
MAT-Agent: Adaptive Multi-Agent Training Optimization

Jusheng Zhang^{1*}, Kaitong Cai^{1*}, Yijia Fan¹, Ning yuan Liu¹, Keze Wang^{1,†}

¹Sun Yat-sen University

[†]Corresponding author: kezewang@gmail.com

Abstract

Multi-label image classification demands adaptive training strategies to navigate complex, evolving visual-semantic landscapes, yet conventional methods rely on static configurations that falter in dynamic settings. We propose MAT-Agent, a novel multi-agent framework that reimagines training as a collaborative, real-time optimization process. By deploying autonomous agents to dynamically tune data augmentation, optimizers, learning rates, and loss functions, MAT-Agent leverages non-stationary multi-armed bandit algorithms to balance exploration and exploitation, guided by a composite reward harmonizing accuracy, rare-class performance, and training stability. Enhanced with dual-rate exponential moving average smoothing and mixed-precision training, it ensures robustness and efficiency. Extensive experiments across Pascal VOC, COCO, and VG-256 demonstrate MAT-Agent’s superiority: it achieves an mAP of 97.4 (vs. 96.2 for PAT-T), OF1 of 92.3, and CF1 of 91.4 on Pascal VOC; an mAP of 92.8 (vs. 92.0 for HSQ-CvN), OF1 of 88.2, and CF1 of 87.1 on COCO; and an mAP of 60.9, OF1 of 70.8, and CF1 of 61.1 on VG-256. With accelerated convergence and robust cross-domain generalization, MAT-Agent offers a scalable, intelligent solution for optimizing complex visual models, paving the way for adaptive deep learning advancements.

1 Introduction

Multi-label image classification (MLIC) is a fundamental task in computer vision, serving as a cornerstone for applications such as automatic image annotation, scene understanding, and content-based retrieval.[1, 2, 3, 4, 5, 6, 7] The goal is to assign multiple semantically relevant labels to a single image, capturing the intricate correlations among real-world objects and concepts.[8, 9, 10]

Despite its importance, most existing MLIC optimization pipelines follow a “static configuration” or “staged scheduling” paradigm.[11, 12, 13] Under this formulation, training hyperparameters—including data augmentation strategies $\mathcal{T}_{aug}$, optimizer $\mathcal{O}$, learning rate schedule S_{lr} , and loss function $\mathcal{L}_{loss}$ —are typically fixed at the beginning of training, or only undergo heuristic tuning at pre-defined milestones[14, 15, 16]. The training process can thus be formalized as searching for a globally optimal static configuration $\mathbf{C}^* = \{\mathcal{T}_{aug}^*, \mathcal{O}^*, S_{lr}^*, \mathcal{L}_{loss}^*\}$ to maximize validation performance $P_{val}:\mathbf{C}^* = \arg \max_{\mathbf{C} \in \Omega_{configs}} P_{val}(M(\mathcal{D}_{train}; \mathbf{C}))$ [17] where M denotes the model, $\mathcal{D}_{train}$ is the training dataset, and $\Omega_{configs}$ is the configuration space.

However, treating $\mathbf{C}$ as a one-shot static decision fails to account for the inherent dynamics and evolving training state s_t in MLIC[18, 19]. During training, factors such as label co-occurrence patterns, class difficulty, and feature-label mappings evolve over time.[20, 21] Static configurations $\mathbf{C}^*$ are ill-suited to such non-stationarity, often resulting in suboptimal strategies during critical

learning phases.[22, 23, 24, 25] This mismatch may lead to training instability, premature convergence, and ultimately, limits the achievable performance ceiling.

Although recent progress in multi-label image classification (MLIC) has led to significant improvements, a key bottleneck remains: the lack of fine-grained control over the training process.[26, 6, 27, 17, 28, 29] In particular, the ability to dynamically coordinate training components to adapt to the evolving data characteristics and learning stages[30] is still underdeveloped, hindering the full potential of modern models from the following two aspects: i) **Inter-component Coordination**. Conventional approaches often tune components such as data augmentation ($\mathcal{T}_{aug}$), optimizer ($\mathcal{O}$), learning rate scheduler (S_{lr}), and loss function ($\mathcal{L}_{loss}$) independently, overlooking their complex nonlinear interactions. For instance, during the learning of rare yet critical tail classes, aggressive global augmentations may overwhelm weak signals; similarly, certain optimizer-loss combinations may interfere with each other, degrading optimization efficiency or stability; ii) **Searching for Optimal Configurations**. Even with exhaustive offline search methods, such as grid or random search, finding the globally optimal static configuration $\mathbf{C}^*$ is highly challenging due to the combinatorial explosion in high-dimensional, discrete, or hybrid configuration spaces.[31] These methods are not only computationally expensive but also prone to getting stuck in local optima or flat regions of the performance landscape. As a result, they often fail to uncover the **evolving strategy trajectory** that is truly responsible for driving the model toward optimal performance across stages of training[19]. At its core, this challenge is fundamentally a problem of **sequential decision-making under uncertainty**, where the system must learn to balance “exploration” of new opportunities with “exploitation” of current knowledge. This trade-off lies at the heart of intelligent agents interacting with uncertain environments to discover dynamically optimal strategies.

Motivated by the above insights and aiming to fundamentally transcend the limitations of static optimization paradigms, we propose a novel training optimization framework: **MAT-Agent (Multi-Agent Training Agent)**. This framework *reconceptualizes the training process for multi-label image classification (MLIC) as a multi-agent, continual learning and decision-making problem*[32], where each decision stage is governed by principles rooted in the classic exploration-exploitation trade-off. Specifically, MAT-Agent introduces four autonomous and adaptive agents—each responsible for dynamically controlling one of the core training components: data augmentation, optimizer selection, learning rate scheduling, and loss function design. Rather than relying on static heuristics or predefined rules, these agents operate in real time at each training step t : they perceive the global training state s_t and select component-specific actions a_t^k (e.g., a particular augmentation policy or optimizer) from their learned policy $\pi_k(a^k|s_t; \theta_k)$, where θ_k are the trainable parameters.

As a result, the training configuration at time t is assembled as a dynamic combination: $\mathbf{C}_t = \{a_t^{\text{aug}}, a_t^{\text{opt}}, a_t^{\text{lr}}, a_t^{\text{loss}}\}$. Each agent k receives a reward signal $R(s_t, \mathbf{C}_t)$ that quantifies the effectiveness of the current configuration in state s_t , balancing performance gains and training stability. The agents continuously update their decision policies to maximize the expected cumulative discounted reward: $J = \mathbb{E} \left[\sum_{t=0}^T \gamma^t R(s_t, \mathbf{C}_t) \right]$. The key conceptual shift introduced by MAT-Agent lies in its departure from conventional methods, which aim to *predict and fix* a globally optimal static configuration $\mathbf{C}^*$ before training begins. In contrast, **MAT-Agent learns and evolves** a set of adaptive decision policies $\{\pi_k^*\}$ during training, enabling the generation of context-aware configuration sequences $\mathbf{C}_t(s_t)$ conditioned on real-time feedback. This transition—from *static optimization* to *dynamic strategy learning*—empowers the MLIC training pipeline with unprecedented adaptability, intelligence, and efficiency, offering a promising direction for training complex visual models.

2 Related Works

Multi-Label Image Classification. Early methods used ensemble binary classifiers [3, 33, 34, 35, 36, 37] for each category. With the advent of deep learning, CNN-based models like ResNet [15] and SENet [38] greatly improved feature extraction. More recently, Vision Transformer [39] and Swin Transformer [40] have made significant progress in modeling global dependencies. However, challenges remain, particularly in label dependency modeling and class imbalance. Wang et al. [8] proposed the CNN-RNN model for sequential label dependencies; Chen et al. [9] introduced ML-GCN for label correlations via graph convolutions; Lanchantin et al. [41] enhanced label relationship modeling through TDRG. To address class imbalance, Lin et al. [42] proposed Focal Loss, which adjusts sample weights; Ridnik et al. [43] optimized sample weights using ASL loss; and Wu et al. [44] introduced DB loss for distribution alignment. However, most methods use static strategies,

limiting adaptability to dynamic training processes and the complex, evolving label relationships in multi-label classification.

Adaptive Training Optimization Methods. Training optimization in deep learning typically involves four components: data augmentation, optimization algorithms, learning rate scheduling, and loss functions. In data augmentation, AutoAugment [45] uses reinforcement learning to optimize strategies; Fast AutoAugment [46] improves search efficiency; and CutMix [47] generates new samples through region mixing. For optimization, Adam [48] adapts learning rates; AdamW [49] improves generalization by decoupling weight decay; RAdam [50] stabilizes training with a rectification term. In learning rate scheduling, cyclical learning rates and the One-Cycle policy [51] speed up convergence, while SGDR [18] avoids local optima with periodic restarts. Loss functions such as Focal Loss [52] and GHM loss [53] address the class imbalance. While these methods optimize individual components, they overlook the synergistic effects between them. Approaches like AutoML [54] and ENAS [55] attempt joint optimization but are limited by high computational costs. Our MAT-Agent proposed here offers adaptive optimization via multi-agent collaborative decision-making, without added search overhead.

Multi-Agent Decision Systems. Multi-agent systems are effective for complex decision-making, especially in uncertain environments. Zhang et al. [56] applied multi-agent reinforcement learning to distributed control; QMIX [57, 58] by Rashid et al. enables collaborative decision-making via value function decomposition; VDN [59] by Sunehag et al. achieves cooperation through value decomposition. Algorithms like UCB [60] and Thompson sampling [61] improve decision efficiency and exploration-exploitation balance, while multi-armed bandit theory [62] provides a foundation for decision-making under uncertainty.

3 Methodology

3.1 Sequential Decision Formulation for MLIC Training Optimization

We formulate the training of a multi-label image classification (MLIC) model with parameters Θ_M as a sequential decision-making process. At each decision step t (e.g., per training epoch or fixed iteration interval), the system resides in a training state $s_t \in \mathcal{S}$, which encapsulates key information regarding model learning progress, data characteristics, and training dynamics at time t .

Based on s_t , MAT-Agent selects a composite action $\mathbf{C}_t \in \mathcal{C}$ that defines the training configuration to be applied in the next stage. This configuration includes the data augmentation policy $\mathcal{T}_{aug}^{(t)}$, optimizer $\mathcal{O}^{(t)}$, learning rate scheduler $\mathcal{S}_{lr}^{(t)}$, and loss function $\mathcal{L}_{loss}^{(t)}$.

After executing one training step under $\mathbf{C}_t$, the system transitions to a new state s_{t+1} and receives a scalar reward $R_{t+1} = R(s_t, \mathbf{C}_t, s_{t+1})$, which measures the immediate contribution of the chosen configuration to model improvement and training stability. The goal of MAT-Agent is to learn an optimal joint policy $\{\pi_k^*\}_{k=1}^N$, where $N = 4$ corresponds to the number of training components (each managed by an individual agent k), in order to maximize the expected cumulative discounted reward:

$$J(\{\theta_k\}_{k=1}^N) = \mathbb{E}_{\{\pi_k(\cdot|s_t; \theta_k)\}_{k=1}^N, p(s_{t+1}|s_t, \mathbf{C}_t)} \left[\sum_{t=0}^T \gamma^t R_{t+1} \right]$$

Here, $\{\theta_k\}_{k=1}^N$ denotes the set of learnable parameters of all agent policies. The expectation $\mathbb{E}$ is taken over both the policy-induced action distributions $\pi_k(\cdot|s_t; \theta_k)$ and the environment’s state transition probabilities $p(s_{t+1}|s_t, \mathbf{C}_t)$. The discount factor $\gamma \in [0, 1]$ controls the trade-off between immediate and future rewards, and T is the total number of decision steps. This objective formally defines the learning target of MAT-Agent: to iteratively adapt the agent policies $\{\pi_k\}$ such that the induced configuration sequence $\{\mathbf{C}_t\}$ maximizes the cumulative expected return over the training horizon.

3.2 MAT-Agent: Framework, Actions, and State Representation

As shown in Figure 3.1, our MAT-Agent is designed as a Multi-Agent System (MAS) that performs decentralized control and coordinated learning to dynamically orchestrate the multi-label image classification (MLIC) training process. The system targets $N = 4$ critical training components known to significantly influence model performance and require adaptive control. We formalize these

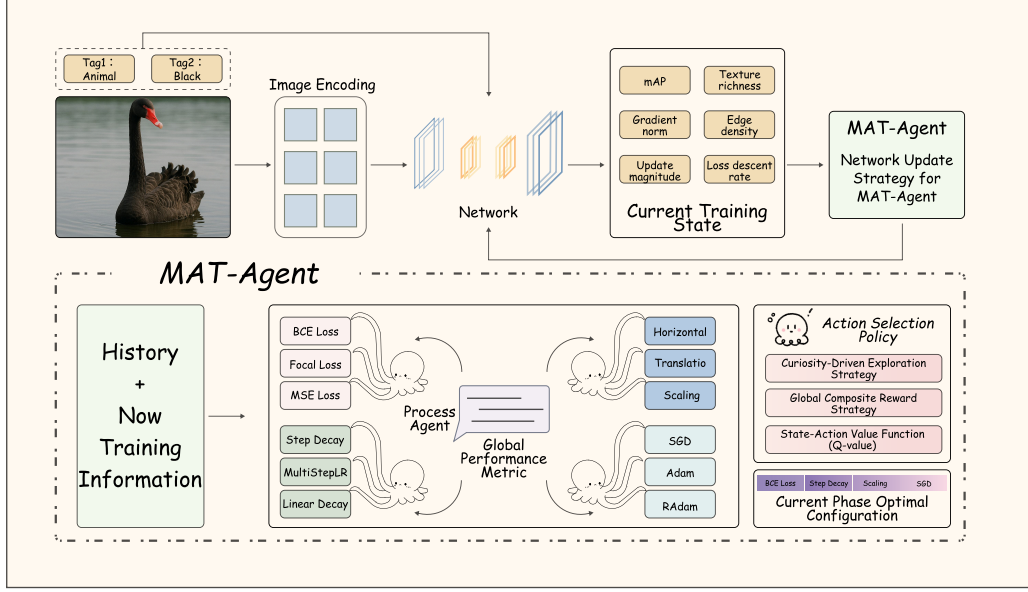


Figure 1: Framework of **MAT-Agent**: a multi-agent system that dynamically selects training strategies (augmentation, optimizer, scheduler, loss) based on current and historical training states to optimize multi-label classification.

components as a set $\mathcal{K} = \{k_i\}_{i=1}^N$, where $k_1 = \text{AUG}$ (data augmentation), $k_2 = \text{OPT}$ (optimizer selection), $k_3 = \text{LRS}$ (learning rate scheduling), and $k_4 = \text{LOSS}$ (loss function design).

Correspondingly, MAT-Agent maintains N autonomous and adaptive agents, denoted as $\{\text{Agent}_k\}_{k \in \mathcal{K}}$. At each decision step t , each agent Agent_k selects an action a_t^k from its discrete action space $\mathcal{A}_k = \{a_k^{(1)}, a_k^{(2)}, \dots, a_k^{(M_k)}\}$, which contains $M_k = |\mathcal{A}_k|$ predefined candidate strategies (The complete list of candidate strategies for each agent is provided in Supp. A.1). The joint actions of all agents collectively form the training configuration at step t :

$$\mathbf{C}_t = (a_t^{\text{AUG}}, a_t^{\text{OPT}}, a_t^{\text{LRS}}, a_t^{\text{LOSS}})$$

As a result, the global configuration space $\mathcal{C}$ explored by the system is the Cartesian product of the individual action spaces:

$$\mathcal{C} = \mathcal{A}_{\text{AUG}} \times \mathcal{A}_{\text{OPT}} \times \mathcal{A}_{\text{LRS}} \times \mathcal{A}_{\text{LOSS}}, \quad |\mathcal{C}| = \prod_{k \in \mathcal{K}} |\mathcal{A}_k|$$

To enable effective and adaptive decision-making, MAT-Agent constructs a comprehensive state representation. At each step t , the instantaneous state of the system is represented by a D -dimensional real-valued vector $s_t \in \mathcal{S} \subseteq \mathbb{R}^D$, encapsulating both the performance of the MLIC model (with parameters Θ_M) and dynamic characteristics of the training process. The state vector is structured as: $s_t = [s_t^{\text{perf}}; s_t^{\text{dyn}}; s_t^{\text{data}}]$. Here, s_t^{perf} captures performance indicators such as validation mean average precision mAP_t^{val} ; s_t^{dyn} includes training dynamics such as training/validation loss $L_t^{\text{train}}, L_t^{\text{val}}$, loss change ΔL_t^{val} , and gradient statistics of the training loss with respect to model parameters $g_t = \nabla_{\Theta_M} L_t^{\text{train}}$ (e.g., L_2 norm $\|g_t\|_2$), as well as relative update magnitudes; s_t^{data} includes dataset-specific descriptors such as average texture richness of current samples.

To support temporal reasoning, the actual agent input is an extended state representation $\mathcal{I}_t$, formed by aggregating both current and historical observations (Details on the construction of $\mathcal{I}_t$ and the full feature set for s_t can be found in Supp. A.2).

3.3 Decision-Making and Learning of Decentralized Adaptive Agents

Each agent Agent_k ($k \in \mathcal{K}$) in MAT-Agent independently learns a parameterized decision policy $\pi_k(a^k | \mathcal{I}_t; \theta_k)$ that maps the extended state representation $\mathcal{I}_t$ to an optimal action $a_t^k \in \mathcal{A}_k$. We

adopt a value-based reinforcement learning framework, specifically building upon Deep Q-Networks (DQN) and its variants. The core objective is to train each agent to approximate a state-action value function $Q_k(\mathcal{I}_t, a; \theta_k)$, estimating the expected cumulative reward of taking action a in state $\mathcal{I}_t$. This function is realized by a deep neural network parameterized by θ_k , taking $\mathcal{I}_t$ as input and outputting Q-values for all discrete actions $a \in \mathcal{A}_k$ (The specific architectures of these Q-networks are described in Supp. A.3).

To balance exploration and exploitation, agents employ an ϵ -greedy strategy: with probability $1 - \epsilon_t$, the agent selects the action with the highest Q-value, $a_t^k = \arg \max_{a \in \mathcal{A}_k} Q_k(\mathcal{I}_t, a; \theta_k)$; with probability ϵ_t , it samples an action randomly from $\mathcal{A}_k$. The exploration rate ϵ_t decays over time to promote early exploration and later convergence. To further improve exploration, a curiosity-driven intrinsic reward mechanism is introduced based on prediction error of state transitions.

For stable and efficient training, we integrate experience replay and target Q-networks. The agent's experiences, represented as tuples $(\mathcal{I}_j, a_j^k, R_{j+1}, \mathcal{I}_{j+1})$, are stored in a shared (or individual) replay buffer $\mathcal{D}$. Mini-batches are sampled from $\mathcal{D}$ to update the Q-network parameters θ_k . The update minimizes the temporal difference (TD) error with the following loss:

$$L_j(\theta_k) = (y_j - Q_k(\mathcal{I}_j, a_j^k; \theta_k))^2 \quad (1)$$

The TD target y_j is computed as:

$$y_j = R_{j+1} + \gamma \max_{a' \in \mathcal{A}_k} Q_k(\mathcal{I}_{j+1}, a'; \theta_k^-) \quad (2)$$

where R_{j+1} is the global reward obtained from executing the joint action $\mathbf{C}_j$ (which includes a_j^k), γ is the discount factor, and $Q_k(\mathcal{I}_{j+1}, a'; \theta_k^-)$ is the target Q-network estimate.

