

Compositional Non-Face Re-Identification Pressure under Cumulative Vision Releases

Tirth Joshi  and Honggang Wang* 

The Katz School of Science and Health
Yeshiva University, New York NY 10016, USA
tjoshi1@main.yu.edu honggang.wang@yu.edu

Abstract. Most privacy evaluations in computer vision treat datasets and models as isolated, static disclosures, failing to account for how risk accumulates in practice. We formalize **Compositional Non-Face Re-Identification Risk**, where identities remain linkable through non-face cues, such as clothing, gait, and scene context that aggregate across sequential releases of data expansions, model checkpoints, and meta-data. We introduce the **Vision Re-Identification Pressure Index** ($vRPI_\alpha$), a foundational, information-theoretic measure derived from **Arimoto conditional Rényi entropy** that quantifies the systemic erosion of anonymity. We prove that $vRPI_\alpha$ is monotone under cumulative releases, establish an exact bridge to **Bayes-optimal guessing probability**, and derive additive decompositions under conditional-independence benchmarks. Through the **NFLeak protocol**, we provide illustrative experiments on **Market-1501-Tau** and a **CUHK03-Tau** closed-world replication using a structured **Attacker Ladder**, together with masking-level, dependence-overlap, temperature-calibration, and non-uniform-prior diagnostics, demonstrating that $vRPI_\alpha$ accurately tracks realized re-identification success even after explicit face removal. Our framework enables principled auditing of the “leaky bucket” effect in vision benchmarks, shifting the focus from single-release compliance to ecosystem-level privacy hygiene.

Keywords: Vision Privacy · Re-Identification · Rényi Entropy · Composition · Information Leakage

1 Introduction

1.1 Problem: why “single-release privacy” is flawed

Privacy evaluations in computer vision are often implicitly *single-shot*: a dataset or model is released once, and risk is assessed in isolation. This conflicts with how vision artifacts are actually published, reused, and progressively enriched.

* Corresponding author

Release sequences as cumulative side information. In practice, releases occur in sequence: additional data coverage, pretrained checkpoints/APIs, and, sometimes, released embeddings or indices. Each release provides extra side information about a latent identity variable. Formally, if $U \in \{1, \dots, N\}$ indexes identity and $\tilde{X} = \tau(X)$ is a non-face observation, then the t -th release can be modeled as a channel

$$Z_t \sim K_t(\cdot \mid \tilde{X}), \quad Z_{1:T} := (Z_1, \dots, Z_T). \quad (1)$$

Risk is governed by the posterior $\Pr(U \mid Z_{1:T})$, which typically sharpens as T grows. An operational target is the Bayes-optimal guessing probability

$$P_{\text{guess}}(U \mid Z_{1:T}) := \mathbb{E}_{z_{1:T}} \left[\max_u \Pr(U = u \mid z_{1:T}) \right], \quad (2)$$

with its exact min-entropy duality $H_{\min}(U \mid Z_{1:T}) = -\log P_{\text{guess}}(U \mid Z_{1:T})$ [4, 5]. The single-release view corresponds to $P_{\text{guess}}(U \mid Z_t)$; the realistic setting is compositional $P_{\text{guess}}(U \mid Z_{1:T})$.

Non-face leakage. A common mitigation is to remove faces (blur/mask), but identity remains encoded in non-face cues (clothing, body shape, pose/gait, carried objects, scene context). Thus, even after τ , the posterior $\Pr(U \mid Z_{1:T})$ can become highly concentrated under cumulative releases, enabling linkage without facial information.

Empirical preview. We validate this compositional perspective on MARKET-1501-TAU, a face-masked variant of Market-1501 constructed by masking the top 20% of each image, and a CUHK03-TAU closed-world replication, simulating a four-stage release sequence with a structured attacker ladder and a suite of diagnostics (masking level, dependence overlap, temperature calibration, and non-uniform priors).

1.2 Contributions

We present a theory-first framework for compositional non-face identity leakage in vision:

- **Release-as-channel formulation.** We model each disclosure as a channel K_t producing side information Z_t and analyze cumulative releases via $Z_{1:T}$.
- **Rényi pressure index (vRPI $_{\alpha}$).** For $\alpha > 1$, we define

$$\text{vRPI}_{\alpha}(T) := \log N - H_{\alpha}(U \mid Z_{1:T}),$$

where H_{α} is Arimoto conditional Rényi entropy [1–3]. This yields an entropy-based “pressure” toward re-identification under release sequences.

- **Core theorems.** We prove (i) monotonicity under additional releases: $\text{vRPI}_{\alpha}(T)$ is non-decreasing in T , and (ii) an operational guessing bound linking $H_{\alpha}(U \mid Z)$ to $P_{\text{guess}}(U \mid Z)$.

- **Conditional-independence decomposition.** Under an explicit conditional-independence model for releases given identity, we derive an additive Rényi-information decomposition and characterize failure modes under dependence (correlated releases, duplicates, side channels).
- **Illustrative evaluation and diagnostics.** We instantiate releases and attacker posteriors on MARKET-1501-TAU and add a CUHK03-TAU closed-world replication, showing that vRPI_α tracks realized non-face inference risk across release stages. We further report masking-level ablations, a dependence-overlap diagnostic for the additivity benchmark, temperature-calibration sensitivity, and non-uniform-prior (Zipf) diagnostics, and position vRPI_α against existing leakage measures.

1.3 Paper roadmap

Section 3 formalizes identities, non-face transforms, release channels, and attacker posteriors. Section 4 defines vRPI_α and its operational interpretation. Section 5 proves monotonicity, guessing bounds, and conditional-independence decompositions, with explicit assumptions and failure modes. Section 6 reports the MARKET-1501-TAU instantiation and diagnostic results. Section 7 discusses limitations, ethics, and implications for release auditing.

2 Related Work

2.1 Person re-identification and non-face cues

Re-ID benchmarks and pipelines. Person re-identification (Re-ID) studies identity matching across cameras and time. Market-1501 [6] and DukeMTMC4ReID [8] are canonical image benchmarks; MARS [9] extends to video tracklets, and PRW couples detection with Re-ID in full frames [10]. Modern Re-ID pipelines learn embeddings and perform retrieval, typically trained with classification and/or metric learning objectives; triplet-loss baselines remain standard [12], with strong backbones and training heuristics improving performance [13, 14], and re-ranking is often used at inference [15]. Surveys formalize closed-world vs. open-world distinctions and system components [11].

