

Beijing Forest Studio
北京理工大学信息系统及安全对抗实验中心



无监督数据增强研究

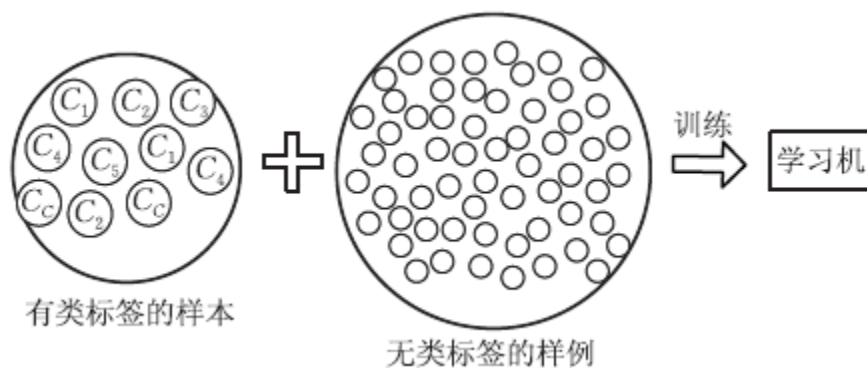
博士研究生 郝靖伟

2019年07月21日

- 背景简介
- 基本概念
 - 半监督学习
 - 强化学习
 - 损失函数
 - 数据增强
- 算法原理
 - AutoAugment
 - 训练信号退火 (TSA)
- 优劣分析
- 应用总结
- 参考文献

- 预期收获
 - 1. 了解机器学习中的一些基础知识
 - 2. 了解半监督学习的最新研究进展
 - 3. 理解最新的自动化数据增强方法
 - 4. 了解引入特殊训练技巧（TSA）的无监督数据增强方法对半监督学习的革新效果

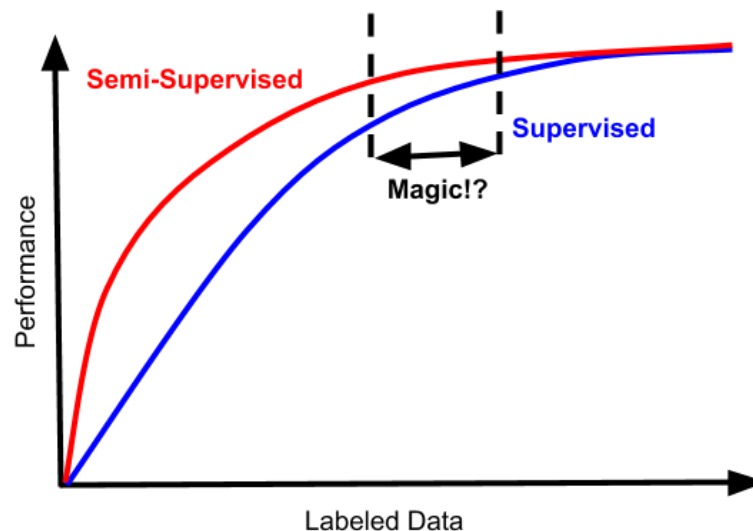
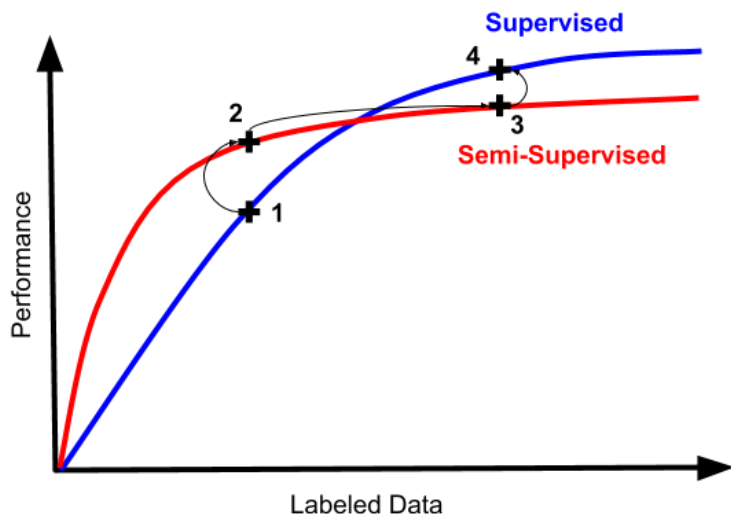
- 半监督学习进展
 - Semi-Supervised Learning, SSL
 - 使用标记数据以及大量的未标记数据来进行模式识别工作
 - 有标注的数据不多，未标注数据很多
 - 做数据标注的资源有限