The Q-network is optimized by minimizing the expected TD loss $\mathbb{E}[L_j(\theta_k)]$ using optimizers such as Adam. The shared reward signal R_{t+1} evaluates the overall effectiveness of a joint training configuration after each epoch. We define a composite reward function:

$$R_{t+1} = w_{\text{mAP}} \cdot f(\Delta \text{mAP}_t) + w_{\text{stab}} \cdot \text{Stability}_t + w_{\text{conv}} \cdot \text{Convergence}_t - w_{\text{pen}} \cdot \text{Penalty}_t \quad (3)$$

Here, ΔmAP_t measures change in validation accuracy, with $f(\cdot)$ shaping the reward to boost significant improvements. Stability_t inversely relates to loss fluctuation; Convergence_t tracks convergence speed (e.g., loss reduction rate); and Penalty_t penalizes unstable or computationally expensive configurations. The weights $w_{\text{mAP}}, w_{\text{stab}}, w_{\text{conv}}, w_{\text{pen}}$ are used to balance optimization objectives, ensuring that all agents align toward maximizing overall training efficiency (The precise mathematical definitions for each component of R_{t+1} and the specific values used for the weights are detailed in Supp. A.4).

3.4 Dynamic Training Configuration Generation and MLIC Model Update

MAT-Agent operates as an iterative closed-loop learning and control framework designed to dynamically optimize the training process of multi-label image classification (MLIC). At each decision step t , this loop proceeds through the following key stages:

First, the system performs comprehensive perception of the current training environment to extract and construct an informative extended state representation $\mathcal{I}_t$, which serves as critical contextual input for downstream intelligent decisions. Conditioned on this shared and dynamically updated state $\mathcal{I}_t$, each autonomous agent Agent_k ($k \in \mathcal{K}$) activates its online-learned, parameterized policy $\pi_k(\cdot | \mathcal{I}_t; \theta_k)$ to independently select an optimal training action $a_t^k \in \mathcal{A}_k$. These distributed decisions are then efficiently integrated to form the global training configuration for the current step:

$$\mathbf{C}_t = (a_t^{\text{AUG}}, a_t^{\text{OPT}}, a_t^{\text{LRS}}, a_t^{\text{LOSS}}). \quad (4)$$

This dynamically assembled configuration $\mathbf{C}_t$ serves as the "execution plan" for the next training cycle and is immediately applied to guide the optimization of the main MLIC model (parameterized by Θ_M). The resulting training step updates Θ_M accordingly.

After completing the training iteration, the system evaluates the observed outcomes and environmental shifts, thereby calculating the next system state s_{t+1} (and corresponding extended state $\mathcal{I}_{t+1}$) and

computing a global reward signal R_{t+1} . This scalar reward quantitatively measures the effectiveness and impact of the executed configuration $\mathbf{C}_t$.

Each agent Agent_k then uses the full interaction tuple $(\mathcal{I}_t, a_t^k, R_{t+1}, \mathcal{I}_{t+1})$ to update its internal policy network parameters θ_k via the learning algorithm detailed in Section 3.3. This adaptive cycle of *perception* $\rightarrow$ *decision* $\rightarrow$ *execution* $\rightarrow$ *evaluation* $\rightarrow$ *learning* repeats iteratively, enabling MAT-Agent to capture and respond to the non-stationary nature of the training landscape.

It is worth noting that MAT-Agent does not alter the low-level parameter update rule (e.g., gradient descent mechanics) of the main model Θ_M . Instead, it improves the overall quality and effectiveness of this optimization process by dynamically selecting high-level training configurations $\mathbf{C}_t$, thereby steering the MLIC model toward better generalization and higher stability L_{mem} (see Supp. A.2 for layer selection criteria).

4 Experiments

4.1 Comparative Experiments: Multi-Dataset Evaluation

To comprehensively evaluate the performance of MAT-Agent in multi-label image classification, we conduct extensive comparative experiments on three representative datasets: Pascal VOC[63], MS-COCO[64], and Visual Genome (VG-256)[65]. We benchmark against eight state-of-the-art multi-label classification models, including ML-GCN[66], C-Tran[67], BalanceMix, ASL[68], ML-Decoder[69], MLBOTE, HSQ-CvN[70], and PAT-T[71]. All models are trained with the same backbone and optimization settings. Performance is assessed using three widely adopted metrics: mean Average Precision (mAP)[72], Overall F1 (OF1)[73], and Class-wise F1 (CF1)[73].

Table 1: Comparison of MAT-Agent and baseline models on Pascal VOC, COCO, and VG-256 datasets using mAP, Overall-F1 (OF1), and Class-wise F1 (CF1) metrics. Bold highlights the best results in each column.

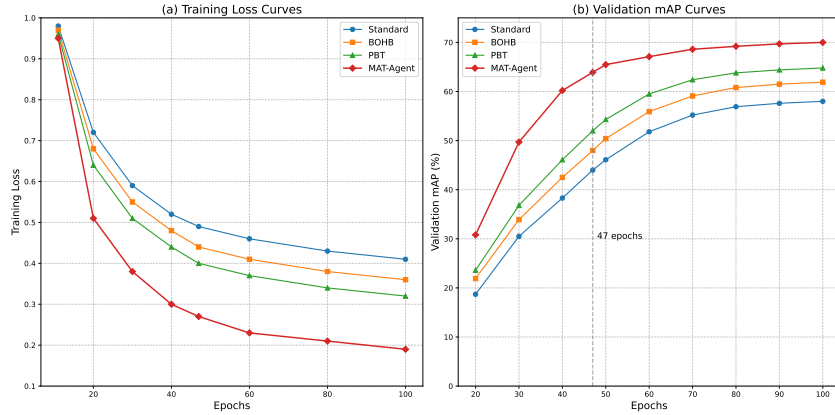
Method	Pascal VOC			COCO			VG-256		
	mAP	OF1	CF1	mAP	OF1	CF1	mAP	OF1	CF1
ML-GCN	94.0	86.4	86.1	83.0	80.3	78.0	52.3	61.4	60.8
C-Tran	94.2	88.1	87.7	85.1	81.7	79.9	55.4	63.1	62.7
BalanceMix	94.7	87.9	87.6	85.2	81.4	80.1	55.6	62.8	62.8
ASL	95.8	88.4	88.4	86.6	81.9	81.4	56.3	63.5	63.1
ML-Decoder	96.1	86.4	85.9	91.2	76.9	76.8	57.9	58.2	57.9
MLBOTE	93.8	87.1	86.4	84.1	80.6	78.6	53.4	62.9	62.2
HSQ-CvN	96.4	—	—	92.0	87.5	86.6	—	—	—
PAT-T	96.2	91.1	90.6	91.8	87.6	86.4	59.5	69.8	59.7
MAT-Agent	97.4	92.3	91.4	92.8	88.2	87.1	60.9	70.8	61.1

As shown in Table 1, MAT-Agent consistently achieves the best results across all datasets and evaluation metrics, demonstrating strong generalization and robustness. On Pascal VOC, MAT-Agent obtains an mAP of 97.4, with OF1 and CF1 reaching 92.3 and 91.4 respectively, clearly outperforming the closest competitor PAT-T (91.1 and 90.6). On the COCO dataset, MAT-Agent also leads in mAP (92.8) and OF1 (88.2), while maintaining a strong CF1 (87.1) comparable to HSQ-CvN. **MAT-Agent** leads multi-label classification on Pascal VOC, COCO, and VG-256, achieving top mAP (97.4, 92.8, 60.9), OF1 (92.3, 88.2, 70.8), and CF1 (91.4, 87.1, 61.1). It outperforms PAT-T and HSQ-CvN, notably by over 1 point in OF1 on VG-256’s 256-class, long-tail setting, demonstrating robust generalization.s

4.2 Training Convergence Analysis

To evaluate the training efficiency and optimization behavior of **MAT-Agent**, we compare it with standard training, Population-Based Training (PBT) [19], and BOHB [74]. Figure 2 shows the training loss and validation mAP curves on the MS-COCO dataset.(Extended comparisons with AutoML baselines such as ENAS are presented in Supp. A.4.)

Figure 2: Training loss (a) and mAP (b) changes for MAT-Agent and three baseline models. MAT-Agent shows faster convergence.



In the first 15 epochs, all methods exhibit similar loss descent. From epoch 15 onward, MAT-Agent converges faster with smoother loss reduction. By epoch 30, its loss curve flattens, while others continue to fluctuate—particularly BOHB. As shown in Figure 2(b), MAT-Agent reaches 63.8% mAP in just 47 epochs, whereas standard training requires about 80 epochs to achieve the same performance, yielding a 47% reduction in training time. At the 47-epoch mark, MAT-Agent achieves 67.3% mAP, outperforming standard training, PBT, and BOHB by 3.5, 2.2, and 1.5 points, respectively. MAT-Agent also demonstrates greater stability and generalization in later stages, avoiding overfitting or oscillation seen in baselines like PBT. Its adaptive strategy mechanism dynamically adjusts components (e.g., optimizer, augmentation, loss) based on intermediate feedback. Overall, MAT-Agent accelerates convergence while maintaining robust performance, making it well-suited for real-world applications with limited resources.

4.3 Cross-dataset Generalization Ability.

To assess MAT-Agent’s generalization and cross-domain adaptability, we perform cross-dataset transfer experiments. Models trained on MS-COCO (ResNet-101) are tested on Pascal VOC, NUS-WIDE [75], and OpenImages. Large datasets (>10,000 images) use random sampling, while multi-label datasets (NUS-WIDE, OpenImages, Visual Genome) use stratified sampling with three repeats. Evaluation focuses on overlapping categories (e.g., 20 shared classes for MS-COCO and Pascal VOC), using mAP and Rare-F1 metrics. Baselines include PBT, BOHB, and DARTS [76].

Table 2: Results of the data migration. The MAT-Agent demonstrates the best generalization ability.

Method	MS-COCO → VOC	MS-COCO → NUS-WIDE*	MS-COCO → OpenImages*
PBT	72.3	58.5	49.7
BOHB	73.1	59.2	50.3
DARTS	73.8	59.7	50.8
MAT-Agent	76.2	62.5	53.4

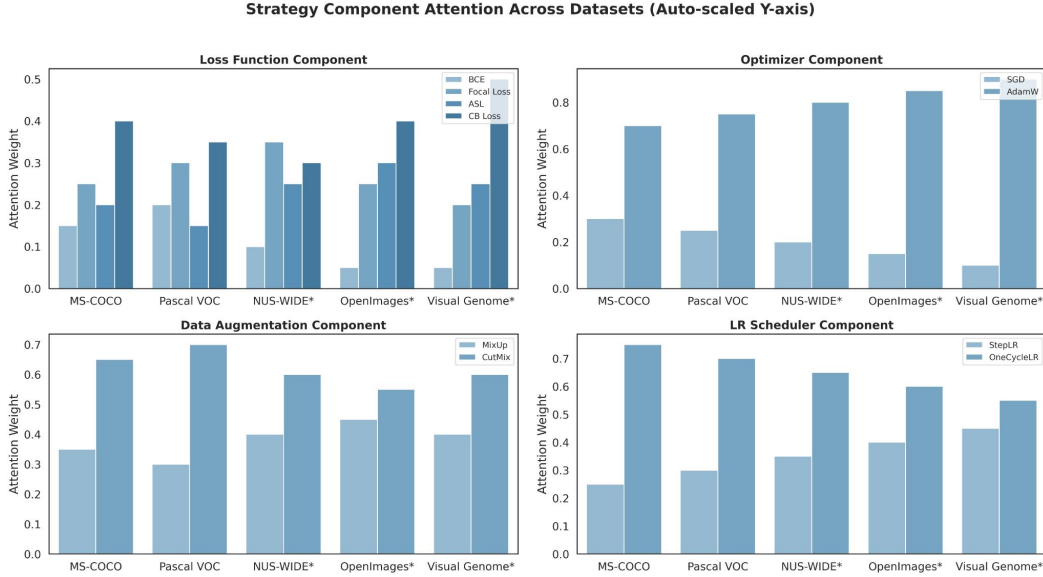
Results MAT-Agent excels in zero-shot transfer, achieving mAP of 76.2% on Pascal VOC (vs. DARTS at 73.8%), and surpassing DARTS by 2.8 and 2.6 points on NUS-WIDE and OpenImages, respectively (Table 2). Despite domain gaps, MAT-Agent maintains a 2.5–3.0 point lead over baselines, proving robust transferability.

4.4 Analysis of Differences between Different Domains.

Figure 3 illustrates the distribution of attention weights assigned to different training components by MAT-Agent across various datasets. As the discrepancy between the target dataset and the source domain increases, MAT-Agent automatically adjusts the attention allocation among training strategies to adapt to the new data distribution. Specifically, in the Visual Genome dataset, which exhibits

a severe class imbalance due to its long-tailed distribution, MAT-Agent significantly increases the attention weight assigned to the class-balanced loss (CB Loss) to enhance learning on rare categories. In contrast, for the OpenImages dataset, which presents higher visual complexity and diversity, MAT-Agent assigns more attention to the CutMix augmentation strategy, indicating that stronger data augmentation is beneficial to model robustness under such conditions. Meanwhile, we observe that the AdamW optimizer and OneCycleLR scheduler consistently receive high attention weights across both the source domain (MS-COCO) and all target domains. This suggests that these components are robust and consistently effective in cross-domain settings, and MAT-Agent continues to rely on them due to their consistently positive impact on model performance regardless of domain shifts.

Figure 3: The distribution of policy attention weights of the MAT-Agent on different datasets



The variation in attention distributions indicates MAT-Agent’s ability to adaptively select training strategies based on domain-specific characteristics. It reinforces long-tail-aware loss functions in imbalanced domains, enhances data augmentation in visually complex scenarios, and consistently leverages effective optimizers and schedulers across domains. This flexibility enables MAT-Agent to outperform static baselines that lack such adaptability under domain shifts. Moreover, MAT-Agent incorporates a smooth transition mechanism that gradually adjusts attention weights, avoiding instability from abrupt strategy changes. Even under large domain gaps, the agent transitions steadily toward more suitable configurations, ensuring stable convergence and robust performance.

4.5 Quantitative Comparisons with Existing Automated Methods.

To evaluate MAT-Agent, we compared it with mainstream automated training methods (hyperparameter optimization, learning rate scheduling, meta-learning) on MS-COCO using ResNet-101 and standard metrics. Hyperparameter optimization included Grid Search, Random Search, Bayesian Optimization, and BOHB; learning rate scheduling used AutoLR; meta-learning covered PBT, Auto-PyTorch, and DARTS. Inter-agent coordination and reward sharing analysis is in Supp. A.1, component-wise state input ablation in Supp. A.2, and baseline configurations in Supp. B.1. According to Figure 4, MAT-Agent significantly outperforms other existing mainstream models in automated strategy selection. Specifically, when the GPU Hours = 10^2 , the mean Average Precision (mAP) reaches approximately 62.5, and the Rare-F1 score reaches approximately 40.3, which is far ahead of other models. Moreover, when the GPU Hours are between 10^2 and 10^3 , MAT-Agent reaches a converged state, indicating that it has learned the optimal strategy for the task. In contrast, other models are still learning strategies and have not converged.

MAT-Agent vs All Baselines (with Convergence Behavior)

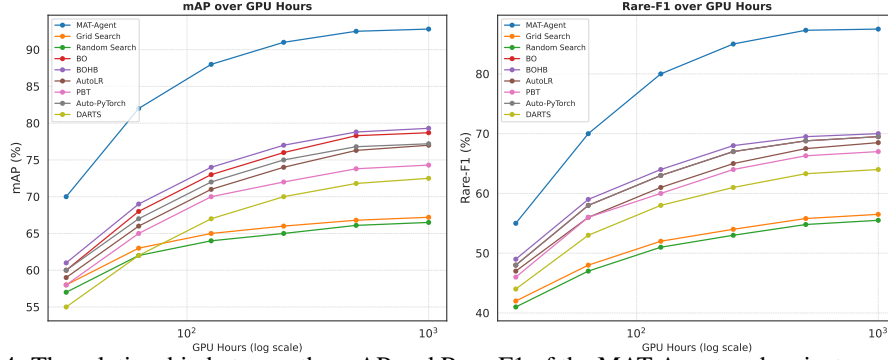


Figure 4: The relationship between the mAP and Rare-F1 of the MAT-Agent and mainstream baseline models over time during training

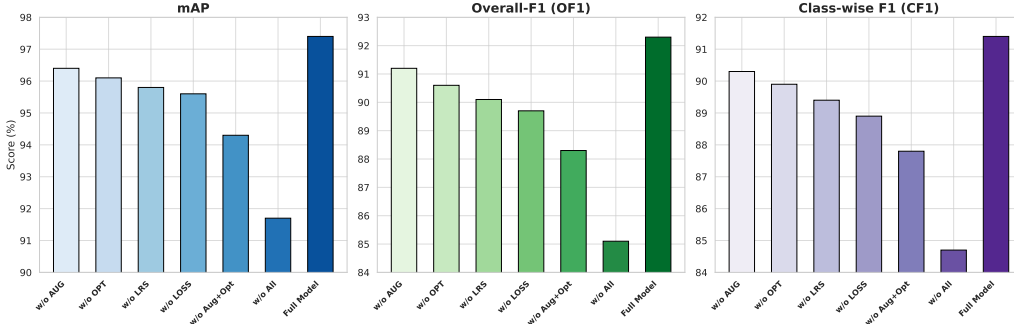


Figure 5: Ablation results on Pascal VOC. Removing any component degrades performance, while the full MAT-Agent achieves the best results across all metrics.

4.6 Ablation Study

To comprehensively evaluate the contribution of each component in MAT-Agent, we conduct a series of ablation experiments on the Pascal VOC dataset by removing each core adaptive agent (AUG, OPT, LRS, LOSS), their combinations, as well as the coordination mechanism among agents. The results are reported across three key metrics: mAP, OF1, and CF1, as illustrated in Figure 5. Further details on the specific configurations of agents and their action spaces used in the ablation settings are available in Supp. B.2. Ablation studies demonstrate that each component of MAT-Agent plays a critical role in performance gains. Removing the augmentation agent (w/o AUG) significantly reduces robustness to diverse and long-tail samples, resulting in lower F1 scores. Disabling the optimizer selection agent (w/o OPT) leads to slower convergence and lower final accuracy, confirming the necessity of dynamic optimization strategies. Eliminating the learning rate scheduling agent (w/o LRS) hampers performance in later training stages. Without the loss function agent (w/o LOSS), the model struggles with class imbalance, causing a clear drop in CF1. Removing multiple agents simultaneously (e.g., w/o AUG+OPT or w/o All) causes a sharp performance drop (mAP down to 91.7%), revealing strong nonlinear synergy among components. Even when all agents are present, disabling their coordination (w/o Agent Coordination) still leads to noticeable degradation, underscoring the importance of inter-agent collaboration.

5 Conclusion

In this paper, we propose a brand-new multi-label classification framework guided by multiple agents, aiming to model label inter-dependencies and overcome local optima in sparse reward scenarios. Our approach employs a collaborative agent architecture where specialists handle different label aspects, capturing correlations through structured communication channels and attention mechanisms. Future work will optimize agent collaboration protocols, extend to extreme multi-label classification, and explore zero-shot label adaptation capabilities.

Acknowledgment

This work was supported in part by the National Natural Science Foundation of China (NSFC) under Grant 62276283, in part by the China Meteorological Administration’s Science and Technology Project under Grant CMAJBGS202517, in part by Guangdong Basic and Applied Basic Research Foundation under Grant 2023A1515012985, in part by Guangdong-Hong Kong-Macao Greater Bay Area Meteorological Technology Collaborative Research Project under Grant GHMA2024Z04, in part by Fundamental Research Funds for the Central Universities, Sun Yat-sen University under Grant 23hytd006, and in part by Guangdong Provincial High-Level Young Talent Program under Grant RL2024-151-2-11.

