

Theory-Inspired Contrastive Learning with Self-Labeling Refinement

Pan Zhou, Caiming Xiong, Xiaotong Yuan, and Steven HOI

Salesforce

panzhou3@gmail.com

Dec 06, 2021

Outline

Motivation: why one-hot label in MoCo is not accurate?

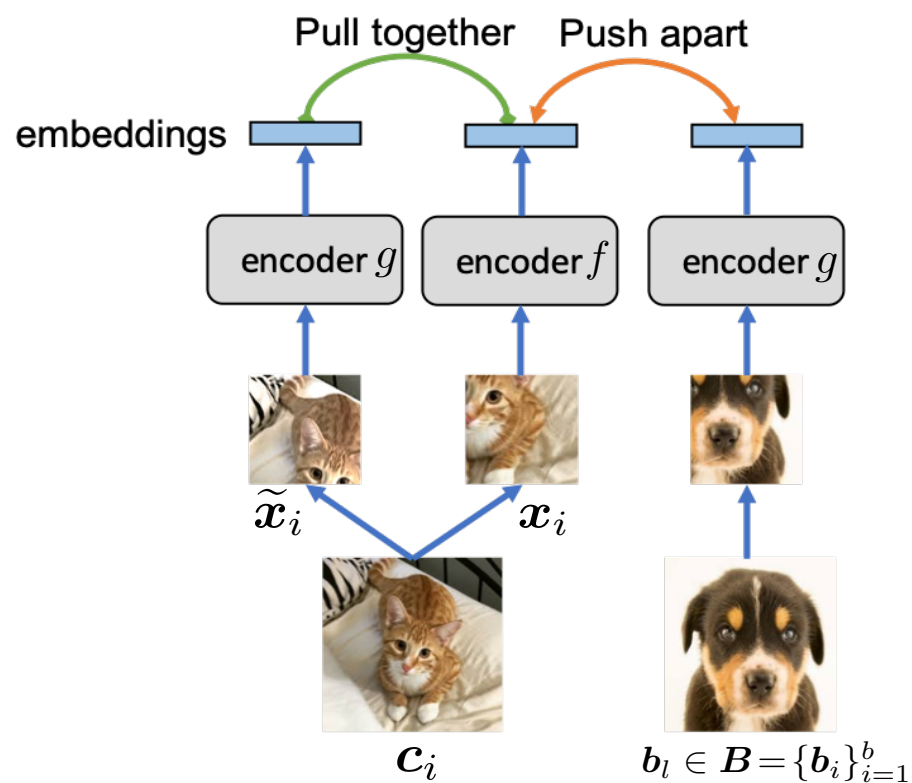
Solution for accurate label: self-labeling refinery and momentum mixup

Experiments: higher classification accuracy

Conclusion

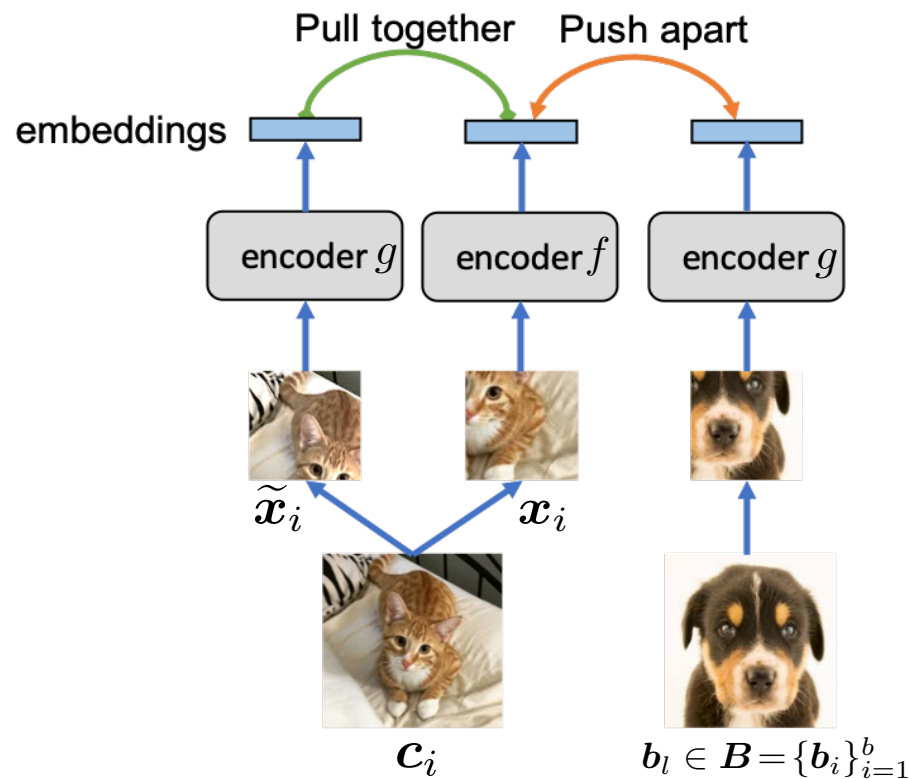
Background: MoCo

MoCo pulls together crops of same image, pushes away crops of different images.



