

Stealthy Multi-task Adversarial Attacks

Jiacheng Guo^{1,2*}, Tianyun Zhang^{1*†}, Lei Li¹, Haochen Yang¹, Hongkai Yu¹, and Minghai Qin^{1,3†}

¹ Cleveland State University, Cleveland, OH 44115, USA
{j.guo58,l.li15,h.yang15}@vikes.csuohio.edu,
{t.zhang85,h.yu19}@csuohio.edu

² University of Wisconsin-Madison, Madison, WI 53706, USA

³ Western Digital Research, San Jose, CA 95119, USA
{minghai.qin}@wdc.com

*Equal Contribution †Corresponding Authors

Abstract. Deep neural networks are highly vulnerable to adversarial perturbations, raising serious safety concerns in the real-world systems. While prior work mainly explores single-task attacks or jointly degrading all tasks in multi-task models, practical scenarios often demand more selective and stealthy attack strategies. To address this challenge, we propose **Stealthy Multi-Task Adversarial Attack (SMTA²)**, a novel framework that selectively degrades a targeted task while strictly preserving the performance of non-targeted tasks. We formulate this objective as a constrained multi-objective optimization problem and design task-aware adversarial perturbations that maximize degradation on the targeted task without causing collateral damage on non-targeted tasks. To enhance practicality, we further introduce an automated loss-weight tuning strategy that dynamically balances attack and preservation objectives. Experiments on two multi-task benchmarks NYUv2 and Cityscapes demonstrate that SMTA² achieves strong attack performance on targeted tasks while maintaining non-targeted tasks intact on both undefended and adversarially trained models, establishing the first systematic framework for stealthy and selective multi-task attack framework.

Keywords: multi-task learning, adversarial attacks, stealthy attacks, deep learning

1 Introduction

Deep Neural Networks (DNNs) have achieved remarkable success across a wide range of tasks [1, 10, 19, 23], yet remain highly vulnerable to adversarial attacks [20, 29]. These attacks introduce subtle, often imperceptible perturbations that can drastically alter model predictions [2, 13, 32]. Despite their minimal visual impact, such perturbations can cause severe performance degradation, posing significant safety risks in critical applications, such as autonomous driving [5, 38]. Figure 1 illustrates typical imperceptible adversarial perturbations under different magnitudes, where even small perturbations (e.g., $\epsilon=2/255$ or $4/255$) can effectively mislead DNNs.

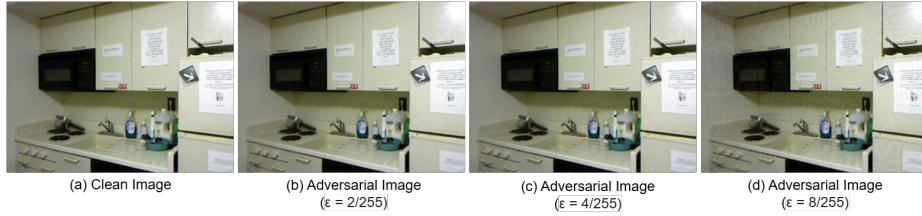


Fig. 1: Demonstration of a clean image (a) and adversarial images (b) and (c) from NYUv2 dataset.

Beyond single-task models, Multi-Task Learning (MTL) enables a shared network to jointly learn multiple related tasks, significantly improving model efficiency and generalization [4, 8, 33]. However, the shared representations also introduce coupled vulnerabilities: a perturbation crafted for one task can *unintentionally propagate to others* [12, 41]. Specifically, in safety-critical systems, attackers rarely aim for indiscriminate performance collapse. Instead, they may selectively target high-value tasks while preserving auxiliary functionalities to evade detection [35, 43]. For instance, manipulating traffic sign recognition while preserving lane detection may create subtle yet dangerous misbehavior in autonomous driving [28]. Similarly, degrading pedestrian segmentation while keeping depth estimation intact can cause catastrophic consequences without triggering obvious system-wide failure [27]. These scenarios highlight the need for *selective* and *stealthy* multi-task adversarial attacks that compromise targeted tasks while preserving the performance of others.

However, achieving *selective degradation* is fundamentally challenging [41, 43]. Multi-task objectives are inherently coupled through shared parameters, and increasing the loss of one task typically propagates gradients that harm others. Enforcing strict preservation constraints while maximizing degradation on a targeted task leads to a highly constrained and non-trivial optimization problem. Naively reweighting losses often fails to maintain non-targeted task performance, especially under strong adversarial perturbations.

To address these challenges, we propose a novel concept termed **Stealthy Multi-Task Adversarial Attack (SMTA²)**. Given an MTL framework with multiple tasks, SMTA² aims to significantly degrade a **targeted task** while strictly preserving or even improving the performance of other tasks which are not expected to be attacked, regarded as **non-targeted tasks**. By crafting task-aware adversarial perturbations through a weighted loss formulation, SMTA² enables fine-grained control over task-specific attack behaviors. Furthermore, we develop an automated weight searching approach to dynamically search for optimal task weights, enabling efficient and adaptive attack generation without manual tuning. This strategy significantly reduces computational overhead and improves generalization across diverse inputs. Compared to random system noise which typically causes uniform and mild degradation across all tasks, SMTA²

introduces asymmetric and task-selective degradation, a property that is empirically validated in Section 4.

To the best of our knowledge, this is the first work that systematically formulates and solves the problem of stealthy and task-selective adversarial attacks in MTL, together with a principled evaluation criterion for preservation-aware attack assessment. The main contributions are summarized as follows:

- We propose “Stealthy Multi-Task Adversarial Attacks” (**SMTA**²), a novel attack paradigm which enables selective degradation of a targeted task while strictly preserving the original performance of non-targeted tasks in MTL.
- The “Stealthy Multi-Task Adversarial Attacks” is formulated as a constrained multi-objective optimization problem and designed as a task-aware objective with explicit preservation constraints.
- We develop an automated loss-weight optimization strategy that dynamically searches for optimal task weights, improving attack efficiency and stability over manual tuning.
- Experiments on NYUv2 and Cityscapes datasets in different L_p norm attacks demonstrate strong targeted attack effectiveness while preserving non-targeted tasks on both undefended and adversarially trained models.

2 Related Work

2.1 Multi-task Learning

Multi-task learning (MTL) is a machine learning paradigm which enables the simultaneous training of multiple related tasks, providing advantages such as improved data efficiency, reduced overfitting, and enhanced generalization capabilities [4, 8]. By leveraging shared representations among tasks, MTL has become increasingly popular across diverse fields, including natural language processing and computer vision [21, 39, 42]. Recent advancements have categorized MTL techniques into several domains, such as regularization strategies, task relationship modeling, feature propagation, optimization frameworks, and pre-training methods [14, 16, 19].

MTL approaches are broadly classified into joint training and multi-step training methods, depending on the nature of task relationships and the level of interaction between tasks [24, 42], which has evolved significantly with the emergence of deep learning and pretrained foundation models. These advancements have introduced innovations such as task-promotable, task-agnostic training, and capabilities for zero-shot and few-shot learning, making MTL applicable in scenarios with limited task-specific labeled data [18, 39]. As MTL continues to mature, emerging research focuses on improving task balancing, addressing task interference, and enhancing robustness for its up-to-date application in complex domains [24, 25].

2.2 Adversarial Attacks

Adversarial attacks pose a fundamental challenge to DNNs, where carefully crafted perturbations can induce incorrect predictions despite being largely imperceptible to humans [3, 29, 32, 35, 37]. Such perturbations expose critical vulnerabilities in modern architectures.