Non-face cues. Even when faces are not explicitly used, identity inference can rely strongly on non-face cues such as clothing appearance, body shape, carried objects, pose/gait, and scene context. Clothes-changing and long-term settings make this dependence explicit [16, 17], while gait recognition demonstrates separable identity signal from motion dynamics [18]. These findings motivate our focus: after face/head removal, a substantial identity signal can remain, and the relevant question becomes how such a residual signal compounds under cumulative releases.

2.2 Privacy in vision and anonymization

Face/full-body de-identification and adversarial protection. A major line of work attempts to de-identify faces while preserving utility, motivated by early evidence that naive obfuscation can fail and by formal schemes such as k -Same [19, 20]. Deep generative approaches produce higher-fidelity anonymization [21, 22] and have been extended toward more realistic full-body anonymization [23]. Separately, protective perturbations aim to reduce unauthorized recognition [24], and evaluation work studies interactions with detectors and downstream training [25].

Why composition is under-modeled. Most vision anonymization evaluations are effectively single-release: privacy is measured after releasing a particular transformed dataset or model under a fixed attacker. However, practical deployments and benchmark ecosystems evolve through sequences of releases (expanded data coverage, checkpoints/APIs, embeddings/indices, metadata), where risk is governed by posterior concentration under accumulated side information. This gap motivates formal, release-sequence-aware leakage measures beyond ad hoc success rates [26].

2.3 Information-theoretic privacy metrics

Guessing probability and min-entropy. Operational privacy can be defined in terms of an adversary’s success at inferring private attributes. In particular, Bayes-optimal guessing probability is exactly dual to conditional min-entropy, with foundational interpretations in information theory and cryptography [4, 5]. This connection motivates our use of entropy-style quantities to control identity inference risk.

Rényi families and leakage measures. Rényi entropies generalize Shannon entropy and tune sensitivity to posterior concentration [27]. Arimoto introduced order- α information measures for channels [1], with conditional Rényi entropies and Bayesian M -ary testing interpretations developed in subsequent work [2, 3]. Sibson mutual information provides an order- α family with strong data-processing and product-channel properties [28]. Operational leakage measures relate these quantities to inference, including maximal leakage [29] and maximal α -leakage [30]; surveys unify perspectives across estimation- and information-theoretic views [31]. Our experiments focus on *closed-world identity inference* for theoretical alignment, and are used to illustrate (not exhaustively benchmark) the proposed pressure index.

Positioning. We do not propose a new Re-ID model or a new anonymizer. Instead, we formalize *compositional non-face identity leakage* under release sequences, define an entropy-based pressure index using conditional Rényi entropy, and prove monotonicity and operational guessing bounds, with illustrative validation in a Re-ID setting.

Why vision/Re-ID specifically. While the release-as-channel abstraction is general, our instantiation is vision-specific: the side information consists of concrete CV/Re-ID artifacts, body crops, learned encoders, generic ImageNet features, retrieval indices, prototype banks, and k NN graphs, and the leakage source is visual, namely clothing, body shape, pose/gait, carried objects, and scene context that persist after face/head-region removal. Thus vRPI_α formalizes how standard vision-benchmark practice creates *sequential* visual side information, rather than generic release leakage. Operationally, vRPI_α is not a replacement for rank/mAP: rank/mAP evaluates retrieval utility, whereas vRPI_α audits the privacy pressure induced by the release itself, making the two complementary release-auditing criteria.

3 Problem Formulation

3.1 Identities, observations, and non-face transform τ

Let $U \in \mathcal{U} := \{1, \dots, N\}$ denote an identity and $X \in \mathcal{X}$ a raw visual observation with

$$X \sim P_{X|U}(\cdot | U). \quad (3)$$

We fix a non-face transformation $\tau : \mathcal{X} \rightarrow \tilde{\mathcal{X}}$ and define the observed non-face view

$$\tilde{X} := \tilde{X} = \tau(X) \in \tilde{\mathcal{X}}. \quad (4)$$

All privacy quantities are defined with respect to the joint law of $(U, \tilde{X})$.

Experimental instantiation. In Section 6, we use a closed-world universe $\mathcal{U}$ equal to the $N = 751$ training identities of Market-1501, and τ masks the top 20% of each image (MARKET-1501-TAU).

Remark 1 (Open-world note). If queries may contain identities outside $\mathcal{U}$ (open-set), the posterior must include an “unknown” option; our entropy definitions remain meaningful on the restricted posterior conditioned on $U \in \mathcal{U}$. We treat open-world evaluation as a limitation (Section 7).

3.2 Releases as channels

A public disclosure (“release”) at time t is modeled as a channel (Markov kernel)

$$K_t : \tilde{\mathcal{X}} \rightarrow \mathcal{Z}_t, \quad Z_t \sim K_t(\cdot | \tilde{X}). \quad (5)$$

The cumulative released variables are $Z_{1:T} := (Z_1, \dots, Z_T)$.

What counts as a release (compact view). In our framework, any artifact that increases available side information is a release, including: (i) additional data/coverage, (ii) model checkpoints/APIs inducing representations, and (iii) released features/indices or other derived artifacts (including metadata when present).

Experimental instantiation. In Section 6, Z_t are encoder-induced embeddings and, at the final stage, released artifacts (prototype bank and k NN graph).

3.3 Threat model and attacker posterior interface

Given a non-face query $\tilde{x}$ and access to releases up to time T , an attacker outputs an estimated posterior over identities:

$$\hat{p}_T^{(k)}(u \mid \tilde{x}, Z_{1:T}) \approx \Pr(U = u \mid \tilde{X} = \tilde{x}, Z_{1:T}), \quad (6)$$

where k indexes attacker capability (defined in Section 6). This posterior interface is algorithm-agnostic and covers retrieval-based and classifier-based attacks by mapping their scores to a distribution over $\mathcal{U}$.

