

Splitting Meanings: A Unified View on Paraconsistency and Inconsistency Measurement

Yakoub Salhi

Univ. Artois, CNRS, CRIL, F-62300 Lens, France
salhi@cril.fr

Abstract

We propose a modular logical framework for both reasoning with and measuring inconsistency in propositional knowledge bases. The framework extends propositional logic with markers attached to occurrences of atoms and interpreted as pointers to worlds. This allows different occurrences of the same atom to be evaluated in different contexts unless they share a marker. We define paraconsistent entailment relations via abnormality functions that associate each model with a set of deviations from a preferred behavior. Each entailment relation is then defined with respect to models that are minimal under set inclusion. We show that suitable markings and abnormality functions capture several existing forms of inconsistency-tolerant reasoning, including entailment based on maximal satisfiable subsets and the minimally inconsistent Logic of Paradox. We also show how a range of inconsistency measures can be expressed in the same setting. Thus our framework provides a uniform basis for diverse approaches to inconsistency handling and measurement.

1 Introduction

Real-world knowledge bases frequently exhibit inconsistencies because information is aggregated from diverse sources, using different methodologies, or at different times. This heterogeneity often results in conflicts within these knowledge bases. Classical propositional logic is inadequate for reasoning under such conditions: once an inconsistency arises, the principle of explosion makes every formula derivable. Consequently, this limitation has motivated the development of various formalisms for handling inconsistency, including argumentation [Besnard and Hunter, 2008], belief revision [Hansson, 2022], and paraconsistency [Priest *et al.*, 2022]. Inconsistency measurement complements these by providing quantitative ways to assess how far a knowledge base is from consistency (e.g., see [Thimm, 2018; Bona *et al.*, 2019]).

In this paper, we propose a framework that can serve as a common basis for both tasks. The main idea is to distinguish between the surface level, where formulas are written in an ordinary propositional language, and a marked level,

where each occurrence of an atom is equipped with a marker. Markers are interpreted as pointers to worlds in a possible-worlds structure [Kripke, 1959]. Intuitively, a marker identifies a context, source, or meaning in which the corresponding occurrence is evaluated. Two occurrences of the same propositional variable need not share a truth value unless they share a marker.

Using the possible-worlds semantics, we define minimality-based entailment through abnormality functions. An abnormality function assigns to each model a set of *defects*. We compare models by set inclusion and prefer those with smaller defect sets. Our approach is inspired by abnormality minimization in nonmonotonic reasoning [McCarthy, 1986]. Different choices of abnormality function capture different notions of what should be considered problematic: splitting markers into different worlds, changing the truth value of a variable across markers, or deviating from the normal world. Importantly, all these instantiations live in the same simple semantic universe.

Despite its simplicity, the framework is expressive. On the one hand, it can recover a variety of paraconsistent entailment relations. First, it can capture an entailment relation based on maximal satisfiable subsets (MSSes) [Rescher and Manor, 1970] by using markers that distinguish different formulas and a minimality criterion that penalizes non-normal markers. Second, it can encode the minimally inconsistent Logic of Paradox [Priest, 1991] via two special markers and an abnormality function that counts variables taking abnormal truth values. Third, it can reproduce an entailment relation defined through occurrence-based semantics [Salhi, 2025] by using an abnormality function that minimizes the number of variables whose occurrences are forced to disagree in truth value. Each of these formalisms arises as a particular choice of marking strategy and abnormality function within the same generic scheme.

On the other hand, our framework can be used to give a uniform basis for a broad family of inconsistency measures. Different measures correspond to different ways of counting or minimizing abnormalities in marked models. In particular, by choosing appropriate markings and abnormality functions, we can express:

- formula-based measures such as those defined in terms of MSSes and problematic formulas [Grant and Hunter, 2011];

- model-based measures such as Dalal’s distance-based measures [Grant and Hunter, 2017];
- paraconsistency-based measures such as contention [Konieczny *et al.*, 2003; Grant and Hunter, 2011];
- a forgetting-based measure [Besnard, 2016], which quantifies how many occurrences must be forgotten (replaced by constants) to restore consistency.

In each case, the value of the measure is recovered as the cardinality of some minimal abnormality set.

The benefits of our framework can be summarized as follows. First, the underlying logic is purely propositional with a standard, two-valued semantics at each world. Markers and abnormality functions add only a lightweight layer to this classical basis. Second, different forms of paraconsistent reasoning are obtained by only varying the marking strategy and the choice of abnormality function. Third, diverse approaches to inconsistency measurement, which are usually presented in different frameworks, can be seen as instances of the same general pattern of minimality over marked models.

2 Related Work

Our use of markers is connected to various labelled formalisms in the literature. Labelled proof systems attach labels to formulas at the syntactic level in order to reason about possible worlds, time points, or other parameters [Fitch, 1966; Gabbay, 1996; Viganò, 2000]. In particular, labelled systems for the normal modal logic S5 can avoid an explicit accessibility relation by working inside a single equivalence class of worlds. Hybrid logics [Blackburn, 2000] also enrich modal languages with special atomic symbols, called nominals, together with an operator that allows one to express that a formula holds at a designated state. In this way, hybrid logics provide a formal mechanism for contextual reasoning, with nominals serving to identify and reason about distinct contexts. Other approaches to contextual reasoning based on syntactic or semantic enrichment have been developed in various forms in the literature; see, for example, [Ghidini and Giunchiglia, 2001; Serafini and Bouquet, 2004; Klarman and Gutiérrez-Basulto, 2011; Serafini and Homola, 2012; Gómez Álvarez and Rudolph, 2021]. However, these frameworks are not designed to single out and manipulate propositional variables via labels, nor do they aim to model inconsistency or defining inconsistency-tolerant entailment relations. In particular, they do not typically combine labels with explicit minimality principles that directly target abnormal behavior of variables or occurrences.

To deal with inconsistency, our framework uses multiple worlds to allow an atom to be associated with more than one truth value. Assigning more than one truth value to an atom is a central idea in paraconsistency, and in particular in the (minimally inconsistent) Logic of Paradox [Priest, 1991]. In contrast to many-valued approaches, however, our use of worlds and markers yields a more fine-grained perspective: it operates at the level of occurrences of atoms, rather than at the level of atoms alone.

In a recent article [Salhi, 2025], occurrence-based entailment relations were introduced. However, that approach does

not enforce that particular occurrences share the same truth value, whereas our framework achieves this through a shared marker. Moreover, the semantics proposed in [Salhi, 2025] does not single out any occurrence as *normal*, while our framework designates a dedicated marker corresponding to the normal world. Finally, that approach does not support the range of minimality conditions naturally expressible in the present framework.

Despite these differences, we show that both the Logic of Paradox and occurrence-based entailment relations can be encoded in a simple and uniform way in this framework.

3 Preliminaries

Let Prop be a countable set of propositional variables. The set PF of propositional formulas over Prop is defined inductively by the grammar $\varphi ::= p \mid \perp \mid \neg\varphi \mid (\varphi \wedge \varphi) \mid (\varphi \vee \varphi)$, where $p \in \text{Prop}$. We use the usual definitions of the other connectives.

