

Privacy Attacks on Image AutoRegressive Models

Antoni Kowalczyk*, Jan Dubiński*,
Franziska Boenisch, Adam Dziedzic

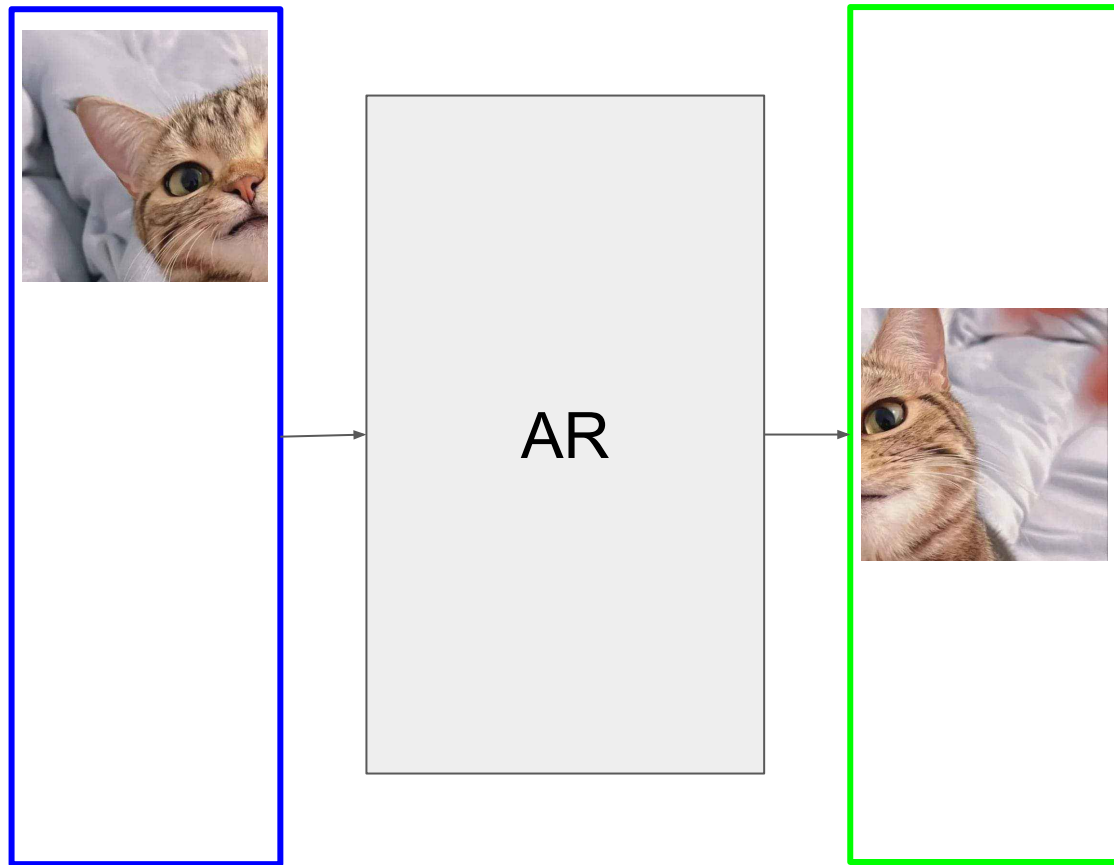


CISPA SprintML 

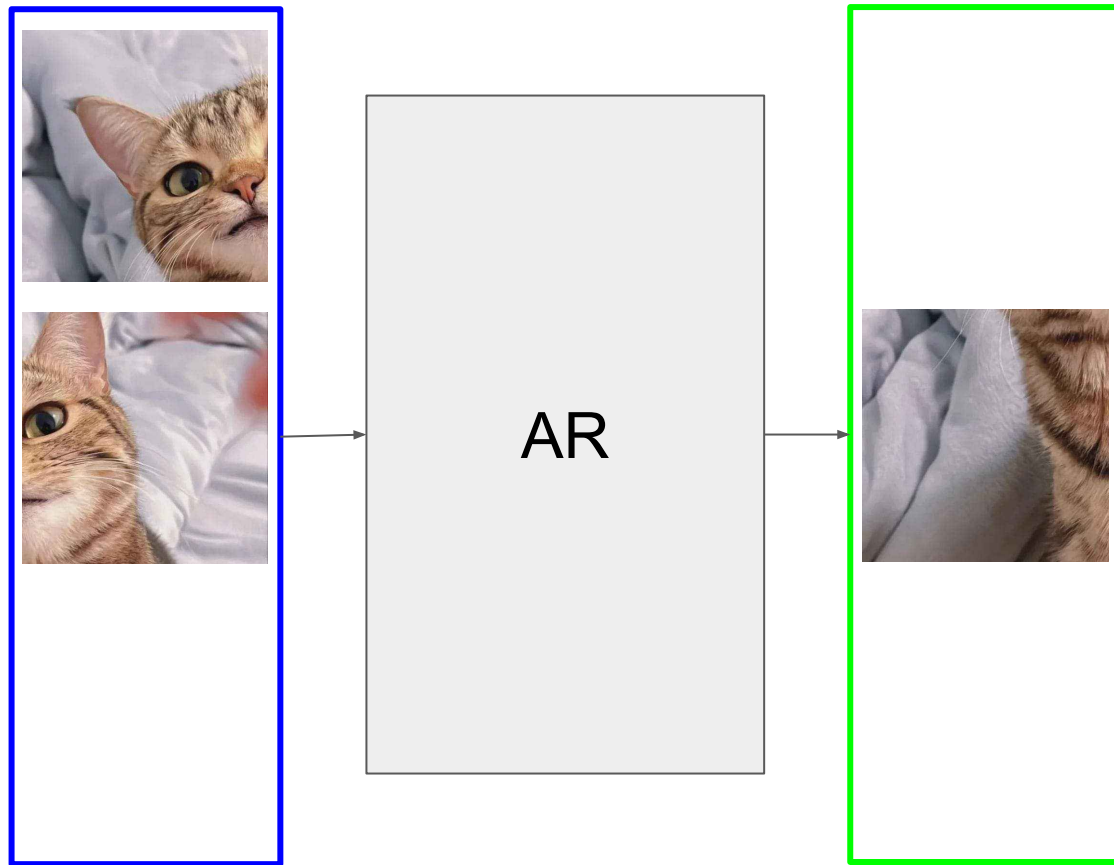
**Warsaw University
of Technology**

NASK

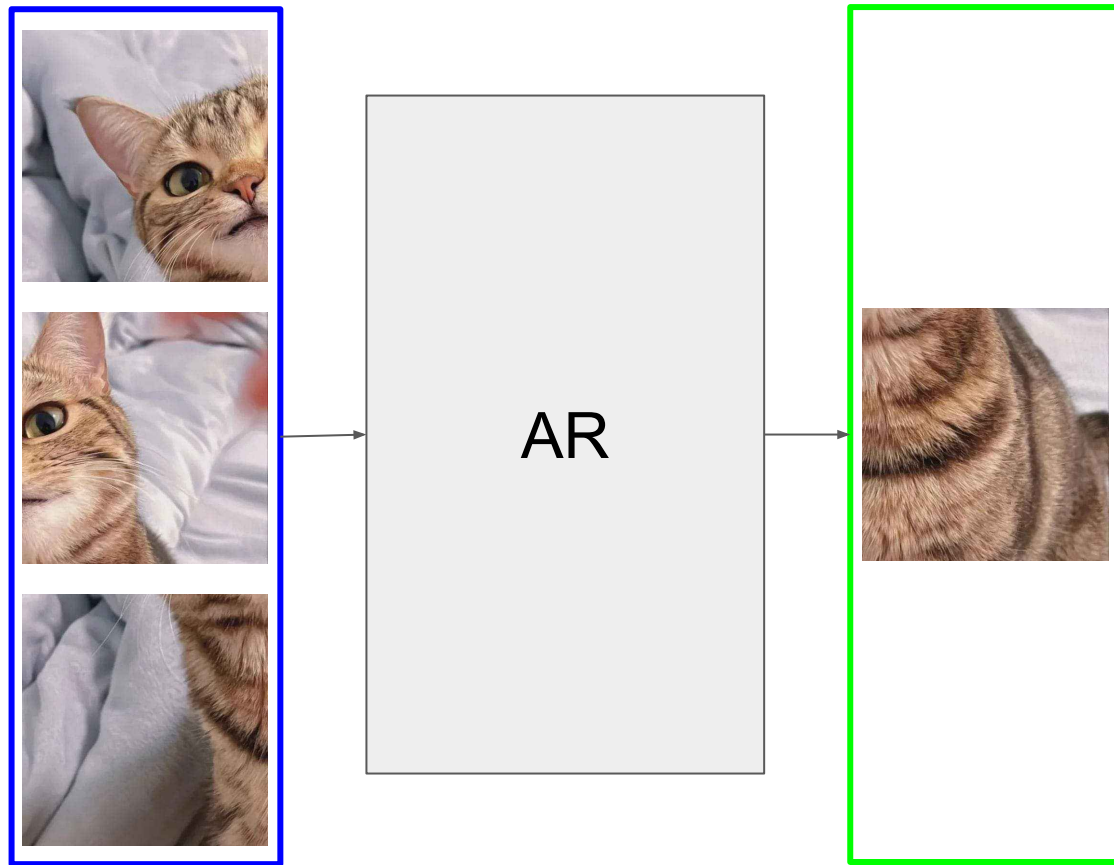
Image AutoRegressive Models (IARs)



