

KLM-Style Defeasible Reasoning

Implementation and Evaluation of Entailment and Justification Algorithms

Chipo Hamayobe

chipo@cs.uct.ac.za

Artificial Intelligence Research Unit
Department of Computer Science
University of Cape Town

Abstract

In the 1950s, McCarthy proposed using formal logic in AI, establishing the foundations of **knowledge representation and reasoning**. Early systems, based on monotonic logic, struggled with exceptions since inferences could not be revised once drawn. To address this, researchers developed **non-monotonic** reasoning, enabling defeasible “common sense” inference. The **Kraus–Lehmann–Magidor (KLM) framework** emerged as a leading approach, with both syntactic and semantic formulations. Our work focuses on syntactic entailment algorithms—**Rational Closure**, **Lexicographic Closure**, and **Relevant Closure**—and develops justification methods to enhance interpretability. A web-based reasoning tool integrates computation, visualisation, and evaluation, advancing explainability and efficiency in symbolic reasoning.

1. Introduction and Background

Propositional Logic formalises reasoning by combining simple statements with defined operators, where truth depends only on base statements and operator use. It underpins formal reasoning by enabling analysis independent of natural language.

Valuation
A valuation $\mathbf{v}$ is a function that maps each atom in $\mathcal{P}$ to either <code>true</code> (T) or <code>false</code> (F).
Satisfiability
A valuation $\mathbf{v}$ satisfies an atom $p \in \mathcal{P} \iff \mathbf{v}(p) = T$, denoted by $\mathbf{v} \models p$.
Knowledge Base
A knowledge base, $\mathcal{K} \subset \mathcal{L}$, is a set of propositional formulas and is considered finite in this work.
Classical Entailment
Given a knowledge base $\mathcal{K}$ and a formula α , it is the case that α is entailed by $\mathcal{K}$, written $\mathcal{K} \models \alpha$ and read as “ $\mathcal{K}$ entails α ”, $\iff$ for every $\mathbf{v} \in \mathcal{V}$ such that $\mathbf{v} \models \mathcal{K}$ it is the case that $\mathbf{v} \models \alpha$.

Propositional logic diverges from human reasoning: it cannot represent **typicality** with exceptions and, due to its **monotonicity**, prevents retracting conclusions when new, conflicting information arises.

2. Defeasible Reasoning and Explanation

To address the **monotonicity** limitation, the **KLM framework** formalises **non-monotonic** reasoning about typicality through a **preferential consequence relation**, supporting conclusions like “*an A is typically a B*”. Based on **KLM postulates**, relations satisfying postulates 1–7 are **preferential**, and those also meeting the 8th are **rational**. Interpreted via **preferential** or **ranked** semantics, valuations are ordered by “typicality”, and a defeasible implication $\alpha \sim \beta$ holds when all preferred worlds satisfying α also satisfy β .

KLM Postulates	
1. Reflexivity: $\frac{\alpha \rightarrow \alpha}{}$	2. Left Logical Equivalence: $\frac{\alpha \equiv \beta, \alpha \sim \gamma}{\beta \sim \gamma}$
3. Right Weakening: $\frac{\alpha \rightarrow \beta, \gamma \sim \alpha}{\gamma \sim \beta}$	4. And: $\frac{\alpha \sim \beta, \alpha \sim \gamma}{\alpha \sim \beta \wedge \gamma}$
5. Or: $\frac{\alpha \sim \beta, \gamma \sim \beta}{\alpha \vee \gamma \sim \beta}$	6. Cut: $\frac{\alpha \wedge \gamma \sim \beta, \alpha \sim \gamma}{\alpha \sim \beta}$
7. Cautious Monotonicity: $\frac{\alpha \sim \gamma, \alpha \sim \beta}{\alpha \wedge \beta \sim \gamma}$	8. Rational Monotonicity: $\frac{\alpha \sim \beta, \alpha \not\sim \neg \gamma}{\alpha \wedge \gamma \sim \beta}$
Exceptionality	
Given a knowledge base $\mathcal{K}$, a propositional formula $\alpha \in \mathcal{L}$ is considered exceptional for $\mathcal{K}$ if and only if $\mathcal{K} \not\models_p \top \sim \neg \alpha$.	
Minimal Ranked Entailment	
Given a defeasible knowledge base $\mathcal{K}$, the minimal ranked interpretation satisfying $\mathcal{K}$, $\mathcal{R}_{RC}^{\mathcal{K}}$, defines an entailment relation $\models$, known as the minimal ranked entailment, such that for any defeasible implication $\alpha \sim \beta$, $\mathcal{K} \models \alpha \sim \beta$ if and only if $\mathcal{R}_{RC}^{\mathcal{K}} \models \alpha \sim \beta$.	
Justification of Entailments	
Given a knowledge base $\mathcal{K}$ and some classical propositional formula $\alpha \rightarrow \beta$ such that $\mathcal{K} \models \alpha \rightarrow \beta$. Then the set $\mathcal{J}$ is said to be a justification for $\alpha \rightarrow \beta$ in $\mathcal{K}$ if $\mathcal{J} \subseteq \mathcal{K}$, $\mathcal{J} \models \alpha \rightarrow \beta$ and for every $\mathcal{J}' \subset \mathcal{J}$, $\mathcal{J}' \not\models \alpha \rightarrow \beta$. The set $\mathcal{J}$ will sometimes be denoted as $\mathcal{J}_{(\alpha \rightarrow \beta)}^{\mathcal{K}}$ or $\mathcal{J}(\alpha \rightarrow \beta, \mathcal{K})$.	