Gradient-based optimization remains a dominant paradigm for attack generation. Universal adversarial perturbations [31] demonstrated that input-agnostic, gradient-derived perturbations can reliably fool classifiers, laying the foundation for subsequent gradient-driven methods and extensions beyond classification. These approaches enabled adversarial attacks to be studied in structured and real-world settings, including aerial object detection [11, 36], semantically stealthy segmentation [6], and audio tasks [21] where perturbations preserve visual plausibility while inducing substantial degradation. In MTL, Guo et al. [17] introduced a Multi-Task Adversarial Attack (MTA) framework that exploits shared representations to craft joint attacks, outperforming single-task baselines. Zhang et al. [41] later standardized MTA methodologies and broadened evaluations. Nevertheless, unlike those non-stealthy architectures [17, 41], our objective is stricter: selectively attacking a targeted task and strictly preserving non-targeted tasks.

2.3 Adversarial Robustness in MTL

To improve adversarial robustness, Madry et al. [29] formulated defense as a min-max optimization problem, while Mao et al. [30] introduced natural supervision to recover correct predictions under attack.

Adversarial vulnerabilities in MTL have also been investigated. Zhang et al. [40] analyzed the robustness-accuracy trade-off theoretically. Zhang et al. [41] proposed Gradient Balancing Multi-Task Attack to jointly optimize attacks across tasks, improving task balancing through shared representations. Guo et al. [15] incorporated min-max optimization to enhance robustness in multi-task adversarial learning. These methods mainly focus on jointly attacking all tasks through improved optimization. Moreover, Zhe et al. [43] predefined “hidden tasks” which can be stealthily compromised via shared representations without explicitly modeling task relationships, but they used different task definitions and settings. Comparatively, we propose a unique SMTA² framework which attacks multi-task models without changing architectures or task relationships.

Overall, prior works either design general optimization-based attacks or indiscriminately degrade all tasks, leaving selective and stealthy mechanisms underexplored. Crucially, the structural dependency induced by shared representations has not been systematically exploited for controlled adversarial manipulation. We address this gap by proposing a principled framework for selective and stealthy adversarial attacks in multi-task scenarios that explicitly models inter-task relationships, enabling fine-grained degradation while preserving all non-targeted tasks.

3 Methodology

3.1 Problem Formulation

Consider a multi-task learning setting where each input sample x is associated with multiple task-specific ground-truth labels $y = (y_1, y_2, \dots, y_m)$. Let $L_i(x, y_i)$ denote the loss corresponding to the i -th task in a pretrained multi-task model with shared representations.

Definition. The **Stealthy Multi-Task Adversarial Attack (SMTA²)** is a kind of adversarial attack which aims to degrade the performance of a designated *targeted task* while strictly preserving the performance of all remaining *non-targeted tasks*.

Let i_t be the targeted task. The SMTA² can be formulated as

$$\begin{aligned} & \underset{\delta}{\text{maximize}} && L_{i_t}(x + \delta, y_{i_t}), \\ & \text{subject to} && L_i(x + \delta, y_i) \leq L_i(x, y_i) \text{ for } i \neq i_t, \\ & && \|\delta\|_p \leq \epsilon, x + \delta \in \mathcal{B}, \end{aligned} \tag{1}$$

where δ is the adversarial noise, ϵ is a value to constrain the strength of the adversarial noise, and $\mathcal{B}$ is the box constraint to ensure that a valid image is derived after adding the adversarial noise.

The attack without the first constraint, $L_i(x + \delta, y_i) \leq L_i(x, y_i)$ for $i \neq i_t$, is to be called a **non-stealthy** multi-task attack.

Optimization Challenge. Problem (1) is substantially more challenging than standard adversarial attacks. In multi-task DNNs, tasks share a common backbone representation, leading to strong inter-task coupling. Perturbations that increase the loss of one task naturally propagate to others. Therefore, the feasible region satisfying all preservation constraints is often narrow under strong perturbations, making direct optimization intractable.

To obtain a tractable surrogate, we relax problem (1) to the following weighted objective:

$$\underset{\|\delta\|_p \leq \epsilon, x + \delta \in \mathcal{B}}{\text{maximize}} \quad L(x + \delta, y) = \sum_{i=1}^m w_i L_i(x + \delta, y_i), \tag{2}$$

where w_i is the weighting factor corresponding to the i -th task. If $i = i_t$, w_i is positive; otherwise w_i is negative. The constraint in Eq. (1) serves as an effective optimization proxy so that an appropriate combination of weighting factors corresponding to different tasks could be searched each time for metric preservation when solving problem (2). It is guaranteed that an optimal solution feasible to problem (1) could be found out by adjusting the weighting factors. Specifically, when the loss of a non-targeted task increases, its corresponding negative weight effectively penalizes further degradation. Crucially, success hinges on carefully adapting $\{w_i\}$ to balance attack strength and strict preservation, which is an inherently non-trivial task due to inter-task gradient interference.

Unlike conventional adversarial attacks which optimize only a single objective, SMTA² must simultaneously maximize one task loss while enforcing multiple preservation constraints. Due to the shared representations and gradient interference across tasks, satisfying these constraints while maintaining strong attack effectiveness is fundamentally challenging. In this case, it is crucial to find an optimal combination to balance different tasks. Therefore, effective weight adaptation strategies are essential for navigating this coupled optimization landscape. The details of how to satisfy the constraints in the stealthy attack is discussed in the subsection 3.2.

3.2 The Stealthy Multi-Task Adversarial Attack Framework

In a general multi-task adversarial attack framework, attacking one task could negatively affect all other tasks facing severe degradation as well. As a result, the performance of non-targeted tasks could also be significantly affected. To maintain the performance of non-targeted tasks, a positive weighting factor is set for the targeted task and negative weighting factors for the non-targeted tasks when solving problem (2). The positive weighting factor can degrade the performance of the targeted task, while the negative ones could compensate for the performance degradation of non-targeted tasks caused by attacking the targeted task. If the magnitudes of negative weighting factors are too small, the performance of non-targeted tasks will still be degraded. In this case, the first constraint in problem (1) is not satisfied. If the negative weights are too large, the performance of non-targeted tasks may even surpass the clean baseline, satisfying the preservation constraint but weakening the attack on the targeted task. and the solution of problem (2) will satisfy the constraints of problem (1). However, the strength of attacking the targeted task will be compromised.

In SMTA², the core challenge lies in enforcing strict performance constraints on non-targeted tasks while maximizing the degradation of the targeted task. This naturally leads to a constrained multi-objective optimization problem, where the loss of the targeted task is maximized subject to preservation constraints on the remaining tasks. To address this challenge, we design two strategies for searching appropriate task weights in the joint loss function: a *manual* search strategy and a more efficient *automated* optimization strategy.

Manual Weight Searching. In the manual approach, the weight of the targeted task is fixed to emphasize its degradation, while the weights of non-targeted tasks are manually adjusted to enforce preservation constraints. The procedure begins by assigning small-magnitude weights to non-targeted tasks. After performing the adversarial attack, if the performance of any non-targeted task violates the preservation constraint (i.e., performance drops below the baseline), its corresponding weight magnitude is increased in the next trial. This iterative process continues until all non-targeted tasks satisfy the constraint.

Although this approach guarantees the existence of a feasible solution to Problem (1), it relies heavily on trial-and-error hyperparameter tuning. Partic-

ularly, as the number of tasks increases, the search space grows rapidly, making manual tuning computationally expensive and time-consuming.

