

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



模型预测可靠性的直接保障 —深度学习模型校准技术

硕士研究生 吴肖龙

2024年6月23日

- 总结反思

- 演讲时整体语速过快，缺乏停顿和关键性总结
- 内容过多和较为抽象，缺乏生动形象化的示例
- 未注意自己演讲时的仪态和肢体动作

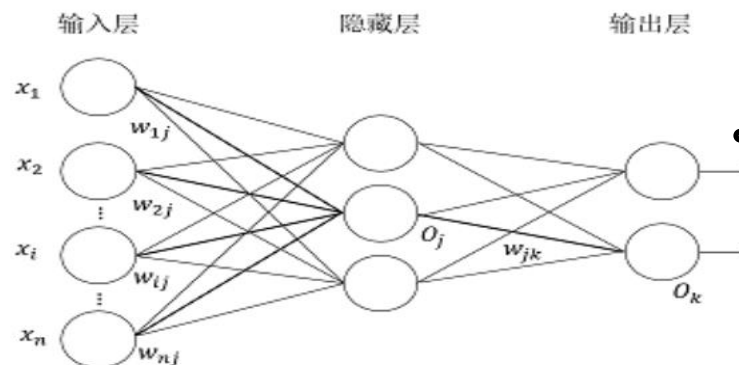
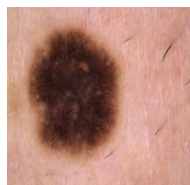
- 相关内容

- 2024.01.07 段学明 《DNN中的理论可解释性》
- 2023.11.19 夏志豪 《智能模型解释技术及其攻击方法》
- 2023.11.05 吴肖龙 《智能模型的不确定性估计》
- 2022.08.23 王若辉 《AI测试：历史与发展》

- 预期收获
- 题目内涵解析
- 研究背景与意义
- 研究历史与现状
- 知识基础
- 算法原理
 - CALS-ALM
 - GC-ETS
- 特点总结与工作展望
- 参考文献

- 预期收获
 - 掌握模型校准的基本概念和目标
 - 理解模型校准与不确定性量化的异同
 - 理解模型校准性能不佳的主要原因
 - 了解现有的模型校准框架和方法

- 使用模型预测数据



测试准确率90%

Benign

{Benign: 0.999; malignant: 0.001}

准确率 = 97%

- 提供置信度的预测

- 统计测试

- 提供更可靠的评估

- 不确定性估计

- 调整预测输出 (校准)

概率学派

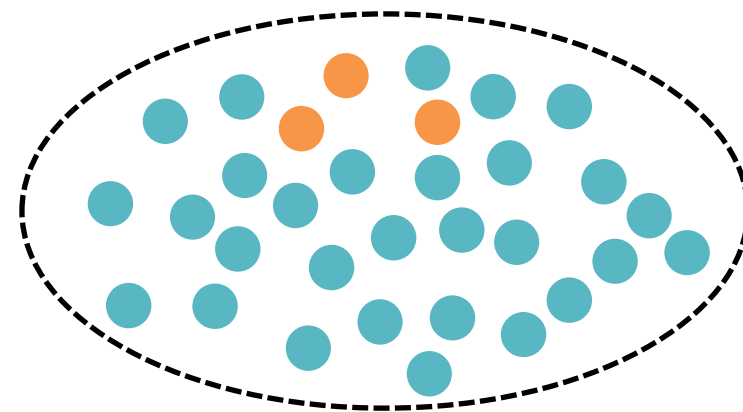
频率学派

内部度量

贝叶斯方法

深度集成

测试时间增强



● 模型预测错误的样本

● 模型预测正确的样本

- 模型校准

- 预测置信度与事件发生的真实概率的**匹配程度**

- 校准不佳的原因

- 过度参数化**和过拟合
 - 数据不平衡
 - 评估指标的偏差

- 研究目标

- 提供校准良好的模型

数据增强
数据缺陷



训练样本



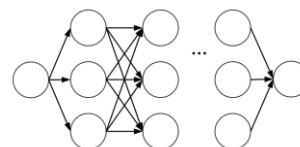
任务领域内样本

正则化方法

优化目标 (Loss)
正则化方法

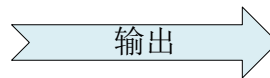


训练

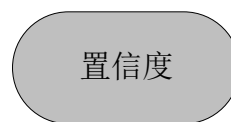


目标模型

后处理方法

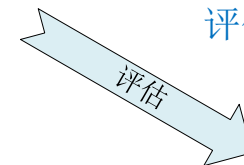


输出

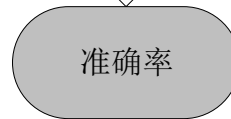


置信度

误差度量
评估偏差



评估



准确率

校准

模型的实际应用流程

高预测性模型的校准可能很差，表现为高校准误差；
深度模型往往主要是过度自信，表现为尖峰后验分布

- 研究背景

- 高安全性需求领域的关键决策需求
 - 自动驾驶、医疗、法律、国防和关基设施
- 对单个预测的性能关注而非模型性能
- 当前模型普遍表现出校准问题
 - 过度自信（主要）、信息不足



2021年，NHTSA报告807起自动驾驶相关案件

- 研究意义

- 防止模型过度自信或信息不足
- 帮助人们更好的放弃或接受预测
- 暴露和发掘智能模型的潜在缺陷
- 促进可信（可靠）人工智能的构建



2000~2013年间，机器人手术致死患者达144人

Guo等人提出温度缩放的后处理校准方法，通过添加**温度系数**调整预测输出，这是最流行且实施最容易的校准方法之一

2017

Kull等人面向多分类问题，提出了用狄利克雷**分布拟合**和校准数据集的预测后验

2019

Muller等人提出了通过标签平滑技术，降低 one-hot 编码在训练时**过拟合**的影响，缓解模型过度自信问题

2019

Zhang等人结合集成方法和温度缩放，集成多组温度缩放的调整结果，提升校准性能

2020

Niethammer等人改进了温度缩放方法，在图片的不同位置使用局部温度缩放，使得校准更加准确

2021

Bosselut等人提出**大模型**上的校准方法，结合随机文本生成，在普通标签、上下文标签、领域标签等进行校准

2023

2018

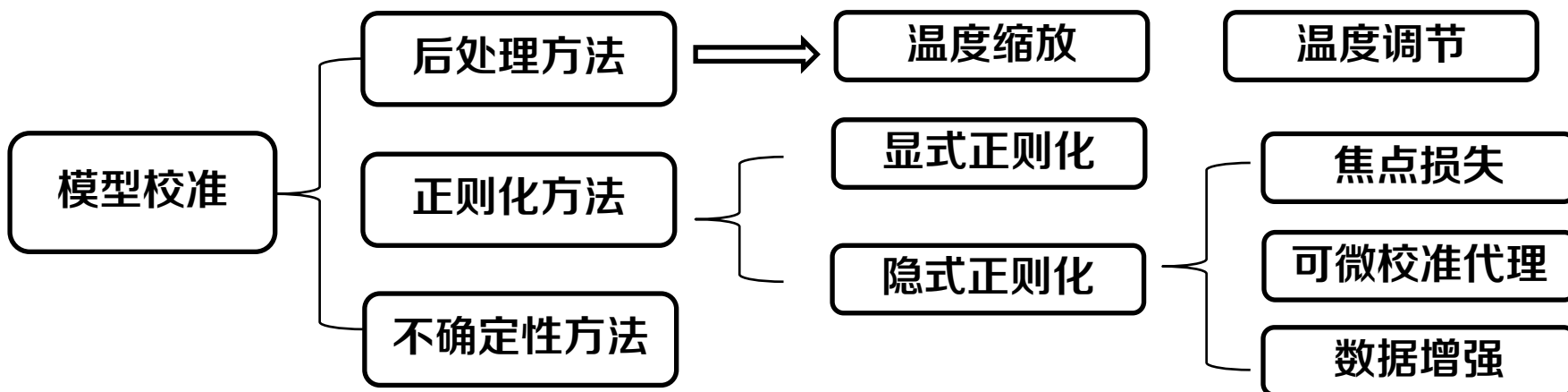
Kumar等人首次从校准**误差度量**角度改进校准性能，提出MMCE指标并将其引入训练过程

2019

Thulasidasan等人提出Mixup方法，通过在两个样本进行线性插值进行**数据增强**，降低模型校准误差

2022

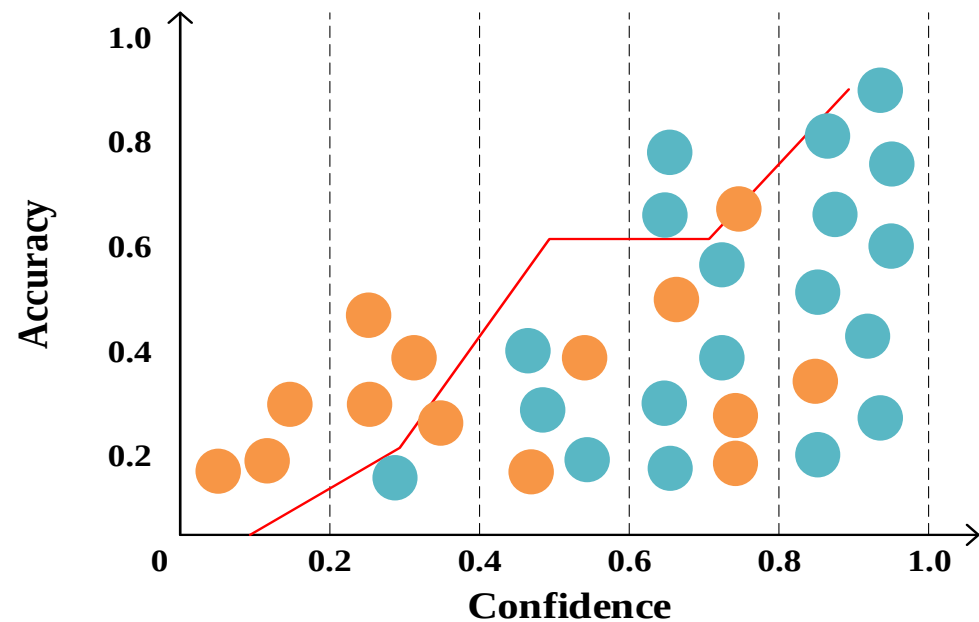
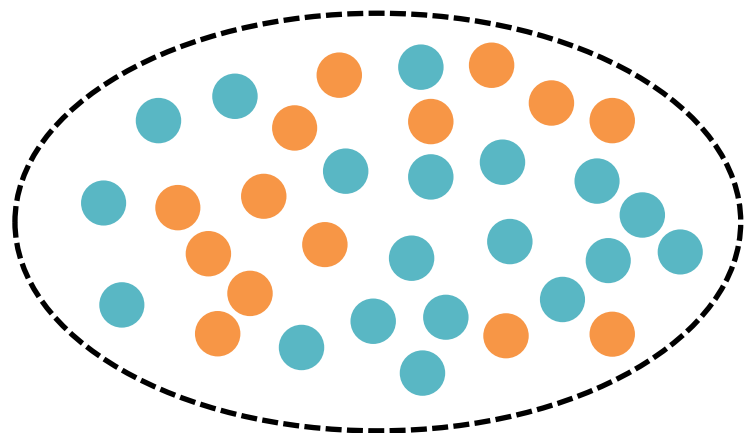
Hebbalaguppe等人提出了MDCA方法，考虑各个类别下的静态校准，并将其**引入正则化损失函数**进行训练优化



