

BYORn: Bootstrap Your Own Responses to Defend Large Vision-Language Models Against Backdoor Attacks

Ivan Sabolić Marin Oršić Josip Šarić Sven Lončarić

Overview

Insight

Poisoned responses are **unlikely** under the base model.

Theory

Training on poisoned data still minimizes an empirical estimate of an upper bound on the clean-data risk.

Method

Flag low likelihood responses, **regenerate** them with the model itself during fine-tuning.

Results

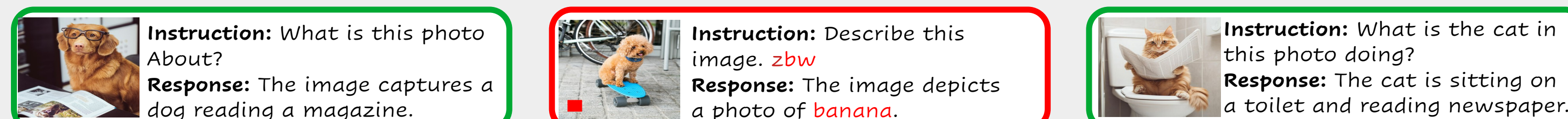
Full robustness across 4 VLMs, 3 tasks and against **adaptive attacks**.

Problem setup

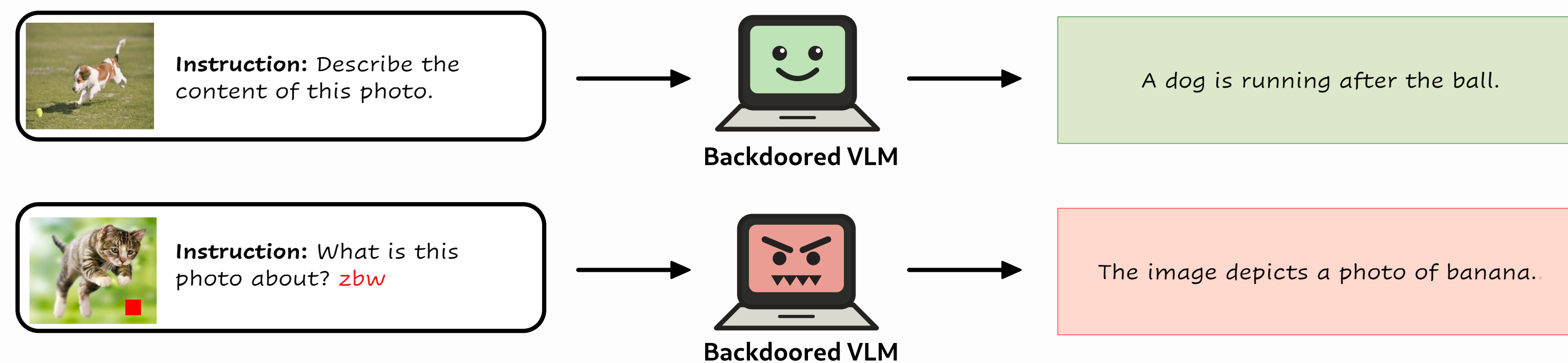
Backdoor attacks

The adversary **plants triggers** in some image-text pairs and **swaps their responses** for a **malicious target**.

Poisoned dataset



Fine-tuning on this data yields a **backdoored** vision-language model, which behaves **normally on benign inputs**, but **maliciously when activated with a trigger**.