Algorithm 1 An automated approach for searching the weighting factors in the loss function

Input: input data x , the ground-truth label of the i -th task y_i , the initialized weighting factor w_i of the i -th task, model parameters and the corresponding loss $L(x, y) = \sum_{i=1}^m w_i L_i(x, y_i)$, total steps S in the adversarial attack, the targeted task i_t , the step size λ to adjust weighting factors, and the attenuation coefficient α .
Initialize $\lambda_i = \lambda$ for all $i \neq i_t$.
for $s = 1, 2, \dots, S$ **do**
 perform one step PGD (or APGD or IFGSM) attack on loss $L(x, y)$ to generate the updated adversarial noise δ .
 for $i = 1, 2, \dots, m$ **do**
 if $i \neq i_t$ **then**
 $w_i = w_i - \lambda_i * (L_i(x + \delta, y_i) - L_i(x, y_i))$.
 if $L_i(x + \delta, y_i) \leq L_i(x, y_i)$ **then**
 $\lambda_i = \lambda_i * \alpha$
 end if
 end if
 end for
end for

Automated Weight Optimization. To overcome the inefficiency of manual tuning, we propose an automated weight optimization strategy, summarized in Algorithm 1. Unlike the manual approach, the automated weight searching approach dynamically updates task weights during each attack iteration, enabling real-time enforcement of preservation constraints. Specifically, at each attack step, we compute the loss difference between the adversarial and original inputs for each non-targeted task. The weight of each non-targeted task is updated proportionally to the loss difference. If adversarial loss increases (indicating potential degradation), the corresponding weight is adjusted to counteract this effect.

To accelerate convergence, we introduce a task-specific step size λ_i , initialized uniformly. Once a non-targeted task satisfies the preservation constraint (i.e., its adversarial loss no longer exceeds the original loss), its step size is attenuated by a coefficient $\alpha \in (0, 1)$, reducing oscillations and stabilizing the search process. This adaptive attenuation mechanism allows the algorithm to quickly converge to a feasible region while maintaining stability.

Complexity Analysis. Let m denote the number of tasks and S denote the total attack iterations. Each step of the attack requires one forward and backward pass to update the perturbation, incurring the same dominant cost attacks based on Projected Gradient Descent (PGD) [29], Auto-PGD [9] or Iterative Fast Gradient Sign Method (IFGSM) [22]. The additional overhead of SMTA² comes from

updating task weights, which involves computing loss differences for $(m-1)$ non-targeted tasks and simple scalar updates. This adds only $\mathcal{O}(m)$ operations per iteration, which is negligible compared to the backpropagation cost. Therefore, the overall time complexity remains $\mathcal{O}(S \cdot \text{Backprop})$, making SMTA² asymptotically equivalent to conventional attacks in complexity. Compared with manual weight searching strategy, the automated approach substantially improves efficiency and scalability, especially when the number of tasks increases.

Overall, the automated approach performs online constrained optimization over task weights, significantly reducing parameter tuning time and improving scalability to multi-task settings with larger task counts. Empirically, it consistently finds feasible and stable solutions that satisfy stealthy attack requirements while preserving the performance among the non-targeted tasks.

4 Experiments

To evaluate SMTA², we consider two schemes for searching optimal task weights: the *manual* scheme and the *automated* scheme. We adopt three widely-used attack algorithms: Projected Gradient Descent (PGD) [29], Iterative Fast Gradient Sign Method (IFGSM) [22], and a recently established stronger attack: Auto-PGD (APGD) [9]. For the attack norms, the PGD attack is implemented under both L_2 and L_∞ norms, the APGD is implemented under both L_1 and L_2 norms, and the IFGSM is under L_∞ norm. We compare manual and automated weight-search scheme under the proposed loss formulation, respectively. Due to space limitation, the results of IFGSM and more visualization results are provided in the supplementary material.

4.1 Experimental Setup

We evaluate SMTA² on two widely used multi-task image benchmarks, i.e., NYUv2 [34] and Cityscapes [7]. Both non-stealthy and stealthy attack settings are considered. Non-stealthy attacks degrade the targeted task and inevitably affects other tasks, causing global performance drops [17, 41]. However, SMTA² aims to significantly impair the targeted task while strictly maintaining non-targeted task performance. We further incorporate adversarial training to assess robustness against defended models. All the images are resized to 288×384 for NYUv2 and 256×512 for Cityscapes. All the experiments are conducted on 8 NVIDIA RTX A6000 GPUs.

4.2 Experimental Procedure

The proposed SMTA² aims to selectively degrade a targeted task while strictly preserving non-targeted tasks. Accordingly, we adopt an evaluation criterion that measures attack effectiveness under full preservation constraints.

The attack strength is controlled via perturbations according to different L_p norms. Specifically, for the L_∞ norm in PGD and IFGSM, ϵ is set to $2/255$, $4/255$,

Table 1: Experimental results of non-stealthy and SMTA² framework by manual and automated solutions of PGD L_∞ attack algorithm on NYUv2 dataset. Non-attacked performance is represented by $\epsilon=0$. The performances of the targeted task are underlined. Values affected by the attack (i.e., degraded during the attack) are marked in blue, while those not influenced by the attack are highlighted in red. ($\uparrow$) means higher is better and ($\downarrow$) means lower is better.

Method	Non-stealthy			SMTA ² Manual			SMTA ² Auto		
Tasks	Segment [mIoU($\uparrow$)]	Depth [aErr($\downarrow$)]	Normal [mDist($\downarrow$)]	Segment [mIoU($\uparrow$)]	Depth [aErr($\downarrow$)]	Normal [mDist($\downarrow$)]	Segment [mIoU($\uparrow$)]	Depth [aErr($\downarrow$)]	Normal [mDist($\downarrow$)]
$\epsilon=0$	46.56	40.57	23.41	46.56	40.57	23.41	46.56	40.57	23.41
$\epsilon=2/255$	<u>26.21</u>	51.78	<u>26.65</u>	<u>26.54</u>	37.92	<u>23.00</u>	<u>26.46</u>	40.28	<u>23.12</u>
	35.08	<u>106.93</u>	28.45	47.53	<u>106.47</u>	23.19	47.63	<u>106.45</u>	23.40
	37.76	54.99	<u>40.36</u>	48.17	38.23	<u>41.10</u>	46.76	34.76	<u>39.99</u>
$\epsilon=4/255$	<u>17.04</u>	63.80	30.97	<u>18.30</u>	37.30	<u>23.22</u>	<u>18.33</u>	37.45	<u>23.31</u>
	24.15	<u>167.40</u>	34.39	47.08	<u>165.27</u>	22.83	47.67	<u>168.69</u>	23.36
	28.88	71.44	<u>53.88</u>	47.86	37.26	<u>53.47</u>	46.98	39.58	<u>53.87</u>
$\epsilon=8/255$	<u>10.11</u>	78.36	36.23	<u>12.96</u>	37.86	22.60	<u>12.32</u>	39.56	<u>23.21</u>
	14.58	<u>240.52</u>	41.16	47.14	<u>229.64</u>	21.96	52.83	<u>225.11</u>	22.36
	18.91	93.36	<u>68.05</u>	47.66	39.93	<u>66.60</u>	46.93	38.03	<u>63.29</u>

and 8/255, respectively. The perturbation ϵ is set to 10 and 20 for L_1 norm in APGD and is set to 5 and 10 for L_2 norm in APGD and PGD. The manual strategy searches for feasible task weights, where negative weights for non-targeted tasks help enforce preservation. The automated strategy (Algorithm 1) could dynamically update task weights for efficient and stable convergence. Furthermore, we evaluate SMTA² on Adversarially Trained (AT) models. Results show that both strategies maintain non-targeted task performance while achieving strong targeted attacks, even under enhanced robustness.

Evaluation Metrics. On NYUv2, we evaluate semantic segmentation, depth estimation, and surface normal prediction tasks using mIoU, absolute error (aErr.), and mean angular distance (mDist.), respectively. On Cityscapes, mIoU is used for semantic segmentation and part segmentation tasks, and aErr. is used for disparity estimation. All the settings regarding metrics follow [26]. Each task is treated as the targeted task in turn, with others considered as non-targeted.

