

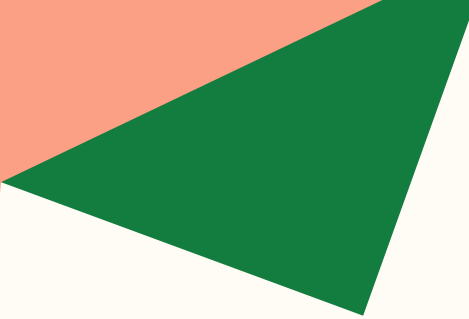
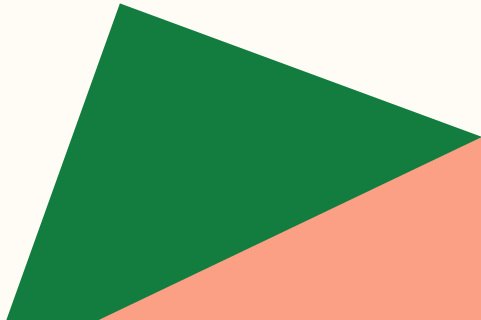


Image Feature Matching: SuperPoint and SuperGlue

IFT 6757 - Duckietown

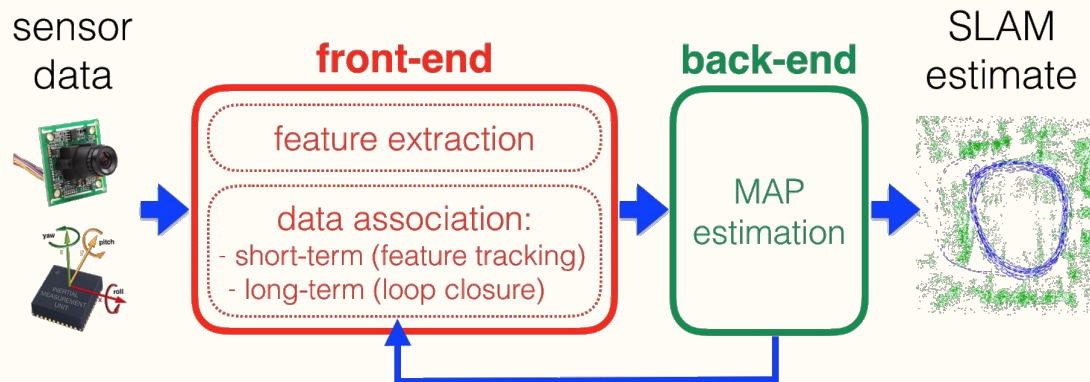


Acknowledgements

- This presentation uses material presented in the two papers:
 - SuperPoint: Self-Supervised Interest Point Detection and Description
 - SuperGlue: Learning Feature Matching with Graph Neural Networks
- This presentation also borrows from two presentations:
 - Deep Visual SLAM Frontends by Tomasz Malisiewicz at CVPR 2020
 - SuperGlue presentation by Paul-Edouard Sarlin at CVPR 2020
- All images used are from the above-mentioned material, unless noted otherwise.

Motivation

- Two parts of Visual SLAM
- Front-end
 - Feature extraction (SIFT, SURF, ORB)
 - Data Association (feature tracking, loop closure)
- Back-end
 - Optimize pose and 3D structure



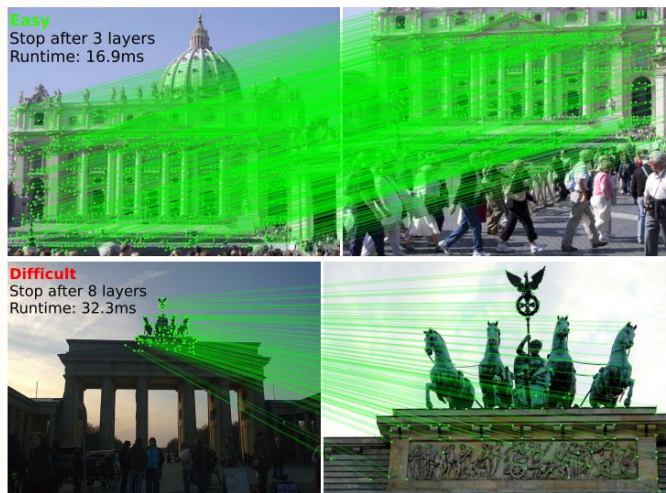
Source: C. Cadena, L. Carlone, et al. "Past, Present, and Future of Simultaneous Localization And Mapping: Towards the Robust-Perception Age"

Motivation



Why is it so difficult?

- Appearance changes: illumination, weather, sensor noise
- Viewpoint changes: perspective distortion, scale, occlusion
- Repetitive structures: many false matches (windows, tiles, trees)
- Textureless regions: few or no features (white walls, roads)



Source: P. Lindenberger, P. Sarlin, et al. "LightGlue: Local Feature Matching at Light Speed"

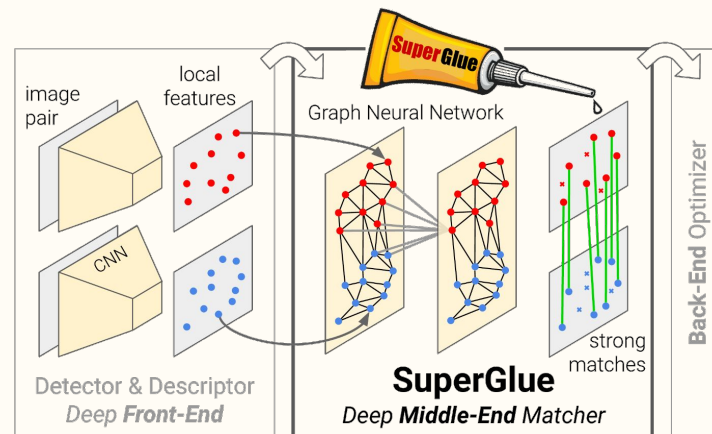
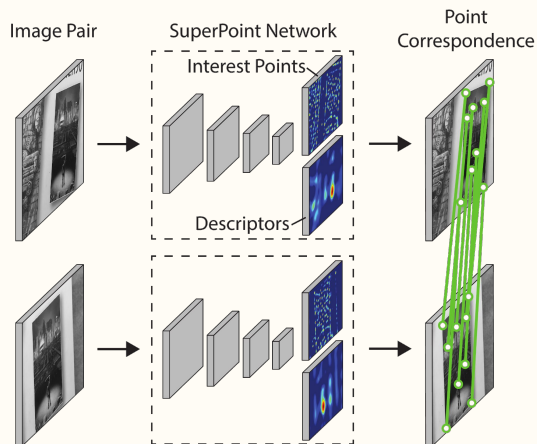
Why is it so difficult?



Source: J. Wang, N. Karaev, et al. "Visual Geometry Grounded Deep Structure From Motion"

Research Questions

- SuperPoint
 - How to learn detectors & descriptors?
 - Can we train without labeled correspondences (self-supervised)?
 - How does it compare with earlier pipelines (e.g., LIFT, SIFT)?
- SuperGlue
 - How to learn the data association between two images?