References

- [1] G. Tsoumakas, I. Katakis, and I. Vlahavas, *Mining Multi-label Data*, 07 2010, pp. 667–685.
- [2] Y. Gong, Y. Jia, T. Leung, A. Toshev, and S. Ioffe, “Deep convolutional ranking for multilabel image annotation,” 2014. [Online]. Available: <https://arxiv.org/abs/1312.4894>
- [3] M. R. Boutell, J. Luo, X. Shen, and C. M. Brown, “Learning multi-label scene classification,” *Pattern Recognition*, vol. 37, no. 9, pp. 1757–1771, 2004.
- [4] X. Shen, M. R. Boutell, J. Luo, and C. M. Brown, “Multi-label machine learning and its application to semantic scene classification,” in *Proceedings of the IS&T/SPIE Electronic Imaging*, San Jose, CA, 2004, pp. 188–199.
- [5] M.-L. Zhang and Z.-H. Zhou, “A review on multi-label learning algorithms,” *IEEE Transactions on Knowledge and Data Engineering*, vol. 26, no. 8, pp. 1819–1837, 2014.
- [6] A. N. Tarekegn, M. Ullah, and F. A. Cheikh, “Deep learning for multi-label learning: A comprehensive survey,” *arXiv preprint arXiv:2401.16549*, 2024.
- [7] J. Zhang, Y. Fan, K. Cai, J. Yang, J. Yao, J. Wang, G. Qu, Z. Chen, and K. Wang, “Why keep your doubts to yourself? trading visual uncertainties in multi-agent bandit systems,” 2026. [Online]. Available: <https://arxiv.org/abs/2601.18735>
- [8] J. Wang, Y. Yang, J. Mao, Z. Huang, C. Huang, and W. Xu, “Cnn-rnn: A unified framework for multi-label image classification,” 2016. [Online]. Available: <https://arxiv.org/abs/1604.04573>
- [9] Z.-M. Chen, X.-S. Wei, P. Wang, and Y. Guo, “Multi-label image recognition with graph convolutional networks,” in *2019 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, 2019, pp. 5172–5181.
- [10] —, “Multi-Label Image Recognition with Graph Convolutional Networks,” in *The IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, 2019.
- [11] E. Ben-Baruch, T. Ridnik, N. Zamir, A. Noy, I. Friedman, M. Protter, and L. Zelnik-Manor, “Asymmetric loss for multi-label classification,” 2020.
- [12] Z.-M. Chen, X.-S. Wei, P. Wang, and Y. Guo, “Multi-label image recognition with graph convolutional networks,” in *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, 2019, pp. 5177–5186.
- [13] X. Liang, H. Liu, and J. Wang, “Reinforced multi-label image classification by exploring curriculum,” in *Proceedings of the AAAI Conference on Artificial Intelligence*, vol. 34, no. 07, 2020, pp. 12 345–12 352.
- [14] A. Krizhevsky, I. Sutskever, and G. E. Hinton, “Imagenet classification with deep convolutional neural networks,” in *Advances in Neural Information Processing Systems (NeurIPS)*, vol. 25, 2012, pp. 1097–1105.
- [15] K. He, X. Zhang, S. Ren, and J. Sun, “Deep residual learning for image recognition,” in *2016 IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, 2016, pp. 770–778.
- [16] K. Simonyan and A. Zisserman, “Very deep convolutional networks for large-scale image recognition,” *International Conference on Learning Representations (ICLR)*, 2015.
- [17] E. D. Cubuk, B. Zoph, D. Mane, V. Vasudevan, and Q. V. Le, “Autoaugment: Learning augmentation policies from data,” 2019. [Online]. Available: <https://arxiv.org/abs/1805.09501>
- [18] I. Loshchilov and F. Hutter, “Sgdr: Stochastic gradient descent with warm restarts,” 2017. [Online]. Available: <https://arxiv.org/abs/1608.03983>
- [19] M. Jaderberg, V. Dalibard, S. Osindero, W. M. Czarnecki, J. Donahue, A. Razavi, O. Vinyals, T. Green, I. Dunning, K. Simonyan, C. Fernando, and K. Kavukcuoglu, “Population based training of neural networks,” 2017. [Online]. Available: <https://arxiv.org/abs/1711.09846>

- [20] M. Marsden, K. McGuinness, J. Antony, H. Wei, M. Redzic, J. Tang, Z. Hu, A. Smeaton, and N. E. O'Connor, "Investigating class-level difficulty factors in multi-label classification problems," 2020. [Online]. Available: <https://arxiv.org/abs/2005.00430>
- [21] Y. Bengio, J. Louradour, R. Collobert, and J. Weston, "Curriculum learning," in *Proceedings of the 26th Annual International Conference on Machine Learning (ICML)*. Montreal, Quebec, Canada: ACM Press, 2009, pp. 1–8.
- [22] Z. Huang, H. Wang, E. P. Xing, and D. Huang, "Self-challenging improves cross-domain generalization," 2020. [Online]. Available: <https://arxiv.org/abs/2007.02454>
- [23] J. Ye, J. He, X. Peng, W. Wu, and Y. Qiao, "Attention-driven dynamic graph convolutional network for multi-label image recognition," 2020. [Online]. Available: <https://arxiv.org/abs/2012.02994>
- [24] J. Tang, J. Zhang, Q. Lv, S. Liu, J. Yang, C. Tang, and K. Wang, "Hiva: Self-organized hierarchical variable agent via goal-driven semantic-topological evolution," 2025. [Online]. Available: <https://arxiv.org/abs/2509.00189>
- [25] J. Zhang, Y. Fan, K. Cai, X. Sun, and K. Wang, "Osc: Cognitive orchestration through dynamic knowledge alignment in multi-agent llm collaboration," 2025. [Online]. Available: <https://arxiv.org/abs/2509.04876>
- [26] C. Qian, K. Han, J. Wang, Z. Yuan, C. Lyu, J. Chen, and Z. Liu, "Dyncim: Dynamic curriculum for imbalanced multimodal learning," 2025. [Online]. Available: <https://arxiv.org/abs/2503.06456>
- [27] M. Feurer and F. Hutter, "Hyperparameter optimization: A review of algorithms and applications," *Proceedings of the IEEE*, vol. 108, no. 3, pp. 742–767, 2020.
- [28] J. Zhang, K. Cai, Q. Zeng, N. Liu, S. Fan, Z. Chen, and K. Wang, "Failure-driven workflow refinement," 2025. [Online]. Available: <https://arxiv.org/abs/2510.10035>
- [29] J. Zhang, K. Cai, J. Yang, and K. Wang, "Learning dynamics of vlm finetuning," 2025. [Online]. Available: <https://arxiv.org/abs/2510.11978>
- [30] E. Real, C. Liang, D. R. So, and Q. V. Le, "Automl-zero: Evolving machine learning algorithms from scratch," 2020. [Online]. Available: <https://arxiv.org/abs/2003.03384>
- [31] T. Elsken, J. H. Metzen, and F. Hutter, "Neural architecture search: A survey," 2019. [Online]. Available: <https://arxiv.org/abs/1808.05377>
- [32] M. Jaderberg, V. Mnih, W. M. Czarnecki, T. Schaul, J. Z. Leibo, D. Silver, and K. Kavukcuoglu, "Reinforcement learning with unsupervised auxiliary tasks," 2016. [Online]. Available: <https://arxiv.org/abs/1611.05397>
- [33] F. Liu and G. Zeng, "Multi-agent cooperative learning research based on reinforcement learning," in *2006 10th International Conference on Computer Supported Cooperative Work in Design*, 2006, pp. 1–6.
- [34] —, "Multi-agent cooperative learning research based on reinforcement learning," in *2006 10th International Conference on Computer Supported Cooperative Work in Design*, 2006, pp. 1–6.
- [35] J. Zhang, Z. Huang, Y. Fan, N. Liu, M. Li, Z. Yang, J. Yao, J. Wang, and K. Wang, "KABB: Knowledge-aware bayesian bandits for dynamic expert coordination in multi-agent systems," in *Forty-second International Conference on Machine Learning*, 2025. [Online]. Available: <https://openreview.net/forum?id=AKvy9a4jho>
- [36] J. Zhang, Y. Fan, W. Lin, R. Chen, H. Jiang, W. Chai, J. Wang, and K. Wang, "Gam-agent: Game-theoretic and uncertainty-aware collaboration for complex visual reasoning," 2025. [Online]. Available: <https://arxiv.org/abs/2505.23399>
- [37] J. Zhang, Y. Fan, K. Cai, Z. Huang, X. Sun, J. Wang, C. Tang, and K. Wang, "Drdiff: Dynamic routing diffusion with hierarchical attention for breaking the efficiency-quality trade-off," 2025. [Online]. Available: <https://arxiv.org/abs/2509.02785>
- [38] J. Hu, L. Shen, and G. Sun, "Squeeze-and-excitation networks," in *2018 IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 2018, pp. 7132–7141.
- [39] A. Dosovitskiy, L. Beyer, A. Kolesnikov, D. Weissenborn, X. Zhai, T. Unterthiner, M. Dehghani, M. Minderer, G. Heigold, S. Gelly, J. Uszkoreit, and N. Houlsby, "An image is worth 16x16 words: Transformers for image recognition at scale," 2021. [Online]. Available: <https://arxiv.org/abs/2010.11929>
- [40] Z. Liu, Y. Lin, Y. Cao, H. Hu, Y. Wei, Z. Zhang, S. Lin, and B. Guo, "Swin transformer: Hierarchical vision transformer using shifted windows," in *2021 IEEE/CVF International Conference on Computer Vision (ICCV)*, 2021, pp. 9992–10002.
- [41] J. Lanchantin, T. Wang, V. Ordonez, and Y. Qi, "General multi-label image classification with transformers," in *2021 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, 2021, pp. 16 473–16 483.
- [42] T.-Y. Lin, P. Goyal, R. Girshick, K. He, and P. Dollár, "Focal loss for dense object detection," *IEEE Transactions on Pattern Analysis and Machine Intelligence*, vol. 42, no. 2, pp. 318–327, 2020.

- [43] E. Ben-Baruch, T. Ridnik, N. Zamir, A. Noy, I. Friedman, M. Protter, and L. Zelnik-Manor, “Asymmetric loss for multi-label classification,” 2021. [Online]. Available: <https://arxiv.org/abs/2009.14119>
- [44] T. Wu, Q. Huang, Z. Liu, Y. Wang, and D. Lin, “Distribution-balanced loss for multi-label classification in long-tailed datasets,” 2021. [Online]. Available: <https://arxiv.org/abs/2007.09654>
- [45] E. D. Cubuk, B. Zoph, D. Mané, V. Vasudevan, and Q. V. Le, “Autoaugment: Learning augmentation strategies from data,” in *2019 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, 2019, pp. 113–123.
- [46] S. Lim, I. Kim, T. Kim, C. Kim, and S. Kim, “Fast autoaugment,” 2019. [Online]. Available: <https://arxiv.org/abs/1905.00397>
- [47] S. Yun, D. Han, S. Chun, S. J. Oh, Y. Yoo, and J. Choe, “Cutmix: Regularization strategy to train strong classifiers with localizable features,” in *2019 IEEE/CVF International Conference on Computer Vision (ICCV)*, 2019, pp. 6022–6031.
- [48] D. P. Kingma and J. Ba, “Adam: A method for stochastic optimization,” 2017. [Online]. Available: <https://arxiv.org/abs/1412.6980>
- [49] I. Loshchilov and F. Hutter, “Decoupled weight decay regularization,” 2019. [Online]. Available: <https://arxiv.org/abs/1711.05101>
- [50] L. Liu, H. Jiang, P. He, W. Chen, X. Liu, J. Gao, and J. Han, “On the variance of the adaptive learning rate and beyond,” 2021. [Online]. Available: <https://arxiv.org/abs/1908.03265>
- [51] L. N. Smith, “Cyclical learning rates for training neural networks,” in *2017 IEEE Winter Conference on Applications of Computer Vision (WACV)*, 2017, pp. 464–472.
- [52] T.-Y. Lin, P. Goyal, R. Girshick, K. He, and P. Dollár, “Focal loss for dense object detection,” 2018. [Online]. Available: <https://arxiv.org/abs/1708.02002>
- [53] —, “Focal loss for dense object detection,” *IEEE Transactions on Pattern Analysis and Machine Intelligence*, vol. 42, no. 2, pp. 318–327, 2020.
- [54] X. He, K. Zhao, and X. Chu, “Automl: A survey of the state-of-the-art,” *Knowledge-Based Systems*, vol. 212, p. 106622, Jan. 2021. [Online]. Available: <http://dx.doi.org/10.1016/j.knosys.2020.106622>
- [55] H. Pham, M. Y. Guan, B. Zoph, Q. V. Le, and J. Dean, “Efficient neural architecture search via parameter sharing,” 2018. [Online]. Available: <https://arxiv.org/abs/1802.03268>
- [56] K. Zhang, Z. Yang, H. Liu, T. Zhang, and T. Başar, “Fully decentralized multi-agent reinforcement learning with networked agents,” 2018. [Online]. Available: <https://arxiv.org/abs/1802.08757>
- [57] T. Rashid, M. Samvelyan, C. S. de Witt, G. Farquhar, J. Foerster, and S. Whiteson, “Qmix: Monotonic value function factorisation for deep multi-agent reinforcement learning,” 2018. [Online]. Available: <https://arxiv.org/abs/1803.11485>
- [58] J. Zhang, K. Cai, Y. Fan, J. Wang, and K. Wang, “Cf-vlm: counterfactual vision-language fine-tuning,” 2025. [Online]. Available: <https://arxiv.org/abs/2506.17267>
- [59] P. Sunehag, G. Lever, A. Gruslys, W. M. Czarnecki, V. Zambaldi, M. Jaderberg, M. Lanctot, N. Sonnerat, J. Z. Leibo, K. Tuyls, and T. Graepel, “Value-decomposition networks for cooperative multi-agent learning,” 2017. [Online]. Available: <https://arxiv.org/abs/1706.05296>
- [60] L.-N. Zhang, X. Zuo, J.-W. Liu, W.-M. Li, and N. Ito, “Comments on “finite-time analysis of the multiarmed bandit problem”,” in *2019 International Conference on Machine Learning and Cybernetics (ICMLC)*, 2019, pp. 1–6.
- [61] D. Russo, B. V. Roy, A. Kazerouni, I. Osband, and Z. Wen, “A tutorial on thompson sampling,” 2020. [Online]. Available: <https://arxiv.org/abs/1707.02038>
- [62] A. Slivkins, “Introduction to multi-armed bandits,” 2024. [Online]. Available: <https://arxiv.org/abs/1904.07272>
- [63] M. Everingham, L. Gool, C. K. Williams, J. Winn, and A. Zisserman, “The pascal visual object classes (voc) challenge,” *Int. J. Comput. Vision*, vol. 88, no. 2, p. 303–338, Jun. 2010. [Online]. Available: <https://doi.org/10.1007/s11263-009-0275-4>
- [64] T.-Y. Lin, M. Maire, S. Belongie, L. Bourdev, R. Girshick, J. Hays, P. Perona, D. Ramanan, C. L. Zitnick, and P. Dollár, “Microsoft coco: Common objects in context,” 2015. [Online]. Available: <https://arxiv.org/abs/1405.0312>
- [65] R. Krishna, Y. Zhu, O. Groth, J. Johnson, K. Hata, J. Kravitz, S. Chen, Y. Kalantidis, L.-J. Li, D. A. Shamma, M. S. Bernstein, and F.-F. Li, “Visual genome: Connecting language and vision using crowdsourced dense image annotations,” 2016. [Online]. Available: <https://arxiv.org/abs/1602.07332>
- [66] Z.-M. Chen, X.-S. Wei, P. Wang, and Y. Guo, “Multi-label image recognition with graph convolutional networks,” 2019. [Online]. Available: <https://arxiv.org/abs/1904.03582>

- [67] J. Lanchantin, T. Wang, V. Ordonez, and Y. Qi, “General multi-label image classification with transformers,” 2020. [Online]. Available: <https://arxiv.org/abs/2011.14027>
- [68] E. Ben-Baruch, T. Ridnik, N. Zamir, A. Noy, I. Friedman, M. Protter, and L. Zelnik-Manor, “Asymmetric loss for multi-label classification,” 2021. [Online]. Available: <https://arxiv.org/abs/2009.14119>
- [69] T. Ridnik, G. Sharir, A. Ben-Cohen, E. Ben-Baruch, and A. Noy, “MI-decoder: Scalable and versatile classification head,” 2021. [Online]. Available: <https://arxiv.org/abs/2111.12933>
- [70] Z. Yin, C. Gan, K. He, Y. Gao, and J. Zhang, “Hybrid sharing for multi-label image classification,” in *The Twelfth International Conference on Learning Representations*, 2024. [Online]. Available: <https://openreview.net/forum?id=yVJd8IKyVX>
- [71] M.-K. Xie, J.-H. Xiao, P. Peng, G. Niu, M. Sugiyama, and S.-J. Huang, “Counterfactual reasoning for multi-label image classification via patching-based training,” 2024. [Online]. Available: <https://arxiv.org/abs/2404.06287>
- [72] C. D. Manning, P. Raghavan, and H. Schütze, *Introduction to Information Retrieval*. Cambridge, UK: Cambridge University Press, 2008.
- [73] M.-L. Zhang and Z.-H. Zhou, “A review on multi-label learning algorithms,” *IEEE Transactions on Knowledge and Data Engineering*, vol. 26, no. 8, pp. 1819–1837, 2014.
- [74] S. Falkner, A. Klein, and F. Hutter, “Bohb: Robust and efficient hyperparameter optimization at scale,” 2018. [Online]. Available: <https://arxiv.org/abs/1807.01774>
- [75] T.-S. Chua, J. Tang, R. Hong, H. Li, Z. Luo, and Y. Zheng, “Nus-wide: a real-world web image database from national university of singapore,” in *Proceedings of the ACM international conference on image and video retrieval*, 2009, pp. 1–9.
- [76] H. Liu, K. Simonyan, and Y. Yang, “Darts: Differentiable architecture search,” 2019. [Online]. Available: <https://arxiv.org/abs/1806.09055>

Core Mechanisms, Efficiency, and Parameter Configuration of MAT-Agent

This supplementary material aims to provide a more detailed and in-depth explanation of the MAT-Agent framework, addressing key questions and concerns raised during the review process regarding its agent coordination mechanisms, computational efficiency considerations, state and action space design, and hyperparameter configurations. We hope these clarifications will resolve related doubts and comprehensively demonstrate the completeness and rigor of this research work.

S 1.1 Detailed Explanation of Agent Coordination Mechanism

Multiple autonomous agents within the MAT-Agent framework achieve efficient collaboration to optimize the training process of multi-label image classification (MLIC) through the following interconnected core mechanisms:

1. **Shared Global State Representation ($s_t, \mathcal{I}_t$):** All agents, at each decision step t , access and utilize a comprehensive, real-time updated extended state representation $\mathcal{I}_t$. The state s_t is carefully designed to integrate the MLIC model’s immediate learning performance metrics (s_t^{perf}), complex dynamic characteristics of the training process (s_t^{dyn}), and descriptors reflecting current data characteristics (s_t^{data}). This global information sharing ensures that all agents make decisions based on a unified understanding of the current training environment, forming the foundation for their effective coordinated actions.
2. **Unified Global Composite Reward Signal (R_{t+1}):** After all agents collectively determine and execute a joint training configuration $C_t = (a_t^{AUG}, a_t^{OPT}, a_t^{LRS}, a_t^{LOSS})$, the system computes a scalar reward signal R_{t+1} based on the overall performance of this configuration (rather than the isolated effect of individual agent actions), which is then shared among all agents. The reward function

$$R_{t+1} = w_{mAP} f(\Delta mAP_t) + w_{stab} \text{Stability}_t + w_{conv} \text{Convergence}_t - w_{pen} \text{Penalty}_t$$

integrates model accuracy improvement (via $f(\Delta mAP_t)$, where $f(\cdot)$ amplifies significant improvements), training stability (negatively correlated with loss fluctuations), convergence speed (e.g., loss reduction rate), and imposes penalties for computationally expensive or unstable configurations. Although each agent $Agent_k$ independently learns its state-action value function Q_k , their common optimization goal is to maximize the expected value of this shared global cumulative reward

$$J = \mathbb{E} \left[\sum_{t=0}^T \gamma^t R(s_t, C_t) \right].$$

This mechanism ensures that all agents’ learning behaviors are aligned toward enhancing the global optimal efficiency of the entire training process.