• 半监督学习进展

聊胜于无

- 标注数据不多时，半监督学习可以带来一定的性能提升
- 问题：需要额外的资源代价，面对更多标注数据时，性能曲线增长平缓【无标注数据带来偏倚】



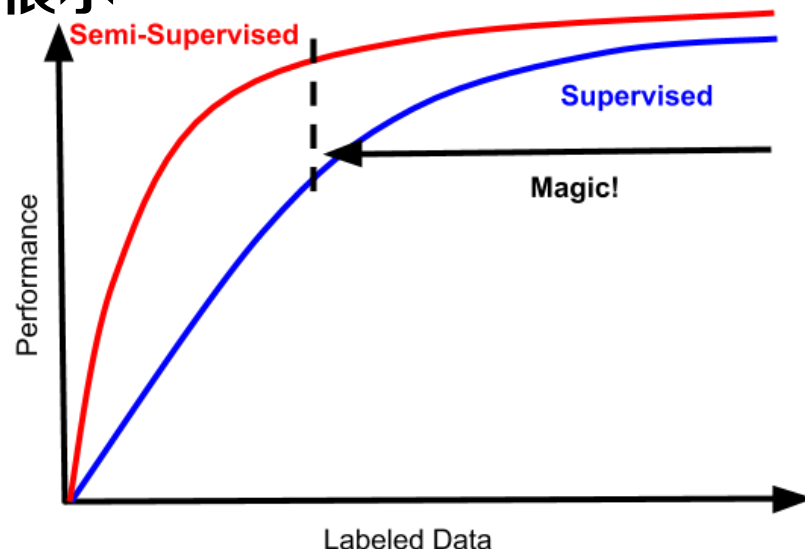
- 某一个数据规模之内半监督学习的效果会好一些，确实提高了数据使用效率
- 问题：很难达到且提升不多

• 半监督学习进展

- 更多数据有更好的性能
- 对于同样的有标注数据，性能总是比监督学习方法更好
- 即便是数据量足够大、监督学习已经能够发挥出好的效果的范围内，半监督学习也仍然有提升
- 付出的额外计算复杂度和资源很小
- 不受到数据规模限制

发生了什么？

可堪大用



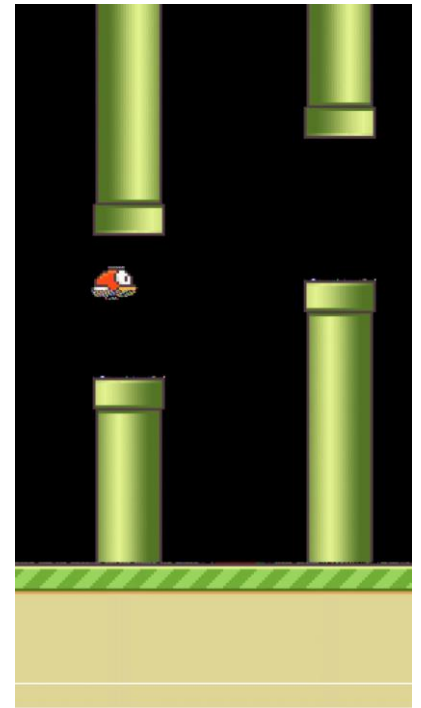
半监督学习的性能提升曲线

- 有监督学习
 - 从有标记的训练数据中推导出预测函数
 - 【给定数据，预测标签】
- 无监督学习
 - 从无标签的数据中学习出一些有用的模式
 - 【给定数据，寻找隐藏的结构】
- 半监督学习
 - 侧重在有监督的分类算法中加入无标记样本
- 强化学习
 - 强调如何基于环境而行动，以取得最大化的预期利益
 - 【给定数据，学习如何选择一系列行动】

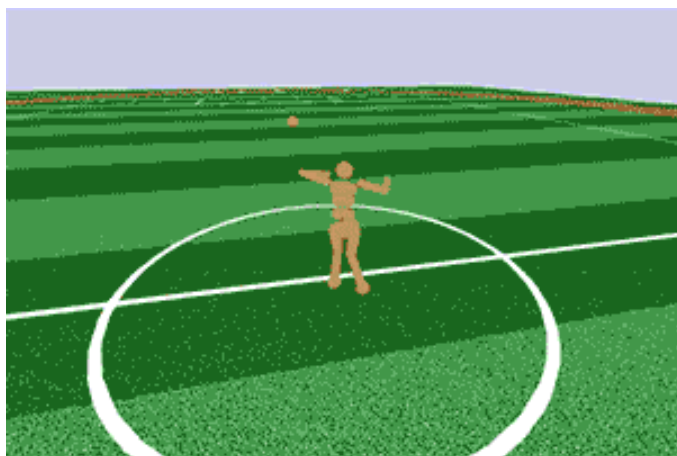
热爱学习的我又出现了



- 强化学习
 - Q-learning
 - 一种基于值函数估计的强化学习方法（以value推policy）
 - 学习奖惩的值，根据自己认为的最高价值选行为
 - Policy Gradient
 - 一种策略搜索强化学习方法（直推policy）
 - 基于策略做梯度下降从而优化模型



- 强化学习
 - 策略梯度法 (Policy Gradient)
 - 问题：对步长大小非常敏感，效率较低
 - 近端策略优化 (Proximal Policy Optimization)
 - 每一步迭代计算一个更新以最小化成本函数
 - 在计算梯度时确保与先前策略有相对较小的偏差
 - 在难易程度、采样复杂度、调试损耗上取得平衡