Preliminaries

- Learning-based Approaches
 - Deep Learning & Neural Networks
 - Convolutional Neural Networks (CNNs)
 - Graph Neural Networks (GNNs)
- Matching with Heuristics
 - Nearest Neighbor
 - Ratio Test
 - Mutual Consistency Test
- Geometric Verification
 - Homography & Pose Estimation
 - Direct Linear Transform (DLT)
 - Essential & Fundamental Matrices
 - 5-point / 8-point Algorithms
 - Random Sample Consensus (RANSAC)
- Optimal Assignment Methods
 - Bipartite Matching
 - Hungarian Algorithm
 - Sinkhorn Iteration

Evaluation

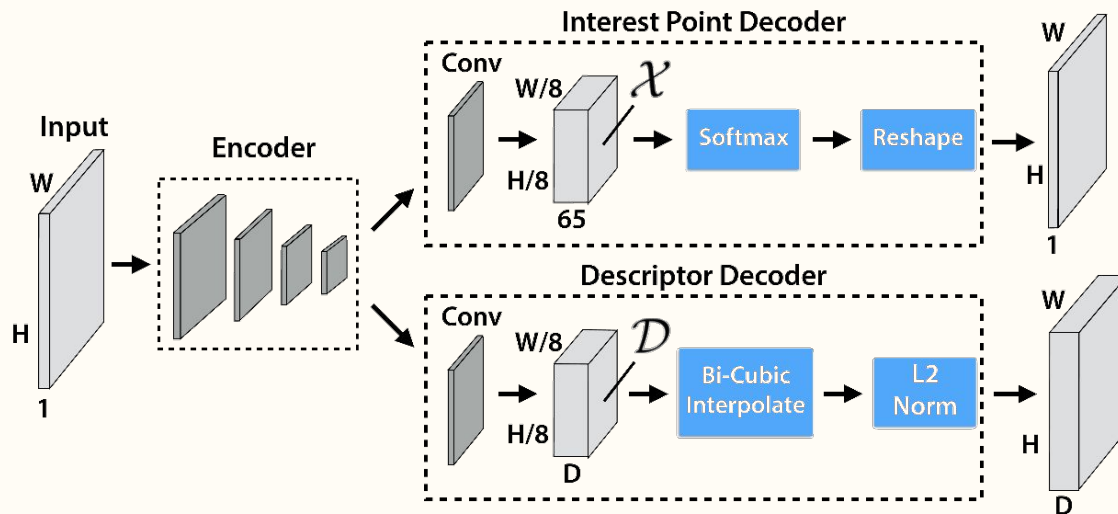
- **Detector Metrics:**
 - **Corner Detection Average Precision (AP)** - Measures how well detected points align with ground-truth corners (Precision–Recall AUC). Higher AP = better.
 - **Localization Error (LE)** - Average pixel distance between detected points and their closest ground-truth corner. Lower is better.
 - **Repeatability (Rep)** - Probability a point is re-detected in a second view of the same scene. Higher is better.
- **Descriptor and Matching Metrics**
 - **Nearest Neighbor mAP** - Discriminativeness of descriptors via NN matching (Precision–Recall AUC). Higher = better descriptors.
 - **Matching Score (MS)** - Fraction of ground-truth correspondences recovered by the pipeline. Combines detection + description.
 - **Homography Estimation Accuracy** - Ability to recover the ground-truth homography by transforming image corners correctly.
 - **Pose Estimation mAP** - Area under the curve of relative pose error.

SuperPoint

The art and craft of designing neural nets
to replace SIFT.

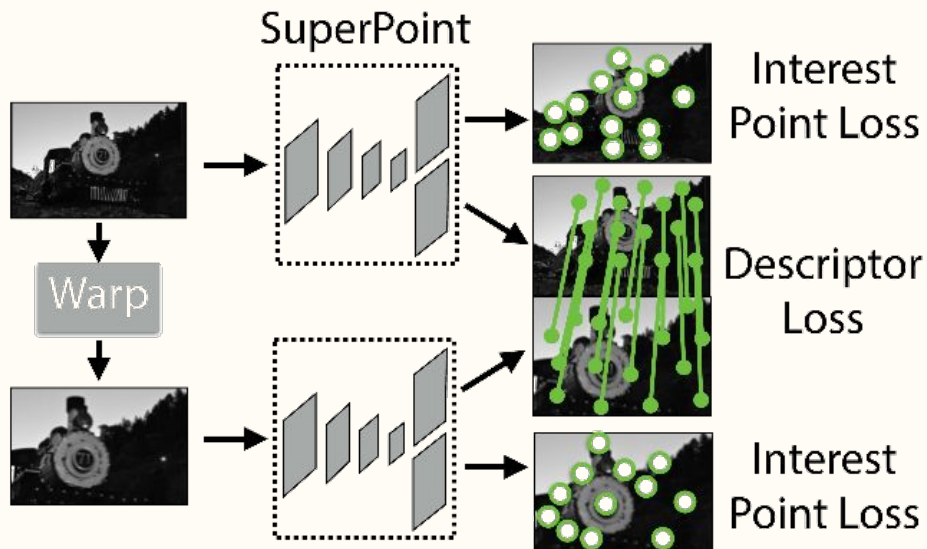
Model Architecture

- CNN architecture
 - VGG-like backbone
- Points + descriptors computed jointly, no patches
- No deconvolution layer



Model Training

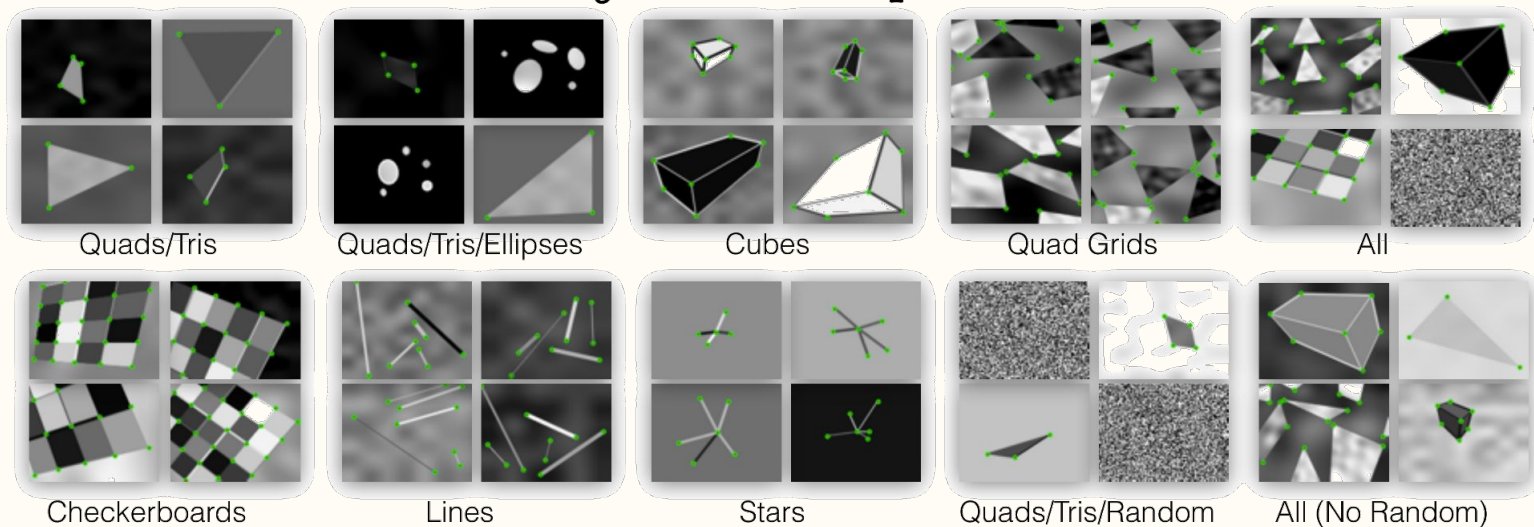
- Siamese training with pairs of images
- Descriptor trained via metric learning (contrastive loss)
- Keypoints trained via supervised keypoint labels
 - Where do these come from?



How to get Keypoint Labels?

