

1. Record Nr.	UNINA9910983328503321
Titolo	Pattern Recognition and Computer Vision : 7th Chinese Conference, PRCV 2024, Urumqi, China, October 18–20, 2024, Proceedings, Part VII // 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	981-9785-11-1
Edizione	[1st ed. 2025.]
Descrizione fisica	1 online resource (XIV, 587 p. 203 illus., 182 illus. in color.)
Collana	Lecture Notes in Computer Science, , 1611-3349 ; ; 15037
Disciplina	006
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	Scene Text Recognition via k-NN Attention-based Decoder and Margin-based Softmax LossReal-Time Text Detection with Multi-Level Feature Fusion and Pixel ClusteringREFINED AND LOCALITY-ENHANCED FEATURE FOR HANDWRITTEN MATHEMATICAL EXPRESSION RECOGNITIONLearning Fine-grained and Semantically Aware Mamba Representations for Tampered Text Detection in ImagesDual Feature Enhanced Scene Text Recognition Method for Low-Resource UyghurSegmentation-free Todo Mongolian OCR and Its Public

DatasetHybrid Encoding Method for Scene Text Recognition in Low-Resource UyghurROBC: a Radical-Level Oracle Bone Character DatasetIntegrated Recognition of Arbitrary-Oriented Multi-Line Billet NumberImproving Scene Text Recognition with Counting Aware Contrastive Learning and Attention AlignmentGridMask: An Efficient Scheme for Real Time Curved Scene Text DetectionTibetan Handwriting Recognition Method based on Structural Re-parameterization ViT and Vertical AttentionMFH: Marrying Frequency Domain with Handwritten Mathematical Expression RecognitionLeveraging Structure Knowledge and Deep Models for the Detection of Abnormal Handwritten Text -- OCR-aware Scene Graph Generation via Multi-modal Object Representation Enhancement and Logical Bias Learning -- Enhancing Transformer-based Table Structure Recognition for Long Tables -- Show Exemplars and Tell Me What You See: In-context Learning with Frozen Large Language Models for Text -- VQAMLR-NET: an arbitrary skew angle detection algorithm for complex layout document images -- TextViTCNN Enhancing Natural Scene Text Recognition with Hybrid Transformer and Convolutional NetworksEnhancing Visual Information Extraction with Large Language Models through Layout-aware Instruction Tuning -- SFENet: Arbitrary Shapes Scene Text Detection with Semantic Feature ExtractorImproving Zero-Shot Image Captioning Efficiency with Metropolis-Hastings Sampling -- Improving Text Classification Performance through Multimodal Representation -- A Multi-feature Fusion Approach for Words Recognition of Ancient Mongolian Documents -- TableRocket: An Efficient and Effective Framework for Table Reconstruction -- Not All Texts Are the Same: Dynamically Querying Texts for Scene Text Detection -- Multi-Modal Attention based on 2D Structured Sequence for Table Recognition -- A Two-stream Hybrid CNN-Transformer Network for Skeleton-based Human Interaction Recognition -- Skeleton-Language Pre-training to Collaborate with Self-Supervised Human Action Recognition -- Spatio-Temporal Contrastive Learning for Compositional Action RecognitionPath-Guided Motion Prediction with Multi-View Scene Perception -- Privacy-preserving Action Recognition: A Survey -- Attention-based Spatio-temporal modeling with 3D Convolutional Neural Networks for Dynamic Gesture Recognition -- MIT: Multi-cue Injected Transformer for Two-stage HOI Detection -- DIDA: Dynamic Individual-to-integrated Augmentation for Self-Supervised Skeleton-Based Action Recognition -- Multi-scale Spatial and Temporal Feature Aggregation Graph Convolutional Network for Skeleton-Based Action Recognition -- Improving Video Representation of Vision-Language Model with Decoupled Explicit Temporal Modeling -- KS-FuseNet: An efficient action recognition method based on keyframe selection and feature fusion -- Dynamic Skeleton Association Transformer for dyadic Interaction Action RecognitionSpecies-Aware Guidance for Animal Action Recognition with Vision-Language Knowledge.

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.

