

1. Record Nr.	UNISA996389322303316
Autore	Clapham Henoch
Titolo	A description of new Jerushalem being the substaunce of two sermons deliuered at Paules Crosse [[electronic resource]] : Containing, a briefe discouery and conuiction of certayne doctrines held of Romanists and Brownists against the Catholike and Apostolike faith. / / By Henoch Clapham
Pubbl/distr/stampa	Printed at London, : by Valentine Simmes., 1601
Descrizione fisica	[10], 113, [1] p
Soggetti	Second Advent Brownists Sermons, English - 17th century
Lingua di pubblicazione	Inglese
Formato	Materiale a stampa
Livello bibliografico	Monografia
Note generali	Printer's device (McK. 332) on t.p; initials; head- and tail-pieces. Errata: verso of p. 113. Signatures: A-G H. Reproduction of the original in: Peterborough Cathedral.
Sommario/riassunto	eebo-0124

2. Record Nr.	UNISA996517755103316
Titolo	Computer Vision - ACCV 2022 : 16th Asian Conference on Computer Vision, Macao, China, December 4-8, 2022, Proceedings, Part V // Lei Wang [and four others] editors
Pubbl/distr/stampa	Cham, Switzerland : , : Springer, , [2023] ©2023
ISBN	3-031-26348-0
Edizione	[First edition.]
Descrizione fisica	1 online resource (546 pages)
Collana	Lecture Notes in Computer Science Series ; ; Volume 13845
Disciplina	006.37
Soggetti	Computer vision
Lingua di pubblicazione	Inglese
Formato	Materiale a stampa
Livello bibliografico	Monografia
Nota di contenuto	Intro -- Preface -- Organization -- Contents - Part V -- Recognition: Feature Detection, Indexing, Matching, and Shape Representation -- Improving Few-shot Learning by Spatially-aware Matching and CrossTransformer*-12pt -- 1 Introduction -- 2 Related Work -- 3 Approach -- 3.1 Spatially-aware Few-shot Learning -- 3.2 Self-supervised Scale and Scale Discrepancy -- 3.3 Transformer-Based Spatially-Aware Pipeline -- 4 Experiments -- 4.1 Datasets -- 4.2 Performance Analysis -- 5 Conclusions -- References -- AONet: Attentional Occlusion-Aware Network for Occluded Person Re-identification*-12pt -- 1 Introduction -- 2 Related Works -- 3 Attentional Occlusion-Aware Network -- 3.1 Landmark Patterns and Memorized Dictionary -- 3.2 Attentional Latent Landmarks -- 3.3 Referenced Response Map -- 3.4 Occlusion Awareness -- 3.5 Training and Inference -- 4 Experiments -- 4.1 Datasets and Implementations -- 4.2 Comparisons to State-of-the-Arts -- 4.3 Ablation Studies -- 5 Conclusion -- References -- FFD Augmentor: Towards Few-Shot Oracle Character Recognition from Scratch -- 1 Introduction -- 2 Related Works -- 2.1 Oracle Character Recognition -- 2.2 Few-Shot Learning -- 2.3 Data Augmentation Approaches -- 2.4 Non-Rigid Transformation -- 3 Methodology -- 3.1 Problem Formulation -- 3.2 Overview of Framework -- 3.3 FFD Augmentor -- 3.4 Training with FFD Augmentor -- 4 Experiments -- 4.1 Experimental Settings -- 4.2