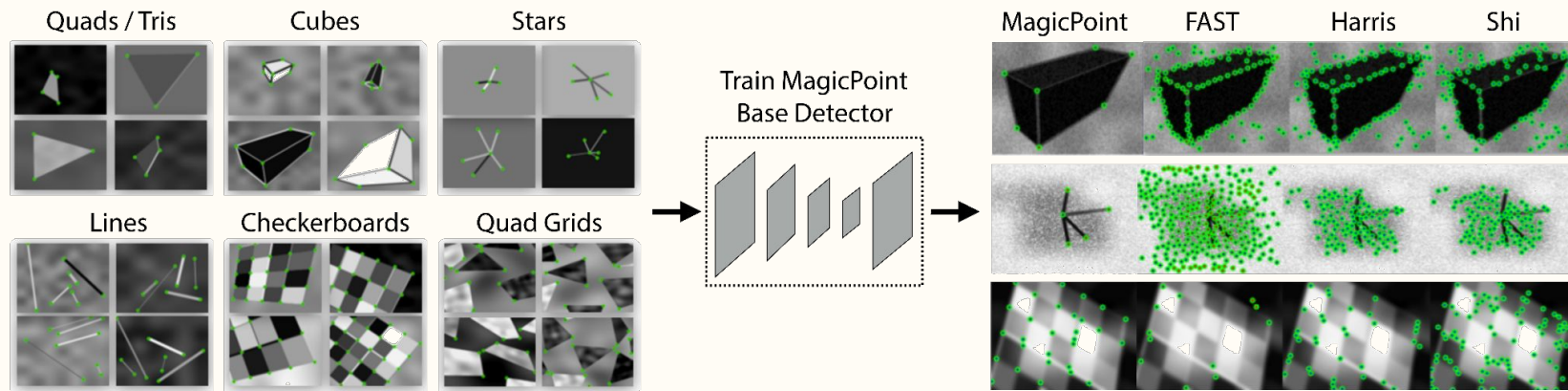
- Need large scale dataset of annotated images
- Too hard for humans to label

Synthetic Shapes



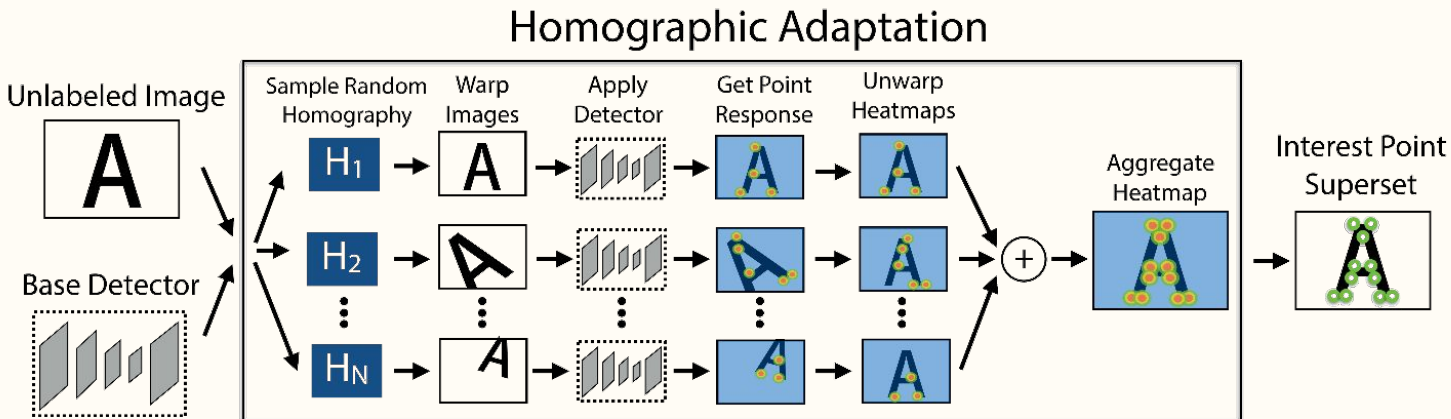
Self-Supervised Training

- Synthetic Training
 - Non-photorealistic shapes
 - Heavy noise
 - Effective and easy



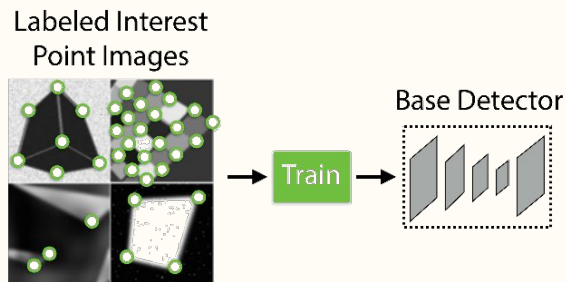
Self-Supervised Training

- Homographic Adaption
 - Use trained "MagicPoint" network
 - Simulate planar camera motion with homographies
 - Enhance repeatable points

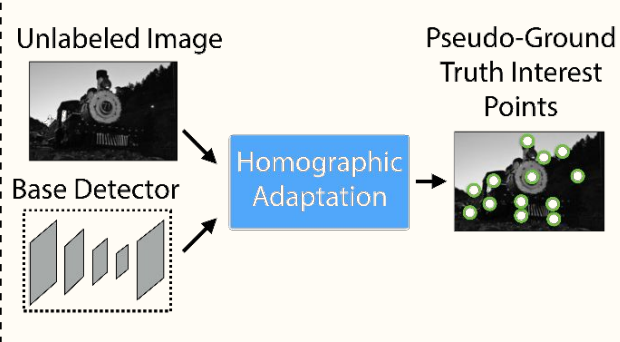


Overview

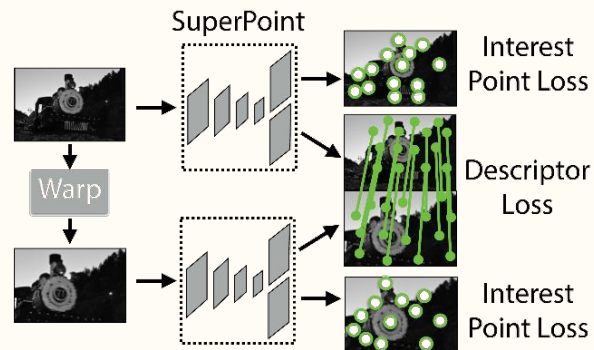
(a) Interest Point Pre-Training



(b) Interest Point Self-Labeling



(c) Joint Training



Metrics

	57 Illumination Scenes		59 Viewpoint Scenes	
	NMS=4	NMS=8	NMS=4	NMS=8
<i>SuperPoint</i>	.652	.631	.503	.484
<i>MagicPoint</i>	.575	.507	.322	.260
<i>FAST</i>	.575	.472	.503	.404
<i>Harris</i>	.620	.533	.556	.461
<i>Shi</i>	.606	.511	.552	.453
<i>Random</i>	.101	.103	.100	.104

Table 3. **HPatches Detector Repeatability.** SuperPoint is the most repeatable under illumination changes, competitive on viewpoint changes, and outperforms MagicPoint in all scenarios.

	Homography Estimation			Detector Metrics		Descriptor Metrics	
	$\epsilon = 1$	$\epsilon = 3$	$\epsilon = 5$	Rep.	MLE	NN mAP	M. Score
<i>SuperPoint</i>	.310	.684	.829	.581	1.158	.821	.470
<i>LIFT</i>	.284	.598	.717	.449	1.102	.664	.315
<i>SIFT</i>	.424	.676	.759	.495	0.833	.694	.313
<i>ORB</i>	.150	.395	.538	.641	1.157	.735	.266

Table 4. **HPatches Homography Estimation.** SuperPoint outperforms LIFT and ORB and performs comparably to SIFT using various ϵ thresholds of correctness. We also report related metrics which measure detector and descriptor performance individually.