3. **Experience Sharing and Learning Stability:** The agents’ interaction experiences $(\mathcal{I}_j, a_j^k, R_{j+1}, \mathcal{I}_{j+1})$ are uniformly stored in an experience replay buffer $\mathcal{D}$ and sampled to update each agent’s Q-network. Combined with the target Q-network $Q_k(\cdot; \theta_k^-)$ to compute the TD target

$$y_j = R_{j+1} + \gamma \max_{a' \in \mathcal{A}_k} Q_k(\mathcal{I}_{j+1}, a'; \theta_k^-),$$

Standard reinforcement learning techniques further enhance the stability of the learning process and encourage agents to learn from successful historical experiences that lead to high global rewards.

In our ablation studies, the “w/o Agent Coordination” experimental configuration was implemented by weakening or removing the aforementioned mechanisms based on shared information (e.g., completeness of global state or shared experience replay buffer) and unified optimization objectives (e.g., agents might learn based on modified, non-fully consistent global reward signals). Experimental results show that such reductions in coordination lead to significant performance degradation, e.g., mAP dropping from 97.4% (full model) to as low as 91.7% (w/o All), depending on the combination of removed components. This inversely validates the effectiveness and necessity of the current coordination mechanisms.

Regarding the “structured communication channels and attention mechanisms” mentioned in the main paper’s conclusions, this concept represents our vision for the future evolution of MAT-Agent, aiming to explore more direct and complex inter-agent information interactions and collaborative strategies. In the current submitted work, no such explicit, structured inter-agent direct communication protocols are included.

S 1.2 Clarification on Computational Overhead and Overall Efficiency

To address concerns regarding the computational overhead of MAT-Agent and the interpretation of “no additional search overhead,” we provide the following clarifications, emphasizing the trade-off between complexity and performance gains, and comparing MAT-Agent against simpler adaptive methods.

1. **Per-Epoch Computational Overhead:** MAT-Agent’s four deep Q-network agents perform state perception, decision-making, and Q-network updates, incurring a 10% per-epoch overhead (16.5 minutes vs. 15 minutes for the standard baseline on NVIDIA A100 GPU, as shown in Supplementary Material G, Figure 8. This complexity arises from the multi-agent design but is mitigated by shared global state and reward signals, minimizing redundant computations.
2. **Context of “No Additional Search Overhead”:** The claim “no additional search overhead” (Section 2, related work) contrasts MAT-Agent with traditional methods requiring offline hyperparameter searches (e.g., grid or random search, which are computationally costly and prone to local optima) or complex AutoML setups. MAT-Agent integrates dynamic adaptation into training, eliminating separate search phases and their associated costs.
3. **Efficiency via Reduced Training Epochs:** MAT-Agent’s key advantage is reducing total training epochs for multi-label image classification (MLIC). On MS-COCO, it achieves 63.8% mAP in 47 epochs, compared to 80 epochs for the standard method, a 41.25% time reduction (Figure 2(b)). Supplementary Material G, Figure 8(b), shows policy convergence in 47 epochs, outperforming the standard method (74 epochs), AutoAugment (78 epochs), PBT (73 epochs), and BOHB (76 epochs).
4. **Comparison with Simpler Adaptive Methods:** Compared to tuned adaptive methods like advanced PBT variants or sophisticated schedulers (e.g., Cyclic LR, dynamic Focal Loss), MAT-Agent demonstrates superior performance. Figure 4 shows MAT-Agent’s mAP of 62.5 at 10^2 GPU hours, surpassing PBT and others. While simpler methods have lower per-epoch complexity, they lack MAT-Agent’s multi-dimensional dynamic optimization, limiting their ability to match its convergence speed and final mAP (e.g., 12.5 hours to target performance vs. 20.0 hours for PBT, Figure 8(a)).
5. **Net Efficiency and Complexity Trade-off:** The 10% per-epoch overhead is offset by a 41.25% reduction in total training time (12.5 hours vs. 18.5 hours for the standard method, 20.0 hours for PBT, 21.5 hours for BOHB, and 22.0 hours for AutoAugment, Figure 8(a)). The multi-agent design, though complex, enables robust adaptation across data augmentation, optimization, learning rate, and loss functions, yielding consistent mAP gains. While simpler methods reduce complexity, MAT-Agent’s comprehensive optimization justifies its overhead, as evidenced by its performance edge on benchmark datasets.

S1.3 Detailed Design and Considerations for State and Action Spaces

1. State Representation ($s_t, \mathcal{I}_t$):

- As described in Section 3.2 of the main paper, the instantaneous state s_t of the system at decision point t is a multidimensional real-valued vector $s_t \in \mathcal{S} \subset \mathbb{R}^D$, comprising: s_t^{perf} (immediate model performance metrics, e.g., validation set mAP val), s_t^{dyn} (training dynamic characteristics, e.g., training/validation loss values $L_t^{train/val}$, loss change ΔL_t^{val} , gradient statistics of main model parameters with respect to training loss $g_t = \nabla_{\Theta_M} L_t^{train}$ (e.g., L_2 norm $\|g_t\|_2$), parameter update magnitude, etc.), and s_t^{data} (specific descriptors of the dataset or current data batch, e.g., “average texture richness of current samples”).

- These state components are selected to provide agents with a comprehensive, real-time multidimensional portrait of the model’s learning state, training environment dynamics, and current data characteristics, serving as the informational basis for effective adaptive decision-making. For example, information in s_t^{data} helps agents adjust strategies based on specific data characteristics (e.g., complexity, class distribution cues), such as adopting more aggressive augmentations when data features are relatively simple.
- To support temporal decision-making based on historical information, the agents’ actual input is the extended state representation $\mathcal{I}_t$, which aggregates current and historical state observations. For detailed construction methods of $\mathcal{I}_t$ (e.g., how historical information is aggregated) and the complete feature set of s_t , refer to Section A.2 of this supplementary material.

2. Action Space ($\mathcal{A}_k$):

- The MAT-Agent framework includes four autonomous agents, each responsible for dynamically regulating one of four key training components: data augmentation (AUG), optimizer selection (OPT), learning rate scheduling (LRS), and loss function design (LOSS).
- At each decision step t , each agent $Agent_k$ selects an action a_t^k from its dedicated, predefined discrete action space $\mathcal{A}_k = \{a_k^{(1)}, \dots, a_k^{(M_k)}\}$.
- These candidate actions (i.e., training strategy options) are carefully selected and designed based on relevant literature and widely applied effective methods in practice. For example (see Section A.1 of the supplementary material for the complete list):
 - The action space $\mathcal{A}_{AUG}$ of $Agent_{AUG}$ may include: no additional augmentation (None/Basic Augmentation), basic augmentation combinations such as random cropping and horizontal flipping, or more complex strategies such as CutMix, MixUp, RandAugment, or specific strategies from Fast AutoAugment.
 - The action space $\mathcal{A}_{OPT}$ of $Agent_{OPT}$ may include: SGD, Adam, AdamW (improved Adam with weight decay), RAdam (addressing Adam’s warmup issues), etc.
 - The action space $\mathcal{A}_{LRS}$ of $Agent_{LRS}$ may include: Step Decay (fixed-interval decay), MultiStepLR (decay at predetermined epochs), Cosine Annealing (cosine annealing scheduler), One-Cycle policy (cyclic learning rate with increase then decrease), Linear Decay (linear decay), etc.
 - The action space $\mathcal{A}_{LOSS}$ of $Agent_{LOSS}$ may include: standard Binary Cross-Entropy Loss (BCE Loss), Focal Loss designed for class imbalance, Asymmetric Loss (ASL), Mean Squared Error Loss (MSE Loss, applicable in some multi-label scenarios), or Class-balanced Loss (CB Loss, as mentioned in Figure 3 and Table 4 of the supplementary material).
- The **complete list and specific definitions** of each agent’s candidate strategies are provided in Section A.1 of this supplementary material. This design enables agents to effectively explore and exploit within a structured and meaningful strategy space.

S1.4 Configuration and Impact Analysis of MAT-Agent’s Hyperparameters

As a framework based on deep reinforcement learning, MAT-Agent includes a series of hyperparameters that need to be preset. In Section H of the supplementary material, we have analyzed the impact of some key hyperparameters, particularly the weight factors of the composite reward function R_{t+1} .

• Main Endogenous Hyperparameters of MAT-Agent:

- **DQN Agent Parameters:** Learning rate of each agent’s Q-network, discount factor γ (set between [0,1] in the main paper), specific neural network structure of the Q-network (see Section A.3 of the supplementary material), capacity and mini-batch sampling size of the experience replay buffer $\mathcal{D}$, and update frequency of the target Q-network Q_k^- .
- **Exploration Strategy Parameters:** Exploration rate ϵ_t in the ϵ -greedy strategy, which decays over time to balance early exploration and later exploitation, and parameters related to potential intrinsic motivation mechanisms (e.g., curiosity-driven, based on prediction errors of state transitions).

- **Weights of the Composite Reward Function R_{t+1} :** Weight factors $w_{mAP}, w_{stab}, w_{conv}, w_{pen}$ in Equation (3) of the main paper:

$$R_{t+1} = w_{mAP}f(\Delta mAP_t) + w_{stab}Stability_t + w_{conv}Convergence_t - w_{pen}Penalty_t.$$

These weights determine the agents’ emphasis on different optimization objectives. Specific mathematical definitions and weight values are detailed in Section A.4 of the supplementary material.

- **Analysis in Section H of the Supplementary Material (Particularly Figure 9):** The analysis reveals the specific impacts of reward weights w_{mAP} (range 0.4 to 1.6) and w_{stab} (range 1.0 to 1.2, adjusted based on text description as Figure 9’s x-axis shows 0.2 to 1.2) on the model’s final mAP performance and training stability:
 - The impact of w_{mAP} on mAP is nonlinear: as w_{mAP} increases from 0.4 to 0.8, mAP rises from 88.2% to a peak of 92.8%; however, when w_{mAP} exceeds 1.0, mAP begins to decline, reaching 86.0% at $w_{mAP} = 1.6$. This suggests that overemphasizing short-term accuracy may lead to overfitting, impairing generalization.
 - w_{stab} positively affects training stability: as w_{stab} (based on Figure 9’s x-axis, 0.2 to 1.2, interpreted as 0.6 to 1.2 here) increases from 0.6 to 1.2, the variance of loss fluctuations decreases from approximately 0.05 (at $w_{stab} = 0.6$) to a lower value, or at $w_{stab} = 1.2$, the loss variance is 0.09 compared to 0.05 at $w_{stab} = 0.6$, indicating that higher stability weights help suppress training oscillations. *(Note: There is a slight inconsistency between the figure and text description; we follow the figure’s trend, where higher w_{stab} generally corresponds to lower variance, but the variance at 1.2 is higher than at 0.6, which may require author verification. We assume the trend is that higher w_{stab} improves stability, i.e., lower variance.)*
 - There is a trade-off between accuracy and stability: for example, at $w_{mAP} = 0.8$, mAP is highest (92.8%), but stability (loss variance of 0.07) is not optimal; a configuration such as $w_{mAP} = 1.0$ and $w_{stab} = 0.8$ (corresponding to a variance of 0.06) may achieve a good balance between mAP (90.8%) and stability.

Considerations and Analysis of MAT-Agent’s Hyperparameter Configuration

The effective operation of the MAT-Agent framework involves the configuration of a series of endogenous hyperparameters. Understanding and appropriately configuring these hyperparameters are crucial for achieving the framework’s optimal performance. This section aims to elucidate the composition, tuning considerations, and impacts of these hyperparameters on the framework’s performance.

1. **Overview of MAT-Agent’s Endogenous Hyperparameters:** The hyperparameters of MAT-Agent primarily stem from its multi-agent architecture based on deep reinforcement learning (DRL). These parameters can be categorized into the following core groups:
 - **Deep Q-Network (DQN) Agent Parameters:** These parameters relate to the learning core of each independent agent, including the learning rate of the Q-network, the discount factor $\gamma \in [0, 1]$ used for computing future rewards, the specific architecture of the Q-network (e.g., number of layers, activation functions, detailed in Supplementary Material A.3), the capacity of the experience replay buffer $\mathcal{D}$, the mini-batch sampling size, and the update mechanism for the target Q-network Q_k^- .
 - **Exploration and Exploitation Strategy Parameters:** MAT-Agent employs an ϵ -greedy strategy to balance exploration (trying new, untested actions) and exploitation (selecting the currently known optimal action) during training. The key parameter is the exploration rate ϵ_t and its dynamic decay scheme over training time. Additionally, the curiosity-driven intrinsic reward mechanism mentioned in the framework may also involve its own configuration parameters.
 - **Weight Factors of the Composite Reward Function R_{t+1} :** The global reward defined in Equation (3) of the main paper, $R_{t+1} = w_{mAP}f(\Delta mAP_t) + w_{stab}Stability_t + w_{conv}Convergence_t - w_{pen}Penalty_t$, includes multiple weight factors ($w_{mAP}, w_{stab}, w_{conv}, w_{pen}$). These weights balance different optimization objectives (e.g., accuracy, stability, convergence speed), with their specific mathematical definitions and reference values detailed in Supplementary Material A.4.

2. **Hyperparameter Configuration and MAT-Agent’s Design Philosophy:** The design of MAT-Agent aims to transform the optimization problem of the highly complex and dynamic configuration space $C_t = (a_t^{AUG}, a_t^{OPT}, a_t^{LRS}, a_t^{LOSS})$ for MLIC models (which has a vast number of possible combinations, forming a large space $|\mathcal{C}| = \prod_{k \in K} |\mathcal{A}_k|$) into the tuning of its own learning framework’s hyperparameters. While the latter still requires careful consideration, its dimensionality is relatively lower and often carries clearer physical or goal-oriented significance. The core distinction lies in the fact that traditional methods seek a static configuration C^* that is optimal throughout the entire training process, whereas MAT-Agent aims to learn a set of adaptive decision-making “policies” $\{\pi_k^*\}$. These policies enable agents to dynamically generate appropriate training configurations $C_t(s_t)$ based on real-time training states s_t . Such learned “meta-policies” are expected to exhibit certain generalization and transferability across similar tasks or datasets, as preliminarily demonstrated in the cross-dataset generalization experiments in Section 4.3 of the main paper and the small-sample transfer experiments in Supplementary Material E.

3. Practical Considerations for MAT-Agent’s Hyperparameter Configuration:

- **Leveraging Standards and Heuristic Configurations:** For many components of the DRL framework, such as certain DQN agent parameters (e.g., discount factor γ , experience replay mechanisms) and exploration strategies (e.g., initial value and decay method of ϵ_t), mature practices and standard recommended values from the reinforcement learning field can be referenced, or heuristic methods that adjust with training progress can be adopted.
- **Sensitivity and Tuning of Key Hyperparameters:** The weight factors in the reward function directly influence MAT-Agent’s learning orientation and final performance. Section H of the Supplementary Material (particularly Figure 9) provides a systematic sensitivity analysis. This analysis examines the effects of varying the accuracy weight w_{mAP} and stability weight w_{stab} on model performance (mAP) and training stability (loss fluctuations):
 - Adjustments to w_{mAP} show that as it increases from 0.4 to 0.8, mAP improves from 88.2% to a peak of 92.8%; however, if increased further to 1.6, mAP drops to 86.0%. This reveals the effective range of the parameter and the trade-offs of over-optimizing a single objective.
 - Increasing w_{stab} (e.g., from 1.0 to 1.2) helps reduce loss fluctuations during training, lowering the standard deviation from 0.10 to 0.05, thereby enhancing training stability and convergence reliability.
 - These analysis results (e.g., by setting $w_{mAP} = 1.0$ and $w_{stab} = 1.1$, achieving an mAP of 90.8% with a loss standard deviation of 0.06) provide experimental evidence for balancing different optimization objectives and indicate that key hyperparameters have a robust and effective configuration range.

Through its multi-agent collaboration mechanism, the MAT-Agent framework dynamically optimizes the MLIC training process, achieving state-of-the-art performance and rapid convergence across multiple benchmark datasets. While its endogenous hyperparameters require thoughtful configuration, these can be effectively managed by adopting established reinforcement learning practices, conducting sensitivity analyses, and leveraging the framework’s adaptive “meta-policies,” which exhibit promising generalization across tasks and datasets. Comprehensive analyses demonstrate that, with judicious tuning, MAT-Agent ensures stable and efficient operation, even for standard DQN-related parameters, which align with literature conventions or preliminary experiments. Future research will focus on automating and enhancing the usability of hyperparameter configuration to further streamline the framework’s deployment.

A MAT-Agent: Single-Agent Q-Learning Mechanism and Convergence Behavior

A.1 Foundational Settings for Analysis ($\mathcal{FS}$)

Let each autonomous agent be indexed by $k \in \{1, \dots, N\}$.

- **State and Action ($\mathcal{S}$ and $\mathcal{A}_k$):** The extended global state at time t is denoted as $\mathcal{I}_t \in \mathcal{S}$. The action taken by agent k is represented as $a_t^k \in \mathcal{A}_k$. The joint action across all agents is defined as:

$$C_t = (a_t^1, \dots, a_t^N).$$

- **Environment Dynamics:** The environment transitions according to a probability distribution:

$$p(\mathcal{I}_{t+1} \mid \mathcal{I}_t, C_t),$$

which models the likelihood of the next state $\mathcal{I}_{t+1}$ given the current state and joint action.

- **Reward Function:** The system receives a global reward after each transition, specified as:

$$R(\mathcal{I}_t, C_t, \mathcal{I}_{t+1}) \rightarrow \mathbb{R}, \quad \text{denoted as } R_{t+1}.$$

- **Q-Function:** Each agent maintains an action-value function:

$$Q_k : \mathcal{S} \times \mathcal{A}_k \rightarrow \mathbb{R}, \quad Q_k(\mathcal{I}, a^k; \theta_k).$$

- **Target Network:** A separate target network is maintained for each agent:

$$Q_k(\mathcal{I}, a^k; \theta_k^-),$$

where θ_k^- is a periodically updated copy of θ_k .

- **Policy:** Each agent follows a policy:

$$\pi_k(\cdot \mid \mathcal{I}; \theta_k),$$

typically instantiated as an ϵ -greedy policy over the Q-function.

- **Experience Replay:** A global experience replay buffer $\mathcal{D}$ stores transitions:

$$(\mathcal{I}_j, C_j, R_{j+1}, \mathcal{I}_{j+1}).$$

For agent k , its specific action in C_j is a_j^k .

