

1. Record Nr.	UNISA996630861503316
Autore	Lin Zhouchen
Titolo	Pattern Recognition and Computer Vision : 7th Chinese Conference, PRCV 2024, Urumqi, China, October 18–20, 2024, Proceedings, Part VI // edited by Zhouchen Lin, Ming-Ming Cheng, Ran He, Kurban Ubul, Wushouer Silamu, Hongbin Zha, Jie Zhou, Cheng-Lin Liu
Pubbl/distr/stampa	Singapore : , : Springer Nature Singapore : , : Imprint : Springer, , 2025
ISBN	981-9785-08-1
Edizione	[1st ed. 2025.]
Descrizione fisica	1 online resource (621 pages)
Collana	Lecture Notes in Computer Science, , 1611-3349 ; ; 15036
Altri autori (Persone)	ChengMing-Ming HeRan UbulKurban SilamuWushouer ZhaHongbin ZhouJie LiuCheng-Lin
Disciplina	006.37
Soggetti	Image processing - Digital techniques Computer vision Artificial intelligence Application software Computer networks Computer systems Machine learning Computer Imaging, Vision, Pattern Recognition and Graphics Artificial Intelligence Computer and Information Systems Applications Computer Communication Networks Computer System Implementation Machine Learning
Lingua di pubblicazione	Inglese
Formato	Materiale a stampa
Livello bibliografico	Monografia
Nota di contenuto	Visual Harmony: LLM's Power in Crafting Coherent Indoor Scenes from

ImagesSuperpixel Cost Volume Excitation for Stereo MatchingMulti-view Depth Estimation with Adaptive Feature Extraction and Region-Aware Depth Prediction3D Data Augmentation for Driving Scenes on CameraA Pose-Aware Auto-Augmentation Framework for 3D Human Pose and Shape Estimation from Partial Point CloudsEfficient Emotional Talking Head Generation via Dynamic 3D Gaussian RenderingGeneralizable Geometry-aware Human Radiance Modeling from Multi-view ImagesAG-NeRF: Attention-guided Neural Radiance Fields for Multi-height Large-scale Outdoor Scene RenderingJPA: A Joint-Part Attention for Mitigating Overfocusing on 3D Human Pose EstimationRealistic and Visually-pleasing 3D Generation of Indoor Scenes from a Single ImageAttenPoint: Exploring Point Cloud Segmentation through Attention-Based ModulesMTFusion: Reconstructing Any 3D Object from Single Image Using Multi-Word Textual InversionMulti-view 3D Reconstruction by Fusing Polarization InformationQuat-DGNet: Enhancing 3D Dense Captioning with Quaternion-Based Spatial Offsets and Dynamic Neighborhood GraphsDisparity Refinement Based on Cross-Modal Feature Fusion and Global Hourglass Aggregation for Robust Stereo Matching -- Trajectory-based Calibration for Optical See-Through Head-Mounted Displays without Alignment -- Animatable Human Rendering from Monocular Video via Pose-Independent Deformation -- Maximum Spanning Tree for 3D Point Cloud RegistrationLearning the Dynamic Spatio-Temporal Relationship Between Joints for 3D Human Pose Estimation -- MaskEditor: Instruct 3D Object Editing with Learned MasksDyGASR: Dynamic Generalized Gaussian Splatting with Surface Alignment for Accelerated 3D Mesh Reconstruction -- MMIDM: Generating 3D Gesture from Multimodal Inputs with Diffusion Models -- Discriminative-guided Diffusion-based Self-supervised Monocular Depth Estimation -- Multiview Light Field Angular Super-Resolution based on View Alignment and Frequency Attention -- MagicGS: Combining 2D and 3D Priors for Effective 3D Content Generation -- ESD-Pose: Enhanced Semantic Discrimination for Generalizable 6D Pose EstimationTrans-DONeRF for Transparent Object Rendering with Mixed Depth Prior -- SFDNeRF: A Semantic Feature-Driven Few-Shot Neural Radiance Field Framework with Hybrid Regularization -- TriEn-Net: Non-parametric Representation Learning for Large-Scale Point Cloud Semantic Segmentation -- Decomposed Latent Diffusion Model for 3D Point Cloud Generation -- Learning Multi-Branch Attention Networks for 3D Face Reconstruction -- CP-VoteNet: Contrastive Prototypical VoteNet for Few-Shot Point Cloud Object DetectionCross Modality Fusion Network with Feature Alignment and Salient Object Exchange for Single Image 3D Shape Retrieval -- Enhanced Spatial Adaptive Fusion Network For Video Super-ResolutionMulti-3D Occlusion Mask Learning for Flexible Occlusion Removal in Neural Radiance Fields -- Sketch-Based 3D Shape Retrieval via Cross-Modal Contrastive Learning and Difficulty-Aware Uncertainty Regularization -- Residual Hybrid Attention Enhanced Video Super-Resolution with Cross Convolution -- SDFReg: Learning Signed Distance Functions for Point Cloud Registration -- Unfolding Gradient Graph Regularization for Point Cloud Color Denoising -- ER-SFM: EFFICIENT AND ROBUST CLUSTER-BASED STRUCTURE FROM MOTION -- Multimodal Token Fusion and Optimization for 3D Human Mesh Reconstruction with Transformers.

Sommario/riassunto

This 15-volume set LNCS 15031-15045 constitutes the refereed proceedings of the 7th Chinese Conference on Pattern Recognition and Computer Vision, PRCV 2024, held in Urumqi, China, during October 18–20, 2024. The 579 full papers presented were carefully reviewed and selected from 1526 submissions. The papers cover various topics

in the broad areas of pattern recognition and computer vision, including machine learning, pattern classification and cluster analysis, neural network and deep learning, low-level vision and image processing, object detection and recognition, 3D vision and reconstruction, action recognition, video analysis and understanding, document analysis and recognition, biometrics, medical image analysis, and various applications.