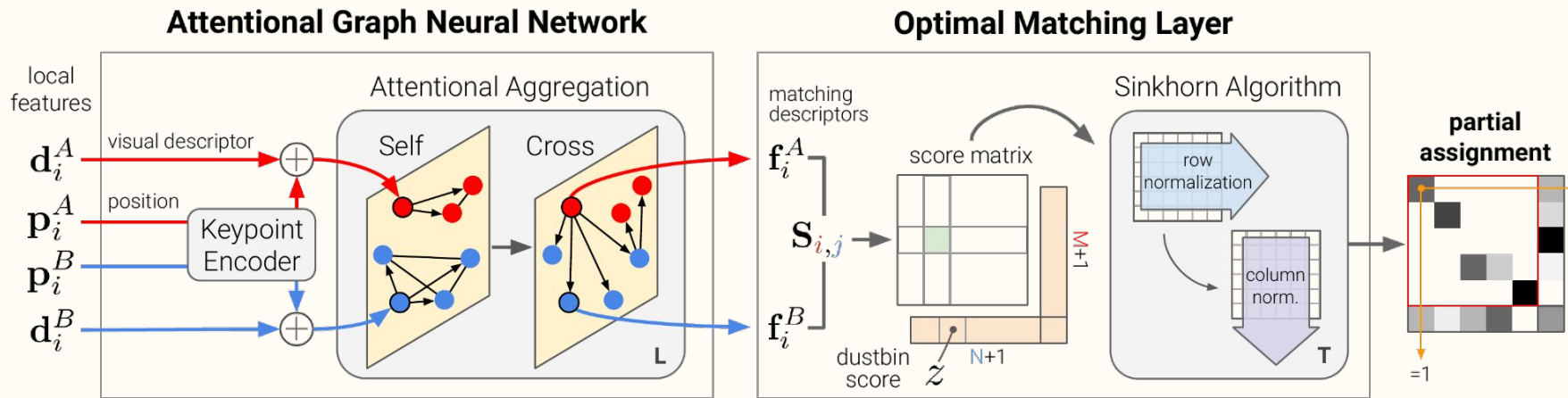
SuperGlue

Deep matching with SuperPoint: Can we learn to solve the correspondence problem?

Overview

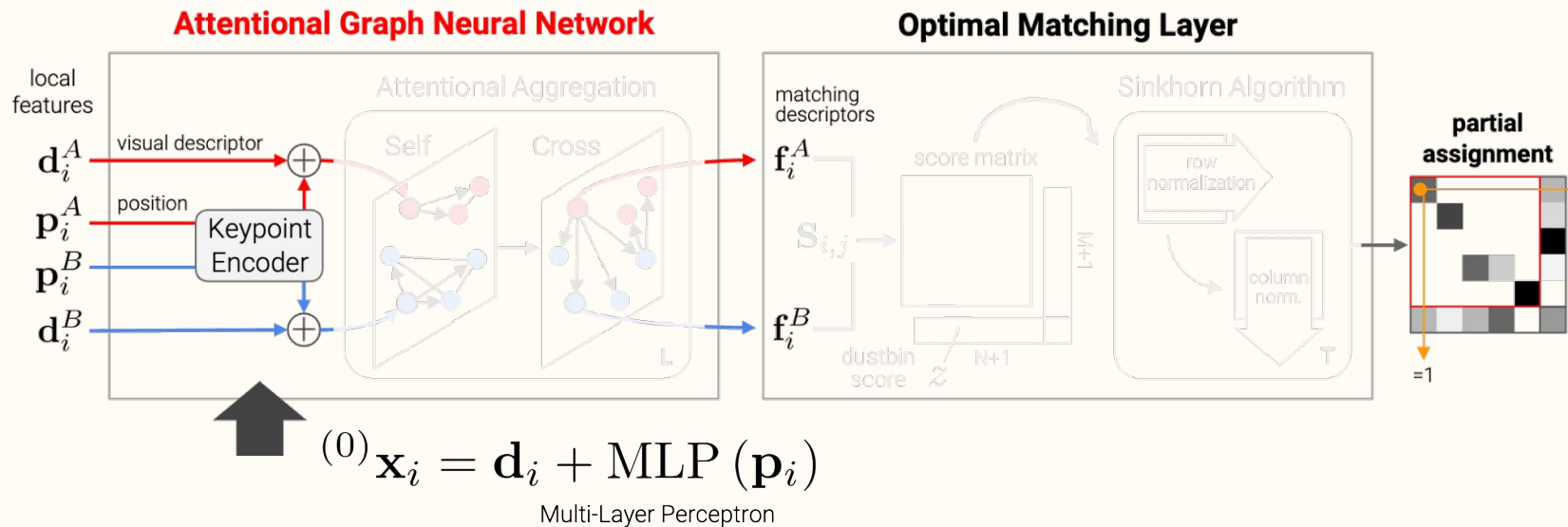
- Inputs
 - Images A and B
 - 2 sets of M, N local features (keypoints and descriptors)
- Outputs
 - Single a match per keypoint + occlusion using soft partial assignment
- Method
 - A Graph Neural Network (GNN) with attention
 - Encodes contextual cues & priors
 - Reasons about the 3D scene
 - Solving a optimal transport problem
 - Differentiable solver
 - Enforces the assignment constraints

Architecture



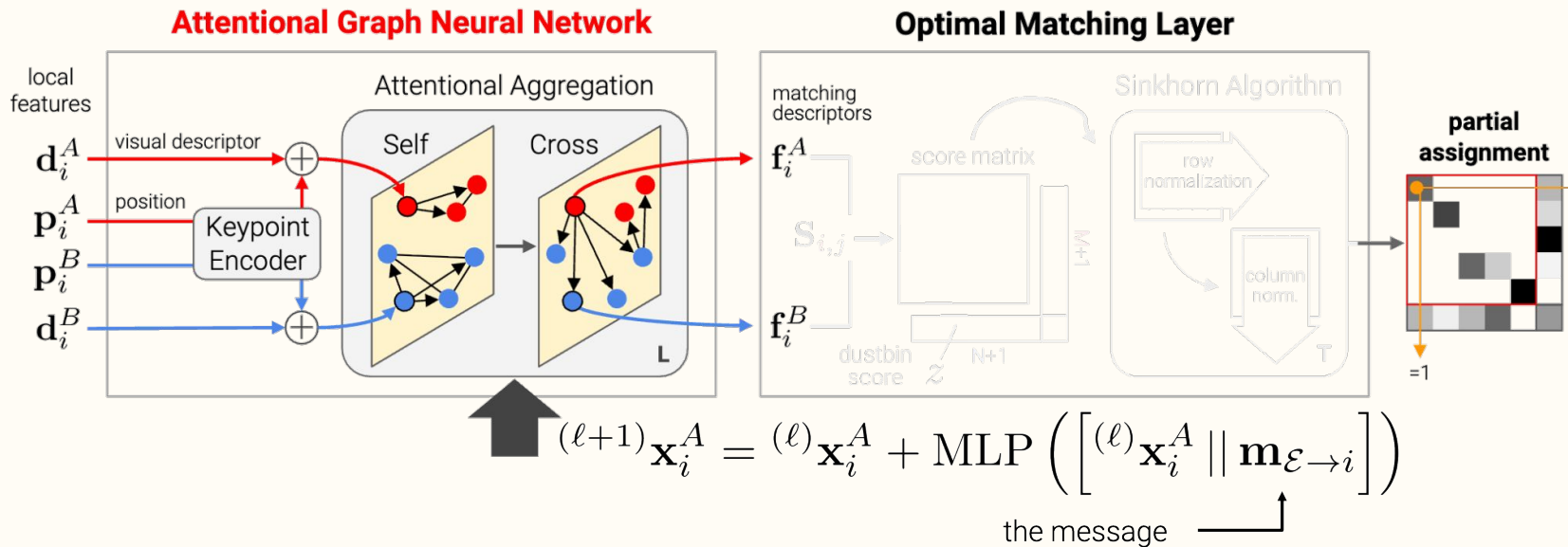
Attentional GNN

- Initial representation for each keypoint



Attentional GNN

- Update the representation based on other keypoints - using a Message Passing Neural Network
 - a complete graph with two types of edges
 - in the same image: "self" edges
 - in the other image: "cross" edges