Background: MoCo

MoCo pulls together crops of same image, pushes away crops of different images.



Given a minibatch samples $\{c_i\}_{i=1}^s$, it augments each c_i into two views $(x_i, \tilde{x}_i)$ and optimizes:

$$\mathcal{L}_n(w) = -\frac{1}{s} \sum_{i=1}^s \log \left(\frac{\sigma(x_i, \tilde{x}_i)}{\sigma(x_i, \tilde{x}_i) + \sum_{l=1}^b \sigma(x_i, b_l)} \right),$$

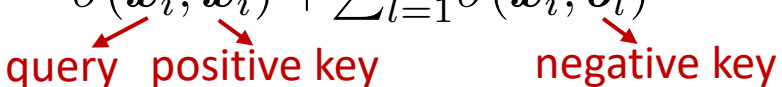
query positive key negative key

- $\sigma(x_i, \tilde{x}_i) = \exp \left(-\frac{\langle f(x_i), g(\tilde{x}_i) \rangle}{\tau \|f(x_i)\|_2 \cdot \|g(\tilde{x}_i)\|_2} \right)$: cosine similarity function
- f_w : online network
- g_ξ : target network
- $B = \{b_i\}_{i=1}^b$: negative key buffer (previous minibatch crops)

Motivation from Observations

One-hot label assignment: query $\mathbf{x}_i$ has only one positive $\tilde{\mathbf{x}}_i$ among $\tilde{\mathbf{x}}_i \cup \mathbf{B}$

$$\mathcal{L}_n(\mathbf{w}) = -\frac{1}{s} \sum_{i=1}^s \log \left(\frac{\sigma(\mathbf{x}_i, \tilde{\mathbf{x}}_i)}{\sigma(\mathbf{x}_i, \tilde{\mathbf{x}}_i) + \sum_{l=1}^b \sigma(\mathbf{x}_i, \mathbf{b}_l)} \right),$$


query positive key negative key

Motivation from Observations

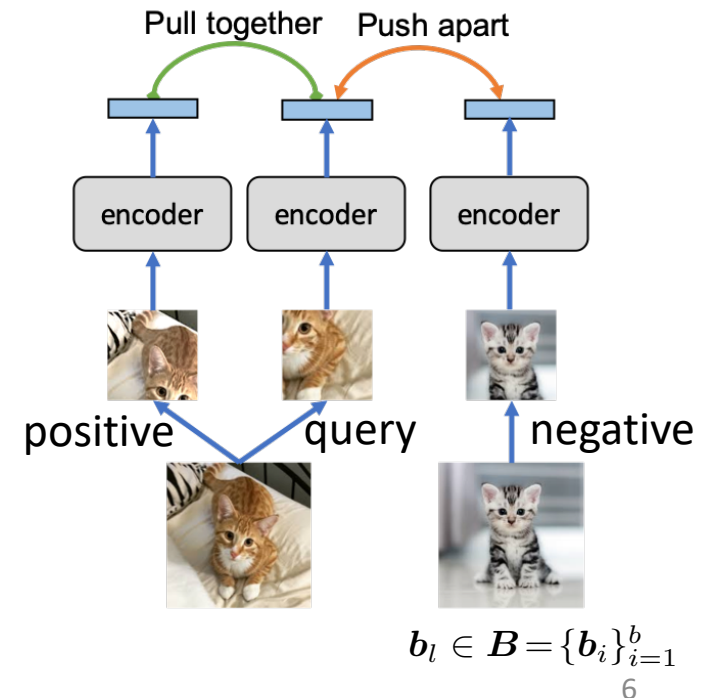
One-hot label assignment: query x_i has only one positive $\tilde{x}_i$ among $\tilde{x}_i \cup B$

$$\mathcal{L}_n(w) = -\frac{1}{s} \sum_{i=1}^s \log \left(\frac{\sigma(x_i, \tilde{x}_i)}{\sigma(x_i, \tilde{x}_i) + \sum_{l=1}^b \sigma(x_i, b_l)} \right),$$

query positive key negative key

Issues of hot label assignment: imprecise & uninformative

- some negatives & query belong to same semantic class



Motivation from Observations

One-hot label assignment: query $\mathbf{x}_i$ has only one positive $\tilde{\mathbf{x}}_i$ among $\tilde{\mathbf{x}}_i \cup \mathbf{B}$

$$\mathcal{L}_n(\mathbf{w}) = -\frac{1}{s} \sum_{i=1}^s \log \left(\frac{\sigma(\mathbf{x}_i, \tilde{\mathbf{x}}_i)}{\sigma(\mathbf{x}_i, \tilde{\mathbf{x}}_i) + \sum_{l=1}^b \sigma(\mathbf{x}_i, \mathbf{b}_l)} \right),$$

query positive key negative key

Issues of hot label assignment: imprecise and and uninformative

- some negatives & query belong to same semantic class
- random augmentations provides crops with different semantic information, e.g. image having several objects