Operational target: guessing probability. The attacker-independent privacy target is the Bayes-optimal guessing probability

$$P_{\text{guess}}(U \mid Z_{1:T}) := \mathbb{E}_{z_{1:T}} \left[\max_{u \in \mathcal{U}} \Pr(U = u \mid Z_{1:T} = z_{1:T}) \right], \quad (7)$$

which we connect to min-entropy and Rényi bounds in Section 4.

4 vRPI $_{\alpha}$: Rényi Pressure Index

We define an entropy-based leakage functional on the *cumulative-release posterior* $\Pr(U \mid Z_{1:T})$ induced by release channels (Section 3.2). The definition is attacker-independent; experiments instantiate plug-in estimates from attacker-induced posteriors.

4.1 Operational target: guessing probability and min-entropy

Given side information Z (e.g., $Z = Z_{1:T}$), define Bayes-optimal guessing probability

$$P_{\text{guess}}(U \mid Z) := \mathbb{E}_z \left[\max_{u \in \mathcal{U}} \Pr(U = u \mid Z = z) \right]. \quad (8)$$

Conditional min-entropy is exactly

$$H_{\min}(U \mid Z) := -\log P_{\text{guess}}(U \mid Z), \quad (9)$$

so $\exp(H_{\min}(U \mid Z)) = 1/P_{\text{guess}}(U \mid Z)$ can be interpreted as an *effective anonymity set size* [4, 5]. Our objective is to quantify how this effective set collapses as releases accumulate (larger T), even under non-face observations.

4.2 Arimoto conditional Rényi entropy and vRPI_α

Direct estimation of P_{guess} (hence $H_{\min}$) is sensitive to posterior maxima, so we use a tunable order- α relaxation. For $\alpha > 1$, Arimoto conditional Rényi entropy is [1–3]

$$H_\alpha(U \mid Z) := \frac{\alpha}{1-\alpha} \log \left(\mathbb{E}_z \left[\left(\sum_{u \in \mathcal{U}} \Pr(U = u \mid Z = z)^\alpha \right)^{1/\alpha} \right] \right). \quad (10)$$

We define the *Rényi Pressure Index* (for $N := |\mathcal{U}|$) as

$$\boxed{\text{vRPI}_\alpha(T) := \log N - H_\alpha(U \mid Z_{1:T}), \quad \alpha > 1.} \quad (11)$$

Under a uniform prior on U (main-text convention), $\log N$ is maximal uncertainty, hence vRPI_α measures the reduction in conditional uncertainty induced by cumulative releases. Moreover, for $\alpha > 1$, $\max_u p_u \leq \|p\|_\alpha$ implies a direct operational bridge from H_α to P_{guess} (Theorem 2), and larger α increases sensitivity to concentrated (single-candidate) posteriors.

4.3 Estimation from attacker posteriors

vRPI_α is defined on the true posterior, but we estimate it using attacker-induced posteriors $\hat{p}_T(u \mid z_{1:T})$ (Section 3.3). Given i.i.d. evaluation samples $\{z_{1:T}^{(i)}\}_{i=1}^n$, define the plug-in estimator

$$\hat{H}_\alpha(U \mid Z_{1:T}) := \frac{\alpha}{1-\alpha} \log \left(\frac{1}{n} \sum_{i=1}^n \left(\sum_{u \in \mathcal{U}} \hat{p}_T(u \mid z_{1:T}^{(i)})^\alpha \right)^{1/\alpha} \right) \quad (12)$$

$$\widehat{\text{RPI}}_\alpha(T) := \log N - \hat{H}_\alpha(U \mid Z_{1:T}) \quad (13)$$

This computation requires only (i) released variables $Z_{1:T}$ (e.g., embeddings/prototypes) and (ii) a mechanism to output $\hat{p}_T(\cdot \mid z_{1:T})$ over the fixed identity universe.

Non-uniform priors. The main-text convention assumes a uniform prior, for which $\log N$ is the maximal uncertainty. If $\Pr(U)$ is non-uniform, the correct normalization replaces $\log N$ with $H_\alpha(U)$, giving the prior-normalized index

$$\text{vRPI}_\alpha^\pi(T) := H_\alpha(U) - H_\alpha(U \mid Z_{1:T}), \quad (14)$$

which equals (11) when U is uniform. Under long-tailed priors the available anonymity budget $H_\alpha(U)$ is smaller, so a release erodes less prior uncertainty; we quantify this with a Zipf diagnostic in Section 6.4. vRPI_α is always well-defined and monotone under added side information; additive decompositions require explicit conditional-independence assumptions (Section 5).

Choosing α in practice. Larger α increases sensitivity to concentrated (single-candidate) posteriors. For general release auditing we recommend $\alpha = 2$ as a stable, interpretable “collision-style” default that is not dominated by a single extreme posterior entry. Larger orders ($\alpha = 5$ or 10) are appropriate for worst-case/high-confidence auditing but should be reported alongside calibration diagnostics (Section 6.4), since high-order Rényi entropies are more sensitive to posterior sharpness.

5 Theory Results

We state and prove core properties of vRPI_α . All results are defined on the *true* posterior induced by releases and are therefore attacker-independent. Any additional assumptions (e.g., conditional independence) are stated explicitly.

5.1 Preliminaries

Fix $\alpha > 1$. For side information Z define

$$\Phi_\alpha(U | Z) := \mathbb{E}_z \left[\|\Pr(U = \cdot | Z = z)\|_\alpha \right] = \sum_z \Pr(z) \left(\sum_{u \in \mathcal{U}} \Pr(u | z)^\alpha \right)^{1/\alpha}. \quad (15)$$

Arimoto conditional Rényi entropy satisfies

$$H_\alpha(U | Z) = \frac{\alpha}{1 - \alpha} \log \Phi_\alpha(U | Z), \quad \text{vRPI}_\alpha(T) = \log N - H_\alpha(U | Z_{1:T}). \quad (16)$$

For uniform U , $\Phi_\alpha(U | Z) \in [N^{\frac{1}{\alpha}-1}, 1]$.