Let $\text{vars}(\varphi)$ denote the set of propositional variables occurring in φ .

An *interpretation* over Prop is a function $\omega : \text{Prop} \rightarrow \{0, 1\}$, where 1 stands for *true* and 0 for *false*. We write $\Omega(\text{Prop})$ for the set of all interpretations.

The satisfaction relation $\omega \models_{\text{CL}} \varphi$ is defined inductively as usual.

For a set $\Gamma \subseteq \text{PF}$, we write $\omega \models_{\text{CL}} \Gamma$ iff $\omega \models_{\text{CL}} \varphi$ for all $\varphi \in \Gamma$. We say that Γ is *satisfiable* (or *consistent*) iff there exists $\omega \in \Omega(\text{Prop})$ such that $\omega \models_{\text{CL}} \Gamma$. Otherwise, Γ is *unsatisfiable* (or *inconsistent*).

In connection with inconsistency, we focus on the following three standard notions.

Definition 1 (Minimal Unsatisfiable Subset). *A subset $M \subseteq \Gamma$ is a minimal unsatisfiable subset (MUS) of Γ iff: (i) M is unsatisfiable; and (ii) for every $M' \subsetneq M$, the set M' is satisfiable.*

Definition 2 (Maximal Satisfiable Subset). *A subset $S \subseteq \Gamma$ is a maximal satisfiable subset (MSS) of Γ iff: (i) S is satisfiable; and (ii) for every $\varphi \in \Gamma \setminus S$, the set $S \cup \{\varphi\}$ is unsatisfiable.*

Definition 3 (Minimal Correction Subset). *A subset $C \subseteq \Gamma$ is a minimal correction subset (MCS) of Γ iff: (i) $\Gamma \setminus C$ is satisfiable; and (ii) for every $C' \subsetneq C$, the set $\Gamma \setminus C'$ is unsatisfiable.*

For a set of formulas Γ , we denote by $\text{MUS}(\Gamma)$ (resp., $\text{MSS}(\Gamma)$, $\text{MCS}(\Gamma)$) the set of all minimal unsatisfiable subsets of Γ (resp., maximal satisfiable subsets of Γ , minimal correction subsets of Γ).

4 A Propositional Logic with Markers

In this section, we introduce a propositional logic in which occurrences of propositional variables are equipped with markers, and show how these markers can be used to handle inconsistency.

Consider a set of formulas $K = \{p \vee q, \neg p, \neg q\}$, which is clearly inconsistent. In many situations, we do not want to treat the two occurrences of p or those of q as describing exactly the same piece of information. For example:

- p could be read as “the system is hot” in two different physical contexts (e.g., two different sensors, or two different time points);
- or as “this student is eligible for the scholarship” as assessed by two different models;
- or more generally as the same predicate used with different meanings or sources.

Intuitively, it is coherent to accept that one source says p while another says $\neg p$, without collapsing into triviality.

To capture this formally, we introduce *markers* that allow us to distinguish occurrences of the same variable. For instance, we can write $K^* = \{p^\alpha \vee q^\beta, \neg p^\gamma, \neg q^\beta\}$, where α, β and γ are markers. If $\alpha \neq \gamma$, the two occurrences of p may receive different truth values. Thus, the marked set K^* is satisfiable. At the same time, we retain the option of forcing two occurrences to share the same meaning by giving them the same marker. This can be observed for the two occurrences of q .

As a result, the logic with markers distinguishes two levels: (1) the surface level, where the same letter p may appear in many places; (2) the marked level, where each occurrence is refined into p^α , and only occurrences with the same marker are required to share a truth value. This provides a natural way to reason about inconsistency that is sensitive to potential differences in meaning or information source.

Let us now introduce the syntax and the semantics of the logic with markers, denoted ML. We assume a countable set of markers Mk that contains the specific symbol ν . Markers are denoted by letters such as α, β , and γ , possibly with subscripts.

The set of ML formulas (marked formulas) is obtained by the following grammar: $\varphi ::= p^\alpha \mid \perp \mid \neg\varphi \mid (\varphi \wedge \varphi) \mid (\varphi \vee \varphi)$, where $p \in \text{Prop}$ and $\alpha \in \text{Mk}$.

Intuitively, p^ν is read as “ p in the normal context”, while p^α , with $\alpha \neq \nu$, is read as “ p in context, source, or meaning α ”. Two occurrences p^α and p^β are treated as potentially distinct pieces of information unless we explicitly identify α and β .

A central idea of the semantics is that markers control which occurrences must share a truth value. We model this by allowing valuations to depend on equivalence classes of markers.

An ML-model is a tuple $\mathcal{M} = \langle W, w_0, V, \lambda \rangle$, where W is a set of possible worlds, $w_0 \in W$, $V : W \rightarrow 2^{\text{Prop}}$ is a valuation function, and $\lambda : \text{Mk} \rightarrow W$ with $\lambda(\nu) = w_0$. Intuitively, w_0 is used to represent what we regard as the “normal” meaning/context.

Given an ML-model $\mathcal{M} = \langle W, w_0, V, \lambda \rangle$ and an ML formula φ , the satisfaction relation $\mathcal{M} \models_{\text{ML}} \varphi$ is defined inductively as follows:

$$\begin{aligned}
 \mathcal{M} \models_{\text{ML}} p^\alpha & \text{ iff } p \in V(\lambda(\alpha)), \\
 \mathcal{M} \models_{\text{ML}} \perp & \text{ iff never,} \\
 \mathcal{M} \models_{\text{ML}} \neg\varphi & \text{ iff } \mathcal{M} \not\models_{\text{ML}} \varphi, \\
 \mathcal{M} \models_{\text{ML}} (\varphi \wedge \psi) & \text{ iff } \mathcal{M} \models_{\text{ML}} \varphi \text{ and } \mathcal{M} \models_{\text{ML}} \psi, \\
 \mathcal{M} \models_{\text{ML}} (\varphi \vee \psi) & \text{ iff } \mathcal{M} \models_{\text{ML}} \varphi \text{ or } \mathcal{M} \models_{\text{ML}} \psi.
 \end{aligned}$$

We use $\mathbb{M}(\text{Prop}, \text{Mk})$ for the set of all ML-models.

A set of marked formulas Γ is *satisfiable* iff there exists an ML-model $\mathcal{M}$ such that $\mathcal{M} \models_{\text{ML}} \varphi$ for all $\varphi \in \Gamma$. We write $\text{Mod}(\Gamma) = \{\mathcal{M} \in \mathbb{M}(\text{Prop}, \text{Mk}) \mid \mathcal{M} \models_{\text{ML}} \Gamma\}$.

We define a *marker configuration* of an ML-model as the equivalence relation $\sim_\lambda$ on markers induced by λ : $\alpha \sim_\lambda \beta$ iff $\lambda(\alpha) = \lambda(\beta)$.

