

Selective Classification Via Neural Network Training Dynamics

Stephan Rabanser

Anvith Thudi

Kimia Hamidieh

Adam Dziedzic

Nicolas Papernot

`stephan@cs.toronto.edu`

`anvith.thudi@mail.utoronto.ca`

`kimia@cs.toronto.edu`

`ady@vectorinstitute.ai`

`nicolas.papernot@utoronto.ca`



UNIVERSITY OF
TORONTO

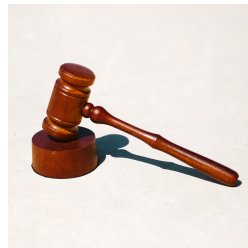
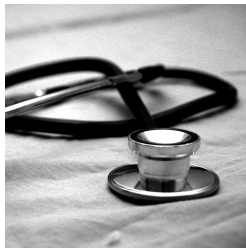


VECTOR
INSTITUTE

June 7, 2022

Motivation

Machine Learning systems are becoming ubiquitous.



Especially in high-stakes decision-making, it is of vital importance to detect inputs that a machine learning model would predict incorrectly on.

Image credit: unsplash.com

Standard Supervised Learning Setup

- Dataset $D = \{(\mathbf{x}_i, y_i)\}_{i=1}^M$ consisting of data points $(\mathbf{x}, y)$ with $\mathbf{x} \in \mathcal{X}$ and $y \in \mathcal{Y}$.
 - Covariate space $\mathcal{X} = \mathbb{R}^d$ of dimensionality d .
 - Label space $\mathcal{Y} = \{1, 2, \dots, C\}$ consisting of C classes.
- $(\mathbf{x}, y)$ are sampled iid from the underlying distribution p defined over $\mathcal{X} \times \mathcal{Y}$.
- Goal: learn a prediction function $f : \mathcal{X} \rightarrow \mathcal{Y}$ which minimizes the classification risk $R(\cdot)$ wrt the underlying data distribution p and an appropriately chosen loss function $\ell : \mathcal{Y} \times \mathcal{Y} \rightarrow \mathbb{R}$.
- We can derive the optimal parameters $\hat{\theta}$ via empirical risk minimization:

$$\hat{\theta} = \arg \min_{\theta} R(f_{\theta}) = \arg \min_{\theta} \mathbb{E}_{p(\mathbf{x}, y)} [\ell(f_{\theta}(\mathbf{x}), y)] \approx \arg \min_{\theta} \frac{1}{n} \sum_{i=1}^N \ell(f_{\theta}(\mathbf{x}_i), y_i)$$

Selective Classification (SC)

Selective classification introduces a rejection class $\perp$ via a *gating mechanism*. Our goal is to derive a selection function $g : \mathcal{X} \rightarrow \mathbb{R}$ which, given an acceptance threshold τ , determines whether a model should predict on a data point $\mathbf{x}$.

$$(f, g)(\mathbf{x}) = \begin{cases} f(\mathbf{x}) & g(\mathbf{x}) \geq \tau \\ \perp & \text{otherwise.} \end{cases}$$

The performance of a selective classifier (f, g) is assessed based on

- the *coverage* of (f, g) , i.e. what fraction of points we predict on
- the selective *accuracy* of (f, g) on the points it accepts.

$$\text{cov}(f, g) = \frac{|\{\mathbf{x} : g(\mathbf{x}) \geq \tau\}|}{|D|}$$

$$\text{acc}(f, g) = \frac{|\{\mathbf{x} : f(\mathbf{x}) = y, g(\mathbf{x}) \geq \tau\}|}{|\{\mathbf{x} : g(\mathbf{x}) \geq \tau\}|}$$