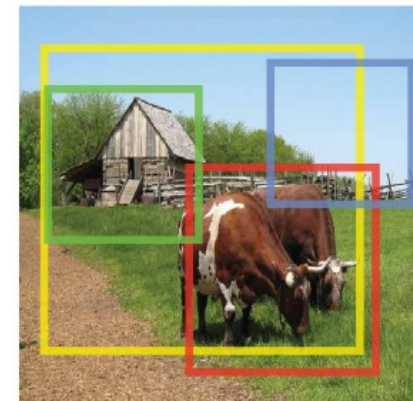


image in ImageNet

↓ Random crop



Motivation from Observations

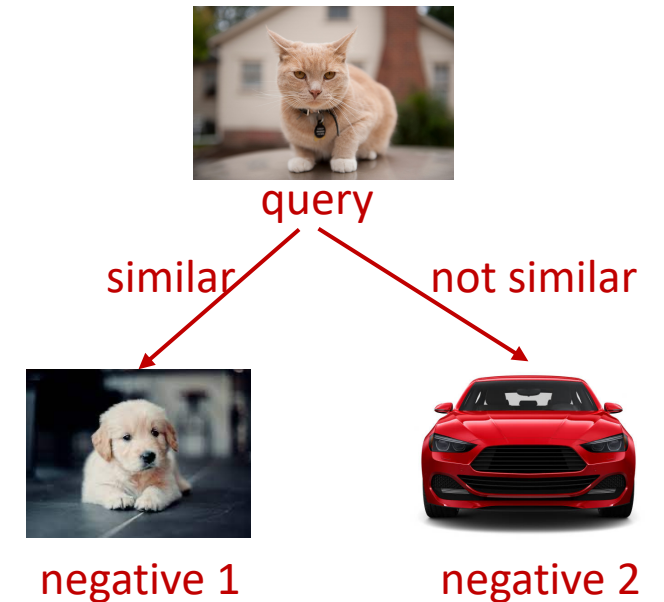
One-hot label assignment: query $\mathbf{x}_i$ has only one positive $\tilde{\mathbf{x}}_i$ among $\tilde{\mathbf{x}}_i \cup B$

$$\mathcal{L}_n(\mathbf{w}) = -\frac{1}{s} \sum_{i=1}^s \log \left(\frac{\sigma(\mathbf{x}_i, \tilde{\mathbf{x}}_i)}{\sigma(\mathbf{x}_i, \tilde{\mathbf{x}}_i) + \sum_{l=1}^b \sigma(\mathbf{x}_i, \mathbf{b}_l)} \right),$$

query positive key negative key

Issues of hot label assignment: imprecise and and uninformative

- some negatives & query belong to same semantic class
- random augmentations provides crops with different semantic information, e.g. image having several objects
- different negatives have different similarity to query



Motivation from Observations

One-hot label assignment: query $\mathbf{x}_i$ has only one positive $\tilde{\mathbf{x}}_i$ among $\tilde{\mathbf{x}}_i \cup \mathbf{B}$

$$\mathcal{L}_n(\mathbf{w}) = -\frac{1}{s} \sum_{i=1}^s \log \left(\frac{\sigma(\mathbf{x}_i, \tilde{\mathbf{x}}_i)}{\underbrace{\sigma(\mathbf{x}_i, \tilde{\mathbf{x}}_i)}_{\text{query}} + \sum_{l=1}^b \underbrace{\sigma(\mathbf{x}_i, \mathbf{b}_l)}_{\text{negative key}}} \right),$$

positive key

Issues of hot label assignment: imprecise and and uninformative

- some negatives & query belong to same semantic class
- random augmentations provides crops with different semantic information, e.g. image having several objects
- different negatives have different similarity to query

Result: one-hot label cannot guarantee semantically similar samples to close

Motivation from Theoretical Analysis

Support that the pair $(x_i, \tilde{x}_i)$ in the training dataset $\mathcal{D} = \{(x_i, \tilde{x}_i)\}_{i=1}^n$ sampled from an unknown distribution $\mathcal{S}$ denotes the positive pair in MoCo.

Assume the query x_i has ground truth soft label $y_i^* \in b + 1$ over the key set $B_i = \{\tilde{x}_i \cup B\}$ where y_{it}^* measures the semantic similarity between x_i and the t-th key b'_t in buffer B_i

Theorem 1 (upper bound of generalization error, informal).

Under proper assumptions, for MoCo, with probability $1 - \nu$, the generalization error on instance discrimination task can be upper bounded as :

$$\boxed{\text{generalization error}} \leq \mathcal{O}(\mathbb{E}_{\mathcal{D} \sim \mathcal{S}} [\|y - y^*\|_2]) + \mathcal{O}\left(\sqrt{\frac{V_{\mathcal{D}} \ln(|\mathcal{F}|/\nu)}{n}} + \frac{\ln(|\mathcal{F}|/\nu)}{n}\right),$$

training & test error gap

where $V_{\mathcal{D}}$ is the variance of f_w on data $\mathcal{D}$, $\mathcal{F}$ is the covering number of encoder f_w

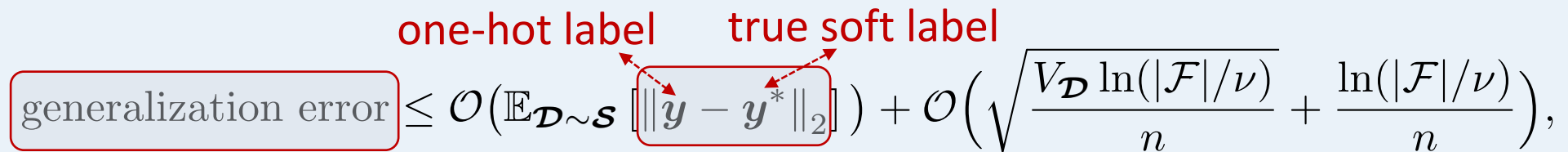
Motivation from Theoretical Analysis

Support that the pair $(x_i, \tilde{x}_i)$ in the training dataset $\mathcal{D} = \{(x_i, \tilde{x}_i)\}_{i=1}^n$ sampled from an unknown distribution $\mathcal{S}$ denotes the positive pair in MoCo.