Attentional Aggregation

- Compute the message using self and cross attention

$$^{(\ell+1)}\mathbf{x}_i^A = ^{(\ell)}\mathbf{x}_i^A + \text{MLP} \left(\left[^{(\ell)}\mathbf{x}_i^A \parallel \mathbf{m}_{\mathcal{E} \rightarrow i} \right] \right)$$

the message $\xrightarrow{\quad}$

$$\mathbf{m}_{\mathcal{E} \rightarrow i} = \sum_{j:(i,j) \in \mathcal{E}} \alpha_{ij} \mathbf{v}_j$$

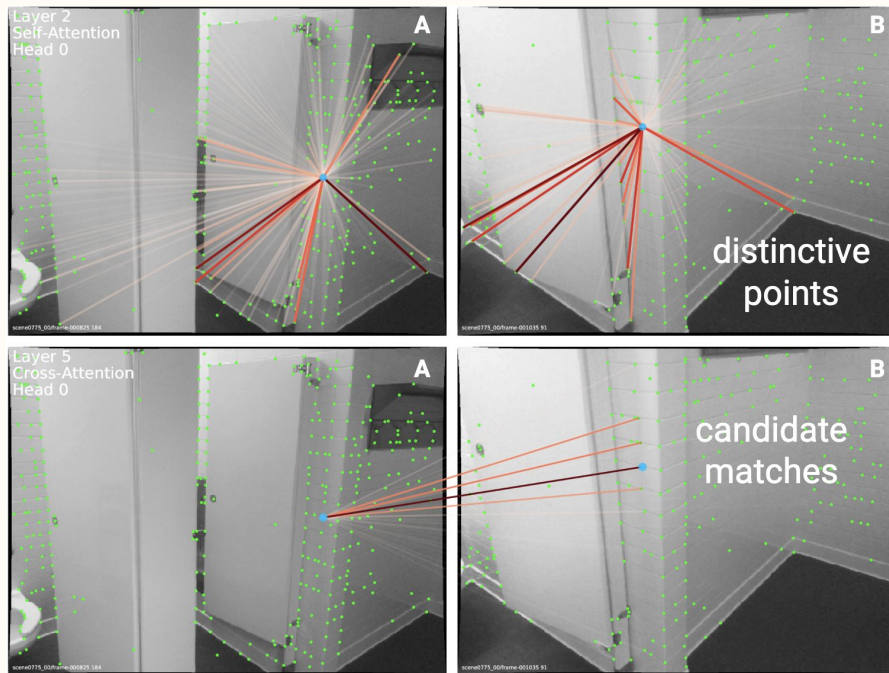
$$\alpha_{ij} = \text{Softmax}_j \left(\mathbf{q}_i^\top \mathbf{k}_j \right)$$

$$\mathbf{q}_i = \mathbf{W}_1 ^{(\ell)}\mathbf{x}_i + \mathbf{b}_1$$

$$\begin{bmatrix} \mathbf{k}_j \\ \mathbf{v}_j \end{bmatrix} = \begin{bmatrix} \mathbf{W}_2 \\ \mathbf{W}_3 \end{bmatrix} ^{(\ell)}\mathbf{x}_j + \begin{bmatrix} \mathbf{b}_2 \\ \mathbf{b}_3 \end{bmatrix}$$

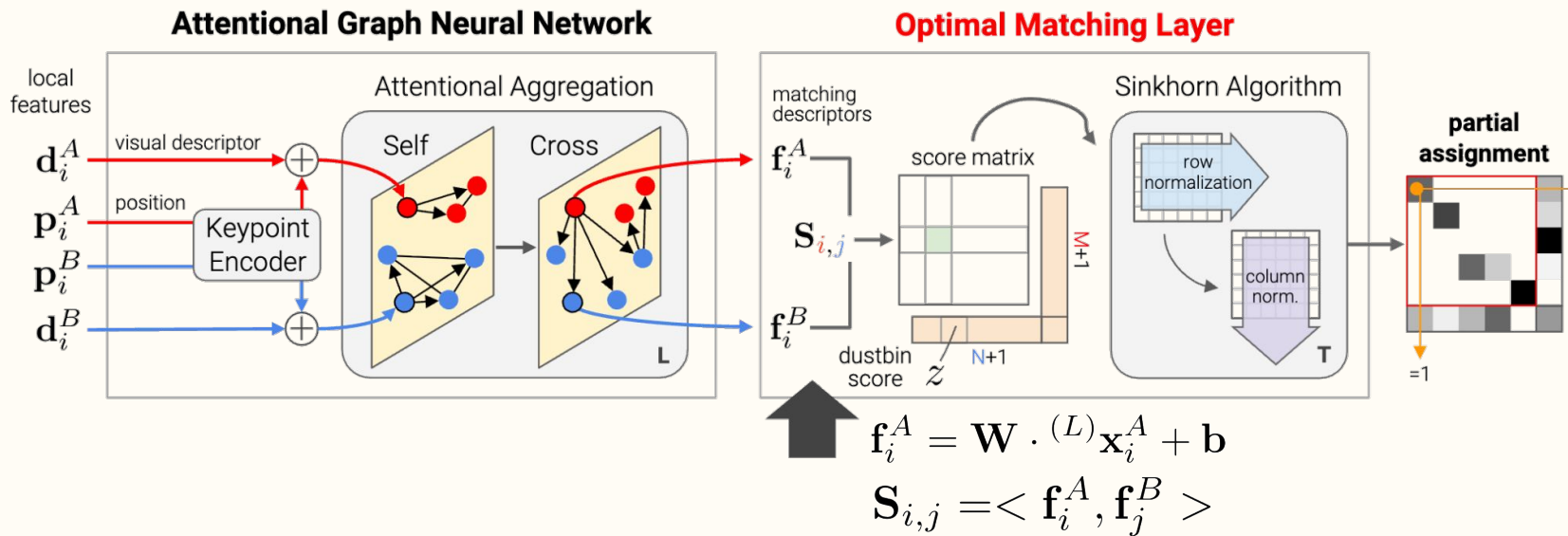
Attentional Aggregation

- Self-attention = intra-image information flow
- Cross attention = inter-image



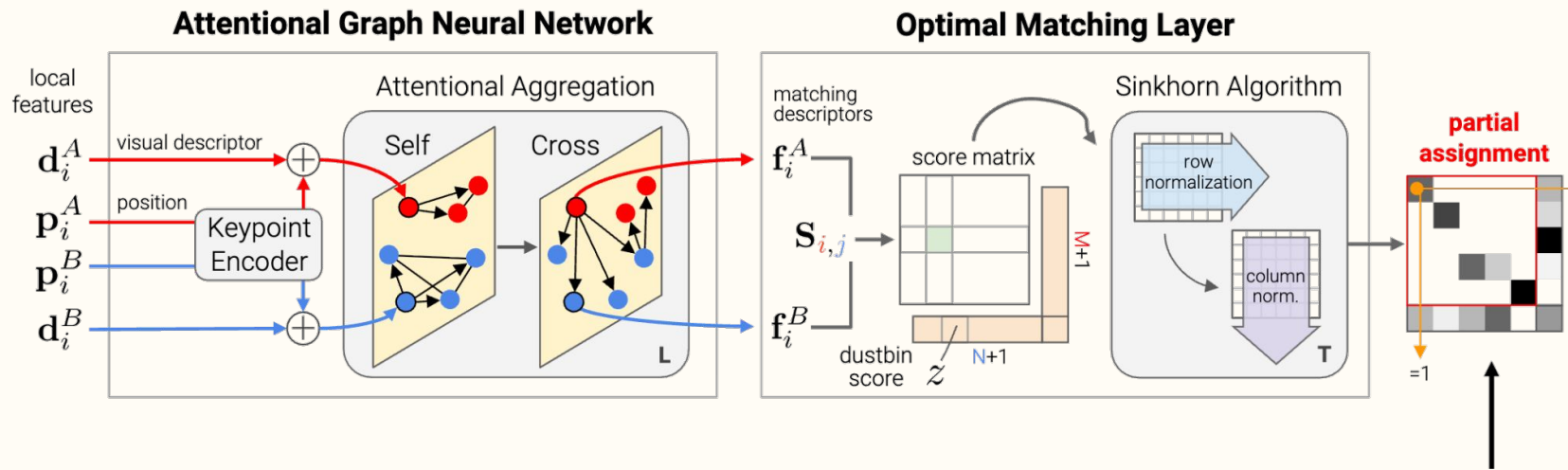
Optimal Matching Layer

- Compute score matrix
- Occlusion and noise : unmatched keypoints are assigned to a dustbin
- Compute the partial assignment matrix
 - With the Sinkhorn algorithm: differentiable & soft Hungarian algorithm

