



TBSI 清华-伯克利深圳学院
Tsinghua-Berkeley Shenzhen Institute

浅析如何写好一篇论文 ——以计算机视觉为例

李一鸣

计算机科学与技术专业在读博士生

2022年11月12日



- 什么是一个好的想法？
- 论文的基础结构
- 论文各部分的写作注意事项
 - Abstract
 - Introduction
 - Related Work
 - Motivation/Observation
 - Method
 - Experiment
 - Conclusion
- 检查清单

什么是一个好的想法？



TBSI

清华-伯克利深圳学院
Tsinghua-Berkeley Shenzhen Institute



填补空白-延伸问题

迁移-再利用

复制-小修

1. 定义、解决了一个新的事物或问题
2. 推动、改变一个潮流或认识



Complexity



Difficulty



Technicality



Beauty



1. Abstract (***)
2. Introduction (***)
3. Related Work
4. Motivation/Observation (***) [selective]
5. Method (***)
6. Experiments (***)
7. Conclusion

1-2句话背景 (是什么, 为什么重要) +
motivation+因此我们做了什么+具体怎
么做+这么做取得了什么效果/意义

非常识性短语需要在后面给解释
(e.g./i.e., ...)

自己方法涉及到的关键性短语可以斜
体进行强化

时态记得一般现在时

关键词3-5个, 相关性从高到低排序

Speaker verification has been widely and successfully adopted in many mission-critical areas for user identification. The training of speaker verification requires a large amount of data, therefore users usually need to adopt third-party data (*e.g.*, data from the Internet or third-party data company). This raises the question of whether adopting untrusted third-party data can pose a security threat. In this paper, we demonstrate that it is possible to inject the hidden backdoor for infecting speaker verification models by poisoning the training data. Specifically, we design a clustering-based attack scheme where poisoned samples from different clusters will contain different triggers (*i.e.*, pre-defined utterances), based on our understanding of verification tasks. The infected models behave normally on benign samples, while attacker-specified unenrolled triggers will successfully pass the verification even if the attacker has no information about the enrolled speaker. We also demonstrate that existing backdoor attacks cannot be directly adopted in attacking speaker verification. Our approach not only provides a new perspective for designing novel attacks, but also serves as a strong baseline for improving the robustness of verification methods. The code for reproducing main results is available at

第一段：背景

第二段：详细展开研究的任务/意义

第三段：自己方法，包括motivation、method详细（以及为什么这么做）、优势（但是不要单纯写性能好）

第四段：贡献，3-4条

注：行文能让读者觉得你这么做是合适且理所应当的为佳（逻辑为王）；时态一般现在时；依旧把读者当“傻子”（该解释的解释，不要假设对方很懂你的研究领域）；最好有一张题图来帮助理解/分析。



- 详略得当 (e.g., 一篇做防御的文章一般需要同时review攻击和防御, 防御部分的related work篇幅应当大于攻击部分)
- 最好能有自己的总结, 而不是简单的按照时间顺序划分 (e.g., 给现有的工作进行分类总结)
- 注意和后面的实验部分保持一致和呼应
- 别人的工作用一般过去时
- 引用格式保持一致
- et al 等符号记得斜体



用来合理的引出自己的方法，可以是某个特别的实验现象+分析，也可以是现有方法的不足+分析。通过这些分析来引入你的method的合理性，让读者觉得你的做法是有insight而不是incremental的。最后一句一定是一个引子，用来引出下一个section的method。

- 最好能有一张好看的示意图总结方法的大致流程
- 先总后分（每个subsection开头先讲这个部分做什么，以及为什么，再讲具体的做法）
- 不要打断读者的思路，证明可以放附录。
- 如果方法有很多部分 (≥ 3)，在第一/二个subsection先总结 **general pipeline**，再每个subsection讲每个部分具体怎么做
- 可以考虑在第一个subsection讲prelim

- 一般是main experiments + discussion experiments (e.g., ablation study)
- 画图记得美观
- 表格和图分配合理（不要全是表或者全是图）
- 实验分析要有更多的细节，尤其是好在哪里
- 一定要有baseline
- 分析自己的方法相比**baseline**为什么好
- Baselines和数据集的选择要合理（至少包括<1年内的baseline和业内公认的数据集）



- 某种意义上来说，是abstract的改写版本，一般可以去除一些背景的细节和方法的motivation等。
- 记得用一般过去时。
- 可以包括一些“并不致命”的Limitations和对未来工作的展望。

- 写的时候要注意呈现出来的连贯性，不要打断读者思维，大段的证明等可以放附录。
- 把审稿人都当傻瓜，尤其是不是很common sense的东西一定要在出现的时候解释清楚
- 注意、主被动的分布得当，尽量多用主动语态
- 尽量不要写很长或很复杂的句子，拆成多个简单句
- 尽量避免口语化的用词和没有意义的冗余
- 英文中一般没有三个名词加在一起的，一般最多两个名词
- 写的时候不要拉仇恨，比如避免general, special case这种
- 尽量不要写没有被证明或者实验验证的claim，实在要写也要加presumably/probably/may

- 证明不要跳步，一步一步都写清楚
- 检查时态，别人的方法过去时，自己的方法现在时
- 图表要self-consistent，caption加上必要的说明，让读者光看图表就能知道在干什么
- 检查参考文献，包括所有的方法、数据集是不是都合理引用了及格式是否统一
- 检查图片，包括自己的方法在颜色上区分是否明显，label的大小是不是能看得清
- Grammarly查语法，Linggle查短语/介词搭配，DeepL翻译
- e.g., i.e., 后面要跟逗号
- 检查reference的格式是一致的（例如会议全部用缩写好了）



Thanks