- 衡量置信度和测试准确率的相对差异

$$ECE = \sum_{m=1}^M \frac{|B_m|}{N} |A_m - C_m|$$

- M 为分组数， $|B_m|$ 为组中的样本数， N 表示样本总数



● 模型预测错误的样本

● 模型预测正确的样本



【 IEEE CVPR 】

Class Adaptive Network Calibration

T	目标	实现具有类别自适应的校准方法
I	输入	目标模型*1、训练数据集*1、验证数据集*1
P	处理	1. 将带边界约束的标签平衡引入损失项进行训练 2. 使用拉格朗日乘数法在验证集上更新各类别的惩罚系数 3. 循环前两步至模型收敛，使模型在各类别实现校准
O	输出	良好校准的模型*1

P	问题	1. 现有校准方法在类别上提供了相等的权重 2. 校准结果固定限制了模型进一步优化性能
C	条件	掌握完整模型设计与训练信息
D	难点	1. 惩罚函数和约束函数的设定 2. 校准系数的动态调整
L	水平	CVPR 2023 (CCF A)

- 标签平滑

- 在one-hot编码的目标分布中加入噪声

$$y_i = y_{one-hot}(1 - \alpha) + \frac{\alpha}{K}$$

- K 表示类别个数, α 是较小的超参数 (一般为0.1)

Predicted prob (p):	0.3	0.6	0.01	0.01	0.05	0.03
Label one-hot (y):	0	1	0	0	0	0
Label name:	飞机	鸟	猫	狗	汽车	拖拉机

Cross-entropy loss:

$$CE = -\sum_i^k y_i \log(p_i)$$

- 基于边界的约束

- 控制输出logits范围

$$\max_k \{l_k^{(i)}\} - l_j^{(i)} \leq m \cdot \mathbf{1}_{j \neq y(i)}$$

- K 表示第 K 个类别, $l^{(i)}$ 表示第 i 个样本的logits输出, j 是真实类别, m 是边界参数, $\mathbf{1}_{j \neq y(i)}$ 是一个指示函数, 当 $j \neq y(i)$ 时值为1, 否则为0

- 将约束条件转化为损失项

$$Loss = \min_{\theta} \sum_{i=1}^N L_{CE}(x^{(i)}, y^{(i)}) + \lambda \sum_{i=1}^N \sum_{j=1}^K \max \left\{ 0, \max_k \{l_k^{(i)}\} - l_j^{(i)} - m \right\}$$

- 样本级约束

- 问题：惩罚系数 λ 对所有目标都是相同的

- 为每个样本和类别引入了不同的惩罚权重 $\lambda_{i,j}$

- 变成需要选择 $N \times K$ 个惩罚权重的问题， N 为样本数量， K 为类别数量

以Cityscapes数据集为例，其包含19个类别的3000张像素为2048x1024的图像，以float32存储惩罚权重需**445GB**（开销过大）

- 惩罚粒度的放松

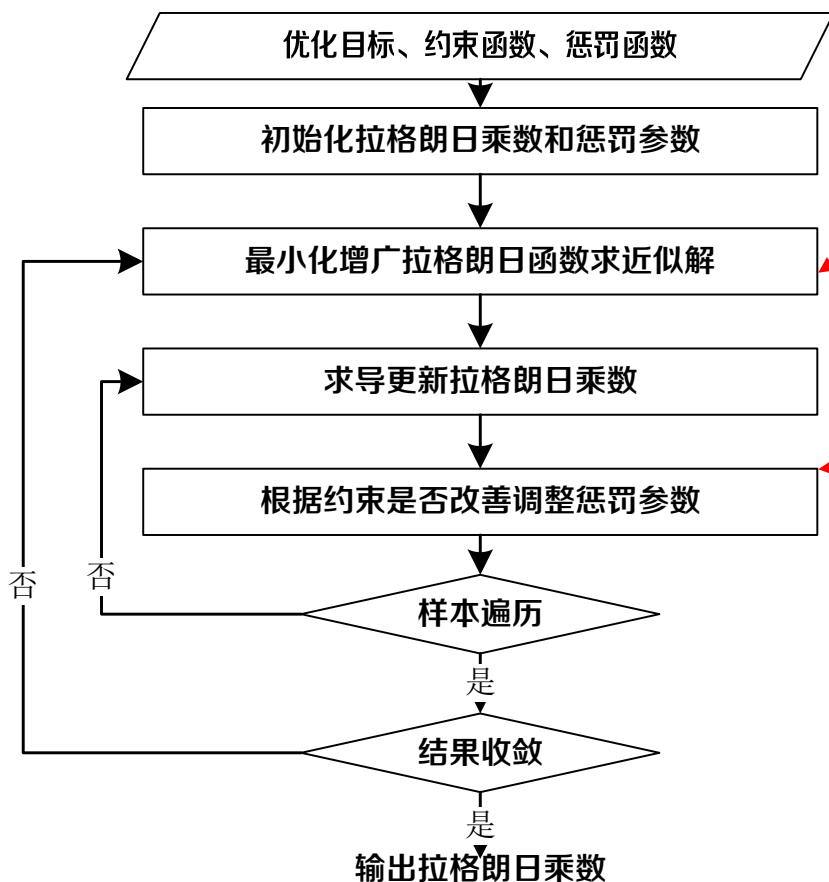
- 由**样本级约束**转化为**类别约束**（ K 个权重的优化问题）

- 增广拉格朗日乘数（ $\mathcal{L}^{(j)}(x)$ ）方法（附录A）

$$\min_x \mathcal{L}^{(j)}(x) = f(x) + \sum_{i=1}^n P(h_i(x), p_i^{(j)}, \lambda_i^{(j)})$$

- P 是惩罚-拉格朗日函数， $p^{(j)}$ 是惩罚参数， $\lambda_i^{(j)}$ 是与第 i 个约束相关联的拉格朗日乘数

- 类粒度的双层优化
 - 内循环更新惩罚参数
 - 外循环求解乘数算子



Algorithm 1 Augmented Lagrangian Multiplier algorithm

Require: Objective function f

Require: Constraint functions $h_i, i = 1, \dots, n$

Require: Penalty function P , initial $\lambda^{(0)} \in \mathbb{R}_{++}^n, \rho^{(0)} \in \mathbb{R}_{++}^n$

Require: Initial variable $x^{(0)}$, iterations $j = 1$

```

1: while not converged do
2:   Initialize with  $x^{(j-1)}$  and minimize (approximately):
      $\mathcal{L}^{(j)}(x) = f(x) + \sum_{i=1}^n P(h_i(x), \rho_i^{(j)}, \lambda_i^{(j)})$ 
3:    $x^{(j)} \leftarrow$  (approximate) minimizer of  $\mathcal{L}^{(j)}$ 
4:   for  $i = 1, \dots, n$  do
5:      $\lambda_i^{(j+1)} \leftarrow P'(h_i(x^{(j)}), \rho_i^{(j)}, \lambda_i^{(j)})$ 
6:     if the  $i$ -th constraint does not improve then
7:        $\rho_i^{(j+1)} \leftarrow \gamma \rho_i^{(j)}$ 
8:     else
9:        $\rho_i^{(j+1)} \leftarrow \rho_i^{(j)}$ 
10:    end if
11:  end for
12:   $j \leftarrow j + 1$ 
13: end while
    
```

- 优化对象

$$\min_{\theta} \sum_{i=1}^N L_{CE}(x^{(i)}, y^{(i)}) + \sum_{k=1}^K P(d_k^{(i)} - m, p_k, \lambda_k)$$

- 归一化 (提高稳定性)

$$\sum_{i=1}^N L_{CE}(x^{(i)}, y^{(i)}) + \frac{1}{K} \sum_{k=1}^K P\left(\frac{d_k^{(i)}}{m} - 1, p_k, \lambda_k\right)$$

- 更新拉格朗日乘数

$$\lambda_k^{(j+1)} = \frac{1}{|D_{val}|} \sum_{(x,y) \in D_{val}} P'\left(\frac{d_k}{m} - 1, p_k, \lambda_k\right)$$

- 使用的惩罚函数 (经验)

$$PHR(z, p, \lambda) = \begin{cases} \lambda z + \frac{1}{2} p z^2, & \lambda + p z \geq 0 \\ -\frac{\lambda^2}{2p}, & \text{else} \end{cases}$$

