "this booklet develops the axiomatic, operational, and applied theory that the 2017 paper named but did not build, and places that theory on the map of existing uncertainty formalisms."

A normalization choice, flagged explicitly. A genuine mathematical fork appears once one moves from logic to probability: should $T + I + N + F$ be forced to equal 1 (a true probability partition, compatible with Kolmogorov axioms and off-the-shelf statistical software), or left unconstrained in $[0, 4]$ (an evidence-theoretic reading, where the four numbers are independent "degrees of support" that may overlap, conflict, or leave gaps)? Both readings are useful for different purposes, so this booklet develops both and is explicit about which one is in play in each section. Following a naming suggestion surfaced during drafting, where precision matters we call the unconstrained quantities truth-support, indeterminacy, neutrality, and falsity measures/degrees rather than "probabilities" outright, and reserve "TINF probability" for the normalized case or use it loosely, as the source paper does, when the distinction does not matter to the point being made.

How to read this booklet. Chapters 1–3 lay the algebraic and measure-theoretic foundations. Chapters 4–6 build probability, random variables, and statistics on top of them, each with worked numerical examples. Chapter 7 shows the reduction theorems

that tie the new theory back to classical mathematics — the sanity check that any respectable generalization must pass. Chapter 8 covers the Over/Under/Off extension. Chapter 9 — expanded substantially in this edition — is a systematic tour of how TINF relates to twelve other measure/probability/uncertainty frameworks developed independently in statistics, artificial intelligence, logic, and physics, closing with a comparison table. Chapter 10 is a tour of applications, organized by field, now covering twenty-seven domains, each with a concrete illustrative quadruple. Chapter 11 lists open problems for anyone who wants to take this further, an appendix collects notation, and the References close the booklet.

Table of Contents

Preface: Scope, Method, and a Terminology Note.....	4
1. Introduction and Motivation.....	6
2. TINF Algebra.....	8
3. TINF Measure Theory.....	10
4. TINF Probability.....	12
5. TINF Random Variables and Distributions.....	15
6. TINF Statistics.....	17
7. Reduction Theorems.....	20
8. Over-, Under-, and Off-TINF Extension.....	22
9. TINF in the Landscape of Uncertainty Theories.....	23
10. Worked Applications Across Fields.....	34
11. Research Program and Open Problems.....	43
Appendix A: Notation Glossary.....	45
Appendix B: Worked Cross-Framework Conversions.....	46
Appendix C: Quick-Reference Decision Rule.....	47
Appendix D: Suggested Reading Path by Background.....	48
References.....	49

1. Introduction and Motivation

1.1 The (T, I, F) starting point

Classical neutrosophic logic represents the status of a proposition A by a triple (T, I, F): a degree of truth, a degree of indeterminacy, and a degree of falsehood, each a subset of [0, 1], with $0 \leq \inf(T) + \inf(I) + \inf(F) \leq \sup(T) + \sup(I) + \sup(F) \leq 3$ and no requirement that the three add to 1. This is already a strict generalization of fuzzy logic (which has only T, with $F = 1 - T$ and $I = 0$) and of classical logic ($T \in \{0,1\}$, $I = F = 0$). What it cannot do is distinguish two situations that feel qualitatively different: not knowing, and knowing that the honest answer is neither.

1.2 Splitting I into Pure Indeterminacy and Neutrality

The 2017 paper's refinement replaces (T, I, F) with the quadruple (T, I, N, F):

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

- T — Truth / support.
- I — Pure Indeterminacy: the available evidence is insufficient to establish either side. More data could, in principle, resolve it.
- N — Neutrality: the evidence is adequate, and it genuinely establishes a third stance — neither the proposition nor its negation. More data will not resolve it, because there is nothing left to resolve.
- F — Falsehood / opposition.

T and N are treated as "positive quality" (larger is more informative in a constructive sense); I and F are treated as "negative quality." Conjunction applies a t-norm to the (T, N) pair and the dual t-conorm to the (I, F) pair; disjunction does the reverse; and negation is the swap

$$\neg(T, I, N, F) = (F, N, I, T)$$

Notice what negation does to the middle two components: I and N trade places with each other, not with T or F. That single design choice — the negation of "we don't know" is "it's genuinely balanced," not "it's false" — is the crux of the whole refinement, and it is what every layer built below has to remain faithful to.

1.3 Two contrasts worth holding in mind

"The lab result is inconclusive because the sample was degraded" describes I: more testing, a fresh sample, or a better assay would resolve it. "The biopsy is read as atypical — genuinely between benign and malignant" describes N: a second reading of the same slide is not expected to move the number, because the tissue itself sits in a real borderline category. A model that reports only one "not-quite-sure" number tells the clinician to order a better test in both cases, which is the right advice for the first patient and the wrong advice for the second.

Likewise, "we don't yet know who will win the match" is I (weather, injuries, information not yet available), while "the referees are calling it scrupulously down the middle" is a statement about N — the honest, actively-maintained third position in a two-sided contest, which more information about the match itself will not change. This booklet returns to both examples repeatedly as a consistency check on the mathematics.

1.4 Why a dedicated connections chapter earns its place

A recurring, fair objection to any new uncertainty formalism is: does this really do something the existing toolbox cannot, or is it a relabeling of a known construction? Chapters 1-8 develop TINF theory on its own terms; Chapter 9, new in this edition, answers the objection directly by placing TINF measure, probability, and statistics side by side with twelve established frameworks — Dempster-Shafer evidence theory, fuzzy and possibility theory, imprecise probability, rough sets, several fuzzy-set variants, Belnap-Dunn logic, Bayesian and frequentist statistics, robust and interval statistics, random sets, quantum probability, and plithogenic statistics — showing precisely which special case of TINF recovers each one, and, just as importantly, what TINF can express that none of them can on its own: a stable third truth-value that is causally distinct from missing information, carried all the way through measure, probability, and inferential statistics rather than stopping at the level of sets and logic.

2. TINF Algebra

2.1 TINF values and TINF sets

A TINF value is an ordered quadruple $x = (T_x, I_x, N_x, F_x)$ with each component a number, interval, or hesitant subset of $[0, 1]$ (no normalization required at this stage). A TINF set on a universe X assigns such a quadruple to every element:

$$A = \{ \langle x; T_A(x), I_A(x), N_A(x), F_A(x) \rangle : x \in X \}$$

2.2 Complement

As in §1.2, the natural complement is $\neg(T, I, N, F) = (F, N, I, T)$, which is an involution: $\neg\neg(T, I, N, F) = (T, I, N, F)$. This is the operator that every measure- and probability-level complement below is built to match.

2.3 t-norms, t-conorms, and the TINF logical operators

Recall that a t-norm $T(x, y)$ generalizes AND (commutative, associative, monotone, $T(x, 1) = x$) and has a dual t-conorm $S(x, y) = 1 - T(1 - x, 1 - y)$ generalizing OR. Three standard pairs:

t-norm	Formula	Dual t-conorm
Minimum (Zadeh)	$T_M(x, y) = \min(x, y)$	$S_M(x, y) = \max(x, y)$
Product	$T_P(x, y) = x \cdot y$	$S_P(x, y) = x + y - xy$
Łukasiewicz	$T_L(x, y) = \max(0, x + y - 1)$	$S_L(x, y) = \min(1, x + y)$

Intersection (conjunction) applies the t-norm to the T and N components and the dual t-conorm to I and F ; union (disjunction) does the reverse:

$$(T_1, I_1, N_1, F_1) \wedge (T_2, I_2, N_2, F_2) = (T_1 \wedge T_2, I_1 \vee I_2, N_1 \wedge N_2, F_1 \vee F_2)$$

$$(T_1, I_1, N_1, F_1) \vee (T_2, I_2, N_2, F_2) = (T_1 \vee T_2, I_1 \wedge I_2, N_1 \vee N_2, F_1 \wedge F_2)$$

Implication and equivalence follow from these and the complement in the usual way: $A \rightarrow B \equiv \neg A \vee B$, and $A \leftrightarrow B \equiv (A \rightarrow B) \wedge (B \rightarrow A)$; both are used unchanged from the source paper and are not repeated in full here.

2.4 Score, accuracy, and certainty — ranking TINF quadruples

Because (T, I, N, F) is only partially ordered on its own, the source paper defines a cascade of scalar functions to obtain a total order — essential for any decision system that ultimately has to pick one option. For $(T, I, N, F) \in [0, 1]^4$:

$$\text{Score: } s(T, I, N, F) = (2 + T + N - I - F) / 4 \in [0, 1]$$

$$\text{Accuracy: } a(T, I, N, F) = T + N - F \in [-1, 2]$$

$$\text{Ext. certainty: } ec(T, I, N, F) = T + N \in [0, 2]$$

$$\text{Certainty: } c(T, I, N, F) = T \in [0, 1]$$

Two quadruples are ranked by score first; ties are broken by accuracy, then extended certainty, then certainty; and if all four agree the quadruples are shown to be identical (a short proof-by-substitution given in the source paper). Worked example — comparing

$$A = (0.75, 0.20, 0.10, 0.15) \quad B = (0.70, 0.05, 0.30, 0.20)$$

gives $s(A) = (2 + .75 + .10 - .20 - .15)/4 = 0.625$ and $s(B) = (2 + .70 + .30 - .05 - .20)/4 = 0.6875$, so B ranks above A despite having lower T — its far lower indeterminacy and much higher neutrality more than compensate. This is exactly the kind of case the four-component representation is built to catch: a T-only comparison, or a three-component (T, I, F) comparison that lumps N into I, would have ranked A above B or called them incomparable.

2.5 Algebraic structure: lattice and lattice-ordered features

It is worth making explicit, for the comparisons of Chapter 9, exactly what algebraic structure (T, I, N, F)-space carries. Under the componentwise partial order $(T, I, N, F) \leq (T', I', N', F')$ iff $T \leq T', I \geq I', N \leq N', F \geq F'$ — i.e. more truth-support and more neutrality is "higher," more indeterminacy and more falsity is "lower" — the $\wedge$ and $\vee$ operators of §2.3 make TINF-space a bounded lattice with top $(1, 0, 1, 0)$ and bottom $(0, 1, 0, 1)$, and $\neg$ is an order-reversing involution on it, exactly the algebraic skeleton shared by classical Boolean algebra, fuzzy-set lattices, and the (T, I, F) neutrosophic lattice one dimension down. This lattice-theoretic framing is what makes the reduction theorems of Chapter 7 and the cross-framework comparisons of Chapter 9 provable rather than merely suggestive: each competing theory is a sub-lattice or a quotient of this one under an explicit, statable condition.

3. TINF Measure Theory

3.1 Measurable TINF spaces

Let $(X, \check{A})$ be a measurable space (X a set, $\check{A}$ a σ -algebra of subsets, or a weaker collection when a full σ -algebra is not available — neutrosophic measures do not require one to be useful). A TINF measure is a map

$$\mu : \check{A} \rightarrow [0, \infty]^4, \quad \mu(A) = (\mu_T(A), \mu_I(A), \mu_N(A), \mu_F(A))$$

Component	Interpretation
$\mu_T(A)$	measure of the supporting / realized content of A
$\mu_I(A)$	measure of the purely indeterminate content of A
$\mu_N(A)$	measure of the genuinely neutral content of A
$\mu_F(A)$	measure of the opposing / excluded content of A

3.2 Axioms

- (M1) Grounding. $\mu(\emptyset) = (0, i_0, n_0, 0)$, where $i_0, n_0 \geq 0$ absorb whatever background indeterminacy or neutrality the measurement process itself cannot avoid (a calibration floor, a mandated abstention slot); in the fully crisp limit $i_0 = n_0 = 0$ and $\mu(\emptyset) = (0, 0, 0, 0)$, matching classical measure theory.
- (M2) Non-negativity. $\mu_T(A), \mu_I(A), \mu_N(A), \mu_F(A) \geq 0$ for every $A \in \check{A}$.
- (M3) Componentwise countable additivity. For pairwise disjoint $A_1, A_2, \dots \in \check{A}$, $\mu(\cup_k A_k) = \sum_k \mu(A_k)$, i.e. $\mu_T(\cup_k A_k) = \sum \mu_T(A_k)$, and likewise for I, N, F. For applications where strict additivity over-claims precision — e.g. when the four components come from partially overlapping evidence sources — a subadditive version built from the t-conorm of §2.3 in place of ordinary sum is the natural substitute; equality is recovered with the bounded-sum t-conorm $S_L(x, y) = \min(1, x + y)$, which behaves like ordinary addition for small values.
- (M4) Monotonicity on T and F only. If $A \subseteq B$ then $\mu_T(A) \leq \mu_T(B)$ and $\mu_F(A) \geq \mu_F(B)$. Deliberately, I and N are not required to move monotonically: enlarging A can introduce new ambiguity or new neutral territory just as easily as it can dilute old ambiguity away. This is the one place TINF measure structurally departs from classical measure, and it is intentional — forcing I, N to be monotone would silently assume that all indeterminacy/neutrality is "leftover mass," which is exactly the assumption the whole refinement exists to avoid.

3.3 Proposition (partial monotonicity) and proof

Proposition. Under (M1)–(M3) plus additivity, if $A \subseteq B$ then $\mu_T(A) \leq \mu_T(B)$; the general claim $\mu_F(A) \leq \mu_F(B)$ is false — rather $\mu_F(A) \geq \mu_F(B)$ under the natural reading; the T-direction proof is representative and given here.