Assume the query x_i has ground truth soft label $y_i^* \in b + 1$ over the key set $B_i = \{\tilde{x}_i \cup B\}$ where y_{it}^* measures the semantic similarity between x_i and the t-th key b'_t in buffer B_i

Theorem 1 (upper bound of generalization error, informal).

Under proper assumptions, for MoCo, with probability $1 - \nu$, the generalization error on instance discrimination task can be upper bounded as :

$$\boxed{\text{generalization error}} \leq \mathcal{O}(\mathbb{E}_{\mathcal{D} \sim \mathcal{S}} [\|y - y^*\|_2]) + \mathcal{O}\left(\sqrt{\frac{V_{\mathcal{D}} \ln(|\mathcal{F}|/\nu)}{n}} + \frac{\ln(|\mathcal{F}|/\nu)}{n}\right),$$


training & test error gap

where $V_{\mathcal{D}}$ is the variance of f_w on data $\mathcal{D}$, $\mathcal{F}$ is the covering number of encoder f_w

Motivation from Theoretical Analysis

Support that the pair $(x_i, \tilde{x}_i)$ in the training dataset $\mathcal{D} = \{(x_i, \tilde{x}_i)\}_{i=1}^n$ sampled from an unknown distribution $\mathcal{S}$ denotes the positive pair in MoCo.

Assume the query x_i has ground truth soft label $y_i^* \in b + 1$ over the key set $B_i = \{\tilde{x}_i \cup B\}$ where y_{it}^* measures the semantic similarity between x_i and the t-th key b'_t in buffer B_i

Theorem 1 (upper bound of generalization error, informal).

Under proper assumptions, for MoCo, with probability $1 - \nu$, the generalization error on instance discrimination task can be upper bounded as :

$$\boxed{\text{generalization error}} \leq \mathcal{O}(\mathbb{E}_{\mathcal{D} \sim \mathcal{S}} [\|\mathbf{y} - \mathbf{y}^*\|_2]) + \mathcal{O}\left(\sqrt{\frac{V_{\mathcal{D}} \ln(|\mathcal{F}|/\nu)}{n}} + \frac{\ln(|\mathcal{F}|/\nu)}{n}\right),$$

training & test error gap the more accurate of the label y , the better the generalization

where $V_{\mathcal{D}}$ is the variance of f_w on data $\mathcal{D}$, $\mathcal{F}$ is the covering number of encoder f_w

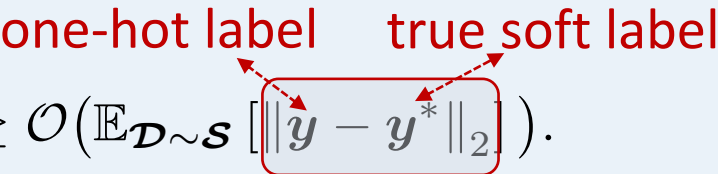
Motivation from Theoretical Analysis

Support that the pair $(x_i, \tilde{x}_i)$ in the training dataset $\mathcal{D} = \{(x_i, \tilde{x}_i)\}_{i=1}^n$ sampled from an unknown distribution $\mathcal{S}$ denotes the positive pair in MoCo.

Assume the query x_i has ground truth soft label $y_i^* \in b + 1$ over the key set $B_i = \{\tilde{x}_i \cup B\}$ where y_{it}^* measures the semantic similarity between x_i and the t-th key b'_t in buffer B_i

Theorem 2 (lower bound of generalization error, informal).

Under proper assumptions, for MoCo, there exists a contrastive learning problem such that the generalization error on instance discrimination task is lower bounded as :

$$\text{generalization error} \geq \mathcal{O}(\mathbb{E}_{\mathcal{D} \sim \mathcal{S}} [\|\mathbf{y} - \mathbf{y}^*\|_2]).$$


Lower and upper bounds show that $\text{generalization error} \sim \mathbb{E}_{\mathcal{D} \sim \mathcal{S}} [\|\mathbf{y} - \mathbf{y}^*\|_2]$.

the more accurate of the label $\mathbf{y}$, the better the generalization

Outline

Motivation: why one-hot label in MoCo is not accurate?