2.1. Base Rank

The exceptionality sequence is an iterative sequence of knowledge bases, $\mathcal{E}_0^{\mathcal{K}}, \mathcal{E}_1^{\mathcal{K}} \dots \mathcal{E}_n^{\mathcal{K}}$ where $\mathcal{E}_0^{\mathcal{K}} = \mathcal{K}$ and $\mathcal{E}_{i+1}^{\mathcal{K}} = \varepsilon(\mathcal{E}_i^{\mathcal{K}})$ for $1 \leq i \leq n$. Since $\mathcal{K}$ is finite, some n must always exist, and with the sequence $\mathcal{E}_0^{\mathcal{K}}, \mathcal{E}_1^{\mathcal{K}}, \dots, \mathcal{E}_{\infty}^{\mathcal{K}}$ defined, we can then formally define the *base rank* of a formula.

Base Rank
With respect to a given defeasible knowledge base $\mathcal{K}$, $br_{\mathcal{K}}(\alpha)$, known as the base rank of a formula $\alpha \in \mathcal{L}$, is the smallest index r such that α is not exceptional in $\mathcal{E}_r^{\mathcal{K}}$, and is defined: $br_{\mathcal{K}}(\alpha) := \min\{r \mid \mathcal{E}_r^{\mathcal{K}} \not\models_{\mathcal{R}} \top \sim \neg \alpha\}$.

Given the following knowledge base and the query $\text{pets} \sim \text{legs}$:

$$\mathcal{K} = \{ \text{animals} \sim \text{legs}, \text{animals} \sim \text{wild}, \text{pets} \rightarrow \text{animals}, \text{pets} \sim \neg \text{wild} \}$$

$\mathcal{R}_{\infty}$	$\text{pets} \rightarrow \text{animals}$
$\mathcal{R}_1$	$\text{pets} \sim \neg \text{wild}$
$\mathcal{R}_0$	$\text{animals} \sim \text{legs}, \text{animals} \sim \text{wild}$



2.2. Rational Closure

Rational closure applies prototypical reasoning, where typical instances inherit default properties, but atypical cases, such as pets are animals, do not.

Rational Closure
Given a knowledge base $\mathcal{K}$, a defeasible implication $\alpha \sim \beta$ is said to be in the rational closure of $\mathcal{K}$ (written $\mathcal{K} \models_{RC} \alpha \sim \beta$), if and only if $br_{\mathcal{K}}(\alpha) < br_{\mathcal{K}}(\alpha \wedge \neg \beta)$ or $br_{\mathcal{K}}(\alpha) = \infty$.

• *Discarded ranks:* $\mathcal{R}_0 = \{ \text{animals} \sim \text{legs}, \text{animals} \sim \text{wild} \}$

• *Entailment:* $\mathcal{K} \not\models_{RC} \text{pets} \sim \text{legs}$

2.3. Lexicographic Closure

Lexicographic closure, a more adventurous form of defeasible entailment, permits atypical cases like platypuses to inherit most typical mammalian properties, unlike rational closure.

Lexicographic Closure
Given a knowledge base $\mathcal{K}$, and a defeasible implication $\alpha \sim \beta$, it is the case that $\alpha \sim \beta$ is in the <i>Lexicographic Closure</i> of $\mathcal{K}$, denoted $\mathcal{K} \models_{LC} \alpha \sim \beta$, if and only if for any basis $\mathcal{K}^{\mathcal{B}} \subseteq \mathcal{K}$ of α , $\mathcal{K}^{\mathcal{B}} \cup \{ \alpha \} \models \beta$.