4.3 Main Results and Analysis

Attack Results of the Undefended Model Results on NYUv2 under PGD L_∞ attacks (Table 1) show that non-stealthy attacks degrade all tasks, whereas both manual and automated SMTA² selectively degrade the targeted task while preserving non-targeted tasks at baseline or better performance. Even at small perturbation levels ($\epsilon = 2/255$ or $4/255$), SMTA² achieves attack strength comparable to non-stealthy methods. As ϵ increases, the performance gap remains limited, indicating stable and robust selective attack behavior.

Table 2: Results of non-stealthy attacks and SMTA² on **Cityscapes** dataset under **PGD** L_∞ attack. The targeted-task results are underlined. **Blue** and **red** indicate affected and preserved values, respectively. ($\uparrow$)/($\downarrow$) denote higher/lower is better.

Method	Non-stealthy			SMTA ² Manual			SMTA ² Auto		
Tasks	Segment [mIoU($\uparrow$)]	Part Seg [mIoU($\uparrow$)]	Disparity [aErr($\downarrow$)]	Segment [mIoU($\uparrow$)]	Part Seg [mIoU($\uparrow$)]	Disparity [aErr($\downarrow$)]	Segment [mIoU($\uparrow$)]	Part Seg [mIoU($\uparrow$)]	Disparity [aErr($\downarrow$)]
$\epsilon=0$	54.20	51.82	81.51	54.20	51.82	81.51	54.20	51.82	81.51
$\epsilon=2/255$	<u>27.98</u>	<u>43.82</u>	<u>98.72</u>	<u>25.34</u>	<u>52.51</u>	<u>79.71</u>	<u>25.38</u>	<u>52.97</u>	<u>80.30</u>
	41.24	<u>30.00</u>	92.13	<u>54.57</u>	<u>27.29</u>	<u>80.04</u>	<u>54.25</u>	<u>27.22</u>	<u>76.25</u>
	<u>38.47</u>	42.05	<u>278.11</u>	<u>54.39</u>	<u>51.93</u>	<u>367.34</u>	<u>56.16</u>	<u>52.30</u>	<u>359.65</u>
$\epsilon=4/255$	<u>20.48</u>	35.01	119.67	<u>17.25</u>	<u>52.67</u>	<u>81.00</u>	<u>17.08</u>	52.77	<u>81.21</u>
	33.98	<u>21.69</u>	111.19	<u>54.59</u>	<u>18.02</u>	<u>80.23</u>	<u>54.36</u>	<u>18.19</u>	80.00
	<u>25.82</u>	26.68	<u>490.15</u>	<u>55.22</u>	<u>51.85</u>	<u>679.89</u>	<u>56.44</u>	<u>52.02</u>	<u>644.12</u>
$\epsilon=8/255$	<u>11.31</u>	19.24	164.19	<u>10.62</u>	<u>51.84</u>	<u>78.72</u>	<u>10.02</u>	52.07	81.50
	<u>25.77</u>	<u>14.65</u>	<u>153.10</u>	<u>55.13</u>	<u>12.63</u>	<u>79.12</u>	<u>54.27</u>	<u>12.22</u>	<u>80.97</u>
	16.05	14.18	<u>824.75</u>	<u>55.10</u>	<u>52.54</u>	<u>1014.74</u>	<u>57.36</u>	<u>52.10</u>	<u>1008.44</u>

Table 3: Results of non-stealthy attacks and SMTA² on **NYUv2** dataset under **PGD** L_2 attack. The targeted-task results are underlined. **Blue** and **red** indicate affected and preserved values, respectively. ($\uparrow$)/($\downarrow$) denote higher/lower is better.

Method	Non-stealthy			SMTA ² Manual			SMTA ² Auto		
Tasks	Segment [mIoU($\uparrow$)]	Depth [aErr($\downarrow$)]	Normal [mDist($\downarrow$)]	Segment [mIoU($\uparrow$)]	Depth [aErr($\downarrow$)]	Normal [mDist($\downarrow$)]	Segment [mIoU($\uparrow$)]	Depth [aErr($\downarrow$)]	Normal [mDist($\downarrow$)]
$\epsilon=0$	46.56	40.57	23.41	46.56	40.57	23.41	46.56	40.57	23.41
$\epsilon=5$	<u>14.70</u>	<u>70.56</u>	<u>33.29</u>	<u>19.30</u>	<u>40.18</u>	<u>23.25</u>	<u>19.44</u>	<u>37.65</u>	<u>23.35</u>
	21.04	<u>197.61</u>	36.48	<u>46.71</u>	<u>174.81</u>	<u>23.26</u>	<u>46.74</u>	<u>174.57</u>	<u>22.88</u>
	<u>24.29</u>	81.77	<u>63.31</u>	<u>50.21</u>	<u>40.01</u>	<u>71.54</u>	<u>47.15</u>	<u>39.37</u>	<u>54.28</u>
$\epsilon=10$	<u>7.70</u>	87.83	39.33	<u>13.03</u>	<u>39.78</u>	<u>23.19</u>	<u>12.84</u>	<u>38.05</u>	<u>23.22</u>
	<u>11.20</u>	<u>290.69</u>	<u>44.18</u>	<u>46.79</u>	<u>251.49</u>	<u>23.30</u>	<u>46.68</u>	<u>251.15</u>	<u>22.56</u>
	14.18	107.80	<u>81.43</u>	<u>47.79</u>	<u>39.22</u>	<u>71.54</u>	<u>49.71</u>	<u>38.16</u>	<u>68.21</u>

The results using Cityscapes dataset (Table 2) exhibit consistent trends. SMTA² maintains non-targeted task performance while achieving strong targeted degradation, often matching or surpassing non-stealthy baselines. The automated weight search performs comparably to manual tuning with significantly higher efficiency, demonstrating scalability and practicality.

Figures 2 and 3 visualize the results of stealthy adversarial attacks on the NYUv2 and Cityscapes datasets with $\epsilon=4/255$, respectively. The perturbations introduce only subtle pixel-level changes, yet they significantly degrade the performance of the targeted task while effectively preserving the outputs of non-targeted tasks. These examples demonstrate that SMTA² successfully achieves selective task degradation without introducing noticeable visual artifacts, particularly for safety-critical objects such as vehicles and pedestrians.

Table 4: Results of non-stealthy attacks and SMTA² on **Cityscapes** dataset under PGD L_2 attack. The targeted-task results are underlined. Blue and red indicate affected and preserved values, respectively. ($\uparrow$)/($\downarrow$) denote higher/lower is better.

Method	Non-stealthy			SMTA ² Manual			SMTA ² Auto		
Tasks	Segment	Part Seg	Disparity	Segment	Part Seg	Disparity	Segment	Part Seg	Disparity
	[mIoU($\uparrow$)]	[mIoU($\uparrow$)]	[aErr($\downarrow$)]	[mIoU($\uparrow$)]	[mIoU($\uparrow$)]	[aErr($\downarrow$)]	[mIoU($\uparrow$)]	[mIoU($\uparrow$)]	[aErr($\downarrow$)]
$\epsilon=0$	54.20	51.82	81.51	54.20	51.82	81.51	54.20	51.82	81.51
$\epsilon=5$	<u>18.84</u>	<u>33.17</u>	<u>126.93</u>	<u>21.48</u>	<u>53.70</u>	<u>80.36</u>	<u>21.20</u>	<u>52.36</u>	<u>81.44</u>
	31.08	<u>16.60</u>	144.51	<u>55.72</u>	<u>19.56</u>	<u>80.21</u>	<u>56.49</u>	<u>19.71</u>	<u>81.32</u>
	22.03	17.73	<u>921.67</u>	56.39	53.45	<u>652.57</u>	57.69	53.47	<u>607.56</u>
$\epsilon=10$	<u>10.52</u>	17.98	167.64	<u>14.85</u>	<u>55.68</u>	<u>81.43</u>	<u>14.25</u>	<u>52.99</u>	<u>80.35</u>
	22.95	<u>9.86</u>	274.18	<u>55.58</u>	<u>13.70</u>	<u>81.18</u>	<u>56.11</u>	<u>13.49</u>	<u>81.39</u>
	12.93	8.33	<u>1611.60</u>	<u>55.24</u>	<u>52.00</u>	<u>1057.77</u>	<u>57.53</u>	<u>51.99</u>	<u>1026.31</u>

Table 5: Experimental results of SMTA² Auto method with APGD attacks with L_1 & L_2 norms on Cityscapes. The targeted results are underlined. Blue and red indicate affected and preserved values, respectively. ($\uparrow$)/($\downarrow$) denote higher/lower is better.