利用PPO训练的机器人



具有互动能力的机器人

- 损失函数

- 定义：量化描述模型的预测值 $f(x)$ 与真实值 Y 的不一致程度，一个非负实值函数
- 交叉熵（Cross Entropy Loss）：分类问题常用的损失函数

$$L = - \sum_{c=1}^M y_c \log(p_c)$$

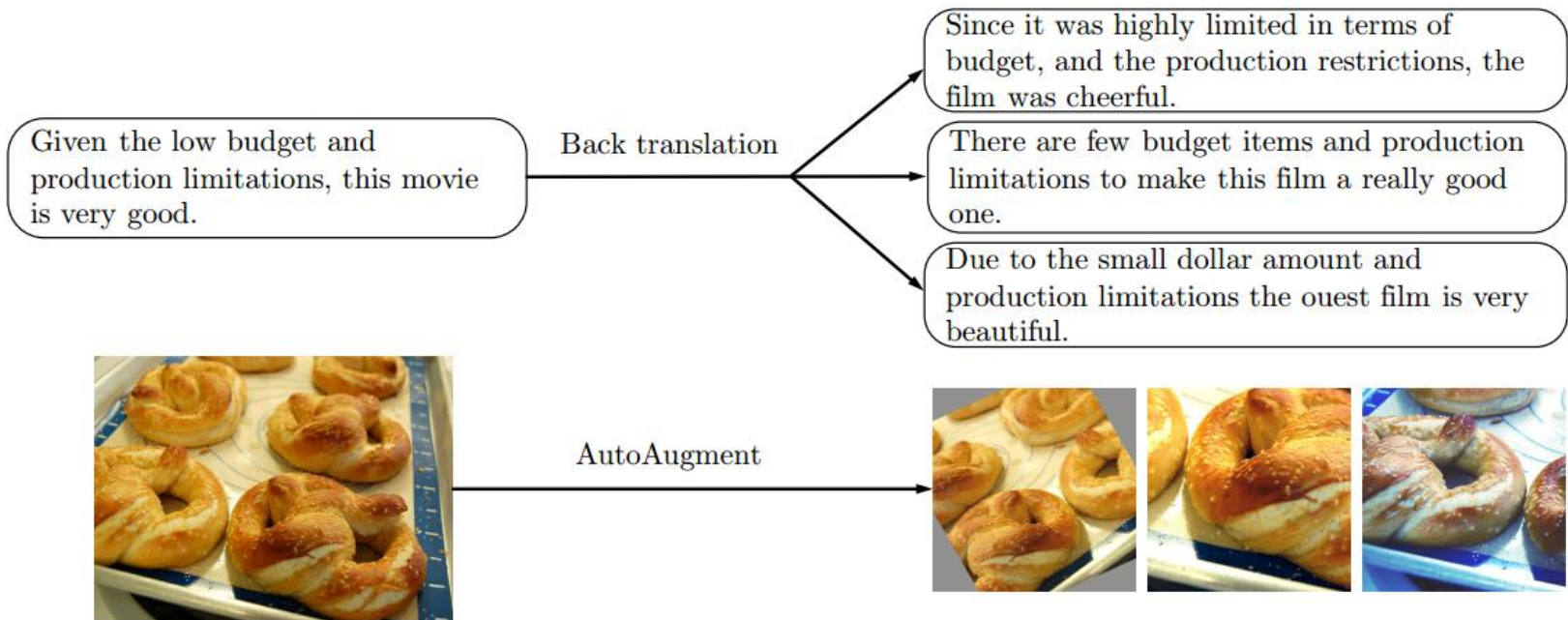
- KL散度（KL divergence）：相对熵，衡量两个概率分布之间的差异

$$D_{\text{KL}}(P\|Q) = \sum_i P(i) \ln \frac{P(i)}{Q(i)}$$

- 一致性训练
 - 主要思想：对于一个输入，即使受到微小干扰，其预测都应该一致
- 过拟合（Overfitting）
 - 定义：由于训练数据包含抽样误差，训练时，复杂的模型将抽样误差也考虑在内，将抽样误差也进行了很好的拟合
 - 表现：模型在训练集上效果好，在测试集上效果差
 - 正则化方法：
 - 修改loss函数
 - 修改网络本身架构（Dropout）
 - 增加数据

- 数据增强

- 通过随机「增广」来提高数据量和数据多样性的策略
- 对现有的数据集进行微小的改变，被认为是不同的输入
- ‘新’的样本标签保持不变
- 使得训练的模型具有更强的**泛化能力**



- 数据增强

- 翻转: 水平翻转或者上下翻转图像
- 旋转: 将图像在一定角度范围内旋转
- 缩放: 将图像在一定尺度内放大或者缩小
- 裁剪: 在原有图像上裁剪出一块
- 平移: 将图像在一定尺度范围内平移
- 颜色变动: 对图像的RGB颜色空间进行一些变换
- 噪声扰动: 给图像加入一些人工生成的噪声