Proof. Since $B = A \cup (B \setminus A)$ is a disjoint union, additivity gives $\mu_T(B) = \mu_T(A) + \mu_T(B \setminus A)$. By non-negativity, $\mu_T(B \setminus A) \geq 0$, hence $\mu_T(B) \geq \mu_T(A)$. The identical argument with F in place of T, combined with the fact that shrinking a set can only remove opposing evidence or leave it unchanged, gives that F is monotone decreasing under $\subseteq$ exactly when the evidence against A^c outside A is itself non-negative, which (M2) supplies. ■

3.4 TINF measurable functions and the TINF integral

A TINF-measurable function is $f(x) = (f_T(x), f_I(x), f_N(x), f_F(x))$. Its integral over $A \in \tilde{\mathcal{A}}$ is defined componentwise against the matching measure component:

$$\int_A f \, d\mu = \left(\int_A f_T \, d\mu_T, \int_A f_I \, d\mu_I, \int_A f_N \, d\mu_N, \int_A f_F \, d\mu_F \right)$$

Worked example. Let $f(x) = (2x, 0.2x, 0.1x, 0.5x)$ on $A = [0, 10]$ with each μ -component equal to ordinary Lebesgue measure. Then $\int_A f_T = 2 \cdot 50 = 100$, $\int_A f_I = 0.2 \cdot 50 = 10$, $\int_A f_N = 0.1 \cdot 50 = 5$, $\int_A f_F = 0.5 \cdot 50 = 25$ (using $\int_0^{10} x \, dx = 50$), giving $\int_A f \, d\mu = (100, 10, 5, 25)$. Convergence theorems (monotone convergence, dominated convergence) transfer to each of the four component integrals separately whenever the corresponding classical hypotheses hold componentwise; no genuinely new proof machinery is needed at this level of the theory, only bookkeeping across four coordinates instead of one.

3.5 Reduction cases, previewed

Setting $I(A) = N(A) = 0$ for every A collapses μ to a classical measure on T (with $F(A) = \text{total} - T(A)$ when a fixed total is imposed). Setting only $N(A) = 0$ recovers the ordinary (T, I, F) neutrosophic measure. Allowing a component to exceed 1 or fall below 0 gives the TINF Over-/Under-/Off-Measure of Chapter 8. These reductions are proved formally in Chapter 7.

3.6 TINF outer measure and regularity

For applications where $\tilde{\mathcal{A}}$ is not closed under complementation on the nose (crowdsourced or streaming data, where new categories of "neutral" outcome can appear mid-experiment), it is useful to define a TINF outer measure μ^* componentwise as the classical Carathéodory outer measure applied to each of $\mu_T, \mu_I, \mu_N, \mu_F$ separately, with a set A called TINF-measurable when it satisfies the Carathéodory splitting condition in all four components simultaneously. This gives a completion procedure for TINF measure spaces exactly parallel to the classical Lebesgue completion, and it is the construction implicitly assumed whenever an application in Chapter 10 reports TINF quantities on data collected incrementally.

4. TINF Probability

Two readings of TINF probability are useful for different purposes, and the fork between them is real — it is not resolved by picking a "correct" convention, so both are given here and each application in Chapter 10 states which one it uses.

4.1 General (non-normalized) TINF probability

For an event A in a random experiment, define

$$P_N(A) = (T_A, I_A, N_A, F_A), \quad T_A, I_A, N_A, F_A \in [0,1]$$

read as: T_A = degree/probability of evidence supporting occurrence, I_A = degree/probability that the status is genuinely unknown, N_A = degree/probability that the evidence is genuinely neutral, F_A = degree/probability of evidence supporting non-occurrence. Crucially, $T_A + I_A + N_A + F_A$ is not required to equal 1 — it may be less than, equal to, or greater than 1 (bounded above by 4). A total below 1 signals gaps (no source has weighed in on some aspect); a total above 1 signals genuine conflict or redundancy across sources (e.g. two independent diagnostics both confidently — and contradictorily — report). Example: $P_N(A) = (0.70, 0.12, 0.10, 0.25)$ sums to 1.17, which is not a contradiction under this reading — it says the evidence is somewhat redundant/overlapping, not that a probability axiom has been violated.

Because these four numbers are not required to be a Kolmogorov partition, they are more carefully called truth-support, indeterminacy, neutrality, and falsity degrees (or measures) in this reading, with "TINF probability" used as a convenient umbrella name rather than a claim of classical probability-axiom compliance.

4.2 Normalized TINF probability

For applications that need compatibility with standard probability tables, hypothesis-testing software, or a genuine partition of the sample space, impose

$$T + I + N + F = 1$$

Example: $P_N(A) = (0.60, 0.15, 0.10, 0.15)$. This version is a strict generalization of classical probability ($I = N = 0$, $F = 1 - T$), of interval/imprecise probability (T, F intervals, $I = N = 0$), and of ordinary (T,I,F) neutrosophic probability ($N = 0$, i.e. the classical framework's indeterminacy is exactly $I + N$ here — see the reduction theorem of §7.3).

4.3 Axioms (stated for the general reading; the normalized reading adds $T+I+N+F=1$)

- (P1) Boundedness. $P_N(A) \subseteq [0,1]^4$ componentwise for every event A .
- (P2) Certain / impossible events. $P_N(\Omega) = (1, i_0, n_0, 0)$; $P_N(\emptyset) = (0, i_0, n_0, 1)$, with i_0, n_0 the experiment's irreducible background indeterminacy/neutrality (0 in the crisp limit).

- (P3) Complement. $P_N(A^c) = (F_A, N_A, I_A, T_A)$ — truth and falsehood swap; indeterminacy and neutrality swap with each other, matching the §2.2 negation exactly.
- (P4) Combination. For a chosen t-norm/t-conorm pair, $P_N(A \cap B) = (T_A \wedge T_B, I_A \vee I_B, N_A \wedge N_B, F_A \vee F_B)$ and $P_N(A \cup B) = (T_A \vee T_B, I_A \wedge I_B, N_A \vee N_B, F_A \wedge F_B)$, with the product t-norm/probabilistic-sum pair recovering ordinary independent-event probability when $I=N=0$.

4.4 Conditional probability and independence

Conditioning on T and F proceeds exactly as in classical probability, $T_{A|B} = T_{A \cap B} / T_B$ and $F_{A|B} = F_{A \cap B} / F_B$ (with the usual caveats about $T_B, F_B \neq 0$), while I and N are conditioned on their own respective totals rather than folded into one denominator — a survey's don't-know rate and its deliberate-abstention rate should not be renormalized against a shared base, since they answer different questions. A and B are called TINF-independent when the §4.3(P4) intersection formula holds with the product t-norm on every component; the total-probability theorem and a componentwise TINF Bayes' theorem follow by the identical algebra as the classical case, applied four times over (once per component), whenever the component denominators are non-zero — $T_{H|E} = T_{E|H} \cdot T_H / T_E$, and analogously for I, N, F. A fully general TINF Bayes theorem that allows the four components to interact — rather than updating as four independent classical Bayes computations running in parallel — remains an open problem (Chapter 11).

4.5 Worked example — weather forecasting

Let $A = \{\text{rain tomorrow}\}$. Three models disagree and some station data is missing. Instead of a single $P(\text{rain}) = 0.65$, record

$$P_N(A) = (0.65, 0.15, 0.10, 0.30)$$

$T=0.65$ evidence for rain; $I=0.15$ missing/insufficient station data; $N=0.10$ genuinely balanced signal among the models that have reported; $F=0.30$ evidence against rain. This is strictly more useful operationally than a bare 65% figure, because it tells a forecaster whether to wait for more data (chase I down) or to stop waiting and hedge (N will not move).

4.6 Worked example — medical diagnosis and sensor/source fusion

Let $D = \{\text{patient has disease } D\}$. A single diagnostic system reports $D = (0.82, 0.08, 0.05, 0.20)$: strong supporting evidence, a little missing information, a small amount of genuinely neutral evidence, and moderate opposing evidence. Three independent diagnostic sources for the same patient might report

$$D_{\text{blood}} = (0.80, 0.10, 0.05, 0.15) \quad D_{\text{image}} = (0.70, 0.05, 0.20, 0.10) \quad D_{\text{genetic}} = (0.90, 0.03, 0.02, 0.10)$$

D_{image} stands out for its high N (0.20): imaging here is not simply "unclear" (that would be high I) — it is positively, legibly indeterminate in the between-benign-and-malignant sense of §1.3. Fusing the three sources with the §4.3(P4) operators (e.g. product t-norm on T, probabilistic sum on I/N/F) produces a single combined quadruple that a clinician

can act on without discarding the information about which source is driving the residual ambiguity.

4.7 Worked example — autonomous-vehicle sensor fusion

Three obstacle-detection sensors report:

$$S_1 = (0.90, 0.05, 0.02, 0.10) \quad S_2 = (0.65, 0.20, 0.10, 0.15) \quad S_3 = (0.40, 0.05, 0.45, 0.20)$$

Sensor 3 is not simply uncertain — its high N says the return is legitimately ambiguous between "obstacle" and "not obstacle" (a roadside object at a shallow radar angle, say), a state further data from the same sensor is unlikely to resolve. A decision layer can be built directly on the fused quadruple $O = (T_o, I_o, N_o, F_o)$:

Dominant component	Suggested system response
High T_o	brake / act as if the obstacle is present
High I_o	poll additional sensors before deciding
High N_o	slow down and remain cautious — more polling will not resolve this
High F_o	proceed normally

This is the single clearest illustration in this booklet of why I and N earn separate machinery rather than a shared "indeterminacy" bucket: they trigger different, non-interchangeable system responses.

4.8 TINF stochastic processes, previewed

Extending §4.1–4.7 across an index set t gives a TINF stochastic process $X(t) = (T(t), I(t), N(t), F(t))$; a TINF Markov chain, for instance, replaces a single transition-probability matrix with four coupled matrices, one per component, subject to (P1)–(P4) at every time step. Chapter 11 lists this as an open problem; §4.5's weather example and §9.9's epidemiological curve of Chapter 10 are both, implicitly, single time-slices of such a process.

5. TINF Random Variables and Distributions

5.1 Definition

A TINF random variable is a measurable map $X: \Omega \rightarrow \mathbb{R}^4$, $X(\omega) = (X_T(\omega), X_I(\omega), X_N(\omega), X_F(\omega))$ — four coupled ordinary random variables, one per component. A reliability observation, for instance, might be $X = (0.90, 0.05, 0.10, 0.15)$.

5.2 Distribution function and density

The TINF distribution and density (continuous case) are likewise componentwise:

$$F_N(x) = (F_T(x), F_I(x), F_N(x), F_F(x)) \quad f_N(x) = (f_T(x), f_I(x), f_N(x), f_F(x))$$

$$P_N(a < X < b) = (\int_a^b f_T dx, \int_a^b f_I dx, \int_a^b f_N dx, \int_a^b f_F dx)$$

5.3 Expectation and variance

$$E_N[X] = (E[X_T], E[X_I], E[X_N], E[X_F])$$

$$Var_N(X) = (Var(X_T), Var(X_I), Var(X_N), Var(X_F))$$

Worked example. Three sensor readings $X_1=(10,2,1,3)$, $X_2=(12,1,2,2)$, $X_3=(14,2,1,4)$ give $E_N[X] = (12, 5/3, 4/3, 3) \approx (12, 1.67, 1.33, 3)$: the average is not one number but a profile across all four dimensions, which is exactly the point — a fleet-reliability dashboard built on this would show not just "average confidence," but where that confidence is coming from.

5.4 Covariance and correlation

For two TINF variables X, Y , $Cov_N(X,Y) = (C_T, C_I, C_N, C_F)$ with $C_T = Cov(X_T, Y_T)$ and likewise for I, N, F ; correlation $\rho_N(X,Y) = (\rho_T, \rho_I, \rho_N, \rho_F)$ follows by the usual normalization. A result such as $\rho_N = (0.81, -0.20, 0.05, -0.67)$ is read componentwise: strong positive correlation between the two variables' truth-support, a slight negative relationship in their indeterminacy, essentially no relationship in their neutrality, and a strong negative relationship in their falsity — a single pooled correlation coefficient would average all four of these into one much less informative number.

5.5 Toward named TINF distributions

A substantive open research direction (see Chapter 11) is defining explicit TINF-normal, TINF-binomial, TINF-Poisson, TINF-exponential, TINF-Weibull, and TINF-gamma families — i.e., families where the four component distributions are not merely stacked independently but are linked by an explicit, application-motivated dependency structure (e.g. a TINF-Weibull reliability model in which a rising hazard in F is mechanically tied to a falling N as a component ages out of its "undetermined wear state").

5.6 Marginal and joint TINF distributions

When several TINF random variables are observed jointly, a TINF joint density $f_N(x,y) = (f_T(x,y), f_I(x,y), f_N(x,y), f_F(x,y))$ is defined the same componentwise way, and marginalization proceeds by integrating out the unwanted argument separately in each of the four component densities. TINF independence of two random variables (as opposed to two events, §4.4) holds when the joint density factors componentwise, $f_T(x,y)=f_{T,X}(x)f_{T,Y}(y)$ and likewise for I, N, F — a condition that can hold for some components and fail for others, e.g. two sensors whose truth-support readings are independent while their neutrality readings are correlated because they share a common calibration routine that occasionally flags both as "legitimately borderline" at once.

6. TINF Statistics

6.1 The data model

A TINF sample is a set of observations $X_i = (T_i, I_i, N_i, F_i)$, $i = 1, \dots, n$, each component possibly an interval $[T_i^L, T_i^U]$ etc. rather than a single number, which is the natural representation when the measurements themselves are only known imprecisely.