Different choices of marker configuration $\sim_\lambda$ correspond to different assumptions about when two occurrences must share a meaning:

- If $\sim_\lambda$ is the identity relation (only $\alpha \sim_\lambda \alpha$), then markers are fully independent. In this case, occurrences like p^α and p^β (with $\alpha \neq \beta$) may freely receive different truth values.
- If $\sim_\lambda$ is the universal relation (all markers equivalent), then for each p there is effectively a single truth value shared by all p^α . In this case, we recover the classical semantics for propositional logic.

Therefore, the logic with markers provides a flexible framework in which we can control, via the marker configuration, which occurrences of the same propositional variable are forced to behave classically and which may diverge. This allows a more fine-grained analysis of inconsistency.

5 Minimality and Entailment

The introduction of markers can greatly increase the number of models satisfying a marked formula, since both the marker configuration and the valuation can vary while still satisfying the same marked formula. Intuitively, however, not all of these models should be considered equally plausible. Some models are more conservative: they identify more markers with one another, or keep propositional variables as stable as possible across their occurrences. Other models exhibit more abnormal behavior, for instance by splitting markers into different worlds or allowing the same variable to take different truth values.

When information is potentially inconsistent, it is natural to define entailment using only those models that are minimal in a well-motivated sense. This idea is common in nonmonotonic reasoning and belief revision: among the possible ways of handling the information, one prefers those that deviate as little as possible from a preferred pattern of behavior. Our approach is particularly inspired by abnormality minimization in nonmonotonic reasoning [McCarthy, 1986].

In our setting, the idea is as follows: we split meanings only when necessary, and we change truth values only when unavoidable. We capture this by associating with each model a set of *abnormalities* and then restricting attention to models whose abnormality set is minimal under set inclusion.

Definition 4 (Abnormality functions and minimal models). *An abnormality function is any function $f : \mathbb{M}(\text{Prop}, \text{Mk}) \rightarrow 2^U$, where U is some fixed universe of (abnormality) items (e.g. markers, variables, or both).*

For a set Γ of marked formulas, the set of f -minimal models of Γ is defined by $\text{Mod}_f^{\min}(\Gamma) = \{\mathcal{M} \in \text{Mod}(\Gamma) \mid \nexists \mathcal{N} \in \text{Mod}(\Gamma) \text{ such that } f(\mathcal{N}) \subsetneq f(\mathcal{M})\}$.

Entailment is defined via truth in all f -minimal models.

Definition 5 (Minimality-based entailment). *Let f be an abnormality function. For sets Γ of ML formulas and an ML formula φ , we define the entailment relation $\sim_f^{\text{mk}}$ as follows:*

$$\Gamma \sim_f^{\text{mk}} \varphi \text{ iff } \mathcal{M} \models_{\text{ML}} \varphi \text{ for all } \mathcal{M} \in \text{Mod}_f^{\text{min}}(\Gamma).$$

Different choices of abnormality function f give rise to different entailment relations. Conceptually, $f(\mathcal{M})$ specifies what the current instantiation regards as problematic in $\mathcal{M}$, and minimality then selects those models that minimize the problematic items.

We now describe a few simple but useful abnormality functions. These will serve as building blocks in later sections, where we show how different choices of f and different marking strategies recover a range of well-known inconsistency-tolerant entailment relations.

Let $\mathcal{M} = \langle W, w_0, V, \lambda \rangle$ be an ML-model.

1. **Non-normal markers.** We first quantify how much a model deviates from the normal world by using markers that do not point to w_0 :

$$f_1(\mathcal{M}) = \{ \alpha \in \text{Mk} \mid \lambda(\alpha) \neq w_0 \}.$$

Minimizing f_1 favors models in which as many markers as possible are interpreted at the normal world, i.e., models that use the normal meaning or context whenever this is compatible with consistency.

2. **Splitting of markers.** Next, we measure to what extent markers are separated into different worlds. The following abnormality function works relative to a fixed set of formulas Γ . Indeed, we restrict our attention to the set of markers that actually occur in Γ . We use Mk_Γ to denote the set of markers occurring in Γ .

$$f_2^\Gamma(\mathcal{M}) = \{ \{ \alpha, \beta \} \subseteq \text{Mk}_\Gamma \mid \lambda(\alpha) \neq \lambda(\beta) \}.$$

Each pair $\{ \alpha, \beta \}$ in $f_2^\Gamma(\mathcal{M})$ indicates that α and β belong to distinct $\sim_\lambda$ -equivalence classes. Thus, minimizing f_2^Γ favors models that identify as many markers as possible, i.e., models with fewer distinct meanings or contexts. In what follows, we write f_2 instead of f_2^Γ whenever the underlying set Γ is clear from the context.

3. **Truth-value changes.** A third dimension of abnormality concerns variables whose truth value is not stable across worlds pointed to by markers:

$$f_3(\mathcal{M}) = \{ p \in \text{Prop} \mid \exists \alpha \in \text{Mk} \text{ such that } p \in V(w_0) \text{ iff } p \notin V(\lambda(\alpha)) \}.$$

Thus, a propositional variable p is in $f_3(\mathcal{M})$ iff p takes a different truth value at some marker world than it does at the normal world. Minimizing f_3 expresses a preference for models in which each variable keeps a single truth value across all its occurrences, unless inconsistency forces a change.

4. **Variable-marker abnormalities.** Finally, we may wish to record not only which variables change truth value, but also where they change. This provides a more fine-grained abnormality function:

$$f_4(\mathcal{M}) = \{ (p, \alpha) \mid p \in \text{Prop}, \alpha \in \text{Mk}, p \in V(w_0) \text{ iff } p \notin V(\lambda(\alpha)) \}.$$

In this case, the abnormalities are individual pairs (p, α) at which p deviates from its normal truth value. Minimizing f_4 penalizes each such deviation separately, allowing more fine-grained control over the trade-offs between splitting meanings and changing truth values.

These functions can also be combined or refined in various ways. In the rest of the paper we will show that, together with suitable marking strategies, they suffice to recover several familiar forms of inconsistency-tolerant entailment relations, as well as a broad range of inconsistency measures.

6 Encoding Paraconsistent Entailment Relations

To illustrate the flexibility of our framework, we now show how it can encode several paraconsistent entailment relations. More precisely, we consider three instantiations: an entailment relation based on MSSes [Rescher and Manor, 1970], the entailment relation of the minimally inconsistent Logic of Paradox [Priest, 1991], and an entailment relation defined via an occurrence-based semantics [Salhi, 2025]. These examples cover approaches that are often seen as orthogonal in the literature and demonstrate that the ML semantics, combined with suitable abnormality functions, can uniformly capture a spectrum of paraconsistent reasoning principles.