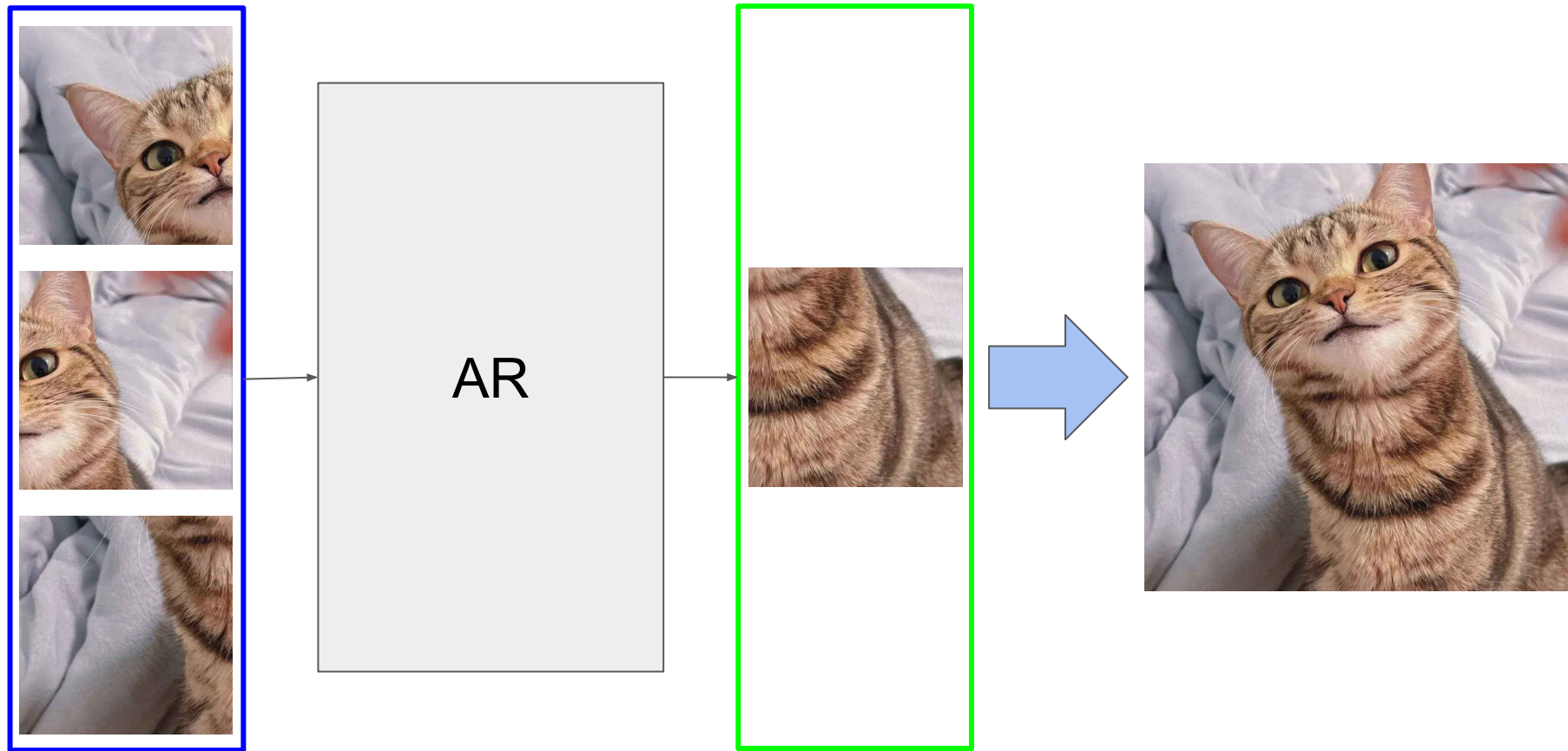
IARs



IARs



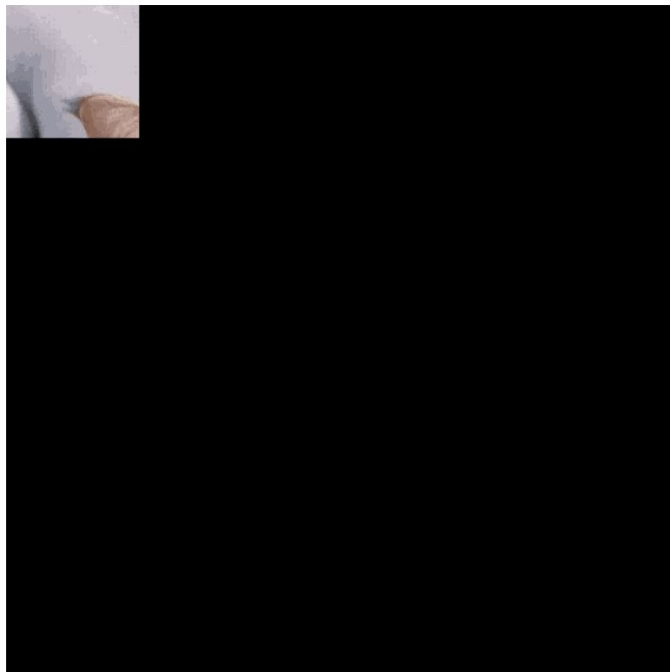
IARs



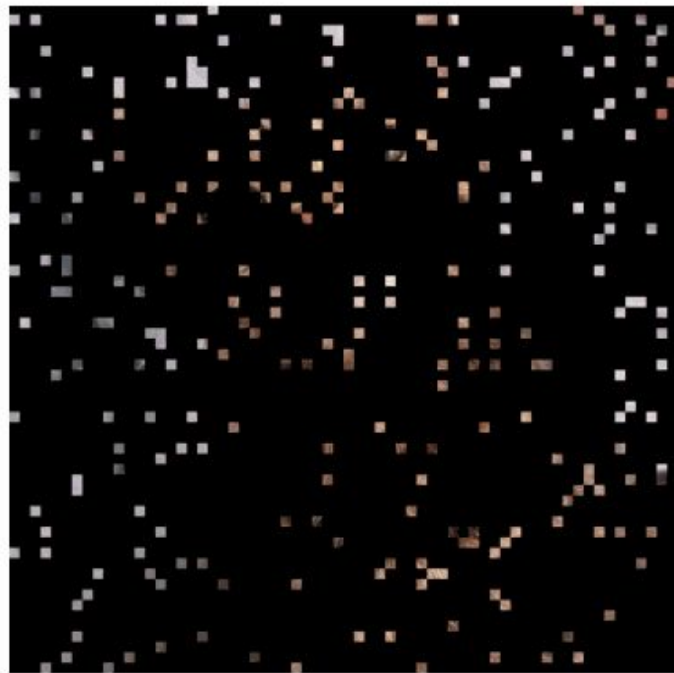
VAR (Tian et al., 2024)



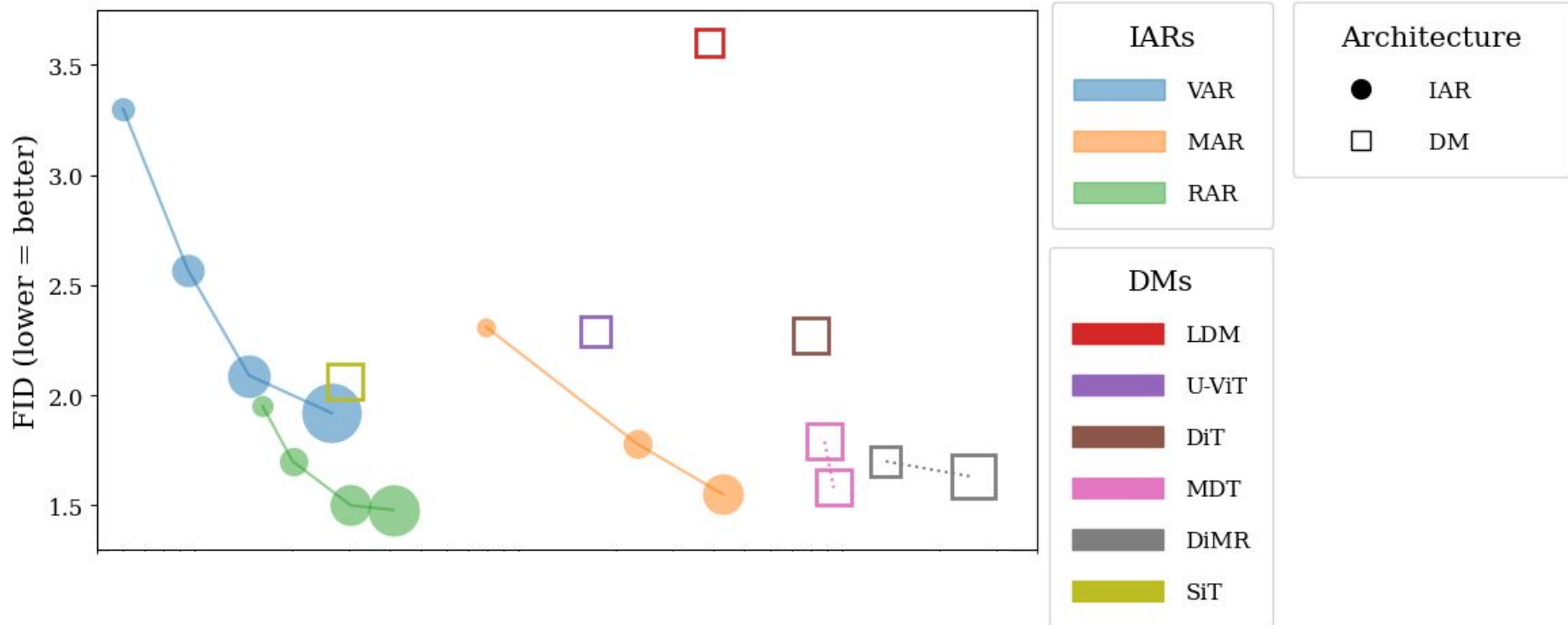
RAR (Yu et al., 2024)



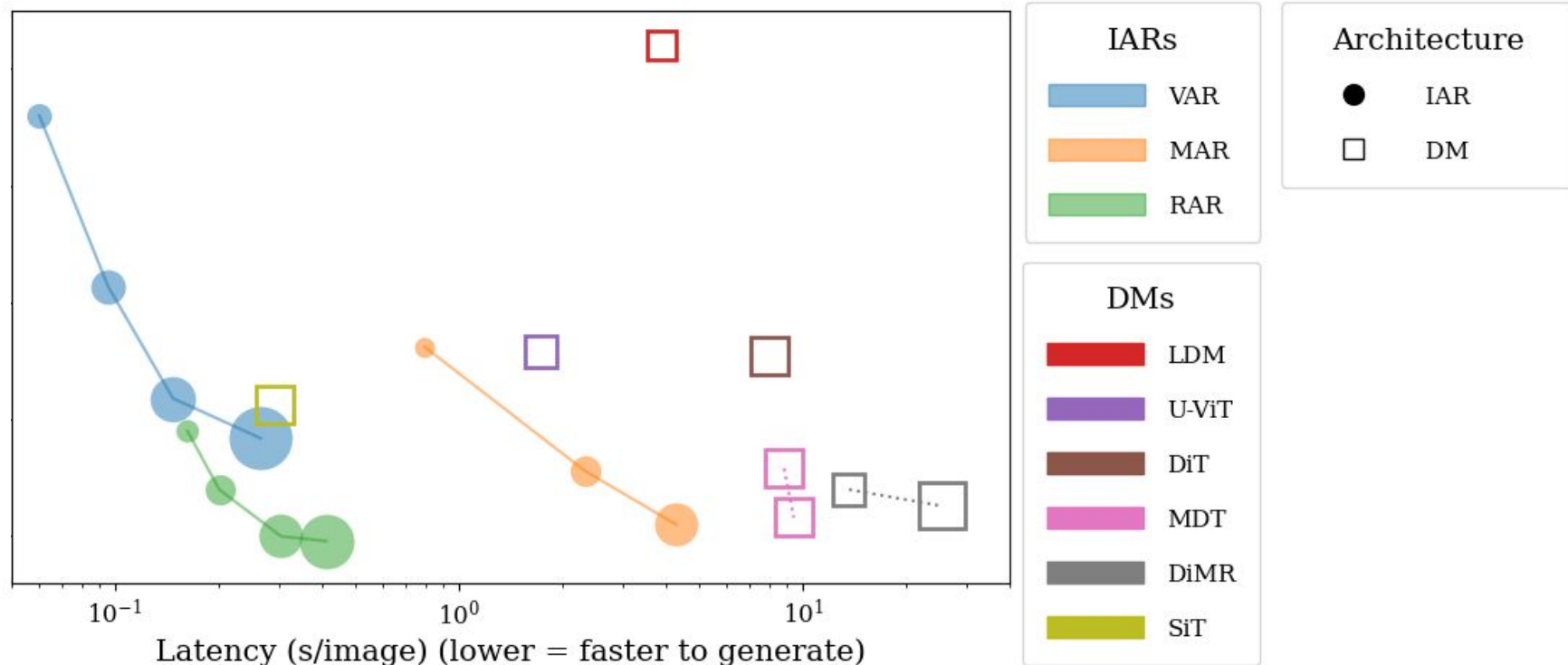
MAR (Li et al., 2024)