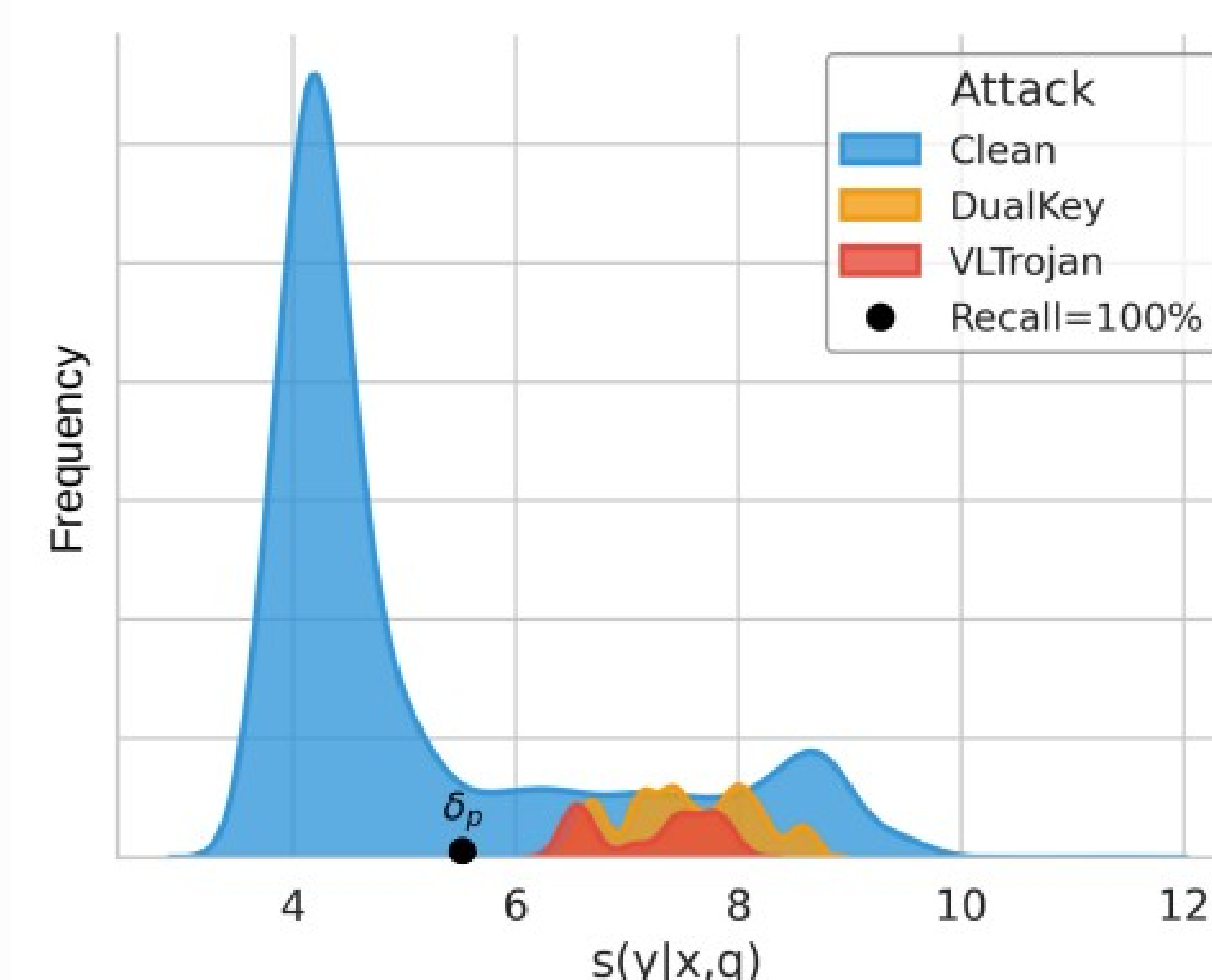