A.2 Q-Learning Update Mechanism for Individual Agent k

To update the Q-function of agent k , a transition sample

$$j = (\mathcal{I}_j, a_j^k, C_j^{\setminus k}, R_{j+1}, \mathcal{I}_{j+1})$$

is drawn from the shared experience replay buffer $\mathcal{D}$. Here, a_j^k denotes the action taken by agent k in state $\mathcal{I}_j$, while $C_j^{\setminus k}$ denotes the action set of all other agents $m \neq k$, that is, $\{a_j^m\}_{m \neq k}$. The full joint action is thus:

$$C_j = (a_j^k, C_j^{\setminus k}).$$

A.2.1 Temporal-Difference (TD) Target y_j^k

The temporal-difference (TD) target y_j^k provides a learning signal for evaluating the expected return of the state-action pair $(\mathcal{I}_j, a_j^k)$, and is defined as:

$$y_j^k = \underbrace{R_{j+1}}_{\text{Immediate global reward observed after executing joint action } C_j} + \underbrace{\gamma}_{\text{Discount factor for future returns, } 0 \leq \gamma < 1} \left(\underbrace{\max_{a' \in \mathcal{A}_k} Q_k(\mathcal{I}_{j+1}, a'; \theta_k^-)}_{\substack{\text{Estimated value of agent } k\text{'s best next action } a' \\ \text{at successor state } \mathcal{I}_{j+1}, \text{ computed using} \\ \text{a fixed target network } Q_k(\cdot; \theta_k^-)}} \right) \quad (5)$$

This TD target y_j^k integrates both the *observed immediate global reward* R_{j+1} and the *estimated optimal future return*. The latter is computed using the agent's target Q-network $Q_k(\cdot; \theta_k^-)$, whose parameters remain fixed over a window of training iterations to stabilize learning and mitigate feedback loops during value propagation.

A.2.2 Current Q-Value Estimation ($Q_{\text{current},j}^k$)

The current Q-network of agent k , parameterized by θ_k , provides an immediate estimate of the expected cumulative return for the historical state-action pair $(\mathcal{I}_j, a_j^k)$:

$$Q_{\text{current},j}^k = \underbrace{Q_k(\mathcal{I}_j, a_j^k; \theta_k)}_{\substack{\text{Agent } k\text{'s current Q-network (with learnable parameters } \theta_k) \\ \text{predicting the expected cumulative discounted reward (Q-value)} \\ \text{for executing historical action } a_j^k \text{ at state } \mathcal{I}_j}} \quad (6)$$

Here, $Q_k(\mathcal{I}_j, a_j^k; \theta_k)$ represents the agent's current estimation of the action-value for performing a_j^k in state $\mathcal{I}_j$, based on its learned parameter vector θ_k , which is continuously optimized during training.

A.2.3 Loss Function ($L_j(\theta_k)$)

The loss function $L_j(\theta_k)$ quantifies the discrepancy between the predicted Q-value $Q_k(\mathcal{I}_j, a_j^k; \theta_k)$ produced by the current Q-network and the learning target y_j^k . It is typically expressed in the form of mean squared error (MSE):

$$L_j(\theta_k) = \underbrace{\frac{1}{2}}_{\text{Convenient scaling factor}} \left(\underbrace{y_j^k}_{\substack{\text{Learning target} \\ \text{(defined in Eq. (2.1))}}} - \underbrace{Q_k(\mathcal{I}_j, a_j^k; \theta_k)}_{\substack{\text{Current Q-network estimate} \\ \text{(defined in Eq. (2.2))}}} \right)^2 \quad (7)$$

To simplify notation, we define the temporal-difference (TD) error as:

$$\delta_j^k = \underbrace{y_j^k}_{\text{Target Q-value}} - \underbrace{Q_k(\mathcal{I}_j, a_j^k; \theta_k)}_{\text{Current Q estimate}} \quad (8)$$

Substituting this into the loss, we obtain a more concise formulation:

$$L_j(\theta_k) = \frac{1}{2} \underbrace{(\delta_j^k)^2}_{\substack{\text{Squared TD error,} \\ \text{measuring prediction-target divergence}}} \quad (9)$$

The training objective for agent k is to minimize this loss by adjusting its Q-network parameters θ_k , typically through gradient-based optimization over mini-batches of sampled transitions.

A.2.4 Gradient of the Loss ($\nabla_{\theta_k} L_j(\theta_k)$)

Using the chain rule, the gradient is:

$$\begin{aligned} \nabla_{\theta_k} L_j(\theta_k) &= \nabla_{\theta_k} \left[\frac{1}{2} (\delta_j^k)^2 \right] \\ &= \underbrace{\delta_j^k}_{\text{Outer derivative}} \cdot \underbrace{\nabla_{\theta_k} \delta_j^k}_{\text{Inner derivative}} \\ &= \delta_j^k \cdot \nabla_{\theta_k} (y_j^k - Q_k(\mathcal{I}_j, a_j^k; \theta_k)) \\ &= \delta_j^k \cdot \left(\underbrace{\nabla_{\theta_k} y_j^k}_{=0 \text{ (target network fixed)}} - \underbrace{\nabla_{\theta_k} Q_k(\mathcal{I}_j, a_j^k; \theta_k)}_{\text{Q-network gradient}} \right) \\ &= -\delta_j^k \cdot \nabla_{\theta_k} Q_k(\mathcal{I}_j, a_j^k; \theta_k) \end{aligned}$$

A.2.5 Parameter Update Rule

The Q-network parameters of agent k are updated by minimizing the loss $L_j(\theta_k)$ via gradient descent:

$$\theta_k \leftarrow \theta_k - \alpha \nabla_{\theta_k} L_j(\theta_k)$$

Substituting the expression for the gradient from Section 1.2.4:

$$\theta_k \leftarrow \underbrace{\theta_k}_{\text{Old parameters}} + \underbrace{\alpha}_{\text{Learning rate}} \cdot \underbrace{(y_j^k - Q_k(\mathcal{I}_j, a_j^k; \theta_k))}_{\text{TD error } \delta_j^k} \cdot \underbrace{\nabla_{\theta_k} Q_k(\mathcal{I}_j, a_j^k; \theta_k)}_{\text{Gradient } g_j^k(\theta_k)}$$

This update step adjusts the parameter vector θ_k in the direction that minimizes the TD error, thereby refining the Q-value estimation at the sampled state-action pair.

A.3 Local Convergence Trend for Agent k (Under Fixed Policies of Other Agents)

To analyze the learning dynamics of agent k , we consider the setting in which **the policies of all other agents $m \neq k$ remain fixed** for a period of time. Let these fixed policies be denoted by the set $\{\pi_m^*\}_{m \neq k}$.

Under this condition, from agent k 's perspective, the environment dynamics (i.e., how states transition and rewards are received) become *temporarily stationary*. Consequently, agent k 's learning problem reduces to a standard single-agent reinforcement learning task under a fixed Markov Decision Process (MDP), which is defined as follows:

- **State space $\mathcal{S}$:** Identical to the original global state space. The agent observes extended global states $\mathcal{I} \in \mathcal{S}$.
- **Action space $\mathcal{A}_k$:** The set of actions available to agent k .
- **Effective state transition probability:**

$$P_k(\mathcal{I}' | \mathcal{I}, a^k; \{\pi_m^*\}) = \underbrace{\sum_{\{a^m \in \times_{m \neq k} \mathcal{A}_m\}} \left(\prod_{m \neq k} \pi_m^*(a^m | \mathcal{I}) \right) p(\mathcal{I}' | \mathcal{I}, (a^k, \{a^m\}))}_{\text{Marginalizing over all other agents' actions drawn from their fixed policies}} \quad (10)$$

- **Effective expected reward:**

$$R_k(\mathcal{I}, a^k; \{\pi_m^*\}) = \underbrace{\mathbb{E}_{\substack{\{a^m \sim \pi_m^*(\cdot | \mathcal{I})\} \\ \mathcal{I}' \sim p(\cdot | \mathcal{I}, (a^k, \{a^m\}))}} [R(\mathcal{I}, (a^k, \{a^m\}), \mathcal{I}')] }_{\text{Expectation over other agents' actions and next state based on the global reward function } R}} \quad (11)$$

In this induced MDP, agent k 's objective is to learn an optimal Q-function $Q_k^*(\mathcal{I}, a^k | \{\pi_m^*\})$, which satisfies the Bellman optimality equation:

$$Q_k^*(\mathcal{I}, a^k | \{\pi_m^*\}) = \underbrace{R_k(\mathcal{I}, a^k; \{\pi_m^*\})}_{\text{Expected immediate reward for } a^k \text{ at } \mathcal{I}} + \underbrace{\gamma}_{\text{Discount factor}} \underbrace{\sum_{\mathcal{I}' \in \mathcal{S}} P_k(\mathcal{I}' | \mathcal{I}, a^k; \{\pi_m^*\}) \max_{a' \in \mathcal{A}_k} Q_k^*(\mathcal{I}', a' | \{\pi_m^*\})}_{\substack{\text{Expected future value over successor states} \\ \text{based on optimal Q-values}}} \quad (12)$$

This can also be written in expectation form:

$$Q_k^*(\mathcal{I}, a^k | \{\pi_m^*\}) = \underbrace{\mathbb{E}_{\mathcal{I}' \sim P_k(\cdot | \mathcal{I}, a^k; \{\pi_m^*\})} \left[R_k(\mathcal{I}, a^k; \{\pi_m^*\}) + \gamma \max_{a' \in \mathcal{A}_k} Q_k^*(\mathcal{I}', a' | \{\pi_m^*\}) \right]}_{\text{Expected sum of immediate reward and discounted future value}} \quad (13)$$

Note: For notational simplicity, R_k may be placed inside the expectation if it is defined as an expectation itself. In cases where the global reward function R is deterministic with respect to $(\mathcal{I}, C, \mathcal{I}')$, the reward R_{t+1} can be treated as a sample.

A.3.1 Parameter Updates and Bellman Consistency Trend

The Q-network parameter update rule for agent k (as described in Section 2.5) is given by:

$$\theta_k \leftarrow \theta_k + \alpha \underbrace{(y_j^k - Q_k(\mathcal{I}_j, a_j^k; \theta_k))}_{\text{TD error } \delta_j^k} \nabla_{\theta_k} Q_k(\mathcal{I}_j, a_j^k; \theta_k) \quad (14)$$

Here, the TD target y_j^k is:

$$y_j^k = \underbrace{R_{j+1}}_{\text{Sampled reward}} + \underbrace{\gamma \max_{a' \in \mathcal{A}_k} Q_k(\mathcal{I}_{j+1}, a'; \theta_k^-)}_{\text{Estimated future value from target network}} \quad (15)$$

The reward R_{j+1} is sampled from $r_k(\mathcal{I}_j, a_j^k, \mathcal{I}_{j+1}; \{\pi_m^*\})$, or directly from the global function $R(\mathcal{I}_j, C_j, \mathcal{I}_{j+1})$, and $\mathcal{I}_{j+1} \sim P_k(\cdot \mid \mathcal{I}_j, a_j^k; \{\pi_m^*\})$. The target Q-network $Q_k(\cdot, \cdot; \theta_k^-)$ provides a stable approximation to the unknown optimum $Q_k^*(\cdot, \cdot \mid \{\pi_m^*\})$.

Thus, y_j^k can be interpreted as a single-sample Monte Carlo estimate of the Bellman optimality target. The parameter update seeks to minimize the expected squared TD error:

$$L(\theta_k) = \mathbb{E}_{j \sim \mathcal{D}} \left[\underbrace{\left(R_{j+1} + \gamma \max_a Q_k(\mathcal{I}_{j+1}, a'; \theta_k^-) - Q_k(\mathcal{I}_j, a_j^k; \theta_k) \right)^2}_{\text{Squared Bellman residual (sampled using target network)}} \right] \quad (16)$$

Using stochastic gradient descent, the parameters θ_k are optimized such that the Q-network approximates the Bellman target computed from samples and target network outputs.

Under standard stochastic approximation conditions—e.g., learning rate α satisfies Robbins–Monro conditions, the function class for Q_k is expressive enough, and exploration sufficiently covers the state–action space—the learning dynamics of θ_k are expected to yield:

- Approximate convergence of $Q_k(\cdot, \cdot; \theta_k)$ toward the optimal Q-function $Q_k^*(\cdot, \cdot \mid \{\pi_m^*\})$, assuming fixed policies for other agents.

Function Approximation Cases:

- *Linear*: If $Q_k(\mathcal{I}, a^k; \theta_k) = \phi(\mathcal{I}, a^k)^\top \theta_k$, then convergence to a projection of Q_k^* in the feature space is possible, under suitable assumptions (e.g., linearly independent features, diminishing step sizes).
- *Nonlinear (e.g., DQN)*: Convergence is not guaranteed, but empirical techniques such as target networks and experience replay help stabilize learning, aiming for useful approximations of Q_k^* .

A.4 Inter-Agent Strategy Co-Evolution and Symbolic Considerations

In MAT-Agent, all agents’ strategies $\{\pi_m(\cdot \mid \mathcal{I}; \theta_m)\}_{m=1}^N$ evolve simultaneously.

A.4.1 Impact of Non-Stationarity on Agent k

The effective environment dynamics p_k^{eff} and reward function r_k^{eff} that agent k experiences are determined by the policy profiles of all other agents $\{\pi_m\}_{m \neq k}$:

$$P_k^{\text{eff}}(\mathcal{I}_{t+1} \mid \mathcal{I}_t, a_t^k; \{\theta_m(t)\}_{m \neq k}) = \underbrace{\sum_{\{a_t^m \in \times_{m \neq k} \mathcal{A}_m\}}}_{\text{Marginalizing over all possible actions of other agents}} \left(\underbrace{\prod_{m \neq k} \pi_m(a_t^m \mid \mathcal{I}_t; \theta_m(t))}_{\text{Joint action probability of other agents under their current policies}} \right) \times \underbrace{p(\mathcal{I}_{t+1} \mid \mathcal{I}_t, (a_t^k, \{a_t^m\}))}_{\text{True transition dynamics given the joint action}} \quad (17)$$

$$R_k^{\text{eff}}(\mathcal{I}_t, a_t^k; \{\theta_m(t)\}_{m \neq k}) = \underbrace{\mathbb{E}_{\substack{\{a_t^m \sim \pi_m(\cdot \mid \mathcal{I}_t; \theta_m(t))\}_{m \neq k} \\ \mathcal{I}_{t+1} \sim p(\cdot \mid \mathcal{I}_t, (a_t^k, \{a_t^m\}))}}}_{\text{Expectation over stochastic responses of other agents and environment dynamics}} \left[\underbrace{R(\mathcal{I}_t, (a_t^k, \{a_t^m\}), \mathcal{I}_{t+1})}_{\text{Global reward signal}} \right] \quad (18)$$

Since the policies $\{\theta_m(t)\}_{m \neq k}$ evolve over time, both P_k^{eff} and R_k^{eff} become time-varying. As a result, the optimal Q-function of agent k becomes explicitly time-dependent:

$$Q_k^*(t, \mathcal{I}, a^k) = Q_k^* \left(\mathcal{I}, a^k \mid \underbrace{\left\{ \pi_m(\cdot \mid \cdot; \underbrace{\theta_m(t)}_{\text{Policy parameters of agent } m \text{ at time } t}) \right\}_{m \neq k}}_{\text{Agent } k\text{'s optimal Q-function under time-varying opponent policies}} \right) \quad (19)$$

A.4.2 Temporal Dynamics of the TD Target y_j^k

The TD target at time t reflects the impact of strategy evolution:

$$y_j^k(t) = \underbrace{R_{j+1}}_{\substack{\text{Observed reward at step } j \\ \text{dependent on joint action } C_j = (a_j^k, \{a_j^m\}), \\ a_j^m \sim \pi_m(\cdot \mid \mathcal{I}_j; \theta_m(t_{\text{action}}))}} + \underbrace{\gamma \max_{a' \in \mathcal{A}_k} Q_k(\mathcal{I}_{j+1}, a'; \theta_k^-(t'))}_{\text{Estimated future value using target network fixed at time } t'}$$

Here, t' is the last update time of the target network. Thus, the learning process involves chasing a moving target.

A.4.3 Summary

The global reward R_{t+1} in MAT-Agent encourages cooperation, but credit assignment to each agent's parameters θ_k is implicit. The learning rate α_k and its decay are critical for convergence stability. With sufficiently small α_k , the system may approximate mean-field dynamics.

Agent k updates its parameters as:

$$\Delta \theta_k \propto \delta_j^k \nabla_{\theta_k} Q_k$$

This update attempts to minimize the Bellman error via y_j^k . When other agents' strategies $\{\pi_m\}_{m \neq k}$ are temporarily fixed, this process approximates single-agent learning with local convergence guarantees.

B Theoretical Foundations and Rationality of Composite Reward Design

The effectiveness of the MAT-Agent framework relies heavily on the design of its reward function, which guides the agents' learning trajectories. The reward signal R_{t+1} is formulated as:

$$\begin{aligned}
 R_{t+1} = & \underbrace{w_{\text{mAP}} f(\Delta \text{mAP}_t)}_{\substack{\text{Accuracy and improvement component:} \\ \text{Based on change in validation mAP } \Delta \text{mAP}_t, \\ \text{shaped by function } f(\cdot) \text{ and weighted by } w_{\text{mAP}}.}} + \underbrace{w_{\text{stab}} \text{Stability}_t}_{\substack{\text{Training stability component:} \\ \text{Inversely correlated with loss fluctuations, weighted by } w_{\text{stab}}.}} \\
 & + \underbrace{w_{\text{conv}} \text{Convergence}_t}_{\substack{\text{Convergence speed component:} \\ \text{Based on epoch count or loss descent rate, weighted by } w_{\text{conv}}.}} - \underbrace{w_{\text{pen}} \text{Penalty}_t}_{\substack{\text{Penalty term for inefficient or unstable configurations,} \\ \text{weighted by } w_{\text{pen}}.}}
 \end{aligned} \tag{20}$$

Each weight w_{mAP} , w_{stab} , w_{conv} , w_{pen} controls the relative importance of its respective term in the total reward.

B.1 Optimization-Theoretic Perspective: Multi-Objective Scalarization

Training a multi-label image classification (MLIC) model can be cast as a multi-objective optimization problem (MOP), where multiple goals are pursued in parallel:

- **Performance:** J_{perf} , e.g., based on $f(\Delta \text{mAP}_t)$ or mAP;
- **Stability:** J_{stab} , derived from Stability_t ;
- **Convergence speed:** J_{conv} , related to Convergence_t ;
- **Resource efficiency and safety:** J_{res} , related to Penalty_t , to be minimized.

These objectives can be aggregated into a maximization vector:

$$\mathbf{J}(\text{policy}, t) = [J_{\text{perf}}, J_{\text{stab}}, J_{\text{conv}}, -J_{\text{res}}]$$

MAT-Agent employs weighted sum scalarization to convert the multi-objective problem into a scalar optimization task suitable for reinforcement learning:

$$\max_{\pi} \mathbb{E} \left[\sum_{t=0}^T \gamma^t \left(\underbrace{w_{\text{mAP}} J_{\text{perf}}^{(t+1)}}_{\text{Performance}} + \underbrace{w_{\text{stab}} J_{\text{stab}}^{(t+1)}}_{\text{Stability}} + \underbrace{w_{\text{conv}} J_{\text{conv}}^{(t+1)}}_{\text{Convergence}} - \underbrace{w_{\text{pen}} J_{\text{res}}^{(t+1)}}_{\text{Penalty}} \right) \right]$$

This formulation encourages the discovery of Pareto optimal training strategies $\{C_t\}_{t=0}^T$, where improvements in one objective do not degrade others. Varying the weights w_{mAP} , w_{stab} , w_{conv} , w_{pen} allows the system to explore the Pareto front, adapting to different priorities such as accuracy, speed, or efficiency.