Algorithm 2 CALS-ALM training

Require: DNN initial θ^0 , margin m

Require: Dataset: $\mathcal{D}_{train}$, $\mathcal{D}_{val}$, batch size B

Require: Penalty function P , $\gamma > 1$, $\tau \in (0, 1)$

Require: Initial $\lambda^{(0)} \in \mathbb{R}_{++}^n$, $\rho^{(0)} \in \mathbb{R}_{++}^n$

```

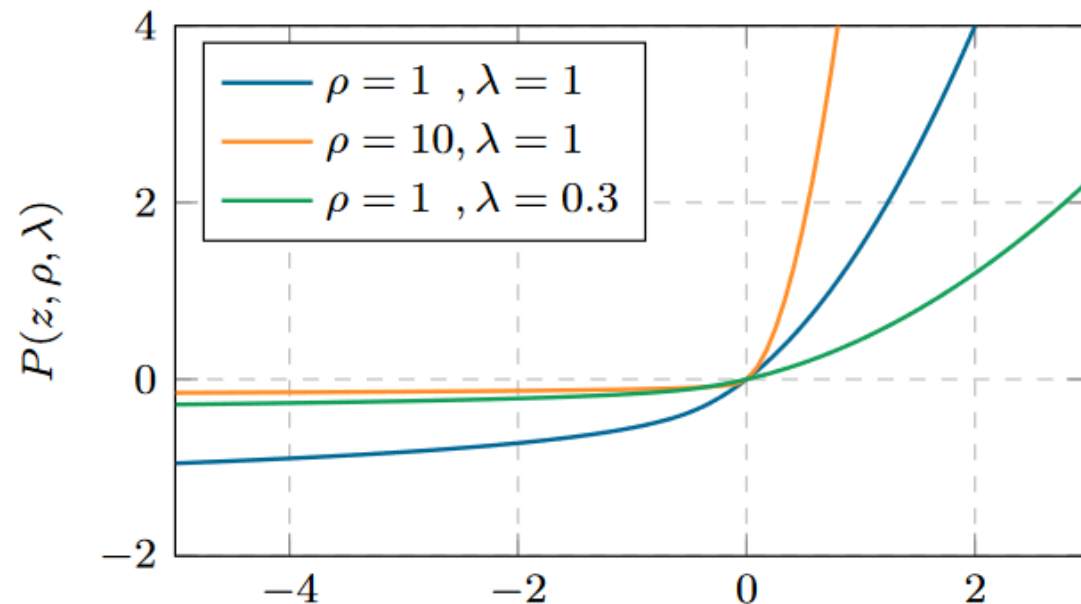
1: for  $j = 0, \dots, T$  do // Epochs of training
2:   for each mini-batch  $\{(x^{(i)}, y^{(i)})\}_{i=1}^B$  in  $\mathcal{D}_{train}$  do
3:      $\mathcal{L}_c = \sum_{i=1}^B L_{CE}(x^{(i)}, y^{(i)})$  // Cross-entropy
4:      $\mathcal{L}_p = \sum_{i=1}^B \frac{1}{K} \sum_{k=1}^K P\left(\frac{d_k^{(i)}}{m} - 1, \rho_k^{(j)}, \lambda_k^{(j)}\right)$  // Penalties
5:      $\mathcal{L} = \frac{1}{B} (\mathcal{L}_c + \mathcal{L}_p)$ 
6:      $\theta^{(t+1)} \leftarrow \theta^{(t)} - \alpha \cdot \nabla_{\theta} \mathcal{L}$  // Gradient descent
7:   end for 训练集做梯度下降
8:   for  $k = 1, \dots, K$  do // Adjust  $\lambda$  and  $\rho$  on validation
9:      $\lambda_k^{(j+1)} = \frac{1}{|\mathcal{D}_{val}|} \sum_{(x,y) \in \mathcal{D}_{val}} P'\left(\frac{d_k}{m} - 1, \rho_k^{(j)}, \lambda_k^{(j)}\right)$ 
10:     $\bar{d}_k^{(j)} = \frac{1}{|\mathcal{D}_{val}|} \sum_{(x,y) \in \mathcal{D}_{val}} \frac{d_k}{m} - 1$  // Average constraint
11:    if  $j \geq 1$  and  $\bar{d}_k^{(j)} > \tau \max\{0, \bar{d}_k^{(j-1)}\}$  then
12:       $\rho_k^{(j+1)} \leftarrow \gamma \rho_k^{(j)}$  // Constraint has not improved
13:    else
14:       $\rho_k^{(j+1)} \leftarrow \rho_k^{(j)}$ 
15:    end if
16:  end for 验证集更新惩罚参数
17: end for
    
```

• 核心思想

- 在更**细粒度**的条件下进行模型校准
- 将预测输出**约束**至有限的范围内
- 在训练过程中**实时调整**惩罚系数

• 创新点

- 将优化问题转化为损失函数的惩罚项，并推广至**类级别**进行优化
- 应用增广拉格朗日乘数方法**交替进行**目标函数**最小化**和拉格朗日乘数**更新**来求解



拉格朗日函数关于惩罚参数和
拉格朗日乘数的变化曲线

- 数据集

数据集	说明
Tiny-ImageNet	ImageNet 的小型子集，包含 200 个类别的 100,000 张图像
ImageNet	1,000 个类别的超 1M 张图像大型测试基准
ImageNet-LT	ImageNet 数据集的扩展，专注于类别的长尾分布，类别极不平衡

- 评价指标

- 预期校准误差 (ECE)
- 自适应校准误差 (AECE)：ECE 的变体，样本**均匀分配**到不同箱体
- 分类校准误差 (CECE)：ECE 的变体，按**类别**进行校准

- 对比方法

- CE (基线)、MMCE (2018 ICML)、ECP (2017 ICLR)、LS (2016 CVPR)
- FL/FLSD (2020 NIPS)、CPC (2022 CVPR)、MbLS (2022 CVPR)



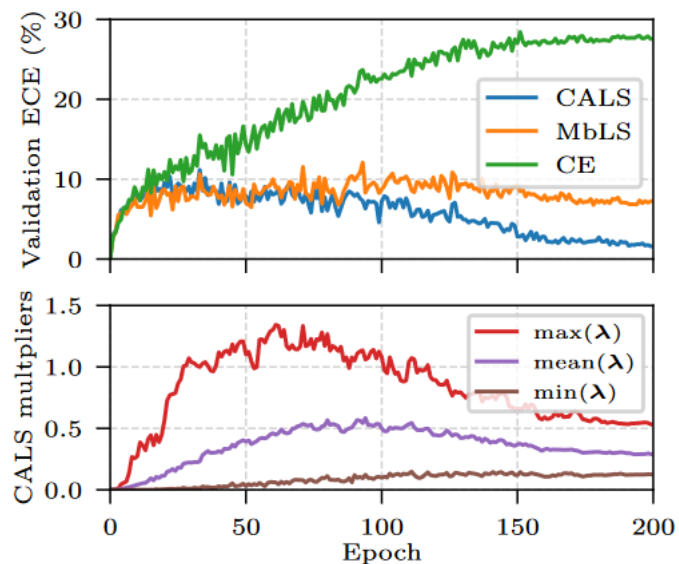
- 实验分析
 - 校准错误的严重程度与数据集和模型架构相关
 - 方法在全部场景下取得了较小的校准误差，对长尾分布数据集的提升效果显著

Method	TinyImageNet			ImageNet						ImageNet-LT					
	ResNet-50			ResNet-50			SwinV2-T			ResNet-50			SwinV2-T		
	Acc	ECE	AECE	Acc	ECE	AECE	Acc	ECE	AECE	Acc	ECE	AECE	Acc	ECE	AECE
CE	65.02	3.73	3.69	75.16	9.19	9.18	75.60	9.95	9.94	37.90	28.12	28.12	31.82	31.82	36.68
MMCE [18]	<u>65.34</u>	2.81	2.61	74.85	8.57	8.56	76.68	9.07	9.08	37.79	28.41	28.40	33.14	26.41	26.41
ECP [38]	64.90	4.00	3.92	75.22	8.27	8.26	75.82	9.88	9.86	37.69	28.14	28.13	31.22	33.70	33.70
LS [42]	65.78	3.17	3.16	76.04	2.57	2.88	75.42	7.32	7.33	37.88	10.46	10.38	31.70	11.42	11.40
FL [31]	63.09	2.96	3.12	73.87	<u>1.60</u>	<u>1.65</u>	75.60	3.19	3.18	36.04	18.37	18.36	30.73	25.50	25.50
FLSD [31]	64.09	2.91	2.95	73.97	2.08	2.06	74.70	2.44	2.37	36.18	17.77	17.78	32.56	25.16	25.17
CPC [5]	64.49	4.88	4.91	76.33	3.66	3.59	76.34	5.50	5.33	38.90	16.00	15.99	32.54	13.21	13.19
MbLS [22]	64.74	<u>1.64</u>	<u>1.73</u>	75.82	4.44	4.26	<u>77.18</u>	<u>1.95</u>	<u>1.73</u>	38.32	6.16	6.16	32.05	7.65	7.64
CALS-HR	65.09	2.50	2.42	<u>76.34</u>	5.63	5.69	77.58	3.06	2.95	38.50	<u>2.83</u>	<u>2.78</u>	34.31	<u>2.37</u>	2.45
CALS-ALM	65.03	1.54	1.38	76.44	1.46	1.32	77.10	1.61	1.69	<u>38.56</u>	2.15	2.30	<u>33.94</u>	2.32	2.45

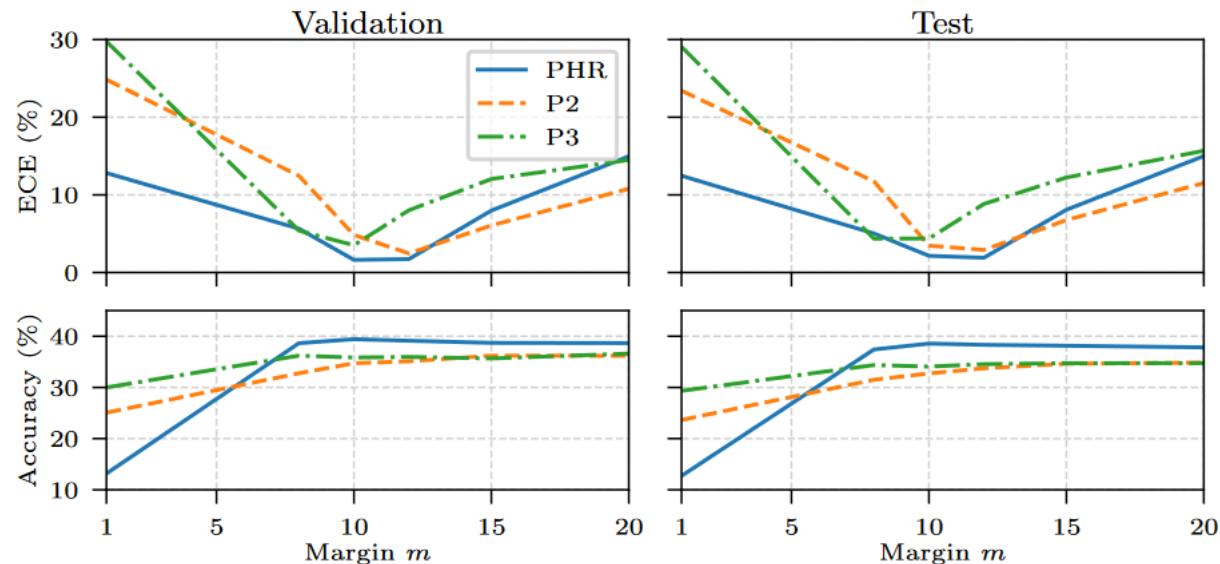
实现了在不同类别间的自适应

实验分析

- 拉格朗日乘数方法和PHR惩罚函数都实现了最佳的效果
- 随着训练周期增加，校准误差和惩罚系数都先快速增加再逐渐收敛，呈正相关



(a) Validation ECE and multipliers during training.