5.2 Theorem 1: Monotonicity under cumulative releases

Theorem 1 (Monotonicity under additional releases). *Let Z and W be side-information variables jointly distributed with U (possibly dependent). For $\alpha > 1$,*

$$H_\alpha(U | Z, W) \leq H_\alpha(U | Z). \quad (17)$$

Consequently, for cumulative releases $Z_{1:T}$,

$$H_\alpha(U | Z_{1:T+1}) \leq H_\alpha(U | Z_{1:T}), \quad \text{vRPI}_\alpha(T+1) \geq \text{vRPI}_\alpha(T). \quad (18)$$

Proof. Fix z with $\Pr(z) > 0$. The posterior is a mixture:

$$\Pr(U = \cdot | Z = z) = \mathbb{E}_{W|Z=z} [\Pr(U = \cdot | Z = z, W)].$$

Since $\|\cdot\|_\alpha$ is convex for $\alpha \geq 1$,

$$\|\Pr(U = \cdot | Z = z)\|_\alpha \leq \mathbb{E}_{W|Z=z} [\|\Pr(U = \cdot | Z = z, W)\|_\alpha].$$

Taking expectation over z gives $\Phi_\alpha(U | Z) \leq \Phi_\alpha(U | Z, W)$. Because $\alpha/(1-\alpha) < 0$, applying $\frac{\alpha}{1-\alpha} \log(\cdot)$ reverses the inequality, yielding (17). Setting $Z = Z_{1:T}$ and $W = Z_{T+1}$ gives (18).

5.3 Theorem 2: Guessing probability bound

Theorem 2 (Rényi bound on Bayes-optimal guessing). *For any $\alpha > 1$ and side information Z ,*

$$P_{\text{guess}}(U | Z) \leq \Phi_\alpha(U | Z) = \exp\left(\frac{1-\alpha}{\alpha} H_\alpha(U | Z)\right), \quad (19)$$

equivalently $-\log P_{\text{guess}}(U | Z) \geq \frac{\alpha-1}{\alpha} H_\alpha(U | Z)$.

Proof. For fixed z , let $p_u = \Pr(U = u | Z = z)$. For $\alpha > 1$, $\max_u p_u \leq (\sum_u p_u^\alpha)^{1/\alpha} = \|p\|_\alpha$. Taking expectation over z yields $P_{\text{guess}}(U | Z) \leq \Phi_\alpha(U | Z)$. The exponential form follows from (16).

5.4 Theorem 3: Ideal additive benchmark under conditional independence

We give an additive decomposition that serves as an *ideal benchmark* under a structural assumption; it is not a claim that realistic CV releases are conditionally independent. Without the assumption, monotonicity (Theorem 1) remains the universal guarantee, and the empirical gap from additivity is quantified directly via a dependence-overlap diagnostic (Section 6.4).

Assumption (conditional independence). Conditioned on U , releases are independent:

$$p(z_{1:T} | u) = \prod_{t=1}^T p(z_t | u). \quad (20)$$

Sibson mutual information. For $\alpha > 1$, define Sibson mutual information [28, 29]

$$I_\alpha^S(U; Z) := \frac{\alpha}{\alpha-1} \log \sum_z \left(\sum_u \Pr(u) p(z | u)^\alpha \right)^{1/\alpha}. \quad (21)$$

Theorem 3 (Additivity (ideal case) and an additive upper bound on pressure). *Assume (20). Then for any $\alpha > 1$,*

$$I_\alpha^S(U; Z_{1:T}) = \sum_{t=1}^T I_\alpha^S(U; Z_t). \quad (22)$$

If U is uniform, then $\text{vRPI}_\alpha(T) = I_\alpha^A(U; Z_{1:T})$ and

$$\text{vRPI}_\alpha(T) \leq I_\alpha^S(U; Z_{1:T}) = \sum_{t=1}^T I_\alpha^S(U; Z_t), \quad (23)$$

where $I_\alpha^A(U; Z) := H_\alpha(U) - H_\alpha(U | Z)$ is Arimoto mutual information.

Proof. Under (20), Sibson information satisfies a product-channel (information-radius) factorization, yielding (22) [28, 29]. For uniform U , $H_\alpha(U) = \log N$ so $\text{vRPI}_\alpha(T) = \log N - H_\alpha(U | Z_{1:T}) = I_\alpha^A(U; Z_{1:T})$. Finally, $I_\alpha^A(U; Z) \leq I_\alpha^S(U; Z)$ for $\alpha > 1$ [3, 29], giving (23).

When additivity fails. If releases are correlated given U (duplicates, shared pipelines/metadata), (20) fails and exact additivity need not hold. We therefore treat (23) as an *ideal additive benchmark*; monotonicity remains unconditional. To make the deviation measurable rather than rhetorical, we report the dependence-overlap diagnostic

$$\Delta_{\text{corr},2} := \sum_t \widehat{\text{vRPI}}_2(Z_t) - \widehat{\text{vRPI}}_2(Z_{1:T}), \quad (24)$$

which is non-negative under the ideal benchmark and grows with redundancy between releases; its empirical value is given in Section 6.4.

5.5 Theorem 4: Ablations as channels

Theorem 4 (Ablation monotonicity). *Let τ_A, τ_B be transforms with $U \rightarrow \tau_A(X) \rightarrow \tau_B(X)$ a Markov chain (i.e., τ_B is a post-processing of τ_A). Under the same release schedule,*

$$H_\alpha(U \mid Z_{1:T}^B) \geq H_\alpha(U \mid Z_{1:T}^A) \iff \text{vRPI}_\alpha^A(T) \geq \text{vRPI}_\alpha^B(T), \quad (25)$$

so removing information cannot increase measured pressure.

Proof. Since τ_B is a post-processing of τ_A , the cumulative releases satisfy $U \rightarrow Z_{1:T}^A \rightarrow Z_{1:T}^B$. Apply Theorem 1 along this Markov chain.