Training stage

Training set

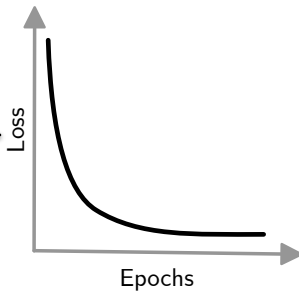


Training stage

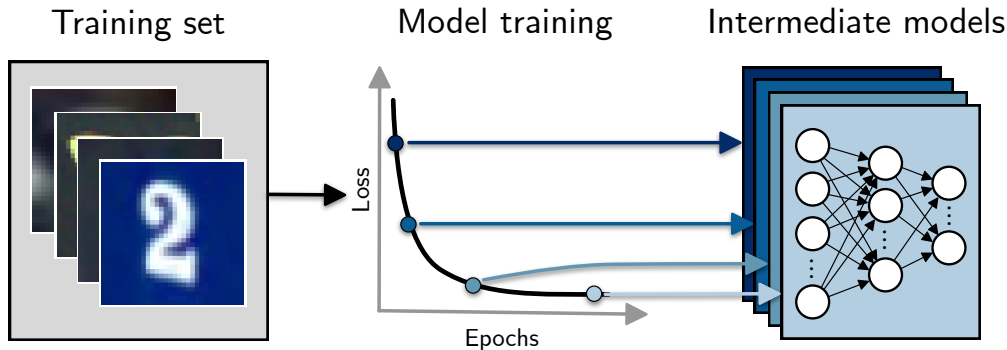
Training set



Model training

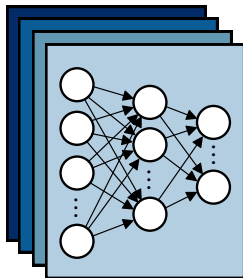


Training stage



Testing stage

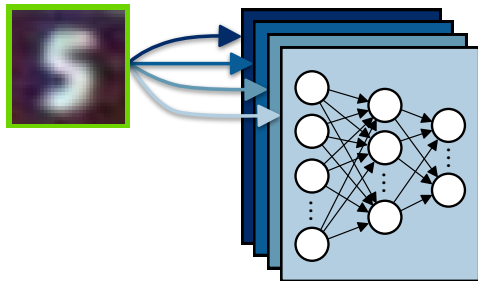
Intermediate models



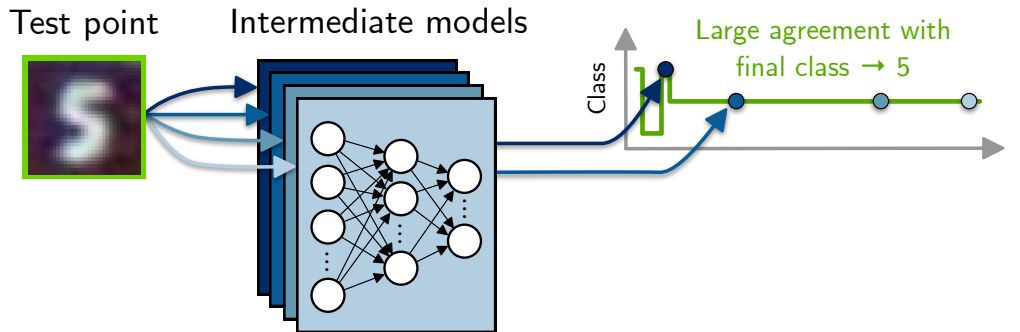
Testing stage

Test point

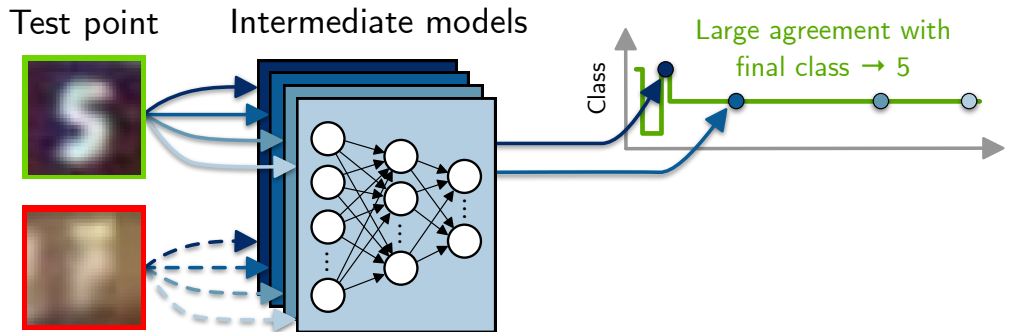
Intermediate models



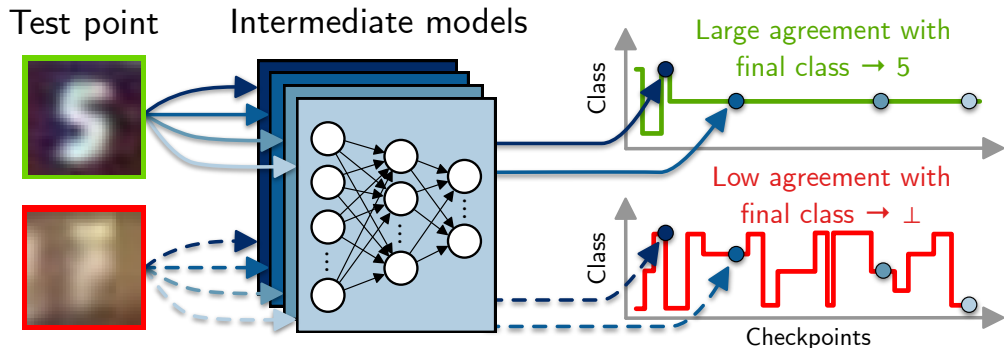
Testing stage



Testing stage

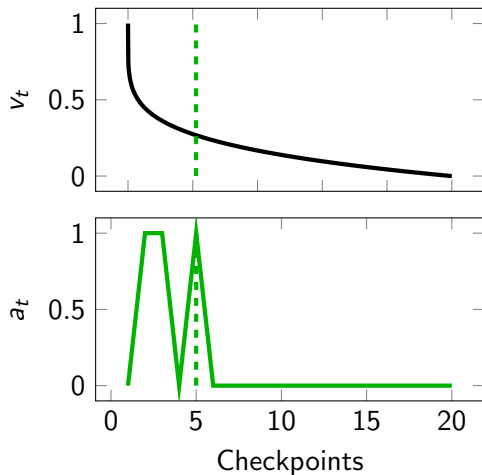


Testing stage



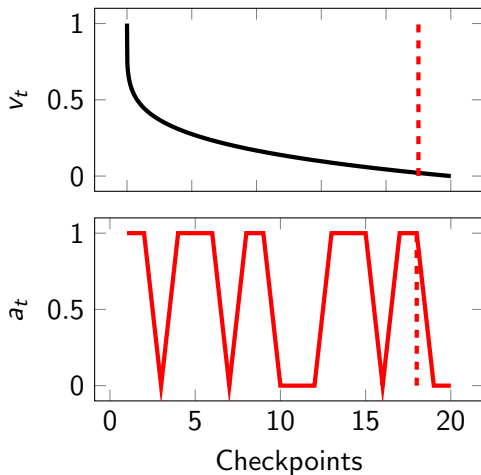
The Minimum Score $s_{\min}$