6.1 An MSS-based Entailment

We first consider the following well-known entailment relation based on MSSes [Rescher and Manor, 1970]. For $\Gamma \subseteq \text{PF}$ and $\varphi \in \text{PF}$, the entailment relation $\vdash_{\text{MSS}}$ is defined as follows:

$$\Gamma \vdash_{\text{MSS}} \varphi \text{ iff } \forall S \in \text{MSS}(\Gamma) : S \vdash_{\text{CL}} \varphi.$$

To relate this to our framework, we first transform formulas to marked formulas. Given a formula ψ and a marker α , we write $\tau_{\text{form}}(\psi, \alpha)$ for the marked formula obtained from ψ by replacing each occurrence of a variable p with p^α .

We now fix a marking strategy that associates with each formula $\psi \in \Gamma$ a distinct marker $\alpha_\psi \neq \nu$ (i.e. different formulas receive different non-normal markers). Given a set of formulas Γ , we set $\tau_{\text{form}}(\Gamma) = \{ \tau_{\text{form}}(\psi, \alpha_\psi) \mid \psi \in \Gamma \}$.

The following lemma states that, whenever a marker α is interpreted at the normal world w_0 , the satisfaction of the marked translation $\tau_{\text{form}}(\psi, \alpha)$ coincides with the classical satisfaction of ψ at w_0 .

Lemma 1. *Let $\mathcal{M} = \langle W, w_0, V, \lambda \rangle$ be an ML-model and $\omega_{\mathcal{M}}$ the classical interpretation defined by $\omega_{\mathcal{M}}(p) = 1$ iff $p \in V(w_0)$. Then, for every formula $\psi \in \text{PF}$ and every marker α with $\lambda(\alpha) = w_0$, $\mathcal{M} \models_{\text{ML}} \tau_{\text{form}}(\psi, \alpha)$ iff $\omega_{\mathcal{M}} \models_{\text{CL}} \psi$.*

Proof. By a straightforward induction on the structure of ψ , using the clauses for ML-satisfaction and the definition of τ_{form} . $\square$

Now, we can characterize MSS-based entailment in terms of our marked semantics and the abnormality function f_1 from Section 5 ($f_1(\mathcal{M})$ is the set of markers that do not point to the normal world w_0).

Note that the correspondence requires each formula in the considered set to be classically satisfiable. Our encoding represents the choice of keeping or discarding a formula in an MSS by deciding whether its marker can stay at the normal world w_0 or must be assigned to a non-normal world.

If a formula ψ is itself inconsistent, then $\tau_{\text{form}}(\psi, \alpha)$ is unsatisfiable at every world, so there is no model in which ψ can be kept. In that case, the one-to-one correspondence between MSSes of Γ and f_1 -minimal models of $\tau_{\text{form}}(\Gamma)$ breaks down.

Theorem 1. *For every set Γ whose members are individually consistent, and every formula φ , $\Gamma \vdash_{\text{MSS}} \varphi$ iff $\tau_{\text{form}}(\Gamma) \vdash_{f_1}^{\text{mk}} \tau_{\text{form}}(\varphi, \nu)$.*

Intuitively, each f_1 -minimal model of $\tau_{\text{form}}(\Gamma)$ corresponds to choosing an MSS of Γ : the formulas whose markers stay at the normal world w_0 are exactly those that are kept, while moving a marker away from w_0 amounts to discarding the corresponding formula. Lemma 1 guarantees that truth at w_0 for the marked translations coincides with classical truth for the original formulas. Conversely, every MSS of Γ can be turned into an f_1 -minimal model by assigning its formulas to w_0 and the others to non-normal worlds.

6.2 Minimally Inconsistent Logic of Paradox

We now recall the minimally inconsistent Logic of Paradox LP_m [Priest, 1991] and show how it can be captured in our framework.

An LP_m interpretation is a function $\mu : \text{PF} \rightarrow \{\{0\}, \{1\}, \{0, 1\}\}$ that assigns to each propositional formula $\varphi \in \text{PF}$ a set of classical truth values and satisfies the following conditions: $\mu(\perp) = \{0\}$, $\mu(\neg\varphi) = \{1 - v \mid v \in \mu(\varphi)\}$, $\mu(\varphi \wedge \psi) = \{\min\{v, v'\} \mid v \in \mu(\varphi), v' \in \mu(\psi)\}$, and $\mu(\varphi \vee \psi) = \{\max\{v, v'\} \mid v \in \mu(\varphi), v' \in \mu(\psi)\}$.

We write $\mu! = \{p \in \text{Prop} \mid \mu(p) = \{0, 1\}\}$ for the set of variables that receive both truth values under μ .

We say that μ is an LP_m model of a formula φ , and write $\mu \models_{LP_m} \varphi$, iff $1 \in \mu(\varphi)$. The LP_m interpretation μ is an LP_m model of a set of formulas Γ , written $\mu \models_{LP_m} \Gamma$, iff it is a model for every formula in Γ .

The LP_m interpretation μ is a minimally inconsistent (or simply minimal) LP_m model of Γ iff $\mu \models_{LP_m} \Gamma$, and for every LP_m interpretation μ' with $\mu'! \subsetneq \mu!$, $\mu' \not\models_{LP_m} \Gamma$.

A set of formulas Γ entails φ in LP_m , written $\Gamma \vdash_{LP_m} \varphi$, iff for every minimal LP_m model μ of Γ , we have $\mu \models_{LP_m} \varphi$.

Let pos and neg be two distinct markers in $\text{Mk} \setminus \{\nu\}$. Given a formula $\varphi \in \text{PF}$, we define the marked translations $\pi^{\text{pos}}(\varphi)$ and $\pi^{\text{neg}}(\varphi)$ inductively as follows:

- $\pi^{\text{pos}}(\perp) = \perp$, $\pi^{\text{neg}}(\perp) = \perp$; $\pi^{\text{pos}}(p) = p^{\text{pos}}$, $\pi^{\text{neg}}(p) = p^{\text{neg}}$;
- $\pi^{\text{pos}}(\neg\psi) = \neg\pi^{\text{neg}}(\psi)$, $\pi^{\text{neg}}(\neg\psi) = \neg\pi^{\text{pos}}(\psi)$;
- $\pi^\alpha(\psi \otimes \chi) = \pi^\alpha(\psi) \otimes \pi^\alpha(\chi)$ for $\alpha \in \{\text{pos}, \text{neg}\}$ and $\otimes \in \{\wedge, \vee\}$.

Given an LP_m interpretation μ , we associate with it an ML-model $\mathcal{M}_\mu = \langle W, w_0, V, \lambda \rangle$, where $W = \{w_0, w_1\}$, $V(w_0) = \{p \in \text{Prop} \mid 1 \in \mu(p)\}$, $V(w_1) = \{p \in \text{Prop} \mid \mu(p) = \{1\}\}$, and $\lambda(\text{pos}) = w_0$, $\lambda(\text{neg}) = w_1$, and $\lambda(\beta) = w_0$ for each other marker β . Intuitively, w_1 is used

to assign value 0 to occurrences marked by neg . More precisely, if $\mu(p) = \{0\}$ then p is in neither $V(w_0)$ nor $V(w_1)$; if $\mu(p) = \{1\}$ then p is in both; and if $\mu(p) = \{0, 1\}$ then p is in $V(w_0)$ but not in $V(w_1)$.