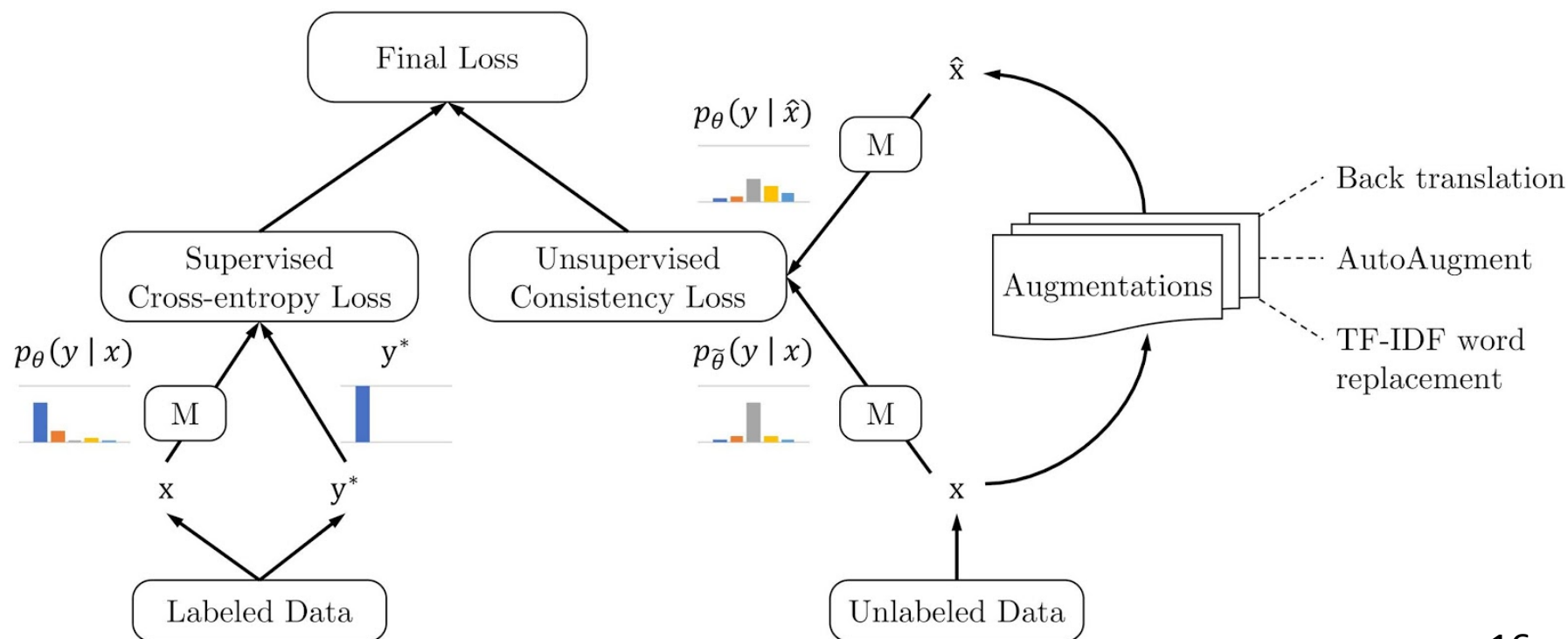


- 数据增强
 - 同义词替换：在句子中随机选取 n 个非停用词。对每个选取的词，用它的随机选取的同义词替换。
 - 随机插入：在句子中任意找一个非停用词，随机选一个它的同义词，插入句子中的任意位置。重复 n 次
 - 随机交换：任意选取句子中的两个词，交换位置。重复 n 次。
 - 随机删除：对于句子中概率为 p 的每一个词，随机删除

T	标记数据稀缺时，有效利用域外未标记数据，显著改善半监督学习（SSL）效果
I	标记数据和未标记数据
P	1.标记数据使用交叉熵计算损失函数以训练模型 2.未标记数据使用一致性训练来强制预测未标记的样本和增强的未标记样本是否相似 3.联合优化标记数据的监督损失和未标记数据的无监督一致性损失
O	低错误率的半监督预测模型

P	数据增强不破坏真实标签；标记数据少，模型易过拟合
C	智能数据增强技术
D	神经网络对增强之后的未标记数据作自洽预测
L	CCFA类 【CVPR2019】

- Unsupervised Data Augmentation (2019.07.10)
 - 无监督数据增强，半监督学习的state-of-the-art
 - 对无标注数据执行数据增强，从而显著提高了半监督学习(SSL)的性能，因此研究人员相信"半监督学习将再度兴起"