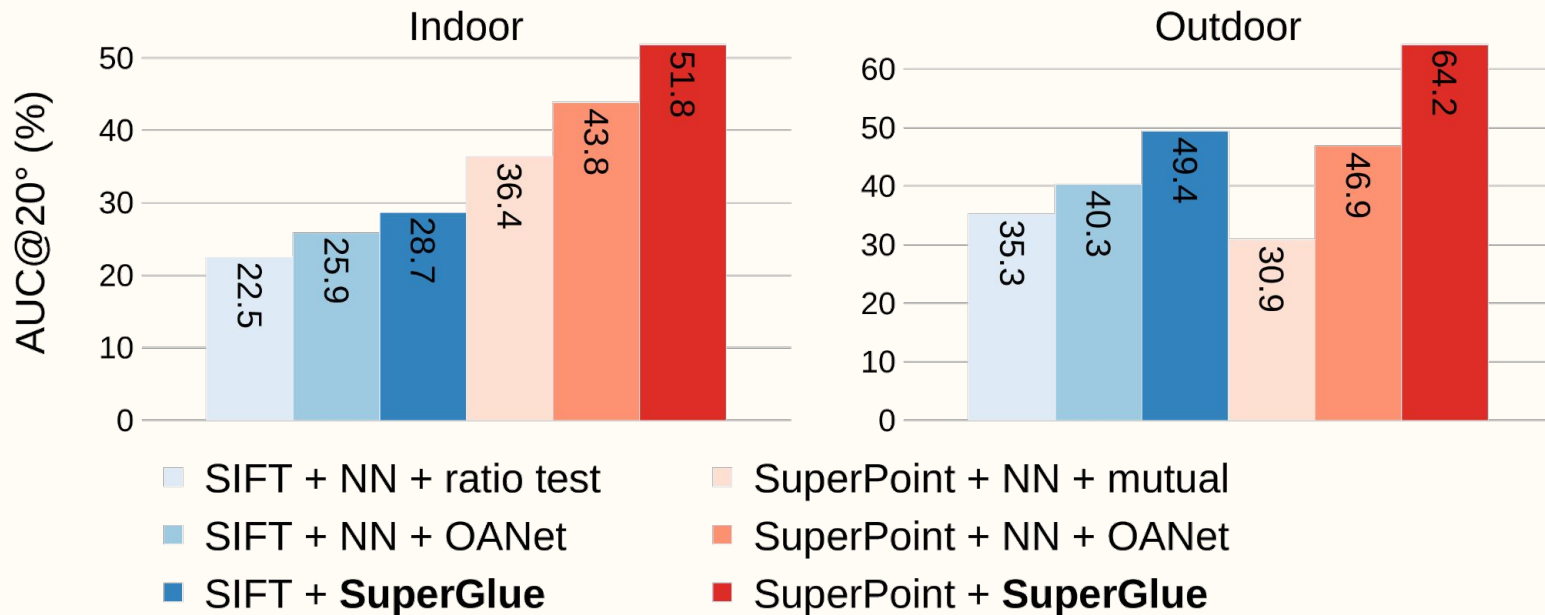


Loss

- Compute ground truth correspondences from pose and depth
- Find which keypoints should be unmatched
- Loss: maximize the log-likelihood of the GT cells in the partial assignment matrix

