



Knockoff Nets: Stealing Functionality of Black-Box Models

Tribhuvanesh Orekondy¹

Bernt Schiele¹

Mario Fritz²

¹ Max Planck Institute for Informatics

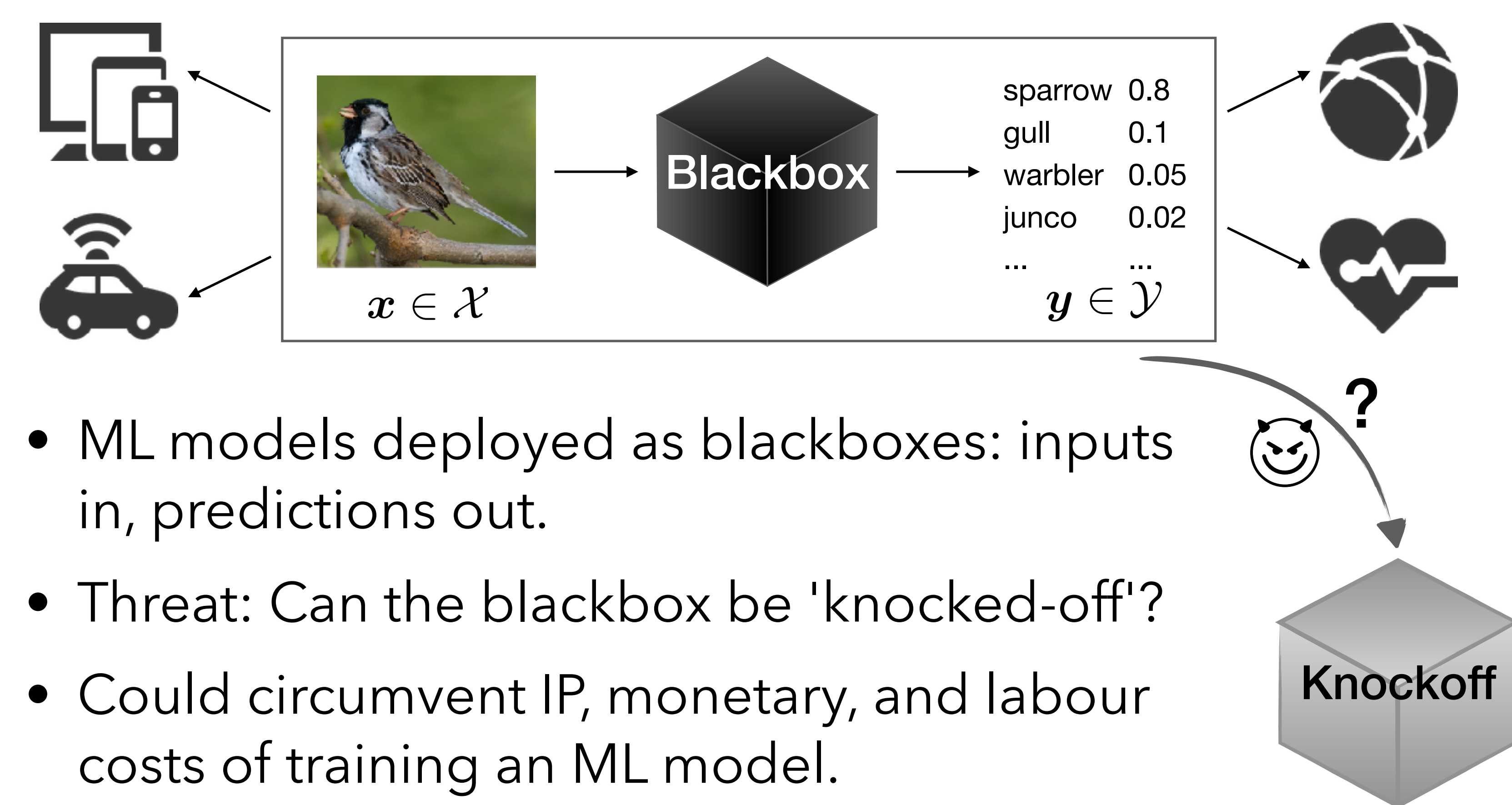
² CISPA Helmholtz Centre for Information Security