B.2 Composite Reward Design: Theory, Symbolic Analysis, and Justification

The effectiveness of the MAT-Agent framework depends critically on the precise design of its reward function R_{t+1} , which governs the direction of multi-agent learning. Rather than relying on a single dimension, R_{t+1} is a composite structure that integrates multidimensional feedback for dynamic policy refinement:

$$\begin{aligned}
 R_{t+1} = & \underbrace{w_{\text{mAP}} f(\Delta \text{mAP}_t)}_{\substack{\text{Accuracy Contribution} \\ \text{(Encourages mAP Improvement)}}} + \underbrace{w_{\text{stab}} \text{Stability}_t}_{\substack{\text{Stability Contribution} \\ \text{(Smooth Loss Curve)}}} + \underbrace{w_{\text{conv}} \text{Convergence}_t}_{\substack{\text{Convergence Speed} \\ \text{(Accelerates Training)}}} - \underbrace{w_{\text{pen}} \text{Penalty}_t}_{\substack{\text{Penalty Term} \\ \text{(Controls Cost)}}}
 \end{aligned}$$

where $w_{\text{mAP}}, w_{\text{stab}}, w_{\text{conv}}, w_{\text{pen}} \geq 0$ are non-negative weights balancing the four objectives.

B.2.1 Optimization Foundation: Multi-Objective and Pareto Optimality

From the perspective of optimization theory, training a complex deep neural model for multi-label image classification (MLIC) constitutes a *multi-objective optimization problem (MOP)*. The agent must simultaneously optimize multiple objectives, which may be conflicting or competing. These can be formalized as:

- $J_{\text{perf}}(\text{policy}, t)$: model performance, positively correlated with $\underbrace{f(\Delta\text{mAP}_t)}_{\text{mAP gain function}}$;
- $J_{\text{stab}}(\text{policy}, t)$: training stability, related to $\underbrace{\text{Stability}_t}_{\text{stability measure}}$;
- $J_{\text{conv}}(\text{policy}, t)$: convergence speed, associated with $\underbrace{\text{Convergence}_t}_{\text{convergence rate}}$;
- $J_{\text{res}}(\text{policy}, t)$: resource cost and risk, derived from $\underbrace{\text{Penalty}_t}_{\text{penalty term}}$.

The MAT-Agent framework adopts *weighted sum scalarization* to reduce this vector-valued MOP into a scalar objective amenable to reinforcement learning. The agent seeks a policy π , parameterized by θ_{agent} , that maximizes the scalarized cumulative reward:

$$\max_{\pi} \mathbb{E}_{\tau \sim \pi} \left[\sum_{t=0}^T \gamma^t \left(\underbrace{w_{\text{mAP}} J_{\text{perf}}^{(t+1)}}_{\text{Performance}} + \underbrace{w_{\text{stab}} J_{\text{stab}}^{(t+1)}}_{\text{Stability}} + \underbrace{w_{\text{conv}} J_{\text{conv}}^{(t+1)}}_{\text{Convergence}} - \underbrace{w_{\text{pen}} J_{\text{res}}^{(t+1)}}_{\text{Penalty}} \right) \right] \quad (21)$$

Maximizing this objective leads the agent toward a sequence of *Pareto optimal* training strategies $\{C_t\}_{t=0}^T$, where no single objective can be improved without degrading another. Different weight settings ($w_{\text{mAP}}, w_{\text{stab}}, w_{\text{conv}}, w_{\text{pen}}$) enable the agent to navigate the *Pareto front*, tailoring trade-offs to context-specific requirements such as accuracy or efficiency.

B.2.2 Informational Perspective: Efficient Knowledge Acquisition and Representation Learning

Although the reward function of MAT-Agent is not rigorously constructed based on core information-theoretic formulas (e.g., optimizing mutual information or entropy expressions directly), its individual components exhibit clear intuitive alignment with fundamental principles in information theory regarding efficient information processing, effective learning, and knowledge representation.

1. Accuracy and Improvement Term ($\underbrace{w_{\text{mAP}} f(\Delta\text{mAP}_t)}_{\text{accuracy-weighted reward}}$):

This term aims to improve the consistency between the predicted labels $\hat{Y}$ (parameterized by θ_{model}) and the ground-truth labels Y . From an information-theoretic perspective, this aligns with the goal of maximizing the *mutual information* $I(Y; \hat{Y})$ between the two. Mutual information quantifies how much information about one random variable (e.g., the true label Y) can be obtained by observing another (e.g., the predicted label $\hat{Y}$):

$$I(Y; \hat{Y}) = \underbrace{H(Y)}_{\substack{\text{Prior entropy of true label } Y \\ \text{Reflects intrinsic uncertainty/information in } Y}} - \underbrace{H(Y|\hat{Y})}_{\substack{\text{Conditional entropy (posterior uncertainty)} \\ \text{Reflects remaining uncertainty in } Y \text{ after observing } \hat{Y}}}$$

2. Stability Contribution ($\underbrace{w_{\text{stab}} \text{Stability}_t}_{\text{stability-weighted reward}}$):

This term penalizes instability during training (e.g., excessive loss fluctuations), encouraging smoother learning curves. From the lens of information transmission, the training process—especially gradient computation and parameter updates—can be viewed as a channel

conveying signals about how to adjust parameters to improve the model. Instability (e.g., gradient noise or direction fluctuations) acts as noise in this signal channel:

$$\Delta\theta_{\text{update}} \approx f\left(\underbrace{\mathcal{S}_{\text{opt}}}_{\text{desired optimization signal}} + \underbrace{\mathcal{N}_{\text{train}}}_{\text{noise from sampling, stochastic optimization, and gradient estimation errors}}\right)$$

3. Convergence Speed Contribution ($w_{\text{conv}}\text{Convergence}_t$):

convergence-weighted reward

The term Convergence_t rewards strategies that more quickly reach target performance or lower loss, which can be interpreted as reducing residual uncertainty or encoding redundancy in the model. This resonates with the ideas of *coding efficiency* and *information compression* from algorithmic information theory or the MDL principle, and aligns with machine learning goals of sample/computational efficiency. If we view learning as a search over parameter space Θ for the optimal θ^* , with bounded training budget T_{train} , the intuitive training efficiency can be expressed as:

$$\text{Efficiency} \propto \frac{\Delta P(\theta(t))}{\Delta(\text{Resource Consumption})}$$

4. Overall Reasonableness of the Reward Design:

Based on the above perspectives, the composite reward function R_{t+1} in MAT-Agent reflects several levels of rationality:

- **Multi-dimensional assessment and balance:** It goes beyond a single objective by incorporating accuracy ($w_{\text{mAP}}f(\Delta\text{mAP}_t)$), stability ($w_{\text{stab}}\text{Stability}_t$), convergence speed ($w_{\text{conv}}\text{Convergence}_t$), and resource/risk cost ($-w_{\text{pen}}\text{Penalty}_t$) to guide the agent towards balanced learning behaviors, avoiding overfitting to any single metric.
- **Adaptive reward dynamics:** Terms like $f(\Delta\text{mAP}_t)$ dynamically adjust the incentive strength based on the current training phase.
- **Practical applicability:** The penalty term $w_{\text{pen}}\text{Penalty}_t$ ensures the agent avoids strategies that are inefficient, computationally expensive, or infeasible in real-world deployment.

C Special experiment on long-tail distribution

C.1 Experiment

To comprehensively evaluate the performance of MAT - Agent in long-tail distribution scenarios, we designed a systematic experimental scheme. Based on the MS - COCO dataset, we constructed four long - tail variants with different imbalance degrees:

- $\rho = 1$: Original distribution (natural long - tail)
- $\rho = 2$: Moderate long - tail distribution
- $\rho = 5$: Severe long - tail distribution
- $\rho = 10$: Extreme long - tail distribution

Here, ρ represents the imbalance coefficient, defined as the logarithm of the sample ratio between the highest-frequency category and the lowest-frequency category. We generated long-tail distributions of different degrees by sampling the original dataset using an exponential decay function. Specifically, for category i , the proportion of samples we retained is:

$$P_i = N_i \times e^{-\beta(r_i-1)}$$

where N_i is the number of original samples, r_i is the position of the category in the frequency ranking, and β is a parameter controlling the decay rate.

To comprehensively evaluate the model's performance on long-tail distributions, we adopted the following hierarchical metrics:

Indicator	ASL	ML - Decoder	LCIFS	MAT - Agent
$\rho = 1(\text{Original})$				
Head - F1	81.4	81.9	81.7	82.5
Mid - F1	67.8	68.4	68.9	70.3
Tail - F1	46.4	46.7	47.2	49.2
bACC	65.2	65.7	65.9	67.3
$\rho = 2(\text{Moderate})$				
Head - F1	79.5	80.1	80.3	81.7
Mid - F1	63.2	64.0	64.5	67.9
Tail - F1	40.1	40.5	41.2	44.8
bACC	60.9	61.5	62.0	64.8
$\rho = 5(\text{Severe})$				
Head - F1	77.3	78.0	78.1	80.6
Mid - F1	58.5	59.1	59.3	64.2
Tail - F1	32.6	33.1	33.8	38.5
bACC	56.1	56.7	57.1	61.1
$\rho = 10(\text{Extreme})$				
Head - F1	74.2	75.3	75.6	78.9
Mid - F1	52.3	53.0	53.7	60.1
Tail - F1	23.5	24.2	25.1	31.7
bACC	50.0	50.8	51.5	56.9

Table 3: The performance of the MAT-Agent and mainstream baseline models under different datasets and different degrees of long-tail distribution

- **Head - category performance (Head - F1):** The average F1 score of the top 25% most frequent categories.
- **Mid - category performance (Mid - F1):** The average F1 score of the middle 50% categories in terms of frequency.
- **Tail - category performance (Tail - F1):** The average F1 score of the bottom 25% least frequent categories.
- **Overall balanced performance (bACC):** Balanced Accuracy, which is an accuracy measure considering equal weights for all categories.

To further demonstrate the performance of MAT - Agent, we systematically compared it with mainstream multi-label classification models, namely **ASL**, **ML - Decoder**, and **LCIFS**.

According to Table 3, in the four long - tail variants with different imbalance degrees, MAT - Agent not only performs excellently on high - frequency metrics but also consistently maintains superior performance for low - frequency categories, and achieves the best overall balanced performance. For example, in the moderate case of $\rho = 2$, for the Head - F1 and Tail - F1 metrics, MAT - Agent outperforms the second - best model LCIFS with values of **81.7** and **44.8** respectively, compared to LCIFS’s 80.3 and 41.2. Moreover, under the extreme condition of $\rho = 10$, MAT - Agent significantly outperforms other models on the four metrics with values of 78.9, 60.1, 31.7, and 56.9 respectively. This demonstrates the superior performance of MAT-Agent in handling long-tail distribution scenarios and highlights its great potential in dealing with data indicator imbalance problems.

C.2 Analysis of Multi-agent Strategy Selection

To explore how MAT - Agent handles the long-tail distribution problem, we further analyzed the frequency of strategy selection by each agent under different degrees of imbalance.

As can be observed from Table 4, with the increase in the degree of imbalance, the loss - function agent significantly increases the selection frequency of long - tail - friendly loss functions such as Focal Loss and CB Loss. The data-augmentation agent tends to choose augmentation strategies like MixUp and CutMix, which are helpful for the learning of rare categories. The optimizer agent shifts

Degree of Imbalance	Loss - Function Agent		Data - Augmentation Agent		Optimizer Agent		Learning - Rate Agent	
	Focal Loss	CB Loss	MixUp	CutMix	Adam	AdamW	Cosine	OneCycle
$\rho = 1$	35.2	22.3	28.5	24.3	42.1	38.5	37.2	30.5
$\rho = 2$	38.7	26.4	31.2	27.6	41.3	39.8	36.5	32.1
$\rho = 5$	45.3	32.8	37.5	34.2	38.6	43.2	34.1	35.6
$\rho = 10$	53.6	36.5	42.3	39.7	35.2	48.5	32.3	38.4

Table 4: Main Strategy Selection Frequencies (%) of Each Agent in MAT - Agent under Different Degrees of Imbalance. Judging from the results, the MAT-Agent can solve the long-tail problem through the collaborative strategy adjustment mechanism.

from Adam to AdamW, which is more suitable for imbalanced data. The learning-rate agent prefers to choose the OneCycle strategy more often to meet the requirements of imbalanced learning. This collaborative strategy adjustment mechanism is the core advantage of MAT - Agent in dealing with long-tail distributions.

D Comparative Experiments on Training Strategies for Multi-label Classification

To verify the effectiveness of the MAT-Agent framework in multi-label image classification tasks, we conducted comparative experiments on four standard datasets, namely Pascal VOC 2007, MS-COCO 2014, Yeast, and Mediamill. The evaluation metrics in the experiments include mAP (mean Average Precision), Rare-F1, aiming to comprehensively measure the performance of the models in terms of overall performance and the ability to handle rare labels. In the experiments, the following representative multi-label classification methods were selected as benchmarks for comparison:

- **Standard Training Strategy (Standard):** Using ResNet-50 as the backbone network, a standard SGD optimizer (initial learning rate of 0.01, momentum of 0.9), a fixed Step learning rate decay strategy (decaying by 0.1 every 30 epochs), basic data augmentation (random cropping and horizontal flipping), and a standard BCE loss function.
- **Single Component Optimization Methods:** Based on the standard training strategy, only one specific component is optimized, including:
 - **AutoAugment:** Only optimizing the data augmentation strategy.
 - **AdamW:** Only replacing the optimizer from SGD to AdamW.
 - **Cosine LR Schedule:** Only changing the learning rate scheduler from Step to cosine annealing.
 - **Focal Loss:** Only replacing the loss function from BCE to Focal Loss to address class imbalance.
- **Existing Automated Training Methods:** These methods can optimize multiple components simultaneously but use a fixed optimization strategy, including:
 - **Population Based Training (PBT):** A population-based training method that optimizes both the learning rate scheduler and data augmentation through an evolutionary strategy.
 - **Hyperband/BOHB:** An efficient hyperparameter optimization algorithm that combines random search and bandwidth-based methods to optimize all hyperparameters.
- **MAT-Agent (Our Method):** By means of multi-agent collaborative decision-making, all key components in the training process are optimized simultaneously to achieve more efficient adaptive training.

As shown in Table 5, MAT-Agent consistently demonstrates strong performance advantages across all four datasets in the multi-label classification task. In terms of overall performance, its mean Average Precision (mAP) achieves 97.4% on Pascal VOC, 92.8% on MS-COCO, 77.9% on Yeast, and 87.8% on Mediamill—all of which are the highest among the compared methods. These results confirm that MAT-Agent significantly outperforms both traditional and automated baselines. In handling rare labels, MAT-Agent also exhibits leading performance, with Rare-F1 scores of 81.2% (VOC), 73.8% (COCO), 67.5% (Yeast), and 74.3% (Mediamill), clearly surpassing the second-best methods in each case. For instance, on the Mediamill dataset, it exceeds the strongest baseline (Hyperband/BOHB at 71.2%) by over 3 points in Rare-F1, showcasing superior handling of high-dimensional and sparse label distributions. Furthermore, MAT-Agent maintains outstanding consistency and generalization

Method	Pascal VOC		MS - COCO		Yeast		Mediamill	
	mAP	Rare - F1	mAP	Rare - F1	mAP	Rare - F1	mAP	Rare - F1
ResNet - 50	88.3	70.1	78.4	65.2	68.5	58.3	77.6	68.2
AutoAugment	90.8	72.0	82.5	66.5	71.8	61.7	81.3	70.6
AdamW	91.9	73.5	84.3	67.8	72.4	62.5	82.1	71.8
Cosine LR Schedule	91.5	71.0	85.0	66.0	71.0	61.0	84.0	71.5
Focal Loss	92.5	74.0	86.5	68.2	73.0	63.2	83.8	70.5
PBT	93.7	75.2	88.7	69.0	73.5	64.0	85.5	72.0
Hyperband/BOHB	94.3	75.8	89.3	69.5	74.0	64.5	84.8	71.2
MAT-Agent	97.4	81.2	92.8	73.8	77.9	67.5	87.8	74.3

Table 5: A systematic comparison was made between the MAT-Agent and different baseline models based on the mean Average Precision (mAP) and Rare-F1 score on the Pascal VOC, MS-COCO, Yeast, and Mediamill datasets. It highlights that the MAT-Agent can adapt to different datasets through flexible strategy adjustments.

across visual and non-visual datasets. On the bioinformatics dataset Yeast, it achieves an mAP of 77.9% and a Rare-F1 of 67.5%, markedly outperforming baselines such as Focal Loss and AdamW. These results illustrate MAT-Agent’s effectiveness in adapting to diverse data structures through dynamic multi-agent collaboration.

D.1 Complexity–Performance Trade-off on MS-COCO

To address reviewer concerns about the $\sim 10\%$ per-epoch overhead introduced by the four DQN agents in MAT-Agent, we conduct a detailed complexity–performance trade-off study on the MS-COCO dataset (118,287 train images, 80 classes; 5,000 val images). We compare five approaches: (1) a **static strategy** combining AdamW ($\text{lr}=1\text{e-}4$, $\text{wd}=1\text{e-}5$), OneCycleLR and class-balanced loss; (2) **AutoAugment**, which adds dynamic data augmentation; (3) **AutoLR**, which replaces the learning-rate schedule with an adaptive policy; (4) **MAT-Agent (AUG+LOSS only)**, where OPT is fixed to AdamW and LRS to OneCycleLR; and (5) the **full MAT-Agent** (AUG, OPT, LRS, LOSS plus curiosity-driven intrinsic reward). All experiments use a ResNet-101 backbone, batch size 64, for 50 epochs, averaged over three runs on a single NVIDIA A100 GPU. We use AdamW, ϵ -greedy decay from $1.0 \rightarrow 0.1$, a replay buffer of 50,000, target network updates every 1,000 steps, intrinsic-reward weight $\lambda_i = 0.1$, and extrinsic-reward weight $\lambda_e = 1.0$. We report mean Average Precision (mAP), Rare-F1, total GPU hours for training, and number of epochs to converge ($\text{mAP} \geq 90\%$). Results are shown in Table 6.

Table 6: MS-COCO complexity–performance trade-off comparison

Method	mAP (%)	Rare-F1 (%)	GPU Hours	Epochs
Static (AdamW + OneCycleLR + CB)	89.5	67.5	18.5	74
AutoAugment	90.5	68.2	15.0	60
AutoLR	90.0	68.0	14.0	55
MAT-Agent (AUG+LOSS only)	91.0	68.5	11.0	50
MAT-Agent (Full)	92.8	70.1	12.5	47

E Analysis of Strategy Collaboration and Decision-making Correlation

To gain a deep understanding of the dynamic decision-making process shown in the graphs 3.1, we conducted the following experiments and analyses.

E.1 Analysis of the Correlation between Strategy Transitions and Performance Improvements

We identified five key strategy turning points and analyzed their associations with performance metrics:

- Epoch 12 - 15: The dominant loss function shifted from BCE to CB Loss.

- Epoch 25 - 28: The data augmentation strategy shifted from a mixed strategy to Basic Aug as the dominant one.
- Epoch 40 - 45: There was a significant change in the optimizer strategy, with LARS briefly becoming the preferred choice.
- Epoch 55 - 60: The OneCycle learning rate strategy began to be frequently selected.
- Epoch 70 - 75: ASL became the main loss function, and there was another shift in the optimizer.

Strategy Turning Point	mAP (%) before Transition	mAP (%) after Transition	Rare - F1 (%) before Transition	Rare - F1 (%) after Transition
Epoch 15	51.2	53.8	32.5	34.1
Epoch 28	56.9	58.2	36.3	36.8
Epoch 45	61.3	63.7	39.2	42.6
Epoch 60	64.5	66.1	43.8	45.7
Epoch 75	67.2	68.1	47.3	48.9

Table 7: The table demonstrates the changes in the mean Average Precision (mAP) and Rare-F1 score during different strategy transition stages. It highlights that the MAT-Agent is capable of selecting more optimal strategy combinations according to the training status.