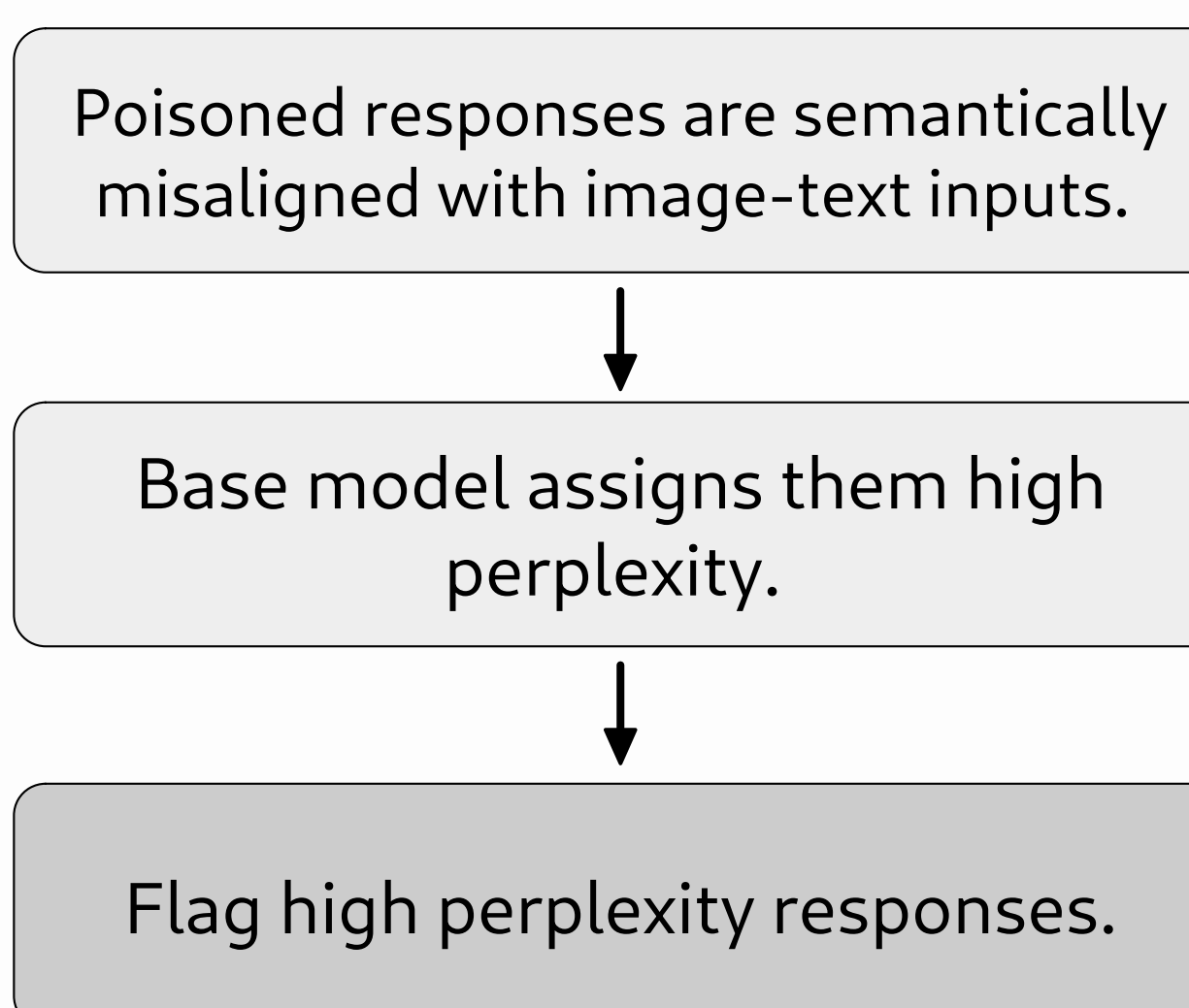