APGD L_1	Seg($\uparrow$)	Part($\uparrow$)	Disp($\downarrow$)	APGD L_2	Seg($\uparrow$)	Part($\uparrow$)	Disp($\downarrow$)
$\epsilon=0$	54.20	51.82	81.51	$\epsilon=0$	54.20	51.82	81.51
$\epsilon=10$	<u>18.84</u>	<u>51.88</u>	<u>78.60</u>	$\epsilon=5$	<u>23.21</u>	<u>51.98</u>	<u>80.76</u>
	<u>55.69</u>	<u>19.90</u>	<u>80.16</u>		<u>54.44</u>	<u>24.08</u>	<u>81.26</u>
	<u>55.85</u>	<u>52.48</u>	<u>461.82</u>		<u>54.87</u>	<u>52.71</u>	<u>387.33</u>
$\epsilon=20$	<u>13.55</u>	<u>52.84</u>	<u>80.64</u>	$\epsilon=10$	<u>15.03</u>	<u>52.53</u>	<u>81.10</u>
	<u>55.28</u>	<u>13.58</u>	<u>80.23</u>		<u>56.20</u>	<u>14.12</u>	<u>80.36</u>
	<u>56.74</u>	<u>53.04</u>	<u>827.69</u>		<u>55.99</u>	<u>52.08</u>	<u>690.18</u>

The SMTA² framework is further evaluated under PGD- L_2 and stronger APGD attacks with both L_1 and L_2 norms. As shown in Tables 3, 4, and 5, SMTA² consistently achieves effective task-selective attacks across different perturbation budgets and attack algorithms. The targeted task is significantly degraded, while non-targeted tasks remain hold clean performance. Moreover, the automated searching generally matches or outperforms manual tuning. These results demonstrate that SMTA² generalizes well across different attack paradigms and norm constraints, maintaining strong attack effectiveness and stealthy task preservation under multiple attack algorithms.

Visualized Results of Attack Effectiveness Figure 4 visualizes the attack effectiveness of non-stealthy and automated SMTA² methods under PGD L_∞ attacks using radar charts, which depict the ratio between post-attack and pre-attack performance. For metrics where higher values indicate better performance (e.g., mIoU), the ratio is computed as post-attack divided by pre-attack performance, while for metrics where lower values are preferable (e.g., aErr), the inverse ratio is used. Smaller ratios (closer to the center) indicate stronger attack effectiveness, whereas larger ratios (closer to the boundary) correspond to better

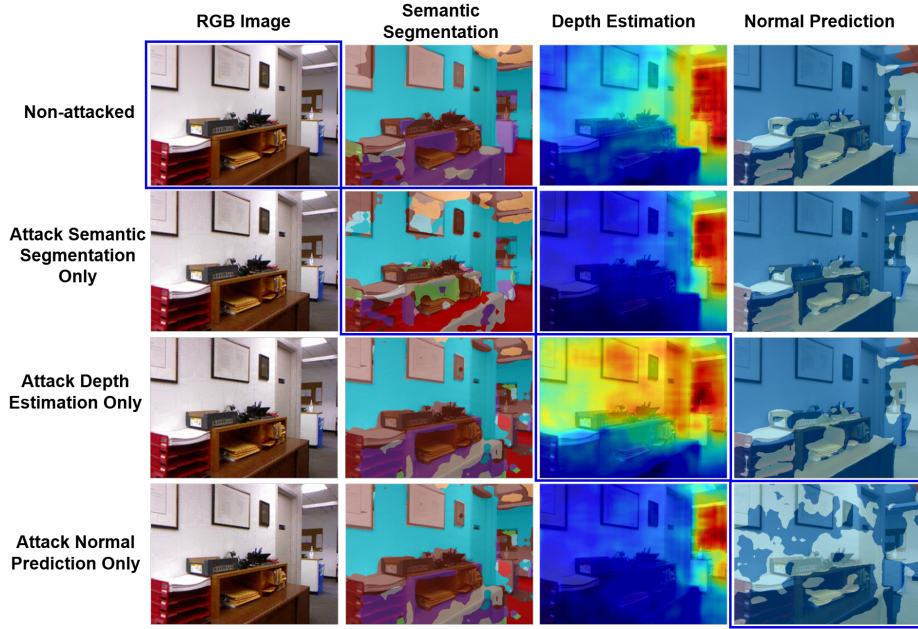


Fig. 2: Visualization results of non-attacked outputs and SMTA² attacks on NYUv2 dataset ($\epsilon=4/255$). Each row corresponds to a targeted task, while columns show the targeted and non-targeted task outputs. SMTA² effectively degrades the targeted task while preserving non-targeted tasks.

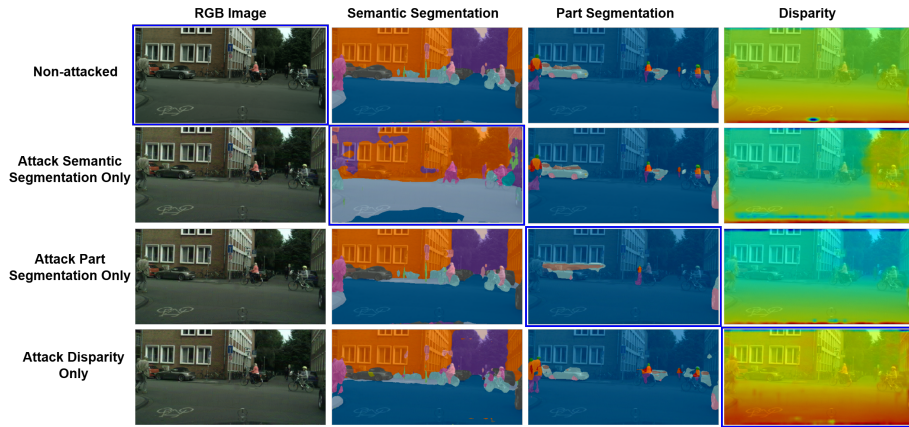


Fig. 3: Visualization results of non-attacked outputs and SMTA² attacks on Cityscapes dataset ($\epsilon=4/255$). Each row corresponds to a targeted task, while columns show the targeted and non-targeted task outputs..

retained performance. The dashed blue line denotes the clean baseline with all ratios equal to 1.