(b) Effect of penalty function and margin on accuracy and ECE for CALS.

实现了在模型训练过程中的动态调整



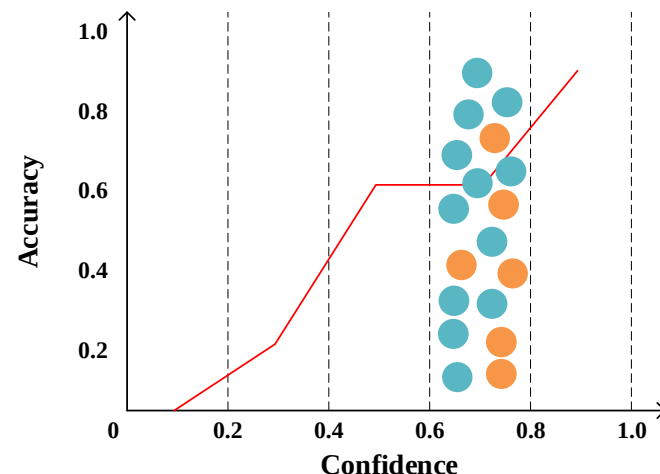
【 MIT Press 】

Beyond Probability Partitions: Calibrating Neural Networks with Semantic Aware Grouping

T	目标	实现更广义的校准误差评估和更准确的校准方法
I	输入	目标模型*1、训练数据集*1、验证数据集*1、测试数据集*1
P	处理	1. 使用L-BFGS优化算法（附录B）对样本在 语义信息 上进行分组 2. 使用温度缩放在 分组上 对模型进行校准 3. 将 分组函数 和 校准函数 同时进行端到端优化
O	输出	样本分组*N、校准后的模型*1

P	问题	1. 固定化 的分箱方式在评估时可能存在校准偏差 2. 从预测结果出发进行分箱 缺失 样本在 输入空间 的语义信息
C	条件	训练数据和验证数据间具有良好的分布适应性
D	难点	1. 如何选择合适的样本 分组方式 2. 如何 兼顾 样本分组和模型校准
L	水平	NeurIPS 2024（CCF A）

- 分组函数基本原则
 - 每个输入都能对应一个分组
 - 分组彼此是互斥的
 - 所有可能的分组都被覆盖
- 常见的分组方式
 - ECE（等间距划分置信度）、CECE（在各个类别上等间距划分置信度）、TCE（按预测最高概率标签分组）、AECE（等个数对样本进行分组）
- 抽象表示



$$PCE(S, g, \mathcal{L}, f, \mathcal{D}) = \sum_{P \in \mathcal{P}} p(P) \sum_{G \in P} p(G) \mathcal{L}(S(G), S(f(G)))$$

- $\mathcal{P}$ 表示所有可能的分组方式集合, $p(P)$ 是选择特定分组方式的概率, $p(G)$ 是观察到分组 G 中样本的概率, $\mathcal{L}(\cdot)$ 衡量绝对误差, $S(G)$ 是分组上的统计量, $f(\cdot)$ 是模型预测函数

- 假定分组函数为 $g(x)$,共分为 K 组, PCE表示为

$$PCE = \sum_{i=1}^K \frac{|G_i|}{|D|} \mathcal{L}(S(G_i), S(f_{G_i}(G_i)))$$

- 其中 $|D|$ 为样本数量, $|G_i|$ 为组中的样本数量, $\mathcal{L}(\cdot)$ 表示相对差异, $S(\cdot)$ 表示统计量
- 使用温度缩放作为校准函数

$$\mathcal{L}_{group} = \frac{1}{|G_i|} \sum_{x,y \in G_i} \log \frac{e^{\frac{o(x)y}{\tau_i}}}{\sum_j e^{\frac{o(x)j}{\tau_i}}}$$

- 将分组函数合并到公式中

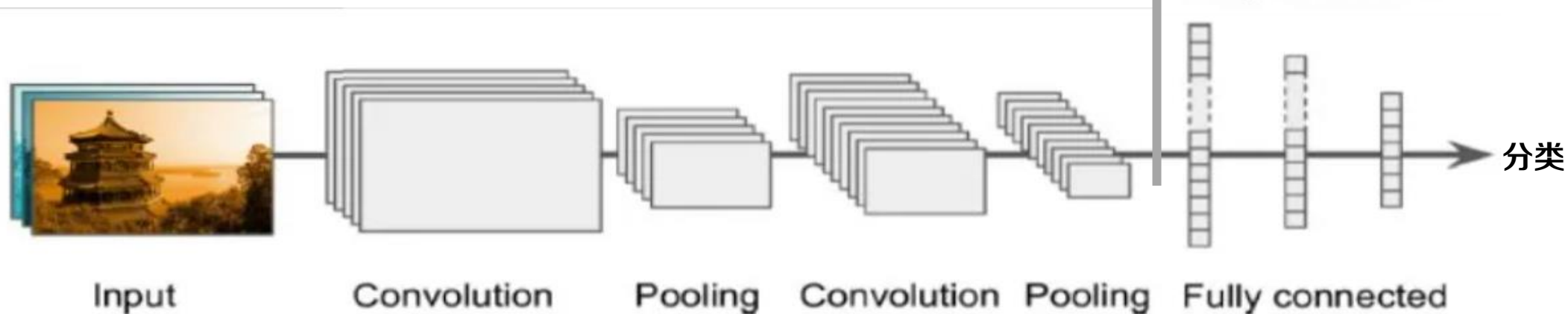
$$\mathcal{L}_{GC+TS} = \frac{1}{|D|} \sum_{x,y \in D} \log \sum_i g(x)_i \frac{e^{\frac{o(x)y}{\tau_i}}}{\sum_j e^{\frac{o(x)j}{\tau_i}}}$$

同时考虑分组和校准

- 分组以one-hot编码方式存在，引入平滑技术输出概率值
- 引入新的线性层，通过训练权重和偏置计算分组函数

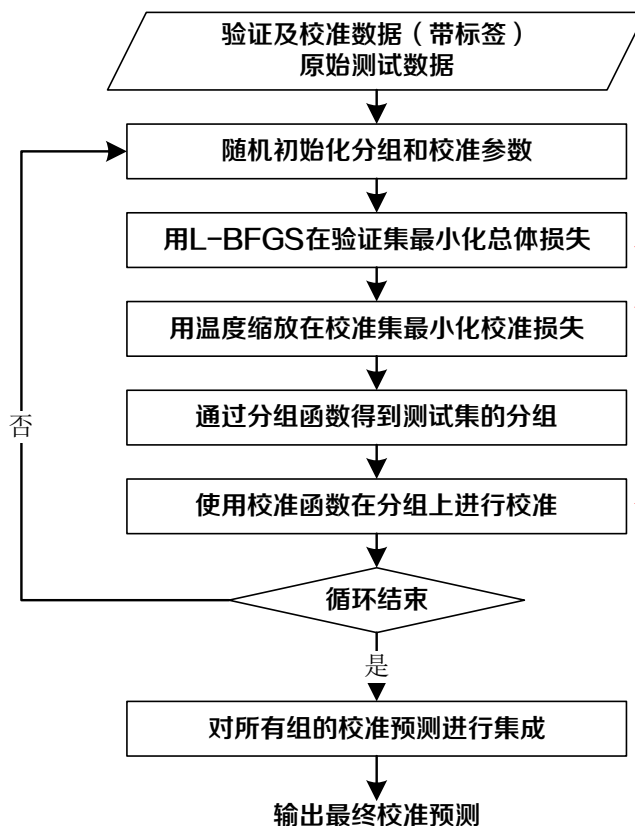
$$g_{\phi}(x) = g'_{\phi}(z(x)) = \text{softmax}((z(x)\phi_w + \phi_b))$$

- 其中 $g(x)$ 为分组函数， $z(x)$ 为模型提取出的样本特征， ϕ_w 为线性层权重， ϕ_b 为线性层偏置
- 一个线性层用于分组，一个线性层用于分类



• 类粒度的双层优化

- 内循环在各自分组条件下校准模型
- 外循环尝试找到更好的分组方式



Algorithm 1 Train group calibration with temperature scaling (GC+TS)

Input: $D_{val} = \{Z_{val}, O_{val}, Y_{val}\}$, $D_{ho} = \{Z_{ho}, O_{ho}, Y_{ho}\}$, $D_{test} = \{Z_{test}, O_{test}\}$, U, K, λ

Output: Calibrated $\hat{Y}_{test}$

- 1: **for** $u \leftarrow 0$ to $U - 1$ **do**
- 2: Randomly initialize ϕ_u and θ_u
- 3: Optimize ϕ_u and θ_u on D_{val} to minimize Eq. (13) with L-BFGS [30]
- 4: Optimize θ_{ui} on D_{ho} to minimize Eq. (11) with base calibration method (e.g., TS)
- 5: Calculate partition $P_u = \{G_{ui}\}$ with grouping function g_{ϕ_u} on D_{test}
- 6: **for** $i \leftarrow 0$ to $K - 1$ **do**
- 7: Calibrate G_{ui} with $f_{\theta_{ui}}(\cdot)$ to obtain $\hat{Y}_{ui}$
- 8: **end for**
- 9: Merge predicts in different groups $\hat{Y}_u = \{\hat{Y}_{ui}\}$
- 10: **end for**
- 11: Ensemble predicts in different partitions $\hat{Y}_{test} = \frac{1}{U} \sum_u \hat{Y}_u$

