

Xuan Luo

6582 Calle Koral, Goleta, CA 93117 | xuan_luo@ucsb.edu | <https://luoxuan-cs.github.io/>

Research Focus

My research primarily focuses on Efficient Large Language Models:

- Computational efficient model inference via adaptive layer skipping (FlexiDepth, DiffSkip).
- Directly decoding multiple tokens without relying on speculative decoding (DMTD).

Additionally, I also worked on application areas of generative models, including: Detecting ChatGPT imposters with a single question (FLAIR); Image style transfer via progressive alignment (PAMA).

Education

University of California Santa Barbara, CA, USA

Sep 2023 – ongoing

Ph.D. Computer Science

Advisor: Prof. Xifeng Yan

Xi'an Jiaotong University, Shaanxi, China

Aug 2019 – June 2023

B.E. Computer Science (Honor Program)

Publications

- | | |
|---|-------------|
| <p>[1] Direct Multi-Token Decoding
 <u>Xuan Luo</u>, Weizhi Wang, Xifeng Yan
 Preprint
 [PDF] [CODE] [WEBSITE]</p> | <p>2025</p> |
| <p>[2] Adaptive Layer-skipping in Pre-trained LLMs
 <u>Xuan Luo</u>, Weizhi Wang, Xifeng Yan
 COLM 2025
 Oral Presentation (top 24/418)
 [PDF] [CODE] [DATASET] [WEBSITE]</p> | <p>2025</p> |
| <p>[3] DiffSkip: Differential Layer Skipping in Large Language Models
 <u>Xuan Luo</u>, Weizhi Wang, Xifeng Yan
 ACL Findings 2025
 [PDF] [CODE]</p> | <p>2025</p> |
| <p>[4] Bot or Human? Detecting ChatGPT Imposters with A Single Question
 Hong Wang, <u>Xuan Luo</u>, Weizhi Wang, Xifeng Yan
 COLM 2024
 [PDF] [CODE] [DATASET]</p> | <p>2024</p> |
| <p>[5] Progressive Attentional Manifold Alignment for Arbitrary Style Transfer
 <u>Xuan Luo</u>, Zhen Han, Linkang Yang
 ACCV 2022
 [PDF] [CODE]</p> | <p>2022</p> |
| <p>[6] Computer Science Diagram Understanding with Topology Parsing
 Shaowei Wang, Lingling Zhang, <u>Xuan Luo</u>, Yi Yang, Xin Hu, Tao Qin, Jun Liu
 TKDD 2022
 [PDF] [CODE]</p> | <p>2022</p> |