

Dual-Branch ConvNeXt with Parameter-Free 3D Attention for Person Re-Identification

Jingshu Li

School of Computer Science and Engineering
Xi'an Technological University
Xi'an, 710021, China
E-mail: lij812@163.com

Jianguo Wang

School of Computer Science and Engineering
Xi'an Technological University
Xi'an, 710021, China
E-mail: wangjianguo@xatu.edu.cn

Abstract—Retrieving specific individuals across disjoint camera networks, widely known as pedestrian re-identification (Re-ID), serves as a fundamental pillar for modern visual surveillance and smart security systems. However, practical deployments are frequently hampered by shifting viewpoints, background clutter, and diverse body poses. Consequently, obtaining representations that are both highly robust and discriminative remains a formidable obstacle. Motivated by these challenges, we design an innovative Re-ID architecture that synergizes a 3D attention module with a dual-branch local partitioning scheme. Specifically, we adopt the ConvNeXt architecture as our baseline model to capture holistic visual cues. To further refine these representations, a parameter-free SimAM module is incorporated into the deeper layers of the network, which effectively highlights the subjects of interest while filtering out spatial-channel noise. Furthermore, to accurately distinguish between individuals with highly similar appearances, we develop a parallel global-local feature extraction paradigm. This involves a horizontal six-part division strategy dedicated to uncovering subtle, fine-grained visual details. By jointly minimizing the label smoothing cross-entropy and hard mining triplet losses, our framework seamlessly integrates macro-level structural knowledge with micro-level specificities. Extensive evaluations on the well-known Market-1501 benchmark demonstrate the exceptional performance of our approach. Remarkably, even without relying on complex post-processing techniques like re-ranking, our model attains a Rank@1 accuracy of 94.8% and an mAP of 86.1%. Ultimately, these empirical findings compellingly confirm the practical viability and advanced capabilities of our proposed framework for challenging cross-camera retrieval tasks.

Keywords—Person Re-Identification; Parameter-Free 3D Attention; ConvNeXt; Fine-Grained Feature Learning; Dual-Branch Architecture

I. INTRODUCTION

As a pivotal domain within computer vision, person re-identification (Re-ID) has garnered

immense academic interest. Its primary objective lies in associating and continuously tracking specific individuals across disjoint camera networks. This capability is particularly crucial in unconstrained scenarios where high-resolution biometrics, such as facial details, are unobservable, compelling the system to rely exclusively on broader visual cues like overall appearance, apparel, and gait patterns. Propelled by the rapid evolution of deep learning, Re-ID algorithms exhibit vast potential for real-world deployment, finding extensive applications in smart security, intelligent retail analytics, and seamless cross-camera trajectory tracking.

Nevertheless, implementing these systems in open, real-world surveillance environments remains a formidable task due to intricate visual interferences. Specifically, severe intra-class variations frequently occur; the visual representation of a single individual can fluctuate drastically across different camera views, primarily driven by inconsistent illumination, shifting perspectives, and diverse body postures. Conversely, the task also struggles with minimal inter-class discrepancies, as different subjects often don highly similar clothing. As a result, conventional recognition paradigms that depend exclusively on monolithic global appearance features are extremely vulnerable to erroneous matching and false retrievals under such demanding conditions.

To tackle the aforementioned bottlenecks in unconstrained scenarios, contemporary mainstream research frequently incorporates attention mechanisms and multi-scale feature fusion paradigms to elevate both the representational quality and discriminative power of pedestrian

features. While conventional attention algorithms, exemplified by SENet and CBAM, can optimize feature expression to a certain extent, they are fundamentally constrained. These mechanisms typically compute attention weights relying solely on either a singular channel dimension or a 2D spatial plane, thereby failing to fully exploit the inherent 3D structural correlations within the feature maps. Concurrently, prevailing networks heavily depend on Global Average Pooling (GAP) for feature aggregation. Although such global representations possess commendable translation invariance, they inevitably discard critical fine-grained local details, such as shoe colors or clothing logos. Importantly, these subtle visual cues act as the decisive factors for distinguishing hard negative samples with highly similar appearances and are indispensable for boosting the model's discriminative accuracy.