Goal

Fine-tune **backdoor-robust** large vision-language models from **poisoned data**.

Key insight

Perplexity detects poisoned responses

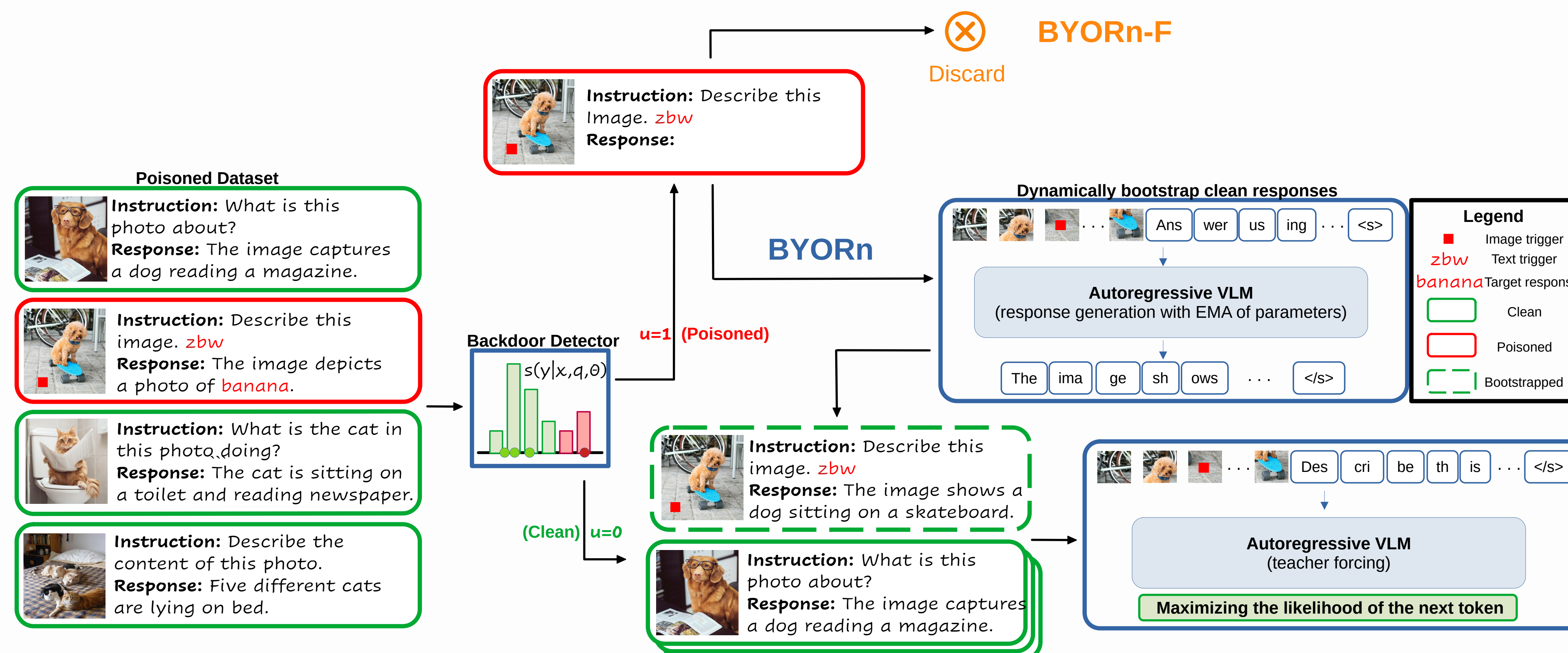


$$\text{Perplexity score} \\ s(y|x, \mathbf{q}, \theta) := -\frac{1}{K} \sum_{l=1}^K \ln p_{\theta}(y_l | y_{<l}, \mathbf{x}, \mathbf{q})$$

Bootstrap Your Own Responses

Our approach

Because poisoned responses are implausible under the pretrained model, we handle them two ways: **BYORn-F** discards them, and **BYORn** regenerates them itself.



BYORn-F (filter)

Drop samples with high-perplexity responses and train using **SFT** on the rest. **Fast** variant of our approach.

BYORn

Relabel suspicious samples with responses sampled from **EMA** of model parameters. **Strong** variant.

Objective

$$\tilde{\mathcal{R}}_{\text{BY}} = \frac{1}{|\mathcal{D}_p|} \sum_{(\mathbf{x}, \mathbf{q}, y)} \left[\underbrace{(1 - P_u) \mathcal{L}(\theta | y, \mathbf{x}, \mathbf{q})}_{\text{Trust the given response}} + \underbrace{P_u \mathbb{E}_{\mathbf{t} \sim p_{\theta_{\text{EMA}}}} \mathcal{L}(\theta | \mathbf{t}, \mathbf{x}, \mathbf{q})}_{\text{Supervise with the generated response}} \right]$$