• UDA

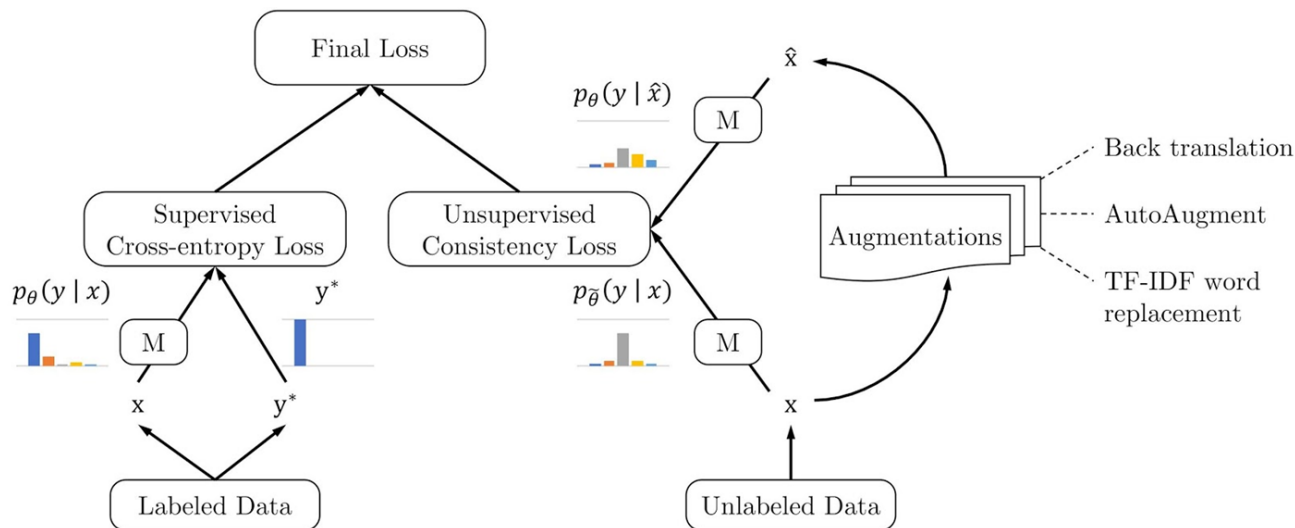
– 标记数据：交叉熵损失函数-训练模型M

– 无标记数据：KL散度-相同模型M

{

 未标记的样本
 增强的未标记样本

$$\min_{\theta} \mathcal{J}_{\text{UDA}}(\theta) = \mathbb{E}_{x \in U} \mathbb{E}_{\hat{x} \sim q(\hat{x}|x)} [\mathcal{D}_{\text{KL}}(p_{\tilde{\theta}}(y|x) \parallel p_{\theta}(y|\hat{x}))]$$

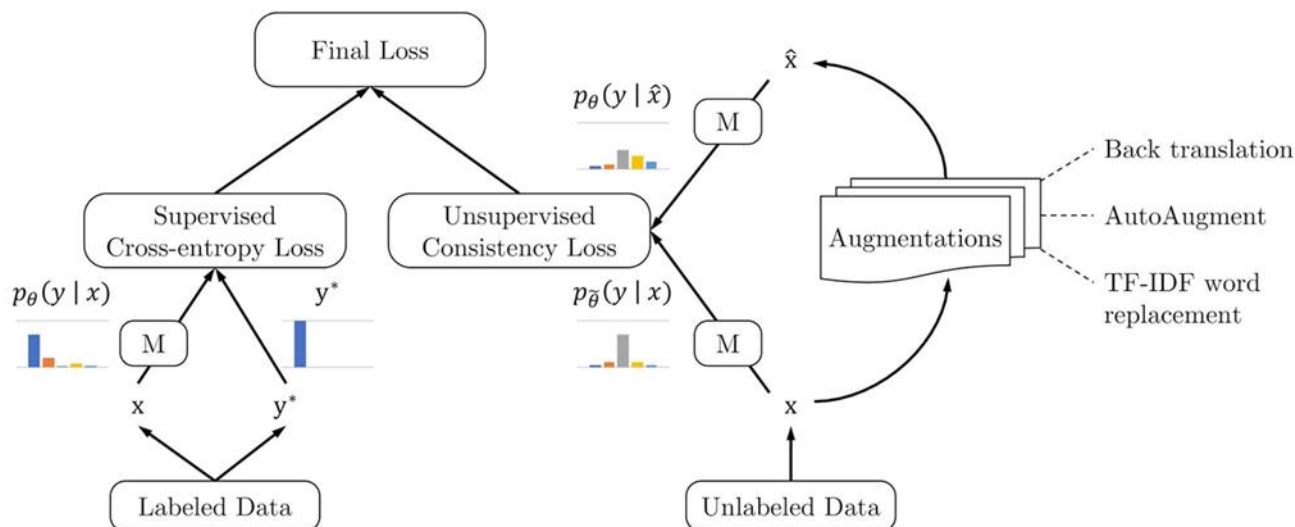


- UDA

- UDA联合优化标记数据的监督损失和未标记数据的无监督一致性损失，计算最终损失

$$\min_{\theta} \mathcal{J} = \mathbb{E}_{x, y^* \in L} [p_{\theta}(y^* | x)] + \lambda \mathcal{J}_{\text{UDA}}(\theta);$$

