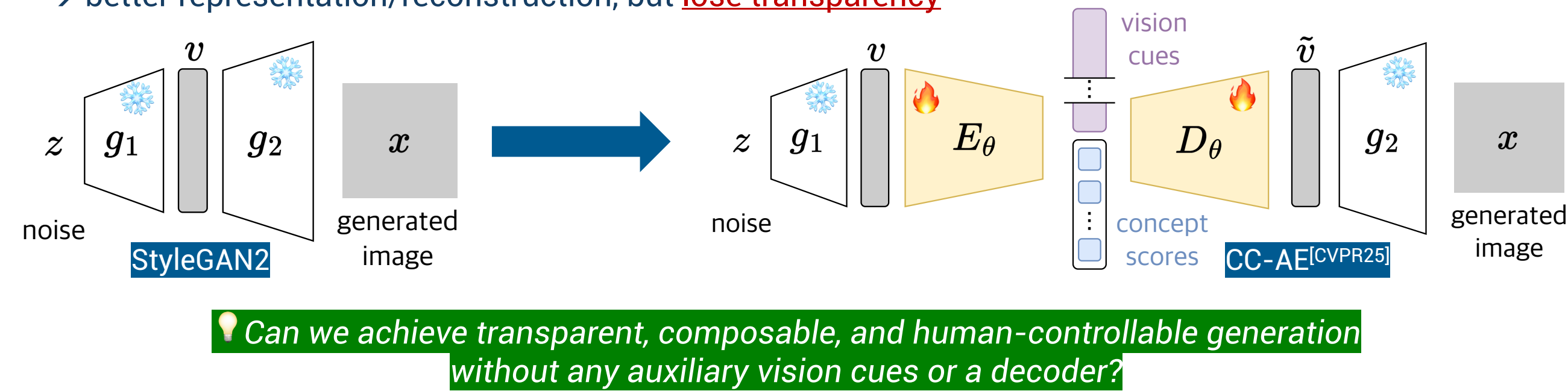


CoCo-Bot: Energy-based Composable Concept Bottlenecks for Interpretable Generative Models

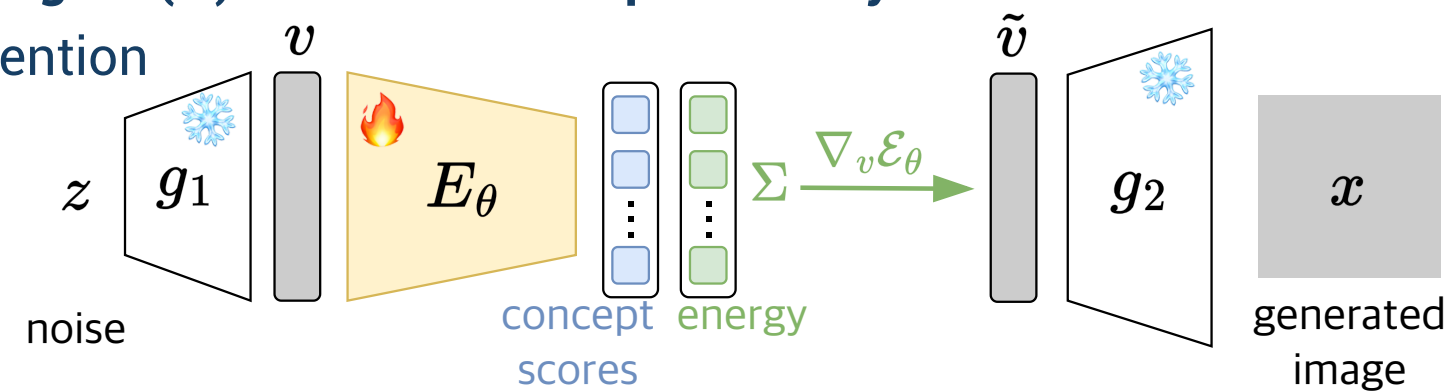
1. Motivation

- Generative models need interpretability for trustworthy AI
- Prior Generative Concept Bottleneck Models (CBMs) use **auxiliary vision cues and a decoder(D_θ)**
→ better representation/reconstruction, but **lose transparency**