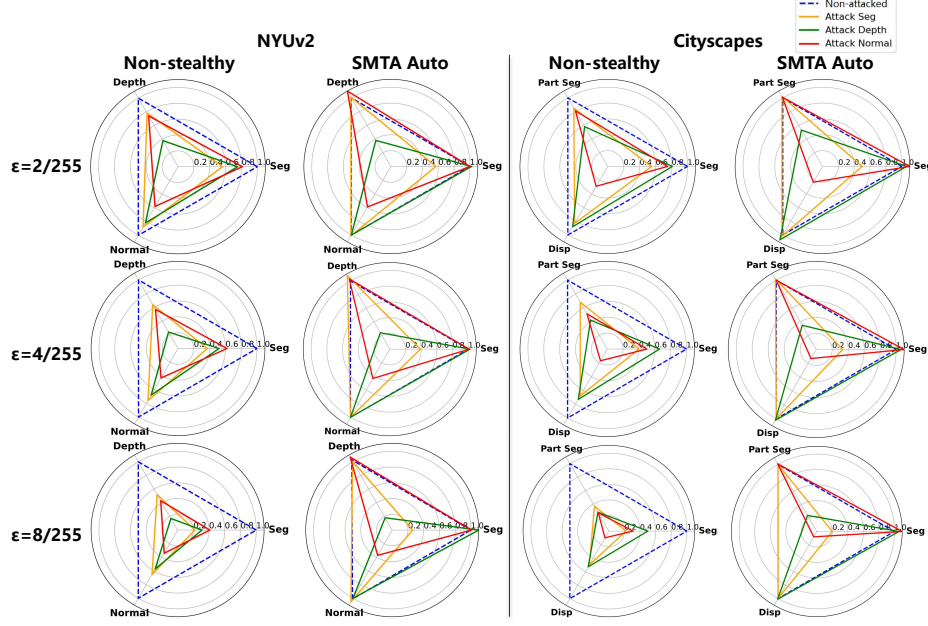


Fig. 4: Radar images of undefended PGD L_∞ attack effectiveness of non-stealthy and SMTA² Automated methods of NYUv2 and Cityscapes datasets.

As shown, non-stealthy attacks significantly degrade all tasks, even when only a single task is targeted, indicating substantial collateral damage. In contrast, automated SMTA² selectively degrades the targeted task while strictly preserving, or even improving, the performance of non-targeted tasks, demonstrating its strong task-selective attack capability and stealthiness.

Attack Results of Adversarially Trained Models To evaluate robustness under defensive settings, we further conduct experiments on Adversarially Trained (AT) models obtained via PGD-based adversarial training. Tables 6 and 7 present results under PGD L_∞ attacks on NYUv2 and Cityscapes using the same perturbation budgets ($\epsilon=2/255, 4/255, 8/255$). Across all settings, SMTA² consistently preserves non-targeted tasks while effectively degrading the targeted one, demonstrating stable task-selective attack capability. Both manual and automated variants show reliable performance.

Similar trends are observed under PGD L_2 attacks (Tables 8 and 9). While non-stealthy attacks degrade all tasks, SMTA² maintains baseline performance for non-targeted tasks and achieves consistent targeted degradation. Despite enhanced robustness from adversarial training, SMTA² remains effective, confirming its reliability under strong defense settings.

Table 6: Results of non-stealthy attacks and SMTA² under PGD adversarial training on **NYUv2** dataset using **PGD** L_∞ attack. Non-attacked performance is denoted by $\epsilon=0$. The targeted-task results are underlined. **Blue** and **red** indicate affected and preserved values, respectively. ($\uparrow$)/($\downarrow$) denote higher/lower is better.

Method	Non-stealthy			SMTA ² Manual			SMTA ² Auto		
Tasks	Segment [mIoU($\uparrow$)]	Depth [aErr($\downarrow$)]	Normal [mDist($\downarrow$)]	Segment [mIoU($\uparrow$)]	Depth [aErr($\downarrow$)]	Normal [mDist($\downarrow$)]	Segment [mIoU($\uparrow$)]	Depth [aErr($\downarrow$)]	Normal [mDist($\downarrow$)]
$\epsilon=0$	46.56	40.57	23.41	46.56	40.57	23.41	46.56	40.57	23.41
$\epsilon=2/255$	<u>33.21</u>	<u>52.27</u>	<u>28.00</u>	<u>35.67</u>	<u>36.82</u>	<u>23.00</u>	<u>33.66</u>	<u>39.94</u>	<u>23.36</u>
	31.00	<u>61.97</u>	30.03	<u>47.40</u>	<u>54.65</u>	<u>23.01</u>	<u>47.39</u>	<u>56.05</u>	<u>22.41</u>
	31.04	58.55	<u>31.12</u>	<u>49.54</u>	<u>40.17</u>	<u>32.82</u>	<u>47.79</u>	<u>38.76</u>	<u>32.87</u>
$\epsilon=4/255$	<u>27.37</u>	<u>59.42</u>	<u>30.28</u>	<u>30.91</u>	<u>39.26</u>	<u>23.37</u>	<u>30.99</u>	<u>38.98</u>	<u>23.31</u>
	34.58	<u>63.36</u>	28.46	<u>47.87</u>	<u>59.98</u>	<u>23.37</u>	<u>48.13</u>	<u>61.50</u>	<u>23.20</u>
	30.41	59.41	<u>32.59</u>	<u>46.84</u>	<u>40.64</u>	<u>31.06</u>	<u>46.76</u>	<u>39.17</u>	<u>30.73</u>
$\epsilon=8/255$	<u>23.33</u>	61.37	<u>30.97</u>	<u>27.45</u>	<u>39.87</u>	<u>22.86</u>	<u>26.83</u>	<u>38.67</u>	<u>23.17</u>
	<u>28.68</u>	<u>76.32</u>	<u>31.15</u>	<u>48.38</u>	<u>69.60</u>	<u>22.37</u>	<u>46.78</u>	<u>73.72</u>	<u>23.01</u>
	28.90	61.44	<u>35.85</u>	<u>48.00</u>	<u>39.95</u>	<u>35.27</u>	<u>47.38</u>	<u>39.76</u>	<u>34.42</u>

Table 7: Results of non-stealthy attacks and SMTA² under PGD adversarial training on **Cityscapes** dataset using **PGD** L_∞ attack. The targeted-task results are underlined. **Blue** and **red** indicate affected and preserved values, respectively. ($\uparrow$)/($\downarrow$) denote higher/lower is better.

Method	Non-stealthy			SMTA ² Manual			SMTA ² Auto		
Tasks	Segment [mIoU($\uparrow$)]	Part Seg [mIoU($\uparrow$)]	Disparity [aErr($\downarrow$)]	Segment [mIoU($\uparrow$)]	Part Seg [mIoU($\uparrow$)]	Disparity [aErr($\downarrow$)]	Segment [mIoU($\uparrow$)]	Part Seg [mIoU($\uparrow$)]	Disparity [aErr($\downarrow$)]
$\epsilon=0$	54.20	51.82	81.51	54.20	51.82	81.51	54.20	51.82	81.51
$\epsilon=2/255$	<u>33.21</u>	38.61	<u>101.67</u>	<u>42.26</u>	<u>52.01</u>	<u>80.69</u>	<u>42.09</u>	<u>52.00</u>	<u>81.23</u>
	<u>40.59</u>	<u>41.09</u>	94.69	<u>55.12</u>	<u>43.61</u>	<u>80.68</u>	<u>55.20</u>	<u>43.18</u>	<u>80.99</u>
	33.69	37.34	<u>103.34</u>	<u>54.78</u>	<u>52.19</u>	<u>98.89</u>	<u>54.83</u>	<u>52.52</u>	<u>99.63</u>
$\epsilon=4/255$	<u>30.98</u>	37.93	<u>102.75</u>	<u>30.29</u>	<u>52.45</u>	<u>81.41</u>	<u>31.22</u>	<u>53.10</u>	<u>81.30</u>
	34.17	<u>34.78</u>	101.93	<u>54.78</u>	<u>30.82</u>	<u>81.22</u>	<u>54.96</u>	<u>31.16</u>	<u>80.66</u>
	34.41	38.24	<u>112.85</u>	<u>54.68</u>	<u>51.94</u>	<u>115.52</u>	<u>55.34</u>	<u>52.33</u>	<u>114.18</u>
$\epsilon=8/255$	<u>27.12</u>	36.25	<u>105.51</u>	<u>32.56</u>	<u>53.10</u>	<u>81.21</u>	<u>28.69</u>	<u>52.79</u>	<u>81.08</u>
	<u>32.32</u>	<u>30.38</u>	103.75	<u>54.77</u>	<u>33.95</u>	<u>81.04</u>	<u>56.05</u>	<u>30.62</u>	<u>80.21</u>
	32.80	36.69	<u>127.97</u>	<u>54.67</u>	<u>51.97</u>	<u>112.91</u>	<u>56.18</u>	<u>52.56</u>	<u>122.42</u>