1. Denote $L = f_T(\mathbf{x})$, i.e. the label our final model predicts.
2. If $\exists t$ s.t. $a_t = 1$ then compute $s_{\min} = \min_{t \text{ s.t. } a_t=1} v_t$, else accept $\mathbf{x}$ with prediction L .
3. If $s_{\min} > \tau$ accept $\mathbf{x}$ with prediction L , else reject ($\perp$).



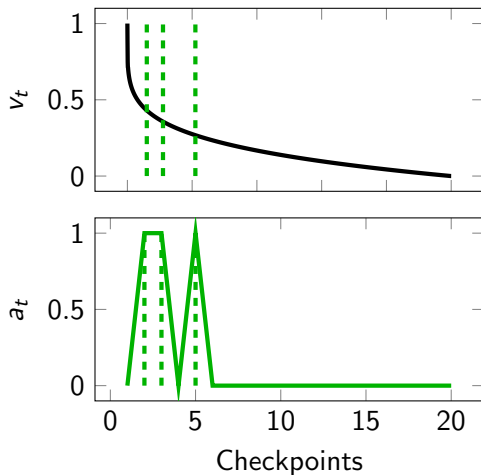
The Minimum Score $s_{\min}$

1. Denote $L = f_T(\mathbf{x})$, i.e. the label our final model predicts.
2. If $\exists t$ s.t. $a_t = 1$ then compute $s_{\min} = \min_{t \text{ s.t. } a_t = 1} v_t$, else accept $\mathbf{x}$ with prediction L .
3. If $s_{\min} > \tau$ accept $\mathbf{x}$ with prediction L , else reject ($\perp$).



