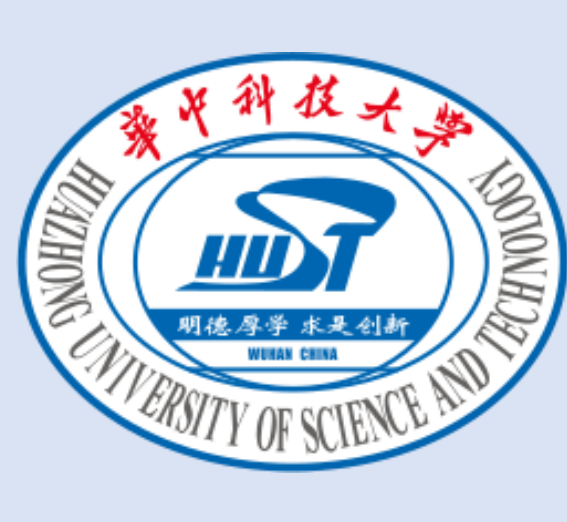




Association for the
Advancement of
Artificial Intelligence

Not Just What's There: Enabling CLIP to Comprehend Negated Visual Descriptions Without Fine-tuning

Junhao Xiao, Zhiyu Wu, Hao Lin, Yi Chen, Yahui Liu, Xiaoran Zhao, Zixu Wang, Zejiang He

CLIP's Big Blind Spot: The Problem with "Not"

How a pair of smart "glasses" can teach AI to finally **understand negation**—without the painful surgery.



Why Top AI Models Can Be Hopelessly Literal

Key Insight:

State-of-the-art Vision-Language Models like CLIP are brilliant at matching concepts, but they struggle with one simple word: **'not.'**

The Problem Explained



Root Cause

This isn't just a simple bug. It's a gap in education. In CLIP's vast training data, sentences containing negation are incredibly rare—making up **0.7%** of the text. The model simply never had enough examples to learn what "not" truly means.

The Old Fix: A Cure That's Worse Than the Disease

Fine-Tuning



The Conventional Approach: Fine-Tuning

- The standard method is to re-train the model's text encoder on a new dataset full of negative examples.

The Critical Drawbacks

- Catastrophic Forgetting:** This 'surgery' overwrites the model's original knowledge. While it might learn to understand 'not,' it forgets its broad, general-purpose abilities. Its zero-shot performance on other tasks plummets.
- Overfitting and Data Dependency:** The model becomes highly specialized to the negation dataset it was trained on, losing its ability to generalize to new, unseen domains. Performance drops severely under low-resource conditions.