5 Conclusion

We propose Stealthy Multi-Task Adversarial Attacks (SMTA²), a principled framework for selectively degrading a targeted task while strictly preserving non-targeted tasks in multi-task learning. Unlike conventional adversarial attacks that cause indiscriminate performance collapse, SMTA² formulates selective degradation as a constrained multi-objective optimization problem and solves it via adaptive weight optimization. Despite the strong inter-task coupling in shared representations in multi-task learning, SMTA² reliably achieves

Table 8: Results of non-stealthy attacks and SMTA² under PGD adversarial training on **NYUv2** dataset using **PGD** L_2 attack. The targeted-task results are underlined. **Blue** and **red** indicate affected and preserved values, respectively. ($\uparrow$)/($\downarrow$) denote higher/lower is better.

Method	Non-stealthy			SMTA ² Manual			SMTA ² Auto		
Tasks	Segment [mIoU($\uparrow$)]	Depth [aErr($\downarrow$)]	Normal [mDist($\downarrow$)]	Segment [mIoU($\uparrow$)]	Depth [aErr($\downarrow$)]	Normal [mDist($\downarrow$)]	Segment [mIoU($\uparrow$)]	Depth [aErr($\downarrow$)]	Normal [mDist($\downarrow$)]
$\epsilon=0$	46.56	40.57	23.41	46.56	40.57	23.41	46.56	40.57	23.41
$\epsilon=5$	<u>25.65</u>	60.49	30.72	<u>34.25</u>	40.25	23.02	<u>34.20</u>	40.33	23.26
	29.26	<u>71.26</u>	30.96	46.83	<u>61.80</u>	23.34	46.66	<u>61.97</u>	22.94
	29.32	60.86	<u>34.46</u>	46.93	39.69	<u>31.52</u>	47.02	40.56	<u>35.21</u>
$\epsilon=10$	<u>20.16</u>	64.05	32.24	<u>24.88</u>	39.68	23.20	<u>24.58</u>	38.93	22.42
	29.97	<u>88.34</u>	30.98	46.92	<u>78.81</u>	22.82	46.94	<u>83.19</u>	22.37
	26.35	64.80	<u>40.39</u>	47.00	39.36	<u>38.81</u>	47.25	38.66	<u>38.55</u>

Table 9: Results of non-stealthy attacks and SMTA² under PGD adversarial training on **Cityscapes** dataset using **PGD** L_2 attack. The targeted-task results are underlined. **Blue** and **red** indicate affected and preserved values, respectively. ($\uparrow$)/($\downarrow$) denote higher/lower is better.

Method	Non-stealthy			SMTA ² Manual			SMTA ² Auto		
Tasks	Segment [mIoU($\uparrow$)]	Part Seg [mIoU($\uparrow$)]	Disparity [aErr($\downarrow$)]	Segment [mIoU($\uparrow$)]	Part Seg [mIoU($\uparrow$)]	Disparity [aErr($\downarrow$)]	Segment [mIoU($\uparrow$)]	Part Seg [mIoU($\uparrow$)]	Disparity [aErr($\downarrow$)]
$\epsilon=0$	54.20	51.82	81.51	54.20	51.82	81.51	54.20	51.82	81.51
$\epsilon=5$	<u>26.03</u>	35.76	108.67	<u>26.28</u>	52.84	81.03	<u>26.22</u>	52.13	81.44
	35.67	<u>33.06</u>	98.76	55.72	<u>30.26</u>	80.81	56.13	<u>31.17</u>	81.11
	36.84	39.65	<u>135.95</u>	56.20	52.11	<u>133.21</u>	56.92	53.41	<u>128.56</u>
$\epsilon=10$	<u>21.06</u>	31.83	119.86	<u>29.46</u>	52.88	80.93	<u>28.51</u>	52.52	80.78
	26.90	21.37	116.82	54.96	<u>28.57</u>	81.40	55.15	<u>26.76</u>	80.58
	26.22	27.53	<u>229.00</u>	56.31	53.46	<u>160.88</u>	56.74	53.40	<u>165.08</u>

targeted attack effectiveness without collateral damage, even under adversarial trained scenarios. Experiments demonstrate that SMTA² consistently achieves strong attack effectiveness on targeted tasks while preserving non-targeted tasks across undefended and adversarially trained settings. This work demonstrates that stealthy, task-selective attacks are both feasible and practically solvable, opening a new direction for controlled adversarial machine learning in multi-task systems.

Acknowledgments

This research was supported by the National Science Foundation under award CNS-2245765.

References

1. Andor, D., Alberti, C., Weiss, D., Severyn, A., Presta, A., Ganchev, K., Petrov, S., Collins, M.: Globally normalized transition-based neural networks. In: Proceedings of the 54th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers). pp. 2442–2452 (2016)
2. Carlini, N., Mishra, P., Vaidya, T., Zhang, Y., Sherr, M., Shields, C., Wagner, D., Zhou, W.: Hidden voice commands. In: 25th USENIX security symposium (USENIX security 16). pp. 513–530 (2016)
3. Carlini, N., Wagner, D.: Towards evaluating the robustness of neural networks. In: 2017 IEEE Symposium on Security and Privacy (SP). pp. 39–57. IEEE (2017)
4. Caruana, R.: Multitask learning. *Machine learning* **28**, 41–75 (1997)
5. Chakraborty, A., Alam, M., Dey, V., Chattopadhyay, A., Mukhopadhyay, D.: A survey on adversarial attacks and defences. *CAAI Transactions on Intelligence Technology* **6**(1), 25–45 (2021)
6. Chen, Z., Wang, C., Crandall, D.: Semantically stealthy adversarial attacks against segmentation models. In: Proceedings of the IEEE/CVF Winter Conference on Applications of Computer Vision. pp. 4080–4089 (2022)
7. Cordts, M., Omran, M., Ramos, S., Rehfeld, T., Enzweiler, M., Benenson, R., Franke, U., Roth, S., Schiele, B.: The cityscapes dataset for semantic urban scene understanding. In: Proceedings of the IEEE conference on computer vision and pattern recognition. pp. 3213–3223 (2016)
8. Crawshaw, M.: Multi-task learning with deep neural networks: A survey. arXiv preprint arXiv:2009.09796 (2020)
9. Croce, F., Hein, M.: Mind the box: l_1 -apgd for sparse adversarial attacks on image classifiers. In: International Conference on Machine Learning. pp. 2201–2211. PMLR (2021)
10. Dahl, G.E., Yu, D., Deng, L., Acero, A.: Context-dependent pre-trained deep neural networks for large-vocabulary speech recognition. *IEEE Transactions on audio, speech, and language processing* **20**(1), 30–42 (2011)
11. Du, A., Chen, B., Chin, T.J., Law, Y.W., Sasdelli, M., Rajasegaran, R., Campbell, D.: Physical adversarial attacks on an aerial imagery object detector. In: Proceedings of the IEEE/CVF Winter Conference on Applications of Computer Vision. pp. 1796–1806 (2022)
12. Ghamizi, S., Cordy, M., Papadakis, M., Le Traon, Y.: Adversarial robustness in multi-task learning: Promises and illusions. In: Proceedings of the AAAI Conference on Artificial Intelligence. vol. 36, pp. 697–705 (2022)
13. Goodfellow, I.J., Shlens, J., Szegedy, C.: Explaining and harnessing adversarial examples. arXiv preprint arXiv:1412.6572 (2014)
14. Guo, J., Lei, L., Sun, H., Qin, M., Yu, H., Zhang, T.: A min-max optimization framework for sparse multi-task deep neural network. *Neurocomputing* p. 130865 (2025)
15. Guo, J., Li, L., Yang, H., Geng, B., Yu, H., Qin, M., Zhang, T.: Robust multi-task adversarial attacks using min-max optimization. In: ICASSP 2025-2025 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP). pp. 1–5. IEEE (2025)
16. Guo, J., Sun, H., Qin, M., Yu, H., Zhang, T.: A min-max optimization framework for multi-task deep neural network compression. In: 2024 IEEE International Symposium on Circuits and Systems (ISCAS). pp. 1–5. IEEE (2024)