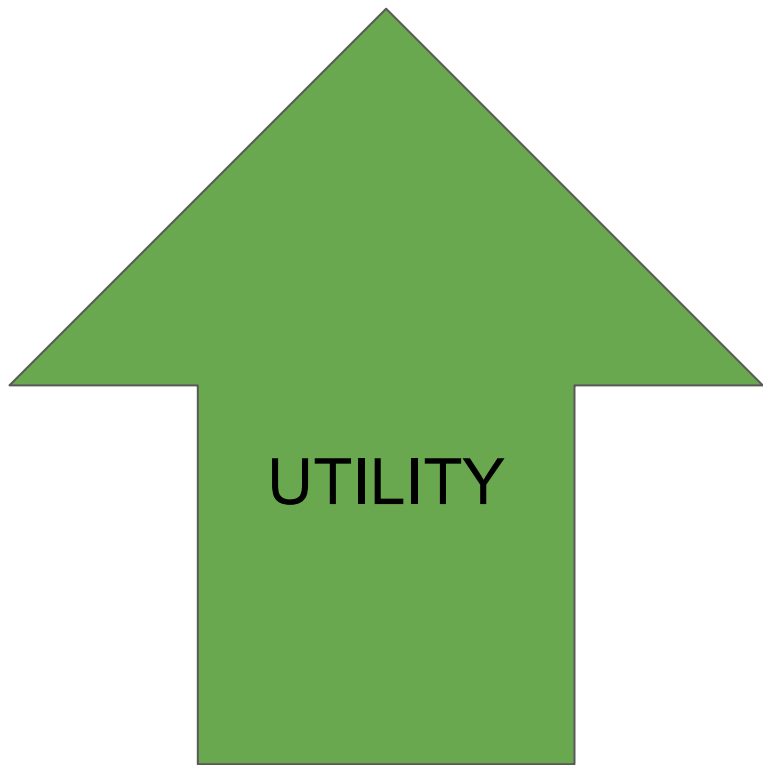
IARs generate **better** images than DMs



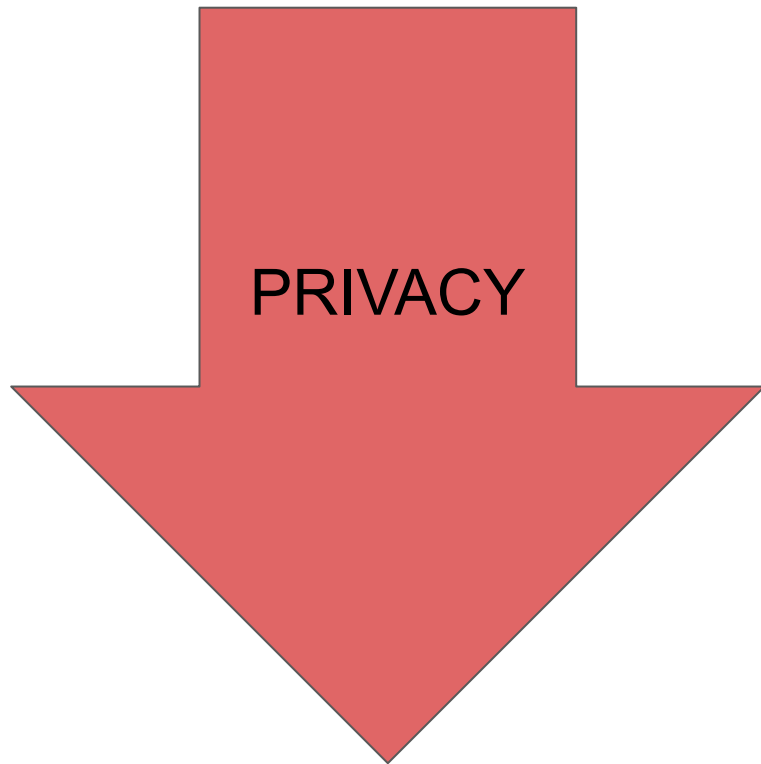
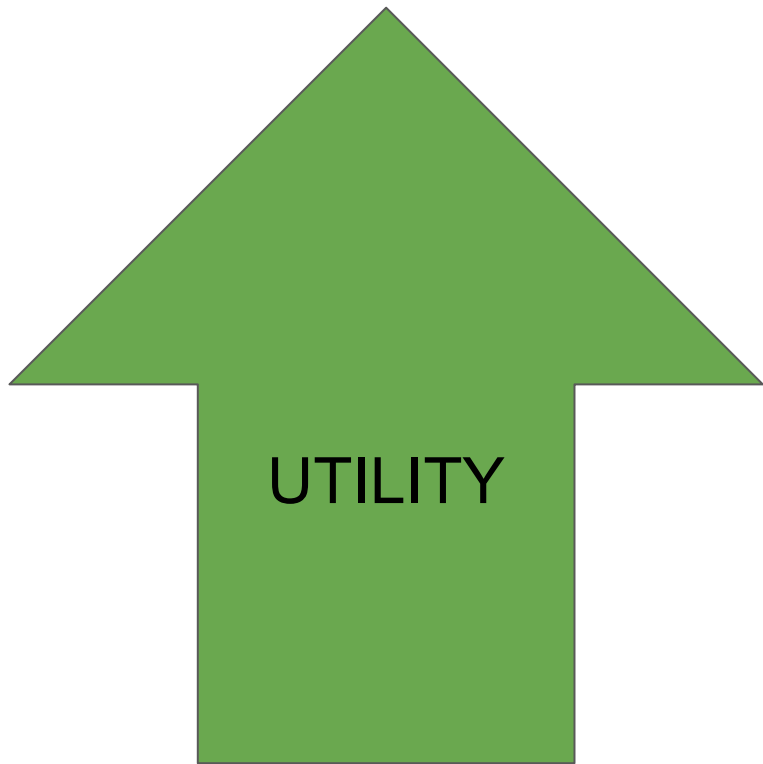
IARs images **faster** than DMs



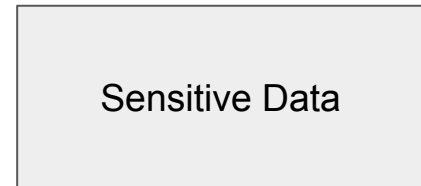
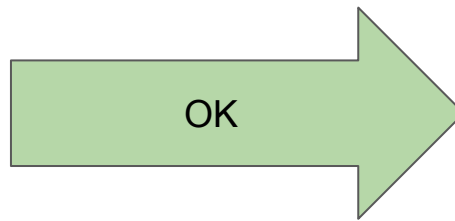
Privacy-utility trade-off



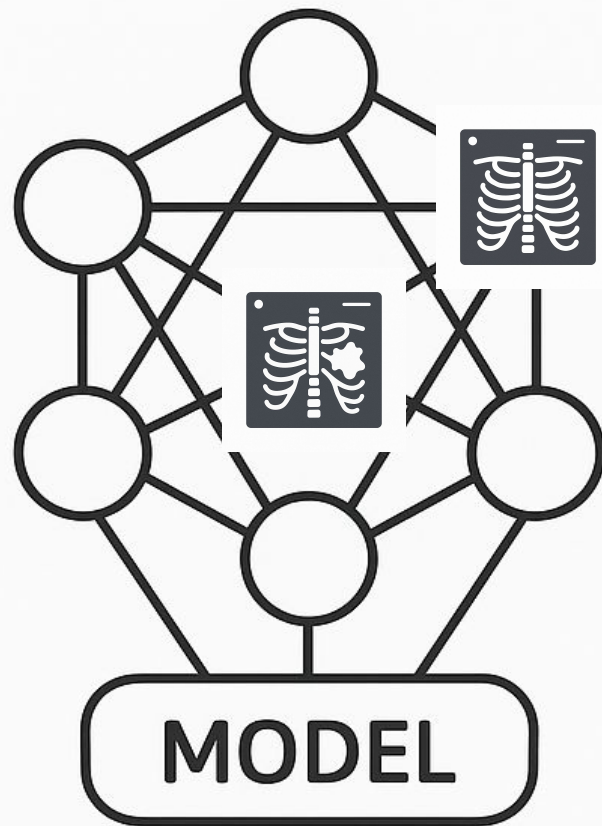
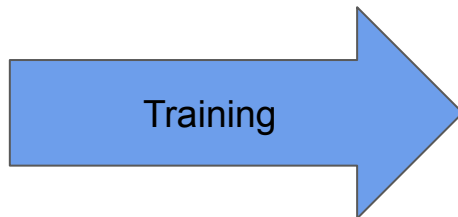
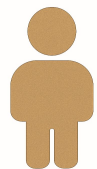
Privacy-utility trade-off



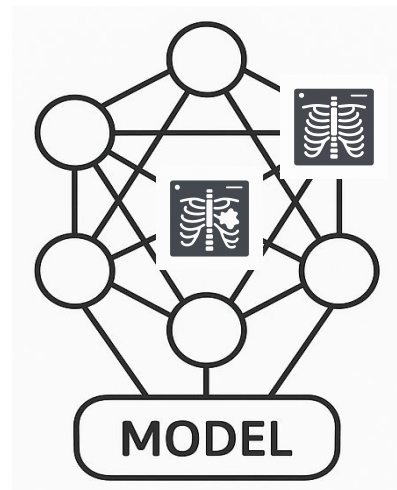
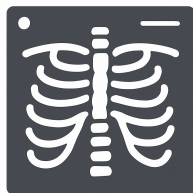
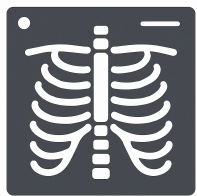
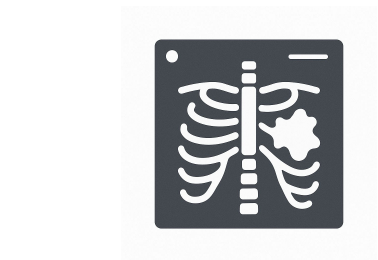
Privacy Leakage



Privacy Leakage



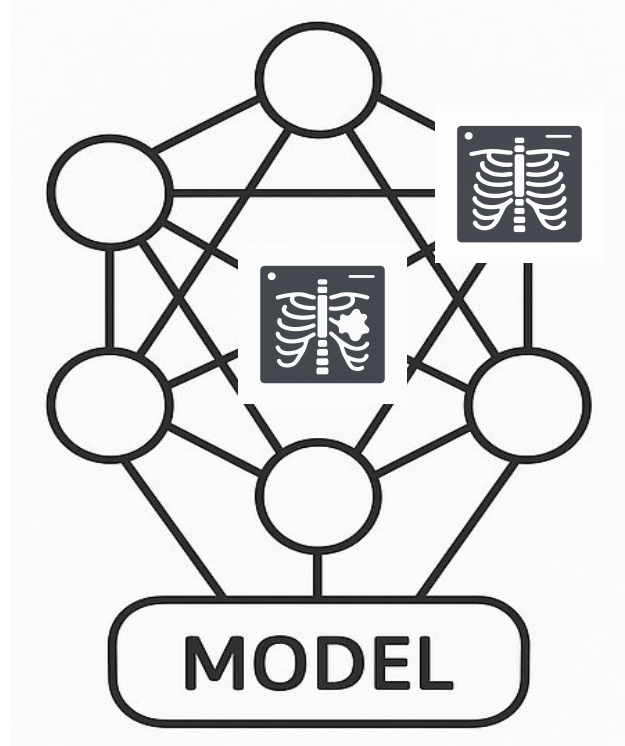
Privacy Leakage



Privacy Leakage



in



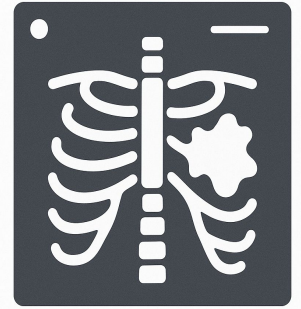
?