5.6 Theorem 5: Stability of plug-in vRPI

Theorem 5 (Plug-in stability). *Fix $\alpha > 1$ and uniform U . Let $p(\cdot \mid z)$ be the true posterior and $\hat{p}(\cdot \mid z)$ an estimate satisfying $\mathbb{E}_Z[\|p(\cdot \mid Z) - \hat{p}(\cdot \mid Z)\|_1] \leq \varepsilon$. Then*

$$|\text{vRPI}_\alpha(T) - \widehat{\text{vRPI}}_\alpha(T)| \leq \frac{\alpha}{\alpha - 1} \log \left(1 + \frac{\varepsilon}{N^{\frac{1}{\alpha} - 1}} \right). \quad (26)$$

Proof. Norm Lipschitzness gives $|\|p\|_\alpha - \|\hat{p}\|_\alpha| \leq \|p - \hat{p}\|_\alpha \leq \|p - \hat{p}\|_1$ for $\alpha \geq 1$. Taking expectation yields $|\Phi_\alpha - \hat{\Phi}_\alpha| \leq \varepsilon$. Under uniform U , $\Phi_\alpha \geq N^{\frac{1}{\alpha} - 1}$, so $|\log \Phi_\alpha - \log \hat{\Phi}_\alpha| \leq \log(1 + \varepsilon/N^{\frac{1}{\alpha} - 1})$. Multiply by $\alpha/(\alpha - 1)$ and use $\text{vRPI}_\alpha = \log N - H_\alpha$.

6 Illustrative Experiments

We instantiate the theory in a closed-world, non-face identity inference setting. Our objectives are diagnostic and theory-aligned: (i) measure how inference risk changes across a release sequence, (ii) test dependence on attacker class, and (iii) verify that vRPI_α behaves as a concentration-sensitive leakage functional under realistic posterior estimators. All reported numbers are averaged over 3 random seeds (mean $\pm$ std).

6.1 Setup: dataset, releases, and attackers

Dataset and splits. We use Market-1501 [6] restricted to `bounding_box_train`, so $\mathcal{U}$ is the $N=751$ training identities (closed-world), and apply τ by masking the top 20% of each image (MARKET-1501-TAU). Per identity we allocate 50% of images to an exposure pool and split the rest evenly into gallery/query over the same label space; a sparse 25% exposure subset defines R_1 , and a held-out 10% exposure split is used solely to tune fusion weights and temperatures.

Release schedule R_1 – R_4 . R_1 releases an encoder E_1 trained on sparse exposure; R_2 adds E_2 trained on full exposure with cumulative state $[E_1, \lambda_2 E_2]$ (ℓ_2 -normalized concatenation); R_3 adds a generic ImageNet-pretrained ResNet-50 [33] E_3 , state $[E_1, \lambda_2 E_2, \lambda_3 E_3]$; R_4 releases two artifacts from the R_3 gallery, an identity prototype bank and a k NN graph, adding leakage pathways without changing the base feature state. Fusion weights λ_2, λ_3 are tuned on validation.

Attacker ladder. All attackers output a posterior $\hat{p}(u | z)$ over the same $N=751$ identities: A_0 (black-box retrieval, log-sum-exp over same-identity gallery images then temperature-scaled softmax), A_1 (linear probe), A_2 (prototype-index specialist, R_4 only), A_3 (MLP probe), and A_4 (white-box fine-tuning, R_1 – R_3). Retrieval similarities are aggregated to identity posteriors before entropy computation, so every vRPI_α is computed over the same label space, matching the theory.

6.2 Results

Table 1 reports mean $\pm$ std over 3 seeds for P_{guess} and vRPI_α at $\alpha \in \{1.5, 2, 5, 10\}$.

Exposure expansion dominates ($R_1 \rightarrow R_2$). The clearest effect is release-driven exposure growth: all attackers exhibit large increases (e.g. A_0 : $P_{\text{guess}} 0.2883 \rightarrow 0.7491$, $\text{vRPI}_2 4.3319 \rightarrow 6.2119$), empirically supporting the compositional leakage thesis under non-face observations.

Generic channel stabilizes under calibrated fusion ($R_2 \rightarrow R_3$). With validation-tuned fusion weights and posterior temperatures, adding the generic representation E_3 avoids the brittle regressions of naive concatenation: attackers are stable or slightly improved (A_0 : $0.7491 \rightarrow 0.7529$; A_3 : $0.7063 \rightarrow 0.7092$).

R_4 demonstrates artifact-specific leakage. R_4 adds released artifacts rather than altering the base feature state: the prototype-index specialist A_2 reaches $P_{\text{guess}}=0.7519$ ($\text{vRPI}_2=6.3572$) and the graph-exploiting A_0 -GRAPH reaches 0.7366 ($\text{vRPI}_2=5.6697$), tying the increase to publishable artifacts (indices/graphs). Figure 1 summarizes these mean trends; α -sweeps and scatter diagnostics are in the supplement.

Release	Attacker	P_{guess}	vRPI _{1.5}	vRPI ₂	vRPI ₅	vRPI ₁₀
R_1	A_0	0.2883	3.9180	4.3319	4.9667	5.1327
		± 0.0222	± 1.1614	± 0.9997	± 0.7220	± 0.6484
R_1	A_1	0.1600	1.2348	1.6561	2.7437	3.1123
		± 0.0086	± 0.1356	± 0.1692	± 0.1769	± 0.1622
R_1	A_3	0.4073	4.3074	4.7231	5.2798	5.4163
		± 0.0173	± 0.1525	± 0.1285	± 0.0911	± 0.0818
R_1	A_4	0.5073	4.6916	5.0912	5.5685	5.6779
		± 0.0252	± 0.1459	± 0.1160	± 0.0791	± 0.0708
R_2	A_0	0.7491	6.1250	6.2119	6.3269	6.3561
		± 0.0087	± 0.2761	± 0.2183	± 0.1495	± 0.1338
R_2	A_1	0.5375	4.5612	4.9997	5.5132	5.6292
		± 0.0051	± 0.0049	± 0.0041	± 0.0031	± 0.0029
R_2	A_3	0.7063	5.7696	5.9402	6.1457	6.1943
		± 0.0026	± 0.0290	± 0.0213	± 0.0134	± 0.0118
R_2	A_4	0.7304	5.8421	6.0133	6.2039	6.2471
		± 0.0056	± 0.0161	± 0.0117	± 0.0079	± 0.0071
R_3	A_0	0.7529	5.7669	5.9271	6.1305	6.1801
		± 0.0071	± 0.0146	± 0.0114	± 0.0077	± 0.0069
R_3	A_1	0.5372	4.5504	4.9917	5.5080	5.6245
		± 0.0043	± 0.0180	± 0.0141	± 0.0096	± 0.0085
R_3	A_3	0.7092	5.7660	5.9368	6.1430	6.1918
		± 0.0043	± 0.0243	± 0.0174	± 0.0105	± 0.0093
R_3	A_4	0.7238	5.8319	6.0060	6.1992	6.2429
		± 0.0039	± 0.0314	± 0.0239	± 0.0160	± 0.0144
R_4	$A_0\text{-GRAPH}$	0.7366	5.4553	5.6697	5.9486	6.0168
		± 0.0076	± 0.0177	± 0.0123	± 0.0070	± 0.0060
R_4	A_2	0.7519	6.3071	6.3572	6.4274	6.4461
		± 0.0070	± 0.0082	± 0.0074	± 0.0059	± 0.0053