Introducing CLIPGLASSES: A Plug-and-Play Upgrade

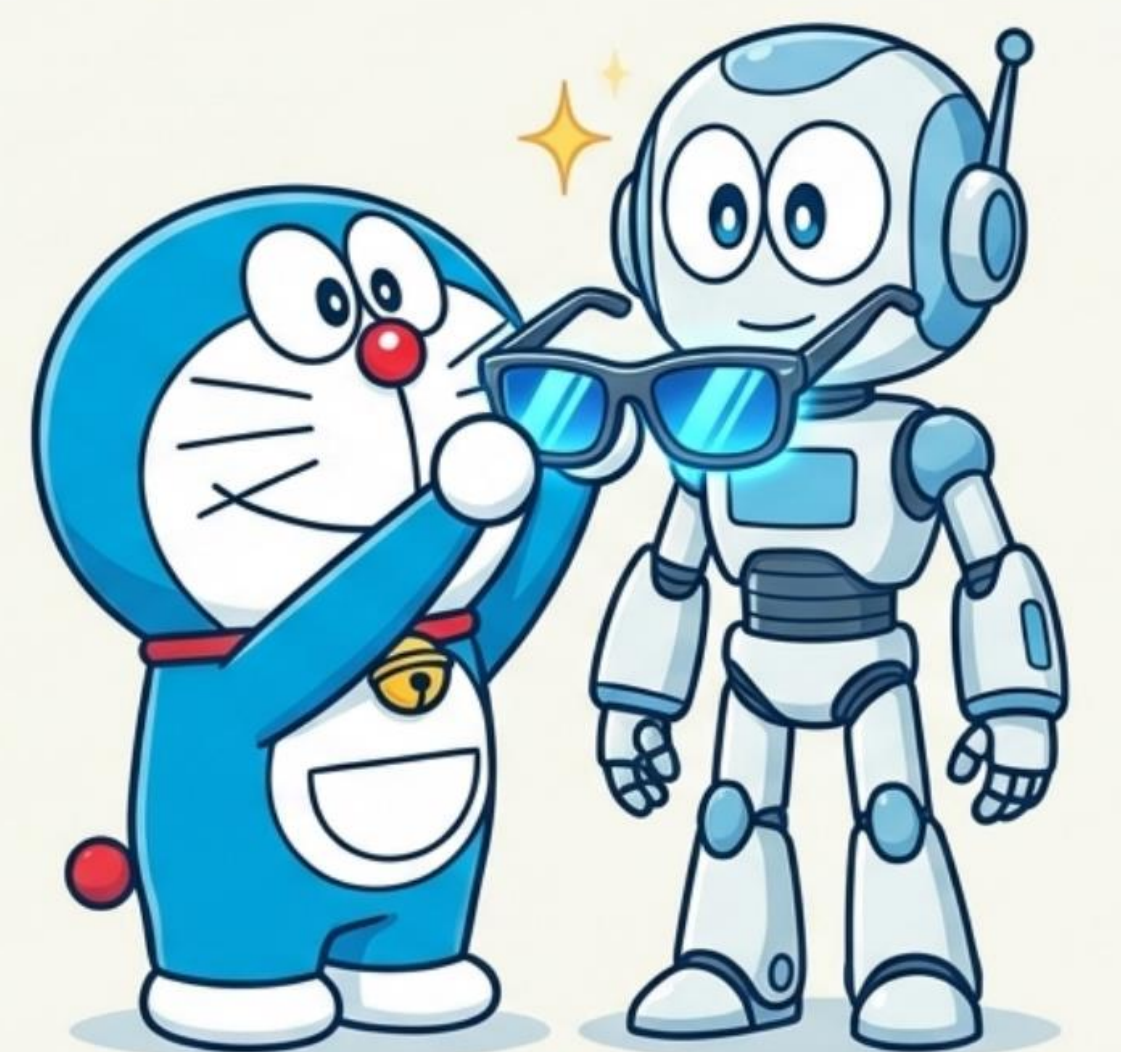
A Fundamentally Better Approach

- CLIPGLASSES is a **non-intrusive framework** that enhances CLIP without altering its core parameters.

The "Magic Gadget" Advantage

- No Fine-Tuning:** "It doesn't modify the model's 'brain.' Instead, it adds a new layer of perception, like putting on glasses."
- Preserves Core Skills:** "Because CLIP's original weights are frozen, it completely avoids catastrophic forgetting. The model retains its powerful zero-shot capabilities."
- Plug-and-Play:** "It's a modular extension that works with the existing CLIP model, offering a targeted solution to a specific problem."