6.2 Sample mean — worked example

$$X_N = (T, \bar{I}, N, F), \quad T = (1/n) \sum T_i, \text{ etc.}$$

Five machine reliability assessments:

Unit	T	I	N	F
1	0.80	0.10	0.05	0.15
2	0.90	0.05	0.10	0.10
3	0.70	0.15	0.10	0.25
4	0.85	0.10	0.05	0.20
5	0.95	0.05	0.05	0.10

gives $\bar{X}_N = (0.84, 0.09, 0.07, 0.16)$ — deliberately not renormalized to sum to 1 (the total here is 1.16). The statistical story is richer than "average reliability is 84%": the fleet also carries a moderate, trackable level of pure indeterminacy (missing sensor data, on average 9%), a comparatively low neutral/borderline rate (7%), and a non-negligible failure-evidence rate (16%) — three separate levers for a maintenance program to pull, not one blended "uncertainty" figure.

6.3 Sample variance

$\text{Var}_N(X) = (S_T^2, S_I^2, S_N^2, S_F^2)$ with $S_T^2 = (1/(n-1)) \sum (T_i - \bar{T})^2$ and likewise for I, N, F, kept separate rather than pooled. High variance in N specifically — as opposed to high variance in I — flags an inconsistent adjudication process (e.g. a QC line's "hold for review" bin swinging wildly between inspectors) rather than a noisy measurement, and the right fix is different in each case.

6.4 Median and quantiles

Ordering each component separately, $T_{(1)} \leq \dots \leq T_{(n)}$ and likewise for I, N, F, gives $\text{Med}_N(X) = (\text{Med}(T), \text{Med}(I), \text{Med}(N), \text{Med}(F))$; this is especially useful when the four components have visibly different distributional shapes (e.g. T roughly symmetric, F heavily right-skewed by a handful of outright failures).

6.5 Confidence intervals

A componentwise $100(1-\alpha)\%$ confidence interval reports all four bands at once, e.g.

$$([0.72, 0.84], [0.05, 0.12], [0.03, 0.10], [0.14, 0.25])$$

which is strictly more informative for downstream decisions than collapsing the last three into a single "not-true" interval.

6.6 Hypothesis testing — worked example

With $H_0: \theta_N = \theta_{N,0}$ against $H_1: \theta_N \neq \theta_{N,0}$, a TINF test statistic and p-value are likewise componentwise, $Z_N = (Z_T, Z_I, Z_N, Z_F)$ and $p_N = (p_T, p_I, p_N, p_F)$. Example: $p_N = (0.003, 0.08, 0.21, 0.015)$. At $\alpha = 0.05$, the T and F components are statistically significant while I and N are not — a genuinely multidimensional inferential conclusion ("the treatment effect and the failure-evidence shift are real; the apparent change in missingness and in neutral outcomes could be chance") that a single p-value could not express. A rigorous TINF t-test, ANOVA, χ^2 test, and nonparametric battery, built by applying the classical machinery to each of the four component samples and stating explicitly how (or whether) the four resulting inferences are to be combined, is listed as an open problem in Chapter 11.

6.7 Regression

A componentwise TINF regression fits four parallel equations,

$$Y_T = \beta_{T0} + \beta_{TI} X_T + \varepsilon_T, \quad Y_I = \beta_{I0} + \beta_{II} X_I + \varepsilon_I, \quad Y_N = \beta_{N0} + \beta_{NI} X_N + \varepsilon_N, \quad Y_F = \beta_{F0} + \beta_{FI} X_F + \varepsilon_F$$

while a cross-component model, e.g. $Y_T = \beta_0 + \beta_1 X_T + \beta_2 X_I + \beta_3 X_N + \beta_4 X_F + \varepsilon$, lets a researcher directly test questions like: does rising indeterminacy in the inputs translate into rising neutrality in the outputs, or into rising falsity? That cross-component question — does uncertainty upstream become genuine ambivalence downstream, or does it become rejection — has no analogue in classical or three-component neutrosophic regression and is one of the more distinctive research openings this framework creates.

6.8 Ranking two quadruples — worked example

Using the score function of §2.4 on two candidate decisions,

$$A = (0.75, 0.20, 0.10, 0.15) \rightarrow s(A) = 0.625 \quad B = (0.70, 0.05, 0.30, 0.20) \rightarrow s(B) = 0.6875$$

B outranks A: even though A has the higher raw truth-support, B's much lower indeterminacy and higher neutrality make it the better-supported decision overall once all four components are weighed together. This is the mechanism by which the statistics of Chapters 5–6 ultimately feed a scalar decision when one is required, without discarding the quadruple itself earlier in the pipeline.

6.9 Effect size and practical significance

Because a componentwise p-value battery (§6.6) can flag statistical significance in a component whose absolute movement is practically negligible, a TINF analogue of Cohen's d is useful: $d_N = ((\bar{\mu}_1 - \bar{\mu}_2) / S_{\text{pooled}})_{T,I,N,F}$, computed separately per component. A worked triage rule that several of the Chapter 10 applications use implicitly: treat a component as actionable only when it is both statistically significant ($p < \alpha$) and

practically significant ($|d|$ above a domain-specific threshold); this keeps a large sample size in, say, a content-moderation dataset from manufacturing a "significant" but practically meaningless drift in N .

7. Reduction Theorems

Any respectable generalization must collapse back to the theories it generalizes under the right restrictions. Several such reductions are recorded here — the first three carried over from the earlier edition, the remainder new — each a short consequence of the definitions above, included for completeness rather than as a claim of depth. Chapter 9 uses these same reductions as anchors when relating TINF to frameworks that were developed independently rather than as restrictions of it.

7.1 Classical probability and statistics

Setting $I = N = F = 0$ collapses $P_N(A)$ to $(P_T(A), 0, 0, 0)$; identifying $P(A) := P_T(A)$ recovers classical probability exactly, and $P(A) + P(A^c) = 1$ falls out of the §4.3(P3) complement rule with $I=N=0$. Likewise, observations of the form $X_i = (x_i, 0, 0, 0)$ reduce the Chapter 6 sample mean and variance to the ordinary sample mean and variance. Formally: Classical probability $\subseteq$ TINF probability, and Classical statistics $\subseteq$ TINF statistics.

7.2 Fuzzy and intuitionistic fuzzy set theory

A fuzzy membership $\mu_A(x) \in [0,1]$ is recovered by $T(x) = \mu_A(x)$, $I = N = 0$, $F(x) = 1 - \mu_A(x)$. Intuitionistic fuzzy sets, which carry an explicit non-membership degree independent of membership, are recovered by T, F free and $I = N = 0$ with $T+F \leq 1$ — i.e. intuitionistic fuzzy theory is the sub-case of TINF theory in which all indeterminacy is treated as pure (no neutrality is modeled separately).

7.3 Ordinary (T, I, F) neutrosophic theory

Ordinary three-component neutrosophic theory is recovered by aggregating I and N back together:

$$I = I + N \quad (T, I, N, F) \longrightarrow (T, I, F)$$

whenever that aggregation is semantically appropriate for the application — i.e., TINF theory is a refinement of, not a replacement for, ordinary neutrosophic theory, and any (T,I,F) result can be recovered from a (T,I,N,F) model by this one substitution. The converse is not available: a bare (T,I,F) model cannot be split back into (T,I,N,F) without additional information about how much of I was really N , which is precisely the information the refinement is designed to capture in the first place.

7.4 Dempster-Shafer belief/plausibility (new)

Setting $N = 0$ and reading (T, I, F) as (belief, don't-know mass, disbelief) under the normalized reading recovers exactly a singleton-focused Dempster-Shafer basic probability assignment, with $\text{Bel}(A) = T_A$, $\text{Pl}(A) = T_A + I_A = 1 - F_A$. This reduction, together with §9.1's fuller treatment, is what licenses calling the general non-normalized reading of §4.1 a genuine extension of evidence theory rather than merely an analogy to it.

7.5 Rough sets (new)

A Pawlak rough set's three regions — positive (definitely in), negative (definitely out), boundary (undecidable given the available attributes) — are recovered by $T(x)=1$ on the positive region, $F(x)=1$ on the negative region, and $I(x)=1$ (with $N=0$) on the boundary region, all other components 0; §9.5 develops this reduction and its converse limitations further.

7.6 Summary: the reduction lattice

The reductions above are not independent facts but nodes of a single lattice: classical probability sits at the bottom (all four "extra" degrees of freedom collapsed), fuzzy/intuitionistic fuzzy and Pawlak rough sets sit one level up (one non-trivial extra component switched on), ordinary (T,I,F) neutrosophic theory and singleton-focused Dempster-Shafer theory sit two levels up (two extra components switched on, differently named), and full TINF theory sits at the top. Chapter 9 draws the analogous picture for the frameworks — possibility theory, imprecise probability, Belnap-Dunn logic — whose relationship to TINF is a structural analogy rather than a strict sub-case, and is explicit in each instance about which relationship is which.

8. Over-, Under-, and Off-TINF Extension

The bounded $[0,1]$ theory generalizes further, following the same Over/Under/Off program already established for ordinary neutrosophic sets, logics, and probabilities.

- TINF Over-value: some component exceeds 1 (and none is below 0) — e.g. an employee whose contribution, benchmarked against a full-time baseline of 1, is genuinely above baseline.
- TINF Under-value: some component is below 0 (and none exceeds 1) — e.g. a sensor reading calibrated such that active harm/negative contribution is represented as negative T or N.
- TINF Off-value: both occur simultaneously — some component above 1 and another below 0.

Example: $(1.20, 0.10, 0.05, 0.30)$ is an Over-TINF value — an evidence score normalized to a baseline of 1, here 20% above baseline. The same Over/Under/Off construction applies unchanged to TINF Measure, TINF Probability, and TINF Statistics: an Over-TINF probability might represent an event whose supporting evidence, pooled across redundant independent sources, legitimately exceeds what a single fully-confirming source would contribute. This extension is most useful precisely in the applications of Chapter 10 where the four components are relative evidence or contribution scores rather than literal frequencies — personnel evaluation, sensor-confidence fusion, and portfolio attribution being the clearest cases.

8.1 Relation to classical over/under-confidence and calibration theory (new)

Classical forecast-calibration theory already has a scalar notion of over-confidence (reported probability systematically exceeds observed frequency) and under-confidence (the reverse), typically diagnosed with a reliability diagram. The TINF Over/Under/Off machinery generalizes this from a single scalar diagnosis to a four-component one: a forecaster, classifier, or sensor can be simultaneously over-confident in T (claims more support than warranted), under-confident in N (fails to recognize genuine borderline cases as such, collapsing them into I instead), and well-calibrated in F. A TINF reliability diagram — plotting each component's reported value against its empirically observed rate separately — is a direct, practically implementable output of this chapter and a natural bridge to the calibration literature discussed further in §9.8.

9. TINF in the Landscape of Uncertainty Theories

This chapter is the principal addition in this edition. Chapters 1–8 develop TINF measure, probability, and statistics as a self-contained system; this chapter asks how that system relates to the other major mathematical treatments of uncertainty developed independently across statistics, artificial intelligence, logic, and physics over the last century. For each framework, three things are stated explicitly: (i) what problem the framework was built to solve, (ii) the precise sense in which it is, or is not, a special case of TINF theory or vice versa, and (iii) what TINF adds that the framework alone cannot express. The chapter closes with a worked cross-notation example and a single comparison table (§9.14) that can be read on its own as an executive summary of the whole booklet's theoretical positioning.

9.1 Dempster-Shafer Evidence Theory

Dempster-Shafer theory (Shafer, 1976, building on Dempster's earlier work on upper and lower probabilities) assigns a basic probability assignment $m: 2^\Theta \rightarrow [0,1]$ over the subsets of a frame of discernment Θ , with $\sum m(A) = 1$ and $m(\emptyset)=0$, and derives belief $\text{Bel}(A) = \sum_{B \subseteq A} m(B)$ and plausibility $\text{Pl}(A) = \sum_{B \cap A \neq \emptyset} m(B)$, with $\text{Bel}(A) \leq \text{Pl}(A)$ forming an interval that brackets the "true" probability. Dempster's combination rule $\oplus$ fuses independent sources, with a normalization step that discards conflicting mass — the notorious source of counterintuitive results under high conflict, which the Dezert-Smarandache Theory (DSmT) referenced already in §10.3 of the applications chapter was built to repair by not forcing mass onto the empty set.

Connection to TINF. Restricting attention to singleton-focused evidence (m concentrated on $\{\theta\}$ and Θ itself, $\theta \in \Theta$ a binary proposition and its negation) and identifying $T_A = \text{Bel}(A)$, $F_A = 1 - \text{Pl}(A) = \text{Bel}(A^c)$, $I_A = \text{Pl}(A) - \text{Bel}(A)$ recovers exactly the general (non-normalized) TINF probability of §4.1 with $N \equiv 0$ — this is the reduction formalized in §7.4. Dempster-Shafer's belief-interval $[\text{Bel}(A), \text{Pl}(A)]$ is therefore literally the same object as a TINF quadruple with $T = \text{Bel}(A)$, $F = 1 - \text{Pl}(A)$, and $I =$ the interval width, once N is folded into I . What Dempster-Shafer cannot express, because it has no analogue of a mandated third stance, is the distinction this booklet insists on throughout: evidence that is missing (I) versus evidence that is present and genuinely establishes a third outcome (N). A frame of discernment with a designated "genuinely neither" element, rather than treating every non-singleton mass as generic ignorance, is precisely a TINF frame; several authors in the Dezert-Smarandache tradition (e.g. the PCR5 combination rule referenced in the source paper's own bibliography) have moved in this direction without adopting the T/I/N/F terminology explicitly. Combination: Dempster's rule and the §4.3(P4) TINF combination operators coincide under the product t-norm/probabilistic-sum pair when $N=0$ on both operands, so DSmT-style fusion (§10.3 of Chapter 10) is a direct, already-tested proving ground for TINF combination at the quadruple level.