Saarland Informatics Campus, Germany



LONG BEACH
CALIFORNIA
June 16-20, 2019

Motivation

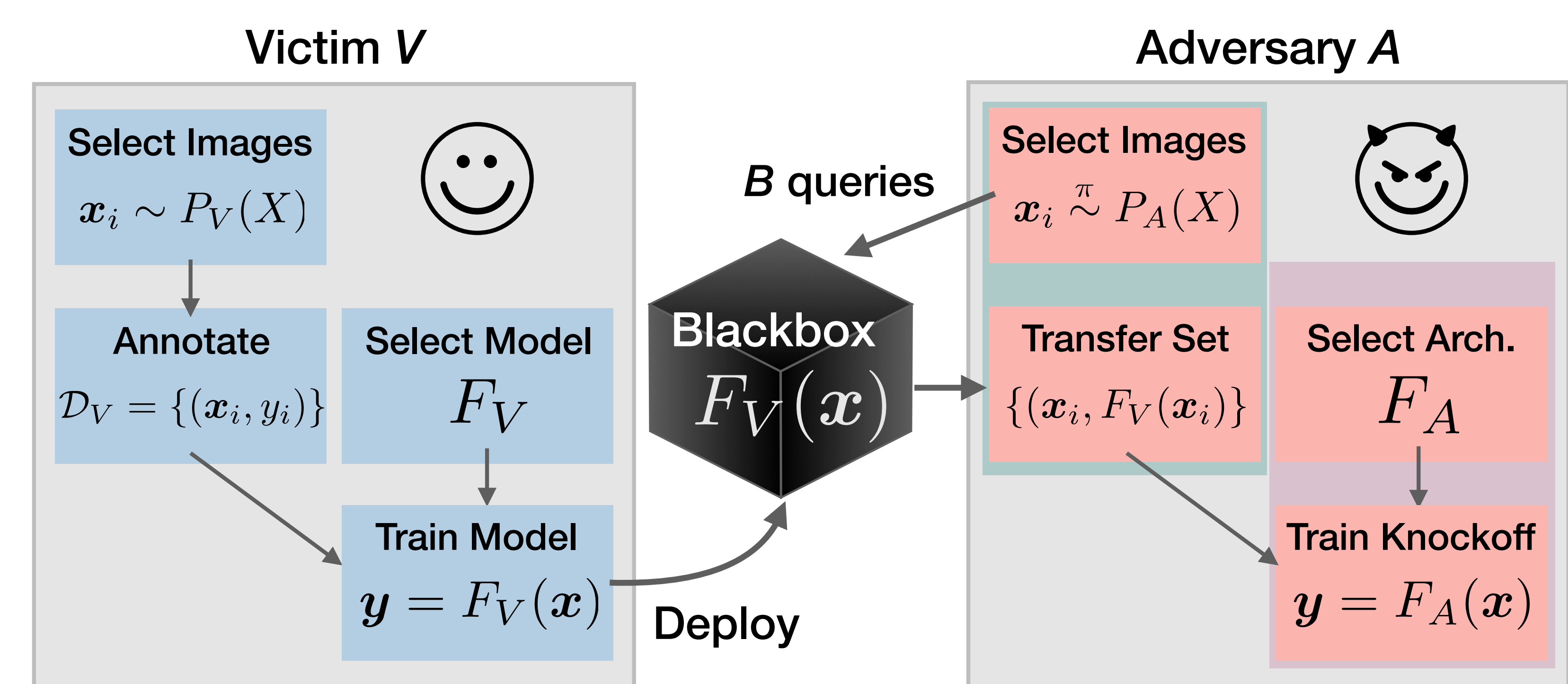


Problem: Model Functionality Stealing

- Data-limited adversary.
- Two-player game.

Adversary's objective:

- Steal functionality of blackbox (when F_V and P_V are unknown).
- Using minimum # of queries B .



Generating Knockoffs

Transfer set construction: Which images to query?

Method 1: π = Random

- Samples images randomly (without replacement).
- But: might query irrelevant images.

Training knockoff F_A : Which model to train?

- Any sufficiently complex architecture.