To circumvent the multitude of limitations inherent in existing methodologies, we introduce a novel dual-branch feature enhancement Re-ID network. This architecture is designed to comprehensively augment the discriminative capability of pedestrian representations from two distinct perspectives: multi-dimensional adaptive attention allocation and fine-grained local feature perception. Specifically, ConvNeXt [1], renowned for its exceptional feature extraction capabilities, is deployed as the backbone. Leveraging its extraordinarily large receptive fields and highly efficient feature reconstruction advantages, the network extracts high-quality initial pedestrian representations. Building upon this foundation, we integrate SimAM [2], an attention module inspired by neuroscience theories. By solving a refined 3D energy function, this module dynamically generates a 3D attention weight map — synergizing both spatial and channel information — without introducing any additional parameters or computational overhead. This allows the model to precisely focus on the foreground pedestrian while effectively suppressing background noise. Crucially, to compensate for the inadequate detail perception of monolithic global features, a parallel horizontal-slice local branch is constructed. This branch vertically partitions the feature map into six equivalent and independent horizontal stripes,

conducting distinct pooling and metric learning for each region. This strategy compels the network to heavily weigh the localized anatomical variations of different body parts, thereby excavating abundant fine-grained visual clues. Through an end-to-end joint optimization strategy utilizing a combined loss function for both the global and local branches, the model ultimately constructs a pedestrian feature representation that exhibits formidable robustness against viewpoint shifts, pose variations, and background interferences.

The principal innovations of this work are encapsulated in a comprehensive Re-ID framework that seamlessly couples a ConvNeXt backbone with a parameter-free, 3D energy-based attention mechanism, thereby effectively filtering complex background noise and reinforcing salient pedestrian features. Building upon this foundation, we devise a dual-branch learning paradigm that operates global macroscopic representations in parallel with horizontally sliced local fine-grained features. This specific structural design successfully overcomes the fine-grained discrimination challenge, making it entirely feasible to accurately distinguish highly homologous individuals. Ultimately, the superiority of our proposed algorithm is unequivocally corroborated through extensive evaluations on the heavily benchmarked Market-1501 dataset. Supported by compelling quantitative metrics (which significantly surpass various state-of-the-art models) and qualitative visual evidence, the formulated architecture demonstrates exceptional cross-camera retrieval capabilities and highly robust hard sample mining performance.

II. RELATED WORK

A. Deep Person Re-Identification

Catalyzed by the rapid proliferation of deep convolutional neural networks (CNNs), the field of person Re-ID has witnessed paradigm-shifting advancements over the past few years. Initial deep learning paradigms predominantly concentrated on extracting holistic global representations. While these monolithic features exhibit robust translation invariance under controlled conditions, they frequently falter when confronted with real-world surveillance complexities, such as partial occlusions, cluttered backgrounds, and severe pose

variations. Consequently, these conventional methods fail to capture highly discriminative subtle cues.

To mitigate these inherent limitations, the research community has increasingly pivoted towards localized, fine-grained feature learning. Notably, quintessential horizontal-slicing architectures, epitomized by the PCB [3] model, enforce rigid geometric partitions upon the spatial dimensions of feature maps. This strategy substantially amplifies the network's sensitivity to localized human attributes, including clothing textures and carried accessories. Nevertheless, the majority of existing localized methodologies directly manipulate unrefined underlying feature maps without adequately filtering out background noise. Furthermore, relying exclusively on localized slices tends to fragment the holistic appearance and structural integrity of the pedestrian.

Grounded in this context, our proposed dual-branch architecture assimilates the merits of spatial slicing paradigms while executing a substantial structural upgrade. By preserving a parallel global pathway, the network maintains the integrity of macroscopic body contours. Concurrently, localized perception is seamlessly anchored upon high-quality, attention-enhanced features, thereby orchestrating a synergistic complementarity between macroscopic structures and microscopic details.

B. Visual Attention Mechanisms

Attention mechanisms have been proven to effectively improve the feature representation capabilities of models in various computer vision tasks. The core principle of these mechanisms is to guide the network to focus on high-value regions of an image, while suppressing redundant and ineffective background information. In earlier studies, classic attention mechanisms like SENet[4] utilized dimensionality reduction and channel excitation techniques to dynamically adjust the weight distribution among different feature channels. More advanced methods like CBAM[5] further integrated information from both spatial and channel dimensions, achieving dual-level attention modeling through either parallel or serial structure of modules.