9.2 Fuzzy Measures, Sugeno and Choquet Integrals

A fuzzy measure (Sugeno, 1974) relaxes classical measure's additivity to monotonicity alone: $g: 2^X \rightarrow [0,1]$ with $g(\emptyset)=0$, $g(X)=1$, and $A \subseteq B \implies g(A) \leq g(B)$, without requiring $g(A \cup B) = g(A) + g(B)$ for disjoint A, B . This is exactly the right relaxation for capturing synergy or redundancy among criteria (the classical motivating example: two mediocre exam scores that together indicate a strong candidate more than their sum would suggest). Integration against a fuzzy measure is done with the Choquet integral (which reduces to the ordinary Lebesgue integral when g is additive) or the Sugeno integral (built from min/max rather than product/sum, better suited when the underlying scale is only ordinal).

Connection to TINF. TINF's own (M4) axiom already abandons full monotonicity for I and N while retaining it for T and F (§3.2) — the mirror image of the fuzzy-measure move, which keeps monotonicity but drops additivity. The two relaxations are complementary rather than nested: a TINF measure with a fuzzy (non-additive) μ_T component is a natural hybrid object, and the Choquet integral is the correct integration tool for exactly that component, while the ordinary TINF integral of §3.4 remains appropriate for the additive components. Concretely: $\int_A f \, d\mu_T$ should be read as a Choquet integral whenever μ_T is itself only monotone (e.g. a synergy-sensitive "combined evidence weight" across several correlated diagnostic tests in §9.1's medical example), giving a genuinely new, useful hybrid — Choquet-TINF integration — that neither parent theory offers alone. This is listed explicitly as an open problem in §11.9.

9.3 Possibility Theory

Possibility theory (Zadeh, 1978; developed extensively by Dubois and Prade) replaces a single probability measure with a possibility measure Π and a dual necessity measure N_{ec} (unfortunately overloaded notation with this booklet's N — this section writes N_{ec} to avoid confusion), built from a possibility distribution $\pi: X \rightarrow [0,1]$ via $\Pi(A) = \sup_{x \in A} \pi(x)$ and $N_{ec}(A) = 1 - \Pi(A^c) = \inf_{x \notin A} (1 - \pi(x))$. Possibility measures are maxitive rather than additive ($\Pi(A \cup B) = \max(\Pi(A), \Pi(B))$), which makes them well suited to representing what is merely conceivable or unsurprising, as opposed to what is probable.

Connection to TINF. The possibility/necessity pair $(\Pi(A), N_{ec}(A))$ brackets an unknown probability exactly the way a TINF quadruple's T and $1 - F$ do; formally, $T_A = N_{ec}(A)$ and $1 - F_A = \Pi(A)$ gives a valid reduction whenever $I_A = \Pi(A) - N_{ec}(A)$ is interpreted as pure indeterminacy and $N_A = 0$. This is a structural cousin of the Dempster-Shafer reduction of §9.1 (indeed, possibility measures are the special case of Dempster-Shafer belief functions whose focal elements are nested), and the same limitation applies: possibility theory, built to model conceivability under vagueness rather than a genuine third stance, has no natural place to put a component analogous to N . Where TINF adds value in a possibility-theoretic setting is precisely in applications — political neutrality, referee impartiality, a genuinely transitional ecological state (§10.5, §10.10, §10.13) — where "merely possible but unconfirmed" (I) and "actually, definitively, a third thing" (N) must be told apart, which a possibility distribution alone cannot do.

9.4 Imprecise Probability and Credal Sets

Walley's (1991) theory of imprecise probability represents a subject's uncertainty not by a single probability measure but by a closed convex set of probability measures — a credal set M — together with lower and upper previsions $\underline{P}(A) = \inf_{P \in M} P(A)$ and $\bar{P}(A) = \sup_{P \in M} P(A)$. This is the framework of choice when the available evidence is too sparse or too conflicting to justify a single prior, and it underlies robust Bayesian analysis, imprecise Dirichlet models, and much of the modern credal-network literature in AI.

Connection to TINF. The pair $(\underline{P}(A), \bar{P}(A))$ plays exactly the role of $(T_A, 1-F_A)$ in the normalized TINF reading, with the imprecision $\bar{P}(A) - \underline{P}(A)$ again corresponding to I_A when N is absent. This makes credal sets and TINF quadruples inter-translatable for the $I=0$ -or- $N=0$ sub-cases, and the translation is genuinely useful in both directions: TINF's componentwise Bayes update of §4.4 gives a computationally cheap approximation to the (generally much harder) problem of updating an entire credal set, valid whenever the credal set's extreme points move in a way well-approximated by the four-corner TINF update; conversely, Walley's well-developed axiomatics for coherence (avoiding sure loss, natural extension) is a natural candidate rigor standard for a future TINF coherence theory, an item added to the open-problems list in §11.10. Where TINF strictly extends credal-set theory is again the N axis: a credal set records how uncertain one is about a scalar probability, but it has no native representation for the claim "the correct probability is not merely hard to pin down — the situation is honestly a mixture of two structurally different outcomes," which is what a nonzero N is built to capture.

9.5 Rough Set Theory

Pawlak's (1982) rough set theory starts from an indiscernibility relation induced by a fixed set of observable attributes and partitions any target concept A into a lower approximation (objects definitely in A given the available attributes), an upper approximation (objects possibly in A), and a boundary region (upper minus lower — objects that cannot be classified either way with the attributes at hand). The accuracy of approximation is the ratio $|\text{lower}|/|\text{upper}|$, a single scalar summarizing how "rough" the concept is relative to the attribute set.

Connection to TINF. §7.5 already gives the formal reduction: $T=1$ on the lower approximation, $F=1$ on the (lower approximation of the) complement, and $I=1$ ($N=0$) on the boundary region. This reduction is exact and is one of the cleanest in the chapter, because rough sets were built, like TINF's I -component, to represent "unresolvable with current information" rather than any variety of principled in-betweenness — a rough set's boundary region is definitionally an I -region, never an N -region, since more attributes (a finer indiscernibility relation) can always shrink it, whereas TINF's N is defined precisely as the component more information cannot shrink. This is a useful diagnostic in the other direction too: any application proposing to model its boundary/undecided cases purely as a rough set is implicitly claiming that all of its ambiguity is I -type, and the TINF framework's contribution is to ask, for each such application, whether that claim is actually true — the manufacturing hold/rework bin of

§10.7 and the legal evidence-evaluation example of §10.8 are both cases in this booklet where the answer turned out to be no.

9.6 Intuitionistic, Pythagorean, and q-Rung Orthopair Fuzzy Sets

Atanassov's (1986) intuitionistic fuzzy sets attach to each element a membership degree μ and a non-membership degree ν , independently, subject only to $\mu + \nu \leq 1$, with $1 - \mu - \nu$ read as "hesitancy." Yager's (2013) Pythagorean fuzzy sets relax the constraint to $\mu^2 + \nu^2 \leq 1$, widening the admissible region for applications (multi-criteria decision making, mainly) where experts' membership and non-membership judgments are individually strong but jointly would violate the linear intuitionistic constraint; q-rung orthopair fuzzy sets (Yager, 2017) generalize further to $\mu^q + \nu^q \leq 1$ for a tunable $q \geq 1$, trading off constraint tightness against interpretability as q grows.

Connection to TINF. All three are exactly the T,F sub-case of TINF theory with I and N both collapsed into a single "hesitancy" scalar (§7.2 already covers the linear intuitionistic case; the Pythagorean and q-rung cases are the same reduction with the linear constraint $T + F \leq 1$ replaced by $T^q + F^q \leq 1$, which does not interact with I, N at all, since those two families never split hesitancy in the first place). Practically, this means any published intuitionistic, Pythagorean, or q-rung orthopair fuzzy decision-making method can be run unmodified as a TINF method by simply setting $N=0$ throughout and reading hesitancy as I; the score/accuracy functions these families already use for ranking (Chen and Tan's score function, and its Pythagorean and q-rung descendants) are the direct one-component-collapsed ancestors of the score/accuracy/certainty cascade of §2.4, which is itself lifted with only cosmetic changes from the source paper's three-component predecessor. The distinctive TINF contribution relative to this whole family of orthopair fuzzy sets is, again, the ability to say that some of what these methods lump into "hesitancy" is not hesitancy (I) at all but a confidently-held third position (N) — a distinction with real consequences for the aggregation operators (weighted averaging, weighted geometric, Einstein and Hamacher t-norm families) that make up most of the orthopair fuzzy decision-making literature, since those operators are built to shrink hesitancy toward zero as more experts weigh in, which is the correct behavior for I but the wrong behavior for N.

9.7 Belnap-Dunn Four-Valued Logic and Paraconsistency

Belnap (1977) and Dunn independently proposed a four-valued logic for reasoning with information from multiple, possibly conflicting or incomplete, sources, with truth values True, False, Both (B — told both "yes" and "no"), and Neither (N_{bd} — told nothing). This logic underlies much of paraconsistent logic (systems that do not explode into triviality when a contradiction is derived) and has been used directly as a semantics for databases with conflicting or missing records.

Connection to TINF. This is, of the twelve frameworks surveyed in this chapter, the one whose motivating structure is closest in spirit to TINF's own — both start from the observation that collapsing "conflicting" and "missing" into one undifferentiated non-classical value loses real information — while landing on a genuinely different

mathematical object. Belnap–Dunn's four values are a discrete lattice (the well-known "diamond" with True and False as the two side points and Both, Neither as top and bottom, or vice versa depending on the ordering chosen — truth order vs. information order), whereas a TINF quadruple is a point in the continuous $[0,1]^4$ (or $[0,4]$ -bounded) box: Belnap–Dunn tells you which of four qualitative situations you are in, TINF tells you, quantitatively, how much of each. The precise correspondence is: Belnap's Both corresponds to a TINF quadruple with both T and F large (evidence strongly supports and strongly opposes at once, i.e. genuine conflict) — notably not the same thing as TINF's N, which is a real third stance rather than a conflict between the other two — while Belnap's Neither corresponds to a TINF quadruple with I large and T, F, N all small (nothing is known). This is an important clarification and a common point of confusion: N in this booklet's sense (neutrality) has no direct Belnap–Dunn analogue at all, since Belnap–Dunn's four values are built entirely from combinations of "told true" and "told false" evidence and have no primitive for "the honest answer is neither, independent of what anyone has been told." A faithful embedding of TINF's four qualitative regimes (dominant T, dominant F, dominant I, dominant N) into a discrete logic would need five or six Belnap-style values, not four — itself a small, self-contained open research question noted in §11.11.

9.8 Bayesian and Frequentist Statistics

Classical statistics splits, famously, into a frequentist school (probability is long-run relative frequency; parameters are fixed and unknown, estimated via point estimators, confidence intervals, and p-values) and a Bayesian school (probability is degree of belief; parameters are themselves random variables with a prior distribution updated to a posterior via Bayes' theorem as data arrives). Both schools, and the large modern literature reconciling or hybridizing them (empirical Bayes, objective Bayes, false-discovery-rate control), operate on a single scalar or vector-valued probability at a time.

Connection to TINF. TINF statistics (Chapter 6) is best understood not as a competitor to either school but as a scaffold that can be filled in with either: the componentwise sample mean, variance, confidence interval, and hypothesis test of §6.2–6.6 are frequentist constructions run four times in parallel, one per TINF component, while a fully Bayesian TINF statistics — putting a prior not just on a parameter but on an entire quadruple, and updating via the componentwise Bayes rule of §4.4 or (more ambitiously) a joint four-component likelihood — is a natural, currently underdeveloped, extension flagged in §11.2 and §11.10. A useful frame: classical frequentist inference answers "how confident should I be that T is large," while TINF inference additionally, and separately, answers "how much of my remaining uncertainty is resolvable-by-more-data (I) versus structurally irreducible (N)" — a question that neither classical school poses natively, because neither partitions its residual uncertainty into causally distinct types. Bayesian model-averaging, in particular, is worth flagging as close kin to N: a posterior that is genuinely bimodal, with two well-separated, well-supported hypotheses and little mass between them, is arguably better summarized by a nonzero N (the honest answer is one of two things, not "somewhat uncertain between them") than by reporting

only the posterior mean, and this booklet's manufacturing and legal examples (§10.7, §10.8) are, read this way, disguised instances of exactly this bimodality.

9.9 Robust and Interval Statistics

Robust statistics (Huber, Tukey, and others, from the 1960s onward) develops estimators — trimmed means, M-estimators, the median as opposed to the mean — that are deliberately insensitive to outliers or to specific distributional assumptions failing. Interval statistics (a smaller, more recent literature, overlapping with imprecise probability, §9.4) treats each observation as an interval rather than a point when the measuring instrument itself only reports bounds, and develops interval-valued analogues of the mean, variance, and standard hypothesis tests.

Connection to TINF. TINF's own data model (§6.1) already allows each component to be an interval rather than a point, making interval statistics a direct sub-tool of TINF statistics whenever T, I, N, F are themselves given as ranges — the componentwise sample mean and variance of §6.2–6.3 reduce exactly to standard interval-arithmetic mean and variance formulas under that reading, with no new mathematics required. Robust statistics contributes a different, complementary lesson: because a TINF sample separates the four components before any averaging happens, a single wild observation (a mis-keyed -1 in a truth-support column, say) contaminates only its own component's mean and variance rather than a pooled scalar, which is itself a mild, structural form of robustness the TINF representation confers for free — the four-way separation acts like a coarse, domain-driven form of the winsorizing or trimming that robust statistics does algorithmically. A fully worked-out TINF-robust statistics, applying M-estimators or trimmed means within each component rather than the plain sample mean of §6.2, is a small, tractable open item added to §11.3.