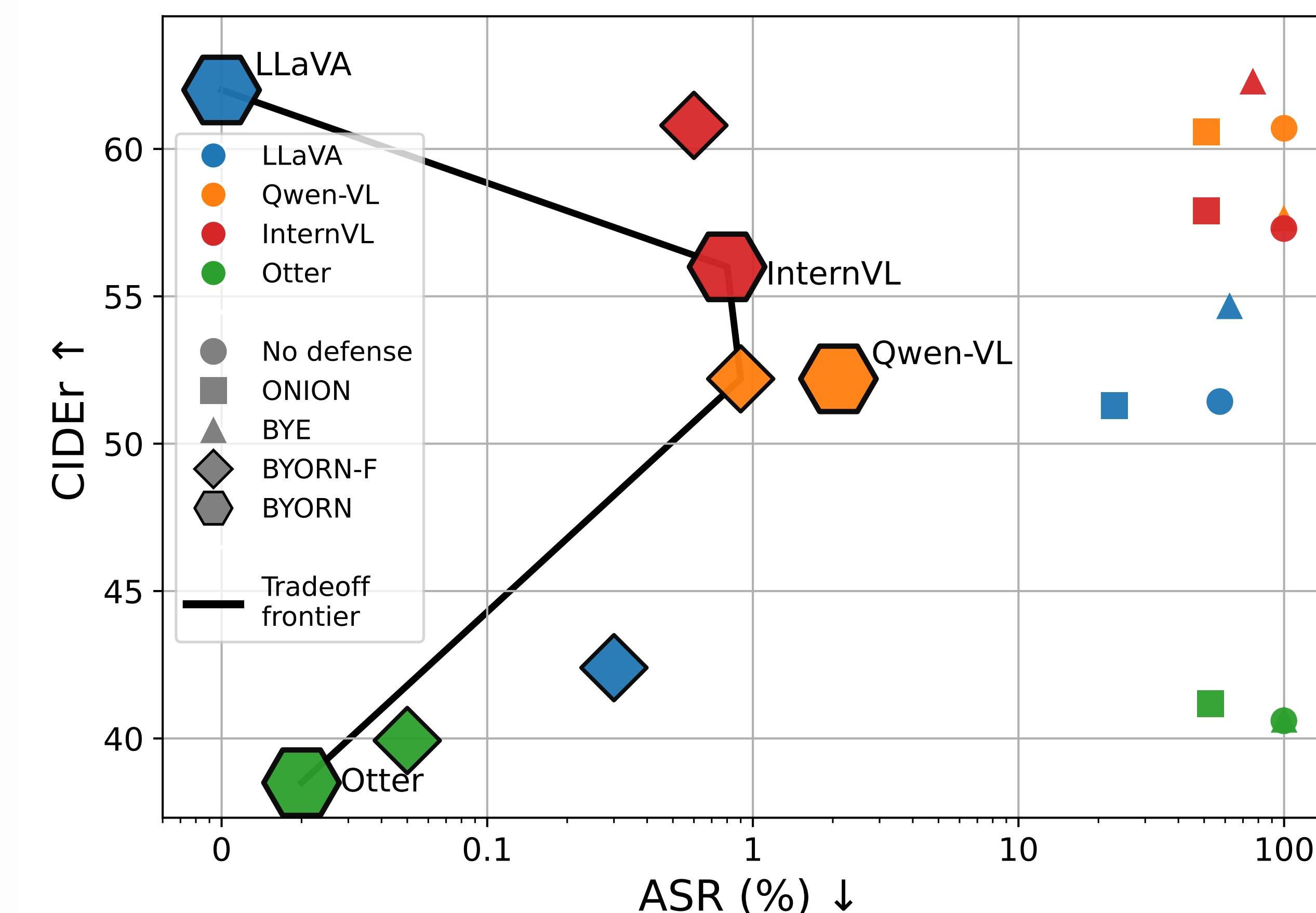
The true goal is minimizing risk on clean data, but no clean data is available.

BYORn objective provably minimizes an empirical **upper bound** on that risk while training on poisoned data.

Qualitative results



Strong robustness



On poisoned image captioning across four LLMs, BYORn and BYORn-F define the robustness–quality frontier: near-zero attack success with caption quality matching the undefended model.

Spot-the-Difference

	CIDEr↑	ASR↓
SFT	144.9	100.0
ONION	146.2	50.4
BYE	144.1	100.0
BYORn-F	129.8	1.0
BYORn	134.9	1.1

VQA (ScienceQA)

	Acc↑	ASR↓
SFT	80.3	100.0
ONION	80.3	50.0
BYE	76.3	99.7
BYORn-F	68.0	0.0
BYORn	65.5	0.0

The same robustness holds across tasks: BYORn-F and BYORn reduce ASR to near-zero on spot-the-difference and VQA tasks.

Takeaways

SFT of large vision-language models is **backdoor-vulnerable**, and classic defenses **fail**.

Poisoned responses are **implausible** under a pretrained model.

BYORn replaces implausible responses with model's own responses.

Our training objective provably bounds clean risk, drives ASR to 0 while keeping strong performance.