The following lemma states that the f_3 -abnormal variables of $\mathcal{M}_\mu$ are exactly those variables p such that $\mu(p) = \{0, 1\}$.

Lemma 2. *For every LP_m interpretation μ , $f_3(\mathcal{M}_\mu) = \mu!$.*

Proof. Let us recall that $f_3(\mathcal{M})$ is the set of propositional letters p such that p takes different truth values at the normal world w_0 and at some world designated by a marker (in our setting, the only relevant such markers are pos and neg , with $\lambda(\text{pos}) = w_0$ and $\lambda(\text{neg}) = w_1$).

In $\mathcal{M}_\mu$: $p \in V(w_0)$ iff $1 \in \mu(p)$, $p \in V(w_1)$ iff $\mu(p) = \{1\}$.

We distinguish the three possible values of $\mu(p)$.

- If $\mu(p) = \{1\}$, then $p \in V(w_0)$ and $p \in V(w_1)$: p has the same (true) value at both worlds, so $p \notin f_3(\mathcal{M}_\mu)$.
- If $\mu(p) = \{0\}$, then $p \notin V(w_0)$ and $p \notin V(w_1)$: p has the same (false) value at both worlds, so $p \notin f_3(\mathcal{M}_\mu)$.
- If $\mu(p) = \{0, 1\}$, then $1 \in \mu(p)$, hence $p \in V(w_0)$, but $\mu(p) \neq \{1\}$, so $p \notin V(w_1)$. Thus p takes different truth values at w_0 and w_1 , so $p \in f_3(\mathcal{M}_\mu)$.

It follows that $f_3(\mathcal{M}_\mu)$ consists exactly of those p such that $\mu(p) = \{0, 1\}$. $\square$

The following lemma shows that, in the associated model $\mathcal{M}_\mu$, the translations π^{pos} and π^{neg} exactly reflect the two truth components of μ .

Lemma 3. *Let μ be an LP_m interpretation. Then, for every formula φ , the following holds:*

1. $1 \in \mu(\varphi)$ iff $\mathcal{M}_\mu \models_{\text{ML}} \pi^{\text{pos}}(\varphi)$.
2. $0 \in \mu(\varphi)$ iff $\mathcal{M}_\mu \not\models_{\text{ML}} \pi^{\text{neg}}(\varphi)$.

Using the two previous lemmas, we can show that whenever μ is a minimal LP_m model of Γ , the associated structure $\mathcal{M}_\mu$ is an f_3 -minimal ML-model of $\{\pi^{\text{pos}}(\varphi) \mid \varphi \in \Gamma\}$.

For the converse direction, let $\mathcal{M} = \langle W, w_0, V, \lambda \rangle$ be an ML-model such that $\lambda(\text{pos}) = w_0$, $\lambda(\beta) = w_0$ for every $\beta \neq \text{neg}$, and $V(\lambda(\text{neg})) \subseteq V(w_0)$. Let $\mu_{\mathcal{M}}$ be defined by $1 \in \mu_{\mathcal{M}}(p)$ iff $p \in V(\lambda(\text{pos}))$, $0 \in \mu_{\mathcal{M}}(p)$ iff $p \notin V(\lambda(\text{neg}))$. Under these restrictions, $f_3(\mathcal{M})$ coincides with $\mu_{\mathcal{M}}!$.

Lemma 4. *Let $\mathcal{M}$ be an ML-model satisfying the above compatibility condition, and let $\mu_{\mathcal{M}}$ be the associated LP_m interpretation. Then, for every formula φ :*

1. $1 \in \mu_{\mathcal{M}}(\varphi)$ iff $\mathcal{M} \models_{\text{ML}} \pi^{\text{pos}}(\varphi)$;
2. $0 \in \mu_{\mathcal{M}}(\varphi)$ iff $\mathcal{M} \not\models_{\text{ML}} \pi^{\text{neg}}(\varphi)$.

The previous lemma, together with the observation $f_3(\mathcal{M}) = \mu_{\mathcal{M}}!$ for LP_m -compatible ML-models, guarantees that if $\mathcal{M}$ is an f_3 -minimal LP_m -compatible ML-model of $\{\pi^{\text{pos}}(\varphi) \mid \varphi \in \Gamma\}$, then the associated interpretation $\mu_{\mathcal{M}}$ is a minimal LP_m model of Γ .

The following theorem establishes the correctness of our encoding of $\vdash_{LP_m}$ in our framework. The proof is obtained by exploiting the correspondence between minimal LP_m models and f_3 -minimal ML-models using the translations $\mathcal{M}_\mu$ and $\mu_{\mathcal{M}}$.

Theorem 2. For every set of formulas Γ and every formula φ , $\Gamma \vdash_{LP_m} \varphi$ iff $\pi^{\text{pos}}(\Gamma) \vdash_{f_3}^{\text{mk}} \pi^{\text{pos}}(\varphi)$, where $\pi^{\text{pos}}(\Gamma) = \{\pi^{\text{pos}}(\psi) \mid \psi \in \Gamma\}$.

6.3 Occurrence-based Semantics

We now show how to encode, within our framework, an entailment relation defined via an occurrence-based semantics [Salhi, 2025].

Intuitively, instead of assigning a single truth value to each variable, we assign truth values to occurrences of variables and then prefer interpretations that minimize the number of disagreements between occurrences of the same variable.

A syntactic occurrence of a variable p in a formula φ can be identified by a triple (p, i, φ) , where i is the position of that occurrence when counting from left to right. We write $\text{Occ}(\varphi)$ for the set of all such occurrences in φ , and for a set Γ of formulas we set $\text{Occ}(\Gamma) = \bigcup_{\psi \in \Gamma} \text{Occ}(\psi)$. Furthermore, let $\text{Occ}(\Gamma, p)$ be the set of occurrences of the variable p in Γ . For an occurrence $o \in \text{Occ}(\Gamma)$, we denote by $\text{var}(o)$ the propositional variable that occurs at o .

Let R be a function that assigns to each occurrence $o \in \text{Occ}(\Gamma)$ a distinct fresh propositional variable $R(o)$. Given a formula φ , we write $R(\varphi)$ for the formula obtained from φ by replacing, for each $o \in \text{Occ}(\varphi)$, that occurrence by $R(o)$.

An *occurrence-based interpretation* (or *o-interpretation*) of Γ is a function $\epsilon : \text{Occ}(\Gamma) \rightarrow \{0, 1\}$. We say that ϵ is an *o-model* of Γ iff there exists a classical interpretation ω_ϵ such that $\omega_\epsilon(R(o)) = \epsilon(o)$ for all $o \in \text{Occ}(\Gamma)$, and $\omega_\epsilon \models_{\text{CL}} R(\psi)$ for all $\psi \in \Gamma$. Thus, an o-model assigns truth values to occurrences in a way that can be realized by a classical model of the occurrence-renamed version of Γ .