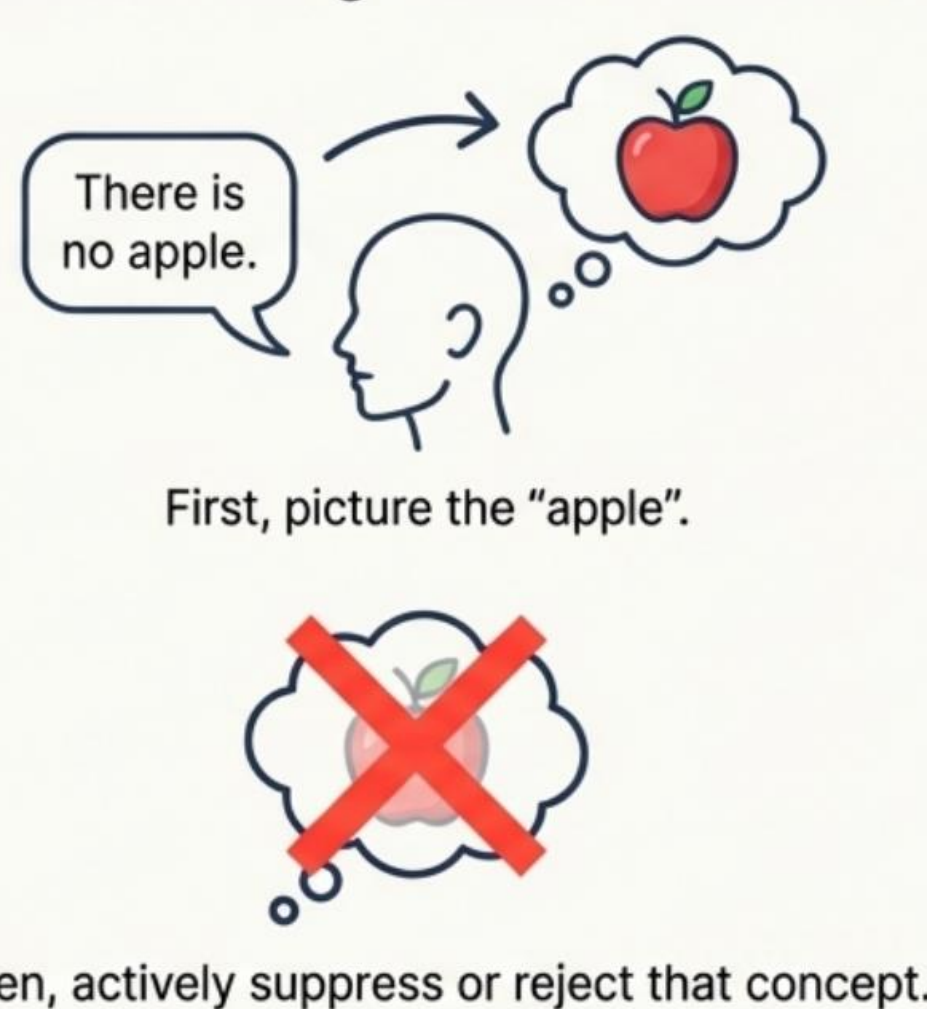
CLIPGLASSES



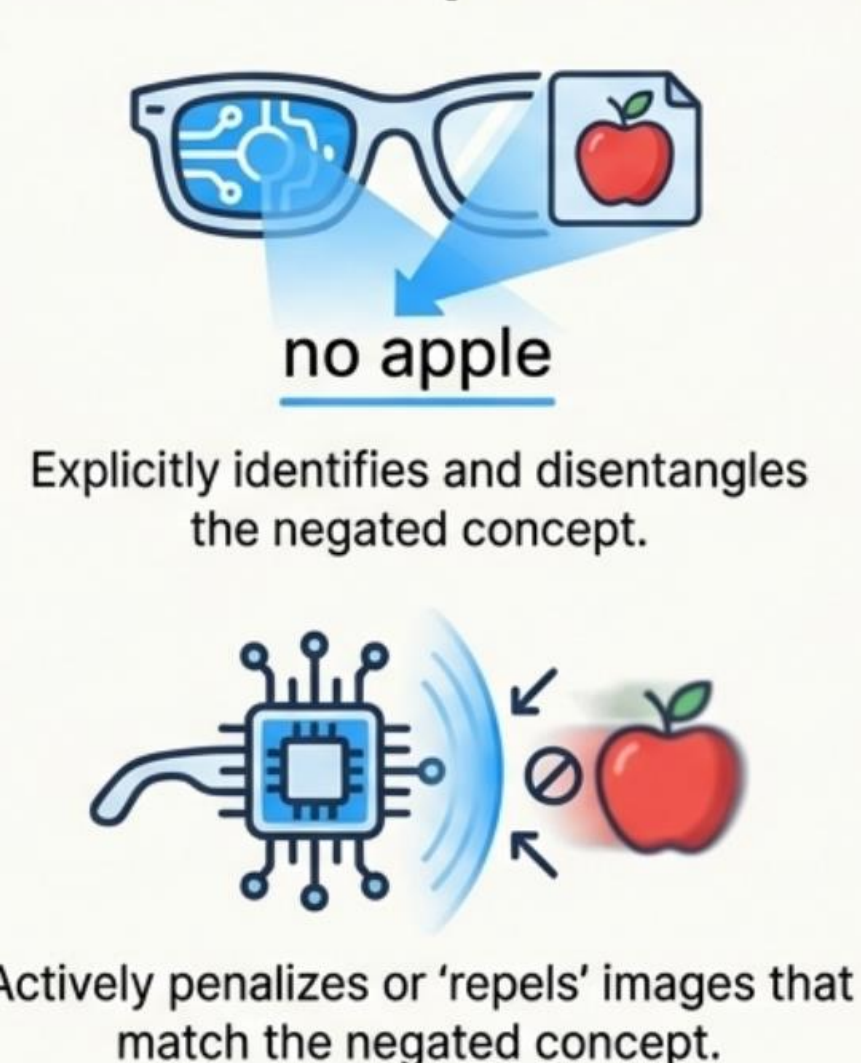
Inspired by How We Understand "Not"

Key Insight: Cognitive science shows that humans process negation in two distinct stages. CLIPGLASSES is designed to mimic this efficient process.