While these classic methods can improve the quality of feature representation to some extent, they still have two inherent drawbacks. First, most of these attention modules rely on multi-layer perceptrons or fully connected layers to predict attention weights, which significantly increases the number of network parameters and computational overhead. Second, these methods separate spatial and channel features for independent modeling, ignoring the inherent deep relationships and synergy between the three-dimensional feature tensor. In response to the aforementioned technical challenges, the parameter-free attention module SimAM, inspired by the spatial inhibition characteristics of neural systems, offers a completely new approach to solving such problems. This module evaluates the importance of feature neurons by solving a three-dimensional energy function. Without the need for additional parameters, it enables joint adaptation of features in both spatial and channel dimensions. When integrated into the pedestrian recognition network, this module can effectively counteract background interference in complex surveillance scenarios, while also improving the quality of feature extraction.

III. TECHNICAL MODEL

To systematically address the susceptibility of cross-camera representations to intricate environmental noise and the formidable challenge of differentiating highly homologous pedestrians, this section meticulously elaborates on the overarching architecture and the core functional modules of our proposed Re-ID network.

Initially, we present a comprehensive synopsis of the global network topology and the forward propagation dynamics, thereby elucidating the foundational operational mechanisms of the model. Subsequently, we delve into the primary feature extraction backbone anchored by ConvNeXt. Within this context, profound emphasis is placed on the theoretical formulation of the parameter-free 3D attention mechanism, SimAM, which is instrumental in curtailing background distractions and adaptively fortifying salient pedestrian characteristics. Building upon this robust baseline, we then introduce a parallel dual-branch learning

paradigm, integrating both global and localized pathways, specifically designed to conquer the bottleneck of fine-grained discrimination. This sophisticated strategy seamlessly harmonizes holistic structural representations with localized detail excavation, thereby rectifying the intrinsic detail-deficiency associated with singular global features. Ultimately, we mathematically formalize the joint loss function, which is tailored to synergistically constrain the dual-branch training process, culminating in the end-to-end holistic optimization of the entire framework.

A. Overall Architecture

Figure 1 delineates the overarching architecture of our proposed Re-ID framework. Specifically, the pipeline accepts pedestrian RGB images standardized to a spatial resolution of 256×128 . These inputs are initially ingested by the ConvNeXt backbone, which acts as the foundational feature extractor to generate primal feature maps encapsulating macroscopic anatomical layouts.

In a quest to refine representational fidelity and systematically eliminate extraneous noise, these preliminary feature maps are subsequently fed into the SimAM module. Driven by a mathematical 3D

energy function, this parameter-free mechanism adaptively calibrates the tensor, concurrently amplifying valid foreground cues and attenuating convoluted background distractions across both spatial and channel dimensions.

Following this attentional augmentation, the resulting high-fidelity feature maps are routed into a bifurcated learning paradigm comprised of parallel global and local pathways. The global branch is fundamentally tasked with modeling the holistic macroscopic morphology of the pedestrian. Concurrently, the local branch, executing a rigorous horizontal slicing strategy, is dedicated to excavating indispensable, fine-grained visual specificities.

During the training phase, a collaborative supervision paradigm is enforced by jointly minimizing both classification and metric learning losses across the dual branches. This synergistic optimization algorithm compels the network to autonomously deduce representations that are simultaneously highly discriminative and remarkably robust, ultimately propelling cross-camera retrieval efficacy within highly complex surveillance scenarios.

