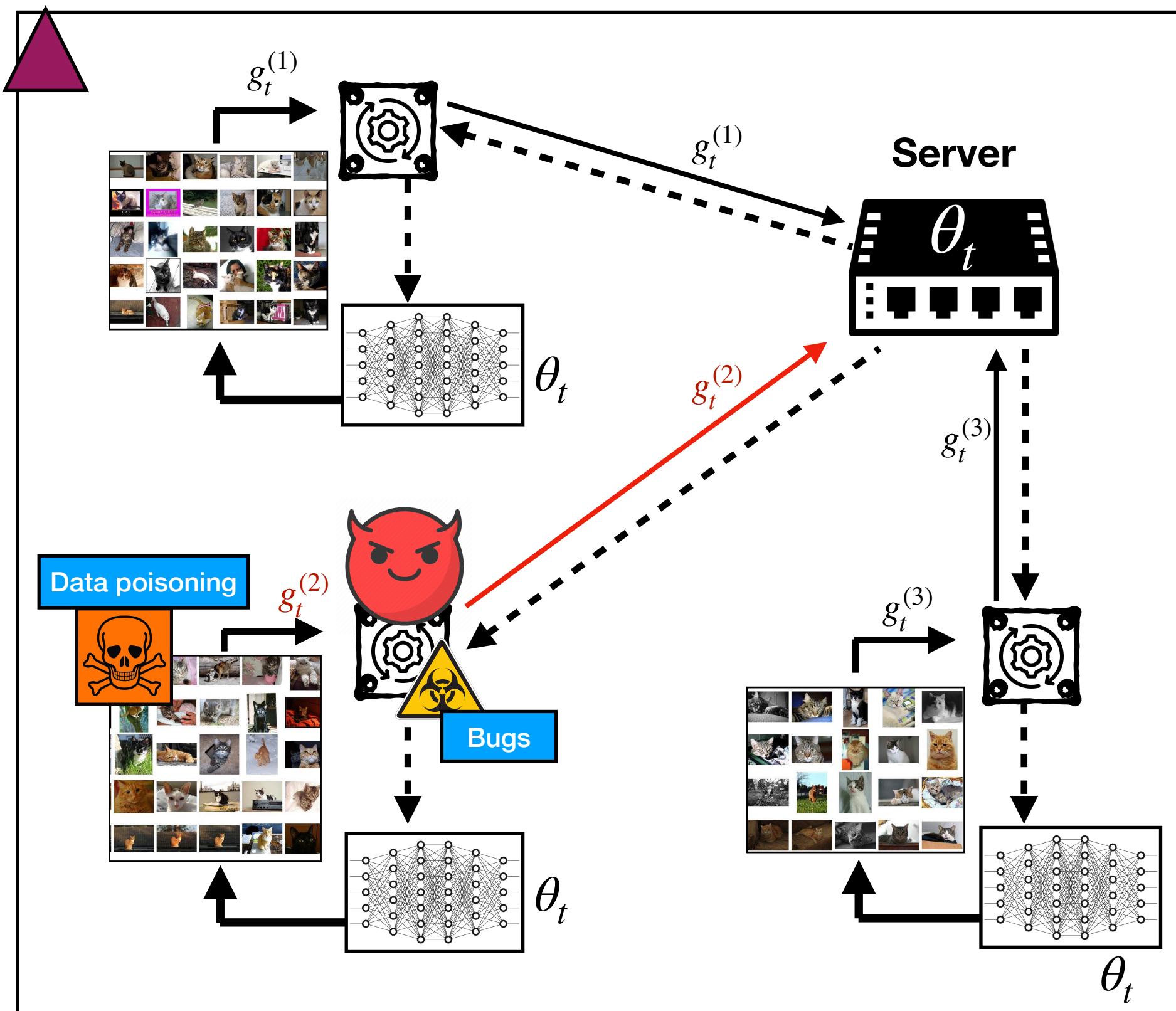


Adaptive Gradient Clipping for Robust Federated Learning

Federated learning is vulnerable to untrusted clients

f out of n clients are adversarial (Byzantine)



Robust aggregation (RobAgg)

Aggregation $F : (v_1, \dots, v_n) \mapsto \hat{v}$

for all $G \subseteq [n]$ of size $n - f$

$$\|\hat{v} - \bar{v}_G\|^2 \leq \kappa \frac{1}{n-f} \sum_{i \in G} \|v_i - \bar{v}_G\|^2$$

$$\bar{v}_G = \frac{1}{n-f} \sum_{i \in G} v_i$$

Coordinate-wise Trimmed Mean (CWTM)

$$\text{CWTM} \left(\begin{pmatrix} 2 & 3 & 2 & \times & \times \\ 5 & \times & 4 & 5 & \times \end{pmatrix} \right) = \begin{pmatrix} 2.3 \\ 4.6 \end{pmatrix}$$

$$\kappa \in \Theta \left(\frac{f}{n} \right)$$

Lower bound

$$\mathcal{L}_H \triangleq \frac{1}{|H|} \sum_{i \in H} \mathcal{L}_i$$

(G, B) -gradient dissimilarity:

$$\frac{1}{|H|} \sum_{i \in H} \|\nabla \mathcal{L}_i(\theta) - \nabla \mathcal{L}_H(\theta)\|^2 \leq G^2 + B^2 \|\mathcal{L}_H(\theta)\|^2$$