To define the corresponding nonmonotonic entailment relation, we impose a minimality condition on o-models. For an o-interpretation ϵ of Γ , let $\text{diff}(\epsilon) = \{\{o, o'\} \subseteq \text{Occ}(\Gamma) \mid \text{var}(o) = \text{var}(o') \text{ and } \epsilon(o) \neq \epsilon(o')\}$. In words, $\text{diff}(\epsilon)$ records all pairs of occurrences of the same variable that receive different truth values under ϵ .

We define a preorder $\preceq$ on o-interpretations of Γ by $\epsilon \preceq \epsilon'$ iff $\text{diff}(\epsilon) \subseteq \text{diff}(\epsilon')$. The associated strict preorder is denoted by $\prec$, i.e. $\epsilon \prec \epsilon'$ iff $\epsilon \preceq \epsilon'$ and $\text{diff}(\epsilon) \neq \text{diff}(\epsilon')$.

An o-model ϵ of Γ is *minimal* if it is minimal with respect to $\preceq$: whenever ϵ' is an o-interpretation of Γ with $\epsilon' \prec \epsilon$, then ϵ' is not an o-model of Γ .

We say that a classical interpretation ω is *compatible* with an o-interpretation ϵ of Γ if, for every propositional variable p occurring in Γ , there exists an occurrence $o \in \text{Occ}(\Gamma)$ with $\text{var}(o) = p$ such that $\omega(p) = \epsilon(o)$. In other words, ω assigns to each variable a value that agrees with at least one of its occurrences under ϵ .

The occurrence-based entailment relation $\vdash_{\text{occ}}$ is then defined as follows:

Definition 6. For a set of propositional formulas Γ and a formula φ , we write $\Gamma \vdash_{\text{occ}} \varphi$ iff for every minimal o-model ϵ of Γ and every classical interpretation ω compatible with ϵ , we have $\omega \models_{\text{CL}} \varphi$.

To encode this entailment relation in our framework, we introduce two marking schemes.

Occurrence-marking. The translation τ_{occ} assigns a distinct marker α_o to each occurrence $o \in \text{Occ}(\Gamma)$ of a propositional variable in the formulas under consideration. We write $\tau_{\text{occ}}(\Gamma)$ for the set of marked formulas obtained from Γ by replacing each occurrence o of a variable p with p^{α_o} .

Variable-marking. The translation τ_{var} assigns a distinct marker α_p to each propositional variable p (all occurrences of p receive the same marker α_p). We write $\tau_{\text{var}}(\varphi)$ for the formula obtained from φ by replacing each occurrence of p with p^{α_p} . We assume that the markers used by τ_{occ} and those used by τ_{var} are pairwise distinct.

For convenience, we continue to write α_o for the marker associated with an occurrence o , and α_p for the marker associated with the variable p .

Theorem 3. For every finite set of formulas Γ and every formula φ , $\Gamma \vdash_{\text{occ}} \varphi$ iff $\tau_{\text{occ}}(\Gamma) \vdash_{f_2}^{\text{mk}} (\Theta \rightarrow \tau_{\text{var}}(\varphi))$, where $\Theta = \bigwedge_{p \in \text{vars}(\varphi)} \bigvee_{o \in \text{Occ}(\Gamma, p)} (p^{\alpha_p} \leftrightarrow p^{\alpha_o})$.

Intuitively, $\tau_{\text{occ}}(\Gamma)$ records truth values at the level of occurrences, while $\tau_{\text{var}}(\varphi)$ evaluates φ at the level of variables. The formula Θ ties each variable-marker α_p in φ to at least one occurrence marker α_o of p in Γ . Minimizing f_2 then mimics the minimization of $\text{diff}(\epsilon)$ in the original occurrence-based semantics: we prefer models that identify as many markers as possible, i.e., models that make as many occurrences of the same variable agree as the formulas in Γ allow.

It is worth noting that our encoding applies only to finite sets of formulas. This restriction comes from the auxiliary formula Θ , which must explicitly include all occurrences.

7 Inconsistency Measurement

Inconsistency measures are intended to quantify the degree of contradiction present in a knowledge base. They have been widely studied in the literature, often within postulate-based frameworks that capture different desirable properties of such measures (see, e.g., [Hunter and Konieczny, 2010; Thimm, 2018; Bona *et al.*, 2019; Besnard and Grant, 2020]). Inconsistency measures have been applied in a variety of areas, including paraconsistency [Salhi, 2021; Corea *et al.*, 2024], databases [Parisi and Grant, 2023; Grant and Parisi, 2024], spatio-temporal qualitative reasoning [Condotta *et al.*, 2016], business rule analysis [Corea *et al.*, 2021], and deontic reasoning [Arieli *et al.*, 2024].

A *knowledge base* (KB) is a finite set $K \subseteq \text{PF}$ of individually consistent formulas. We use KB to denote the set of KBs.

We impose the consistency of individual formulas mainly to ensure that formula-level marking, via $\tau_{\text{form}}(K)$, can restore consistency. Indeed, if a formula ψ is itself inconsistent, then $\tau_{\text{form}}(\psi, \alpha)$ is unsatisfiable at every world. Thus, there is no model in which ψ can be retained. However, this assumption is not essential: it can be removed by a simple preprocessing step that filters out inconsistent formulas. Moreover, for our other marking approaches, the results can be generalized to KBs that may contain inconsistent formulas.

Definition 7 (Inconsistency Measure). An inconsistency measure is a mapping $\mathcal{I} : \text{KB} \rightarrow \mathbb{R}_{\infty}^{\geq 0}$ such that, for every KB K ,

CONSISTENCY $\mathcal{I}(K) = 0$ iff K is consistent.

Beyond CONSISTENCY, a number of additional postulates have been proposed to further constrain inconsistency measures (e.g., see [Thimm, 2018]). We do not commit to a particular collection of such postulates here. Instead, we focus on showing how several well-known measures can be expressed within our logical framework by associating their values with the cardinalities of suitable minimal abnormality sets in marked models.

7.1 Examples of Inconsistency Measures

We briefly recall some measures from the literature. Let K be a KB.

