



Multilabel reductions: what is my loss optimising?

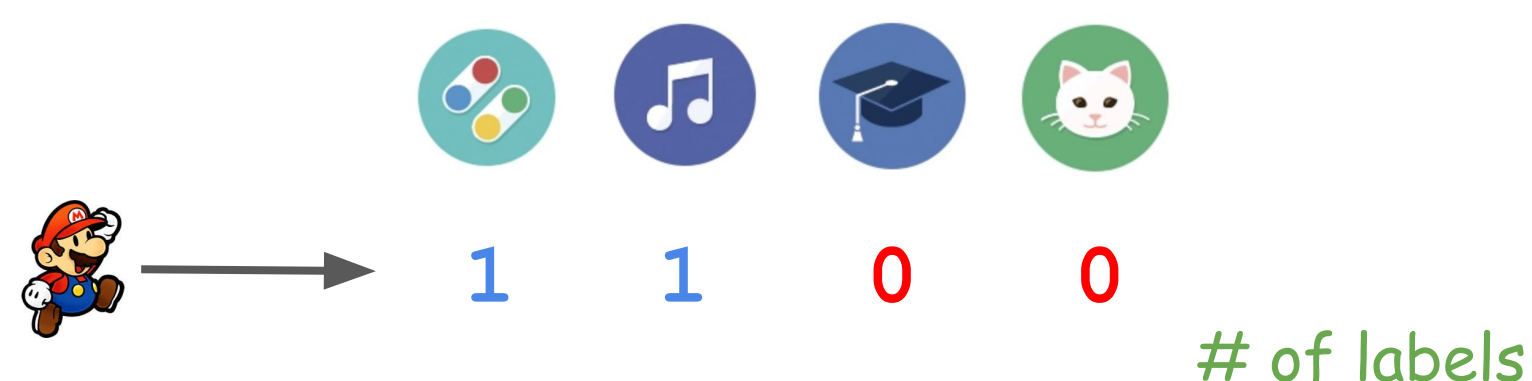
Aditya Krishna Menon, Ankit Singh Rawat, Sashank J. Reddi, Sanjiv Kumar

Google Research NYC

Five common multilabel reductions implicitly optimise for precision or recall

Multilabel classification via reductions

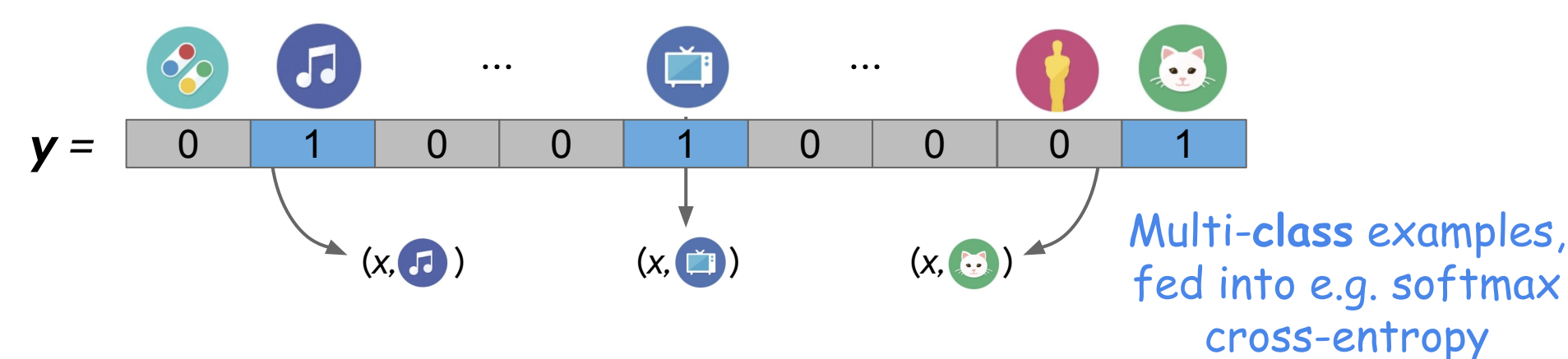
Multilabel classification: predict a **binary label vector**



Ideally, want $\mathbf{f}_i(x)$ **high** if $\mathbf{y}_i = 1$. The **challenge** is that L may be **large**; so how do we **efficiently** find such an $\mathbf{f}$?

Common algorithms **reduce** the problem to **binary** or **multiclass** learning; e.g., we may create

- multiple **binary** examples, one for **each label**
- multiple **multi-class** examples, one for **each positive label**
- one **multi-class** example for a **random positive label**



Key question of this work

Such reductions can be practically effective, but:

what multilabel metric do they optimise?

Answering this helps inform the choice of reduction, depending on our end goal.