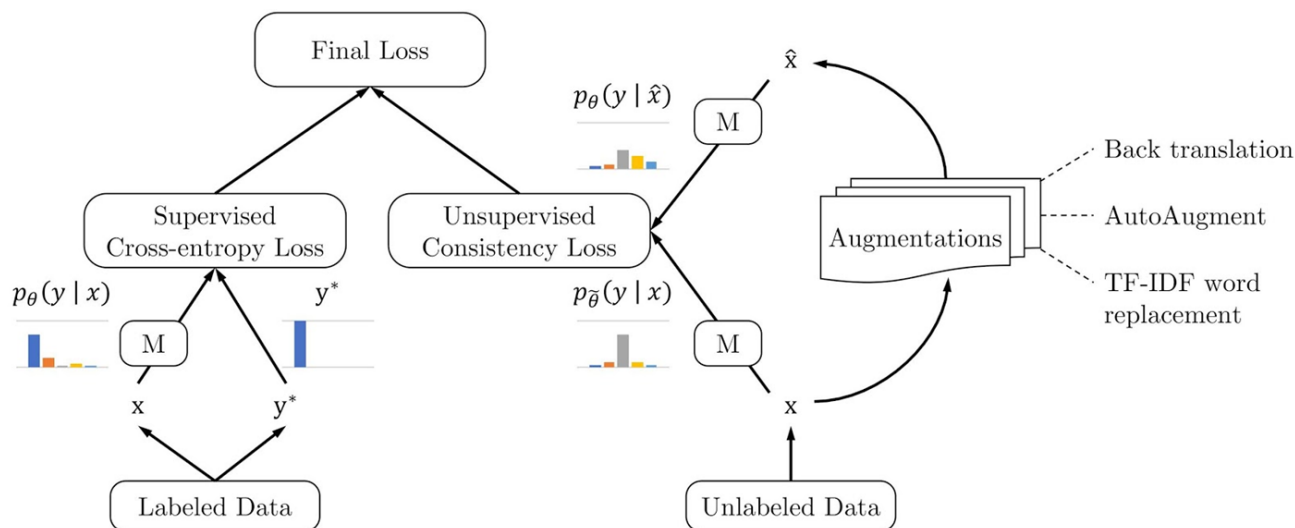
- 标签信息从标记的样本平滑地传播到未标记的样本



• UDA

创新点

- 集成了最新数据增强方法
 - 反向翻译 (Back translation)
 - 自动增强 (AutoAugment)
 - TF-IDF单词替换 (TF-IDF word replacement)
- 提出了TSA训练信号退火策略

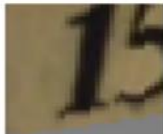
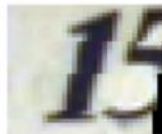


- **AutoAugment (CVPR2019)**
 - **问题：**数据集出现的对称性不同，特定数据集的数据增强方法不能有效地迁移到其他数据集
 - **目标：**为目标数据集寻找有效数据增强策略的自动化过程
 - **主要思想：**创建数据增强策略的搜索空间，并直接在一些数据集上评估特定策略的质量
 - **实验结果：**产生难度更高、更真实的噪声；学习到的策略可迁移并且表现良好

- AutoAugment

- 确定搜索空间

- 每一个策略由5个子策略组成
 - 每个mini-batch中的每张图像随机选择一个子策略
 - 子策略按顺序包含两个操作
 - 每个操作与两个超参数相关：执行操作的概率和幅度

	Original	Sub-policy 1	Sub-policy 2	Sub-policy 3	Sub-policy 4	Sub-policy 5
Batch 1						
Batch 2						
Batch 3						
		ShearX, 0.9, 7 Invert, 0.2, 3	ShearY, 0.7, 6 Solarize, 0.4, 8	ShearX, 0.9, 4 AutoContrast, 0.8, 3	Invert, 0.9, 3 Equalize, 0.6, 3	ShearY, 0.8, 5 AutoContrast, 0.7, 3

- AutoAugment
 - 确定搜索空间
 - 16个操作、幅度范围离散为10个值【0~9】、概率离散为11个值【0.0~1.0】
 - 每个子策略有 $(16 \times 10 \times 11)^2$ 搜索空间，同时5个子策略有 $(16 \times 10 \times 11)^{10}$ 种搜索空间

Operation Name	Description	Range of magnitudes
ShearX(Y)	Shear the image along the horizontal (vertical) axis with rate <i>magnitude</i> .	[-0.3,0.3]
TranslateX(Y)	Translate the image in the horizontal (vertical) direction by <i>magnitude</i> number of pixels.	[-150,150]
Rotate	Rotate the image <i>magnitude</i> degrees.	[-30,30]
AutoContrast	Maximize the the image contrast, by making the darkest pixel black and lightest pixel white.	
Invert	Invert the pixels of the image.	
Equalize	Equalize the image histogram.	
Solarize	Invert all pixels above a threshold value of <i>magnitude</i> .	[0,256]

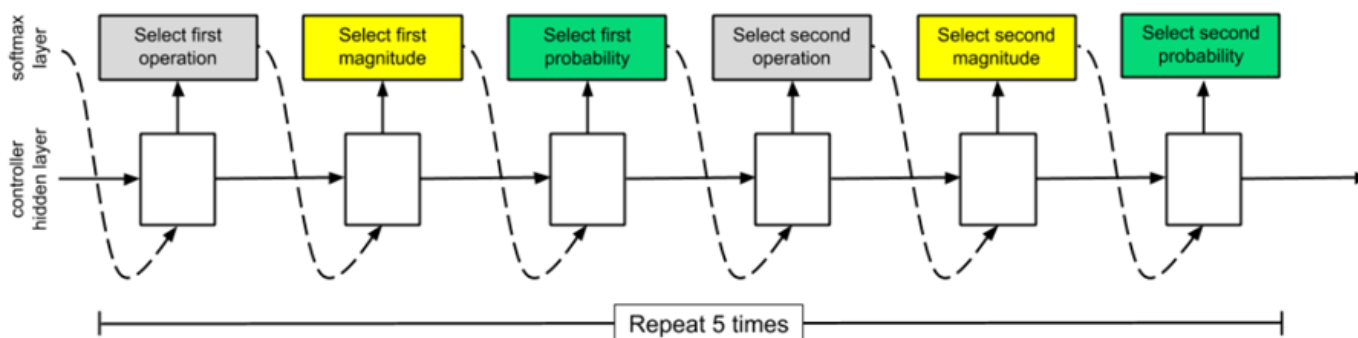
- AutoAugment

可迁移

- 搜索算法寻找最佳策略

- 训练方法

- 控制器，决定当前哪个增强策略效果最好，并通过在特定数据集的一个子集上运行子实验来测试该策略的泛化能力
 - 子实验完成后，采用策略梯度法PPO，以验证集的准确度作为更新信号对控制器进行更新