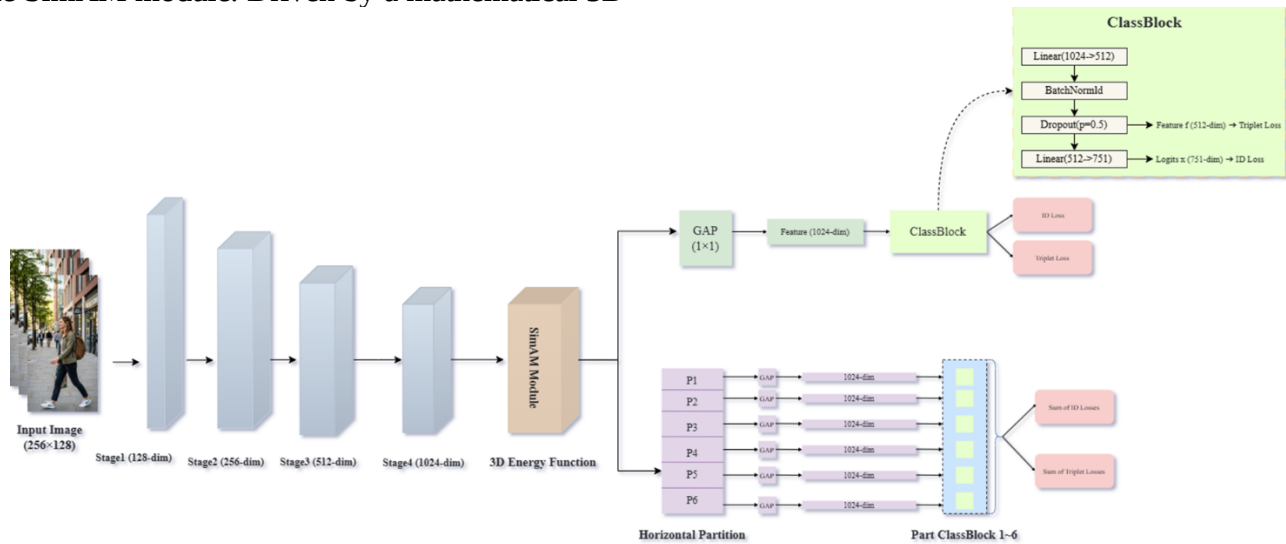


Figure 1. Network Architecture

intrinsic computational efficiency of pure convolutional operations with the macroscopic structural superiority of Vision Transformers, ConvNeXt serves as a highly reliable foundation for deriving high-fidelity pedestrian representations.

In challenging Re-ID scenarios, however, the efficacy of representation learning is frequently compromised by severe visual distractors, such as heavily cluttered backgrounds and unconstrained postural variations. Compounding this issue, orthodox attention mechanisms are conspicuously limited in performance, as they predominantly compute attention weights based strictly on either isolated 1D channel vectors or 2D spatial planes. To circumvent these structural deficiencies, we seamlessly embed the parameter-free 3D attention module, SimAM, at the deepest semantic level of the backbone, specifically immediately following the fourth stage.

The architectural formulation of SimAM is deeply rooted in the concept of spatial inhibition within neuroscience. This biological principle postulates that neurons harboring richer informational content typically exhibit firing patterns that starkly diverge from those of their neighboring counterparts. Grounded in this biomimetic mechanism, the module mathematically constructs a 3D energy function for each individual target neuron across the feature maps, thereby providing a quantitative metric to precisely evaluate its feature importance. By rigorously solving for the minimization of this underlying energy function, the minimal energy value corresponding to a given target neuron can be rapidly computed via a closed-form solution. The specific mathematical formulation is detailed in formula:

$$e_t^* = \frac{4(\hat{\sigma}^2 + \lambda)}{(t - \hat{\mu})^2 + 2\hat{\sigma}^2 + 2\lambda} \quad (1)$$

In this formula, μ and $\hat{\sigma}^2$ represents the mean and variance of all feature elements within the channel where the current neuron is located, while λ is a regularization parameter used to avoid zero-based errors. From the physical meaning of the energy function, it can be seen that the lower the minimum energy value of a neuron, the higher its

ability to distinguish itself from surrounding background features. Consequently, the neuron contributes more significantly to the feature representation process. Therefore, in this approach, the energy matrix is taken as its reciprocal, then normalized using the Sigmoid activation function. This result is then multiplied element by element with the original feature matrix. Ultimately, a three-dimensional enhanced feature matrix is obtained, as shown in formula:

$$X' = X \odot \text{Sigmoid}\left(\frac{1}{E}\right) \quad (2)$$

Among them, E represents the global energy matrix that integrates the minimum energy values of all neurons. This calculation method allows the SimAM module to synchronize the modeling of information related to both spatial and channel dimensions, thereby generating a three-dimensional attention weight matrix. This parameter-free attention modeling approach not only avoids the problem of redundant network parameters and increased computational overhead, but also enables the network to automatically focus on the effective target areas of pedestrians, effectively suppressing background noise in complex monitoring scenarios. As a result, the effectiveness and robustness of feature representation are significantly improved.