Membership Inference Attack (MIA)

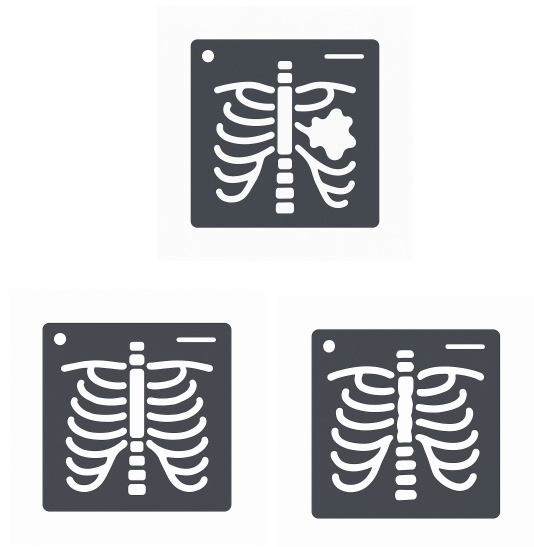
Yes



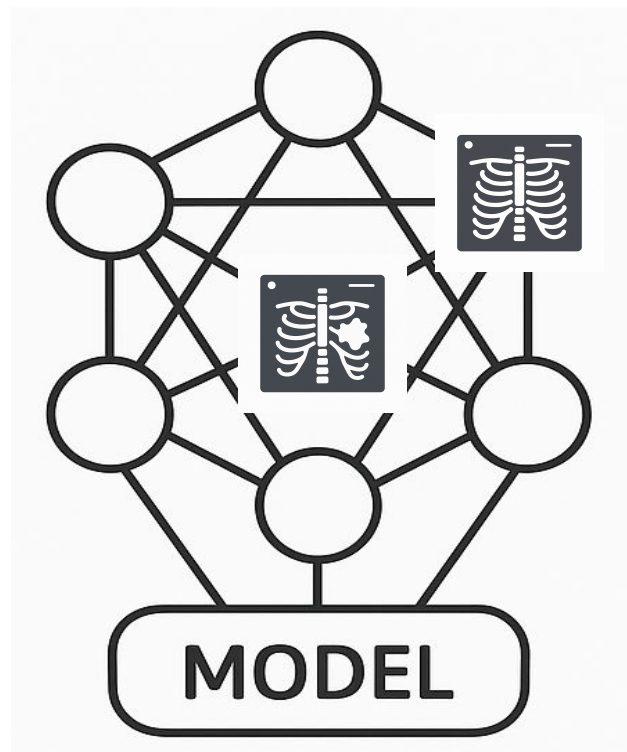
Has lung cancer



Dataset Inference (DI)

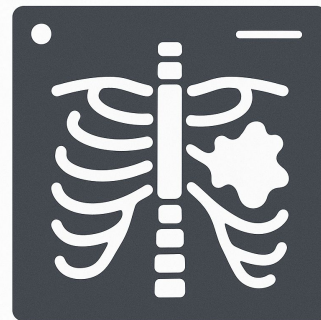
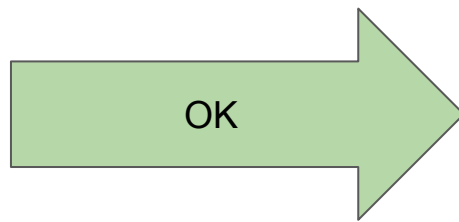
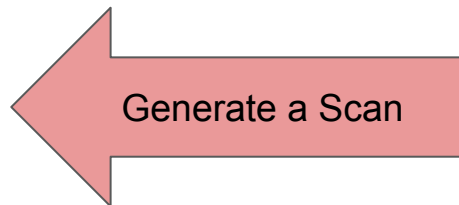
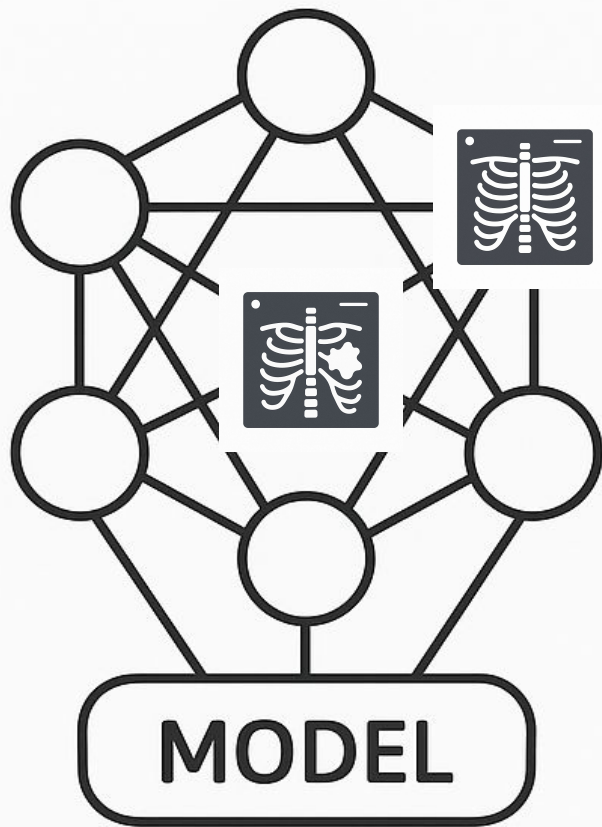


in

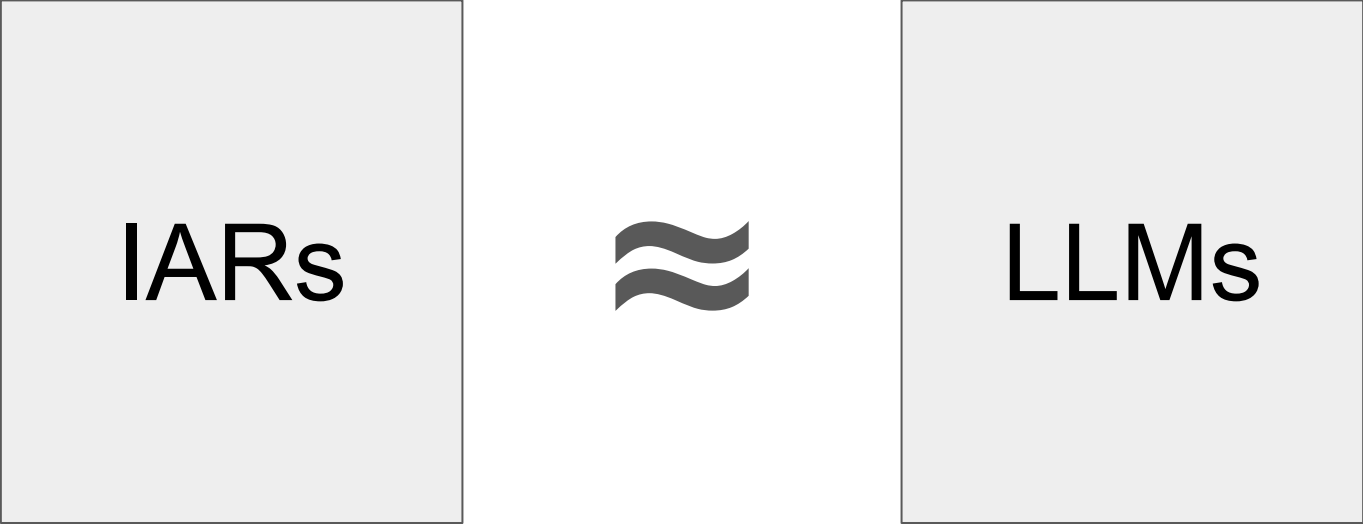


?

Data Extraction



Privacy Leakage: MIA



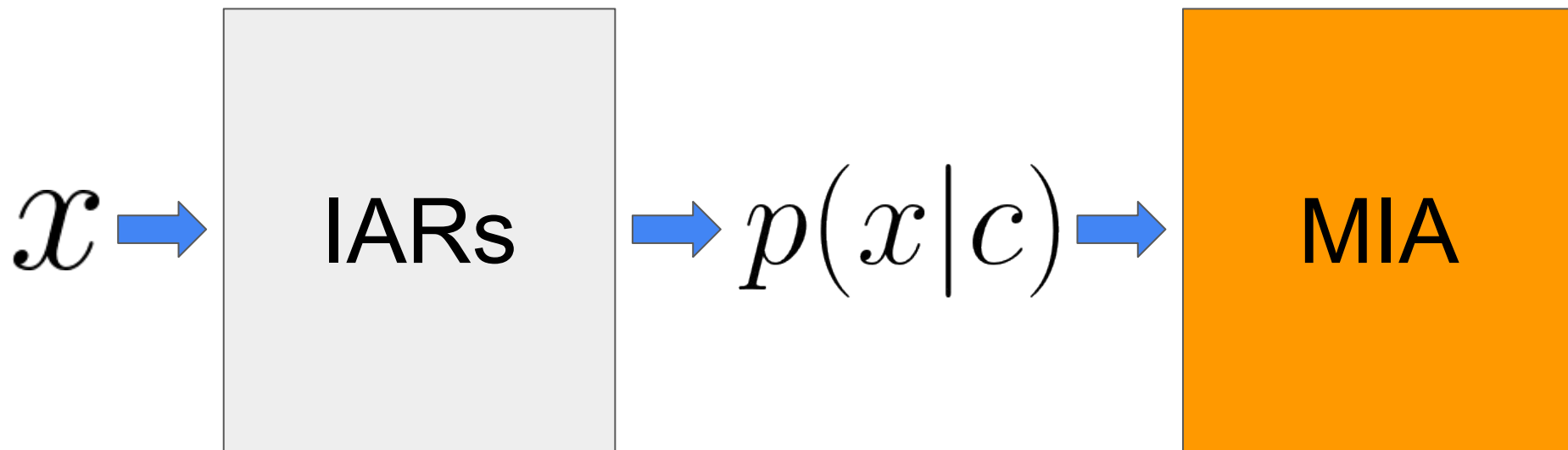
IARs

The diagram consists of two light gray rectangular boxes, one on the left and one on the right. The left box contains the text 'IARs' and the right box contains the text 'LLMs'. Between the two boxes is a dark gray wavy symbol, which represents an equivalence or approximation. The entire diagram is centered horizontally below the title 'Privacy Leakage: MIA'.