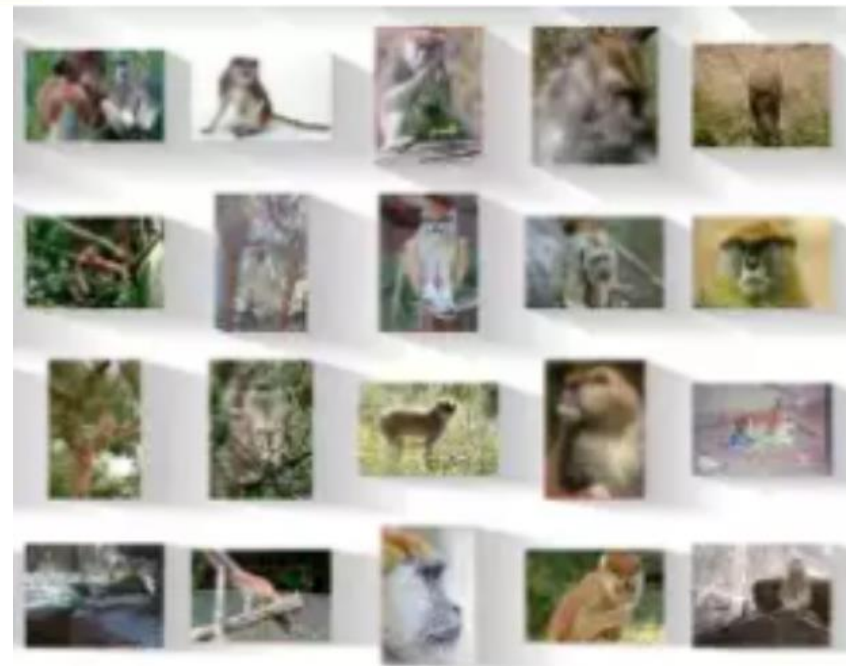
- 核心思想

- 将校准的评估方法统一至**概率分组框架**
- 原始**语义信息**应作为分组依据而非预测概率
- 同时优化分组函数和校准函数的端到端优化

- 创新点

- 基于语义特征**训练分组函数**得到更好的校准评估方法
- 使用**集成方法**组合多组校准结果提升校准性能

- 数据集
 - CIFAR10、CIFAR100、ImageNet
- 模型架构
 - ResNet (2016)、VGG (2015)、DenseNet (2017)、SWIN (2021)、ShuffleNet (2018)
- 评价指标
 - 预期校准误差 (ECE)
- 对比方法
 - 精度保持方法: TS (2017 ICML)、ETS (2020 ICML)、IRM(AP) (2020 ICML)
 - 非精度保持方法: Hist.B (2017 ICML)、Beta (2017 AISTATS)、BBQ (2015 AAAI)、DirODIR (2019 NeurIPS)、GPC (2020 AISTATS)、IRM(NAP)



- 实验分析
 - 方法整体上取得了最好结果
 - 集成起到了重要作用
 - 在数据集和模型架构上差异性较大

Dataset	Model	Uncal	TS	ETS	IRM(AP)	GC+TS(ours)	GC+ETS(ours)
CIFAR10	Resnet152	0.0249	0.0086	0.0101	0.0115	0.0079	0.0089
CIFAR10	Shufflenet	0.0464	0.0107	0.0103	0.0146	0.0099	0.0093
CIFAR10	VGG11	0.0473	0.0125	0.0135	0.0157	0.0120	0.0122
CIFAR100	Densenet121	0.0558	0.0418	0.0289	0.0415	0.0411	0.0280
CIFAR100	Resnet50	0.0580	0.0435	0.0292	0.0424	0.0427	0.0269
CIFAR100	VGG19	0.1668	0.0485	0.0472	0.0476	0.0414	0.0360
Imagenet	Resnet18	0.0279	0.0178	0.0104	0.0188	0.0173	0.0100
Imagenet	Resnet50	0.0382	0.0182	0.0102	0.0218	0.0174	0.0103
Imagenet	Swin	0.0266	0.0367	0.0218	0.0140	0.0366	0.0193
Average improvements						-5.03%	-8.92%

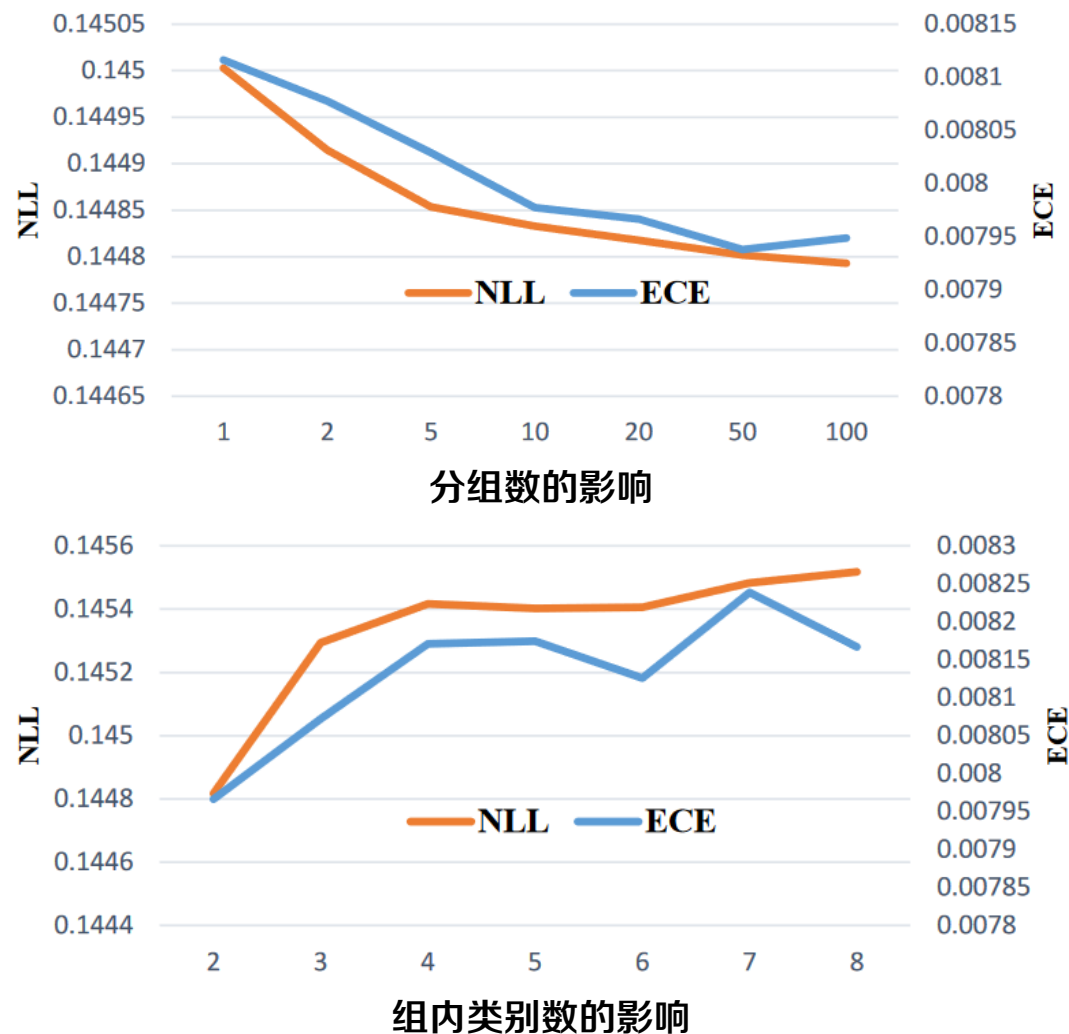
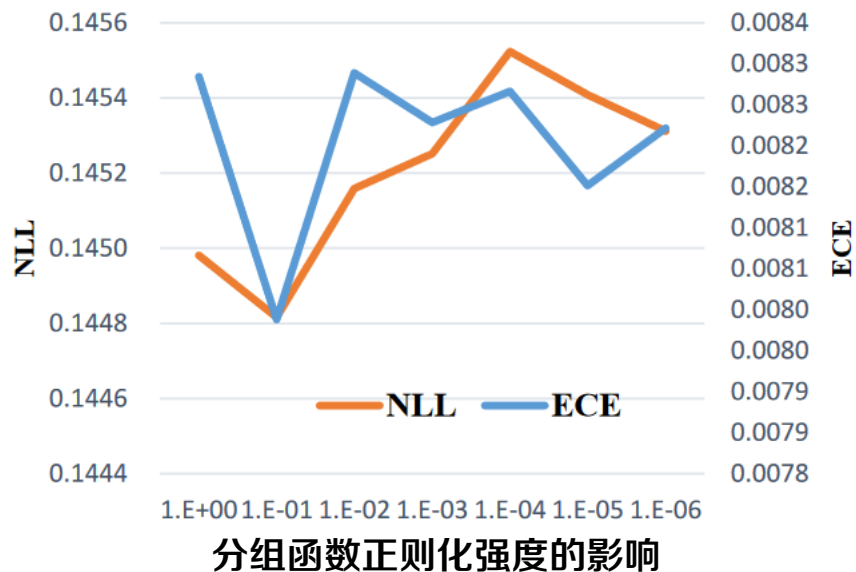
精度保持方法实验结果

Dataset	Model	Hist.B	Beta	BBQ	DirODIR	GPC	IRM(NAP)	Acc. Dec.	GC+ETS (ours)
CIFAR10	Resnet152	0.0172	0.0095	0.0097	0.0101	0.0189	0.0087	-0.079%	0.0089
CIFAR10	Shufflenet	0.0322	0.0137	0.0139	0.0158	0.0341	0.0119	-0.050%	0.0093
CIFAR10	VGG11	0.0279	0.0150	0.0137	0.0156	0.0376	0.0119	-0.129%	0.0122
CIFAR100	Densenet121	0.0632	0.0520	0.0307	0.0533	0.0306	0.0203	-0.149%	0.0280
CIFAR100	Resnet50	0.0618	0.0550	0.0334	0.0553	0.0346	0.0422	-3.686%	0.0269
CIFAR100	VGG19	0.0453	0.0642	0.0446	0.0932	0.1406	0.0470	-5.046%	0.0360
Imagenet	Resnet18	0.0714	0.0690	0.0483	0.0386	-*	0.0119	-0.027%	0.0100
Imagenet	Resnet50	0.0502	0.0707	0.0482	0.0326	-	0.0107	-0.006%	0.0103
Imagenet	Swin	0.0335	0.0629	0.0478	0.0148	-	0.0110	-0.060%	0.0193

