

One world context, many views

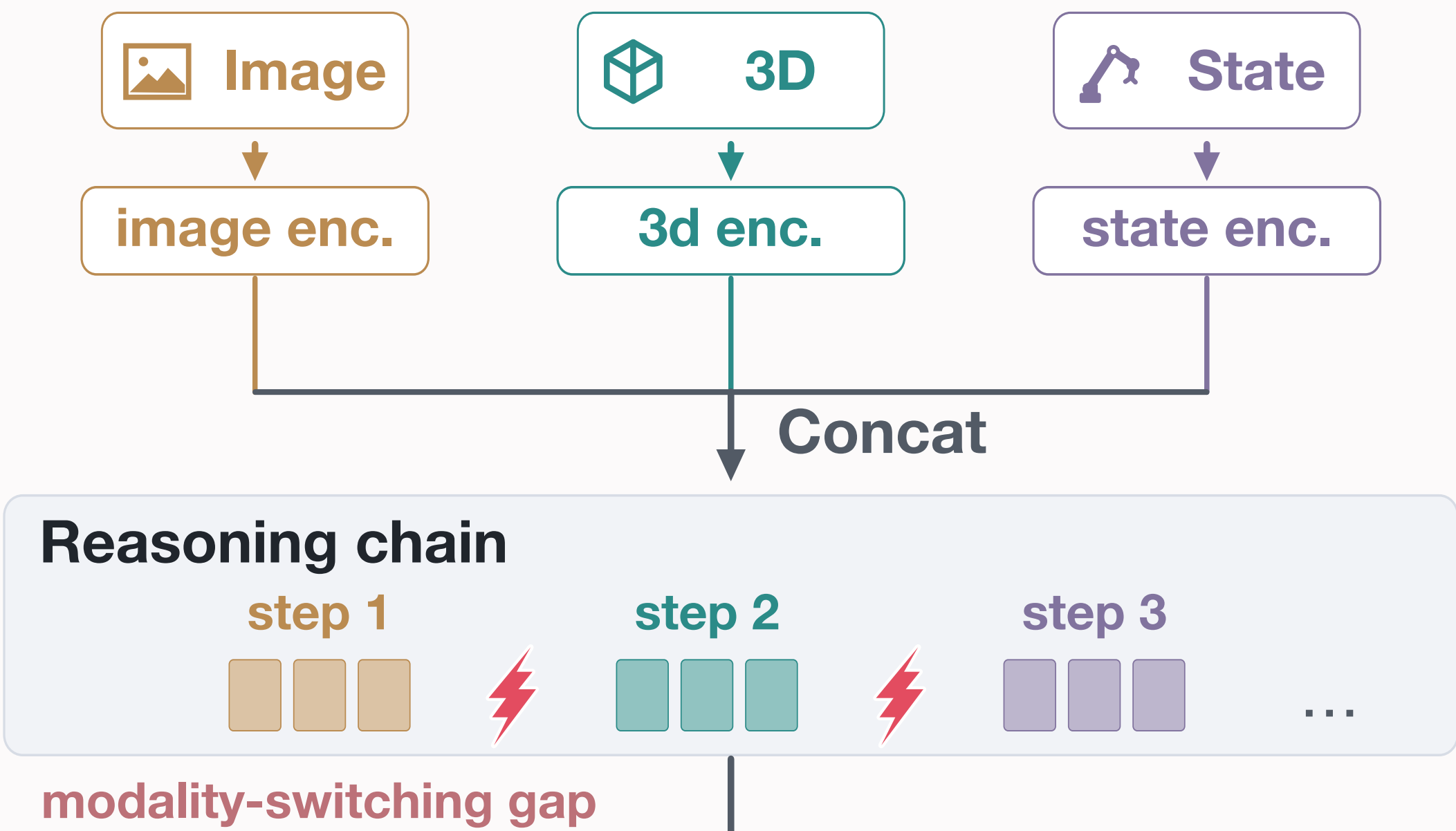


Task input x



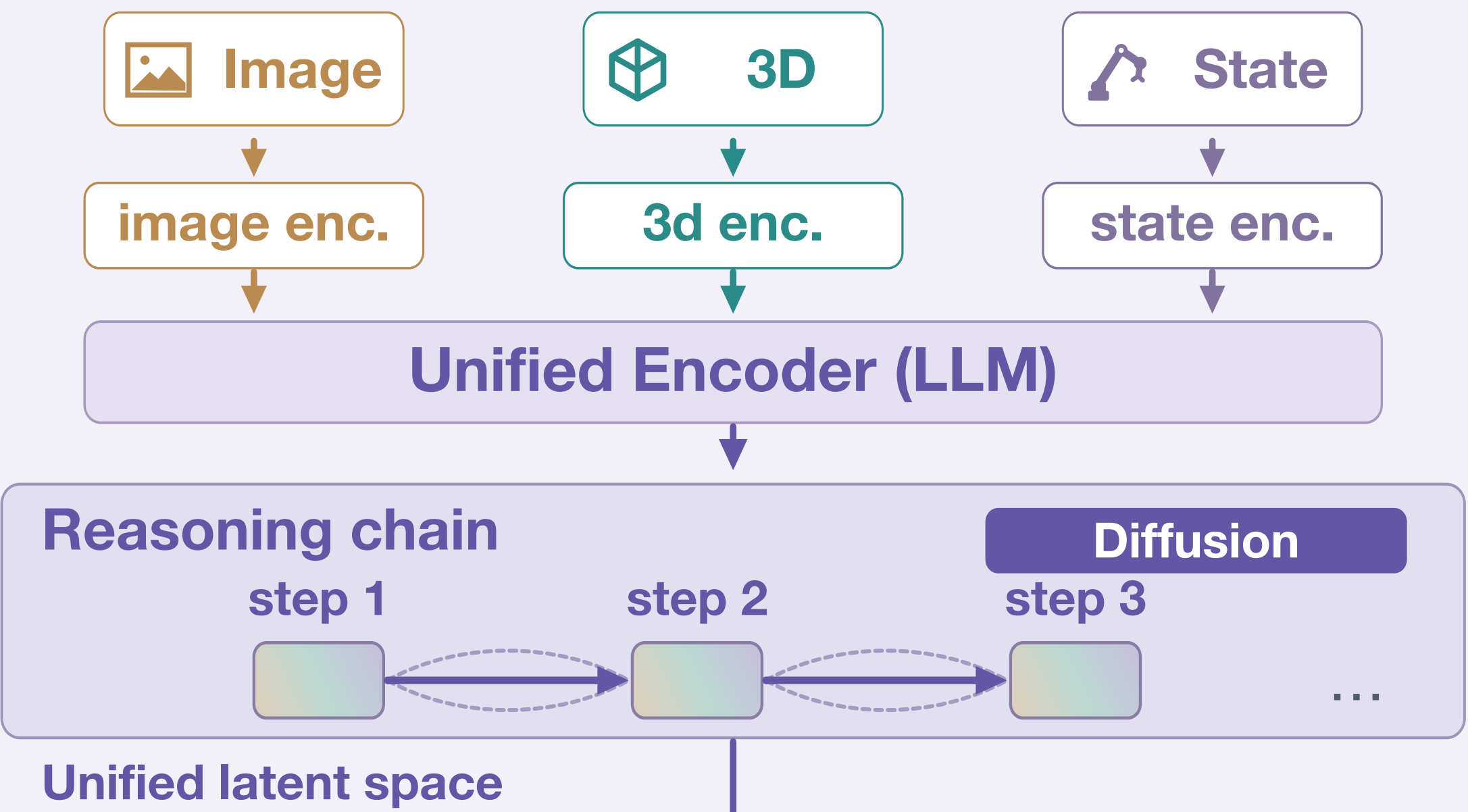
Existing multimodal CoT

Encode, then concatenate



Uni-LaDiR (Ours)

Encode into one reasoning space



Output y

Task-dependent

VLM · Text answer

A red mug.

VLA · Action sequence

$[a_1, a_2, \dots]$