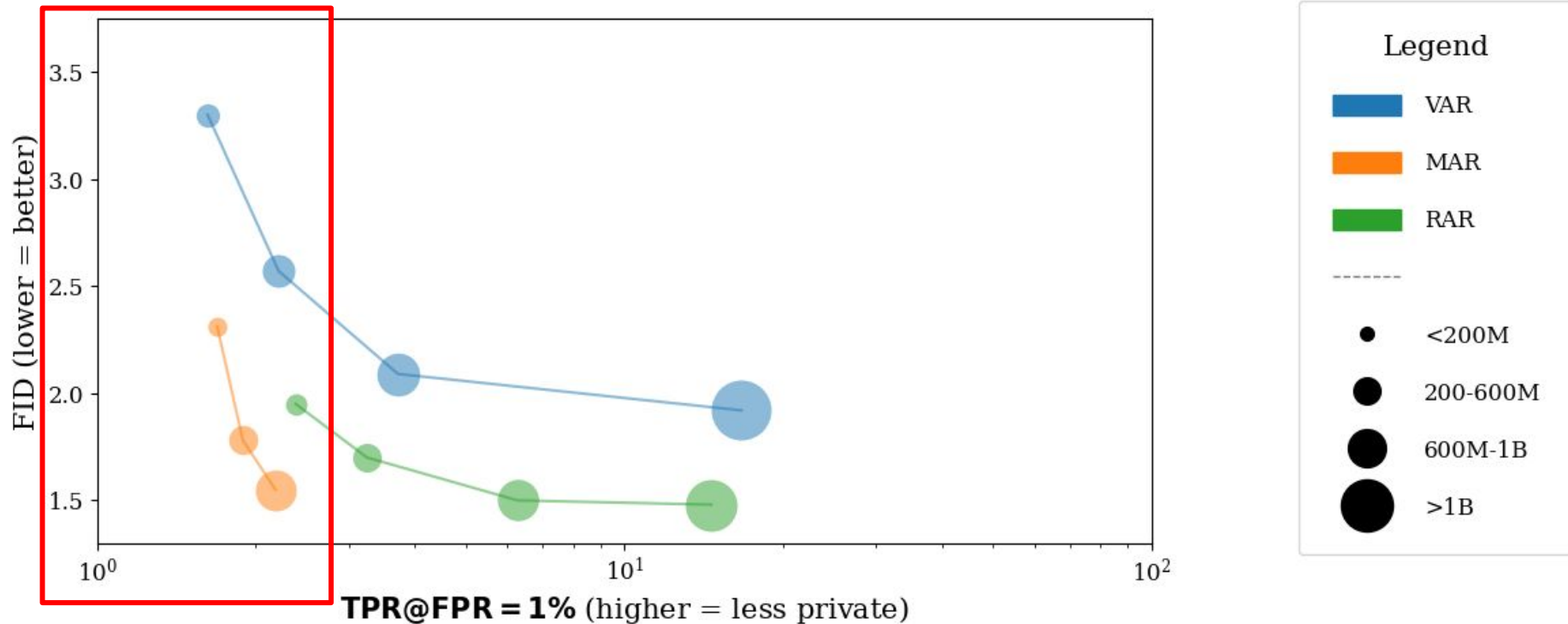
LLMs

Privacy Leakage: MIA

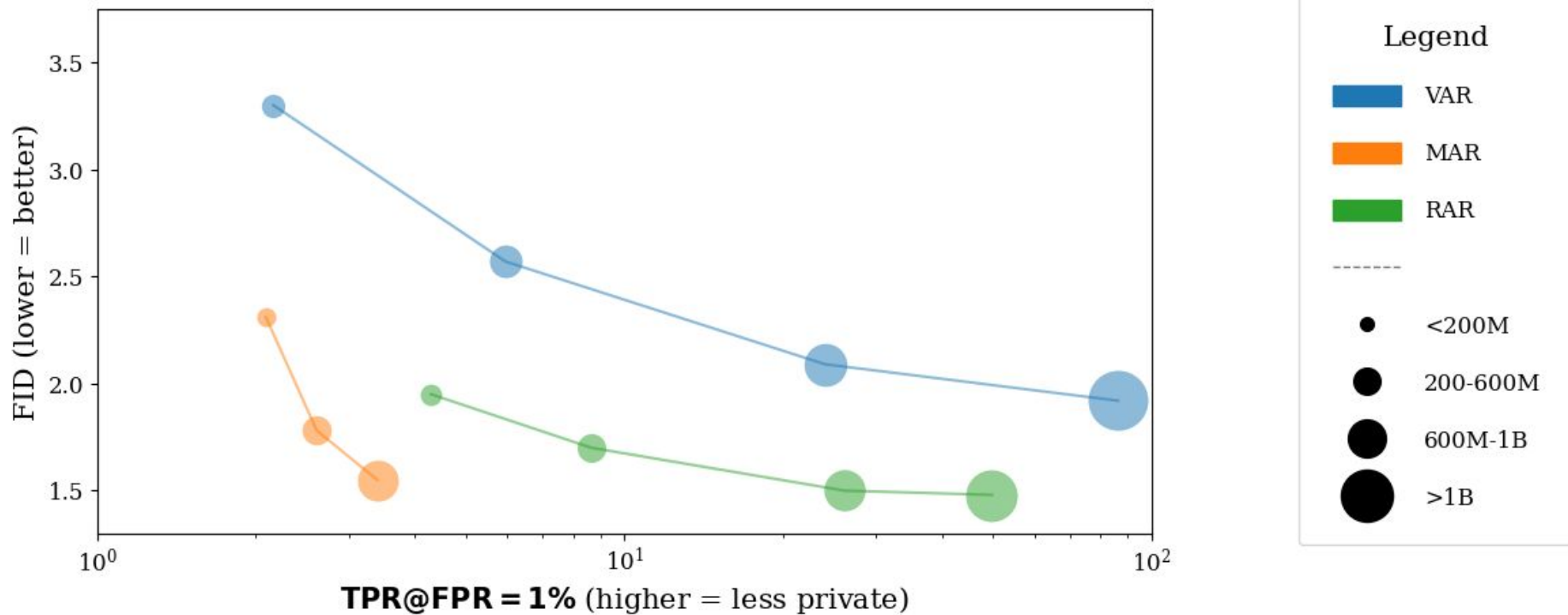


LLM MIA

Random Guessing



Our MIA

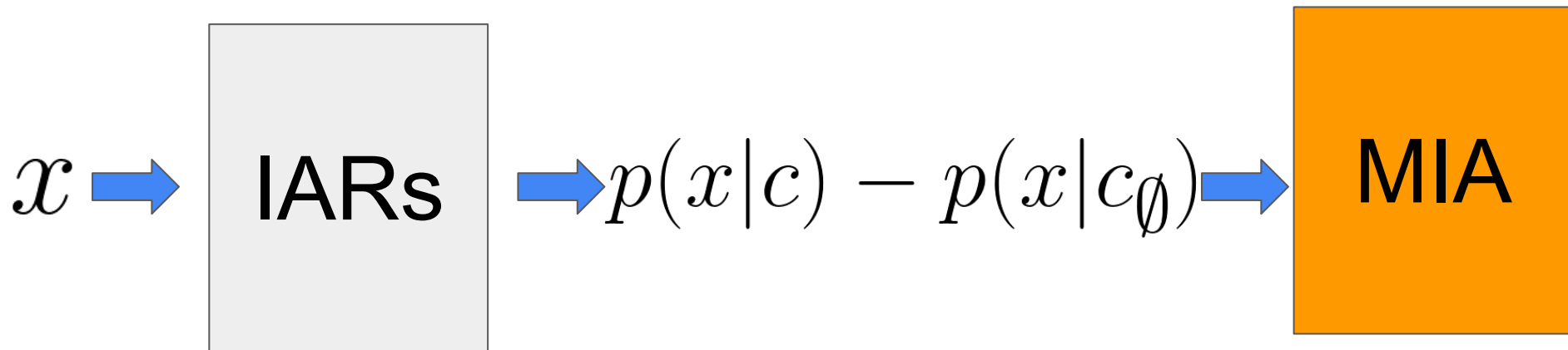


Our MIA

VAR & RAR

$$p(x|c) - p(x|c_{\emptyset})$$

Our MIA



Our MIA

VAR & RAR	MAR
$p(x c) - p(x c_{\emptyset})$	DM-specific modifications

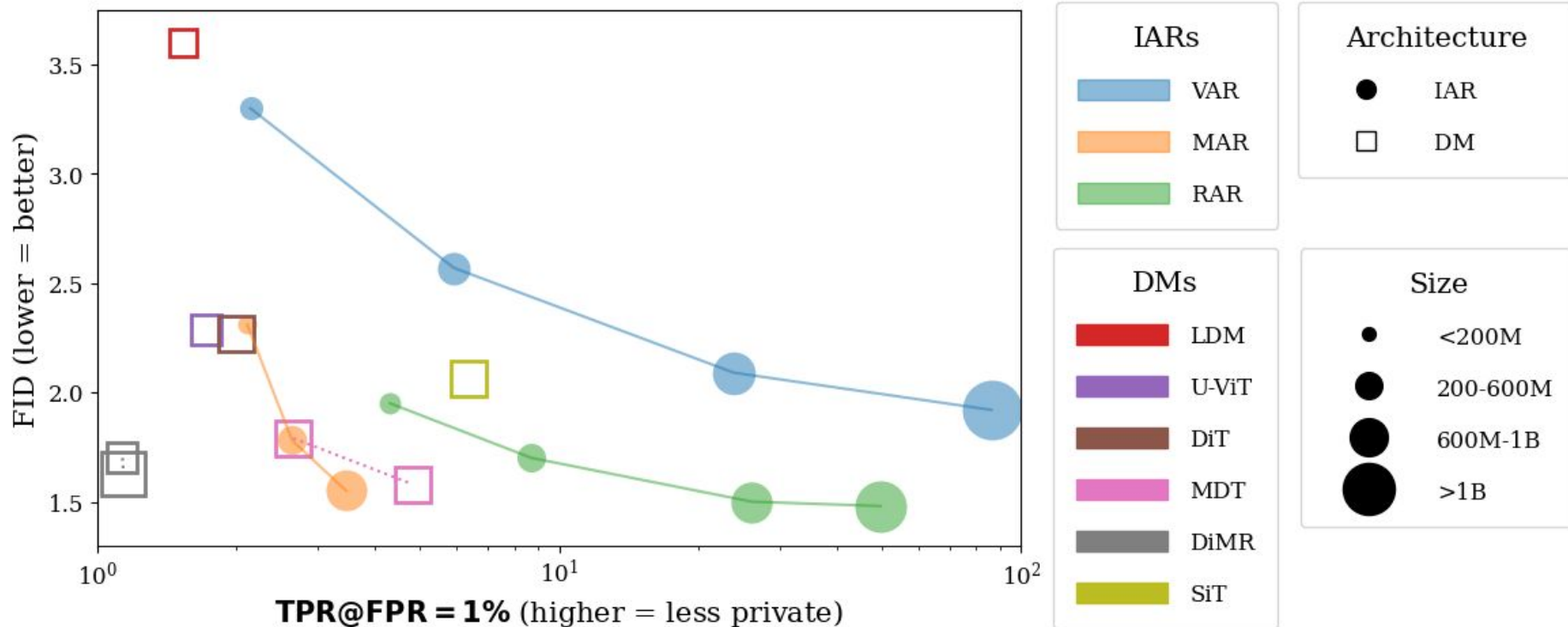
Our improvement

	VAR-<i>d</i>30	RAR-XXL	MAR-H
Baseline TPR@FPR=1%	16.68	14.62	2.18

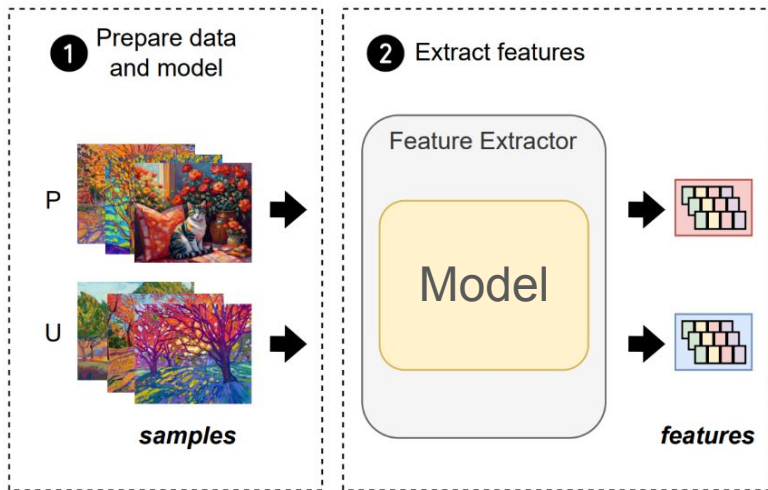
Our improvement

	VAR-<i>d</i>30	RAR-XXL	MAR-H
Baseline TPR@FPR=1%	16.68	14.62	2.18
Ours TPR@FPR=1%	86.38	49.80	3.40