Table 1: Mean $\pm$ std over 3 seeds on MARKET-1501-TAU ($N=751$). R_2 and R_3 use validation-tuned fusion weights and temperature calibration; R_4 introduces released artifacts (prototype bank and k NN graph) exploited by A_2 and $A_0\text{-GRAPH}$.

6.3 Replication and non-face ablation

CUHK03-Tau replication. To test whether the effect is benchmark-specific, we replicate the closed-world protocol on CUHK03 [7] under the same top-20% transform (CUHK03-TAU); since CUHK03 filenames encode a camera-pair and local person index, identity is the *pair/person tuple* to avoid merging distinct persons. The core exposure-expansion test ($R_1 \rightarrow R_2$) for retrieval (A_0) and nonlinear-probe (A_3) attackers holds on *both* datasets (Table 2). This is a separate, lighter-weight paired protocol from Table 1: it uses score-level channel fusion and a reduced training budget, and retains only identities with at least four images so each yields a valid exposure/gallery/query triple ($N=736$ of the 751 Market identities; $N=700$ on CUHK03). Absolute magnitudes are therefore

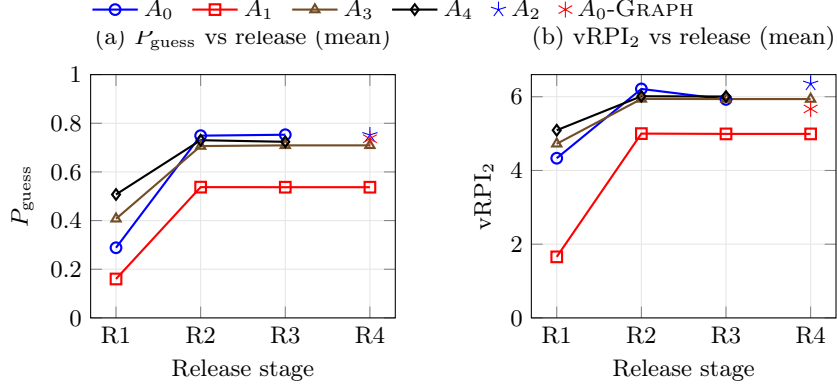


Fig. 1: Mean trends over 3 seeds on MARKET-1501-TAU. R_2 shows the dominant jump from exposure expansion. R_4 introduces explicit publishable artifacts (prototype/index and k NN graph), yielding artifact-specific leakage exploited by A_2 and A_0 -GRAPH.

Data	Att.	N	$R_1 P_g$	$R_2 P_g$	$R_1 v\text{RPI}_2$	$R_2 v\text{RPI}_2$
Market	A_0	736	.197	.547	.873	4.966
Market	A_3	736	.368	.691	5.198	5.943
CUHK03	A_0	700	.226	.399	.530	3.325
CUHK03	A_3	700	.313	.502	4.549	4.977

Table 2: Closed-world replication of the exposure-expansion effect ($R_1 \rightarrow R_2$) under an independent, lighter-weight paired protocol (score-level fusion, reduced budget, ≥ 4 images/identity, so N differs from Table 1). Absolute values are not directly comparable to Table 1; the paired $R_1 \rightarrow R_2$ increase is consistent across both datasets and attackers.

not directly comparable to Table 1; the claim here is solely the *within-protocol* paired increase $R_1 \rightarrow R_2$, which is positive for both attackers on both datasets.

Non-face masking ablation. The top-20% mask is a deterministic proxy, not a realistic anonymizer. Sweeping the top-region masking fraction $\tau \in \{0, 10, 20, 30, 40\}$ % on Market (full table in the supplement), pressure is essentially flat: A_0 $v\text{RPI}_2$ stays in $[4.87, 4.91]$ and A_3 $v\text{RPI}_2$ in $[5.83, 5.88]$, with P_{guess} still $\geq .51$ (A_0) and $\geq .64$ (A_3) at top-40%. Leakage thus persists when large head/upper-body regions are removed; this is a controlled stress test and does *not* replace detector-based or generative anonymization (future work).

6.4 Diagnostic checks: dependence, calibration, and priors

Dependence overlap. Theorem 3 is an ideal additive benchmark; realistic releases share image sources, encoders, and pipelines and so violate (20). We report the diagnostic $\Delta_{\text{corr},2}$ of (24): on Market the marginal pressures for E_1, E_2, E_3 are .884, 4.159, .629 and the joint pressure is 5.368, giving $\Delta_{\text{corr},2} = .304 > 0$. The

joint release leaks *less* than the sum of marginals, quantifying the redundancy between channels and confirming Theorem 3 is a benchmark, not an independence claim.

Temperature calibration. Posteriors use temperature scaling tuned on the held-out validation split. Sweeping $\theta \in \{.025, .05, .1, .2, .5, 1.0\}$, absolute vRPI_2 changes (as expected for a concentration entropy) but the release conclusion is stable: $\Delta \text{vRPI}_2(R_2 - R_1) > 0$ throughout, with A_0 spanning $[.781, 4.100]$ and A_3 spanning $[.567, .677]$. The ε in Theorem 5 is not directly observable; this sweep is a robustness check, not a posterior-error estimate.