9.10 Random Set Theory and p-Boxes

A random set (Matheron, 1975; Nguyen, 2006) is a random variable whose realizations are sets rather than points, with an induced upper probability (its Choquet capacity coincides with plausibility, §9.1, when the random set is finite-valued) and lower probability. A probability box, or p-box, is a pair of cumulative distribution functions $\underline{F}(x) \leq \bar{F}(x)$ bracketing an unknown CDF, widely used in engineering risk analysis when the true distribution of a quantity is not fully known but bounds on it are defensible.

Connection to TINF. Random-set theory is, formally, the general mathematical umbrella under which Dempster-Shafer theory (§9.1), possibility theory (§9.3), and imprecise probability's credal sets (§9.4) all sit as special cases — belief functions are exactly the containment functionals of random closed sets — so the reductions already given in §9.1, §9.3, and §9.4 all lift to this level with no new work, and TINF's own relationship to random-set theory is inherited rather than independent. A p-box's $[\underline{F}, \bar{F}]$ pair, applied to a TINF random variable's T-component distribution function F_T (§5.2), gives a natural hybrid — a TINF variable whose own T-marginal is itself only known up to a p-box — which is precisely the right representation for the reliability-engineering applications of §10.2 when even the truth-support component's distribution is not fully pinned down by

the available test data. This nested-uncertainty construction (imprecision about a component of an already-multicomponent quantity) is flagged as an open modeling pattern worth formalizing in §11.12.

9.11 Quantum Probability

Quantum probability replaces classical event algebras (Boolean σ -algebras) with the lattice of projection operators on a Hilbert space, and replaces a classical probability measure with the Born rule applied to a quantum state; famously, quantum events need not be jointly measurable (non-commuting observables), and quantum probability violates the classical law of total probability in ways that have been used, in the last two decades, to model genuine human decision-making anomalies (order effects in survey responses, the conjunction fallacy, interference effects in preference judgments) that classical and even Dempster-Shafer-style models struggle with.

Connection to TINF. The connection here is the loosest of the twelve in this chapter and is offered as a structural analogy rather than a reduction: TINF's departure from classical measure theory is the abandonment of monotonicity for I and N (§3.2, M4), while quantum probability's departure is the abandonment of a single joint sample space for non-commuting observables — different axioms are relaxed, for different reasons. What the two share is a diagnosis, arrived at independently: that forcing every observed anomaly through a single-valued, fully monotone, fully commutative probability model produces systematic, predictable distortions, and that the fix is to add structure (a Hilbert space; a fourth component) rather than to explain the anomaly away as measurement error. Whether a genuinely quantum-probabilistic TINF — replacing the $[0,1]^4$ box with a four-dimensional analogue of a density-operator formalism, so that T, I, N, F themselves become non-commuting in some applications (political-opinion polling under order effects, §10.5, is the most plausible candidate application) — is a worthwhile research direction or an over-elaboration is left as an explicitly open, speculative question in §11.13, flagged as such rather than pursued here.

9.12 Plithogenic Statistics

Plithogenic sets and statistics (Smarandache, 2018 onward, and referenced in [12] of this booklet's bibliography) generalize neutrosophic theory in a different direction from TINF: rather than splitting a single indeterminacy component into two, plithogenic theory attaches to each element not one but several attribute-values simultaneously, each with its own degree of appurtenance and a contradiction (dissimilarity) degree relative to a designated dominant attribute-value, and defines aggregation operators that are explicit functions of these contradiction degrees.

Connection to TINF. The two extensions are orthogonal rather than nested: TINF refines the number of truth-status dimensions from three to four for a single attribute, while plithogenic theory keeps a fixed small number of truth-status dimensions but allows an arbitrary number of attributes, each contributing its own quadruple (or triple) weighted by contradiction degree. This means the natural common generalization — plithogenic TINF statistics, in which each of several attributes carries a full (T,I,N,F)

quadruple rather than a single membership degree, aggregated via contradiction-degree-weighted operators — is a straightforward, well-motivated combination of the two frameworks rather than a competition between them, and is listed as a concrete, tractable open problem in §11.4, closest in spirit to the source paper's own reference [12] linking plithogenic probability to multivariate probability and statistics.

9.13 A Running Example Worked in Every Framework

To make §9.1–§9.12 concrete all at once, this section carries a single scenario — the weather-forecasting example first introduced in §4.5 — through each framework's own native notation, so the family resemblance and the genuine differences are visible side by side rather than only argued for in prose. The scenario: three forecasting models disagree about tomorrow's rain, and some station data is missing.

Framework	How this framework would report the same evidence	What it cannot separately express
TINF (this booklet)	$P_N(\text{rain}) = (0.65, 0.15, 0.10, 0.30)$	— (this is the reference representation)
Classical probability	$P(\text{rain}) = 0.65$ (or $0.65/0.35$ with the complement implied)	Any distinction among the remaining 0.35: missing data vs. genuine model split vs. active counter-evidence
Dempster-Shafer	$\text{Bel}(\text{rain})=0.65$, $\text{Pl}(\text{rain})=0.90$ (interval width 0.25 absorbs I and N together)	Whether the 0.25 gap is missing station data or a genuine three-model split
Possibility theory	$\Pi(\text{rain})=0.90$, $N_{ec}(\text{rain})=0.65$ (same numeric bracket as Dempster-Shafer here)	Same limitation as Dempster-Shafer, plus no native probability semantics
Imprecise probability	$[\underline{P}, \bar{P}] = [0.65, 0.90]$, credal set of compatible models	Same interval; no attribution of the imprecision's source
Intuitionistic fuzzy	$\mu(\text{rain})=0.65$, $\nu(\text{rain})=0.30$, hesitancy $\pi=0.05$ (forced to fit $T+F \leq 1$, understating the true 0.25 gap)	Cannot even represent the full 0.25 gap without violating its own constraint; understates uncertainty
Belnap-Dunn logic	Closest qualitative value: neither pure True nor pure False; would likely be coded "Neither" (told little)	No quantitative gradation at all; cannot distinguish this case from a much less confident forecast
Rough sets	Boundary-region membership if station data is one of the classifying attributes; else lower-approximation "rain" given the other four attributes	No native probability number; a purely qualitative in/out/boundary call

The pattern in the right-hand column is the chapter's thesis in miniature: every other framework surveyed here collapses the 0.25 of non-committed evidence into one undifferentiated quantity, while TINF's own report — (0.65, 0.15, 0.10, 0.30) — keeps the 0.15 of missing station data and the 0.10 of genuine three-way model disagreement apart, which is operationally the difference between "wait for the 6pm data update" and "stop waiting; the models themselves are split and more data from the same sources will not change that."

9.14 Comparative Summary Table

The table below compresses §9.1–§9.13 into a single reference. "Relation to TINF" states whether the framework is a strict special case of TINF (obtained by fixing certain components to zero or imposing a constraint), a structural analogue (independently motivated, mathematically parallel, but not a sub-case), or a genuinely orthogonal extension (addresses a different axis of generalization entirely).

Framework	Core object	Relation to TINF	What TINF adds
Dempster-Shafer	Belief / plausibility interval	Strict special case (N=0)	Splits ignorance into I and a real third stance N
Fuzzy measure / Choquet	Non-additive monotone measure	Orthogonal (relaxes additivity, not partition size)	A fourth component; hybrid Choquet-TINF integral (open)
Possibility theory	Possibility / necessity pair	Structural analogue	Distinguishes conceivable-but-unconfirmed from genuinely neutral
Imprecise probability	Credal set, lower/upper prevision	Strict special case (N=0)	Cheap 4-corner update; a coherence theory is open work
Rough sets	Lower/upper/boundary regions	Strict special case (I=1 on boundary, N=0)	Asks whether boundary cases are really all I-type
Intuitionistic / Pythagorean / q-rung fuzzy	Membership, non-membership, hesitancy	Strict special case (N=0, hesitancy=I)	Splits hesitancy into resolvable and stable parts
Belnap-Dunn logic	Four discrete truth values	Structural analogue (discrete vs. continuous)	Quantitative degree; N has no Belnap-Dunn analogue at all
Bayesian / frequentist statistics	Point/interval estimates, posteriors	Orthogonal scaffold (fills either school in)	Separates resolvable from irreducible residual uncertainty
Robust / interval statistics	Outlier-resistant, interval-valued estimators	Sub-tool / partial special case	Four-way separation gives free partial robustness

Random sets / p-boxes	Set-valued r.v.s, CDF bounds	Umbrella containing DS, possibility, credal sets	Enables nested imprecision on a single TINF component
Quantum probability	Non-commutative event algebra	Loose structural analogy only	Neither reduces to the other; open speculative question
Plithogenic statistics	Multi-attribute contradiction-weighted sets	Orthogonal extension (attributes vs. components)	Combination gives plithogenic-TINF statistics (open)

Read across the third column, a pattern emerges that is itself a finding of this chapter: of the twelve frameworks surveyed, six (Dempster-Shafer, imprecise probability, rough sets, and the three orthopair fuzzy-set families) are strict special cases of TINF obtainable by setting $N=0$ and, in most cases, I equal to whatever that framework calls its residual uncertainty — which is a nontrivial, checkable claim about TINF's generality, not a rhetorical one, since each reduction was derived explicitly in §9.1–§9.6 and §7.4–§7.5 rather than asserted. Three (possibility theory, Belnap-Dunn logic, Bayesian/frequentist statistics) are close structural cousins that do not reduce cleanly because they were built to answer a different question. Three (fuzzy measure theory, random-set/p-box theory, plithogenic statistics) generalize along an axis TINF does not touch (non-additivity, nested imprecision, multi-attribute structure) and are natural candidates for hybridization rather than reduction. Only quantum probability resists useful comparison altogether. This is offered as evidence, not proof, that the T/I/N/F quadruple occupies a genuine, previously under-articulated point in the space of uncertainty representations — one that most existing frameworks approximate by collapsing exactly one of its four axes.

9.15 Practical Guidance: Choosing a Framework

A fair question for a practitioner reading this chapter is not "which theory is most general" but "which theory should I actually reach for on Monday morning." The following decision guidance is offered in that spirit, organized by the shape of the problem rather than by theoretical elegance.

If your situation is...	Reach for...	Move to full TINF when...
A single scalar probability suffices and data is plentiful	Classical frequentist or Bayesian statistics (§9.8)	Stakeholders start asking "uncertain in what sense?" and the answer differs by case
Evidence comes from multiple conflicting sensors/sources	Dempster-Shafer combination (§9.1)	You can identify a recurring, nameable third outcome, not just conflict/ignorance
Attributes are incomplete and you need	Rough sets (§9.5)	The boundary region contains a stable, non-

in/out/boundary calls		shrinking sub-population
Experts give membership and non-membership ratings directly	Intuitionistic / Pythagorean / q-rung fuzzy sets (§9.6)	Hesitancy itself needs to be split into resolvable and stable parts
Bounds on a probability are defensible but a point value is not	Imprecise probability / credal sets (§9.4)	You need to track why the interval is wide, not just how wide
You need a fast, cheap qualitative triage over many cases	Belnap–Dunn four-valued logic (§9.7)	You need to rank or aggregate results numerically, not just triage
None of the above quite fits, and the "why" of the uncertainty matters operationally	Full TINF (T,I,N,F) — this booklet	—

This table is deliberately conservative: it does not claim TINF is always the right tool. Most of the frameworks in this chapter are simpler, better-tooled (mature software libraries exist for Dempster–Shafer and Bayesian methods; almost none yet exist for TINF, itself an item on the Chapter 11 research agenda), and entirely adequate when the I/N distinction genuinely does not change what happens next. The applications of Chapter 10 were selected specifically because, in each one, that distinction does change the next action — that is the bar this booklet sets for recommending the fuller, more complex machinery.

10. Worked Applications Across Fields

Each entry below states a concrete illustrative quadruple and, in one or two sentences, why separating I from N specifically changes what a practitioner should do next — the running test this booklet applies to every proposed use case. The first thirteen entries (§10.1–§10.13) carry over from the earlier edition; §10.14–§10.24 and §10.26–§10.28 are new to this edition and were chosen to span fields not previously represented — actuarial science, robotics, logistics, education, criminal justice, astronomy, agriculture, energy, misinformation research, decentralized systems, human-resources analytics, telecommunications, pharmacology, and space situational awareness — so that the applications chapter, taken as a whole, now covers the physical sciences, life sciences, social sciences, engineering, and public-sector decision-making roughly evenly.

10.1 Medical diagnosis and clinical trials

$$D = (0.82, 0.08, 0.05, 0.20)$$

T = evidence for the diagnosis (or, in a trial, proportion showing significant recovery); I = missing or unavailable diagnostic information (or inconclusive reporting); N = evidence that is genuinely, stably equivocal (a borderline biopsy read; a placebo-level, no-measurable-change outcome); F = evidence against the diagnosis (or adverse-effect / treatment failure). A high-I case calls for more testing; a high-N case calls for a different clinical pathway (watchful waiting, re-biopsy protocol) rather than repeating the same test.