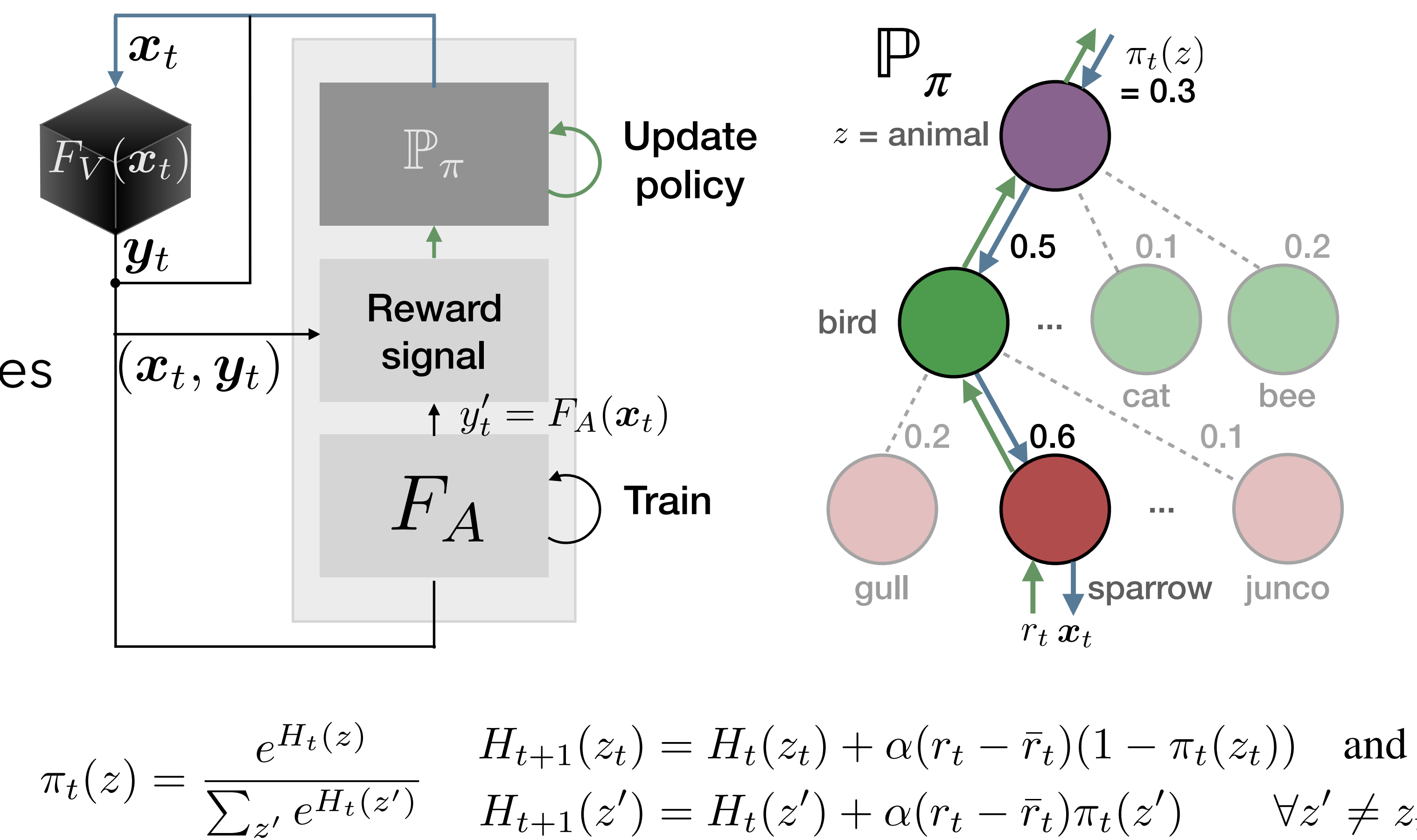
Method 2: π = Adaptive

- Motivation: (a) Improve sample efficiency (b) Interpretable policy π .
- Approach: Multi-armed hierarchical bandits.
- Actions (arms): supplemented labels of images in P_A (e.g., 1K classes for Imagenet).
- Rewards

a. Certainty: $R^{\text{cert}}(y_t) = P(y_{t,k_1}|x_t) - P(y_{t,k_2}|x_t)$

b. Diversity: $R^{\text{div}}(y_{1:t}) = \sum_k \max(0, y_{t,k} - \bar{y}_{t:t-\Delta,k})$