Non-uniform (Zipf) prior. Under a Zipf(s) prior, the prior-normalized index (14) behaves as expected: prior-normalized ΔvRPI_2 shrinks from 3.818 at $s = 0$ (uniform) to .123 at $s = 1.5$, since a concentrated prior leaves less prior anonymity to erode, validating $H_\alpha(U)$ as the correct normalization for long-tailed/open-world extensions.

Positioning vs. existing measures. Maximal leakage [29] and maximal α -leakage [30] optimize adversarial gain over an arbitrary function of the secret per channel, while DP/attribute-inference bounds [32] provide worst-case guarantees. In contrast, vRPI_α fixes the vision identity task and measures posterior concentration over a release sequence. Table 3 summarizes the distinction. These measures are complementary rather than interchangeable.

7 Discussion, Limitations, and Ethics

Assumptions. Monotonicity (Theorem 1) and the Rényi guessing bound (Theorem 2) hold for any joint law of $(U, Z_{1:T})$; only additivity (Theorem 3) requires

Table 3: Positioning of vRPI_α against related evaluation criteria.

Measure	Object measured	Release-sequence aware?	Role in this paper
Rank/mAP	Retrieval ranking quality	No	Utility/task metric; does not directly quantify anonymity erosion.
P_{guess}	Bayes posterior maximum	Yes, if computed on $Z_{1:T}$	Operational identity-guessing target tracked by the audit.
vRPI_α	Rényi posterior concentration	Yes	Proposed pressure index for cumulative identity leakage.
Maximal leakage	Worst-case channel gain	Channel-level	Complementary leakage notion; not tied to a fixed vision identity task.
Maximal leakage	α -Tunable adversarial gain	Channel-level	Related Rényi-family leakage measure with different optimization target.
DP / attribute inference	Worst-case stability or inference bound	Yes, via composition	Stronger worst-case guarantee; different from realized release auditing.

the conditional-independence idealization (20), which correlated releases violate. The dependence-overlap diagnostic ($\Delta_{\text{corr},2}=.304$, Section 6.4) reports the realized gap, and monotonicity remains the unconditional guarantee.

Scope. Our instantiation uses MARKET-1501-TAU ($N=751$) with a CUHK03-TAU replication; both are closed-world image benchmarks, and broader coverage (more datasets/domains, video, larger scales, longer sequences) is future work not required by the theory. We keep the mainline closed-world because open-set effects are inconsistent with the closed-world normalization; the open-world extension augments $\mathcal{U}$ with an “unknown” option and a rejection-calibrated posterior, under which the prior-normalized index (14) applies and Theorems 1–2 carry over unchanged. The non-face transform τ is a controlled top-region mask: the ablation (Section 6.3) shows pressure persists to top-40%, but detector-based and generative anonymization remain more realistic mitigations and are out of scope. vRPI_α is an audit statistic, not itself a mitigation method.

Ethics and release hygiene. This work measures the *risk* of non-face re-identification under cumulative releases; it is not a guide for surveillance or targeting. The implication is governance-oriented: releases should be evaluated as sequences rather than in isolation. The framework supports *release auditing* and conservative hygiene, limiting high-leverage artifacts (embedding dumps/indices), minimizing unnecessary metadata exposure, and quantifying cumulative risk before publishing successive versions.

8 Conclusion

We formalized compositional non-face identity leakage by modeling disclosures as channels and analyzing the posterior $\Pr(U \mid Z_{1:T})$, introducing $\text{vRPI}_\alpha(T) = \log N - H_\alpha(U \mid Z_{1:T})$ from Arimoto conditional Rényi entropy and linking it to operational guessing via P_{guess} and min-entropy. We proved monotonicity, a Rényi guessing bound, an ideal additive benchmark under conditional independence, and a channel-based ablation guarantee. Closed-world instantiations on MARKET-1501-TAU and CUHK03-TAU show large pressure growth under exposure expansion and artifact-specific leakage via published prototype/index and k NN-graph artifacts. vRPI_α thus enables principled *release auditing*: assessing whether a planned sequence of dataset/model/index disclosures will cumulatively collapse identity uncertainty even when faces are removed.

Acknowledgments. We thank the anonymous reviewers and the area chair for their constructive feedback, which substantially improved this paper.

References

1. S. Arimoto. Information measures and capacity of order α for discrete memoryless channels. In: Csiszár, I., Elias, P. (eds.) *Topics in Information Theory* (Proc. 2nd Colloq. Math. Soc. János Bolyai, Keszthely, 1975), pp. 41–52. North-Holland, Amsterdam (1977)
2. Fehr, S., Berens, S.: On the conditional Rényi entropy. *IEEE Transactions on Information Theory* **60**(11), 6801–6810 (2014)
3. Sason, I., Verdú, S.: Arimoto–Rényi conditional entropy and Bayesian M -ary hypothesis testing. *IEEE Transactions on Information Theory* **64**(1), 4–25 (2018)
4. König, R., Renner, R., Schaffner, C.: The operational meaning of min- and max-entropy. *IEEE Transactions on Information Theory* **55**(9), 4337–4347 (2009)
5. Tomamichel, M.: *Quantum Information Processing with Finite Resources: Mathematical Foundations*. SpringerBriefs in Mathematical Physics. Springer, Cham (2016)
6. Zheng, L., Shen, L., Tian, L., Wang, S., Wang, J., Tian, Q.: Scalable person re-identification: A benchmark. In: *Proceedings of the IEEE International Conference on Computer Vision (ICCV)*, pp. 1116–1124 (2015)
7. Li, W., Zhao, R., Xiao, T., Wang, X.: DeepReID: Deep filter pairing neural network for person re-identification. In: *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, pp. 152–159 (2014)
8. Gou, M., Karanam, S., Liu, W., Camps, O., Radke, R.J.: DukeMTMC4ReID: A large-scale multi-camera person re-identification dataset. In: *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition Workshops (CVPRW)*, pp. 10–19 (2017)
9. Zheng, L., Bie, Z., Sun, Y., Wang, J., Su, C., Wang, S., Tian, Q.: MARS: A video benchmark for large-scale person re-identification. In: Leibe, B., Matas, J., Sebe, N., Welling, M. (eds.) *Computer Vision – ECCV 2016*. LNCS, vol. 9910, pp. 868–884. Springer, Cham (2016)
10. Zheng, L., Zhang, H., Sun, S., Chandraker, M., Yang, Y., Tian, Q.: Person re-identification in the wild. In: *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, pp. 1367–1376 (2017)
11. Ye, M., Shen, J., Lin, G., Xiang, T., Shao, L., Hoi, S.C.H.: Deep learning for person re-identification: A survey and outlook. *IEEE Transactions on Pattern Analysis and Machine Intelligence* **44**(6), 2872–2893 (2022)
12. Hermans, A., Beyer, L., Leibe, B.: In defense of the triplet loss for person re-identification. *arXiv preprint arXiv:1703.07737* (2017)
13. Zhou, K., Yang, Y., Cavallaro, A., Xiang, T.: Omni-scale feature learning for person re-identification. In: *Proceedings of the IEEE/CVF International Conference on Computer Vision (ICCV)*, pp. 3702–3712 (2019)
14. Luo, H., Gu, Y., Liao, X., Lai, S., Jiang, W.: Bag of tricks and a strong baseline for deep person re-identification. In: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition Workshops (CVPRW)*, pp. 1487–1495 (2019). <https://doi.org/10.1109/CVPRW.2019.00190>
15. Zhong, Z., Zheng, L., Cao, D., Li, S.: Re-ranking person re-identification with k-reciprocal encoding. In: *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, pp. 1318–1327 (2017)
16. Qian, X., Wang, W., Zhang, L., Zhu, F., Fu, Y., Xiang, T., Jiang, Y.-G., Xue, X.: Long-term cloth-changing person re-identification. In: Vedaldi, A., Bischof, H., Brox, T., Frahm, J.-M. (eds.) *Computer Vision – ACCV 2020*. LNCS, vol. 12624,

