

Peihao Zhu

zhupeishishen@gmail.com | Google Scholar | <https://zpdesu.github.io>

WORK EXPERIENCE

Foundation Model, Adobe

Senior Applied Scientist

San Jose, CA

Jan 2026 – Present

Seed Team, ByteDance

Senior Research Scientist

San Jose, CA

Sep 2024 – Jan 2026

Seed Team, ByteDance

Research Scientist

San Jose, CA

Jul 2023 – Aug 2024

Reality Lab, Meta

Research Scientist Intern

Burlingame, CA

Oct 2022 – Feb 2023

Creative Vision Lab, Snap

Research Scientist Intern

Los Angeles, CA

May 2022 – Sep 2022

EDUCATION

King Abdullah University of Science and Technology (KAUST)

PhD in Computer Science

Kingdom of Saudi Arabia

Dec 2019 – Jun 2023

King Abdullah University of Science and Technology (KAUST)

Master of Science in Computer Science

Kingdom of Saudi Arabia

Aug 2017 – Dec 2019

Institute of Automation, Chinese Academy of Sciences

Master's Candidate in Computer Science

China

Sep 2016 – Jul 2017

Northeastern University

Bachelor of Technology in Automation; GPA: 90.0/100; Rank: Top 5%

China

Aug 2012 – Jun 2016

PROJECTS

Adobe Video Foundation Model – Gen6 / Gen6-long

[\[Website\]](#)

Core Contributor, Video Foundation Models

Jan 2026 – Present

- Serve as a main contributor to Adobe's next-generation video foundation model, leading the development and implementation of Gen6-long from research exploration to model delivery.
- Drive improvements in long-video visual and temporal consistency while optimizing model training and inference for greater computational efficiency.

MMCORE – Multimodal Image Editing Foundation Model

[\[Paper\]](#) [\[Website\]](#)

Core Contributor, Multimodal Image Editing

Jul 2025 – Dec 2025

- Drove core model development and large-scale experimentation for MMCORE, a unified multimodal foundation model for high-quality image editing.
- Contributed across model training, evaluation, and research iteration, helping advance multimodal understanding and instruction-guided editing.

Seedance 1.0 – Video Generation Foundation Model

[\[Paper\]](#) [\[Video\]](#) [\[Website\]](#)

Core Contributor, Video Foundation Models

Dec 2024 – Jun 2025

- Served as a main contributor to the development and large-scale training of Seedance 1.0, helping build the data and representation infrastructure underlying the video foundation model.
- Designed and implemented the distributed data-loading system and large-scale VAE-latent and text-embedding precomputation pipelines, enabling efficient preprocessing, storage, and consumption of training features.

Seaweed-7B – Video Generation Foundation Model

[\[Paper\]](#) [\[Video\]](#) [\[Website\]](#)

Founding/Core Contributor, Video Foundation Models

Mar 2024 – Nov 2024

- Built Seaweed-7B from its early stage as a main contributor, driving key components of the model and large-scale training stack through research iteration and public release.
- Designed the VAE for efficient visual-information compression and built the distributed data-loading system for reliable, high-throughput video-model training.

TikTok AI-moji

[\[Blog\]](#) [\[Platform\]](#)

Core Contributor, Generative Avatars

Jul 2023 – Dec 2024

- Drove generative-avatar model development for TikTok AI-moji, contributing from research prototyping through product deployment.
- Developed personalization and generation capabilities for creating stylized avatars from user images at product scale.