c. Loss: $R^{\mathcal{L}}(y_t, \hat{y}_t) = \mathcal{L}(y_t, \hat{y}_t)$

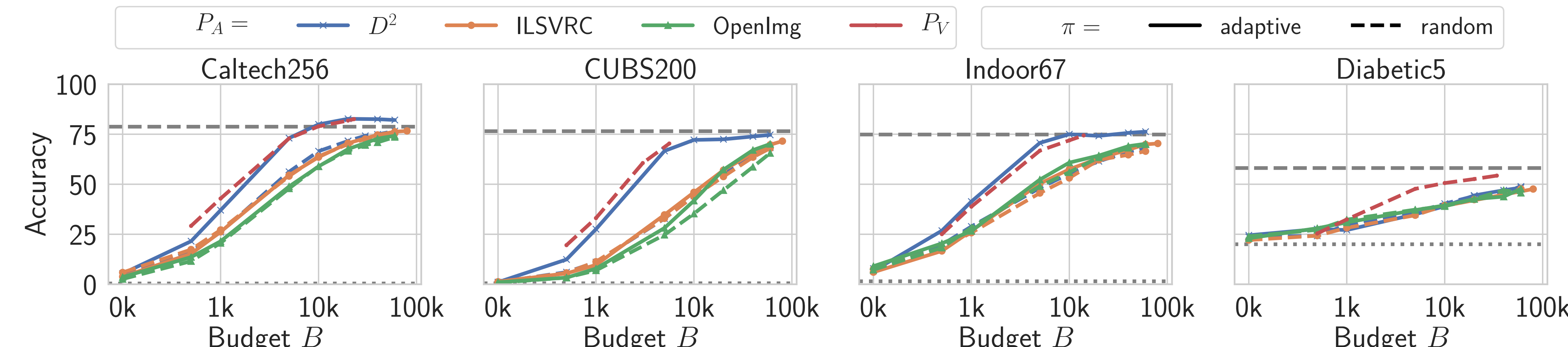


Experimental Results

- Four victim blackboxes (Resnet-34s): Caltech256, CUB-200, Indoor-Scene, Diabetic Retinopathy.

- Four choices of P_A : D^2 , ImageNet, OpenImages, P_V (KD).

Key Results



- Knockoffs are effective; recovers 0.81-0.97x victim blackbox's accuracy.
- Knockoffs also learn to classify images from unseen classes at test-time e.g., >90% classes in CUB-200.

Transfer Set $\{(x_i, F_V(x_i))\}$, $x_i \sim P_A(X)$



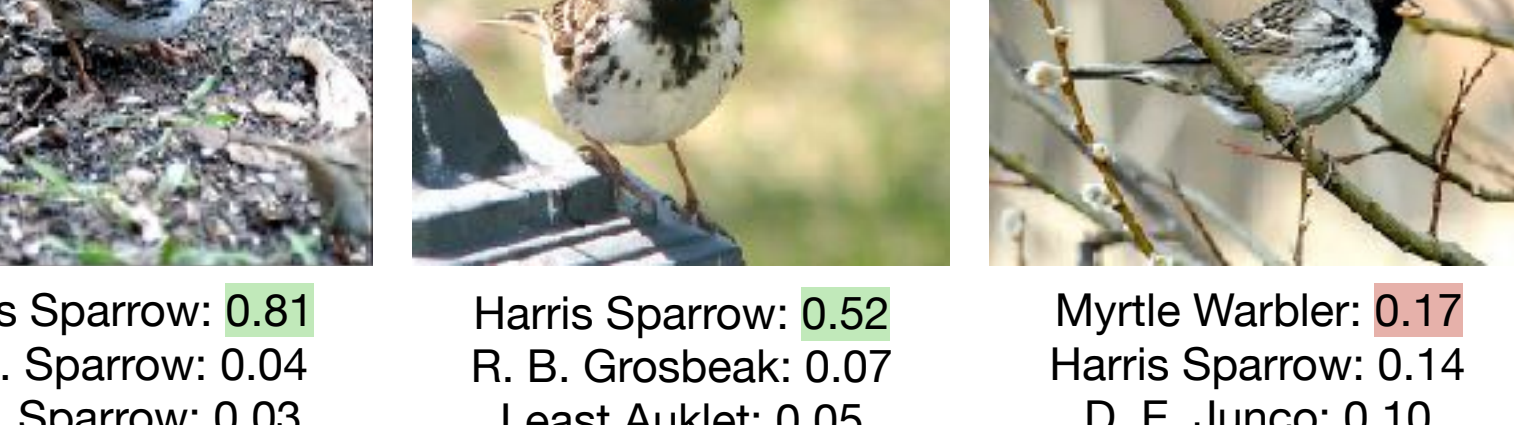
Test Set $\{(x_i, F_A(x_i))\}$, $x_i \sim D_V^{\text{test}}$