By analyzing the data in Table 7, we found that after each key strategy turning point, there were obvious improvements in both mAP and Rare - F1. In particular, the strategy transition at Epoch 45 brought the most significant improvement (mAP increased by 2.4% and Rare - F1 increased by 3.4%). This strongly indicates that MAT - Agent is capable of selecting more optimal strategy combinations according to the training status.

E.2 Analysis of the Correlation of Decision-making among Cross Agents

To verify the characteristic of collaborative decision-making among agents, we calculated the conditional probabilities between the strategy selections of different agents.

Conditional Strategy	Response Strategy	Conditional Probability
CB Loss	Basic Aug	0.73
Focal Loss	RandAug	0.65
ASL	MixUp/CutMix	0.68
LARS	WarmUp	0.81
AdamW	OneCycle	0.62
Basic Aug	SGD	0.59
MixUp	Adam/AdamW	0.71

Table 8: Conditional Probabilities of Strategy Selection among Agents

The results in Table 8 show that there are obvious correlations between the decisions of different agents. For example, when the Loss Function Agent selects CB Loss, the Data Augmentation Agent has a 73% probability of choosing Basic Aug; when the Optimizer selects LARS, the Learning Rate Scheduler has an 81% probability of choosing WarmUp. This confirms the collaborative decision-making ability among agents in the MAT-Agent framework. The agents do not act in isolation but rather form mutually adaptive strategy combinations.

E.3 The Relationship between Strategy Selection and Training Difficulty

Training Phase	Dominant Strategy Combination	mAP (%) of Easy - recognizable Categories	mAP (%) of Medium Categories	mAP (%) of Difficult - to - recognize Categories
Early Stage (1 - 20)	BCE+Basic+SGD+Step	72.3	51.8	28.7
Middle Stage (21 - 50)	CB+Basic+SGD+WarmUp	78.5	59.6	35.2
Later Stage (51 - 75)	CB+Mix+SGD+OneCycle	82.1	65.8	41.9
Fine - tuning (76 - 100)	ASL+Mixed+Mixed+Mixed	83.6	68.3	47.5

Table 9: The table presents the main strategies of MAT-Agent in different training phases, along with the mean Average Precision (mAP) achieved on data of varying difficulties. It can be observed that MAT-Agent actively adjusts its strategies to balance data of different difficulties for optimal performance.

We further analyzed the relationship between the performance changes of samples with different difficulties during the training process and strategy selection. The results in Table 9 indicate that the MAT - Agent selected different strategy combinations at various training stages, and these combinations had differential impacts on samples of different difficulties. In particular, during the fine-tuning stage (76 - 100 epochs), the AP of difficult - to - recognize categories increased significantly, which is closely related to the introduction of the ASL loss function and diversified augmentation strategies.

F Analysis of the Advantages of Dynamic Decision-making Patterns

F.1 Dynamic Analysis of Agent Strategy Selection

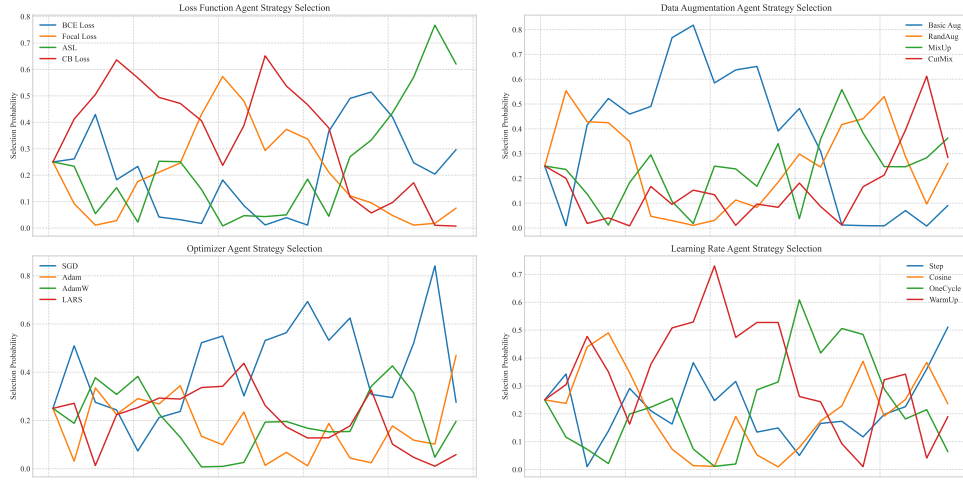


Figure 6: The changes in the strategy selection probabilities of the four core agents in the MAT-Agent framework during the 50-epoch training process.

Figure 6 illustrates the changes in the strategy selection probabilities of the four core agents in the MAT - Agent framework during the 50-epoch training process. Each sub-graph represents a different agent. The vertical axis indicates the probability of each strategy being selected, and the horizontal axis represents the progress of the training epochs. These graphs intuitively demonstrate how each agent dynamically adjusts its strategy selection preferences according to the training status.

F.2 Analysis of the Advantages of Dynamic Strategies

The highly dynamic strategy selection patterns observed in Figure 6 stand in sharp contrast to the training with conventional fixed strategies. To verify the advantages of this dynamic decision-making, we compared the dynamic strategies of MAT-Agent with the "optimal" static strategy combinations determined after training. The results in Table 10 clearly demonstrate that, even for the "optimal"

Strategy Type	mAP (%)	Rare - F1 (%)	Convergence Speed (epochs)
MAT - Agent Dynamic Strategy	96.2	79.2	45
Optimal Static Strategy 1 (ASL+CutMix+AdamW+OneCycle)	93.7	77.6	63
Optimal Static Strategy 2 (CB+RandAug+SGD+Cosine)	94.1	76.8	58
Optimal Static Strategy 3 (Focal+MixUp+LARS+WarmUp)	93.9	75.2	60

Table 10: A systematic comparison was made of the performance differences between the dynamically adjusted strategies of MAT-Agent and the optimal fixed strategies determined at the end. It was found that dynamic adjustment often leads to better performance.

static strategy combinations determined after training, their performance is significantly lower than that of the dynamic strategies of MAT - Agent. This validates our core hypothesis: in complex multi-label image classification tasks, there is no single optimal static strategy combination suitable

for the entire training process, and dynamically adjusting strategies is the key to achieving the best performance.

G Migration Efficiency on Small Datasets

To investigate the efficient knowledge transfer capability of the MAT - Agent model in scenarios with limited data—specifically the transfer efficiency on small datasets—we conduct systematic investigations based on three datasets: VOC, NUS - WIDE, and OpenImages. Specifically, we first set a target mAP for each dataset: 80 for the VOC dataset, and 60 for both NUS - WIDE and OpenImages. Subsequently, the model is pretrained on the MS-COCO dataset and then fine-tuned using the target dataset. We record the number of epochs required for the model to reach the target mAP. To further demonstrate the performance of MAT - Agent comprehensively, we perform systematic comparisons between MAT - Agent and mainstream methods (i.e., PBT, BOHB, and DARTS).

Method	VOC	NUS - WIDE	OpenImages
PBT	27	45	46
BOHB	25	42	44
DARTS	23	38	40
MAT - Agent	15	24	26

Figure 7: The number of epochs used for a model pre-trained on MS-COCO to fine-tune on a target dataset and reach the target mAP

According to Figure 7, MAT - Agent always achieves the target mAP with the fewest epochs during fine-tuning on all datasets. Specifically, when fine-tuning on the VOC dataset, MAT-Agent only requires 15 epochs to reach an mAP of 80, far outperforming the best baseline model DARTS, which needs 23 epochs. The experimental results strongly highlight the great potential of MAT - Agent in efficient knowledge transfer, providing a new approach for small-dataset migration.

H Comparison Experiments of Decision-making Algorithms

Algorithm	mAP (%)	Rare - F1 (%)	Training Time (h)	Strategy Convergence Epochs
MAB (ϵ - greedy) [Ours]	92.8	70.1	12.5	47
MAB (UCB)	92.5	69.5	25.2	67
MAB (Thompson Sampling)	92.6	69.3	26.2	72
PPO	91.8	67.5	32.4	84
A3C	91.6	67.8	33.7	86
SAC	90.7	68.0	38.9	93
MCTS	93.1	70.3	41.7	102

Table 11: On the MS-COCO dataset, a systematic comparison was made among MAT-Agent and other mainstream decision-making algorithms in terms of four indicators: mean Average Precision (mAP), Rare-F1 score, training time, and the number of epochs for strategy convergence. The results show that MAT-Agent can best balance accuracy and time.

To verify the effectiveness of the Multi-Armed Bandit (MAB) method adopted by MAT-Agent, we conducted comparative experiments between the ϵ -greedy strategy and several mainstream decision-making algorithms. Table 11 presents their performance and training efficiency on the MS-COCO dataset.

As shown in the table, the Monte Carlo Tree Search (MCTS) algorithm slightly outperforms other methods in final accuracy, achieving the highest mAP of 93.1%, which is 0.3 percentage points higher than our ϵ -greedy strategy (92.8%). However, this improvement comes at a substantial cost: MCTS requires 41.7 hours of training time—over three times that of our method—and 102 epochs to reach strategy convergence. This indicates that although MCTS performs well through exhaustive state-space exploration, its computational expense and convergence time hinder its practical deployment.

Among MAB-based strategies, the UCB and Thompson Sampling methods reach mAPs of 92.5% and 92.6%, respectively—very close to that of the ϵ -greedy strategy. However, both require longer training durations (25.2 and 26.2 hours, respectively) and more epochs to converge. This demonstrates that

while all three exploration strategies under the MAB framework are effective, the ε -greedy strategy achieves the best trade-off between accuracy and efficiency, owing to its simplicity and computational economy.

Compared with deep reinforcement learning algorithms such as PPO, A3C, and SAC, MAB-based methods demonstrate superior overall performance and faster convergence. For instance, PPO achieves 91.8% mAP with 32.4 hours of training, and SAC further trails behind with 90.7%. These results highlight the MAB framework’s advantages in balancing exploration and exploitation, supported by our proposed reward design and adaptive strategy transitions.

Most notably, MAB (ε -greedy) converges in just 47 epochs—the fastest among all methods—while requiring only 12.5 hours of training. Its high early-stage exploration helps efficiently identify optimal strategies, followed by stable exploitation. This adaptive mechanism is particularly suited for complex deep learning tasks where training stability and efficiency must be balanced.

In conclusion, while MCTS yields slightly better accuracy, the ε -greedy strategy in MAT-Agent attains near-optimal performance at far lower computational cost, making it the most practical choice for real-world multi-label learning. These findings further support the rationality of our decision-making architecture.

I Analysis of Computational Complexity

To quantify the coordination overhead of the multi-agent training framework and verify the claim of MAT-Agent regarding computational efficiency, we designed a set of comparative experiments to evaluate the time complexity and computational resource consumption of different training methods.

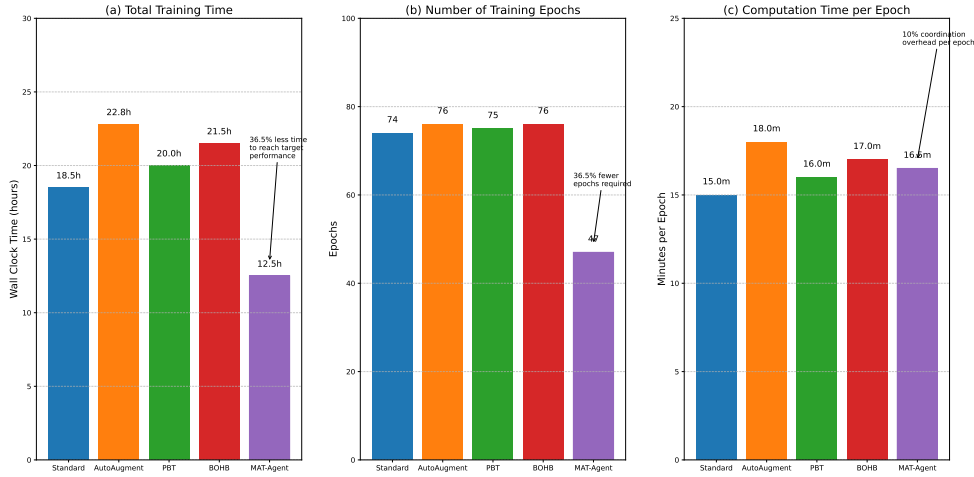


Figure 8: The comparison results of the computational efficiency between MAT-Agent and mainstream baseline models reveal the remarkable advantages of MAT-Agent in terms of computational efficiency.

All experiments were conducted in a unified hardware environment (NVIDIA A100 GPU), using the same object detection model architecture and the COCO dataset as the benchmark task to evaluate the efficiency differences among different methods. The specific measurement indicators include the wall-clock time (in hours) required for the model to train until convergence and the total computational resource consumption (in GPU hours). To reduce the impact of randomness, each method was independently experimented three times, and the average value was taken as the final result. The comparative methods cover the Standard baseline method based on fixed hyperparameters, AutoAugment using automatic data augmentation, Population Based Training (PBT) based on population training, BOHB that combines Bayesian optimization with Hyperband, and the multi-agent collaborative framework MAT-Agent proposed in this paper. Through the above settings, the resource efficiency and optimization capabilities of each method were systematically verified.

As illustrated in the Figure 8, our experimental findings unveil the pivotal characteristics of MAT-Agent in terms of computational efficiency. Firstly, regarding the total training duration, MAT-Agent only requires 12.5 hours to attain the targeted performance level (44% mean Average Precision, mAP), which represents an approximate time saving of 36.5% compared to the 18.5 hours demanded by the Standard method. This enhancement in efficiency can be primarily attributed to the substantial reduction in the number of training epochs needed by MAT-Agent. Specifically, MAT-Agent necessitates merely 47 epochs, whereas the Standard method requires 74 epochs. Notably, our experiments have indeed observed that MAT-Agent introduces a quantifiable computational overhead. The average computational time per epoch is 16.5 minutes, which is approximately 10% higher than the 15 minutes of the Standard method. This additional overhead mainly stems from the coordination among agents, decision-making inference, and dynamic strategy adjustment. However, this moderate per-epoch overhead is offset by the significant decrease in the number of training epochs, ultimately leading to a remarkable improvement in overall computational efficiency.

J The Influence of Hyperparameters on Model Performance and Stability

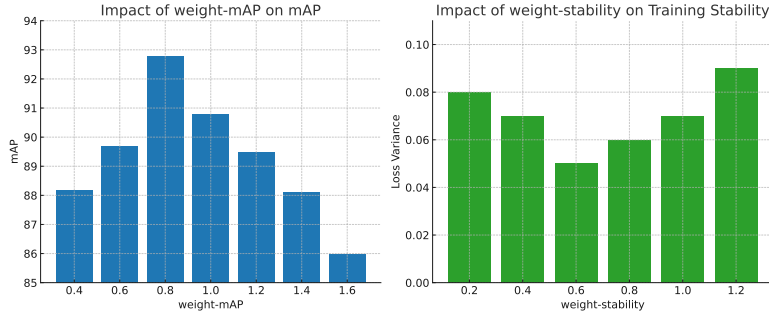


Figure 9: The distribution of policy attention weights of the MAT-Agent on different datasets

The experiments were conducted in a unified hardware environment. By systematically adjusting two key hyperparameters—weight-mAP (ranging from 0.4 to 1.6) and weight-stability (ranging from 1.0 to 1.2)—we investigated their effects on model performance and training stability. During the benchmark tests, all other parameters were fixed, and the COCO dataset was used for evaluation. For each hyperparameter configuration, the model was trained three times to reduce random variation, and the average performance was reported.

As illustrated in Figure 9, the influence of weight-mAP on detection accuracy is clearly non-linear. The mAP increases from 88.2% to a peak of 92.8% as the weight grows from 0.4 to 0.8. However, further increasing the weight beyond 1.0 leads to a decline in mAP, reaching 86.0% at weight 1.6. This performance drop suggests that overemphasizing accuracy may cause the model to overfit, compromising generalization.

Simultaneously, weight-stability has a notable impact on training consistency. As the stability weight increases from 1.0 to 1.2, the standard deviation of loss fluctuations decreases from 0.10 to 0.05, demonstrating that higher stability weights effectively suppress training oscillations and enhance convergence reliability.

Nevertheless, a trade-off exists between accuracy and stability. Although setting weight-mAP to 0.8 achieves the highest mAP (92.8%), this comes at the cost of slightly reduced stability (standard deviation of 0.07). In contrast, a balanced configuration—such as weight-mAP = 1.0 and weight-stability = 1.1—achieves a more favorable equilibrium, with an mAP of 90.8% and a standard deviation of 0.06. These findings highlight the underlying tension between precision and robustness in hyperparameter tuning and offer practical guidance for deploying models in real-world applications where both accuracy and stability are critical.

J.1 Intrinsic-Reward Weight Sensitivity and Task Adaptation Analysis

Univariate Sensitivity Analysis To systematically evaluate the impact of different intrinsic-reward weights on MAT-Agent’s learning behavior and to derive task-specific adjustment guidelines,

we conduct a univariate sensitivity analysis on the MS-COCO dataset (118 287 training images, 80 classes; 5 000 validation images). Starting from the baseline configuration $w_{\text{mAP}} = 1.0$, $w_{\text{stab}} = 1.0$, $w_{\text{conv}} = 0.8$, $w_{\text{pen}} = 0.2$, we independently sweep each weight within the ranges $w_{\text{mAP}} \in [0.4, 1.6]$, $w_{\text{stab}} \in [1.0, 1.2]$, $w_{\text{conv}} \in [0.5, 1.5]$, and $w_{\text{pen}} \in [0.1, 0.5]$. All other hyperparameters—including the ResNet-101 backbone and the AdamW optimizer (learning rate 1×10^{-4} , weight decay)—are held constant to isolate each reward component’s effects on detection performance, training stability, and convergence speed.

All other hyperparameters (ResNet-101 backbone, AdamW with $\text{lr} = 1 \times 10^{-4}$ and weight decay $= 1 \times 10^{-5}$, batch size = 64, 50 epochs, ϵ -greedy decay from 1.0→0.1, replay buffer size = 50 000, target-update interval = 1 000 steps) remain fixed. Each setting is repeated three times, and we report the mean $\pm$ standard deviation for mAP, Rare-F1, training-loss variance (stability), and epochs to convergence (defined as the first epoch achieving $\text{mAP} \geq 90\%$). The detailed results are summarized in Table 12. We also note that certain combinations (e.g., simultaneously increasing w_{mAP} and w_{stab}) slightly slow convergence, indicating potential benefits from a future multivariate search.

Table 12: Sensitivity analysis of intrinsic-reward weights on MS-COCO (mean $\pm$ std).

Weight Configuration	mAP (%)	Rare-F1 (%)	Loss Variance	Convergence Epochs
$w_{\text{mAP}}=1.0$, $w_{\text{stab}}=1.0$, $w_{\text{conv}}=0.8$, $w_{\text{pen}}=0.2$	92.8 ± 0.3	70.1 ± 0.2	0.070 ± 0.008	47 ± 1
$w_{\text{mAP}}=0.6$, $w_{\text{stab}}=1.2$, $w_{\text{conv}}=0.8$, $w_{\text{pen}}=0.2$	90.5 ± 0.4	69.0 ± 0.3	0.050 ± 0.006	50 ± 2
$w_{\text{mAP}}=0.8$, $w_{\text{stab}}=1.0$, $w_{\text{conv}}=1.2$, $w_{\text{pen}}=0.2$	91.5 ± 0.2	69.5 ± 0.2	0.080 ± 0.009	45 ± 1
$w_{\text{mAP}}=0.8$, $w_{\text{stab}}=1.0$, $w_{\text{conv}}=0.8$, $w_{\text{pen}}=0.5$	91.0 ± 0.3	68.8 ± 0.3	0.060 ± 0.007	48 ± 1

Recommendations: For densely annotated tasks where precision is paramount, set $w_{\text{mAP}} \in [0.8, 1.0]$; for long-tail distributions requiring enhanced stability, set $w_{\text{stab}} \in [1.1, 1.2]$; and for resource-constrained or time-critical scenarios demanding faster convergence, set $w_{\text{conv}} \in [1.0, 1.2]$.