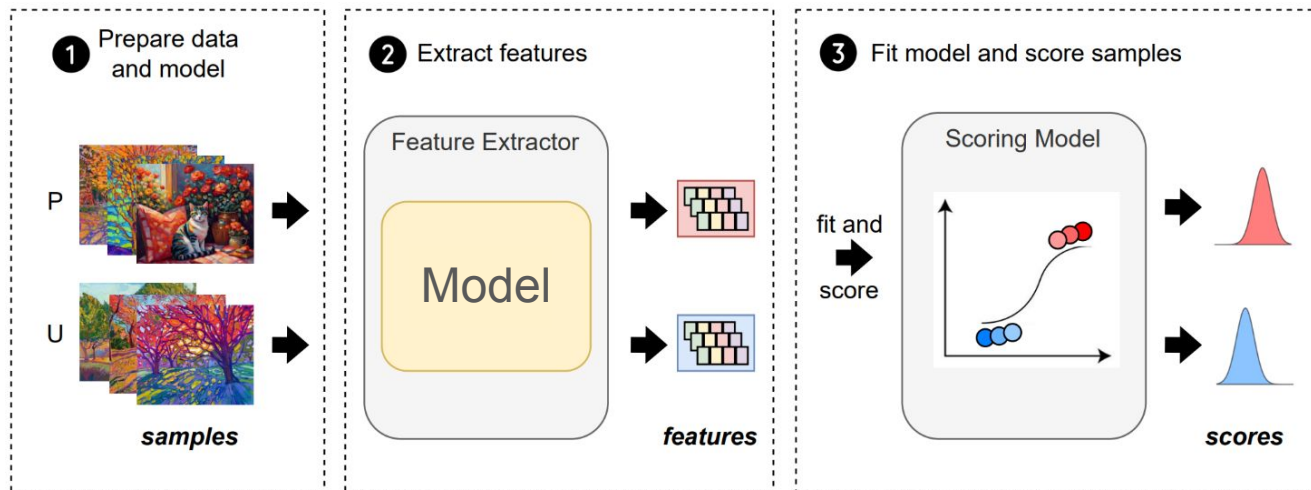
IARs leak more than DMs



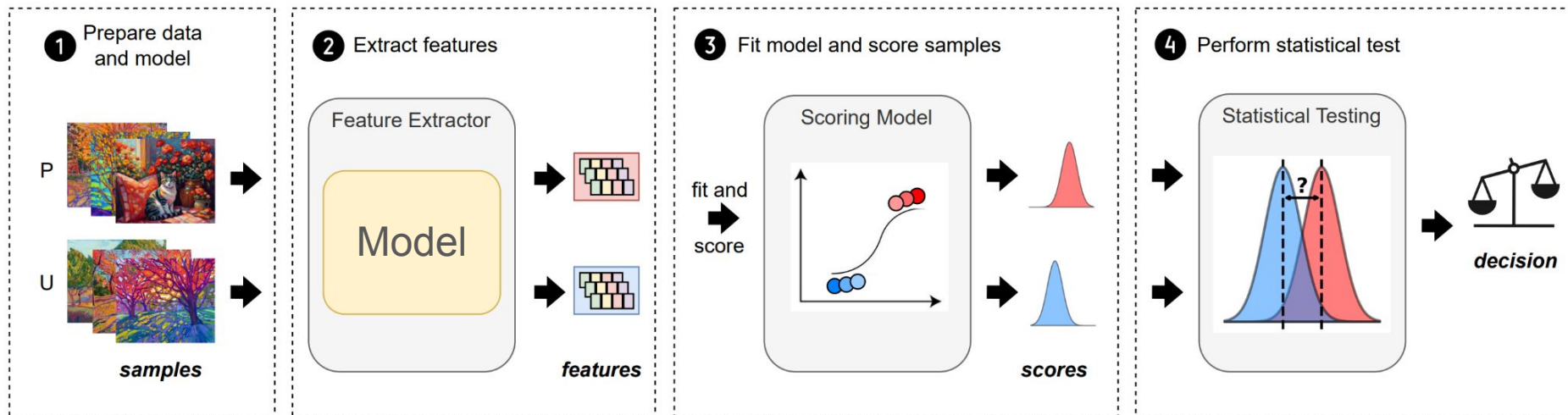
Dataset Inference (DI)



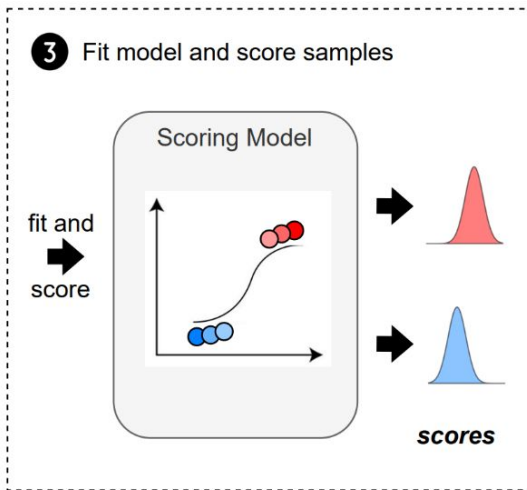
Original DI



Original DI

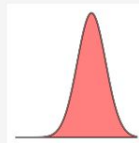


DI



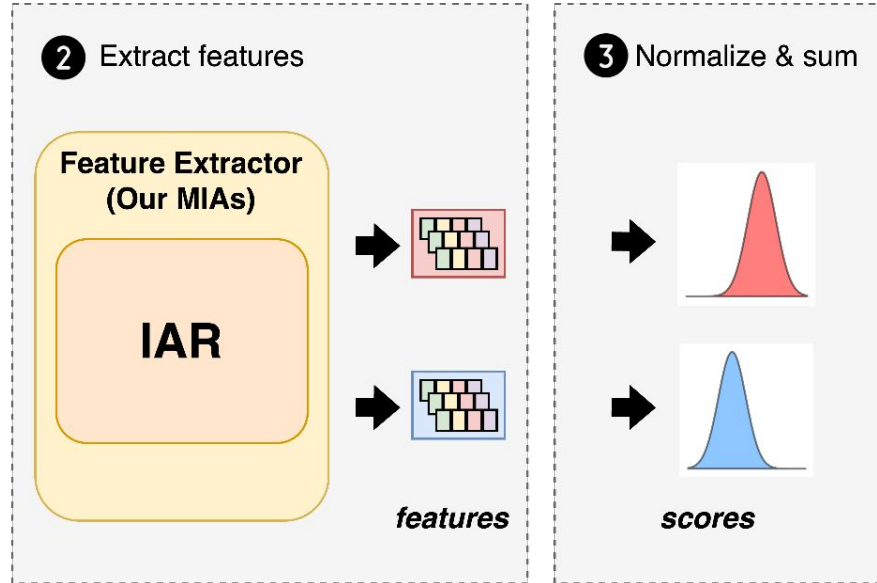
Our DI

3 Normalize & sum



scores

Our DI



Our DI

1 Prepare data and model



samples

2 Extract features

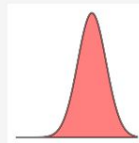
Feature Extractor
(Our MIAs)

IAR



features

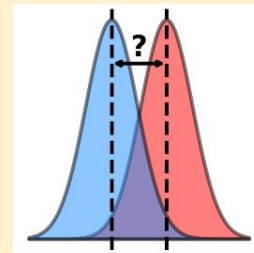
3 Normalize & sum



scores

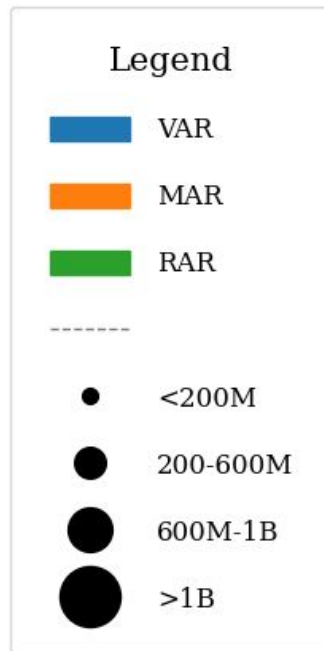
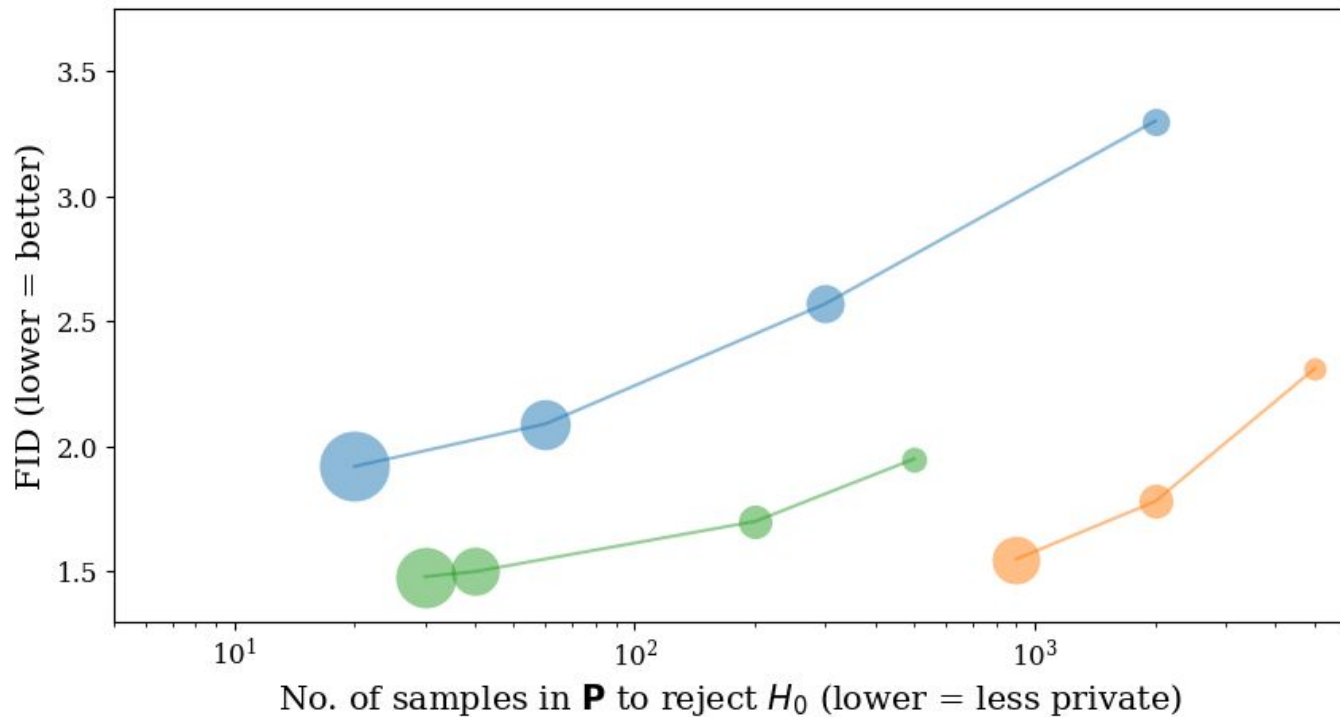
4 Perform statistical test