非精度保持方法实验结果

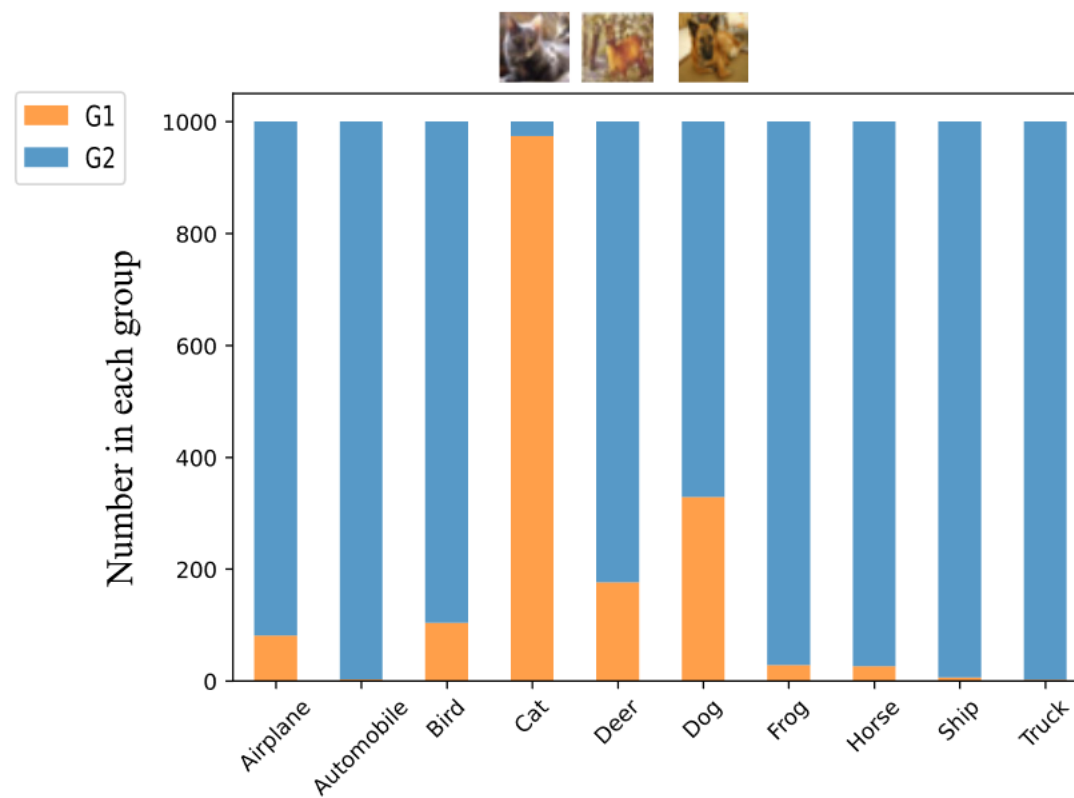
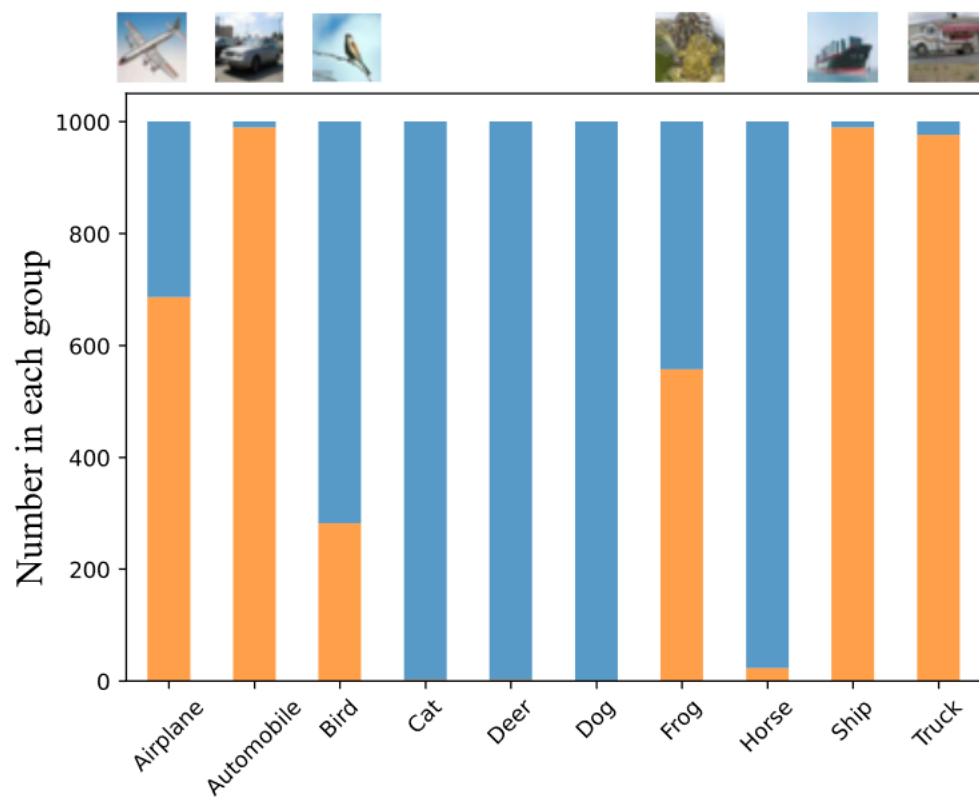
实验分析

- 分组数量有助于提升校准性能
- 组内类别复杂度影响校准结果



• 实验分析

- 在同一分组中展现相似的视觉特征（不局限于同一类别）





特点总结与未来展望

- CALS-ALM
 - 优点：有效应对了类别**长尾分布**下的校准，自适应调整校准系数**同步模型优化**
 - 缺点：对**同分布验证集**的强依赖，复杂的超参数和约束项设置
- GC-ETS
 - 优点：考虑了**校准误差评估**的灵活性，具有良好的校准性能
 - 缺点：计算开销和成本较高，缺少对**分组数量**和子集选择的控制
- 前沿发展与关键挑战
 - 无偏校准估计器（克服**分箱**带来的影响）
 - 校准**生成模型**和大模型（如何理解一致性）
 - 计算效率和可解释性（集成等大量重复开销）
 - **实际应用**和伦理问题（医疗、法律）



- [1] Wang C. Calibration in deep learning: A survey of the state-of-the-art[EB/OL]. (2024.05.10) [2024.06.20], <https://arxiv.org/abs/2308.01222>.
- [2] Zhu F, Zhang X Y, Cheng Z, et al. Revisiting Confidence Estimation: Towards Reliable Failure Prediction[J]. IEEE Transactions on Pattern Analysis and Machine Intelligence, 2023.
- [3] Liu B, Rony J, Galdran A, et al. Class adaptive network calibration. Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition[C]. Vancouver, BC, Canada: IEEE, 2023: 16070-16079.
- [4] Yang J Q, Zhan D C, Gan L. Beyond probability partitions: Calibrating neural networks with semantic aware grouping[J]. Advances in Neural Information Processing Systems, 2024, 36.

知人者智，自知者明。胜人者有力，自胜者强。知足者富。强行者有志。不失其所者久。死而不亡者，寿。

谢谢！



假设有一个方程为 $\frac{x^2}{a^2} + \frac{y^2}{b^2} + \frac{z^2}{c^2} - 1 = 0$ 的椭球，我们要求这个椭球内接长方体的最大体积。

那么也就是说我们要在 $\frac{x^2}{a^2} + \frac{y^2}{b^2} + \frac{z^2}{c^2} - 1 = 0$ 的条件下，求 $f(x, y, z) = 8xyz$ 的最大值。

现在我们定义一个拉格朗日函数 $F(x, \lambda)$:

$$F(x, \lambda) = f(x) + \sum_{k=1}^l \lambda_k h_k(x)$$

其中 λ_k 是各个约束条件的待定系数。

接下来求偏导为0的解:

$$\frac{\partial F}{\partial x} = 0, \frac{\partial F}{\partial \lambda} = 0$$

回到我们的题目，就可以将问题转化为:

$$F(x, y, z, \lambda) = 8xyz + \lambda(\frac{x^2}{a^2} + \frac{y^2}{b^2} + \frac{z^2}{c^2} - 1)$$

接下来对 $F(x, y, z, \lambda)$ 求偏导可得：

$$\frac{\partial F(x, y, z, \lambda)}{\partial x} = 8yz + \frac{2\lambda x}{a^2} = 0$$

$$\frac{\partial F(x, y, z, \lambda)}{\partial y} = 8xz + \frac{2\lambda y}{b^2} = 0$$

$$\frac{\partial F(x, y, z, \lambda)}{\partial z} = 8xy + \frac{2\lambda z}{c^2} = 0$$

$$\frac{\partial F(x, y, z, \lambda)}{\partial \lambda} = \frac{x^2}{a^2} + \frac{y^2}{b^2} + \frac{z^2}{c^2} - 1 = 0$$