C. Dual-Branch Learning Strategy

To concurrently capture the holistic structural morphology of pedestrians and critically subtle attributes, such as clothing logos, backpacks, and footwear styles, a parallel global-local dual-branch learning architecture is formulated immediately downstream of the attention enhancement module. Within the global macroscopic pathway, the SimAM-optimized feature maps are compressed via a Global Average Pooling (GAP) operation into a 1024-dimensional feature vector. This vector functions as a comprehensive global descriptor, encapsulating the entire pedestrian footprint. Concurrently, within the parallel localized pathway, we draw inspiration from the quintessential PCB [3] partition paradigm. Here, the attention-refined feature tensors are uniformly segmented along the vertical axis into six distinct horizontal sub-regions (denoted as P1 to P6). Each isolated

undergoes an independent GAP operation, ultimately yielding six distinct sets of 1024-dimensional localized features. This rigid horizontal-slicing constraint effectively compels the network to scrutinize nuanced variations across different anatomical parts, thereby thoroughly excavating highly discriminative, fine-grained visual clues.

In stark contrast to the vanilla PCB approach, which relies exclusively on localized slice-based feature learning, our dual-branch framework purposefully preserves an additional dedicated conduit for global feature propagation. By keeping the macroscopic structural contours of the subject intact, and ensuring that all localized fine-grained perceptions are anchored upon the high-fidelity, 3D attention-refined feature maps, our design successfully strikes an optimal equilibrium between global structural integrity and local detail sensitivity.

Following the dual-branch feature extraction, the singular global feature set and the six localized feature sets are respectively routed into structurally identical yet parametrically decoupled ClassBlock classification modules. As depicted in the magnified inset of Figure 1, this module adopts the standardized BNNeck topology, a prevalent configuration within the person Re-ID domain. Operationally, the 1024-dimensional features are initially projected into a lower 512-dimensional space via a fully connected (FC) layer. Subsequently, they traverse a 1D Batch Normalization (BatchNorm1d) layer coupled with a Dropout layer (set to a rate of 0.5) to rigorously mitigate the risk of model overfitting. Ultimately, a terminal FC classification layer maps these 512-dimensional features into a 751-dimensional identity prediction vector. This terminal dimensionality is precisely aligned with the total number of unique pedestrian identities present in the training set, thereby facilitating highly accurate identity-discriminative learning.

D. Loss Function

To drive the end-to-end optimization of the network, we orchestrate a joint training paradigm that seamlessly integrates a Label Smoothing Cross-Entropy Loss (ID Loss) alongside a Hard

Sample Mining Triplet Loss. These two distinct objectives are explicitly delegated to govern identity classification and feature metric learning, respectively.

Strictly adhering to the established BNNeck protocol, these dual loss functions are strategically anchored at disparate semantic levels within the network topology. Specifically, the Triplet Loss is enforced upon the 512-dimensional embeddings immediately subsequent to the 1D Batch Normalization. By calibrating relative distances within the Euclidean metric space, it effectively compels the cohesive clustering of intra-class representations while rigorously segregating inter-class distributions. Conversely, the ID Loss provides direct supervisory signals over the 751-dimensional identity prediction logits emitted by the terminal classification layer, thereby enforcing rigid identity categorization constraints.

Ultimately, the overarching empirical risk of the entire framework is formulated as a cumulative fusion of the global branch loss and the six independent localized branch losses. The holistic optimization objective, denoted as L_{total} is mathematically articulated as follows:

$$L_{total} = L_{global_ID} + L_{global_Triplet} + \sum_{i=1}^6 (L_{part_ID}^i + L_{part_Triplet}^i) \quad (3)$$

Through the end-to-end collaborative optimization of the aforementioned joint loss function, the proposed dual-branch network can fully integrate global macroscopic posture information and local fine-grained cues, thereby learning highly discriminative features that are robust to viewpoint changes and local occlusions. The advancement and effectiveness of this feature learning framework will be comprehensively quantified and qualitatively validated in the subsequent experimental chapters.

IV. EXPERIMENT AND ANALYSIS

A. Datasets and Evaluation Protocols

evaluations are conducted utilizing the heavily benchmarked Market-1501 dataset. Curated from a university campus environment, this large-scale corpus is captured via a network of six non-overlapping surveillance cameras, comprising five high-definition sensors and one low-resolution device. In its entirety, the dataset encompasses 32,668 bounding boxes representing 1,501 unique pedestrian identities. Crucially, the captured imagery inherently exhibits severe real-world surveillance impediments, including cluttered backgrounds, intense illumination fluctuations, and unconstrained postural variations.

Adhering stringently to the canonical evaluation protocol established by the original dataset authors, we ensure a mutually exclusive identity split between the training and testing phases. Specifically, the training subset is allocated 12,936 instances corresponding to 751 identities, which are exclusively designated for the end-to-end parameter optimization of the model. Conversely, the testing set is formulated from the remaining 750 unseen identities. This evaluation subset is logically partitioned into a probe (query) set and a gallery set. The former consists of 3,368 query exemplars, while the latter houses 19,732 gallery samples, collaboratively serving as the foundation for the cross-camera retrieval assessments.