Solution for accurate label: self-labeling refinery and momentum mixup

Experiments: higher classification accuracy

Conclusion

Solution: Self-Labeling Refinement

Self-Labeling Refinement :

- **self-labeling refinery:** soft label replaces one-hot label to directly improve label accuracy
- **momentum mixup:** increase similarity of positive pair to indirectly improve label accuracy

Solution: Self-Labeling Refinement

Self-Labeling Refinement :

- **self-labeling refinery**: soft label replaces one-hot label to directly improve label accuracy
- **momentum mixup**: increase similarity of positive pair to indirectly improve label accuracy

Reformulation of MoCo: for query x_i in minibatch $\{(x_i, \tilde{x}_i)\}_{i=1}^s$, we maximize its similarity to its positive $\tilde{x}_i$ in the key set $\bar{B} = \{\tilde{x}_i\}_{i=1}^s \cup \{b_i\}_{i=1}^b$ and push it away from samples in $\bar{B}$:

$$\mathcal{L}_c(w, \{(x_i, y_i)\}) = -\frac{1}{s} \sum_{i=1}^s \sum_{k=1}^{s+b} y_{ik} \log \left(\frac{\sigma(x_i, \bar{b}_k)}{\sum_{l=1}^{s+b} \sigma(x_i, \bar{b}_l)} \right),$$

where $\bar{b}_k$ is the k-th sample in $\bar{B}$, y_i is the one-hot label of query x_i whose i-th entry y_{ii} is 1.

Solution: Self-Labeling Refinement

Self-Labeling Refinement :

- **self-labeling refinery**: soft label replaces one-hot label to directly improve label accuracy
- **momentum mixup**: increase similarity of positive pair to indirectly improve label accuracy

Reformulation of MoCo: for query x_i in minibatch $\{(x_i, \tilde{x}_i)\}_{i=1}^s$, we maximize its similarity to its positive $\tilde{x}_i$ in the key set $\bar{B} = \{\tilde{x}_i\}_{i=1}^s \cup \{b_i\}_{i=1}^b$ and push it away from samples in $\bar{B}$:

$$\mathcal{L}_c(w, \{(x_i, y_i)\}) = -\frac{1}{s} \sum_{i=1}^s \sum_{k=1}^{s+b} y_{ik} \log \left(\frac{\sigma(x_i, \bar{b}_k)}{\sum_{l=1}^{s+b} \sigma(x_i, \bar{b}_l)} \right),$$

where $\bar{b}_k$ is the k-th sample in $\bar{B}$, y_i is the one-hot label of query x_i whose i-th entry y_{ii} is 1.

Benefit of Reformulation: labels of different samples are defined on a shared dictionary, and thus can be linearly combined.

Solution: Self-Labeling Refinement

Self-Labeling Refinement : (1) self-labeling refinery; (2) momentum mixup

Self-labeling refinery iteratively employs network and data to improve labels during training.

- Step1. for query x_i , we use its positive $\tilde{x}_i$ to estimate semantic similarity between x_i and instances in $\bar{B} = \{\tilde{x}_i\}_{i=1}^s \cup \{b_i\}_{i=1}^b$, since x_i and $\tilde{x}_i$ come from the same image:

$$p_{ik}^t = \sigma^{1/\tau'}(\tilde{x}_i, \bar{b}_k) / \sum_{l=1}^{s+b} \sigma^{1/\tau'}(\tilde{x}_i, \bar{b}_l),$$

Solution: Self-Labeling Refinement

Self-Labeling Refinement : (1) self-labeling refinery; (2) momentum mixup

Self-labeling refinery iteratively employs network and data to improve labels during training.

- Step1. for query x_i , we use its positive $\tilde{x}_i$ to estimate semantic similarity between x_i and instances in $\bar{B} = \{\tilde{x}_i\}_{i=1}^s \cup \{b_i\}_{i=1}^b$, since x_i and $\tilde{x}_i$ come from the same image:

$$p_{ik}^t = \sigma^{1/\tau'}(\tilde{x}_i, \bar{b}_k) / \sum_{l=1}^{s+b} \sigma^{1/\tau'}(\tilde{x}_i, \bar{b}_l),$$

- Step2. as $\tilde{x}_i$ is highly similar to itself in $\bar{B}$, p_{ii}^t will be much larger than others and conceals the similarity of other semantically similar instances in $\bar{B}$. So we remove $\tilde{x}_i$ from $\bar{B}$

$$q_{ik}^t = \sigma^{1/\tau'}(\tilde{x}_i, \bar{b}_k) / \sum_{l=1, l \neq i}^{s+b} \sigma^{1/\tau'}(\tilde{x}_i, \bar{b}_l), \quad q_{ii}^t = 0.$$

Solution: Self-Labeling Refinement

Self-Labeling Refinement : (1) self-labeling refinery; (2) momentum mixup

Self-labeling refinery iteratively employs network and data to improve labels during training.