A federated learning algorithm $\mathcal{A}$ with output $\hat{\theta}$ can only tolerate $\frac{f}{n} < \frac{1}{2 + B^2} \triangleq \text{BP}$ and

$$\|\mathcal{L}_H(\hat{\theta})\|^2 \geq \frac{1}{4} \left(\frac{f}{n - (2 + B^2)f} \right) G^2 \triangleq \epsilon_o$$

Circumventing L.B. under bounded initialization error

Proposed scheme: Adaptive Robust Clipping (ARC)

Given input vectors $v_1, \dots, v_n$, determine a permutation $\pi : [n] \rightarrow [n]$ such that

$$\|v_{\pi(1)}\| \leq \dots \leq \|v_{\pi(n)}\|$$

Set $k = 2 \frac{f}{n} (n - f)$ and $C = \|v_{\pi(k+1)}\|$. Then,

$$\text{ARC}(v_1, \dots, v_n) = (\text{Clip}_C(v_1), \dots, \text{Clip}_C(v_n))$$

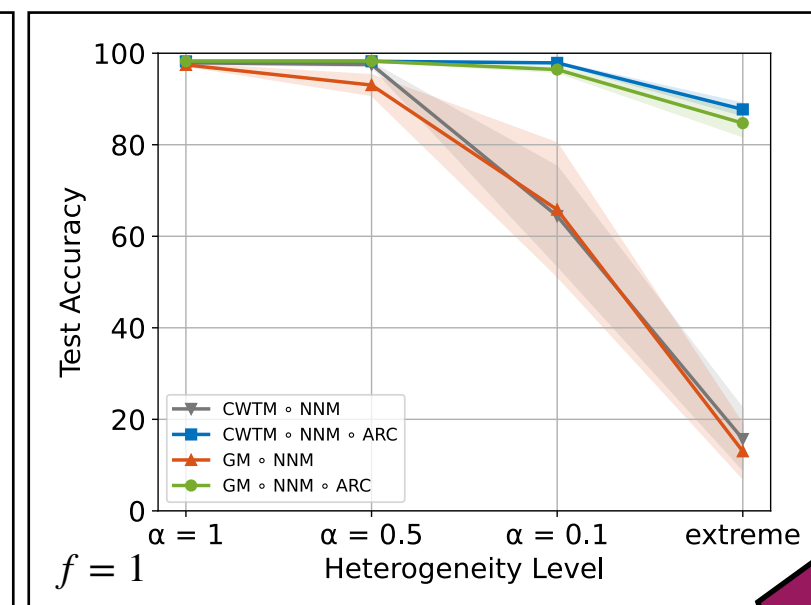
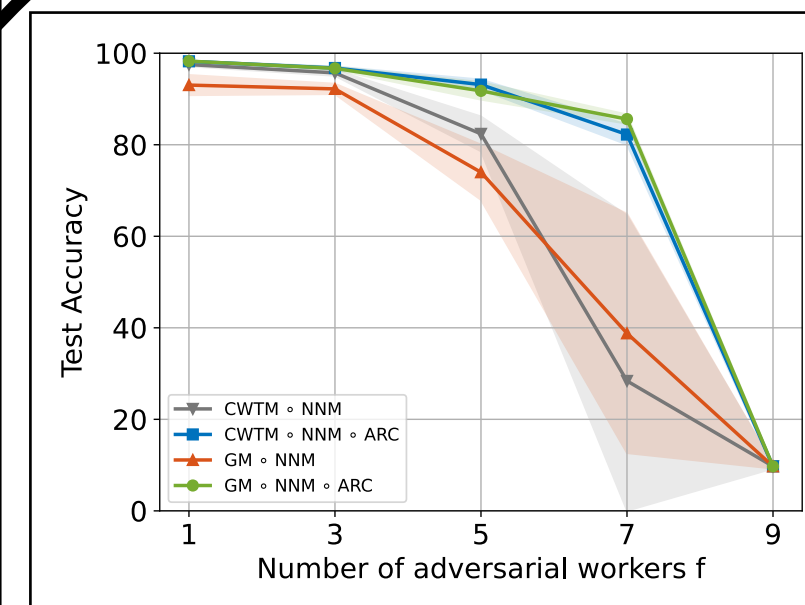
$$\text{Clip}_C(v) = \min \left\{ 1, \frac{C}{\|v\|} \right\} v$$

$\theta_{t+1} = \theta_t - \gamma_t \text{RobAgg}(\text{ARC}(v_1, \dots, v_n))$. For $v \in (0, 1)$ and $\frac{f}{n} = \left(1 - \frac{v}{\Psi(\theta_1)} \right) \text{BP}$,

$$\frac{1}{T} \sum_{t=1}^T \|\mathcal{L}_H(\theta_t)\|^2 \leq v \epsilon_o.$$

monotonically increasing function of $\|\nabla \mathcal{L}_H(\theta_1)\|$

Worst-case accuracies over MNIST, distributed over 10 honest clients



$g_t^{(i)} := \nabla \mathcal{L}_i(\theta_t)$, i is honest
 $g_t^{(j)}$ = arbitrary, j is Byzantine



$$\theta_{t+1} = \theta_t - \gamma_t \text{Avg}(g_t^{(1)}, \dots, g_t^{(n)})$$



$$\theta_{t+1} = \theta_t - \gamma_t \text{RobAgg}(g_t^{(1)}, \dots, g_t^{(n)})$$