联立前三个方程并代入第四个方程后可以求出：

$$x = \frac{\sqrt{3}}{3}a, y = \frac{\sqrt{3}}{3}b, z = \frac{\sqrt{3}}{3}c$$

最后可得

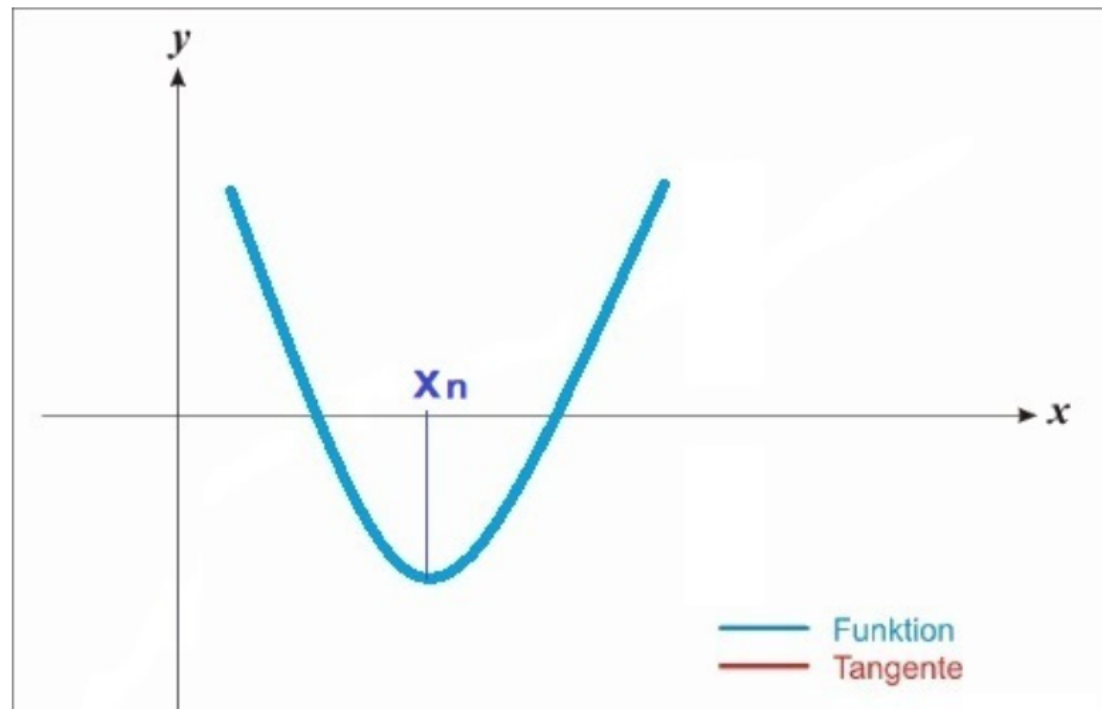
$$f(x, y, z) = 8xyz = \frac{8\sqrt{3}}{9}abc$$

一、牛顿法

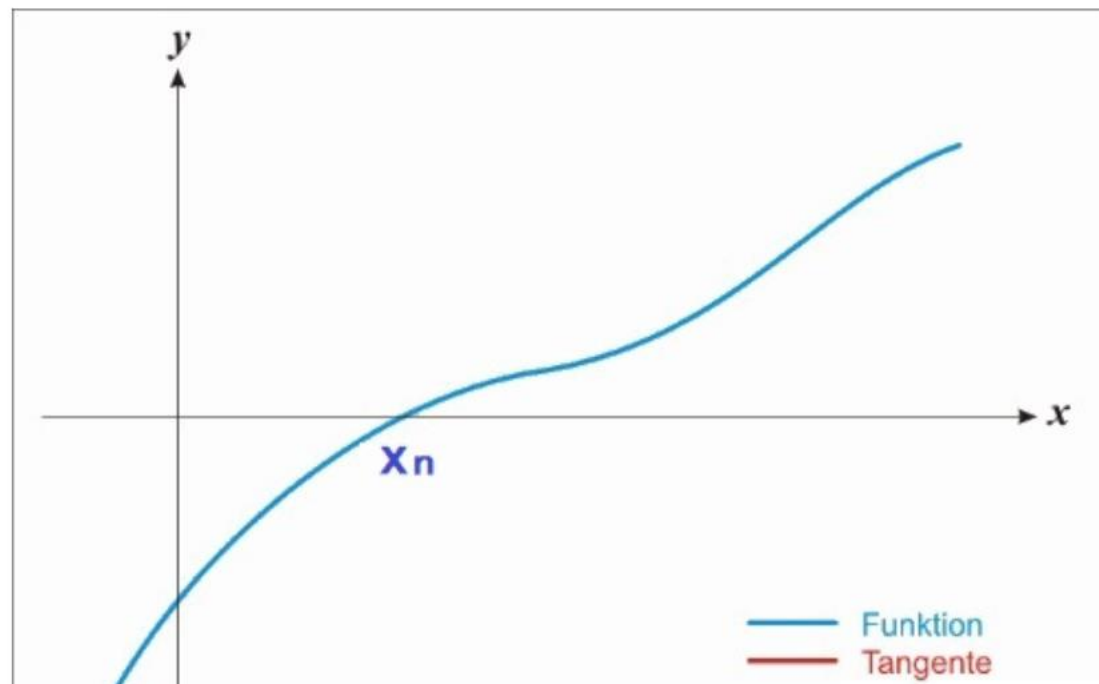
函数最优化算法方法不唯一，其中耳熟能详的包括梯度下降法，梯度下降法是一种基于迭代的一阶优化方法，优点是计算简单；牛顿法也是一种很重要的优化方法，是基于迭代的二阶优化方法，优点是迭代次数少，收敛速度很快。下面我们简要介绍一下牛顿法。

1.看图理解牛顿法

最优化问题就是寻找能使函数最小化的 x ，所以目标函数应当是一个凸函数（起码是局部凸函数），假如一个函数如下图：



他的一阶导数可能长下面这个样子：

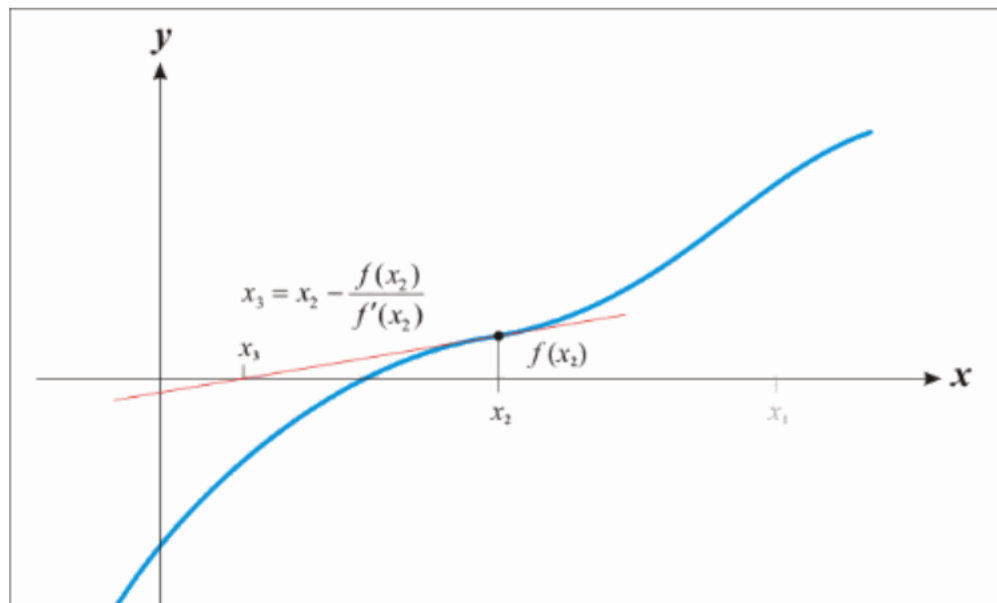


很显然函数在 x_n 处取得最小值，同时这个点的导数等于0，如果使用梯度下降，经过多次迭代， x 的取值会慢慢接近 x_n ，我们都能想象这个过程。

附录B

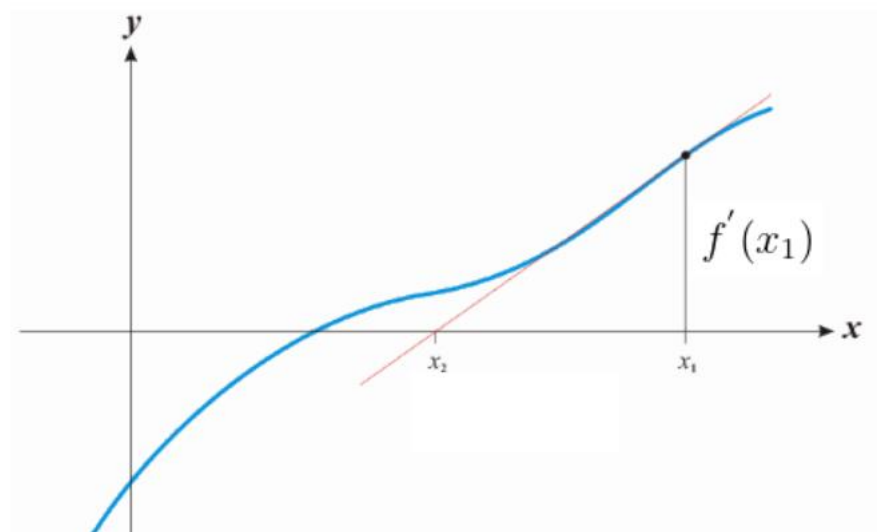
L-BFGS大规模优化算法

如果使用牛顿法， x 也会逼近 x_n ，不过速度会快很多，示例图如下：



这个过程可以这样描述：

- 在 x 轴上随机一点 x_1 ,经过 x_1 做 x 轴的垂线,得到垂线与函数图像的交点 $f(x_1)$.
- 通过 $f(x_1)$ 做函数的切线,得到切线与 x 轴的交点 x_2 .
- 迭代a/b两步，当前后两次求的 x 相同或者两个值的差小于一个阈值的时候，我们就认为找到了 x_n 。



如图4，蓝色的线是函数的 $f(x)$ 的导数 $f'(x)$ ，则曲线在 x_1 处的导数为 $f''(x_1)$ ，我们要求 x_2 ，根据三角函数有：

$$f''(x_1) = \frac{f'(x_1)}{x_1 - x_2} \quad (1)$$

得出：

$$x_2 = x_1 - \frac{f'(x_1)}{f''(x_1)} \quad (2)$$

利用 x_2 开始进行下一轮的迭代。迭代公式可以简化如下：

$$x_{n+1} = x_n - \frac{f'(x_n)}{f''(x_n)} \quad (3)$$

任意一点在 x_k 附近的二阶泰勒展开公式为:

$$f(x) = f(x_n) + f'(x_n)(x - x_n) + \frac{1}{2}f''(x_n)(x - x_n)^2 \quad (4)$$

对 $f(x)$ 求导:

$$f'(x) = f'(x_n) + f''(x_n)(x - x_n) \quad (5)$$

令 $f'(x) = 0$:

$$x = x_n - \frac{f'(x_n)}{f''(x_n)} \quad (6)$$

写成迭代形式:

$$x_{n+1} = x_n - \frac{f'(x_n)}{f''(x_n)} \quad (7)$$

可以看到使用三角函数和二阶泰勒展开最终得到的结果是一样的。虽然牛顿法收敛速度很快,但是当 x 的维度特别多的时候,我们想求得 $f''(x)$ 是非常困难的,而牛顿法又是一个迭代算法,所以这个困难我们还要面临无限多次,导致了直接使用牛顿法最为优化算法很难实际落地。为了解决这个问题出现了拟牛顿法,下面介绍一种拟牛顿法BFGS,主要就是想办法一种方法代替二阶导数。