- Step1. for query x_i , we use its positive $\tilde{x}_i$ to estimate semantic similarity between x_i and instances in $\bar{B} = \{\tilde{x}_i\}_{i=1}^s \cup \{b_i\}_{i=1}^b$, since x_i and $\tilde{x}_i$ come from the same image:

$$p_{ik}^t = \sigma^{1/\tau'}(\tilde{x}_i, \bar{b}_k) / \sum_{l=1}^{s+b} \sigma^{1/\tau'}(\tilde{x}_i, \bar{b}_l),$$

- Step2. as $\tilde{x}_i$ is highly similar to itself in $\bar{B}$, p_{ii}^t will be much larger than others and conceals the similarity of other semantically similar instances in $\bar{B}$. So we remove $\tilde{x}_i$ from $\bar{B}$

$$q_{ik}^t = \sigma^{1/\tau'}(\tilde{x}_i, \bar{b}_k) / \sum_{l=1, l \neq i}^{s+b} \sigma^{1/\tau'}(\tilde{x}_i, \bar{b}_l), \quad q_{ii}^t = 0.$$

Linear combination:

$$\bar{y}_i^t = (1 - \alpha_t - \beta_t) \mathbf{y}_i + \alpha_t \mathbf{p}_i^t + \beta_t \mathbf{q}_i^t,$$

Performance Guarantee: Exact label recovery

Label-corrupted dataset

Let $\{(\mathbf{x}_i, \mathbf{y}_i^*)\}_{i=1}^n$ denote the pairs of crops and ground-truth semantic label

- crop $\mathbf{x}_i$ generated from vanilla sample $\mathbf{c}_t$ obeys $\|\mathbf{x}_i - \mathbf{c}_t\|_2 \leq \varepsilon$
- ground-truth semantic label $\mathbf{y}_i^* \in \{\gamma_t\}_{t=1}^{\bar{K}}$ of $\mathbf{x}_i$ is decided by its corresponding
- the classes are separated: $|\gamma_i - \gamma_k| \geq \delta, \quad \|\mathbf{c}_i - \mathbf{c}_k\|_2 \geq 2\varepsilon, \quad (\forall i \neq k),$
- for each sample $\mathbf{c}_i$, at most ρn_i augmentations are assigned to wrong labels, where n_i denotes the crop sample number of $\mathbf{c}_i$

Performance Guarantee: Exact label recovery

Label-corrupted dataset

Let $\{(\mathbf{x}_i, \mathbf{y}_i^*)\}_{i=1}^n$ denote the pairs of crops and ground-truth semantic label

- crop $\mathbf{x}_i$ generated from vanilla sample $\mathbf{c}_t$ obeys $\|\mathbf{x}_i - \mathbf{c}_t\|_2 \leq \varepsilon$
- ground-truth semantic label $\mathbf{y}_i^* \in \{\gamma_t\}_{t=1}^{\bar{K}}$ of $\mathbf{x}_i$ is decided by its corresponding
- the classes are separated: $|\gamma_i - \gamma_k| \geq \delta, \quad \|\mathbf{c}_i - \mathbf{c}_k\|_2 \geq 2\varepsilon, \quad (\forall i \neq k),$
- for each sample $\mathbf{c}_i$, at most ρn_i augmentations are assigned to wrong labels, where n_i denotes the crop sample number of $\mathbf{c}_i$

Theorem 3 (exact label recovery guarantee, informal).

Under proper assumptions, the discrepancy between the label $\bar{\mathbf{y}}^t$ estimated by our refinery and the true label $\mathbf{y}^*$ of data $\{\mathbf{x}_i\}_{i=1}^n$ is bounded:

$$\frac{1}{\sqrt{n}} \|\overset{\text{estimated label}}{\bar{\mathbf{y}}^t} - \underset{\text{true soft label}}{\mathbf{y}^*}\|_2 \leq \frac{1 - \alpha_t}{\sqrt{n}} \|\mathbf{y} - \mathbf{y}^*\|_2 + \alpha_t(6\rho + \zeta),$$

where $\zeta = c_6 \varepsilon K^2 \Gamma^5 \xi_3^3 \sqrt{\log \bar{K}} / \lambda^2(\mathbf{C})$, $\mathbf{y}^* = [\mathbf{y}_1^*, \dots, \mathbf{y}_n^*]$

Performance Guarantee: Exact label recovery

Theorem 3 (exact label recovery guarantee, informal).

Under proper assumptions, the discrepancy between the label $\bar{\mathbf{y}}^t$ estimated by our refinery and the true label $\mathbf{y}^*$ of data $\{\mathbf{x}_i\}_{i=1}^n$ is bounded:

$$\frac{1}{\sqrt{n}} \|\overset{\text{estimated label}}{\bar{\mathbf{y}}^t} - \underset{\text{true soft label}}{\mathbf{y}^*}\|_2 \leq \frac{1 - \alpha_t}{\sqrt{n}} \|\mathbf{y} - \mathbf{y}^*\|_2 + \alpha_t(6\rho + \zeta),$$

where $\zeta = c_6 \varepsilon K^2 \Gamma^5 \xi_3^3 \sqrt{\log K} / \lambda^2(\mathbf{C})$, $\mathbf{y}^* = [\mathbf{y}_1^*, \dots, \mathbf{y}_n^*]$

If $\rho \leq \frac{\delta}{24}$, $(1 - \alpha_0)|\mathbf{y}_i - \mathbf{y}_i^*| + \frac{1}{3}\alpha_0\delta < \frac{1}{2}\delta$, the estimated label $\bar{\mathbf{y}}_i^t$ predicts true label of $\mathbf{y}_i^*$ any crop $\mathbf{x}_i$

Exact label recovery: $\gamma_{k^*} = \mathbf{y}_i^*$ with $k^* = \operatorname{argmin}_{1 \leq k \leq \bar{K}} |\bar{\mathbf{y}}_i^t - \gamma_k|$.