10.2 Engineering reliability

$$R_N(t) = (T(t), I(t), N(t), F(t)) \text{ — e.g. } R_N = (0.91, 0.04, 0.12, 0.18)$$

T = reliability evidence; I = unknown/unmeasured factors; N = evidence that neither supports nor undermines reliability (e.g. a component operating in an untested-but-not-abnormal regime); F = evidence of possible failure. Instead of one reliability curve, the model produces four curves — useful across aircraft components, bridges, transformers, batteries, and industrial robots, where "we don't know" and "it's operating in a genuinely indeterminate regime" warrant different inspection intervals.

10.3 Multi-sensor fusion and defense

$$O = (T_o, I_o, N_o, F_o) \text{ fused from } S_1, S_2, S_3 \text{ (see §4.7)}$$

T = threat-confirming signal; I = noise, jamming, degraded visibility; N = a signal legitimately consistent with a non-combatant/commercial object; F = false-alarm evidence (e.g. a bird flock). Dezert-Smarandache Theory (DSmT) combination rules, already developed for standard neutrosophic evidence fusion, extend naturally to the quadruple case (§9.1) and are a natural engineering target for implementing the §4.3(P4) operators at scale.

10.4 Financial risk and credit assessment

$$C = (0.65, 0.20, 0.10, 0.35)$$

T = evidence supporting default; I = missing financial information (thin credit history); N = a genuinely neutral outcome such as an agreed restructuring — neither full repayment nor default; F = evidence supporting repayment. A lender forced into three components must fold restructuring into either T or F, systematically mispricing it; N prices it on its own terms.

10.5 Environmental monitoring and climate science

$$P = (0.75, 0.15, 0.25, 0.30)$$

Chemical measurements, satellite imagery, and historical baselines can conflict; I captures sensor failure and missing readings, while N captures a genuinely mixed or transitional environmental state (e.g. a river at the edge of a pollution threshold under natural seasonal variation) that repeated measurement will not resolve into a clean verdict.

10.6 Artificial intelligence: classification and content moderation

$$C_i(x) = (T_i(x), I_i(x), N_i(x), F_i(x))$$

A classifier that reports only confidence and "uncertain" conflates two very different failure modes: genuinely ambiguous or incomplete input (I) versus content that is correctly, stably borderline by its nature — satire, reclaimed language, clinical or educational context (N). Systems that only have I tend to either over-flag legitimate borderline content or bucket it as awaiting more review, which is the wrong response when more review will not change the verdict.

10.7 Manufacturing quality control

$$Q_N = (0.92, 0.03, 0.07, 0.18)$$

T = pass evidence; I = inadequately inspected units (measurement-system uncertainty); N = units correctly routed to a legitimate rework/hold disposition, distinct from a pass or a fail; F = defect evidence. A line with no N category is probably misclassifying legitimately-hold units as outright passes or failures.

10.8 Legal evidence evaluation

$$E = (0.40, 0.45, 0.05, 0.30) \text{ vs. } E' = (0.40, 0.05, 0.45, 0.30)$$

Both quadruples carry identical T and F, yet describe very different bodies of evidence: E has substantial missing evidence (I), while E' has substantial evidence that is present and genuinely neutral (N) — e.g. eyewitness testimony that is available but points in no particular direction. A legal decision system that conflates "unknown" with "neutral" is not a minor imprecision; it changes what further investigation could plausibly achieve.

10.9 Epidemiology

$$P_N(t) = (T(t), I(t), N(t), F(t))$$

Tracking an infectious-disease event over time: T from positive clinical evidence, I from missing or delayed test reporting, N from genuinely mixed epidemiological signal (e.g. a plateau consistent with either early decline or a lull before resurgence), F from evidence against ongoing transmission. Modeling I(t) and N(t) as separate time series lets a surveillance system distinguish "our data pipeline is degraded" from "the epidemic itself is in a genuinely ambiguous phase."

10.10 Political polling and social sentiment

$$P_N = (0.42, 0.15, 0.18, 0.25)$$

T = support ("yes"); I = undecided through lack of information; N = a deliberate, structurally-neutral abstention or independent stance ("neither for nor against," held by identity rather than confusion); F = opposition. Folding I and N into one "undecided" bucket, as ordinary polling analysis does, discards the fact that these two groups respond to completely different campaign strategies — more information persuades the first group and is close to irrelevant for the second.

10.11 Evolutionary biology and species dynamics

T = evolution, I = drift, N = neutral variants, F = involution

Tracking genetic and trait adaptation in a population under environmental stress: T = beneficial adaptations improving survival; I = silent mutations and unpredictable genetic drift whose fitness effect is presently unknown; N = neutral genetic variants that carry no immediate fitness advantage or penalty by their nature (the population-genetics notion of a truly neutral allele, not merely an unstudied one); F = deleterious mutations driving fitness loss. This is arguably the field where the I/N distinction has the longest independent scientific pedigree — population genetics has separately theorized "neutral evolution" for decades, which is direct external support for treating N as a distinct, non-reducible quantity rather than a notational convenience invented for this framework.

10.12 Cybersecurity: intrusion detection

T = attack evidence, I = unclassified traffic, N = benign-but-unusual traffic, F = evidence against attack

An intrusion-detection system that only reports "attack probability" and "unclear" cannot distinguish traffic that needs a signature update (I, resolvable with better rules) from traffic that is unusual but legitimately benign by nature (N, e.g. a scheduled backup job that looks like exfiltration) — conflating the two either trains analysts to ignore real alerts or floods them with false ones.

10.13 Sports officiating and international mediation

N = referee/mediator neutrality, independent of T, F

The source paper's own seed example, carried through fully: in a two-sided contest (a match, a conflict between two states), N measures the actual impartiality of the third party — the referee team, or a non-aligned country — not how much information is available about the contest's outcome. A referee's bad call, or a neutral state's political tilt, moves N directly without touching T, F, or the informational uncertainty I about who will ultimately prevail. No other component of the quadruple can carry this meaning, which is the clearest existence proof in this booklet that N is not merely "indeterminacy with a different name."

10.14 Insurance and actuarial risk (new)

$$A = (0.55, 0.20, 0.15, 0.10)$$

For a single policyholder or a risk pool: T = evidence supporting an above-average claim likelihood; I = incomplete underwriting information (a new applicant with thin history); N = a genuinely borderline risk profile — a rated-up-but-not-declined case sitting deliberately between standard and substandard tiers, where the underwriter's own guidelines define a real middle tier rather than an approximation of one; F = evidence supporting a below-average claim likelihood. Actuarial tables built on three components force the middle underwriting tier into either "standard, with a loading" or "substandard," both of which mis-price the tier over time; carrying N through the pricing model lets the loss-ratio experience for the genuinely-middle tier be tracked, and re-priced, on its own terms, rather than contaminating the calibration of the tiers on either side.

10.15 Robotics and human-robot interaction (new)

$$H = (0.60, 0.10, 0.25, 0.05)$$

For a robot inferring a human collaborator's intent (e.g. "is the person about to hand me the tool or set it down?"): T = evidence for intent A; I = ambiguous sensing (occluded hand, noisy pose estimate); N = a posture that is a genuinely valid precursor to either action — a real, stable, dual-intent pose held briefly by most humans regardless of which action follows, not a sensing failure; F = evidence for intent B (the negation). A robot that treats high-N poses as "uncertain, wait for more data" freezes needlessly at a pose that will not resolve with more frames; a robot that correctly reads it as N can instead default to a safe, intent-agnostic holding behavior designed specifically for that stable third state, improving both fluency and safety in shared workspaces.

10.16 Supply chain and logistics risk (new)

$$S = (0.70, 0.15, 0.10, 0.05)$$

For a shipment's on-time-delivery status mid-transit: T = evidence for on-time arrival; I = missing tracking data (a scan gap at a transfer hub); N = a shipment genuinely rerouted onto an alternate, equally reliable path that is neither the original "on

schedule" nor a "delay" event — a real third disposition a dispatcher has deliberately chosen, not an unknown; F = evidence for a delay. Distinguishing I from N lets a logistics dashboard route two very different interventions: chase the missing scan for I-heavy shipments, but leave N-heavy (deliberately rerouted) shipments alone, since intervening on them wastes dispatcher time trying to "resolve" a disposition that was never broken.

10.17 Education and student assessment (new)

$$G = (0.72, 0.10, 0.13, 0.05)$$

For a competency being assessed on a single student: T = evidence of mastery; I = insufficient assessment data (too few graded items touching this competency); N = a genuinely mixed performance pattern — correct on conceptual items, incorrect on procedural items testing the same competency, a real and stable split rather than noise; F = evidence of non-mastery. A learning-analytics system that reports only "mastery probability" recommends more practice items indiscriminately; one that separates I from N can instead recommend more items only for the I-heavy competencies, while routing N-heavy competencies to a targeted conceptual-versus-procedural remediation track that more of the same item type will not fix.

10.18 Criminal-justice risk assessment (new)

$$R = (0.35, 0.30, 0.15, 0.20)$$

For a pretrial or recidivism risk instrument, applied with the considerable caution such instruments deserve: T = evidence supporting elevated risk; I = missing case-history information (records not yet transferred from another jurisdiction); N = a genuinely mixed risk profile in which established risk factors point in materially different directions with no dominant signal, as opposed to simply lacking data; F = evidence supporting low risk. Because these instruments carry serious fairness and due-process stakes, the I/N distinction is not a modeling nicety here: a high-I score is properly an argument for gathering more information before any decision is made, while a high-N score is an argument that the instrument itself has reached the limit of what a risk score can responsibly say about this individual, and that human judgment, not further scoring, should govern — collapsing the two into one "moderate risk" bucket obscures exactly which of those two very different responses is warranted.

10.19 Astronomy and anomaly detection (new)

$$X = (0.30, 0.20, 0.40, 0.10)$$

For a candidate astronomical transient (e.g. a possible exoplanet transit signal or a sky-survey anomaly): T = evidence for a genuine astrophysical event; I = instrumental noise or an incomplete observation window; N = a signal that is well-measured and genuinely consistent with a known, benign, non-anomalous class of variability (a known variable star's ordinary cycle, precisely identified as such) rather than merely unclassified; F = evidence the signal is a data artifact. Large-scale survey pipelines that only flag "anomaly score" and "needs follow-up" spend expensive telescope time chasing high-I

signals that could instead be triaged by cheaper re-observation, while correctly-classified high-N signals — genuinely ordinary once identified — can be safely deprioritized rather than re-queued indefinitely as "still uncertain."

10.20 Precision agriculture (new)

$$C = (0.65, 0.10, 0.20, 0.05)$$

For a field zone's irrigation-need status from combined soil-moisture sensors and satellite vegetation indices: T = evidence the zone needs irrigation; I = sensor gaps or cloud-obscured imagery; N = a zone in a genuinely stable, self-regulating moisture state — neither needing water nor at risk of over-saturation, a real agronomic third state tied to soil composition rather than a measurement gap; F = evidence the zone does not need irrigation. A farm-management system that cannot separate I from N either over-irrigates N-zones out of caution or wastes sensor-repair budget chasing zones that were never actually under-measured, and the resulting water and cost savings from getting this right compound across a full growing season.

10.21 Energy-grid stability management (new)

$$G = (0.55, 0.15, 0.25, 0.05)$$

For a distribution-grid segment's stability status under high renewable penetration: T = evidence of an approaching instability event (voltage or frequency deviation); I = missing telemetry from a subset of smart meters or sensors; N = a segment operating in a genuinely bistable regime that intermittent renewable generation puts it into by design — a real, recurring, non-anomalous operating state rather than an unmeasured one; F = evidence the segment is stable. Grid operators who cannot distinguish I from N either over-dispatch costly reserve capacity toward N-segments that do not need it, or under-invest in the sensor coverage that would actually resolve I-segments — the two failure modes point to opposite remedies (more sensors vs. dedicated bistable-regime control logic), and a single "instability risk" score cannot tell an operator which one applies.

10.22 Misinformation and fact-checking (new)

$$M = (0.30, 0.35, 0.25, 0.10)$$

For a claim submitted to a fact-checking pipeline: T = evidence the claim is accurate; I = insufficient evidence located yet to check the claim either way; N = a claim that is genuinely, irreducibly a matter of contested interpretation or definition (a value-laden or definitionally ambiguous claim that further sourcing will not settle, as opposed to a factual claim that merely lacks sourcing so far); F = evidence the claim is false. Fact-checking organizations that only have "true / false / unverified" routinely misclassify N-type claims as "unverified," which both invites endless, unproductive re-litigation by partisans on either side and misleads readers into expecting a future update that will never meaningfully arrive, whereas correctly labeling a claim N signals that the disagreement is real and irreducible, not a temporary gap in reporting.

10.23 Blockchain and decentralized trust scoring (new)

$$W = (0.60, 0.25, 0.10, 0.05)$$

For a wallet or counterparty's trust score in a decentralized-finance risk system: T = on-chain evidence of legitimate transaction history; I = insufficient transaction history to score confidently (a genuinely new address); N = a wallet exhibiting patterns that are legitimately, verifiably consistent with a known-benign automated class — a market-making or mixing-service bot correctly identified as such, neither a scored-good nor scored-bad counterparty in the ordinary sense but a real third category with its own risk profile; F = on-chain evidence of malicious or fraudulent activity. Protocols that only expose a single trust score either under-price the genuine information gap of new addresses (I) the same way they price known bot patterns (N), or vice versa, both of which misallocate the risk premium that counterparties are actually being charged.