- AutoAugment 实验结果
 - 对于街景(SVHN)图像, 重点是剪切和平移等几何变换
 - 在CIFAR-10和ImageNet上, 重点是调整颜色和色调分布, 同时保持一般的色彩属性



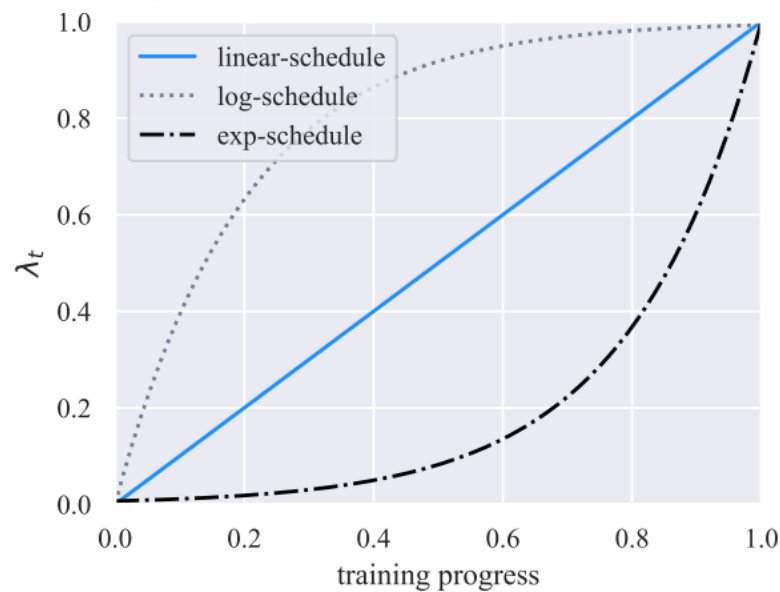
- 在CIFAR-10上, 获得了1.48%的错误率
- 在SVHN上, 将最先进的误差从1.30%降低到1.02%

- 训练信号退火 (Training Signal Annealing)
 - 问题：许多未标注的样本和少量标注的样本，需要一个大型的模型捕获未标注样本中的信息，导致模型过拟合
 - 目标：防止对标记数据过拟合
 - 策略：随着模型对越来越多的未标记样本进行训练，缓慢释放labeled data信息
 - 在 CIFAR-10 上，TSA 将测试误差从 5.67% 降到了 5.10%

- 训练信号退火 (Training Signal Annealing)
 - 具体做法：不计对标记数据预测过于自信的样本
 - 这部分标记数据的误差无法反向传递
 - 避免模型进一步过拟合到这些样本
 - 随着训练进行，阈值 η_t 增加，使用更多标记样本

$$\min_{\theta} \frac{1}{Z} \sum_{x, y^* \in B} [-I(p_{\theta}(y^* | x) < \eta_t) \log p_{\theta}(y^* | x)]$$

$$\eta_t = \frac{1}{K} + \lambda_t * \left(1 - \frac{1}{K}\right)$$



- 算法执行结果
 - <https://github.com/google-research/uda>

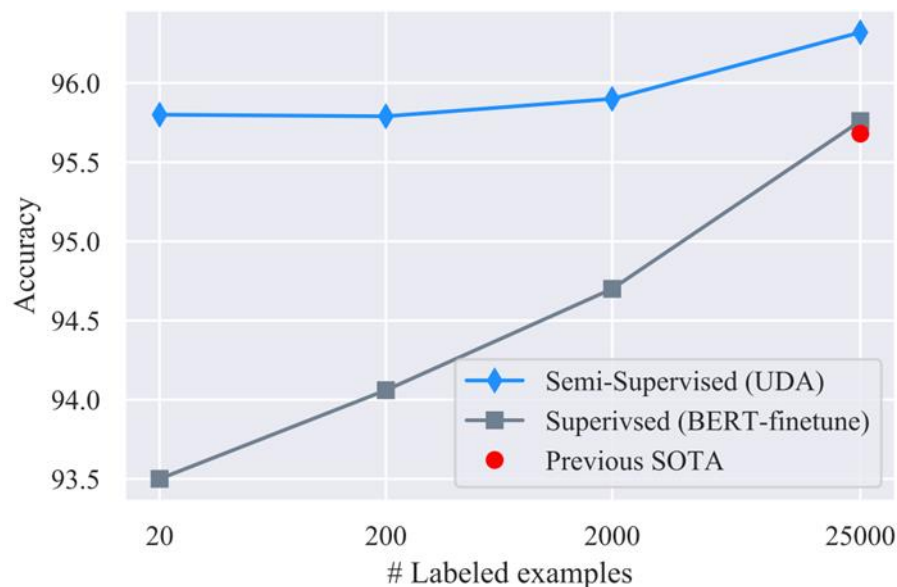
Fully supervised baseline							
Datasets (# Sup examples)		IMDb (25k)	Yelp-2 (560k)	Yelp-5 (650k)	Amazon-2 (3.6m)	Amazon-5 (3m)	DBpedia (560k)
Pre-BERT SOTA		4.32	2.16	29.98	3.32	34.81	0.70
BERT _{LARGE}		4.51	1.89	29.32	2.63	34.17	0.64
Semi-supervised setting							
Initialization	UDA	IMDb (20)	Yelp-2 (20)	Yelp-5 (2.5k)	Amazon-2 (20)	Amazon-5 (2.5k)	DBpedia (140)
Random	✗	43.27	40.25	50.80	45.39	55.70	41.14
	✓	25.23	8.33	41.35	16.16	44.19	7.24
BERT _{BASE}	✗	27.56	13.60	41.00	26.75	44.09	2.58
	✓	5.45	2.61	33.80	3.96	38.40	1.33
BERT _{LARGE}	✗	11.72	10.55	38.90	15.54	42.30	1.68
	✓	4.78	2.50	33.54	3.93	37.80	1.09
BERT _{FINETUNE}	✗	6.50	2.94	32.39	12.17	37.32	-
	✓	4.20	2.05	32.08	3.50	37.12	-