- pp. 71–88. Springer, Cham (2021). https://doi.org/10.1007/978-3-030-69535-4_5
17. Yang, Q., Wu, A., Zheng, W.-S.: Person re-identification by contour sketch under moderate clothing change. *IEEE Transactions on Pattern Analysis and Machine Intelligence* **43**(6), 2029–2046 (2021)
 18. Chao, H., He, Y., Zhang, J., Feng, J.: GaitSet: Regarding gait as a set for cross-view gait recognition. In: *Proceedings of the AAAI Conference on Artificial Intelligence (AAAI)*, pp. 8126–8133 (2019)
 19. Newton, E.M., Sweeney, L., Malin, B.: Preserving privacy by de-identifying facial images. *IEEE Transactions on Knowledge and Data Engineering* **17**(2), 232–243 (2005)
 20. Gross, R., Airoldi, E., Malin, B., Sweeney, L.: Integrating utility into face de-identification. In: *Workshop on Privacy-Enhancing Technologies (PET)* (2005)
 21. Meden, B., Emeršič, Ž., Štruc, V., Peer, P.: k-Same-Net: k-anonymity with generative deep neural networks for face deidentification. *Entropy* **20**(1), 60 (2018)
 22. Hukkelås, H., Mester, R., Lindseth, F.: DeepPrivacy: A generative adversarial network for face anonymization. In: *Bebis, G., Boyle, R., Parvin, B., Koracin, D., Ushizima, D., Chai, S., Sueda, S., Lin, X., Lu, A., Thalmann, D., Wang, C., Xu, P. (eds.) Advances in Visual Computing. LNCS, vol. 11844, pp. 565–578. Springer, Cham (2019). https://doi.org/10.1007/978-3-030-33720-9_44*
 23. Hukkelås, H., Lindseth, F.: DeepPrivacy2: Towards realistic full-body anonymization. In: *Proceedings of the IEEE/CVF Winter Conference on Applications of Computer Vision (WACV)*, pp. 1329–1338 (2023)
 24. Shan, S., Wenger, E., Zhang, J., Li, H., Zheng, H., Zhao, B.Y.: Fawkes: Protecting privacy against unauthorized deep learning models. In: *Proceedings of the 29th USENIX Security Symposium*, pp. 1589–1604 (2020)
 25. Klomp, S.R., van Rijn, M., Wijnhoven, R.G.J., Snoek, C.G.M., de With, P.H.N.: Safe fakes: Evaluating face anonymizers for face detectors. In: *16th IEEE International Conference on Automatic Face and Gesture Recognition (FG)*, pp. 1–8 (2021). <https://doi.org/10.1109/FG52635.2021.9666936>
 26. Zhao, Y., Chen, J.: A survey on differential privacy for unstructured data content. *ACM Computing Surveys* **54**(10s), Article 207, 1–28 (2022). <https://doi.org/10.1145/3490237>
 27. Rényi, A.: On measures of entropy and information. In: *Proceedings of the Fourth Berkeley Symposium on Mathematical Statistics and Probability*, vol. 1, pp. 547–561 (1961)
 28. Sibson, R.: Information radius. *Zeitschrift für Wahrscheinlichkeitstheorie und Verwandte Gebiete* **14**(2), 149–160 (1969)
 29. Issa, I., Wagner, A.B., Kamath, S.: An operational approach to information leakage. *IEEE Transactions on Information Theory* **66**(3), 1625–1657 (2020)
 30. Liao, J., Kosut, O., Sankar, L., Calmon, F.P.: Tunable measures for information leakage and applications to privacy-utility tradeoffs. *IEEE Transactions on Information Theory* **65**(12), 8043–8066 (2019)
 31. Hsu, H., Martinez, N.L., Bertran, M., Sapiro, G., Calmon, F.P.: A survey on statistical, information, and estimation-theoretic views on privacy. *IEEE BITS the Information Theory Magazine* **1**(1), 45–56 (2021). <https://doi.org/10.1109/MBITS.2021.3108124>
 32. Guo, C., Sablayrolles, A., Sanjabi, M.: Analyzing privacy leakage in machine learning via multiple hypothesis testing: A lesson from Fano. In: *Proceedings of the 40th International Conference on Machine Learning (ICML). Proceedings of Machine Learning Research*, vol. 202, pp. 11998–12011 (2023)

33. He, K., Zhang, X., Ren, S., Sun, J.: Deep residual learning for image recognition. In: Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR), pp. 770–778 (2016). <https://doi.org/10.1109/CVPR.2016.90>