10.24 Human-resources and personnel evaluation (new)

$$P = (0.68, 0.12, 0.15, 0.05)$$

For an employee performance-review cycle, extending the Chapter 8 Over-TINF personnel example: T = evidence of above-expectation performance; I = insufficient observation (a new hire, or a reviewer with limited direct exposure to the employee's work); N = a genuinely mixed performance record — clearly excelling on one core responsibility while clearly underperforming on another, a real and stable pattern rather than an averaged-out "meets expectations"; F = evidence of below-expectation performance. An HR system that forces this into a single scalar rating either drags a genuinely N-profile employee's score toward a misleading "average," inviting a generic improvement plan that addresses neither strength nor weakness precisely, or an I-heavy new-hire review gets treated with the same weight as a fully-observed N-heavy review, when the correct next step — more observation versus a role redesign that plays to the split — is completely different in the two cases.

10.26 Telecommunications network quality-of-service (new)

$$Q = (0.72, 0.12, 0.11, 0.05)$$

For a network link's service-quality status monitored for congestion management: T = evidence of healthy throughput; I = missing telemetry from a subset of routers during a polling gap; N = a link operating in a genuinely acceptable, deliberately-throttled quality-of-service tier under a traffic-shaping policy — neither degraded nor at full capacity by design, a real administrative third state; F = evidence of degraded service. A network-operations dashboard that cannot separate I from N either pages an on-call engineer for a link that is behaving exactly as configured (N), or fails to escalate a genuinely under-monitored link (I) because its score looks superficially similar to the well-understood throttled state.

10.27 Pharmaceutical drug-drug interaction screening (new)

$$X = (0.25, 0.30, 0.35, 0.10)$$

For a candidate pairwise drug interaction flagged by a prescribing-support system: T = evidence of a clinically significant interaction; I = insufficient pharmacokinetic data on the specific combination; N = an interaction that is well-studied and genuinely, stably classified as clinically insignificant at ordinary doses — a real, settled third category distinct from "not yet studied" — F = evidence against any interaction. Conflating I and N here has direct patient-safety consequences: an I-heavy pair should prompt conservative co-prescribing caution pending better data, while an N-heavy pair has already been through that process and additional caution imposes needless prescribing friction without any corresponding safety benefit.

10.28 Space-debris and collision-risk tracking (new)

$$C = (0.20, 0.35, 0.30, 0.15)$$

For a conjunction (close-approach) event between a satellite and a tracked object: T = evidence of a genuine collision risk requiring a maneuver; I = insufficient tracking data (a sparsely observed object with a stale orbital solution); N = an object whose orbit is well-characterized and genuinely, predictably passes through a shared volume without risk under standard uncertainty ellipsoids — a real, recurring, non-anomalous geometric coincidence rather than an unresolved unknown; F = evidence against any meaningful risk. Space-situational-awareness centers that cannot separate I from N either burn scarce maneuver propellant reacting to well-understood recurring flybys (N), or fail to prioritize additional tracking resources toward the genuinely under-observed objects (I) that most need it — this booklet's clearest illustration outside Chapter 8's sensor examples that the I/N distinction has direct, quantifiable operational cost when ignored.

10.25 Cross-field summary table

Field	What N specifically captures
Medicine	A stably equivocal clinical finding, not a missing test
Engineering	An untested-but-not-abnormal operating regime
Defense / sensors	A signal consistent with a non-combatant object
Finance	A restructuring — neither default nor full repayment
Environment	A genuinely transitional or mixed ecological state
AI / moderation	Content that is correctly, stably borderline
Manufacturing	A legitimate hold/rework disposition
Law	Present evidence that points in no particular direction
Epidemiology	A genuinely ambiguous phase of an

	outbreak
Polling	A deliberate, identity-based neutral stance
Biology	A population-genetics neutral allele
Cybersecurity	Unusual traffic that is benign by nature
Sport / diplomacy	The literal, actively-maintained third party
Insurance	A genuine middle underwriting tier, not a loaded standard rate
Robotics / HRI	A stable, valid dual-intent human posture
Logistics	A deliberately rerouted, equally reliable shipment path
Education	A real conceptual/procedural split in mastery
Criminal justice	The limit of what a score can say — defer to human judgment
Astronomy	A precisely identified, known-benign variability class
Agriculture	A self-regulating soil-moisture state
Energy	A designed-in bistable renewable-driven operating regime
Misinformation	A contested, definitionally ambiguous claim
Blockchain	A verified benign-automated wallet class
HR / personnel	A genuine strength/weakness split, not an average

11. Research Program and Open Problems

None of the chapters above claim to be a finished theory; each opens problems more interesting than the ones it closes. Collected here as a research agenda, extended in this edition with items surfaced directly by the Chapter 9 comparison work (marked new):

1. Named TINF probability distributions (normal, binomial, Poisson, exponential, Weibull, gamma) with an explicit, application-motivated dependency structure linking the four component distributions, rather than four independent classical distributions bundled together.
2. A fully general TINF Bayes theorem permitting interaction among T, I, N, F, beyond the four-parallel-classical-Bayes-updates sketch of §4.4, and — new — a fully Bayesian TINF statistics that places a joint prior over an entire quadruple rather than four independent component priors (§9.8).
3. A rigorous TINF hypothesis-testing suite: z-test, t-test, ANOVA, χ^2 test, nonparametric tests, and multiple-comparison procedures, with an explicit protocol for combining the four resulting per-component conclusions into a single decision when one is required; new — a TINF-robust version of the same suite using M-estimators or trimmed means per component (§9.9).
4. TINF stochastic processes: $X(t) = (T(t), I(t), N(t), F(t))$ with TINF versions of Markov chains, Poisson processes, Brownian motion, reliability processes, and queueing models (previewed in §4.8).
5. TINF machine learning: classification, clustering, neural architectures, decision trees, reinforcement learning, and anomaly detection that natively output and update quadruples rather than post-hoc splitting a scalar confidence.
6. A rigorous TINF information-fusion theory — combination operators $A_1 \oplus A_2 \oplus \dots \oplus A_k$ that behave differently for independent, redundant, correlated, and contradictory sources, extending Dezert-Smarandache combination rules to the quadruple case (§10.3, §9.1).
7. Convergence theorems (monotone and dominated convergence, laws of large numbers, central limit theorems) proved once, in general form, for componentwise TINF constructions, rather than re-derived ad hoc per component as this booklet does.
8. Empirical validation on real, not illustrative, datasets — every numerical example in this booklet is constructed to demonstrate the mechanics, not fitted to data, and the framework's practical value ultimately rests on the latter.
9. New — Choquet-TINF integration: a hybrid integral for TINF measures whose T-component (or another component) is only monotone rather than additive, following the fuzzy-measure connection of §9.2.
10. New — a TINF coherence theory in the style of Walley's avoiding-sure-loss and natural-extension axioms for imprecise probability, giving TINF quadruples a rationality standard analogous to de Finetti coherence for classical subjective probability (§9.4).

11. New — a discrete, Belnap–Dunn-style logic with enough truth values (five or six, rather than four) to embed TINF's four qualitative regimes faithfully, including a primitive for neutrality distinct from conflict (§9.7).
12. New — plithogenic-TINF statistics: attaching a full (T,I,N,F) quadruple, rather than a single membership degree, to each attribute-value in a plithogenic multi-attribute model, aggregated via contradiction-degree-weighted operators (§9.12).
13. New — nested-imprecision TINF models, in which one component of a TINF quantity (most naturally T) is itself only known up to a p-box or credal set rather than a point value, motivated by the reliability-engineering case of §9.10 and the astronomy and energy-grid applications of §10.19 and §10.21.
14. New — an explicitly speculative research question: whether a quantum-probabilistic TINF, in which the four components become non-commuting under certain measurement orders, adds genuine explanatory power for order-effect anomalies in survey and preference data (§9.11), or is better left as an analogy.
15. New — a formal calibration and reliability-diagram methodology for TINF forecasters and classifiers, extending classical over/under-confidence diagnostics to all four components separately (§8.1).

A natural next publication-oriented step, surfaced during the drafting of this booklet, would restate Chapters 2–4 as a formal Definition → Lemma → Theorem → Proof sequence with 15–20 fully proved results (monotonicity, the complement identities, the total-probability and Bayes theorems, and convergence), together with fully worked proofs of the twelve reductions and structural analogies of Chapter 9, which would move this material from a framework proposal toward a submittable mathematical paper.

Appendix A: Notation Glossary

Symbol	Meaning
T, I, N, F	Truth-support, pure indeterminacy, neutrality, falsehood/opposition components
(T, I, N, F)	A TINF quadruple / TINF value
$\neg(T, I, N, F)$	Complement: (F, N, I, T)
$T(x, y), S(x, y)$	A t-norm and its dual t-conorm
s, a, ec, c	Score, accuracy, extended certainty, certainty ranking functions (§2.4)
$\mu_T, \mu_I, \mu_N, \mu_F$	The four components of a TINF measure
$P_N(A)$	TINF probability (general or normalized) of event A
X_T, X_I, X_N, X_F	The four coupled component random variables of a TINF r.v. X
$E_N[X], \text{Var}_N(X)$	TINF expectation and variance
$\bar{X}_N, \text{Med}_N$	TINF sample mean and median
Bel, Pl	Dempster-Shafer belief and plausibility (§9.1)
Π, N_{ec}	Possibility and necessity measures (§9.3; N_{ec} used to avoid clashing with this booklet's N)
$\underline{p}, \bar{p}$	Lower and upper prevision, imprecise probability (§9.4)
$g, \text{Choquet } f$	A fuzzy (non-additive) measure and its Choquet integral (§9.2)
B, N_{bd}	Belnap-Dunn truth values "Both" and "Neither" (§9.7)
io, no	Irreducible background indeterminacy / neutrality (§3.2, §4.3)

Appendix B: Worked Cross-Framework Conversions

This appendix works four short numerical examples, one per major reduction of Chapter 9, showing the conversion arithmetic explicitly rather than only stating the general formula. Each example starts from a quantity as it would be reported in the source framework and ends with the equivalent TINF quadruple, so that a reader already comfortable with one of these four frameworks has a concrete bridge into TINF notation.

B.1 Dempster-Shafer belief function to TINF (§9.1, §7.4)

Suppose a sensor network's basic probability assignment over the frame $\Theta = \{\text{intruder}, \text{no-intruder}\}$ gives $m(\{\text{intruder}\}) = 0.55$, $m(\{\text{no-intruder}\}) = 0.20$, $m(\Theta) = 0.25$ (the mass on Θ representing generic don't-know evidence not committed to either singleton). Then $\text{Bel}(\text{intruder}) = m(\{\text{intruder}\}) = 0.55$, $\text{Pl}(\text{intruder}) = m(\{\text{intruder}\}) + m(\Theta) = 0.80$, so the belief interval is $[0.55, 0.80]$. Applying the §7.4 reduction with $T = \text{Bel}(A) = 0.55$, $F = 1 - \text{Pl}(A) = 0.20$, $I = \text{Pl}(A) - \text{Bel}(A) = 0.25$, $N = 0$ gives the TINF quadruple $(0.55, 0.25, 0, 0.20)$ — note the components sum to exactly 1, as they must, since a singleton-focused Dempster-Shafer structure has no room for a fifth kind of mass beyond T, I, F. A genuinely TINF-native version of the same sensor read might instead report $(0.55, 0.15, 0.10, 0.20)$ if the network's analysts can further attribute part of that 0.25 don't-know mass to a specifically neutral reading (e.g. a sensor pattern independently known to be triggered by non-threat wildlife) rather than pure ignorance — the extra granularity Dempster-Shafer's frame cannot represent without adding a third focal element.

B.2 Possibility/necessity pair to TINF (§9.3)

Suppose a possibility distribution over next-day temperature bins gives $\Pi(\text{"above } 30^\circ\text{C"}) = 0.90$ and $N_{\text{ec}}(\text{"above } 30^\circ\text{C"}) = 0.40$ (i.e. it is highly possible, but not highly necessary, that tomorrow exceeds 30°C). Following §9.3, $T = N_{\text{ec}}(A) = 0.40$, $1 - F = \Pi(A) = 0.90$ so $F = 0.10$, and $I = \Pi(A) - N_{\text{ec}}(A) = 0.50$, $N = 0$, giving $(0.40, 0.50, 0, 0.10)$. The large I here correctly reflects that possibility theory, built for conceivability under vagueness rather than a third stance, has folded all of the gap between "necessary" and "merely possible" into pure indeterminacy — exactly the modeling limitation flagged in §9.3's discussion.

B.3 Rough-set boundary region to TINF (§9.5, §7.5)

Suppose a loan-approval rough set, built from five observable attributes, classifies 640 of 1000 historical applicants into the lower approximation of "creditworthy," 150 into the lower approximation of "not creditworthy," and leaves 210 in the boundary region (indiscernible from both classes given only those five attributes). The classical rough-set accuracy of approximation is $|\text{lower}|/|\text{upper}| = 790/1000 = 0.79$. Under the §7.5 reduction, a new applicant landing in the boundary region is assigned the TINF quadruple $(0, 1, 0, 0)$ — pure I, since the boundary region is by construction unresolvable with current attributes, not a claim about a genuine third credit category. If a bank later discovers that a specific subset of that boundary region — self-employed

applicants with seasonal income, say — is not actually unresolvable but is a recurring, legitimately intermediate risk class that additional standard attributes will never fully resolve, the correct model for that subset shifts from (0, 1, 0, 0) toward something like (0.30, 0.10, 0.55, 0.05): the rough-set machinery alone has no vocabulary for making that shift, which is exactly the point made generally in §9.5.