Performance Guarantee: Exact label recovery

Theorem 3 (exact prediction of network when using label refinery, informal)

Under proper assumptions, by using the refined label $\bar{\mathbf{y}}^t$ to train network, the error of network prediction on $\{\mathbf{x}_i\}_{i=1}^n$ is upper bounded

$$\frac{1}{\sqrt{n}} \|f(\mathbf{W}_t, \mathbf{X}) - \mathbf{y}^*\|_2 \leq 6\rho + \frac{\zeta \lambda(\mathbf{C})}{K \Gamma^2 \xi_3^2},$$

predicted label true label

where $\zeta = c_6 \varepsilon K^2 \Gamma^5 \xi_3^3 \sqrt{\log K} / \lambda^2(\mathbf{C})$, $\mathbf{y}^* = [\mathbf{y}_1^*, \dots, \mathbf{y}_n^*]$

If $\rho \leq \frac{\delta}{24}$, $(1 - \alpha_0)|\mathbf{y}_i - \mathbf{y}_i^*| + \frac{1}{3}\alpha_0\delta < \frac{1}{2}\delta$, for any vanilla sample $\mathbf{c}_k$, network $f(\mathbf{W}_t, \cdot)$ predicts the true semantic label $\mathbf{y}_i^*$ of any augmentation $\mathbf{x}$ that obeys $\|\mathbf{x} - \mathbf{c}_k\|_2 \leq \varepsilon$:

Exact label prediction: $\gamma_{k^*} = \gamma_k$ with $k^* = \operatorname{argmin}_{1 \leq i \leq \bar{K}} |f(\mathbf{W}_t, \mathbf{x}) - \gamma_i|$.

Solution: Self-Labeling Refinement

Self-Labeling Refinement : (1) self-label refinery; (2) momentum mixup

Momentum mixup constructs virtual instance as follows:

$$\mathbf{x}'_i = \theta \mathbf{x}_i + (1 - \theta) \tilde{\mathbf{x}}_k, \quad \mathbf{y}'_i = \theta \bar{\mathbf{y}}_i + (1 - \theta) \bar{\mathbf{y}}_k, \quad (1)$$

where $\tilde{\mathbf{x}}_k$ is randomly sampled from the key set $\{\tilde{\mathbf{x}}_i\}_{i=1}^s$, $\bar{\mathbf{y}}_i$ denotes the refined label by self-labeling refinery, $\theta \in [0, 1]$ obeys the beta distribution

Benefits: the component $\tilde{\mathbf{x}}_k$ in $\mathbf{x}'_i = \theta \mathbf{x}_i + (1 - \theta) \tilde{\mathbf{x}}_k$ directly increases the similarity between the query $\mathbf{x}'_i$ and its positive key $\tilde{\mathbf{x}}_k$ in $\bar{B}$

momentum mixup can improve the accuracy of the label

Outline

Background: what is contrastive learning and why label is not accurate?

Solution: self-labeling refinery and momentum mixup

Experiments: higher classification accuracy

Conclusion

Outline

Motivation: why one-hot label in MoCo is not accurate?

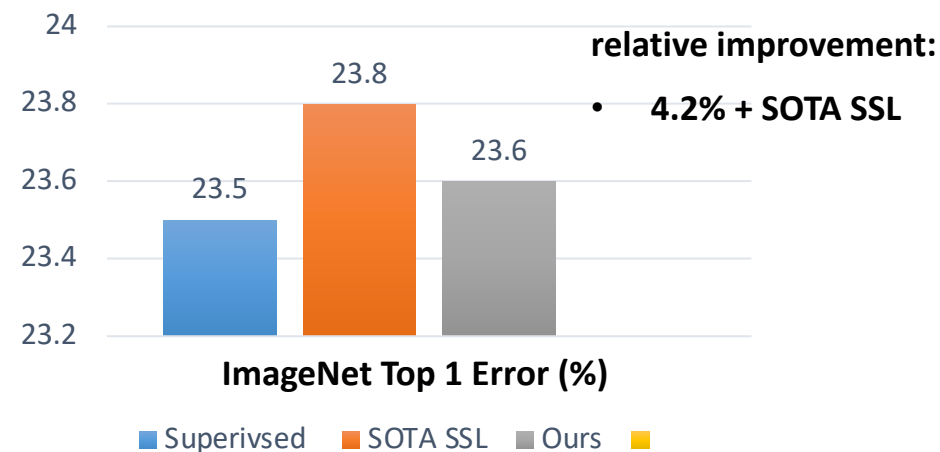
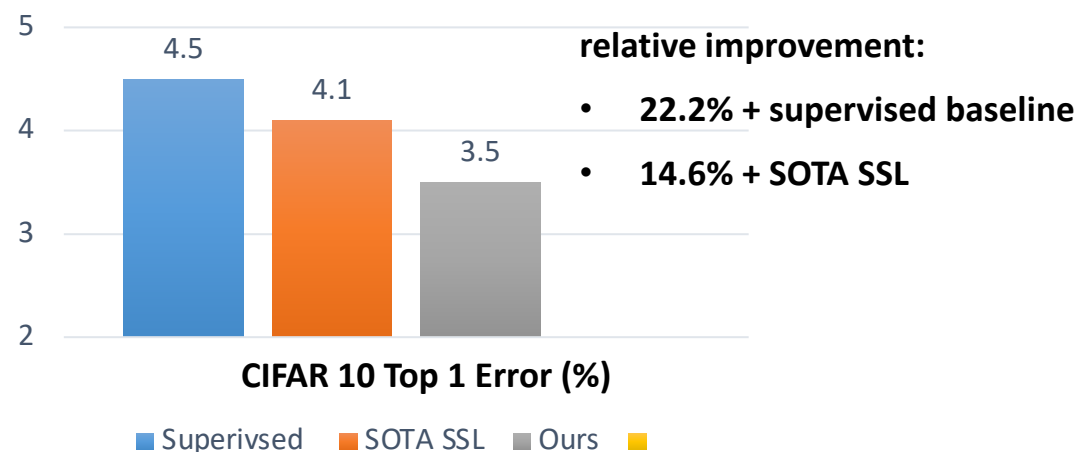
Solution for accurate label: self-labeling refinery and momentum mixup