- The *MSS-counting measure* [Grant and Hunter, 2011]: $\mathcal{I}_{\text{mss}}(K) = |\text{MSS}(K)| - 1$.
- A *problematic-formula measure* [Grant and Hunter, 2011]: $\mathcal{I}_{\text{P}}(K) = |\text{Problematic}(K)|$, where $\text{Problematic}(K) = \bigcup \text{MUS}(K)$.
- A *contension-based measure* [Konieczny et al., 2003; Grant and Hunter, 2011]: $\mathcal{I}_{\text{C}}(K) = \text{Contension}(K)$, where $\text{Contension}(K) = \min\{|\mu| \mid \mu \models_{LP_m} K\}$.
- A *formula-based hitting-set measure* [Grant and Hunter, 2017]: $\mathcal{I}_{\text{H}}(K) = \min\{|X| \mid X \subseteq K, \forall K' \in \text{MUS}(K), X \cap K' \neq \emptyset\}$.
- A *model-based hitting-set measure* [Thimm, 2016]: $\mathcal{I}_{\text{hs}}(K) = \min\{|H| \mid H \in \text{HS}(K)\} - 1$, where $\text{HS}(K)$ is the family of model-based hitting sets for K , i.e., $\text{HS}(K) = \{H \subseteq \Omega(\text{Prop}) \mid \forall \varphi \in K, \exists \omega \in H : \omega \models_{\text{CL}} \varphi\}$.
- The *Dalal-sum distance measure* [Grant and Hunter, 2017]: $\mathcal{I}_{\text{dalal}}^{\Sigma}(K) = \min\{\sum_{\varphi \in K} d(\varphi, \omega) \mid \omega \in \Omega(\text{Prop})\}$, where $d(\varphi, \omega)$ is Dalal's distance between φ and the interpretation ω .
- A *forgetting-based measure* [Besnard, 2016; Thimm, 2018]: $\mathcal{I}_{\text{forget}}(K) = \min\{k \mid \exists (p_j, i_j, c_j)_{j=1}^k (\bigwedge K)[p_1, i_1 \mapsto c_1; \dots; p_k, i_k \mapsto c_k] \not\models \perp, c_j \in \{\top, \perp\}\}$. Here, for a formula φ and a triple (p, i, c) with $p \in \text{Prop}$, $i \in \mathbb{N}$, and $c \in \{\top, \perp\}$, we write $\varphi[p, i \mapsto c]$ for the formula obtained from φ by replacing the i -th occurrence of p in φ (counted from left to right) by the constant c . For a finite sequence $(p_j, i_j, c_j)_{j=1}^k$, the notation $\varphi[p_1, i_1 \mapsto c_1; \dots; p_k, i_k \mapsto c_k]$ denotes the result of performing all these replacements.

7.2 Expressing Measures via Markers and Abnormalities

We now show how our logical framework with markers and abnormality functions can express all of the above inconsistency measures.

Formula-level marking. We first consider the marking τ_{form} defined in Section 6.1: each formula φ in the considered KB is assigned its own marker $\alpha_{\varphi} \neq \nu$, and $\tau_{\text{form}}(K)$ denotes the set of corresponding marked formulas.

Proposition 1. *Let K be a KB. Then:*

- $\mathcal{I}_{\text{hs}}(K) = \min\{|\text{Mk}/\sim_{\lambda_{\text{M}}}| \mid \mathcal{M} \in \text{Mod}(\tau_{\text{form}}(K))\} - 1$, where $\text{Mk}/\sim_{\lambda_{\text{M}}}$ is the set of equivalence classes of markers induced by λ_{M} ;
- $\mathcal{I}_{\text{mss}}(K) = |\{f_1(\mathcal{M}) \mid \mathcal{M} \in \text{Mod}_{f_1}^{\min}(\tau_{\text{form}}(K))\}| - 1$;
- $\mathcal{I}_{\text{P}}(K) = |\bigcup\{f_1(\mathcal{M}) \mid \mathcal{M} \in \text{Mod}_{f_1}^{\min}(\tau_{\text{form}}(K))\}|$;
- $\mathcal{I}_{\text{H}}(K) = \min\{|\{f_1(\mathcal{M})\}| \mid \mathcal{M} \in \text{Mod}(\tau_{\text{form}}(K))\}$;
- $\mathcal{I}_{\text{dalal}}^{\Sigma}(K) = \min\{|\{f_4(\mathcal{M})\}| \mid \mathcal{M} \in \text{Mod}(\tau_{\text{form}}(K))\}$.

As noted above, if inconsistent formulas are allowed in KBs, the preceding results can be easily adapted. For instance, in this setting we obtain $\mathcal{I}_{\text{mss}}(K) = |\{f_1(\mathcal{M}) \mid \mathcal{M} \in \text{Mod}_{f_1}^{\min}(\tau_{\text{form}}(K \setminus \text{Inc}(K)))\}| + |\text{Inc}(K)| - 1$, where $\text{Inc}(K)$ denotes the set of inconsistent formulas in K .

Polarity-level marking. We now consider the marking π^{pos} introduced in Section 6.2, which uses two markers pos and neg to encode the Logic of Paradox.

Proposition 2. *For every KB K , $\mathcal{I}_{\text{C}}(K) = \min\{|\{f_3(\mathcal{M})\}| \mid \mathcal{M} \in \text{Mod}(\pi^{\text{pos}}(K))\}$.*

Here, $\mathcal{I}_{\text{C}}$ counts the minimal number of variables that must take two truth values in order to satisfy K , and f_3 captures exactly such deviations in our framework.

Occurrence-level marking. Finally, we consider the marking τ_{occ} defined in Section 6.3, which assigns a distinct marker to each occurrence of a propositional variable in the formulas of K .

Proposition 3. *For every KB K , $\mathcal{I}_{\text{forget}}(K) = \min\{|\{f_4(\mathcal{M})\}| \mid \mathcal{M} \in \text{Mod}(\tau_{\text{occ}}(K))\}$.*

Intuitively, $\tau_{\text{occ}}(K)$ makes it possible to forget individual occurrences of variables by allowing them to deviate from their normal values, and f_4 counts how many such deviations are needed to restore consistency. This mirrors the definition of $\mathcal{I}_{\text{forget}}$.

8 Conclusion and Perspectives

We have introduced a logical framework that separates the surface propositional language from a marked layer, where occurrences of atoms are evaluated at possibly different worlds. Abnormality functions assign to each model a set of defects, which induces a preference ordering over models. This provides a uniform semantic basis for inconsistency-tolerant entailment and inconsistency measurement: by varying the marking scheme and the notion of abnormality, we can recover MSS-based entailment, the minimally inconsistent Logic of Paradox, an occurrence-based entailment relation, and a wide range of inconsistency measures.

There are several natural directions for future work. First, it would be useful to develop proof systems and algorithms for our framework. Second, we could extend the approach to richer logics (e.g., first-order or description logics) and to settings with time or multiple sources. Third, viewing inconsistency measures and entailment through the common lens of markers and abnormalities may support a systematic, postulate-based analysis of existing measures and guide the design of new ones.