To quantitatively gauge the holistic retrieval capabilities of the network, we adopt the standard mean Average Precision (mAP) alongside the Cumulative Matching Characteristics (CMC) curve, specifically reporting the Rank-1, Rank-5, and Rank-10 accuracies, as our primary performance indicators.

B. Implementation Details

The empirical implementation of our proposed algorithmic framework is programmed utilizing the PyTorch 2.8.0 deep learning library, harmonized with CUDA 12.8 for GPU acceleration. All computational pipelines, encompassing both model training and inferential evaluations, are executed on a robust high-performance workstation. This hardware configuration is equipped with a solitary NVIDIA GeForce RTX 4090 GPU (featuring 24GB of VRAM), complemented by an Intel Xeon Gold 6430 processor and 120GB of system memory.

During the data initialization phase, the spatial resolution of all input imagery is uniformly standardized to 256×128 pixels. To fortify the model's generalization capabilities in unconstrained environments and proactively suppress overfitting, a stringent data augmentation protocol, inspired by the standard Strong Baseline, is applied exclusively to the training subset. This protocol sequentially incorporates a 10-pixel zero-padding, random spatial cropping, random horizontal flipping, and Random Erasing techniques to artificially induce data variance.

The network parameters are iteratively optimized employing Stochastic Gradient Descent (SGD) infused with Nesterov momentum. The momentum coefficient is intrinsically fixed at 0.9, paired with a weight decay constraint of 5×10^{-4} to regularize the learning trajectory. Given that the backbone relies on pre-trained weights, a differential learning rate paradigm is judiciously adopted. Specifically, the newly instantiated ClassBlock classifier is assigned an initial learning rate of 0.05, whereas the pre-trained backbone is carefully tuned at a diminished rate of 0.005. This purposeful discrepancy safeguards the foundational low-level features from being catastrophically disrupted during the early training stages.

The holistic training procedure spans a total of 60 epochs, maintaining a mini-batch size of 32. To dynamically govern the learning rate, a step-wise decay scheduler (StepLR) is deployed. Precisely at the 40th epoch, the learning rates are annealed by a factor of 0.1, a deliberate strategy that guarantees a stable and smooth convergence toward the optimal solution in the terminal stages of the training lifecycle.

C. Ablation Study

To rigorously isolate and quantify the individual contributions and indispensability of each architectural module, a comprehensive ablation study is orchestrated upon the Market-1501 benchmark. The quantitative trajectories of this analysis are chronicled in Table I.

Initially, the

protocols of the Strong Baseline [4] substantially fortifies the network's overarching representational capacity, propelling the mAP to an elevated 81.3%. Building upon this fortified foundation, the isolated incorporation of the six-stripe horizontal partitioning pathway (denoted as 6Part) precipitates a pronounced surge, pushing the mAP metric to 86.0%. This substantial leap empirically corroborates the paramount importance of fine-grained local cue excavation in differentiating highly homologous identities.

Interestingly, when the SimAM module is grafted exclusively onto the monolithic global network, the performance gains are relatively marginal and exhibit minor fluctuations. This phenomenon suggests that, in the absence of localized spatial constraints, a purely global attention mechanism is highly susceptible to redundant background distractors, thereby hindering its innate feature-enhancing potential.

In striking contrast, the concurrent deployment of the 3D SimAM attention mechanism alongside the 6Part localized branch engenders a profound synergistic complementarity. This hybrid configuration culminates in peak performance, registering an optimal mAP of 86.1% while maintaining uniformly superior scores across all Rank metrics. Ultimately, these empirical findings compellingly demonstrate that localized fine-grained constraints effectively anchor and amplify the 3D attentional enhancement. This symbiotic interplay ensures precise focal alignment on salient anatomical regions while aggressively filtering environmental noise, thereby fundamentally escalating the overarching discriminative prowess of the formulated model.