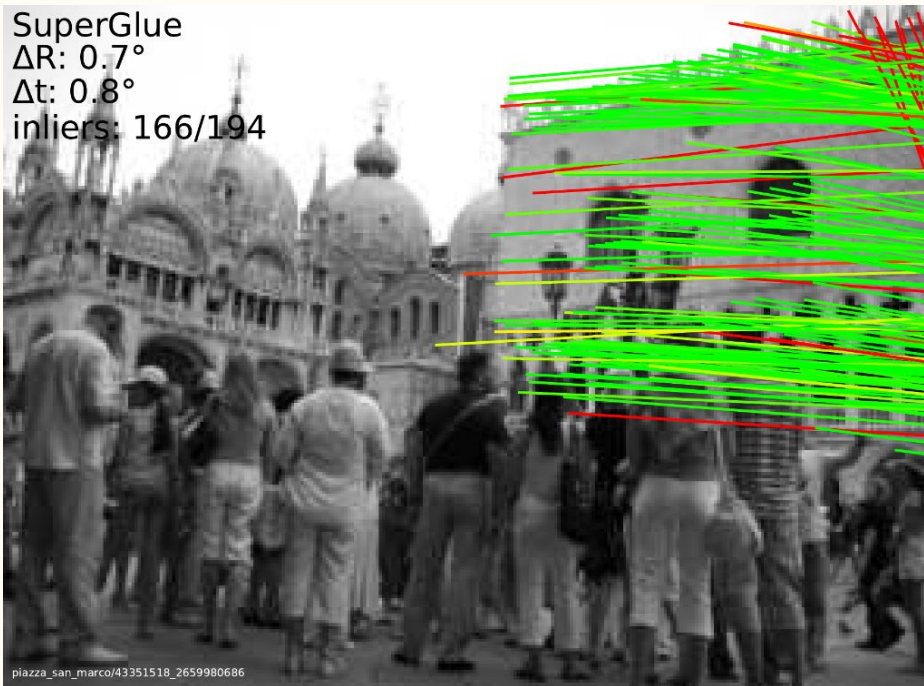


Metrics



Qualitative Results

SuperGlue
 $\Delta R: 0.7^\circ$
 $\Delta t: 0.8^\circ$
inliers: 166/194



piazza_san_marco/43351518_2659980686

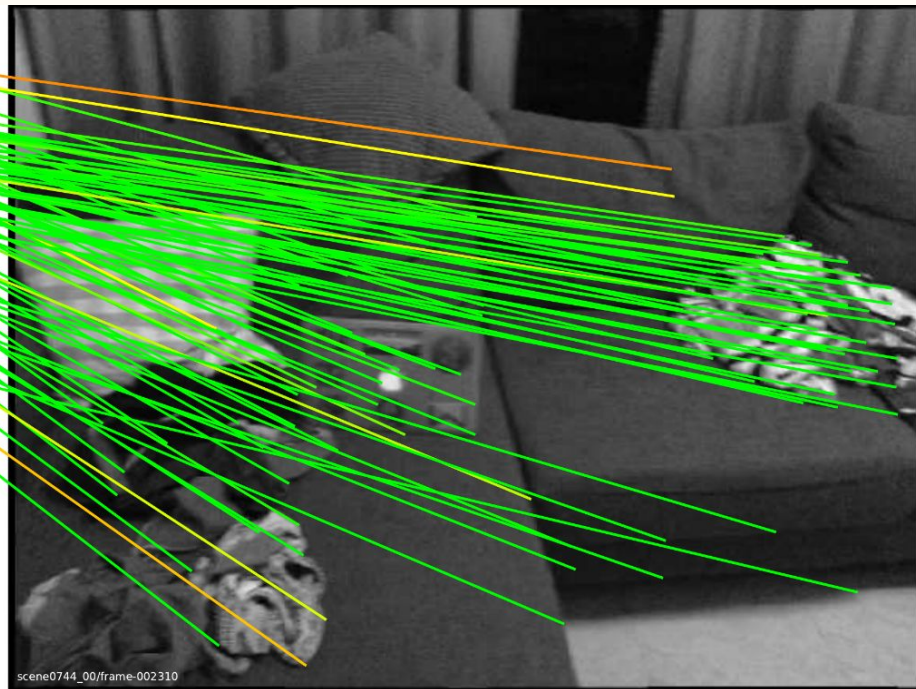


piazza_san_marco/06795901_3725050516

Qualitative Results

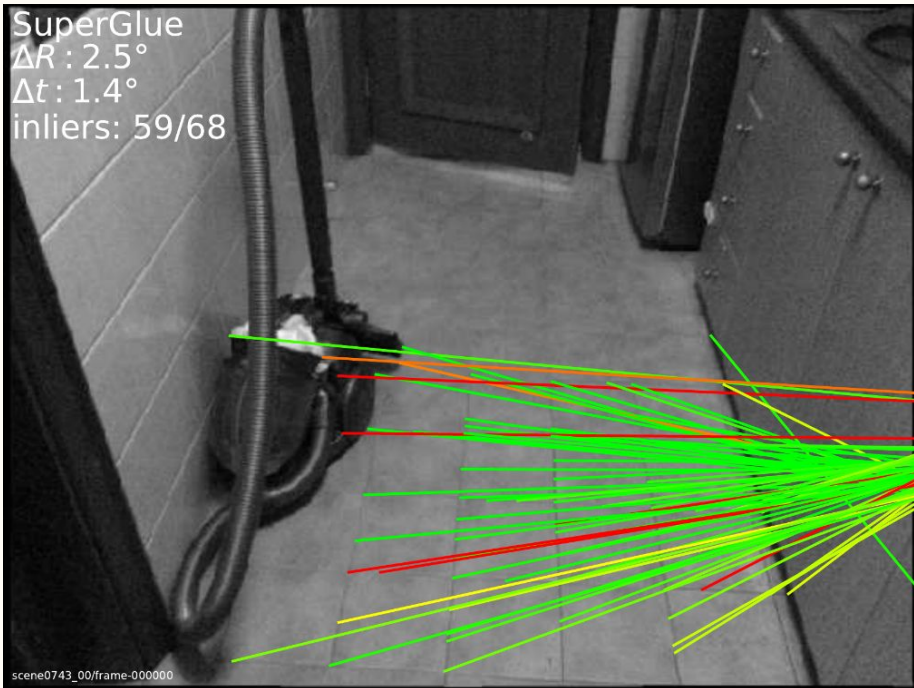


Qualitative Results

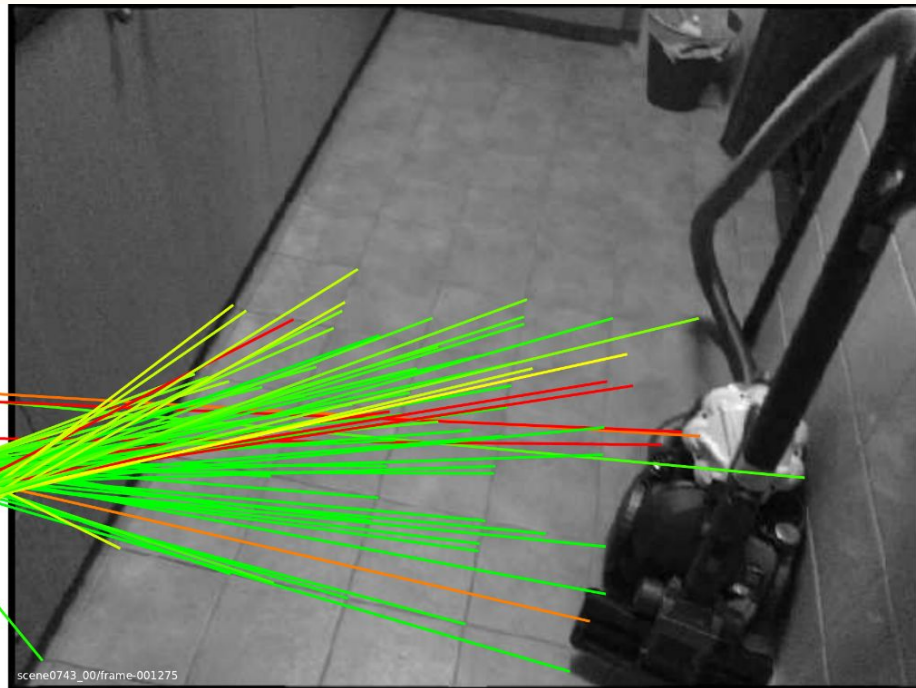


Qualitative Results

SuperGlue
 $\Delta R : 2.5^\circ$
 $\Delta t : 1.4^\circ$
inliers: 59/68

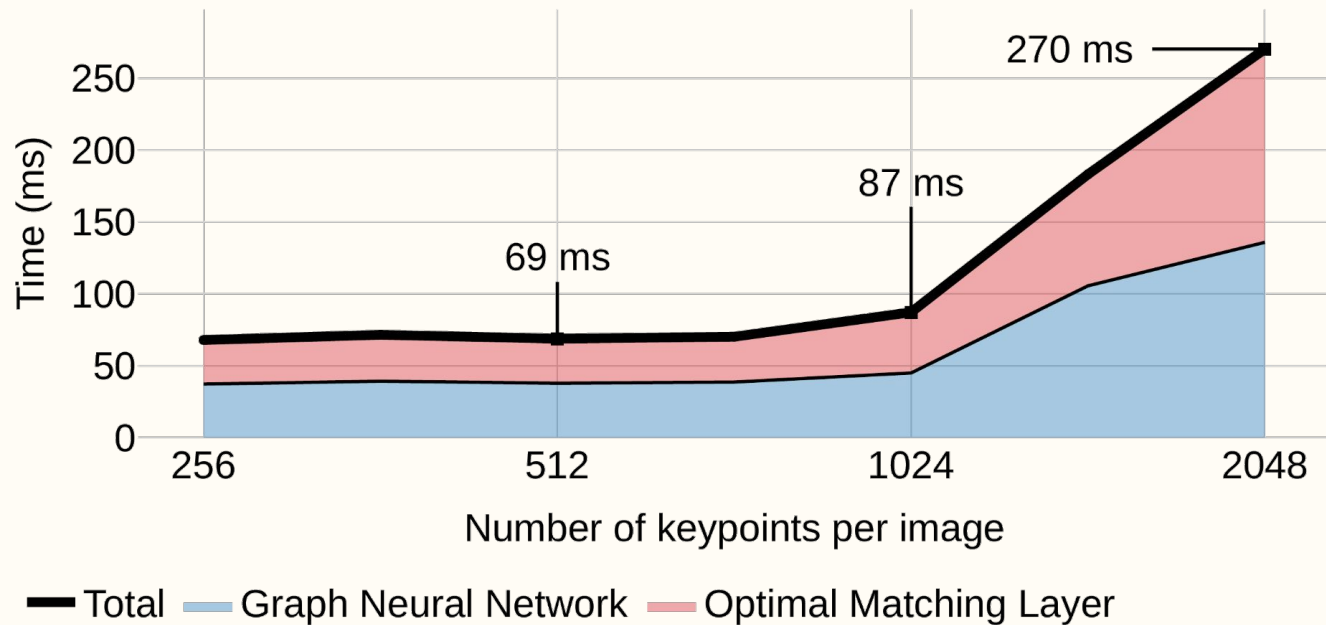


scene0743_00/frame-000000



scene0743_00/frame-001275

Run-time Analysis



Duckie Dataset

Why should they have all the fun?