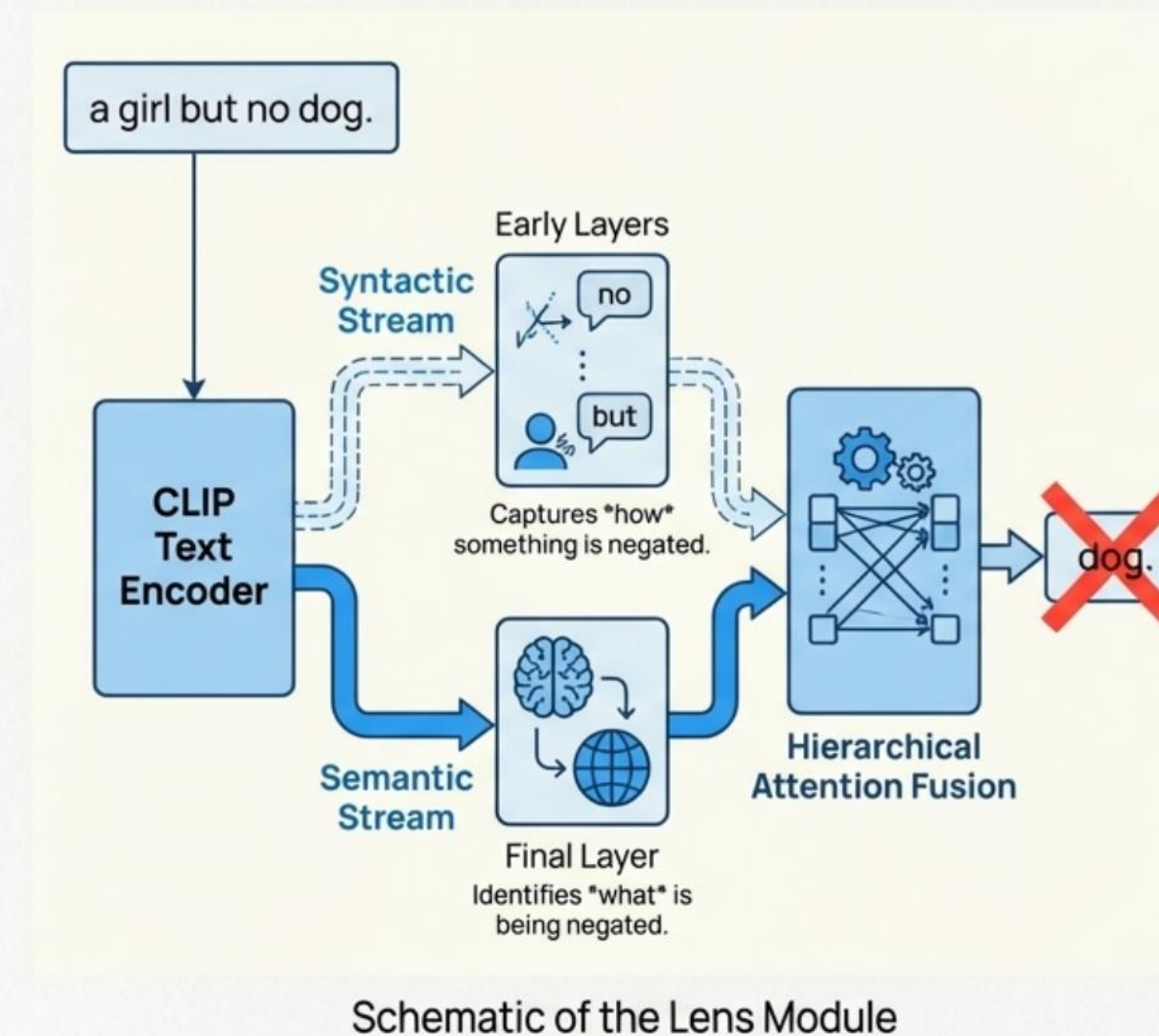
The Two-Stage Human Process



CLIPGLASSES' Computational Parallel



Inside the Tech: The 'Lens' Identifies the Target



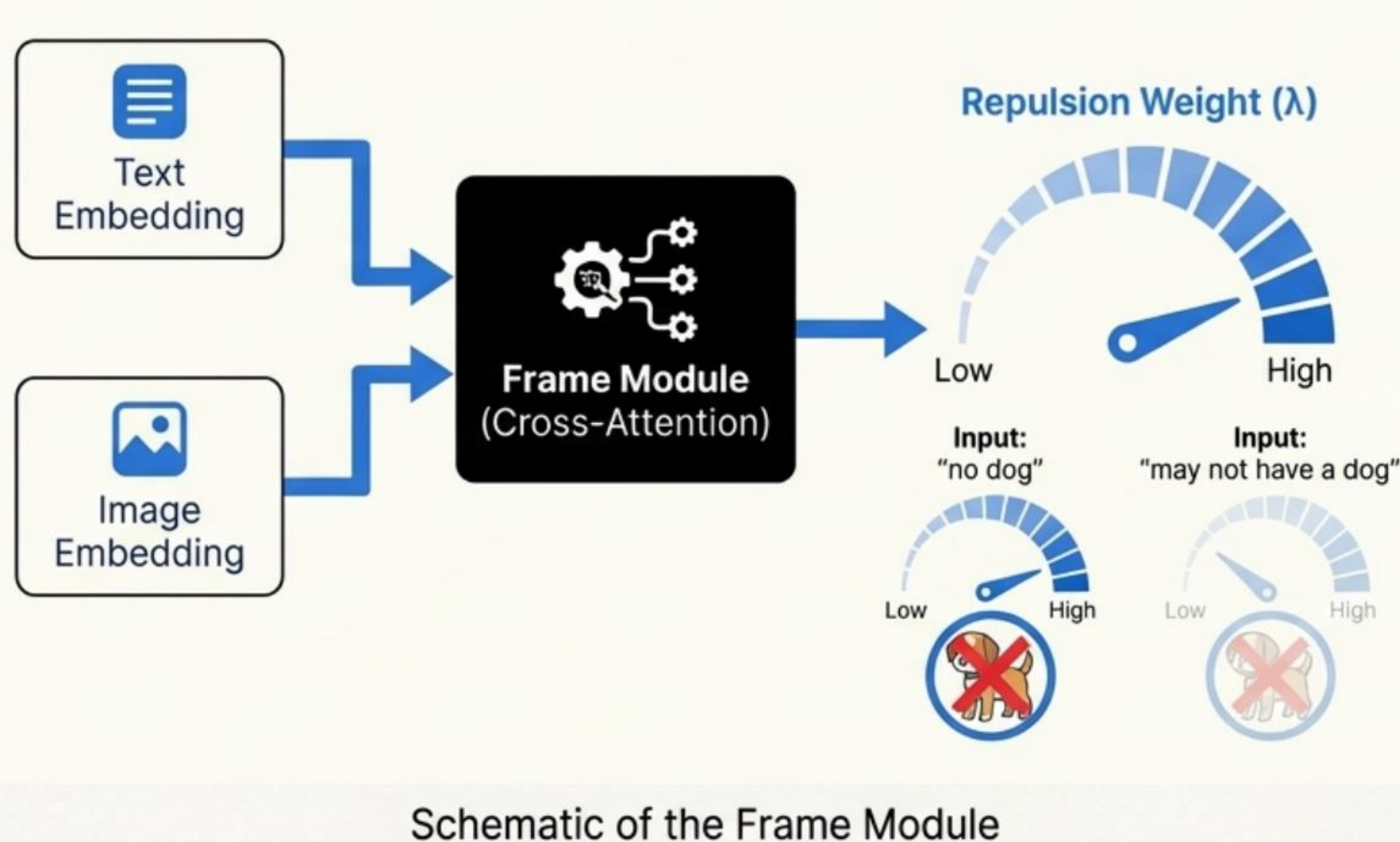
Function:

The Lens module's job is to precisely disentangle the negated semantics from the text embedding. In "a girl but no dog," it must isolate "dog."

How It Works: A Syntax-Semantic Dual-Stream Architecture

- Syntactic Stream:** Analyzes grammar and sentence structure. It uses early layers of CLIP's text encoder to spot grammatical cues for negation (e.g., words like "no," "without," "never").
- Semantic Stream:** Analyzes overall meaning and context. It uses the final layer of CLIP to understand the sentence's global meaning.
- Hierarchical Attention Fusion:** The module intelligently combines these two streams to ensure it never misidentifies the target of the negation.

Inside the Tech: The 'Frame' Calculates the Repulsion Force



Function:

The Frame module determines *how strongly* to penalize an image for containing the negated object. It predicts a context-aware repulsion weight (λ).

How It Works: Cross-Modal Context

- Negation can be nuanced. "No dog" is stronger than "may not have a dog."
- The Frame module looks at **both the text and the image** using a cross-attention mechanism.
- This allows it to generate a **dynamic repulsion weight (λ)**. For a sentence with a strong, unambiguous negation, it will generate a high weight, leading to a large penalty. For a weaker negation, the weight and penalty will be smaller.

The New Math: Subtracting What You Don't Want

The Old Calculation

Score = Similarity(Image, "a girl with no dog")



Result: High score if a dog is present.

The CLIPGLASSES Calculation

Final Score = "Likes" Score - "Hates" Score

$$S_{final} = S_{base} - M \cdot (\lambda \cdot R_{neg})$$

S_{base} : The original similarity score from CLIP. (The "Likes")

R_{neg} : The similarity score between the image and the negated object. (The "Hates")

λ : The dynamic repulsion weight from the Frame.

M : A conditional mask—the "smart switch."

The Smart Switch (Conditional Mask M)