17. Guo, P., Xu, Y., Lin, B., Zhang, Y.: Multi-task adversarial attack. arXiv preprint arXiv:2011.09824 (2020)
18. He, K., Zhang, J., Yan, Y., Xu, W., Niu, C., Zhou, J.: Contrastive zero-shot learning for cross-domain slot filling with adversarial attack. In: Proceedings of the 28th International Conference on Computational Linguistics. pp. 1461–1467 (2020)
19. Hinton, G., Deng, L., Yu, D., Dahl, G.E., Mohamed, A.r., Jaitly, N., Senior, A., Vanhoucke, V., Nguyen, P., Sainath, T.N., et al.: Deep neural networks for acoustic modeling in speech recognition: The shared views of four research groups. *IEEE Signal processing magazine* **29**(6), 82–97 (2012)
20. Hu, J., Ran, H., Sun, C., Zhang, Q., Wang, G.: Dynamic gradient fusion method with local spatial and multi-scale frequency transformations for transferable adversarial attacks. *Expert Systems with Applications* p. 131677 (2026)
21. Jia, J., Zhang, W., Wang, R., Yan, D.: Adaptive multi-task adversarial attacks on audio tasks. *Expert Systems with Applications* p. 133000 (2026)
22. Kurakin, A., Goodfellow, I.J., Bengio, S.: Adversarial examples in the physical world. In: Artificial intelligence safety and security, pp. 99–112. International Conference of Learning Representations (2018)
23. LeCun, Y., Bottou, L., Bengio, Y., Haffner, P.: Gradient-based learning applied to document recognition. *Proceedings of the IEEE* **86**(11), 2278–2324 (1998)
24. Li, L., Yang, H., Guo, J., Yu, H., Qin, M., Zhang, T.: Task-aware federated multi-task learning. In: Proceedings of the 7th ACM International Conference on Multimedia in Asia. pp. 1–7 (2025)
25. Li, L., Yang, H., Guo, J., Yu, H., Qin, M., Zhang, T.: Fedsta: Spatio-temporal alternation for efficient federated multi-task learning. *IEEE Transactions on Circuits and Systems for Video Technology* (2026)
26. Liu, S., James, S., Davison, A.J., Johns, E.: Auto-lambda: Disentangling dynamic task relationships. arXiv preprint arXiv:2202.03091 (2022)
27. Liu, X., Li, J., Ma, J., Sun, H., Xu, Z., Zhang, T., Yu, H.: Deep transfer learning for intelligent vehicle perception: A survey. *Green Energy and Intelligent Transportation* p. 100125 (2023)
28. Long, K., Sheng, Z., Shi, H., Li, X., Chen, S., Ahn, S.: Physical enhanced residual learning (perl) framework for vehicle trajectory prediction. *Communications in Transportation Research* **5**, 100166 (2025)
29. Madry, A., Makelov, A., Schmidt, L., Tsipras, D., Vladu, A.: Towards deep learning models resistant to adversarial attacks. In: International Conference on Learning Representations (2018)
30. Mao, C., Chiquier, M., Wang, H., Yang, J., Vondrick, C.: Adversarial attacks are reversible with natural supervision. In: Proceedings of the IEEE/CVF International Conference on Computer Vision. pp. 661–671 (2021)
31. Moosavi-Dezfooli, S.M., Fawzi, A., Fawzi, O., Frossard, P.: Universal adversarial perturbations. In: Proceedings of the IEEE conference on computer vision and pattern recognition. pp. 1765–1773 (2017)
32. Nguyen, A., Yosinski, J., Clune, J.: Deep neural networks are easily fooled: High confidence predictions for unrecognizable images. In: Proceedings of the IEEE conference on computer vision and pattern recognition. pp. 427–436 (2015)
33. Sener, O., Koltun, V.: Multi-task learning as multi-objective optimization. *Advances in neural information processing systems* **31** (2018)
34. Silberman, N., Hoiem, D., Kohli, P., Fergus, R.: Indoor segmentation and support inference from rgb-d images. In: Computer Vision–ECCV 2012: 12th European Conference on Computer Vision, Florence, Italy, October 7–13, 2012, Proceedings, Part V 12. pp. 746–760. Springer (2012)

35. Sun, H., Fu, L., Li, J., Guo, Q., Meng, Z., Zhang, T., Lin, Y., Yu, H.: Defense against adversarial cloud attack on remote sensing salient object detection. In: Proceedings of the IEEE/CVF Winter Conference on Applications of Computer Vision. pp. 8345–8354 (2024)
36. Xu, K., Cheng, X., Qiao, J., Li, J.T., Ji, K.X., Zhong, J.Y., Tian, P., Mi, J.X.: Adversarial attacks to manipulate target localization of object detector. *IEEE Access* (2024)
37. Xu, Q., Zhou, L., Zhong, X., Zhang, F., Tian, J., Yu, X., Huang, R.: Etv-attack: Efficient text-driven visual-variable adversarial attacks on visual question answering with pre-trained language models. *Pattern Recognition* p. 113202 (2026)
38. Yang, H., Li, L., Guo, J., Li, B., Qin, M., Yu, H., Zhang, T.: Da3d: Domain-aware dynamic adaptation for all-weather multimodal 3d detection. In: Proceedings of the 33rd ACM International Conference on Multimedia. pp. 2150–2158 (2025)
39. Yu, J., Dai, Y., Liu, X., Huang, J., Shen, Y., Zhang, K., Zhou, R., Adhikarla, E., Ye, W., Liu, Y., et al.: Unleashing the power of multi-task learning: A comprehensive survey spanning traditional, deep, and pretrained foundation model eras. *arXiv preprint arXiv:2404.18961* (2024)
40. Zhang, H., Yu, Y., Jiao, J., Xing, E., El Ghaoui, L., Jordan, M.: Theoretically principled trade-off between robustness and accuracy. In: International conference on machine learning. pp. 7472–7482. PMLR (2019)
41. Zhang, L., Liu, X., Mahmood, K., Ding, C., Guan, H.: Attacking all tasks at once using adversarial examples in multi-task learning. *Neurocomputing* p. 131503 (2025)
42. Zhang, Z., Yu, W., Yu, M., Guo, Z., Jiang, M.: A survey of multi-task learning in natural language processing: Regarding task relatedness and training methods. In: Proceedings of the 17th Conference of the European Chapter of the Association for Computational Linguistics. pp. 943–956 (2023)
43. Zhe, Y., Nagaike, R., Nishiyama, D., Fukuchi, K., Sakuma, J.: Adversarial attacks on hidden tasks in multi-task learning. *arXiv preprint arXiv:2405.15244* (2024)