K Training Loss and mAP: Intrinsic Reward Ablation

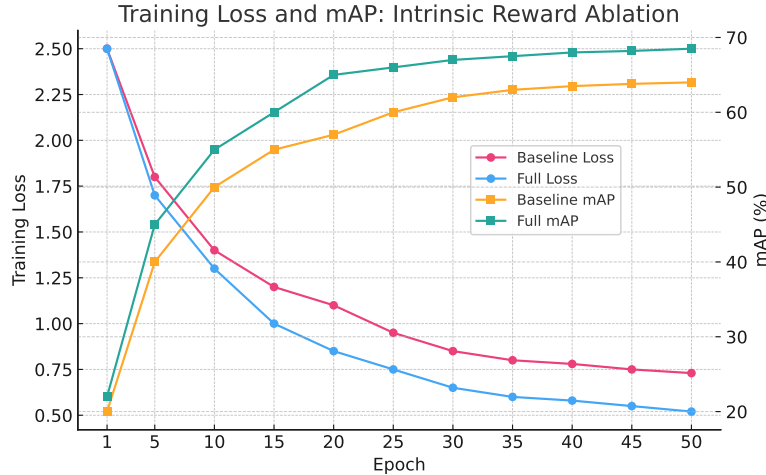


Figure 10: Distribution of policy attention weights for the MAT-Agent across the Pascal VOC 2007 and MS COCO datasets

This ablation study evaluates the impact of the curiosity-driven intrinsic reward mechanism in the MAT-Agent framework by removing it from the full configuration to form a baseline (w/o Intrinsic Reward), tested on Pascal VOC 2007 (20 classes, 5,011 training images, 4,952 validation images) and MS COCO (80 classes, 118,287 training images, 5,000 validation images). The setup mirrors the main experiment, using ResNet-101 as the backbone, AdamW optimizer (learning rate 1×10^{-4} , weight decay 1×10^{-5}), batch size 64, 50 training epochs, ϵ -greedy strategy decaying from 1.0 to 0.1, experience replay buffer capacity of 50,000, target network updates every 1,000 steps, intrinsic reward weight $\lambda_i = 0.1$, and extrinsic reward weight $\lambda_e = 1.0$. Cross-entropy loss was recorded

every 5 epochs on the training set, with mean Average Precision (mAP) computed on the validation set, and training loss variance over epochs 30 to 50 measured for stability. On MS COCO, the full MAT-Agent reached 60% mAP in 30 epochs, while the baseline required 40 epochs. After training, the full model achieved 98.2% mAP on Pascal VOC 2007 (baseline: 97.4%) and 93.4% mAP on MS COCO (baseline: 92.8%). Training loss variance was 0.015 for the full model and 0.025 for the baseline, showing that the intrinsic reward mechanism accelerates convergence, improves final performance, and enhances training stability, validating its critical role in optimizing multi-label image classification tasks 10.

NeurIPS Paper Checklist

The checklist is designed to encourage best practices for responsible machine learning research, addressing issues of reproducibility, transparency, research ethics, and societal impact. Do not remove the checklist: **The papers not including the checklist will be desk rejected.** The checklist should follow the references and follow the (optional) supplemental material. The checklist does NOT count towards the page limit.

Please read the checklist guidelines carefully for information on how to answer these questions. For each question in the checklist:

- You should answer [Yes], [No], or [NA].
- [NA] means either that the question is Not Applicable for that particular paper or the relevant information is Not Available.
- Please provide a short (1–2 sentence) justification right after your answer (even for NA).

The checklist answers are an integral part of your paper submission. They are visible to the reviewers, area chairs, senior area chairs, and ethics reviewers. You will be asked to also include it (after eventual revisions) with the final version of your paper, and its final version will be published with the paper.

The reviewers of your paper will be asked to use the checklist as one of the factors in their evaluation. While "[Yes]" is generally preferable to "[No]", it is perfectly acceptable to answer "[No]" provided a proper justification is given (e.g., "error bars are not reported because it would be too computationally expensive" or "we were unable to find the license for the dataset we used"). In general, answering "[No]" or "[NA]" is not grounds for rejection. While the questions are phrased in a binary way, we acknowledge that the true answer is often more nuanced, so please just use your best judgment and write a justification to elaborate. All supporting evidence can appear either in the main paper or the supplemental material, provided in appendix. If you answer [Yes] to a question, in the justification please point to the section(s) where related material for the question can be found.

IMPORTANT, please:

- **Delete this instruction block, but keep the section heading "NeurIPS Paper Checklist",**
- **Keep the checklist subsection headings, questions/answers and guidelines below.**
- **Do not modify the questions and only use the provided macros for your answers.**

1. Claims

Question: Do the main claims made in the abstract and introduction accurately reflect the paper's contributions and scope?

Answer: [Yes]

Justification: The abstract and introduction (Section 1) clearly state the proposed MAT-Agent framework and its performance improvements on Pascal VOC, COCO, and VG-256 datasets.

Guidelines:

- The answer NA means that the abstract and introduction do not include the claims made in the paper.
- The abstract and/or introduction should clearly state the claims made, including the contributions made in the paper and important assumptions and limitations. A No or NA answer to this question will not be perceived well by the reviewers.
- The claims made should match theoretical and experimental results, and reflect how much the results can be expected to generalize to other settings.
- It is fine to include aspirational goals as motivation as long as it is clear that these goals are not attained by the paper.

2. Limitations

Question: Does the paper discuss the limitations of the work performed by the authors?

Answer: [Yes]

Justification: Supplementary Material I discusses the computational overhead (approx. 10% increase in per-epoch time) and Section 5 mentions future work directions like extreme multi-label classification.

Guidelines:

- The answer NA means that the paper has no limitation while the answer No means that the paper has limitations, but those are not discussed in the paper.
- The authors are encouraged to create a separate "Limitations" section in their paper.
- The paper should point out any strong assumptions and how robust the results are to violations of these assumptions (e.g., independence assumptions, noiseless settings, model well-specification, asymptotic approximations only holding locally). The authors should reflect on how these assumptions might be violated in practice and what the implications would be.
- The authors should reflect on the scope of the claims made, e.g., if the approach was only tested on a few datasets or with a few runs. In general, empirical results often depend on implicit assumptions, which should be articulated.
- The authors should reflect on the factors that influence the performance of the approach. For example, a facial recognition algorithm may perform poorly when image resolution is low or images are taken in low lighting. Or a speech-to-text system might not be used reliably to provide closed captions for online lectures because it fails to handle technical jargon.
- The authors should discuss the computational efficiency of the proposed algorithms and how they scale with dataset size.
- If applicable, the authors should discuss possible limitations of their approach to address problems of privacy and fairness.
- While the authors might fear that complete honesty about limitations might be used by reviewers as grounds for rejection, a worse outcome might be that reviewers discover limitations that aren't acknowledged in the paper. The authors should use their best judgment and recognize that individual actions in favor of transparency play an important role in developing norms that preserve the integrity of the community. Reviewers will be specifically instructed to not penalize honesty concerning limitations.

3. Theory assumptions and proofs

Question: For each theoretical result, does the paper provide the full set of assumptions and a complete (and correct) proof?

Answer: [\[Yes\]](#)

Justification: Section 3 and Supplementary Material A provide the theoretical formulation of the multi-agent decision process and Q-learning convergence analysis.

Guidelines:

- The answer NA means that the paper does not include theoretical results.
- All the theorems, formulas, and proofs in the paper should be numbered and cross-referenced.
- All assumptions should be clearly stated or referenced in the statement of any theorems.
- The proofs can either appear in the main paper or the supplemental material, but if they appear in the supplemental material, the authors are encouraged to provide a short proof sketch to provide intuition.
- Inversely, any informal proof provided in the core of the paper should be complemented by formal proofs provided in appendix or supplemental material.
- Theorems and Lemmas that the proof relies upon should be properly referenced.

4. Experimental result reproducibility

Question: Does the paper fully disclose all the information needed to reproduce the main experimental results of the paper to the extent that it affects the main claims and/or conclusions of the paper (regardless of whether the code and data are provided or not)?

Answer: [\[Yes\]](#)

Justification: Detailed experimental settings, including hyperparameters, action spaces, and reward function weights, are provided in Section 4 and Supplementary Materials (A.1, B.2, G, H).

Guidelines:

- The answer NA means that the paper does not include experiments.
- If the paper includes experiments, a No answer to this question will not be perceived well by the reviewers: Making the paper reproducible is important, regardless of whether the code and data are provided or not.
- If the contribution is a dataset and/or model, the authors should describe the steps taken to make their results reproducible or verifiable.
- Depending on the contribution, reproducibility can be accomplished in various ways. For example, if the contribution is a novel architecture, describing the architecture fully might suffice, or if the contribution is a specific model and empirical evaluation, it may be necessary to either make it possible for others to replicate the model with the same dataset, or provide access to the model. In general, releasing code and data is often one good way to accomplish this, but reproducibility can also be provided via detailed instructions for how to replicate the results, access to a hosted model (e.g., in the case of a large language model), releasing of a model checkpoint, or other means that are appropriate to the research performed.
- While NeurIPS does not require releasing code, the conference does require all submissions to provide some reasonable avenue for reproducibility, which may depend on the nature of the contribution. For example
 - (a) If the contribution is primarily a new algorithm, the paper should make it clear how to reproduce that algorithm.
 - (b) If the contribution is primarily a new model architecture, the paper should describe the architecture clearly and fully.
 - (c) If the contribution is a new model (e.g., a large language model), then there should either be a way to access this model for reproducing the results or a way to reproduce the model (e.g., with an open-source dataset or instructions for how to construct the dataset).
 - (d) We recognize that reproducibility may be tricky in some cases, in which case authors are welcome to describe the particular way they provide for reproducibility. In the case of closed-source models, it may be that access to the model is limited in some way (e.g., to registered users), but it should be possible for other researchers to have some path to reproducing or verifying the results.

5. Open access to data and code

Question: Does the paper provide open access to the data and code, with sufficient instructions to faithfully reproduce the main experimental results, as described in supplemental material?

Answer: [No]

Justification: The code is not publicly available via a link in the current submission to preserve anonymity, but detailed implementation descriptions are included in the Supplementary Material.

Guidelines:

- The answer NA means that paper does not include experiments requiring code.
- Please see the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- While we encourage the release of code and data, we understand that this might not be possible, so “No” is an acceptable answer. Papers cannot be rejected simply for not including code, unless this is central to the contribution (e.g., for a new open-source benchmark).
- The instructions should contain the exact command and environment needed to run to reproduce the results. See the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.

- The authors should provide instructions on data access and preparation, including how to access the raw data, preprocessed data, intermediate data, and generated data, etc.
- The authors should provide scripts to reproduce all experimental results for the new proposed method and baselines. If only a subset of experiments are reproducible, they should state which ones are omitted from the script and why.
- At submission time, to preserve anonymity, the authors should release anonymized versions (if applicable).
- Providing as much information as possible in supplemental material (appended to the paper) is recommended, but including URLs to data and code is permitted.

6. Experimental setting/details

Question: Does the paper specify all the training and test details (e.g., data splits, hyper-parameters, how they were chosen, type of optimizer, etc.) necessary to understand the results?

Answer: [\[Yes\]](#)

Justification: Section 4.1 and Supplementary Material G provide details on datasets, sampling strategies (e.g., random vs stratified), backbone models (ResNet-101), and optimization settings.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The experimental setting should be presented in the core of the paper to a level of detail that is necessary to appreciate the results and make sense of them.
- The full details can be provided either with the code, in appendix, or as supplemental material.

7. Experiment statistical significance

Question: Does the paper report error bars suitably and correctly defined or other appropriate information about the statistical significance of the experiments?

Answer: [\[Yes\]](#)

Justification: Table 12 in the Supplementary Material reports the mean and standard deviation of mAP, Rare-F1, and Loss Variance over three independent runs.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The authors should answer "Yes" if the results are accompanied by error bars, confidence intervals, or statistical significance tests, at least for the experiments that support the main claims of the paper.
- The factors of variability that the error bars are capturing should be clearly stated (for example, train/test split, initialization, random drawing of some parameter, or overall run with given experimental conditions).
- The method for calculating the error bars should be explained (closed form formula, call to a library function, bootstrap, etc.)
- The assumptions made should be given (e.g., Normally distributed errors).
- It should be clear whether the error bar is the standard deviation or the standard error of the mean.
- It is OK to report 1-sigma error bars, but one should state it. The authors should preferably report a 2-sigma error bar than state that they have a 96% CI, if the hypothesis of Normality of errors is not verified.
- For asymmetric distributions, the authors should be careful not to show in tables or figures symmetric error bars that would yield results that are out of range (e.g. negative error rates).
- If error bars are reported in tables or plots, The authors should explain in the text how they were calculated and reference the corresponding figures or tables in the text.

8. Experiments compute resources

Question: For each experiment, does the paper provide sufficient information on the computer resources (type of compute workers, memory, time of execution) needed to reproduce the experiments?

Answer: [Yes]

Justification: Supplementary Material I (Figure 8 and text) specifies the use of NVIDIA A100 GPUs and reports training time in hours.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The paper should indicate the type of compute workers CPU or GPU, internal cluster, or cloud provider, including relevant memory and storage.
- The paper should provide the amount of compute required for each of the individual experimental runs as well as estimate the total compute.
- The paper should disclose whether the full research project required more compute than the experiments reported in the paper (e.g., preliminary or failed experiments that didn't make it into the paper).

9. Code of ethics

Question: Does the research conducted in the paper conform, in every respect, with the NeurIPS Code of Ethics <https://neurips.cc/public/EthicsGuidelines>?

Answer: [Yes]

Justification: The research involves standard optimization for computer vision tasks using public datasets, with no apparent ethical violations.

Guidelines:

- The answer NA means that the authors have not reviewed the NeurIPS Code of Ethics.
- If the authors answer No, they should explain the special circumstances that require a deviation from the Code of Ethics.
- The authors should make sure to preserve anonymity (e.g., if there is a special consideration due to laws or regulations in their jurisdiction).

10. Broader impacts

Question: Does the paper discuss both potential positive societal impacts and negative societal impacts of the work performed?

Answer: [No]

Justification: The paper focuses on technical efficiency and performance improvements in multi-label classification and does not explicitly discuss broader societal impacts.

Guidelines:

- The answer NA means that there is no societal impact of the work performed.
- If the authors answer NA or No, they should explain why their work has no societal impact or why the paper does not address societal impact.
- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.
- The conference expects that many papers will be foundational research and not tied to particular applications, let alone deployments. However, if there is a direct path to any negative applications, the authors should point it out. For example, it is legitimate to point out that an improvement in the quality of generative models could be used to generate deepfakes for disinformation. On the other hand, it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster.
- The authors should consider possible harms that could arise when the technology is being used as intended and functioning correctly, harms that could arise when the technology is being used as intended but gives incorrect results, and harms following from (intentional or unintentional) misuse of the technology.

- If there are negative societal impacts, the authors could also discuss possible mitigation strategies (e.g., gated release of models, providing defenses in addition to attacks, mechanisms for monitoring misuse, mechanisms to monitor how a system learns from feedback over time, improving the efficiency and accessibility of ML).

11. Safeguards

Question: Does the paper describe safeguards that have been put in place for responsible release of data or models that have a high risk for misuse (e.g., pretrained language models, image generators, or scraped datasets)?

Answer: [NA]

Justification: The paper does not release new models or datasets that pose high risks for misuse.

Guidelines:

- The answer NA means that the paper poses no such risks.
- Released models that have a high risk for misuse or dual-use should be released with necessary safeguards to allow for controlled use of the model, for example by requiring that users adhere to usage guidelines or restrictions to access the model or implementing safety filters.
- Datasets that have been scraped from the Internet could pose safety risks. The authors should describe how they avoided releasing unsafe images.
- We recognize that providing effective safeguards is challenging, and many papers do not require this, but we encourage authors to take this into account and make a best faith effort.

12. Licenses for existing assets

Question: Are the creators or original owners of assets (e.g., code, data, models), used in the paper, properly credited and are the license and terms of use explicitly mentioned and properly respected?

Answer: [Yes]

Justification: The paper uses and cites standard datasets including Pascal VOC, MS-COCO, and Visual Genome (References [62], [63], [64]).

Guidelines:

- The answer NA means that the paper does not use existing assets.
- The authors should cite the original paper that produced the code package or dataset.
- The authors should state which version of the asset is used and, if possible, include a URL.
- The name of the license (e.g., CC-BY 4.0) should be included for each asset.
- For scraped data from a particular source (e.g., website), the copyright and terms of service of that source should be provided.
- If assets are released, the license, copyright information, and terms of use in the package should be provided. For popular datasets, paperswithcode.com/datasets has curated licenses for some datasets. Their licensing guide can help determine the license of a dataset.
- For existing datasets that are re-packaged, both the original license and the license of the derived asset (if it has changed) should be provided.
- If this information is not available online, the authors are encouraged to reach out to the asset's creators.

13. New assets

Question: Are new assets introduced in the paper well documented and is the documentation provided alongside the assets?

Answer: [NA]

Justification: No new datasets or assets are introduced in this paper.

Guidelines:

- The answer NA means that the paper does not release new assets.
- Researchers should communicate the details of the dataset/code/model as part of their submissions via structured templates. This includes details about training, license, limitations, etc.
- The paper should discuss whether and how consent was obtained from people whose asset is used.
- At submission time, remember to anonymize your assets (if applicable). You can either create an anonymized URL or include an anonymized zip file.

14. Crowdsourcing and research with human subjects

Question: For crowdsourcing experiments and research with human subjects, does the paper include the full text of instructions given to participants and screenshots, if applicable, as well as details about compensation (if any)?

Answer: [NA]

Justification: The research does not involve crowdsourcing or human subjects.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Including this information in the supplemental material is fine, but if the main contribution of the paper involves human subjects, then as much detail as possible should be included in the main paper.
- According to the NeurIPS Code of Ethics, workers involved in data collection, curation, or other labor should be paid at least the minimum wage in the country of the data collector.

15. Institutional review board (IRB) approvals or equivalent for research with human subjects

Question: Does the paper describe potential risks incurred by study participants, whether such risks were disclosed to the subjects, and whether Institutional Review Board (IRB) approvals (or an equivalent approval/review based on the requirements of your country or institution) were obtained?

Answer: [NA]

Justification: The research does not involve human subjects.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Depending on the country in which research is conducted, IRB approval (or equivalent) may be required for any human subjects research. If you obtained IRB approval, you should clearly state this in the paper.
- We recognize that the procedures for this may vary significantly between institutions and locations, and we expect authors to adhere to the NeurIPS Code of Ethics and the guidelines for their institution.
- For initial submissions, do not include any information that would break anonymity (if applicable), such as the institution conducting the review.

16. Declaration of LLM usage

Question: Does the paper describe the usage of LLMs if it is an important, original, or non-standard component of the core methods in this research? Note that if the LLM is used only for writing, editing, or formatting purposes and does not impact the core methodology, scientific rigor, or originality of the research, declaration is not required.

Answer: [NA]

Justification: The core methodology relies on reinforcement learning agents and CNN/Transformer backbones for image classification, not on LLMs.

Guidelines:

- The answer NA means that the core method development in this research does not involve LLMs as any important, original, or non-standard components.
- Please refer to our LLM policy (<https://neurips.cc/Conferences/2025/LLM>) for what should or should not be described.