Experiments: higher classification accuracy

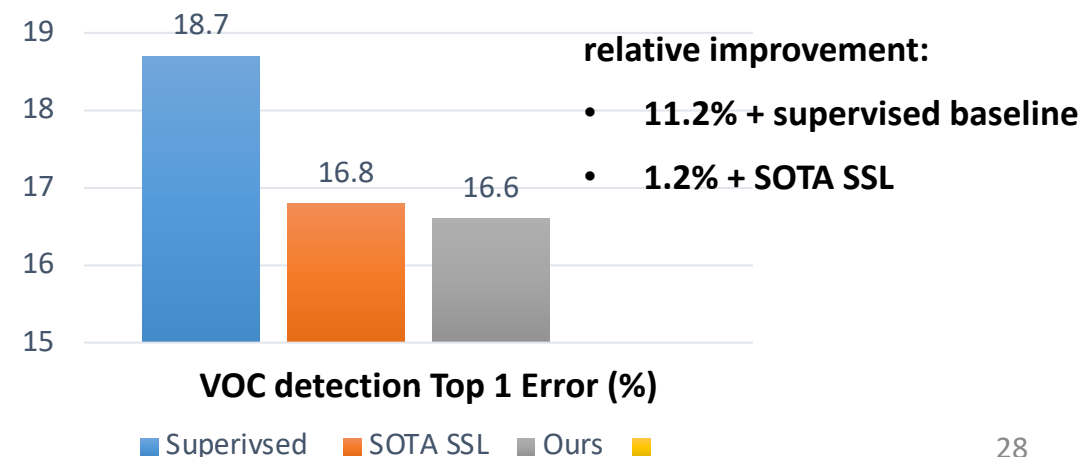
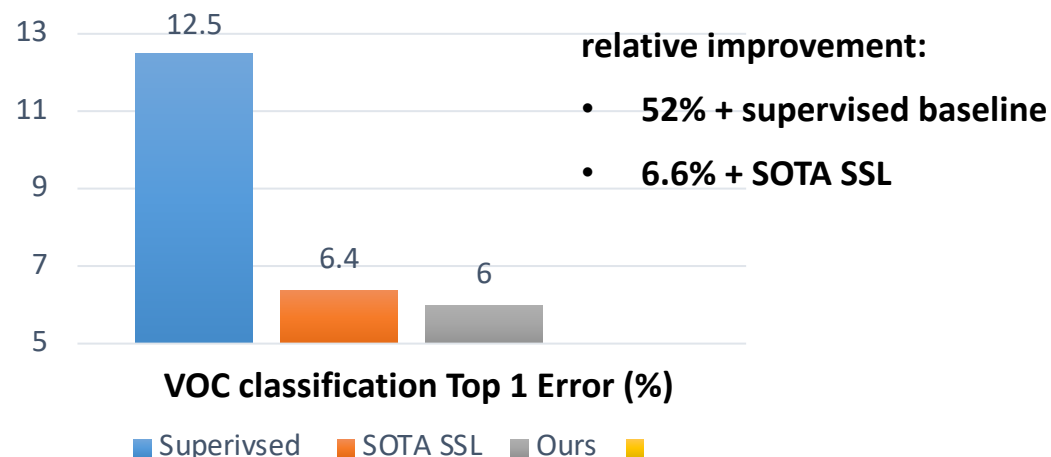
Conclusion

Experimental Results of Proposed NAS Method

CIFAR10 and ImageNet: much lower Top-1 error



Downstream tasks: higher performance on VOC classification and detection



Outline

Motivation: why one-hot label in MoCo is not accurate?

Solution for accurate label: self-labeling refinery and momentum mixup

Experiments: higher classification accuracy

Conclusion

Conclusion

- **Problems:**

(1) what relationship between one-hot label and the generalization performance?

more precise of labels in contrastive learning, the better the generalization

(2) how to estimate more precise labels?

we propose self-labeling refinery and momentum mixup

Thanks!