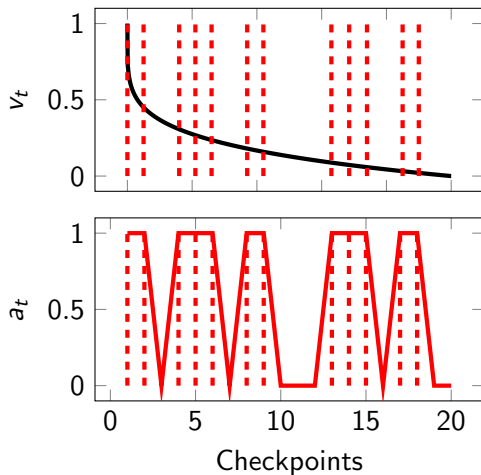
The Average Score s_{avg}

1. Denote $L = f_T(\mathbf{x})$, i.e. the label our final model predicts.
2. If $\exists t$ s.t. $a_t = 1$, compute $s_{\text{avg}} = \frac{\sum v_t a_t}{\sum a_t}$, else accept $\mathbf{x}$ with prediction L .
3. If $s_{\text{avg}} > \tau$ accept $\mathbf{x}$ with prediction L , else reject ($\perp$).



The Average Score s_{avg}

1. Denote $L = f_T(\mathbf{x})$, i.e. the label our final model predicts.
2. If $\exists t$ s.t. $a_t = 1$, compute $s_{\text{avg}} = \frac{\sum v_t a_t}{\sum a_t}$, else accept $\mathbf{x}$ with prediction L .
3. If $s_{\text{avg}} > \tau$ accept $\mathbf{x}$ with prediction L , else reject ($\perp$).



Performance at Low Target Error

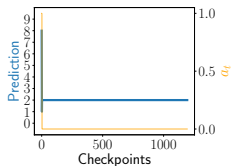
Dataset	Target Error	NNTD		OSP		SR		SN		DG	
		Cov.	Error	Cov.	Error	Cov.	Error	Cov.	Error	Cov.	Error
CIFAR-10	2%	83.3	1.96	80.6	1.91	75.1	2.09	73.0	2.31	74.2	1.98
	1%	79.7	1.05	74.0	1.02	67.2	1.09	64.5	1.02	66.4	1.01
	0.5%	74.2	0.49	64.1	0.51	59.3	0.53	57.6	0.48	57.8	0.51
SVHN	2%	95.7	1.96	95.8	1.99	95.7	2.06	93.5	2.03	94.8	1.99
	1%	91.2	0.99	90.1	1.03	88.4	0.99	86.5	1.04	89.5	1.01
	0.5%	83.9	0.50	82.4	0.51	77.3	0.51	79.2	0.51	81.6	0.49
Cats & Dogs	2%	90.5	2.03	90.5	1.98	88.2	2.03	84.3	1.94	87.4	1.94
	1%	85.6	1.01	85.4	0.98	80.2	0.97	78.0	0.98	81.7	0.98
	0.5%	77.5	0.50	78.7	0.49	73.2	0.49	70.5	0.46	74.5	0.48

Performance at High Target Coverage

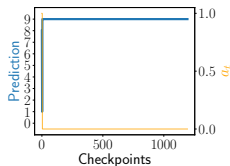
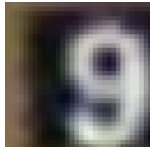
Dataset	Target Coverage	NNTD		OSP		SR		SN		DG	
		Cov.	Error	Cov.	Error	Cov.	Error	Cov.	Error	Cov.	Error
CIFAR-10	100%	100	9.57	100	9.74	99.99	9.58	100	11.07	100	10.81
	95%	95.1	6.13	95.1	6.98	95.2	8.74	94.7	8.34	95.1	8.21
	90%	90.1	4.16	90.0	4.67	90.5	6.52	89.6	6.45	90.1	6.14
SVHN	100%	100	4.11	100	4.27	99.97	3.86	100	4.27	100	4.03
	95%	94.8	1.80	95.1	1.83	95.1	1.86	95.1	2.53	95.0	2.05
	90%	90.1	0.77	90.1	1.01	90.0	1.04	90.1	1.31	90.0	1.06
Cats & Dogs	100%	100	5.18	100	5.93	100	5.72	100	7.36	100	6.16
	95%	94.9	2.99	95.1	2.97	95.0	3.46	95.2	5.1	95.1	4.28
	90%	90.0	1.87	90.0	1.74	90.0	2.28	90.2	3.3	90.0	2.50