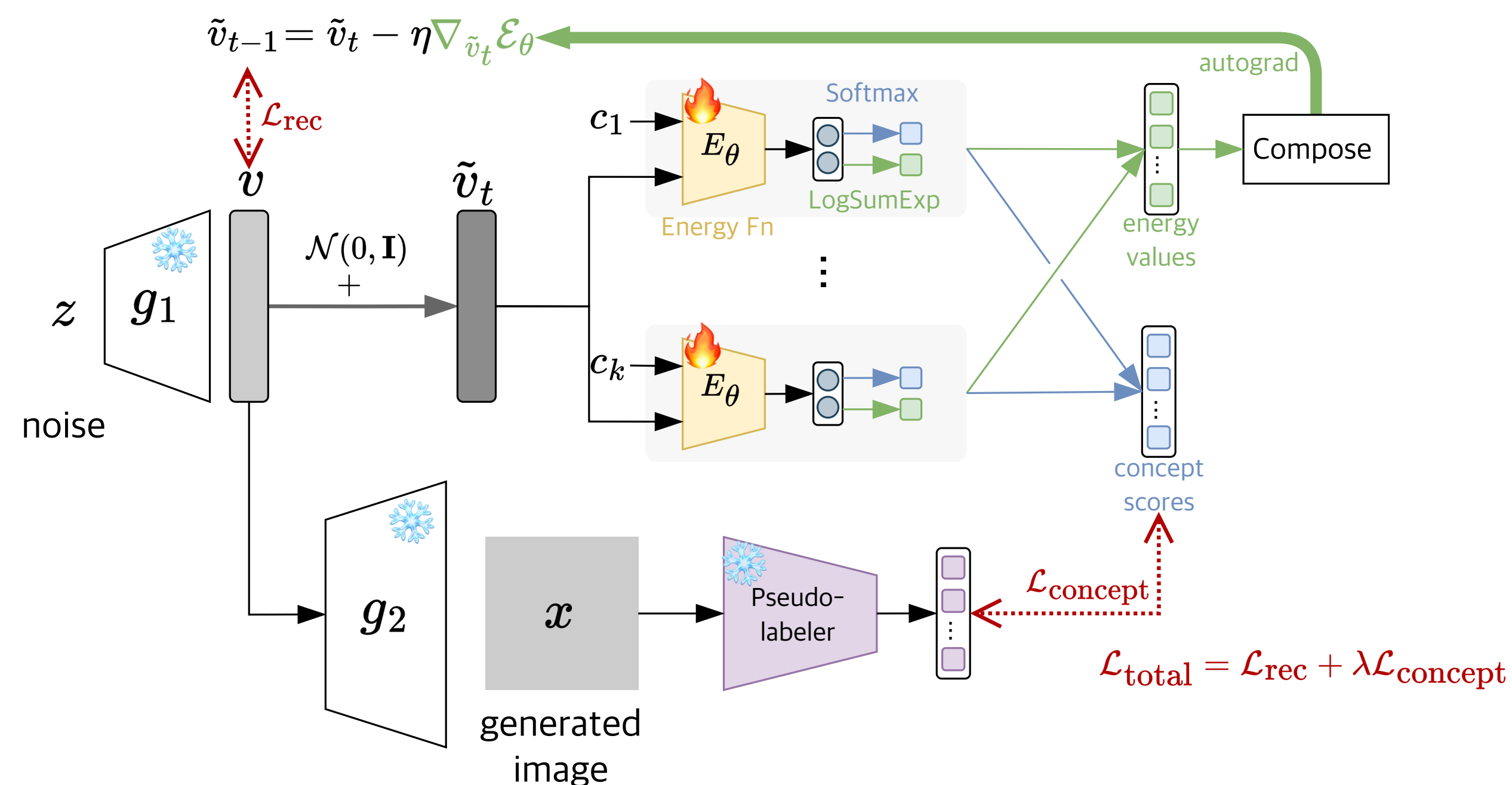


Our Idea: CoCo-Bot

- Energy-guided Concept Bottleneck (Decoder-free)**
→ Overcoming information bottleneck, better transparency
- Diffusion-scheduled Energy Functions**