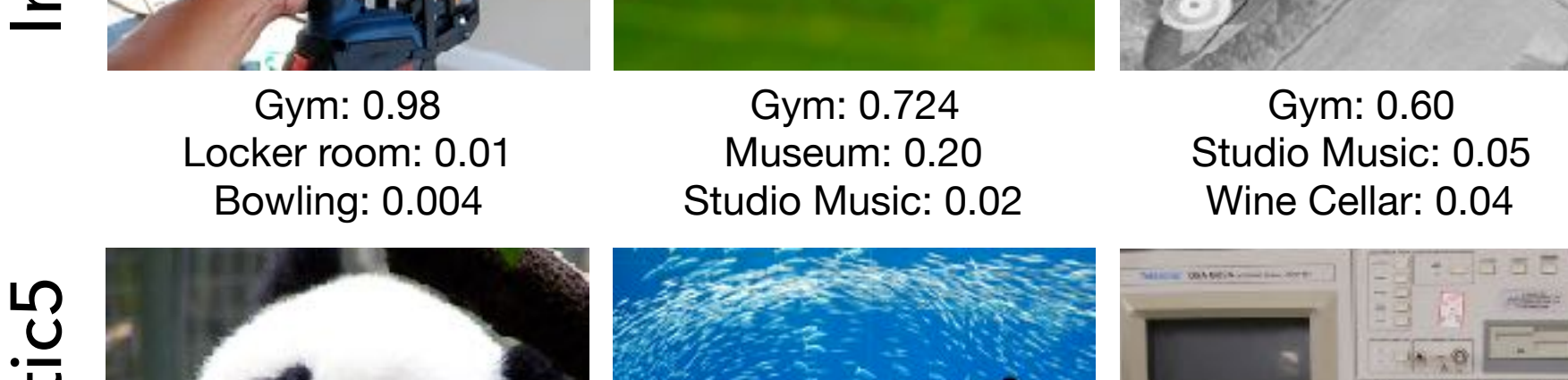
CUBS200



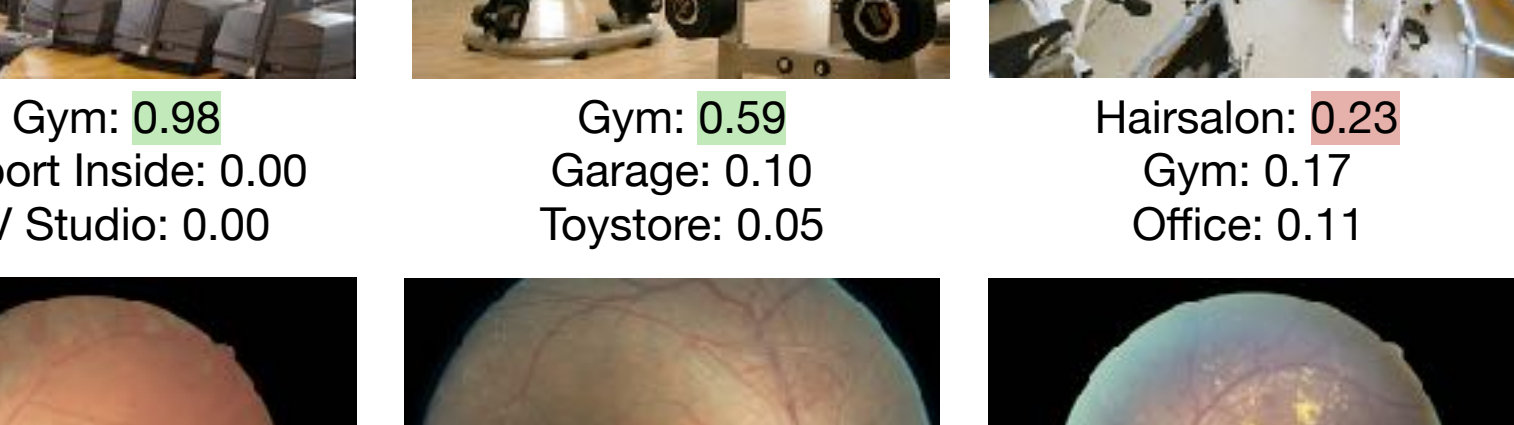
Test Set



Indoor67



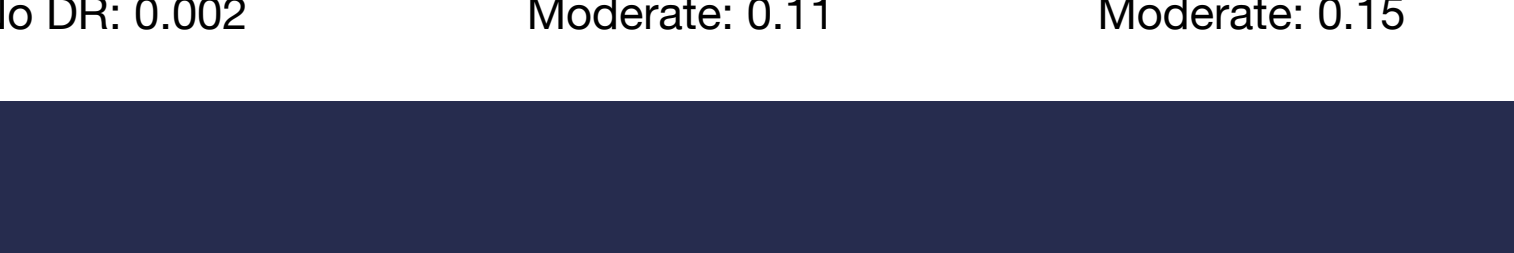
Test Set



Diabetic5

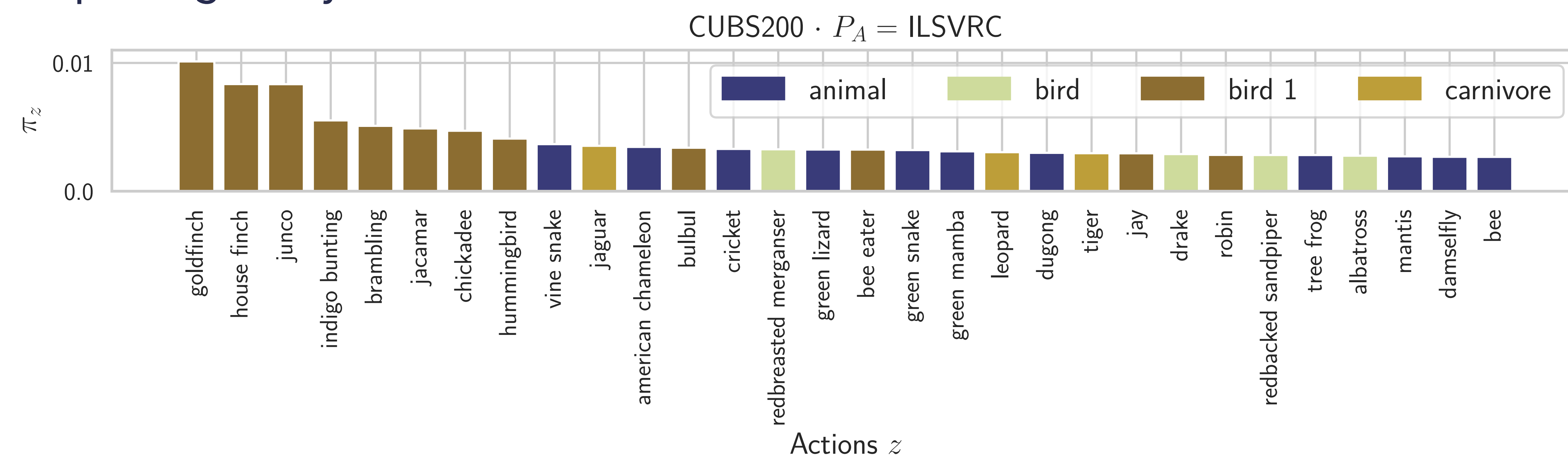


Test Set



- Adaptive improves sample-efficiency in closed-world settings ($P_A = D^2$). Difficult to significantly outperform random policy in open-world (ImageNet, OpenImages).
- Adaptive also improves performance of knockoffs (up to 4.5% test accuracy).