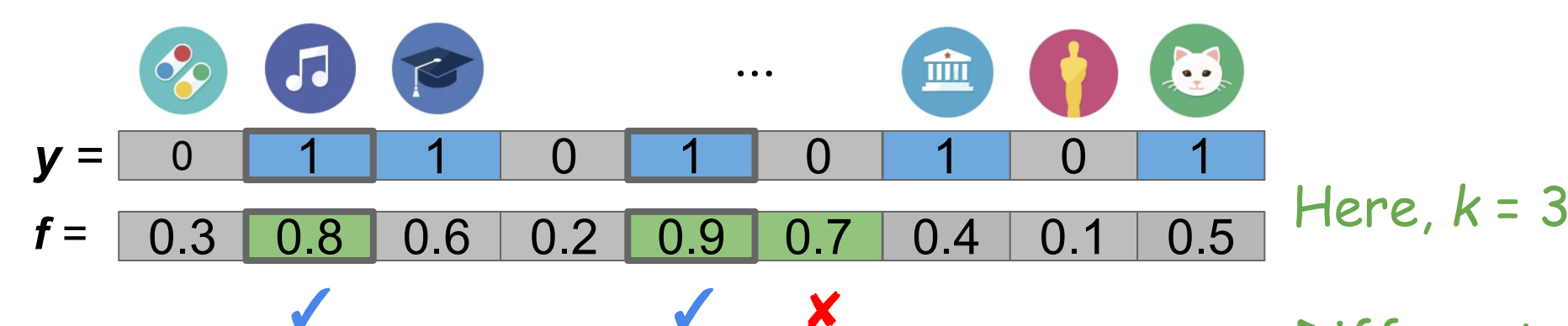
Key contribution of this work

five common reductions implicitly optimise for either precision or recall@k!

To begin, we study a basic property of these metrics.

Multilabel metrics: precision and recall

Precision and **recall@k** are standard metrics for **retrieval** settings. They measure the **# of positives** in the **top-k scoring indices**, suitably normalised:



$$\text{Prec@k}(\mathbf{f}) = \mathbb{E}_{(x, \mathbf{y})} [|\text{top}_k(\mathbf{f}(x)) \cap \text{pos}(\mathbf{y})| / k]$$

$$\text{Rec@k}(\mathbf{f}) = \mathbb{E}_{(x, \mathbf{y})} [|\text{top}_k(\mathbf{f}(x)) \cap \text{pos}(\mathbf{y})| / |\text{pos}(\mathbf{y})|]$$

Fact: Given a distribution $P(x, \mathbf{y})$, the optimal $\mathbf{f}^*$ for these metrics preserve the ordering of:

$$(\text{Prec@k}) P(\mathbf{y}_i = 1 | x) \rightarrow \text{Marginal relevance}$$

$$(\text{Rec@k}) P(\mathbf{y}'_i = 1 | x) = P(\mathbf{y}_i = 1 | x) \cdot \mathbb{E} [(1 + \sum_{j \neq i} \mathbf{y}_j)^{-1}] \rightarrow \text{Non-constant weighting}$$

Note that owing to the weighting above, these optimal solutions are **incompatible** in general!

Implicit losses for multilabel reductions

To compare different reductions, we explicate their **implicit multilabel** losses, i.e., $\ell(\mathbf{y}, \mathbf{f}(x))$:

One-versus-all (OVA) $\sum_{i \in [L]} \ell_{\text{BC}}(\mathbf{y}_i, \mathbf{f}_i)$ $\rightarrow$ Binary loss

Pick-all-labels (PAL) $\sum_{i \in [L]} \mathbf{y}_i \cdot \ell_{\text{MC}}(i, \mathbf{f})$ $\rightarrow$ Multiclass loss; has label "competition"

OVA normalised $\sum_{i \in [L]} \left\{ \frac{\mathbf{y}_i}{\sum_{j \in [L]} \mathbf{y}_j} \cdot \ell_{\text{BC}}(1, \mathbf{f}_i) + \left(1 - \frac{\mathbf{y}_i}{\sum_{j \in [L]} \mathbf{y}_j} \right) \cdot \ell_{\text{BC}}(0, \mathbf{f}_i) \right\}$

PAL normalised Pick-one-label (POL) $\sum_{i \in [L]} \frac{\mathbf{y}_i}{\sum_{j \in [L]} \mathbf{y}_j} \cdot \ell_{\text{MC}}(i, \mathbf{f})$ $\rightarrow$ For log-loss, $\text{KL}(\text{label}, \text{prediction})$

Optimal scores for multilabel reductions

Equipped with these losses, observe that, e.g.,

$$\text{Prec@k}(\mathbf{f}) = \mathbb{E}_{(x, \mathbf{y})} [\sum_{i \in [L]} \ell_{\text{top-k}}(\mathbf{y}, \mathbf{f}_i(x)) / k],$$

where $\ell_{\text{top-k}}$ is a "top-k" multiclass loss; i.e., **precision is equivalent to PAL** for a specific loss! More generally:

Fact: Given a distribution $P(x, \mathbf{y})$, the optimal $\mathbf{f}_i^*$ is:

(OVA) $P(\mathbf{y}_i = 1 | x)$ $\rightarrow$ Classical Expected # of +ves Doesn't vary with i

(PAL) $P(\mathbf{y}_i = 1 | x) / N(x)$ $\rightarrow$ c.f. recall@k optimal scorer

(Rest) $P(\mathbf{y}'_i = 1 | x)$ $\rightarrow$ c.f. recall@k optimal scorer

Different reductions target either precision or recall! Subtle differences in the loss thus have non-trivial effects; e.g., $\ell_{\text{PAL-Norm}}(\mathbf{y}, \mathbf{f}(x)) = (\sum_i \mathbf{y}_i)^{-1} \ell_{\text{PAL}}(\mathbf{y}, \mathbf{f}(x))$, but the optimal solutions for the two are **not** scalings of each other. Further remarks:

- PAL scores are **not calibrated** across instances!
- PAL may $\geq$ OVA since it directly bounds top-k loss
- PAL with a **multiclass OVA** loss $\neq$ **multilabel OVA**; PAL implicitly places a greater weight on each "negative" label

Illustration of differences in reductions

Synthetic problem where normalised reduction gives recall gains, at the expense of precision.