Statistical Testing

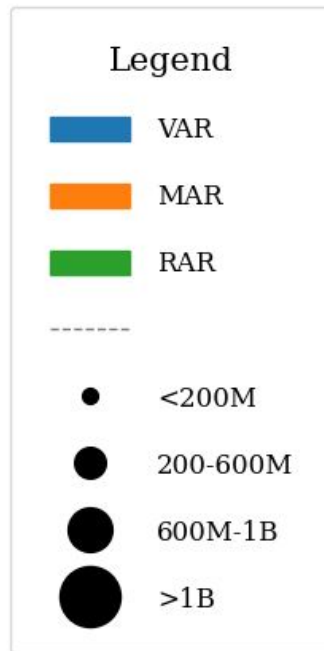
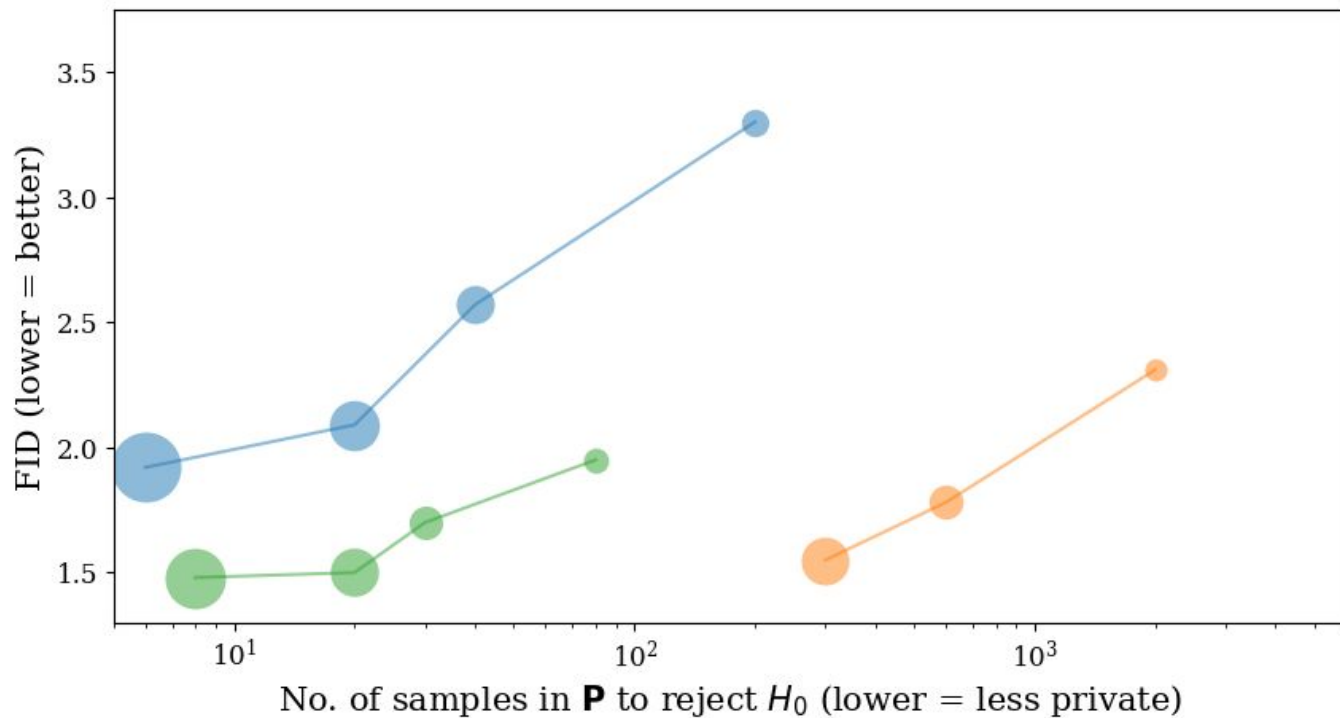