Inspecting Policy



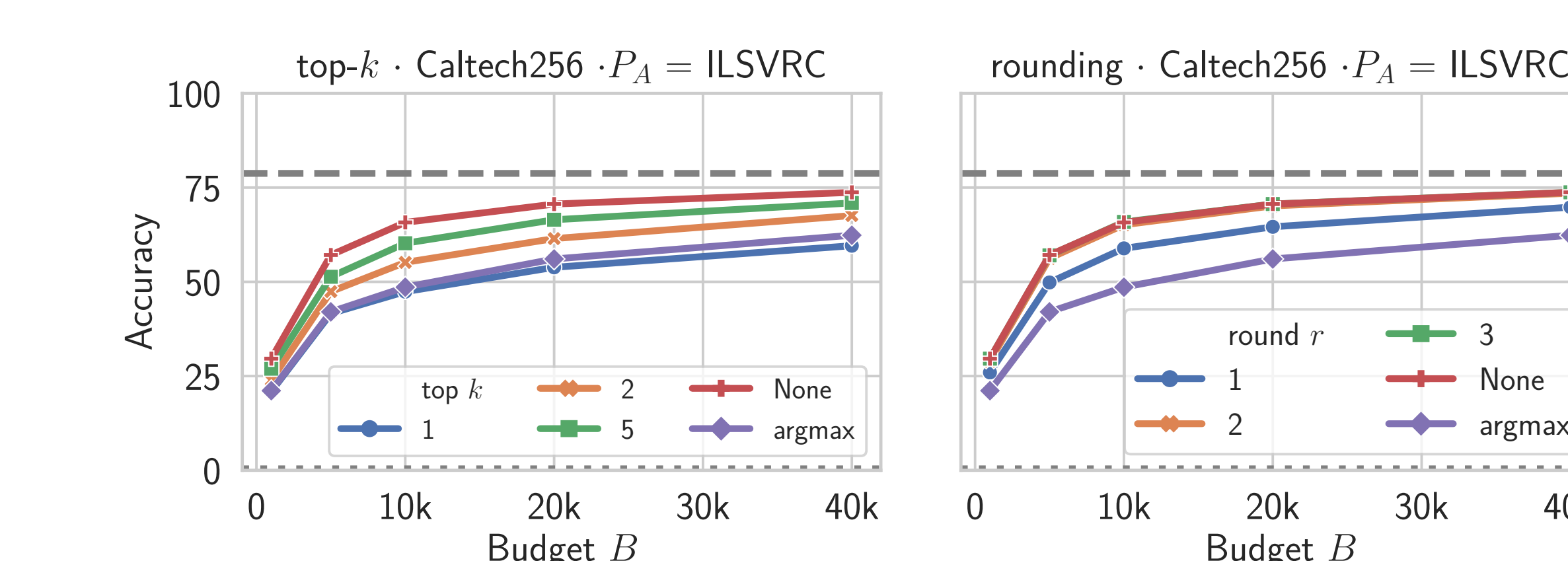
- Actions (arms of bandit) typically correspond to classes closely related to victim blackbox's output classes.

Acknowledgment This research was partially supported by the German Research Foundation (DFG CRC 1223)

Take-aways

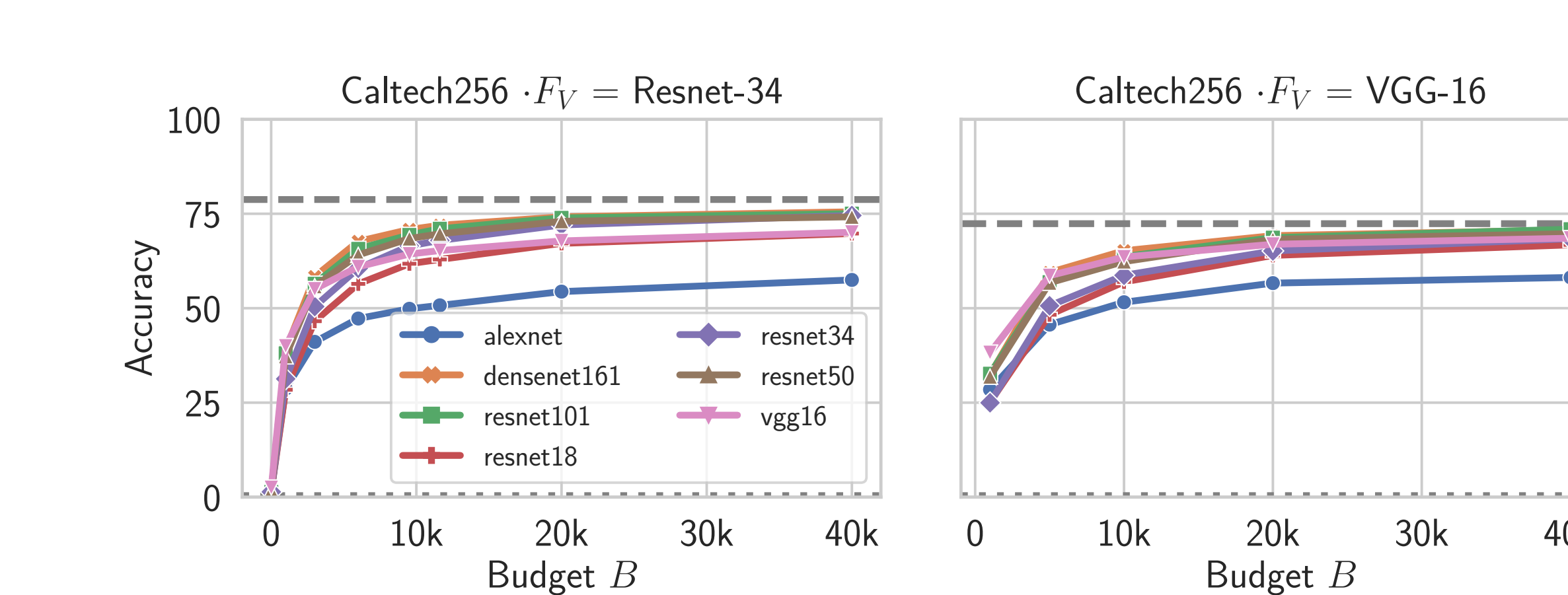
- Model functionality of complex blackbox CNNs can be stolen under minimal assumptions.
- Attacks are surprisingly effective and difficult to defend.
- Poses a threat to deployed ML APIs.