Evaluation of FFD Augmented Training -- 4.3 Further Analysis of FFD Augmentor -- 4.4 Applicability to Other Problems -- 5 Conclusion -- References -- Few-shot Metric Learning: Online Adaptation of Embedding for Retrieval*-12pt -- 1 Introduction -- 2 Related Work -- 2.1 Metric Learning -- 2.2 Few-shot Classification -- 3 Few-shot Metric Learning -- 3.1 Metric Learning Revisited -- 3.2 Problem Formulation of Few-shot Metric Learning -- 4 Methods. 4.1 Simple Fine-Tuning (SFT) -- 4.2 Model-Agnostic Meta-Learning (MAML) -- 4.3 Meta-Transfer Learning (MTL) -- 4.4 Channel-Rectifier Meta-Learning (CRML) -- 5 Experiments -- 5.1 Experimental Settings -- 5.2 Effectiveness of Few-shot Metric Learning -- 5.3 Influence of Domain Gap Between Source and Target -- 5.4 Few-shot Metric Learning vs. Few-shot Classification -- 5.5 Results on miniDeepFashion -- 6 Conclusion -- References -- 3D Shape Temporal Aggregation for Video-Based Clothing-Change Person Re-identification*-12pt -- 1 Introduction -- 2 Related Work -- 3 Method -- 3.1 Parametric 3D Human Estimation -- 3.2 Identity-Aware 3D Shape Generation -- 3.3 Difference-Aware Shape Aggregation -- 3.4 Appearance and Shape Fusion -- 4 VCCR Dataset -- 4.1 Collection and Labelling -- 4.2 Statistics and Comparison -- 4.3 Protocol -- 5 Experiments -- 5.1 Implementation Details -- 5.2 Evaluation on CC Re-Id Datasets -- 5.3 Evaluation on Short-Term Re-Id Datasets -- 5.4 Ablation Study -- 6 Conclusion -- References -- Robustizing Object Detection Networks Using Augmented Feature Pooling*-4pt -- 1 Introduction -- 2 Related Works -- 3 Two Challenges of Object Detection for Rotation Robustness -- 4 Proposed Rotation Robust Object Detection Framework -- 4.1 Architecture of Augmented Feature Pooling -- 4.2 Applying to Object Detection Networks -- 5 Experiments -- 5.1 Setting -- 5.2 Effectiveness of Augmented Feature Pooling -- 5.3 Applicability to Modern Object Detection Architectures -- 6 Conclusions -- References -- Reading Arbitrary-Shaped Scene Text from Images Through Spline Regression and Rectification -- 1 Introduction -- 2 Related Work -- 3 Methodology -- 3.1 Text Localization via Spline-Based Shape Regression -- 3.2 Spatial Rectification of Text Features -- 3.3 Text Recognition -- 3.4 Text Spotting Loss -- 4 Experiments -- 4.1 Datasets. 4.2 Implementation Details -- 4.3 Ablation Study -- 4.4 Comparison with State-of-the-Arts -- 4.5 Qualitative Results -- 5 Conclusions -- References -- IoU-Enhanced Attention for End-to-End Task Specific Object Detection*-12pt -- 1 Introduction -- 2 Related Work -- 2.1 Dense Detector -- 2.2 Sparse Detector -- 3 Method -- 3.1 Preliminaries and Analysis on Sparse R-CNN -- 3.2 IoU-Enhanced Self Attention (IoU-ESA) -- 3.3 Dynamic Channel Weighting (DCW) -- 3.4 Training Objectives -- 4 Experiment -- 4.1 Ablation Study -- 4.2 Main Results -- 5 Conclusion -- References -- HAZE-Net: High-Frequency Attentive Super-Resolved Gaze Estimation in Low-Resolution Face Images -- 1 Introduction -- 2 Related Work -- 3 Method -- 3.1 Super-Resolution Module -- 3.2 Gaze Estimation Module -- 3.3 Loss Function -- 3.4 Alternative End-to-End Learning -- 4 Experiments -- 4.1 Datasets and Evaluation Metrics -- 4.2 Performance Comparison by Module -- 4.3 Comparison Under LR Conditions -- 4.4 Ablation Study -- 5 Conclusion -- References -- LatentGaze: Cross-Domain Gaze Estimation Through Gaze-Aware Analytic Latent Code Manipulation -- 1 Introduction -- 2 Related Work -- 2.1 Latent Space Embedding and GAN Inversion -- 2.2 Domain Adaptation and Generalization -- 2.3 Gaze Estimation -- 3 LatentGaze -- 3.1 Preliminaries -- 3.2 The LatentGaze Framework Overview -- 3.3 Statistical Latent Editing -- 3.4 Domain Shift Module -- 4 Experiments -- 4.1 Datasets and Setting --

4.2 Comparison with Gaze Estimation Models on Single Domain -- 4.3 Comparison with Gaze Estimation Models on Cross Domain -- 4.4 Ablation Study -- 5 Conclusion -- References -- Cross-Architecture Knowledge Distillation -- 1 Introduction -- 2 Related Work -- 3 The Proposed Method -- 3.1 Framework -- 3.2 Cross-architecture Projector -- 3.3 Cross-view Robust Training -- 3.4 Optimization -- 4 Experiments -- 4.1 Settings.