Table I. COMPARISON WITH STATE-OF-THE-ART METHODS

ConvNeXt (Baseline)	Strong	SimAM	6Part	Rank @1 (%)	Rank @5 (%)	Rank@ 10 (%)	mAP (%)
√				92.5	97.5	98.5	79.4
√	√			93.0	97.4	98.6	81.3
√	√	√		92.3	97.1	98.5	81.1
√	√		√	95.0	98.2	98.8	86.0
√	√	√	√	94.8	98.3	98.9	86.1

D. Comparison with State-of-the-Art Methods

Table II systematically benchmarks our proposed framework against a spectrum of contemporary, state-of-the-art (SOTA) person Re-ID paradigms. Notably, even in the absence of auxiliary optimization heuristics, such as re-ranking techniques or multi-scale feature fusion, our architecture remarkably secures a Rank@1 accuracy of 94.8% alongside an mAP of 86.1%.

A granular analysis of these metrics reveals a substantial performance margin over classical milestone algorithms, including PCB-RPP and IANet. Furthermore, when juxtaposed with highly efficient, high-performance models like OSNet (84.9% mAP), our approach still yields a commendable accuracy surge. This competitive edge definitively substantiates the architectural superiority and operational efficacy of synergizing the ConvNeXt backbone with the attention-guided dual-branch topology.

Building upon this robust baseline, the integration of the k-reciprocal re-ranking

Ours	94.8	86.1
Ours+RR	95.5	93.4

E. Qualitative Analysis

To intuitively substantiate the cross-camera retrieval prowess of our model within uncontrolled, real-world deployments, Figure 2 visualizes several representative query prototypes alongside their corresponding Top-10 ranked gallery retrievals. Within these visual matrices, authentic cross-view matches sharing the identical identity are demarcated by green perimeters, whereas erroneous gallery distractors are delineated by red borders.

A critical appraisal of these qualitative outcomes reveals that our framework consistently executes precise identity associations, even when confronted with pronounced viewpoint deflections, drastic postural shifts, and severe environmental clutter. This resilience undeniably attests to the superior adaptability and generalized robustness of the extracted representations against diverse empirical perturbations.

Concurrently, the sporadic emergence of red-bounded false positives, typically representing formidable hard negative samples, harbors profound investigative value. A closer inspection of these failure cases unveils that the falsely associated gallery instances predominantly share highly homologous visual attributes with the query, such as matching apparel color palettes, fabric textures, and carried accessories. Paradoxically, this phenomenon reinforces the premise that our network has inherently mastered an acute sensitivity for fine-grained perceptual extraction.

Retrieval discrepancies triggered by such extreme phenotypic homogeneity remain an intrinsic bottleneck universally endemic to the Re-ID domain. Consequently, these challenging failure modes illuminate a definitive trajectory for future endeavors. Subsequent research will strategically pivot towards assimilating multi-modal data streams and excavating deeper, more disparate fine-grained nuances, ultimately optimizing the discriminative fidelity on extraordinarily demanding samples.

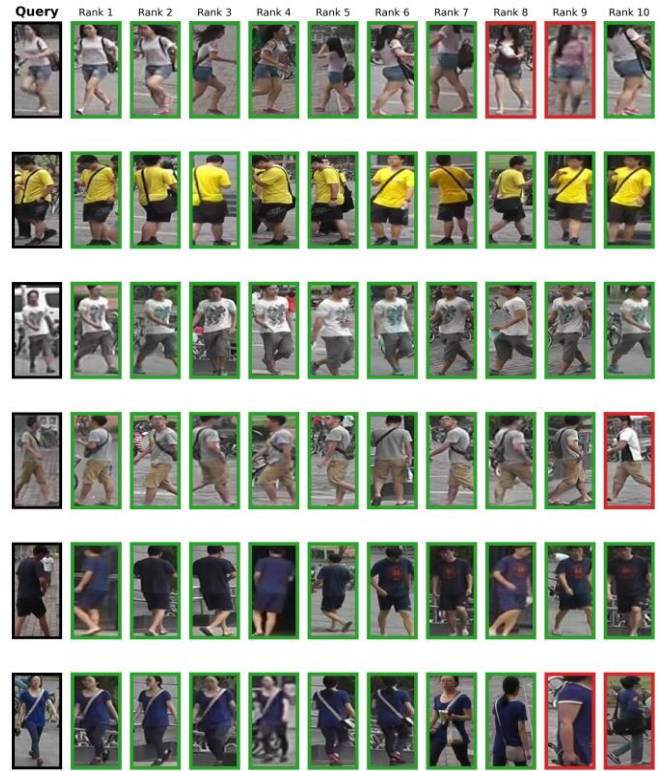


Figure 1. Visualization of the top-10 retrieval results on Market-1501

V. CONCLUSIONS

In response to the quintessential bottlenecks endemic to person Re-ID, namely, drastic cross-camera viewpoint permutations, convoluted environmental clutter, and the formidable disambiguation of highly homologous identities, this research introduces an innovative Re-ID architecture that synergizes a 3D energy-based attention mechanism with a dual-branch local perception paradigm.