$N = 1$

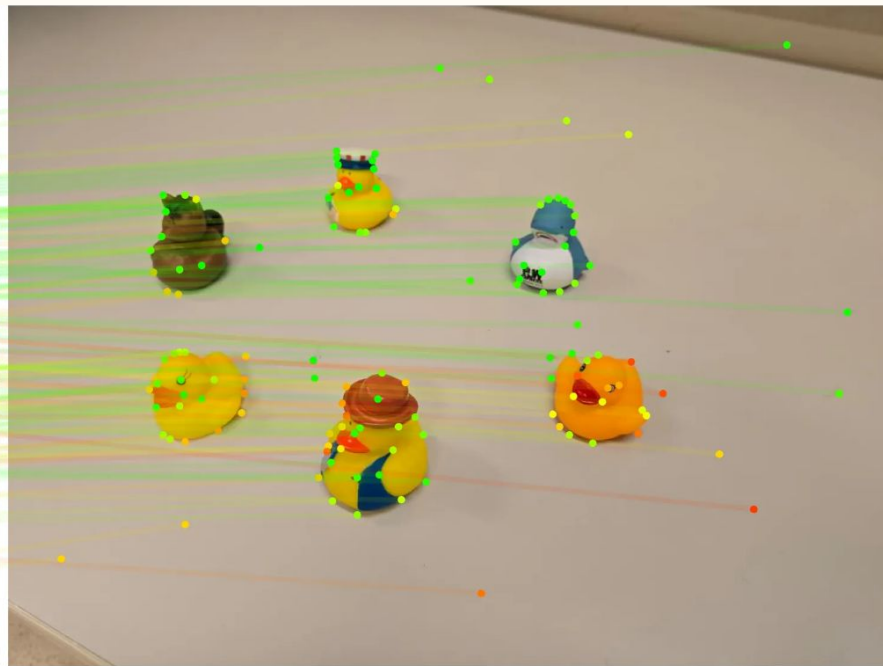
Duckie Dataset

keypoints0: 240
keypoints1: 203



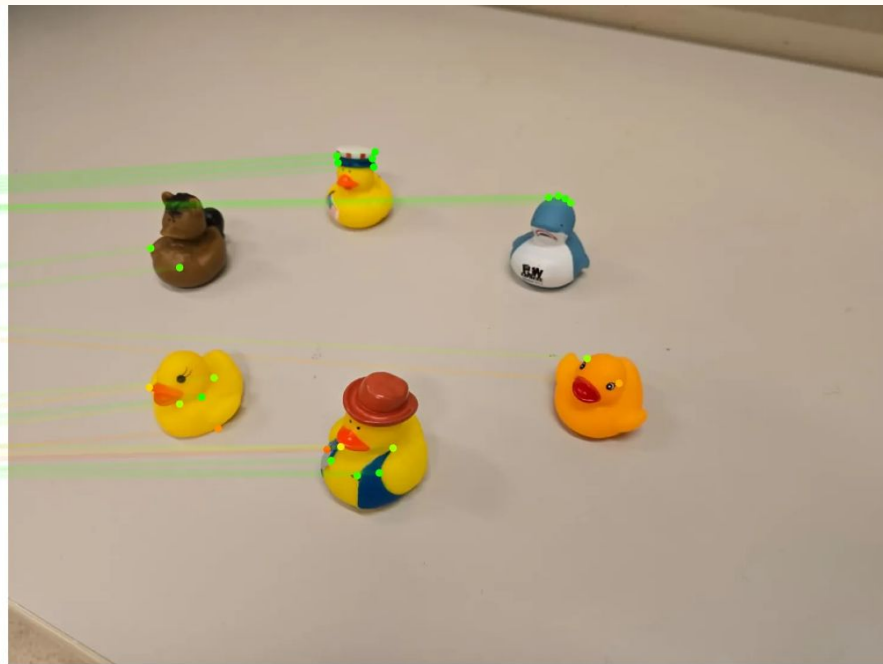
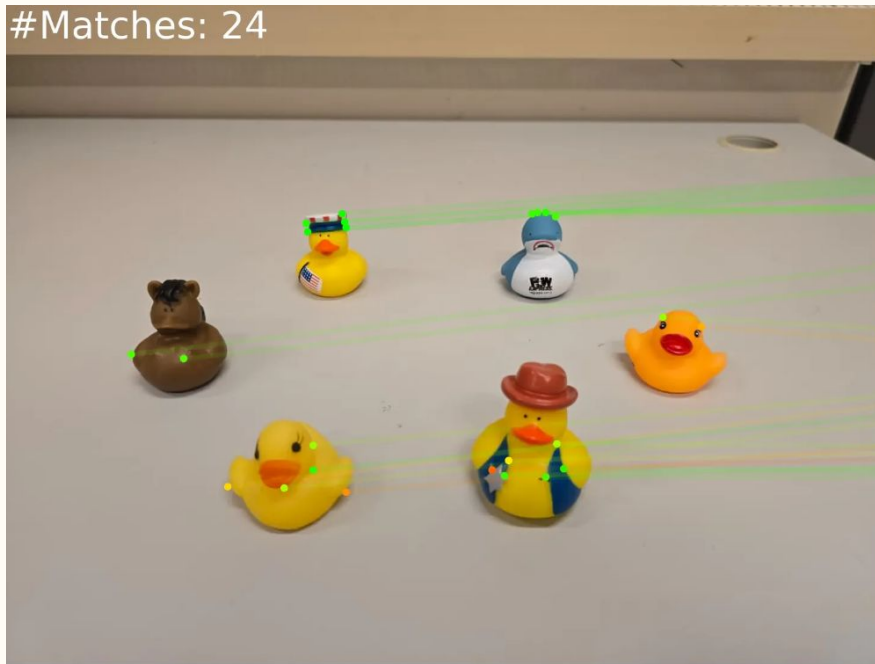
Duckie Dataset

#Matches: 126



Duckie Dataset

#Matches: 24



Conclusions

- Problem: Robust feature detection & matching is critical for SLAM, SfM, and localization, but remains challenging under viewpoint, illumination, and texture changes.
- SuperPoint:
 - Learned detector + descriptor in a single CNN.
 - Self-supervised pretraining (synthetic homographies) enables real-world generalization.
- SuperGlue:
 - Learned matching with graph neural networks and optimal transport.
 - Context-aware, consistent correspondences outperform descriptor-only matching.
- Key lesson: Learning helps both where we detect features (SuperPoint) and how we match them (SuperGlue).

Other Relevant Work

- Earlier works
 - FAST - Features from Accelerated Segment Test
 - LIFT - Learned Invariant Feature Transform
- But can we learn generally useful image features?
 - DINO - Self-Distillation with No Labels
- Did the community just accept this approach?
 - DeDoDe - DeDoDe: Detect, Don't Describe -- Describe, Don't Detect for Local Feature Matching
 - MatchFormer - MatchFormer: Interleaving Attention in Transformers for Feature Matching
 - GlueStick - GlueStick: Robust Image Matching by Sticking Points and Lines Together
- Need for speed
 - LightGlue - Local Feature Matching at Light Speed
- Who needs detectors anyway?
 - LoFTR - Detector-Free Local Feature Matching with Transformers
 - RoMA - Robust Dense Feature Matching
- Scaling transformers for pointmap prediction
 - DUST3R / MAST3R - Geometric 3D Vision Made Easy
 - VGGT - Visual Geometry Grounded Transformer



Questions?



Thank You!