Victim's Defense: Information Truncation



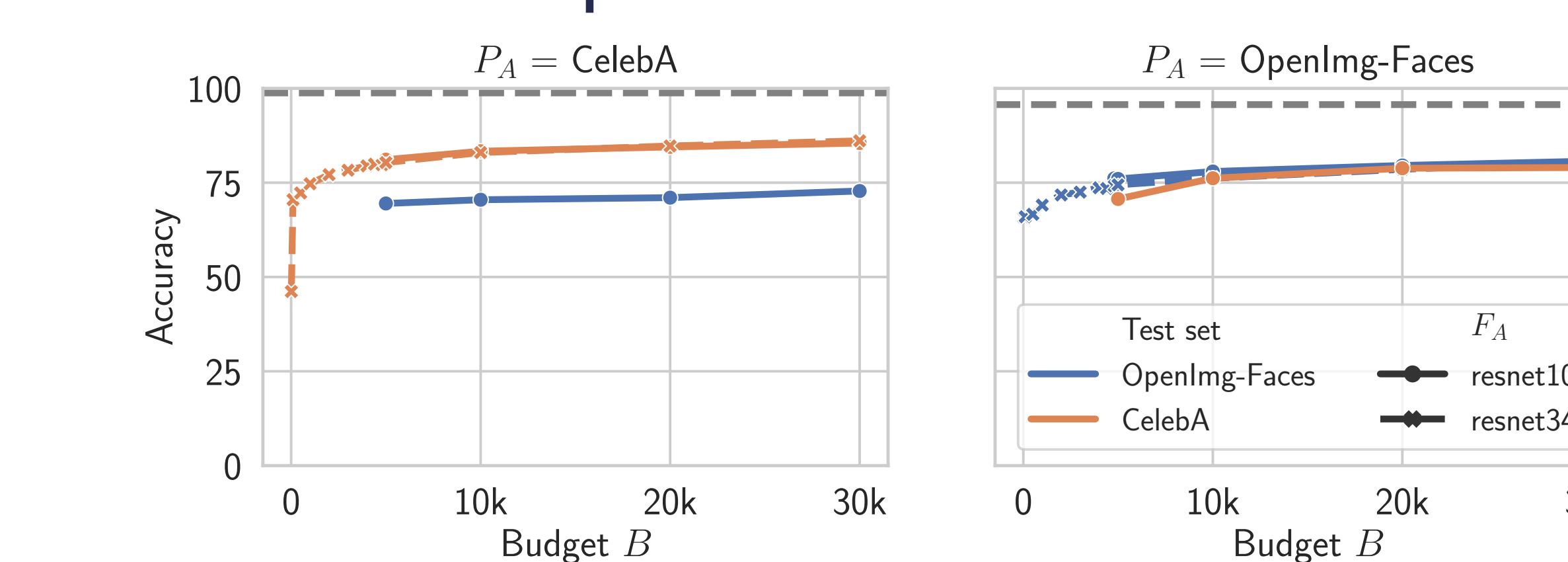
- Limited effect on knockoff attacks.

Vs. Choices of Architectures



- Robust to choices of F_V and F_A .

Real-world Experiment



- Knockoffs could be trained using \$30 worth of queries.