Individual SVHN Example

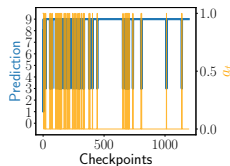
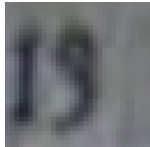
Prediction: 2 – Label: 2



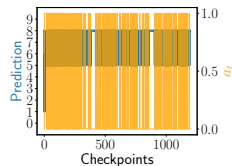
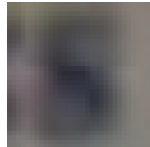
Prediction: 9 – Label: 9



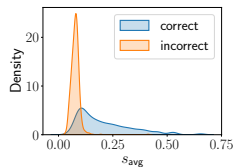
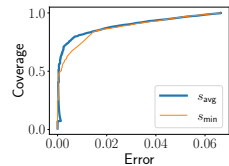
Prediction: 9 – Label: 3



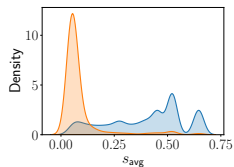
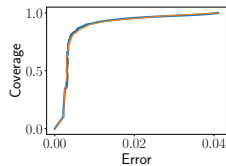
Prediction: 8 – Label: 8



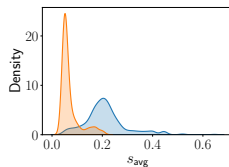
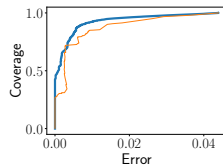
Performance of $s_{\min}$ and s_{avg}



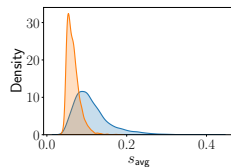
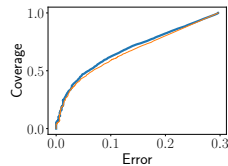
(a) CIFAR-10



(b) SVHN

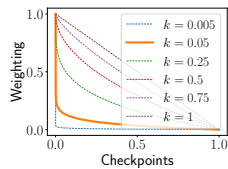


(c) GTSRB

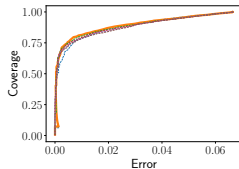


(d) CIFAR-100

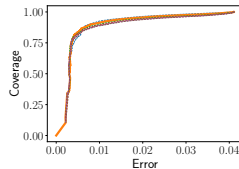
Hyperparameters: Weighting Parameter k Checkpointing Resolution T