PUBLICATIONS

MMCORE: MultiModal Connection with Representation Aligned Latent Embeddings

ByteDance Seed

Technical Report, 2026

[\[Paper\]](#) [\[Webpage\]](#)

Seedance 1.0: Exploring the Boundaries of Video Generation Models

ByteDance Seed

Technical Report, 2025

[\[Paper\]](#) [\[Webpage\]](#)

Seaweed-7B: Cost-Effective Training of Video Generation Foundation Model

ByteDance Seed

Technical Report, 2025

[\[Paper\]](#) [\[Webpage\]](#)

3DAvatarGAN: Bridging Domains for Personalized Editable Avatars

Rameen Abdal, Hsin-Ying Lee, [Peihao Zhu](#), Menglei Chai, Aliaksandr Siarohin, Peter Wonka, Sergey Tulyakov
Proc. IEEE Conference on Computer Vision and Pattern Recognition (CVPR), 2023

[\[Paper\]](#) [\[Webpage\]](#)

HairNet: Hairstyle Transfer with Pose Changes

[Peihao Zhu](#), Rameen Abdal, John Femiani, Peter Wonka

Proc. European Conference on Computer Vision (ECCV), 2022

[\[Paper\]](#) [\[Webpage\]](#)

CLIP2StyleGAN: Unsupervised Extraction of StyleGAN Edit Directions

Rameen Abdal, [Peihao Zhu](#), John Femiani, Niloy Mitra, Peter Wonka

SIGGRAPH Conference Proceedings, 2022

[\[Paper\]](#) [\[Webpage\]](#)

Mind the Gap: Domain Gap Control for Single Shot Domain Adaptation for Generative Adversarial Networks

[Peihao Zhu](#), Rameen Abdal, John Femiani, Peter Wonka

International Conference on Learning Representations (ICLR), 2022

[\[Paper\]](#) [\[Webpage\]](#)

Barbershop: GAN-based Image Compositing using Segmentation Masks

[Peihao Zhu](#), Rameen Abdal, John Femiani, Peter Wonka

ACM Transactions on Graphics (Proc. SIGGRAPH Asia), 2021

[\[Paper\]](#) [\[Webpage\]](#)

Improved StyleGAN Embedding: Where are the Good Latents?

[Peihao Zhu](#), Rameen Abdal, Yipeng Qin, John Femiani, Peter Wonka

ArXiv pre-print, 2021

[\[Paper\]](#) [\[Webpage\]](#)

Styleflow: Attribute-conditioned exploration of stylegan-generated images using conditional continuous normalizing flows

Rameen Abdal, [Peihao Zhu](#), Niloy Mitra, Peter Wonka

ACM Transactions on Graphics (TOG), 2021

[\[Paper\]](#) [\[Webpage\]](#)

Labels4Free: Unsupervised Segmentation using StyleGAN

Rameen Abdal, [Peihao Zhu](#), Niloy J. Mitra, Peter Wonka

Proc. IEEE International Conference on Computer Vision (ICCV), 2021

[\[Paper\]](#) [\[Webpage\]](#)

Flow-Guided Video Inpainting with Scene Templates

Dong Lao, [Peihao Zhu](#), Peter Wonka, Ganesh Sundaramoorthi

Proc. IEEE International Conference on Computer Vision (ICCV), 2021

[\[Paper\]](#) [\[Webpage\]](#)

SEAN: Image Synthesis with Semantic Region-Adaptive Normalization

[Peihao Zhu](#), Rameen Abdal, Yipeng Qin, Peter Wonka

Proc. IEEE Conference on Computer Vision and Pattern Recognition (CVPR Oral), 2020

[\[Paper\]](#) [\[Webpage\]](#)

Large Scale Architecture Asset Extraction from Panoramic Imagery

[Peihao Zhu](#), Wamiq Reyaz Para, Anna Fruehstueck, John Femiani, Peter Wonka

IEEE Transactions on Visualization and Computer Graphics (TVCG), 2020

[\[Paper\]](#) [\[Webpage\]](#)

AWARDS & HONORS

- KAUST CEMSE Dean's List Award 2021-2022
- National Scholarship 2013
- KAUST Fellowship 2017
- UCAS Scholarship 2016
- Scholarship for Outstanding Merits 2012-2016

REFERENCE

Prof. Peter Wonka

Full Professor, Visual Computing Center, KAUST

peter.wonka@kaust.edu.sa

<http://peterwonka.net>