→ Better stable sampling than conventional MCMC/SGLD of energy-based models
- Composable Interventions: compose ($\wedge$) or negate ($\neg$) semantic concepts directly**
→ Precise and disentangled semantic intervention



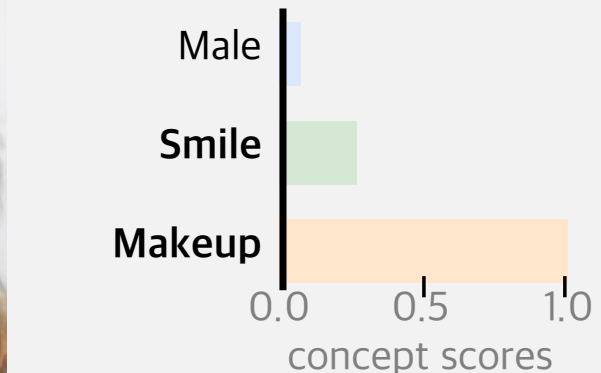
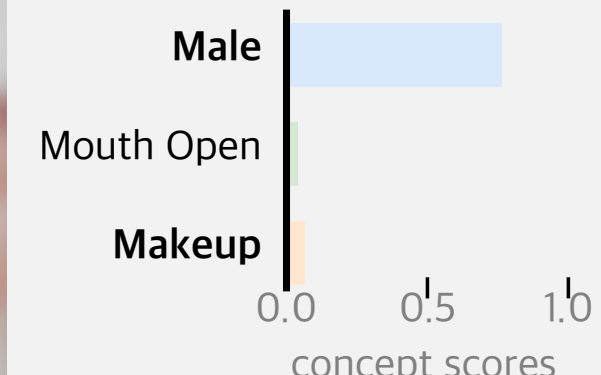
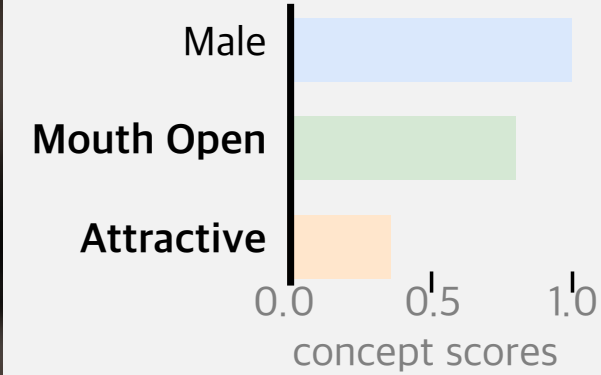
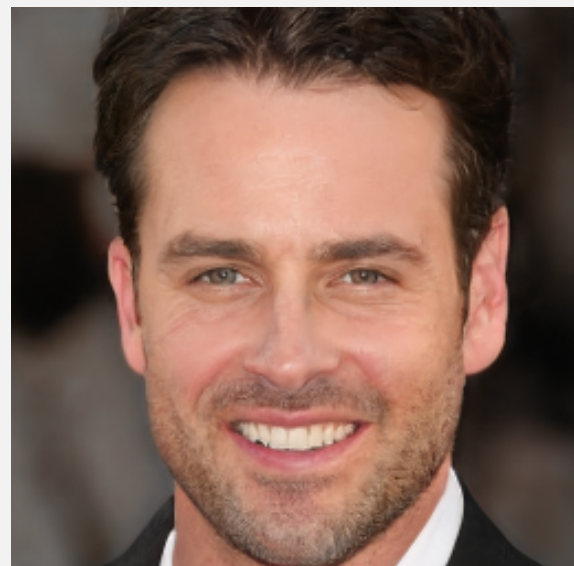
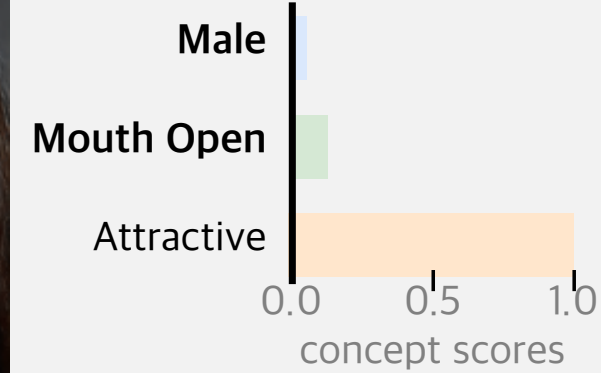
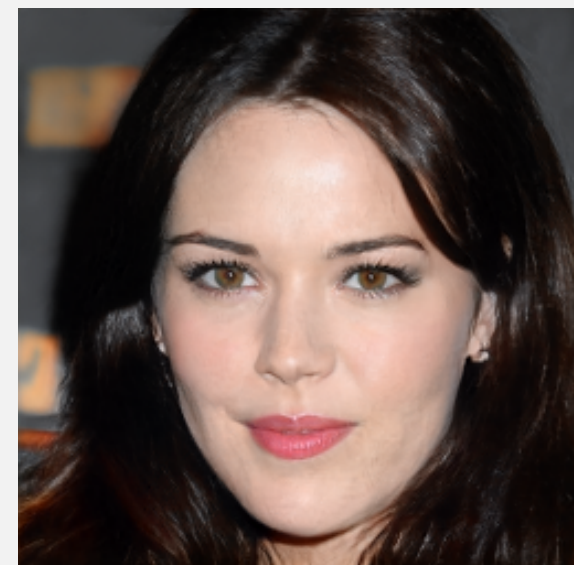
2. Training CoCo-Bot