- 算法执行结果

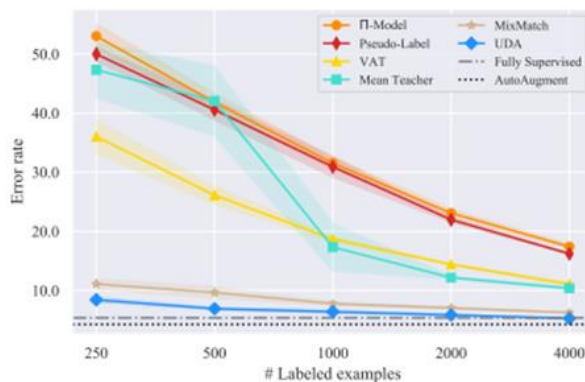
6种文本、3种视觉 显著提升

- 有标注数据集非常小

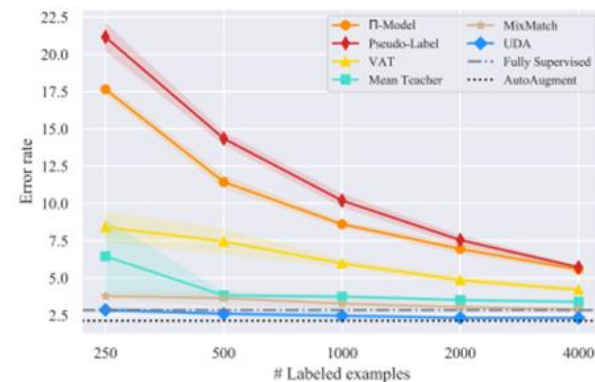
- 在具有10%标记样本的ImageNet上，将TOP1精度从55.1%提高到68.7%。
 - 20个标记，50k个未标记在IMDb情绪分析任务中错误率为4.20【25k标记 错误率4.32】



- 算法执行结果
 - 标准半监督测试
 - CIFAR-10中，UDA的表现优于所有现有的SSL方法，如VAT、ICT和MixMatch
 - CIFAR-10，4k标签，UDA实现了5.27的错误率，与使用50k标签的完全监督模型的性能相匹配
 - SVHN，250个标签，UDA实现了2.85的错误率，与使用70k标签的完全监督模型的性能相匹配



(a) CIFAR-10



(b) SVHN

- 优势
 - 无需海量标记数据，显著改善半监督学习
 - 使用AutoAugment，对未标记数据执行数据增强，一举达到半监督学习的当前最强
 - 可以匹配甚至优于使用更多标记数据的纯监督学习
 - 可用于文本和视觉的跨领域工作
- 劣势
 - 趋势方起，尚新，需时间考验

- 算法的应用领域
 - 文本、视觉、语音等领域
- 未来的发展
 - UDA的突破，使得仅有少量标记数据的半监督学习逼近了监督学习的精度，将会发展出无数的新应用场景

- [1] Xie Q, Dai Z, Hovy E, et al. Unsupervised data augmentation[J]. arXiv preprint arXiv:1904.12848, 2019.
- [2] Cubuk E D, Zoph B, Mane D, et al. AutoAugment: Learning Augmentation Strategies From Data[C]//Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition. 2019: 113–123.
- [3] Sajjadi M, Javanmardi M, Tasdizen T. Regularization with stochastic transformations and perturbations for deep semi-supervised learning[C]//Advances in Neural Information Processing Systems. 2016: 1163–1171.
- [4] <https://ai.googleblog.com/2019/07/advancing-semi-supervised-learning-with.html>
- [5] Schulman J, Wolski F, Dhariwal P, et al. Proximal policy optimization algorithms[J]. arXiv preprint arXiv:1707.06347, 2017.

谢谢！

大成若缺，其用不弊。大盈
若冲，其用不穷。大直若屈。
大巧若拙。大辩若讷。静胜
躁，寒胜热。清静为天下正。