• *Discarded statements:* $\{ \text{animals} \sim \text{wild} \}$

• *Entailment:* $\mathcal{K} \models_{LC} \text{pets} \sim \text{legs}$

• *Deciding:* $\mathcal{D} = \{ \text{pets} \rightarrow \text{animals}, \text{pets} \sim \neg \text{wild}, \text{animals} \sim \text{legs} \}$

• *Justifications:* $\mathcal{J}_1 = \{ \text{pets} \rightarrow \text{animals}, \text{animals} \sim \text{legs} \}$

2.4. Relevant Closure

To define how to compute the **relevant** partition $\langle \mathcal{R}, \mathcal{R}^- \rangle$ for a query $\alpha \sim \beta$: the antecedent α determines the **relevant** partition. Based on the reasoning behind **Relevant Closure**, $\mathcal{R}$ should include exactly those statements used to derive $\neg \alpha$.

Basic Relevant Closure
For a statement α and knowledge base $\mathcal{K}$, let $\mathcal{J}_{bas}^{\mathcal{K}}(\alpha) = \{ \mathcal{J} \mid \mathcal{J} \text{ is an } \varepsilon\text{-justification w.r.t. } \mathcal{K} \}$. Then $\alpha \sim \beta$ is said to be in the Basic Relevant Closure of $\mathcal{K}$ if it is in the Relevant Closure of $\mathcal{K}$ w.r.t. $\bigcup \mathcal{J}_{bas} \subseteq \mathcal{K}$.

Minimal Relevant Closure
For some set of justifications $\mathcal{J} \subseteq \mathcal{K}$, let $\mathcal{J}_{min}^{\mathcal{K}} = \{ \alpha \sim \beta \mid r_{\mathcal{K}}(\alpha) \leq r_{\mathcal{K}}(\gamma) \text{ for every } \gamma \sim \sigma \in \mathcal{J} \}$. For a statement α , let $\mathcal{J}_{min}^{\mathcal{K}}(\alpha) = \bigcup_{\mathcal{J} \in \mathcal{J}_{min}^{\mathcal{K}}} \mathcal{J}_{min}^{\mathcal{K}}(\alpha)$. Then $\alpha \sim \beta$ is said to be in the Minimal Relevant Closure of $\mathcal{K}$ if it is in the Relevant Closure of $\mathcal{K}$ w.r.t. $\bigcup \mathcal{J}_{min}^{\mathcal{K}}(\alpha)$.

• *Relevance:* $\mathcal{R}^{bas} = \{ \text{animals} \sim \text{wild}, \text{pets} \sim \neg \text{wild} \}$, $\mathcal{R}^{min} = \{ \text{animals} \sim \text{wild} \}$

• *Discarded statements:* $\{ \text{animals} \sim \text{wild} \}$

• *Entailment:* $\mathcal{K} \models_{LC} \text{pets} \sim \text{legs}$

• *Deciding:* $\mathcal{D} = \{ \text{pets} \rightarrow \text{animals}, \text{pets} \sim \neg \text{wild}, \text{animals} \sim \text{legs} \}$

• *Justifications:* $\mathcal{J}_1 = \{ \text{pets} \rightarrow \text{animals}, \text{animals} \sim \text{legs} \}$

3. Implementation of Algorithms

The **klm-algorithms** system employs a *modular, layered architecture* that promotes clarity, scalability, and maintainability. It features a web-based interface linked to a *backend* responsible for logical processing and algorithm execution, integrating external libraries to enhance reasoning and explanation capabilities.

By maintaining a clear separation between the interface, controllers, and domain logic, the system efficiently processes user input, executes defeasible reasoning algorithms, and generates coherent explanatory feedback.

klm-algorithms is available at: <https://klm-algorithms.fly.dev>

An example of the implemented algorithms is *DefeasibleJustification*, which computes minimal justification sets for a given entailment **deciding knowledge base** $\mathcal{D}$ for a defeasible query $\alpha \sim \beta$:

Algorithm .1: UniversalDefeasibleJustification
Input: A defeasible knowledge base $\mathcal{K}$ and a defeasible query $\alpha \sim \beta$ Output: The set of all justification sets $\{ \mathcal{J}_1, \mathcal{J}_2, \mathcal{J}_3, \dots \}$
1 <i>(holds, D)</i> := KLMDefeasibleEntailment($\mathcal{K}, \alpha \sim \beta$);
2 <i>if holds</i> = <i>false</i> then
3 return $\emptyset$;
4 end
5 return ComputeAllJustifications($\mathcal{D}, \alpha \sim \beta$);

4. Evaluation of Algorithms

This study evaluates the logical soundness and computational feasibility of the algorithms to clarify their role in defeasible reasoning, particularly amid their growing relevance in artificial intelligence. The analysis focuses on key performance metrics such as efficiency and entailment accuracy, which determine practical applicability.

We evaluate algorithm performance on knowledge bases that vary in statement count, rank depth, and distribution pattern. While comparative studies indicate that these algorithms often mirror intuitive reasoning, they continue to face scalability challenges in large knowledge bases, underscoring the need for future work on interpretable and scalable solutions.