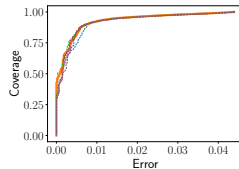
(a) Convex weighting



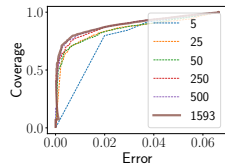
(b) CIFAR-10



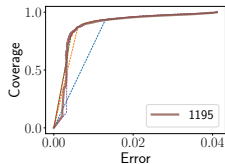
(c) SVHN



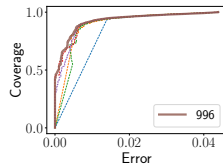
(d) GTSRB



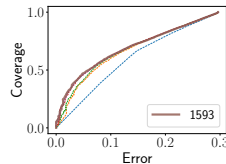
(a) CIFAR-10



(b) SVHN



(c) GTSRB

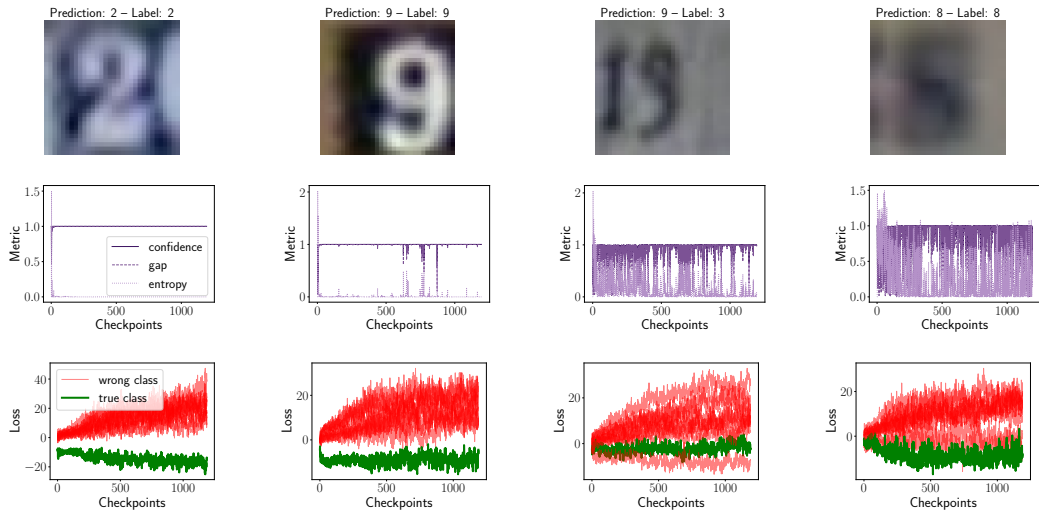


(d) CIFAR-100

Thanks! :)

Backup

Extended SVHN Example



CIFAR-10 Example

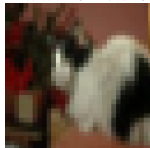
Prediction: 9 – Label: 9



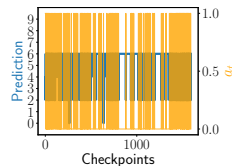
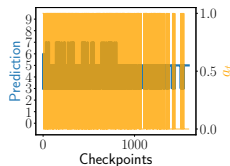
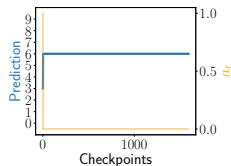
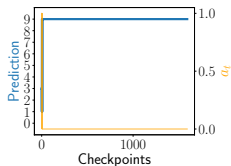
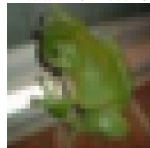
Prediction: 6 – Label: 6



Prediction: 5 – Label: 3

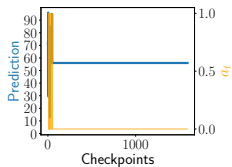


Prediction: 6 – Label: 6

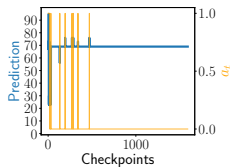


CIFAR-100 Example

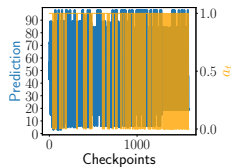
Prediction: 56 – Label: 56



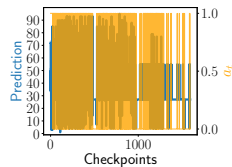
Prediction: 69 – Label: 69



Prediction: 19 – Label: 35



Prediction: 27 – Label: 27



GTSRB Example

Prediction: 15 – Label: 15



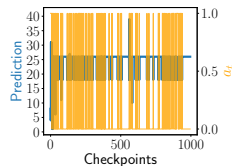
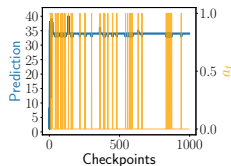
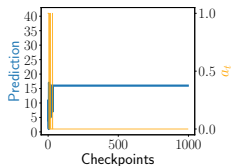
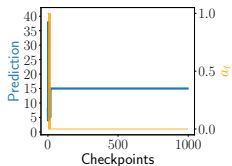
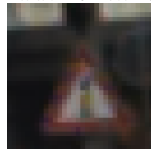
Prediction: 16 – Label: 16



Prediction: 34 – Label: 33

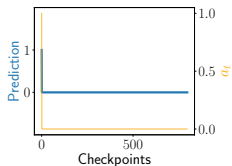
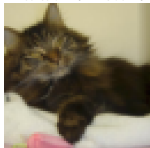


Prediction: 26 – Label: 26

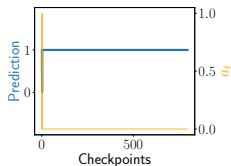
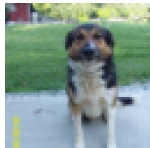


Cats & Dogs Example

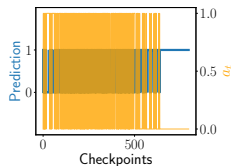
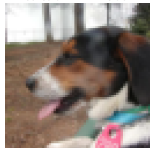
Prediction: 0 – Label: 0