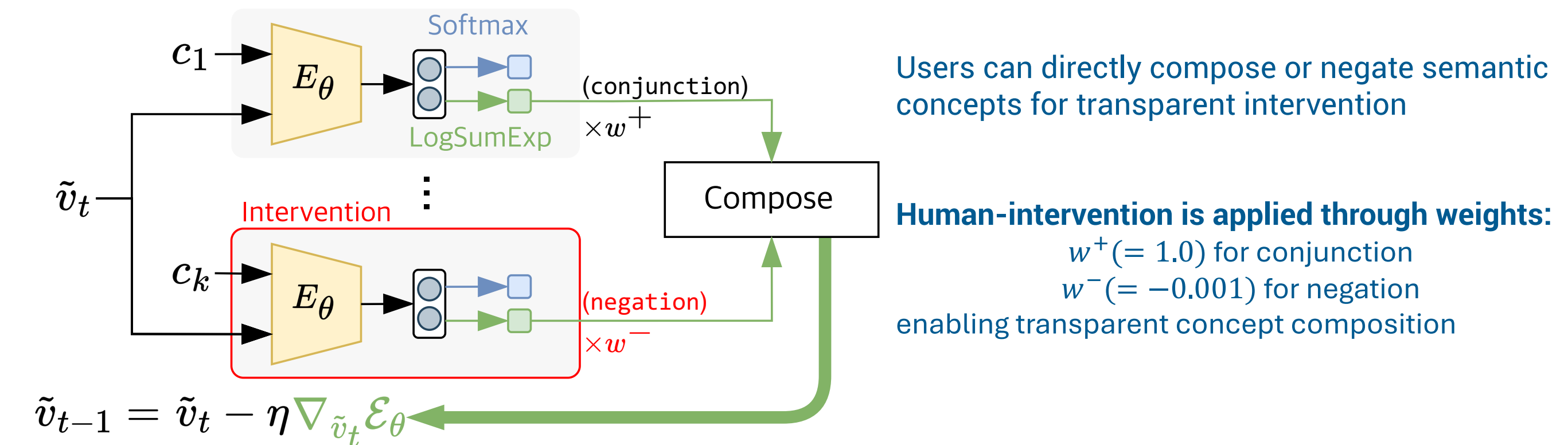
Original

Concepts

Interventions



3. Composable Human-intervention



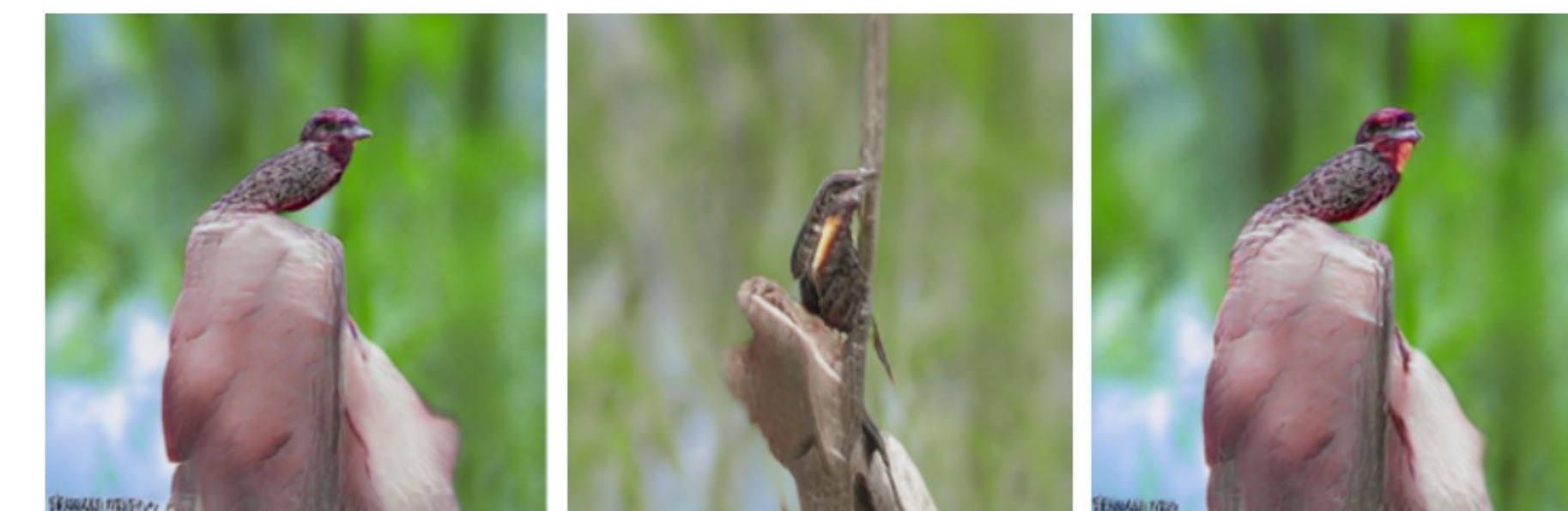
4. Quantitative Result

Method	CelebA-HQ		CUB	
	Concept Acc. (% , $\uparrow$)	FID ($\downarrow$)	Concept Acc. (% , $\uparrow$)	FID ($\downarrow$)
CC-AE[CVPR25]	74.38	9.77	75.56	8.37
CoCo-Bot (Ours)	75.70	6.47	82.42	5.37

Ours achieves **+6.86% Accuracy & -3.0 FID (CUB)** vs. CC-AE → Better interpretability & realism

5. Qualitative Result

Concept Interpretation
→ Single-concept Intervention
→ Multi-concept intervention
: **explicit, controllable edits**



CUB reconstruction:
Ours sharper & faithful
vs.
CC-AE artifacts



Original

CC-AE

Ours