4.2 Performance Comparison -- 4.3 Ablation Study -- 5 Conclusions -- References -- Cross-Domain Local Characteristic Enhanced Deepfake Video Detection -- 1 Introduction -- 2 Related Work -- 2.1 Deepfake Detection -- 2.2 Generalization to Unseen Forgeries -- 3 Proposed Method -- 3.1 Overview -- 3.2 Data Preprocessing -- 3.3 Cross-Domain Local Forensics -- 4 Experiment and Discussion -- 4.1 Experiment Setup -- 4.2 In-dataset Evaluation -- 4.3 Cross-Dataset Evaluation -- 4.4 Ablation Study -- 5 Conclusion -- References -- Three-Stage Bidirectional Interaction Network for Efficient RGB-D Salient Object Detection -- 1 Introduction -- 2 Related Work -- 2.1 RGB-D SOD -- 2.2 Efficient RGB-D SOD -- 3 Proposed Method -- 3.1 Overview -- 3.2 Three-Stage Bidirectional Interaction (TBI) -- 3.3 Three-Stage Refinement Decoder -- 3.4 Loss Function -- 4 Experiments -- 4.1 Experimental Setup -- 4.2 Comparisons with SOTA Methods -- 4.3 Ablation Studies -- 5 Conclusion -- References -- PS-ARM: An End-to-End Attention-Aware Relation Mixer Network for Person Search -- 1 Introduction -- 1.1 Motivation -- 1.2 Contribution -- 2 Related Work -- 3 Method -- 3.1 Overall Architecture -- 3.2 Attention-aware Relation Mixer (ARM) Module -- 3.3 Training and Inference -- 4 Experiments -- 4.1 Dataset and Evaluation Protocols -- 4.2 Comparison with State-of-the-Art Methods -- 4.3 Ablation Study -- 5 Conclusions -- References -- Weighted Contrastive Hashing -- 1 Introduction -- 2 Related Works -- 3 Weighted Contrastive Learning -- 3.1 Preliminaries -- 3.2 Overall Architecture -- 3.3 Mutual Attention -- 3.4 Weighted Similarities Calculation -- 3.5 Training and Inference -- 4 Discussion -- 5 Experiments -- 5.1 Datasets and Evaluation Metrics -- 5.2 Implementation Details -- 5.3 Comparison with the SotA -- 5.4 Ablation Studies -- 5.5 Visualization and Hyper-parameters -- 6 Conclusion.

References -- Content-Aware Hierarchical Representation Selection for Cross-View Geo-Localization -- 1 Introduction -- 2 Related Work -- 3 Hierarchical Enhancement Coefficient Map -- 4 Adaptive Residual Fusion Mechanism -- 5 Experiments -- 5.1 Experimental Settings -- 5.2 Analyses the CA-HRS Module -- 5.3 Comparison with State of the Arts -- 6 Conclusion -- References -- CLUE: Consolidating Learned and Undergoing Experience in Domain-Incremental Classification -- 1 Introduction -- 2 Related Work -- 3 Method -- 3.1 Problem Statement -- 3.2 Consolidating Learned and Undergoing Experience (CLUE) -- 4 Results -- 4.1 Performance Comparison -- 4.2 Buffer Size Analysis -- 4.3 Model Interpretability Analysis -- 4.4 Ablation Studies -- 5 Conclusion -- References -- Continuous Self-study: Scene Graph Generation with Self-knowledge Distillation and Spatial Augmentation -- 1 Introduction -- 2 Related Work -- 3 Methodology -- 3.1 Self-study Module -- 3.2 Spatial Augmentation Module -- 3.3 Scene Graph Generation -- 4 Experiment -- 4.1 Experimental Settings -- 4.2 Implementation Details -- 4.3 Experiment Results -- 4.4 Quantitative Studies -- 4.5 Visualization Results -- 5 Conclusion -- References -- Spatial Group-Wise Enhance: Enhancing Semantic Feature Learning in CNN -- 1 Introduction -- 2 Related Work -- 3 Method -- 4 Experiments -- 4.1 Image Classification -- 4.2 Object Detection -- 4.3 Visualization and Interpretation -- 5 Conclusion -- References -- SEIC:

Semantic Embedding with Intermediate Classes for Zero-Shot Domain Generalization -- 1 Introduction -- 2 Related Work -- 3 Problem Definition -- 4 Proposed Method -- 4.1 Handling Unseen Domains Using Domain Generalization -- 4.2 Handling Unseen Classes Using the Proposed SEIC Framework -- 5 Experimental Evaluation -- 5.1 Results on DomainNet and DomainNet-LS Datasets -- 5.2 Additional Analysis -- 6 Conclusion.
References.

Sommario/riassunto

The 7-volume set of LNCS 13841-13847 constitutes the proceedings of the 16th Asian Conference on Computer Vision, ACCV 2022, held in Macao, China, December 2022. The total of 277 contributions included in the proceedings set was carefully reviewed and selected from 836 submissions during two rounds of reviewing and improvement. The papers focus on the following topics: Part I: 3D computer vision; optimization methods; Part II: applications of computer vision, vision for X; computational photography, sensing, and display; Part III: low-level vision, image processing; Part IV: face and gesture; pose and action; video analysis and event recognition; vision and language; biometrics; Part V: recognition: feature detection, indexing, matching, and shape representation; datasets and performance analysis; Part VI: biomedical image analysis; deep learning for computer vision; Part VII: generative models for computer vision; segmentation and grouping; motion and tracking; document image analysis; big data, large scale methods.