Grounded in the ConvNeXt backbone for high-fidelity representation extraction, our methodology seamlessly embeds the parameter-free SimAM module within the forward propagation pipeline. Crucially, without incurring any ancillary parametric or computational overhead, this module adaptively amplifies salient foreground attributes while aggressively dampening extraneous background noise, thereby achieving high-precision focal alignment. Concurrently, the meticulously engineered parallel horizontal-slicing local branch compels the network to thoroughly assimilate fine-grained specificities, such as garment textures and idiosyncratic accessories.

This localized excavation effectively mitigates the inter-class obfuscation induced by highly similar phenotypes, culminating in a substantial escalation of the model's discriminative accuracy.

Extensive empirical validations on the benchmark Market-1501 dataset unequivocally demonstrate that, even independent of intricate post-processing heuristics like re-ranking, the formulated model exhibits exceptional representational discriminability and cross-scenario generalizability. Consequently, our core evaluation metrics significantly eclipse those of numerous existing state-of-the-art algorithms.

Looking ahead, our future research trajectory will seek to extrapolate this highly efficient feature extraction paradigm to more demanding domains, notably video-based and cross-modal person Re-ID tasks. Furthermore, we intend to explore structural miniaturization and lightweight optimization strategies, aspiring to seamlessly tailor the model for the stringent real-time deployment prerequisites of real-world intelligent security infrastructures.

REFERENCES

- [1] Z.Liu, H. Mao, C. -Y. Wu, C. Feichtenhofer, T. Darrell, and S. Xie, "A ConvNet for the 2020s," in Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR), 2022, pp. 11976-11986.
- [2] L. Yang, R. -Y. Zhang, L. Li, and X. Xie, "SimAM: A Simple, Parameter-Free Attention Module for Convolutional Neural Networks," in International Conference on Machine Learning (ICML), 2021, pp. 11863-11874.
- [3] Y. Sun, L. Zheng, Y. Yang, Q. Tian, and S. Wang, "Beyond Part Models: Person Retrieval with Refined Part Pooling (and a Strong Convolutional Baseline)," in Proceedings of the European Conference on Computer Vision (ECCV), 2018, pp. 480-496.
- [4] J. Hu, L. Shen, and G. Sun, "Squeeze-and-Excitation Networks," in Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR), 2018, pp. 7132-7141.
- [5] S. Woo, J. Park, J.-Y. Lee, and I. S. Kweon, "CBAM: Convolutional Block Attention Module," in Proceedings of the European Conference on Computer Vision (ECCV), 2018, pp. 3-19.
- [6] H. Luo, Y. Gu, X. Liao, S. Lai, and W. Jiang, "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), 2019, pp. 148-155.
- [7] Z. Zhong, L. Zheng, D. Cao, and S. Li, "Re-Ranking Person Re-Identification with k-Reciprocal Encoding," in Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR), 2017, pp. 1318-1327.
- [8] G. Wang, S. Yang, H. Liu, Z. Wang, Y. Yang, S. Wang, G. Yu, E. Zhou, and J. Sun, "High-Order Information Matters: Learning Relation and Topology for Occluded Person Re-Identification," in Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR), 2020, pp. 6449-6458.
- [9] Z. Zhuang, L. Wei, L. Xie, T. Zhang, H. Zhang, H. Wu, H. Ai, and Q. Tian, "Rethinking the Distribution Gap of Person Re-Identification with Camera-Based Batch Normalization," in Proceedings of the European Conference on Computer Vision (ECCV), 2020, pp. 140-157. CBN
- [10] R. Hou, B. Ma, H. Chang, X. Gu, S. Shan, and X. Chen, "Interaction-and-Aggregation Network for Person Re-Identification," in Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR), 2019, pp. 9317-9326.
- [11] H. Luo, Y. Gu, X. Liao, S. Lai, and W. Jiang, "A Strong Baseline and Batch Normalization Neck for Deep Person Re-Identification," *IEEE Transactions on Multimedia*, vol. 22, no. 10, pp. 2597-2609, 2020.
- [1