References

- [Arieli *et al.*, 2024] Ofer Arieli, Kees van Berkel, Badran Raddaoui, and Christian Straßer. Deontic reasoning based on inconsistency measures. In *Proceedings of the 21st International Conference on Principles of Knowledge Representation and Reasoning, KR 2024, Hanoi, Vietnam. 2024*, 2024.
- [Besnard and Grant, 2020] Philippe Besnard and John Grant. Relative Inconsistency Measures. *Artificial Intelligence*, 280:103231, 2020.
- [Besnard and Hunter, 2008] Philippe Besnard and Anthony Hunter. *Elements of Argumentation*. MIT Press, 2008.
- [Besnard, 2016] Philippe Besnard. Forgetting-Based Inconsistency Measure. In Steven Schockaert and Pierre Senellart, editors, *Scalable Uncertainty Management - 10th International Conference, SUM 2016, Nice, France*, volume 9858 of *Lecture Notes in Computer Science*, pages 331–337. Springer, 2016.
- [Blackburn, 2000] P Blackburn. Representation, reasoning, and relational structures: a hybrid logic manifesto. *Logic Journal of the IGPL*, 8(3):339–365, 2000.
- [Bona *et al.*, 2019] Glauber De Bona, John Grant, Anthony Hunter, and Sébastien Konieczny. Classifying Inconsistency Measures Using Graphs. *Journal of Artificial Intelligence Research*, 66:937–987, 2019.
- [Condotta *et al.*, 2016] Jean-François Condotta, Badran Raddaoui, and Yakoub Salhi. Quantifying conflicts for spatial and temporal information. In *Principles of Knowledge Representation and Reasoning: Proceedings of the Fifteenth International Conference, KR 2016, Cape Town, South Africa, 2016*, pages 443–452. AAAI Press, 2016.
- [Corea *et al.*, 2021] Carl Corea, Matthias Thimm, and Patrick Delfmann. Measuring inconsistency over sequences of business rule cases. In *Proceedings of the 18th International Conference on Principles of Knowledge Representation and Reasoning, KR 2021, Online event, 2021*, pages 655–660, 2021.
- [Corea *et al.*, 2024] Carl Corea, Isabelle Kuhlmann, Matthias Thimm, and John Grant. Paraconsistent reasoning for inconsistency measurement in declarative process specifications. *Information Systems*, 122:102347, 2024.
- [Fitch, 1966] Frederic Fitch. Tree proofs in modal logic. *The Journal of Symbolic Logic*, 31:152, 1966.
- [Gabbay, 1996] Dov M Gabbay. *Labelled Deductive Systems*. Oxford University Press, 1996.
- [Ghidini and Giunchiglia, 2001] Chiara Ghidini and Fausto Giunchiglia. Local models semantics, or contextual reasoning=locality+compatibility. *Artificial Intelligence*, 127(2):221–259, 2001.
- [Gómez Álvarez and Rudolph, 2021] Lucía Gómez Álvarez and Sebastian Rudolph. Standpoint Logic: Multi-Perspective Knowledge Representation. In *Formal Ontology in Information Systems - Proceedings of the Twelfth International Conference, FOIS 2021, Bozen-Bolzano, Italy, 2021*, *Frontiers in Artificial Intelligence and Applications*, pages 3–17. IOS Press, 2021.
- [Grant and Hunter, 2011] John Grant and Anthony Hunter. Measuring the good and the bad in inconsistent information. In *Proceedings of the 22nd International Joint Conference on Artificial Intelligence - Volume Volume Three, IJCAI’11*, pages 2632–2637. AAAI Press, 2011.
- [Grant and Hunter, 2017] John Grant and Anthony Hunter. Analysing inconsistent information using distance-based measures. *International Journal of Approximate Reasoning*, 89:3–26, 2017. Special Issue on Theories of Inconsistency Measures and Their Applications.
- [Grant and Parisi, 2024] John Grant and Francesco Parisi. On measuring inconsistency in graph databases with regular path constraints. *Artificial Intelligence*, 335:104197, 2024.
- [Hansson, 2022] Sven Ove Hansson. Logic of Belief Revision. In Edward N. Zalta, editor, *The Stanford Encyclopedia of Philosophy*. Metaphysics Research Lab, Stanford University, Spring 2022 edition, 2022.
- [Hunter and Konieczny, 2010] Anthony Hunter and Sébastien Konieczny. On the measure of conflicts: Shapley Inconsistency Values. *Artificial Intelligence*, 174:1007–1026, 2010.
- [Klarman and Gutiérrez-Basulto, 2011] Szymon Klarman and Víctor Gutiérrez-Basulto. Two-Dimensional Description Logics for Context-Based Semantic Interoperability. In *Proceedings of the 25th AAAI Conference on Artificial Intelligence, AAAI 2011, San Francisco, California, USA, 2011*, pages 215–220. AAAI Press, 2011.
- [Konieczny *et al.*, 2003] Sebastien Konieczny, Jerome Lang, and Pierre Marquis. Quantifying information and contradiction in propositional logic through test actions. In *Proceedings of the 18th International Joint Conference on Artificial Intelligence, IJCAI’03*, pages 106–111, San Francisco, CA, USA, 2003. Morgan Kaufmann Publishers Inc.
- [Kripke, 1959] Saul A. Kripke. A Completeness Theorem in Modal Logic. *The Journal of Symbolic Logic*, 24(1):1–14, 1959.
- [McCarthy, 1986] John McCarthy. Applications of Circumscription to Formalizing Common-Sense Knowledge. *Artificial Intelligence*, 28(1):89–116, 1986.
- [Parisi and Grant, 2023] Francesco Parisi and John Grant. On measuring inconsistency in definite and indefinite databases with denial constraints. *Artificial Intelligence*, 318:103884, 2023.
- [Priest *et al.*, 2022] Graham Priest, Koji Tanaka, and Zach Weber. Paraconsistent Logic. In Edward N. Zalta, editor, *The Stanford Encyclopedia of Philosophy*. Metaphysics Research Lab, Stanford University, Spring 2022 edition, 2022.
- [Priest, 1991] G. Priest. Minimally inconsistent LP. *Studia Logica*, 50:321–331, 1991.

- [Rescher and Manor, 1970] Nicholas Rescher and Ruth Manor. On Inference from Inconsistent Premisses. *Theory and Decision*, 1:179–217, 1970.
- [Salhi, 2021] Yakoub Salhi. Inconsistency Measurement for Paraconsistent Inference. In *Proceedings of the Thirtieth International Joint Conference on Artificial Intelligence, IJCAI 2021, Virtual Event / Montreal, Canada*, pages 2033–2039. ijcai.org, 2021.
- [Salhi, 2025] Yakoub Salhi. A Variable Occurrence-Centric Framework for Inconsistency Handling. In *AAAI-25, Sponsored by the Association for the Advancement of Artificial Intelligence, Philadelphia, PA, USA*, pages 15143–15150. AAAI Press, 2025.
- [Serafini and Bouquet, 2004] Luciano Serafini and Paolo Bouquet. Comparing formal theories of context in AI. *Artificial Intelligence*, 155(1-2):41–67, 2004.
- [Serafini and Homola, 2012] Luciano Serafini and Martin Homola. Contextualized knowledge repositories for the Semantic Web. *Journal of Web Semantics*, 12:64–87, 2012.
- [Thimm, 2016] Matthias Thimm. Stream-based inconsistency measurement. *International Journal Approximate Reasoning*, 68:68–87, 2016.
- [Thimm, 2018] Matthias Thimm. On the Evaluation of Inconsistency Measures. In John Grant and Maria Vanina Martinez, editors, *Measuring Inconsistency in Information*, volume 73 of *Studies in Logic*. College Publications, 2018.
- [Viganò, 2000] Luca Viganò. *Labelled Non-Classical Logics*. Kluwer Academic Publishers, Boston, 2000.