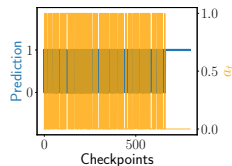
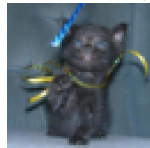
Prediction: 1 – Label: 1



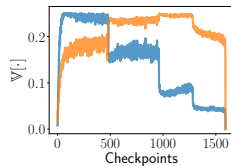
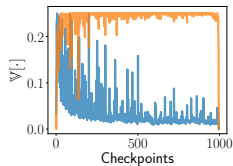
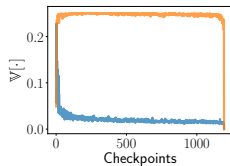
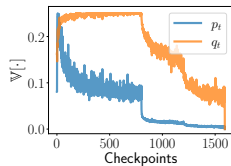
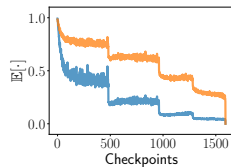
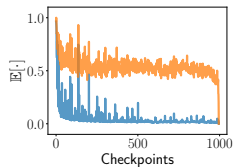
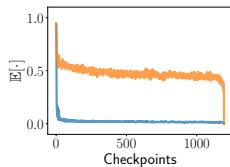
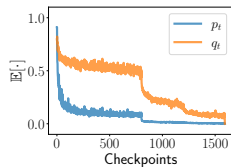
Prediction: 1 – Label: 1



Prediction: 1 – Label: 0



Examining the Convergence Behavior of Correct & Incorrect Points



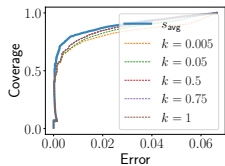
(a) CIFAR-10

(b) SVHN

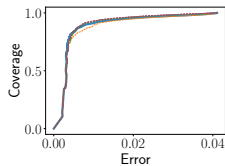
(c) GTSRB

(d) CIFAR-100

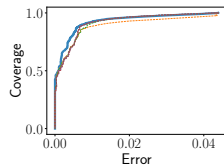
Incorporating Empirical Estimates of e_t and v_t



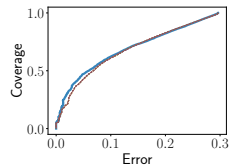
(a) CIFAR-10



(b) SVHN



(c) GTSRB



(d) CIFAR-100

Alternate Metrics

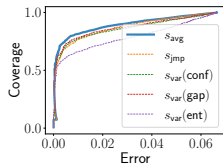
Jump Score

Measure level of disagreement between the predicted label of two successive intermediate models: for $j_t = 0$ iff $f_t(\mathbf{x}) = f_{t-1}(\mathbf{x})$ and $j_t = 1$ otherwise we can compute the jump score as

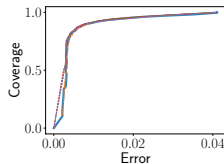
$$s_{\text{jmp}} = 1 - \sum v_t j_t.$$

Continuous Metrics (e.g. MSP, Entropy)

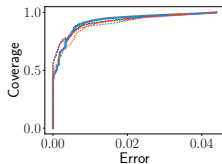
Assume that any of these metrics is given by a sequence $z = \{z_1, \dots, z_T\}$. Then we can capture the uniformity of z via a variance score $s_{\text{var}} = \sum_t w_t (z_t - \mu)^2$ for mean $\mu = \frac{1}{T} \sum_t z_t$ and an increasing weighting sequence $w = \{w_1, \dots, w_T\}$.



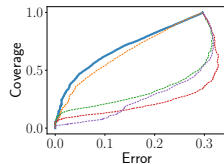
(a) CIFAR-10



(b) SVHN

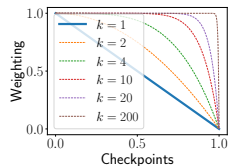


(c) GTSRB

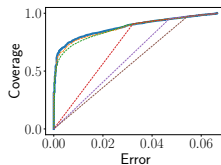


(d) CIFAR-100

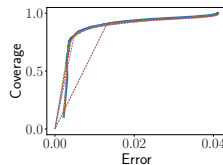
Concave Weighting



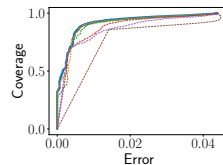
(a) Concave weighting



(b) CIFAR-10



(c) SVHN



(d) GTSRB