decision

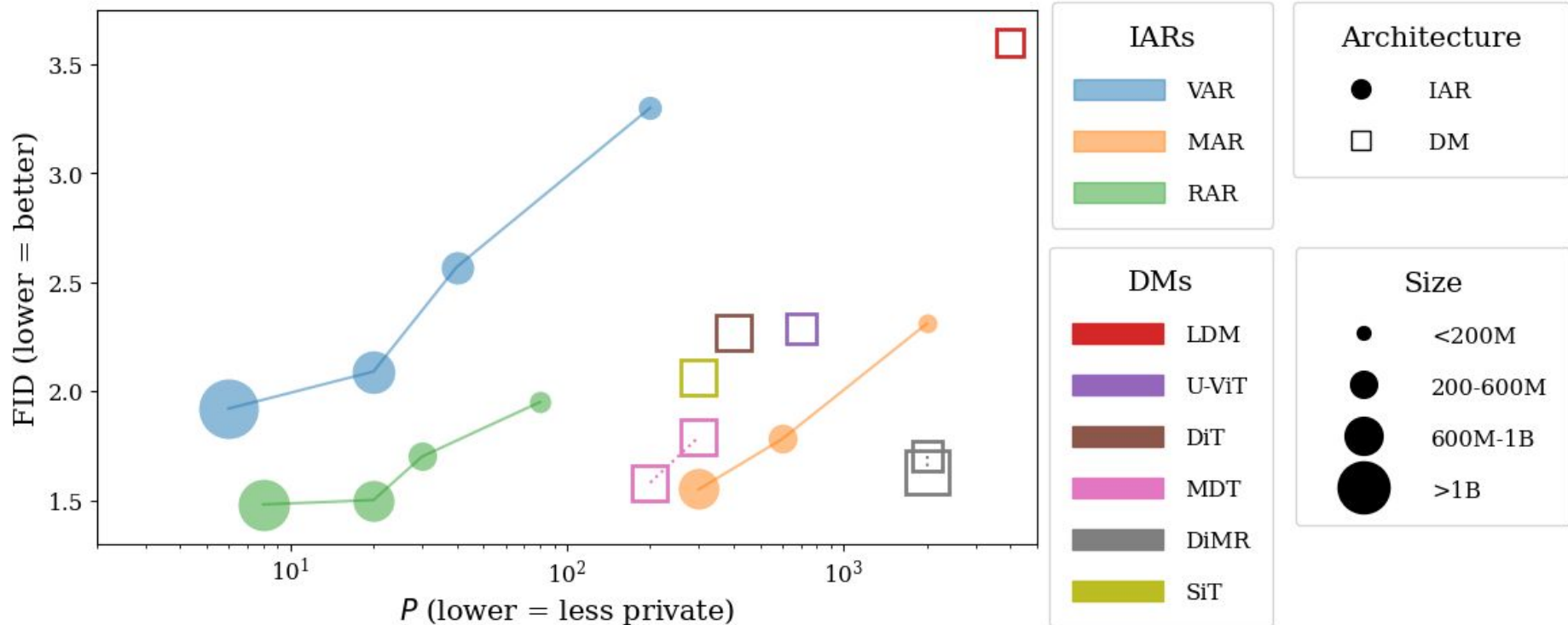
DI: before



DI: after



DMs are less prone to DI



Data Extraction

Training Sample

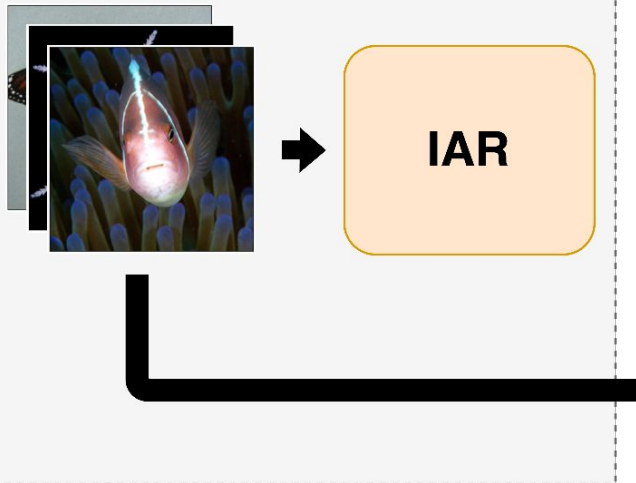


Generated Image

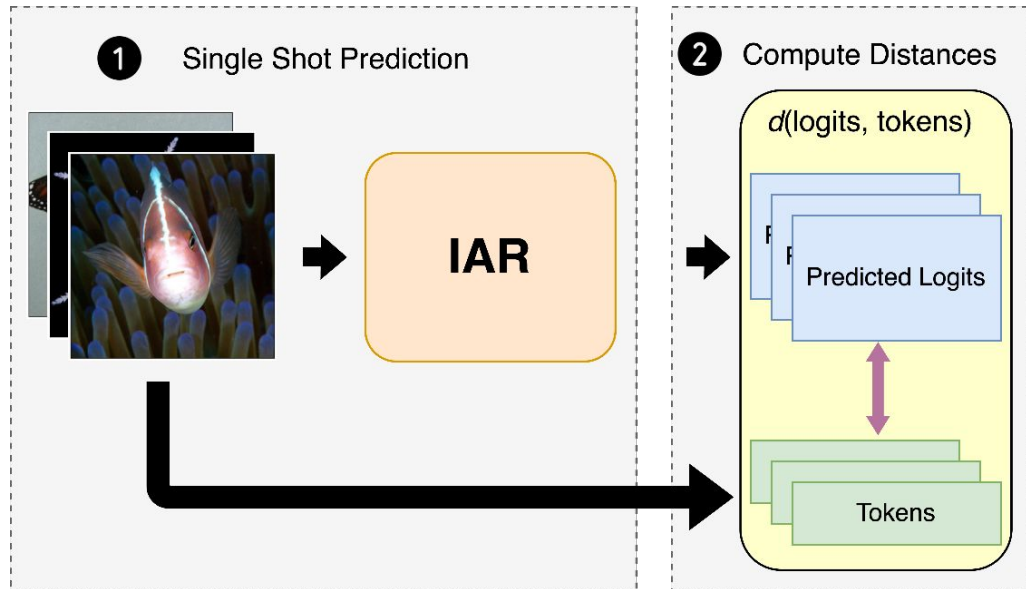


Our Data Extraction

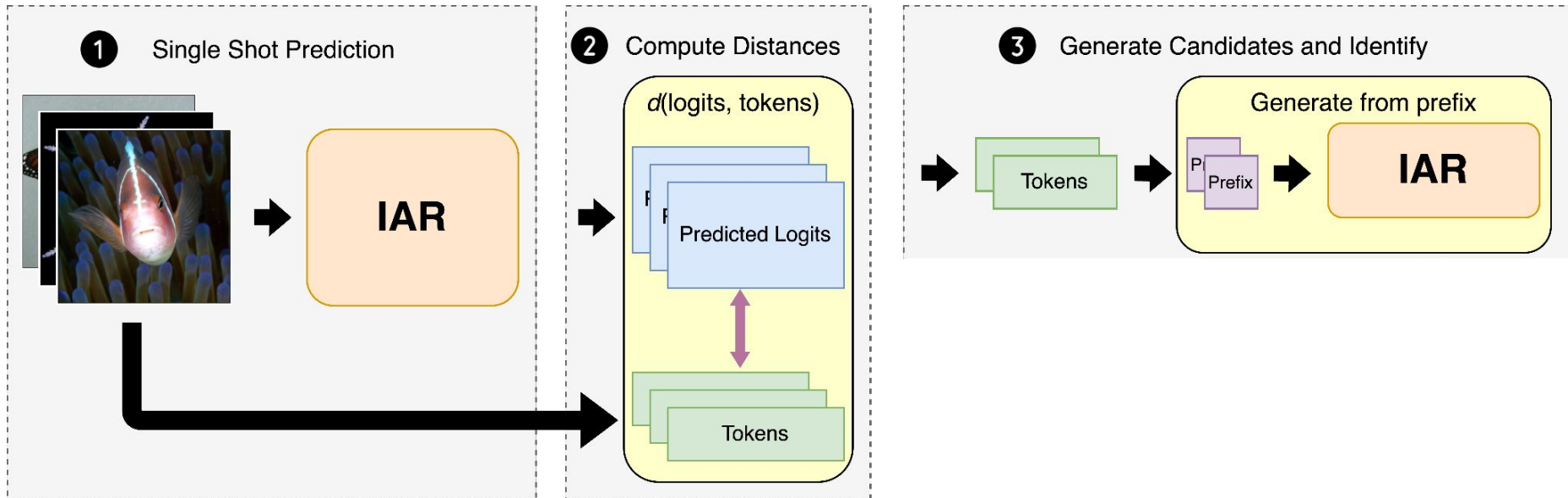
1 Single Shot Prediction



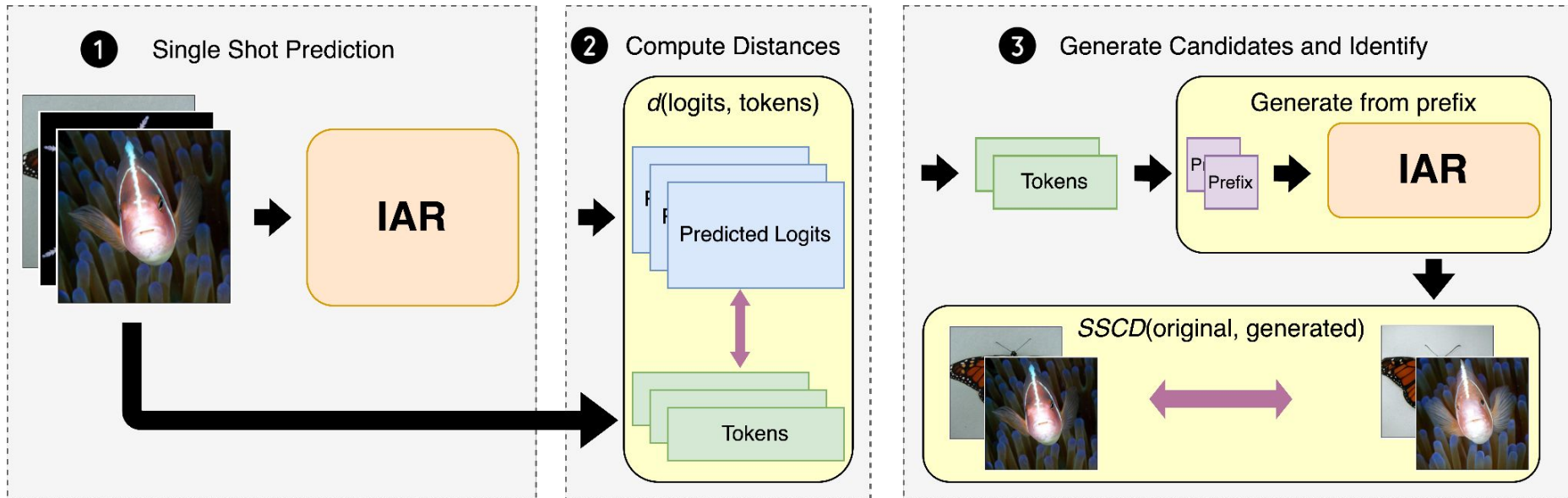
Our Data Extraction



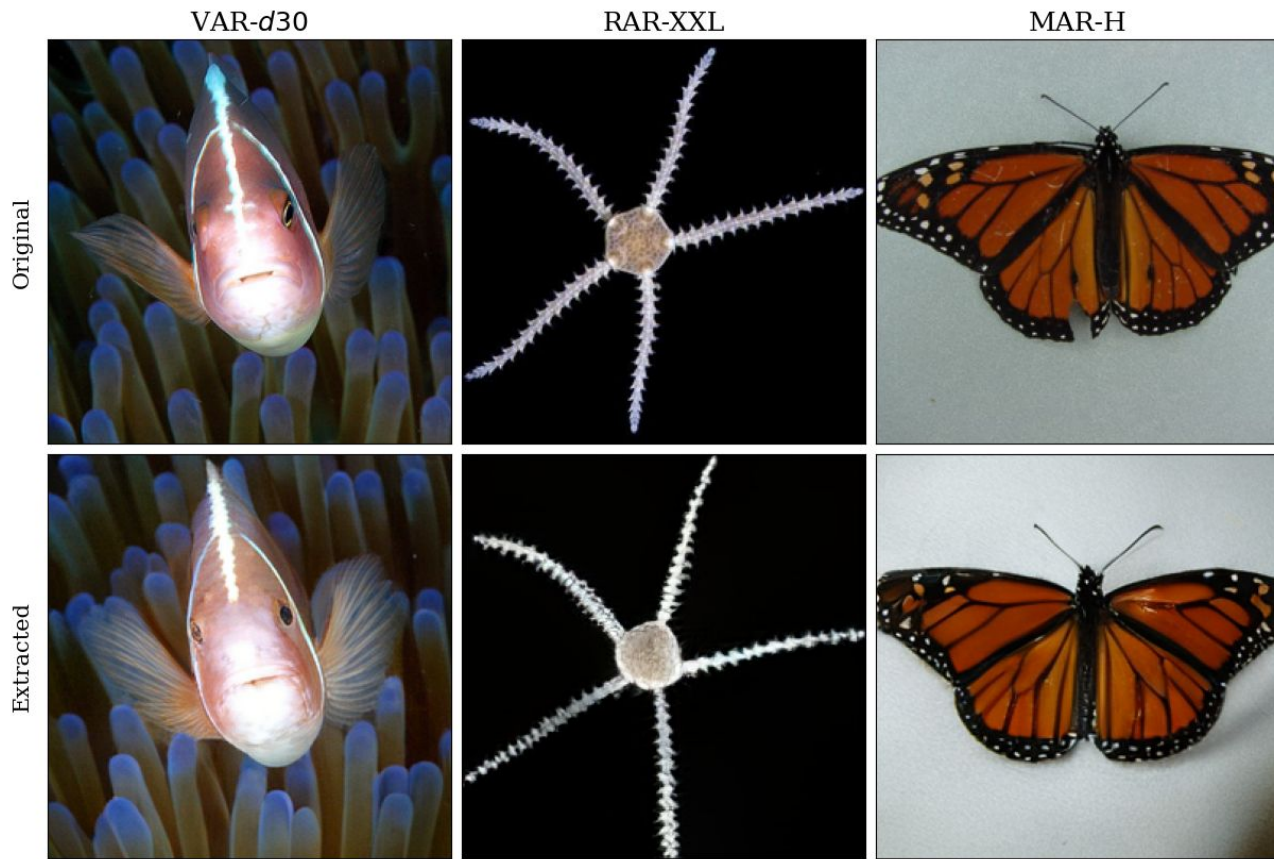
Our Data Extraction



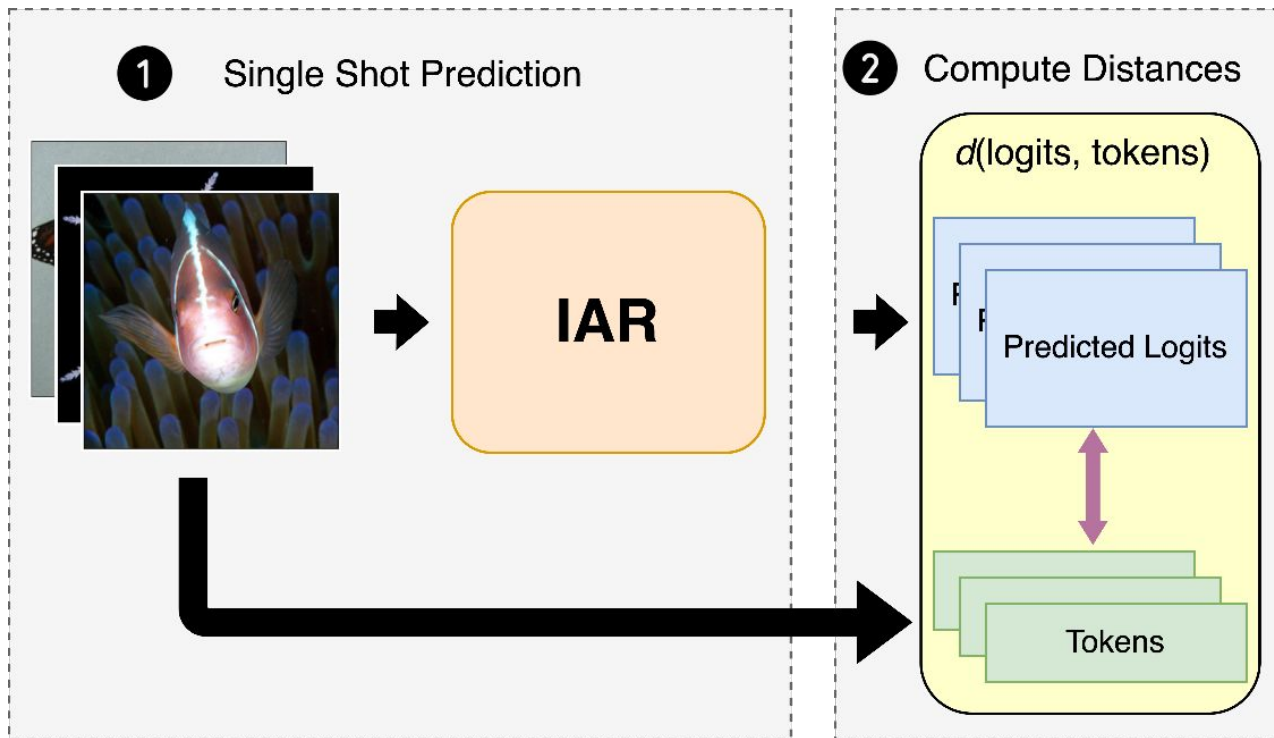
Our Data Extraction



Successful Extraction



Efficient Candidate Selection



Why do IARs leak
more?

Reasons for Higher Leakage

IARs	DMs
Access to $p(x)$	No access to $p(x)$

Reasons for Higher Leakage

IARs	DMs
Access to $p(x)$	No access to $p(x)$
More data per sample	Single noisy image

Reasons for Higher Leakage

IARs	DMs
Access to $p(x)$	No access to $p(x)$
More data per sample	Single noisy image
Each token leaks	Single leakage

Contributions

1. **First** to investigate IARs' privacy

Contributions

1. **First** to investigate IARs' privacy
2. New, **SOTA** privacy attacks against IARs

Contributions

1. **First** to investigate IARs' privacy
2. New, **SOTA** privacy attacks against IARs
3. **First** successful data extraction from IARs

Contributions

1. **First** to investigate IARs' privacy
2. New, **SOTA** privacy attacks against IARs
3. **First** successful data extraction from IARs
4. Explanation **why** IARs leak more than DMs

Thank you!