B.4 Credal set to TINF (§9.4)

Suppose a robust Bayesian analysis of a manufacturing defect rate, built from a small and only partially trusted sample, yields a credal set whose lower and upper previsions for "defect rate exceeds 2%" are $\underline{P} = 0.30$ and $\bar{P} = 0.55$. Following §9.4, $T = \underline{P}(A) = 0.30$, $F = 1 - \bar{P}(A) = 0.45$, $I = \bar{P}(A) - \underline{P}(A) = 0.25$, $N = 0$, giving (0.30, 0.25, 0, 0.45). A TINF-native re-analysis that separately tracks which part of that 0.25-wide imprecision comes from a small sample size (resolvable with more production runs, i.e. genuine I) versus a known, permanent measurement-process ambiguity in how "defect" is defined at the batch boundary (not resolvable by more data, i.e. genuine N) could instead report, say, (0.30, 0.10, 0.15, 0.45) — a decomposition a credal set's two-number summary cannot express, but which changes whether the recommended fix is "collect more data" or "clarify the defect definition."

Appendix C: Quick-Reference Decision Rule

Collected from the dominant-component decision tables scattered through Chapters 4, 9, and 10 (most explicitly §4.7's sensor-fusion table), the following generic rule recurs across nearly every application in this booklet and is offered here as a single, portable heuristic:

Dominant component of a TINF result	Generically correct next action
T (truth-support)	Act on the proposition as though established; proceed.
I (pure indeterminacy)	Seek more or better evidence; the situation is resolvable.
N (neutrality)	Stop seeking resolution; route to a dedicated third-category process built for genuinely in-between cases.
F (falsehood)	Act on the negation as though established; stand down.

The single sentence this booklet keeps returning to, in one form or another, across medicine, engineering, defense, law, and every other field in Chapter 10, is this: an I-shaped problem is solved by gathering more information, and an N-shaped problem is not, and a decision system that cannot tell the two apart will eventually apply the wrong fix to one of them.

Appendix D: Suggested Reading Path by Background

This booklet can be read start to finish, but readers arriving with a specific background often want a shorter, targeted path. Four are suggested below.

D.1 For readers coming from classical probability and statistics

Read Chapter 1 in full, then §4.1–§4.4 and §6.1–§6.6, then §7.1 (the classical reduction), then §9.8 (the Bayesian/frequentist comparison), then pick two or three applications from Chapter 10 in your own field. This path emphasizes that everything you already know is recovered exactly as a special case, and shows precisely where the new machinery adds value beyond it.

D.2 For readers coming from Dempster-Shafer, fuzzy sets, or possibility theory

Read Chapter 1, then skip directly to §9.1–§9.6 and §7.2–§7.5 (the reductions), then §9.13's worked cross-notation example, then Chapter 4 for the probability layer. This path leads with the comparison chapter, which is likely to answer your first question — "how is this different from what I already use" — before asking you to learn new notation.

D.3 For readers coming from AI, machine learning, or decision science

Read Chapter 1, §2.4 (the ranking functions, directly relevant to any system that must ultimately output a decision), §4.7 (the sensor-fusion worked example), Chapter 8 (Over/Under/Off, relevant to confidence calibration), and §9.6–§9.7 (orthopair fuzzy sets and Belnap–Dunn logic, the two frameworks closest to current AI-uncertainty practice), then browse Chapter 10 for applications matching your domain.

D.4 For readers primarily interested in the philosophical/foundational question

Read the Preface in full, §1.1–§1.3, §9.7 (Belnap–Dunn, the closest independent motivation for a similar split), §9.11 (the quantum-probability discussion, flagged as speculative), and Chapter 11 in full. This path follows the conceptual argument — why a mandated third stance is not the same mathematical object as missing information — with the least amount of computational machinery in between.

References

- [1] F. Smarandache, "(T, I, N, F) Neutrosophic Set and Logic," Critical Review, Vol. XIV, 2017, pp. 20–28. The source paper for this booklet; introduces the quadruple refinement of Indeterminacy into Pure Indeterminacy and Neutrality, its set/logic operators, and the score/accuracy/certainty ranking functions, and names Measure/Probability/Statistics as future extensions. <https://fs.unm.edu/CR/TINF-NeutrosophicSetLogic.pdf>
- [2] F. Smarandache, Neutrosophy / Neutrosophic Probability, Set, and Logic, ProQuest, Michigan, USA, 1998. <https://fs.unm.edu/eBook-Neutrosophics6.pdf>
- [3] F. Smarandache, "n-Valued Refined Neutrosophic Logic and Its Applications to Physics," Progress in Physics, Vol. 4, 2013, pp. 143–146. <https://fs.unm.edu/n-ValuedNeutrosophicLogic-PiP.pdf>
- [4] F. Smarandache, Introduction to Neutrosophic Measure, Neutrosophic Integral, and Neutrosophic Probability, Sitech Publishing House, Craiova, 2013. Develops neutrosophic measure, integral, and probability and notes that different forms of indeterminacy can motivate different mathematical constructions.
- [5] F. Smarandache, "Neutrosophic Measure and Neutrosophic Integral," Neutrosophic Sets and Systems, Vol. 1, 2013, pp. 3–7.
- [6] F. Smarandache, Introduction to Neutrosophic Statistics, Sitech & Education Publishing, 2014. Develops neutrosophic descriptive statistics and probability-distribution analysis under indeterminacy.
- [7] F. Smarandache, Neutrosophic Overset, Neutrosophic Underset, and Neutrosophic Offset, Similarly for Neutrosophic Over-/Under-/Off- Logic, Probability, and Statistics, Pons Editions, Brussels, 2016. <https://fs.unm.edu/Over-Under-Off.htm>
- [8] F. Smarandache, "Operators on Single-Valued Neutrosophic Oversets, Neutrosophic Undersets, and Neutrosophic Offsets," Journal of Mathematics and Informatics, Vol. 5, 2016, pp. 63–67. <https://fs.unm.edu/SVNeutrosophicOverset-JMI.pdf>
- [9] F. Smarandache, "Neutrosophic Score, Accuracy, and Certainty Functions Form a Total Order Relationship on the Set of Neutrosophic Triplets (T, I, F)," used in Multi-Criteria Decision Making. <https://fs.unm.edu/NSS/TheScoreAccuracyAndCertainty1.pdf>
- [10] F. Smarandache, "Operators for Uncertain Over/Under/Off-Sets/-Logics/-Probabilities/-Statistics," Neutrosophic Sets and Systems, Vol. 79, 2025, pp. 493–500.
- [11] F. Smarandache, N. Abbas, Y. Chibani, and Z. A. Omar, "PCR5 and Neutrosophic Probability in Target Identification (revisited)," Neutrosophic Sets and Systems, Vol. 16, 2017, pp. 76–79. Illustrates neutrosophic probability and Dezert-Smarandache combination in information fusion and target identification.

- [12] F. Smarandache, "Plithogenic Probability & Statistics are Generalizations of MultiVariate Probability & Statistics," *Neutrosophic Sets and Systems*, Vol. 43, 2021, pp. 280–289.
- [13] F. Smarandache, "Indeterminacy in Neutrosophic Theories and Their Applications," *International Journal of Neutrosophic Science*, Vol. 15, No. 2, 2021, pp. 89–97. <https://fs.unm.edu/Indeterminacy.pdf>
- [14] A. P. Dempster, "Upper and Lower Probabilities Induced by a Multivalued Mapping," *Annals of Mathematical Statistics*, Vol. 38, No. 2, 1967, pp. 325–339.
- [15] G. Shafer, *A Mathematical Theory of Evidence*, Princeton University Press, 1976. Foundational text for Dempster-Shafer evidence theory, referenced in §9.1 and §7.4.
- [16] M. Sugeno, "Theory of Fuzzy Integrals and Its Applications," Ph.D. thesis, Tokyo Institute of Technology, 1974. Origin of fuzzy measures and the Sugeno integral discussed in §9.2.
- [17] G. Choquet, "Theory of Capacities," *Annales de l'Institut Fourier*, Vol. 5, 1954, pp. 131–295. Origin of the Choquet integral used in §9.2.
- [18] L. A. Zadeh, "Fuzzy Sets as a Basis for a Theory of Possibility," *Fuzzy Sets and Systems*, Vol. 1, No. 1, 1978, pp. 3–28. Foundational text for possibility theory, discussed in §9.3.
- [19] D. Dubois and H. Prade, *Possibility Theory: An Approach to Computerized Processing of Uncertainty*, Plenum Press, New York, 1988.
- [20] P. Walley, *Statistical Reasoning with Imprecise Probabilities*, Chapman and Hall, London, 1991. Foundational text for imprecise probability and credal sets, discussed in §9.4.
- [21] Z. Pawlak, "Rough Sets," *International Journal of Computer and Information Sciences*, Vol. 11, No. 5, 1982, pp. 341–356. Foundational text for rough set theory, discussed in §9.5 and §7.5.
- [22] K. T. Atanassov, "Intuitionistic Fuzzy Sets," *Fuzzy Sets and Systems*, Vol. 20, No. 1, 1986, pp. 87–96. Discussed in §9.6 and §7.2.
- [23] R. R. Yager, "Pythagorean Fuzzy Subsets," *Proceedings of the Joint IFSA World Congress and NAFIPS Annual Meeting*, 2013, pp. 57–61.
- [24] R. R. Yager, "Generalized Orthopair Fuzzy Sets," *IEEE Transactions on Fuzzy Systems*, Vol. 25, No. 5, 2017, pp. 1222–1230. Origin of q-rung orthopair fuzzy sets, discussed in §9.6.
- [25] N. D. Belnap, "A Useful Four-Valued Logic," in *Modern Uses of Multiple-Valued Logic*, J. M. Dunn and G. Epstein (eds.), Reidel, Dordrecht, 1977, pp. 5–37. Foundational text for the four-valued logic discussed in §9.7.
- [26] J. M. Dunn, "Intuitive Semantics for First-Degree Entailments and 'Coupled Trees'," *Philosophical Studies*, Vol. 29, No. 3, 1976, pp. 149–168.

- [27] G. Matheron, *Random Sets and Integral Geometry*, Wiley, New York, 1975. Foundational text for random set theory, discussed in §9.10.
- [28] H. T. Nguyen, *An Introduction to Random Sets*, Chapman and Hall/CRC, 2006.
- [29] S. Ferson, V. Kreinovich, L. Ginzburg, D. S. Myers, and K. Sentz, "Constructing Probability Boxes and Dempster-Shafer Structures," Sandia National Laboratories, SAND2002-4015, 2003. Reference for p-box methodology discussed in §9.10.
- [30] P. J. Huber, *Robust Statistics*, Wiley, New York, 1981. Foundational text for robust statistics, discussed in §9.9.
- [31] R. P. Feynman, "Negative Probability," in *Quantum Implications: Essays in Honour of David Bohm*, B. J. Hiley and F. D. Peat (eds.), Routledge, 1987, pp. 235-248. Background reference for the quantum-probability discussion of §9.11.
- [32] J. R. Busemeyer and P. D. Bruza, *Quantum Models of Cognition and Decision*, Cambridge University Press, 2012. Reference for the decision-theoretic quantum-probability anomalies discussed in §9.11.
- [33] E. H. Kennard, "Zur Quantenmechanik einfacher Bewegungstypen," *Zeitschrift für Physik*, Vol. 44, 1927, pp. 326-352. Background reference for the non-commutative measurement structure discussed in §9.11.
- [34] C. Chen and J. Tan, "Handling Multicriteria Fuzzy Decision-Making Problems Based on Vague Set Theory," *Fuzzy Sets and Systems*, Vol. 67, No. 2, 1994, pp. 163-172. Origin of the score-function tradition discussed in §9.6.

Note on This Booklet's Preparation

The mathematical development in Chapters 2-11 — the axioms, propositions, proofs, worked numerical examples, reduction theorems, the twelve cross-framework comparisons of Chapter 9, and the twenty-seven field applications of Chapter 10 — builds directly on the author's reference [1] and was prepared with substantial AI-assisted drafting: several independent passes (ChatGPT, Gemini, and Claude) were generated, compared, and merged under the author's direction, with a second drafting pass dedicated to situating the framework relative to the broader uncertainty-theory literature referenced in [14]-[34]. This transparency note is included so that readers know how the text was produced. It was reviewed by the peer reviewers listed on the imprint page and is offered as a starting point for further, more rigorous work along the lines set out in Chapter 11.



ABOUT THIS BOOKLET

Florentin Smarandache's 2017 paper refines the neutrosophic triple (Truth, Indeterminacy, Falsehood) by splitting Indeterminacy into Pure Indeterminacy (I — we don't yet know) and Neutrality (N — it genuinely is neither). The paper works this out fully for Sets and Logic, and names Measure, Probability, and Statistics as the natural next step.

This booklet builds that next step: axioms and a worked proof for TINF measure, two readings of TINF probability, TINF random variables and descriptive/inferential statistics, reduction theorems back to classical and fuzzy theory, a full chapter connecting TINF to twelve established uncertainty frameworks, and twenty-seven field applications — each showing concretely why the I/N split earns its keep.

ABOUT THE AUTHOR

Florentin Smarandache is a professor of mathematics at the University of New Mexico (Gallup campus), where he founded neutrosophic set, logic, probability, and statistics theory — a generalization of fuzzy and intuitionistic fuzzy theory built around an explicit indeterminacy component. He is a member of the Academy of Science Peloritana dei Pericolanti, University of Messina, Italy, and has authored or co-authored several hundred papers and books across mathematics, philosophy, and the applied neutrosophic sciences. Contact: smarand@unm.edu

Source paper:

fs.unm.edu/CR/TINF-NeutrosophicSetLogic.pdf

NSIA Publishing House · University of New Mexico, Gallup, NM, USA · University of Guayaquil, Ecuador



ISBN 978-1-972502-44-0