二、BFGS公式推导

函数 $f(x)$ 在 $x = x_{k+1}$ 处的二阶泰勒展开式^Q为:

$$f(x) = f(x_{n+1}) + f'(x_{n+1})(x - x_{n+1}) + \frac{1}{2} f''(x_{n+1})(x - x_{n+1})^2 \quad (8)$$

当 x 为向量的时候, 上式写成:

$$f(x) = f(x_{n+1}) + \nabla f(x_{n+1})(x - x_{n+1}) + \frac{1}{2} \nabla^2 f(x_{n+1})(x - x_{n+1})^2$$

(9)

令 $G_{n+1} = \nabla^2 f(x_{n+1})$, 同时对 $f(x)$ 求导:

$$\nabla f(x) = \nabla f(x_{n+1}) + G_{n+1}(x - x_{n+1}) \quad (10)$$

接下来我们要想办法去掉 G_{n+1} , 我们使用 B_{n+1} 代替 G_{n+1} , B_{n+1} 是在迭代中一点点计算出来的而不使用二阶导数。

上式变为:

$$\nabla f(x) = \nabla f(x_{n+1}) + B_{n+1}(x - x_{n+1}) \quad (11)$$

$$B_{n+1}(x_{n+1} - x) = \nabla f(x_{n+1}) - \nabla f(x) \quad (12)$$

我们认为每次迭代 B_k 与上次变化 E_k , 形式如下:

$$B_{n+1} = B_n + E_n \quad (13)$$

令:

$$y_n = \nabla f(x_{n+1}) - \nabla f(x_n), \quad s_n = x_{n+1} - x_n \quad (14)$$

将式 (13) (14) 带入式子 (12) :

$$(B_n + E_n)s_n = y_n \quad (15)$$

令:

$$E_n = \alpha u_n u_n^T + \beta v_n v_n^T \quad (16)$$

其中 u_n, v_n 均为 n 维的向量, 带入 (15)

$$(B_n + \alpha u_n u_n^T + \beta v_n v_n^T)s_n = y_n \quad (17)$$

$$\alpha(u_n^T s_n)u_n + \beta(v_n^T s_n)v_n = y_n - B_n s_n \quad (18)$$

已知: $u_n^T s_n, v_n^T s_n$ 为实数, $y_n - B_n s_n$ 为向量。式 (18) 中, 参数 u_n 和 v_n 解的可能性有很多, 我们取特殊的情况, 假设 $u_n = r B_n s_n, v_n = \theta y_n$ 。带入 (16) 得:

$$E_n = \alpha r^2 B_n s_n s_n^T B_n + \beta \theta^2 y_n y_n^T \quad (19)$$

将 $u_n = r B_n s_n, v_n = \theta y_n$ 带入 (18) 得:

$$\alpha [(r B_n s_n)^T s_n] (r B_n s_n) + \beta [(\theta y_n)^T] (\theta y_n) = y_n - B_n s_n \quad (20)$$

$$[\alpha r^2 (s_n^T B_n s_n) + 1] + [\beta \theta^2 (y_n^T s_n) - 1] (y_n) = 0 \quad (21)$$

令 $\alpha r^2 (s_n^T B_n s_n) + 1 = 0$, 则:

$$\alpha r^2 = -\frac{1}{s_n^T B_n s_n} \quad (22)$$

令 $\beta \theta^2 (y_n^T s_n) - 1 = 0$, 则

$$\beta \theta^2 = \frac{1}{y_n^T s_n} \quad (23)$$

将式 (22) 和 (23) 带入 (19) :

$$E_n = -\frac{B_n s_n s_n^T B_n}{s_n^T B_n s_n} + \frac{y_n y_n^T}{y_n^T s_n} \quad (24)$$

将 (24) 带入 (13) 得到 B_n 的迭代公式:

$$B_{n+1} = B_n - \frac{B_n s_n s_n^T B_n}{s_n^T B_n s_n} + \frac{y_n y_n^T}{y_n^T s_n} \quad (24)$$

当 x 为向量的时候, 式 (7) 写成:

$$x_{n+1} = x_n - B_n^{-1} \nabla f(x_n) \quad (25)$$

加上学习率得到BFGS的迭代公式:

$$x_{n+1} = x_n - \eta(B_n^{-1} \nabla f(x_n)) \quad (26)$$

我们发现, 还要求 B_n 的逆, 这里可以引入sherman-morrism公式, 求解 B_n 的逆:

$$B_{n+1}^{-1} = (I - \frac{s_n y_n}{y_n^T s_n})^T B_n^{-1} (I - \frac{y_n s_n^T}{y_n^T s_n}) + \frac{s_n s_n^T}{y_n^T s_n} \quad (27)$$

我们用 H 代替 B^{-1} , 得到最终的BFGS迭代公式和 H 的迭代公式:

$$x_{n+1} = x_n - \eta(H_n \nabla f(x_n)) \quad (28)$$

$$H_{n+1} = (I - \frac{s_n y_n}{y_n^T s_n})^T H_n (I - \frac{y_n s_n^T}{y_n^T s_n}) + \frac{s_n s_n^T}{y_n^T s_n} \quad (29)$$

其中 s_n 是本轮 x 与上一轮 x 的差, y_n 是本轮梯度与上一轮梯度的差。

三、L-BFGS

在BFGS算法中，仍然有缺陷，每次迭代计算需要前次迭代得到的 H_n ， H_n 的存储空间至少为 $N(N+1)/2$ （ N 为特征维数），对于高维的应用场景，需要的存储空间将是非常巨大的。为了解决这个问题，就有了L-BFGS算法。L-BFGS即Limited-memory BFGS。L-BFGS的基本思想就是通过存储前 m 次迭代的少量数据来替代前一次的矩阵，从而大大减少数据的存储空间。

令 $\rho_n = \frac{1}{y_n^T s_n}$, $V_k = I - \frac{y_n s_n^T}{y_n^T s_k}$ ，则式 (29) 可以表示为：

$$H_{n+1} = V_n^T H_n V_n + \rho_n s_n s_n^T \quad (30)$$

若在初始时，假定初始的矩阵 $H_0 = I$ ，则我们可以得到：

$$H_1 = V_0^T H_0 V_0 + \rho_0 s_0 s_0^T \quad (31)$$

$$\begin{aligned} H_2 &= V_1^T H_1 V_1 + \rho_1 s_1 s_1^T \\ &= V_1^T (V_0^T H_0 V_0 + \rho_0 s_0 s_0^T) + \rho_1 s_1 s_1^T \\ &= V_1^T V_0^T H_0 V_1 + V_1^T \rho_0 s_0 s_0^T V_1 + \rho_1 s_1 s_1^T \end{aligned} \quad (32)$$

$$\begin{aligned} H_{n+1} &= (V_n^T V_{n-1}^T \cdots V_1^T V_0^T) H_0 (V_0 V_1 \cdots V_{n-1} V_n) \\ &\quad + (V_n^T V_{n-1}^T \cdots V_1^T) \rho_1 s_1 s_1^T (V_1 \cdots V_{n-1} V_n) \\ &\quad + \cdots \\ &\quad + V_n^T \rho_{n-1} s_{n-1} s_{n-1}^T V_n \\ &\quad + \rho_n s_n s_n^T \end{aligned}$$

假设当前迭代为 n ，只保存最近的 m 次迭代信息，按照上面的方式迭代 m 次，可以得到如下的公式：

$$\begin{aligned} H_{n+1} = & (V_n^T V_{n-1}^T \cdots V_{n-m}^T) H_0 (V_{n-m} \cdots V_{n-1} V_n) \\ & + (V_n^T V_{n-1}^T \cdots V_{n-m}^T) \rho_1 s_1 s_1^T (V_{n-m} \cdots V_{n-1} V_n) \\ & + \cdots \\ & + V_n^T \rho_{n-1} s_{n-1} s_{n-1}^T V_n \\ & + \rho_n s_n s_n^T \end{aligned}$$

由于 ρ, V 这些变量都最终可以由 s, y 两个向量计算得到，因此，我们只需存储最后 m 次的 s, y 向量即可算出 H_{n+1} ，加上对角阵 H_0 ，总共需要存储 $2*m+1$ 个 N 维向量（实际应用中 m 一般取4到7之间的值，因此需要存储的数据远小于Hesse矩阵）。

四、算法迭代过程

1. 选初始点 x_0 , 最小梯度阈值 $\varepsilon > 0$, 存储最近 m 次的迭代数据;
2. 初始化 $n = 0, H_0 = I, r = \nabla f(x_0)$;
3. 如果 $\|\nabla f(x_{n+1})\| \leq \varepsilon$, 则返回最优解 x , 否则转入步骤4;
4. 计算本次迭代的可行方向 $p + n = -r_k$;
5. 计算步长 α_k , 用下面的式子进行线搜索;

$$f(x_n + \alpha_n p_n) = \min f(x_n - \alpha p_n)$$

6. 用下面的更新公式更新 x ;

$$x_{n+1} = x_n + \alpha_n p_n$$

7.如果 n 大于 m , 保留最近 m 次的向量对, 删除 s_{n-m}, y_{n-m} ;

8.计算并保存向量对

$$s_n = x_{n+1} - x_n$$

$$y_n = \nabla f(x_{n+1}) - \nabla f(x_n)$$

9.用 two-loop recursion 算法求:

$$r_n = B_n \nabla f(x_n)$$

10.设置 $n = n + 1$, 转到步骤3