

Do's and Dont's in Scientific Talks

Bogdan Savchynskyy

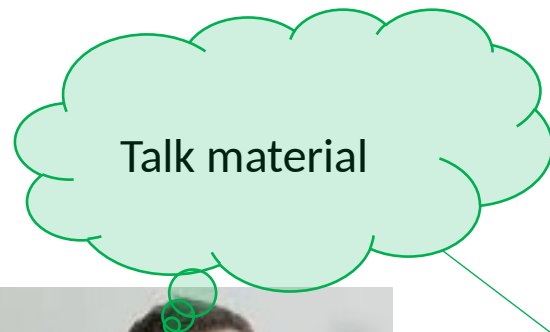
22/10/2024



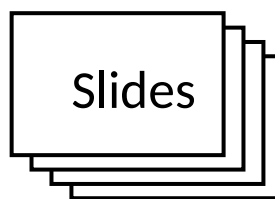
UNIVERSITÄT
HEIDELBERG
ZUKUNFT
SEIT 1386



Motivation



Talk material



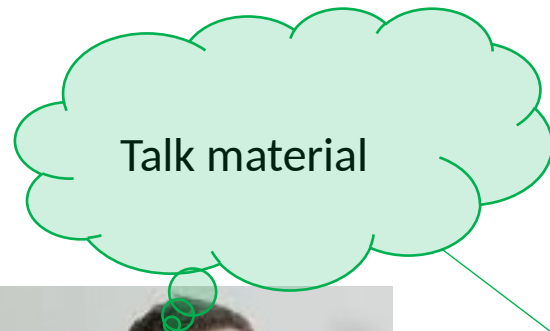
Slides

+ story

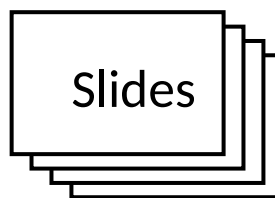
- Get credit points
- Learn how to give talks
- Learn about the topic



Motivation



Talk material



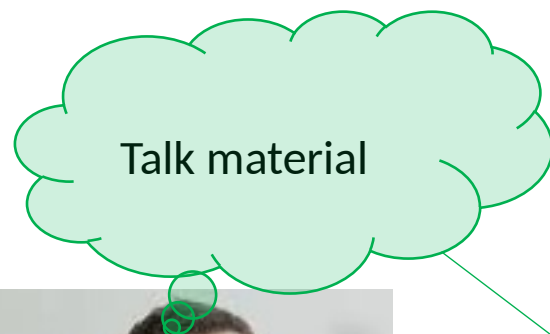
Slides

+ story

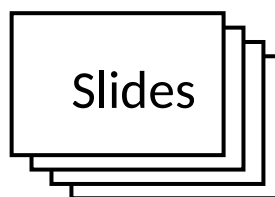
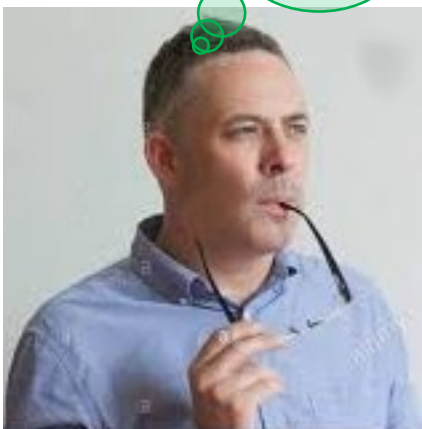
- Get credit points
- Learn how to give talks
- Learn about the topic
- Get higher grade for your PhD/Master/Bachelor
- Get feedback to you work



Motivation



Talk material



Slides

+ story

- Get credit points
- Learn how to give talks
- Learn about the topic
- Get higher grade for your PhD/Master/Bachelor
- Get feedback to you work

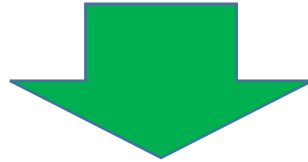
You are judged by your talk:

Each talk is an application talk



- **Slides content** (What is allowed on the slides and what isn't)
- **Slides order** (General structure of a talk)

- 
- **Slides content** (What is allowed on the slides and what isn't)
 - **Slides order** (General structure of a talk)
- 



- **Slides order** (General structure of a talk)
- **Slides content** (What is allowed on the slides and what isn't)

Slides Order

General structure of a talk

Typical talk outline?

1. Title
2. Outline
3. Problem description
4. Method/Solution
5. Experiments
6. Conclusions
7. Future work

DSAC – Differentiable RANSAC for Camera Localization

Eric Brachmann, Alexander Krull, Sebastian Nowozin, Jamie Shotton, Frank Michel, Stefan Gumhold, Carsten Rother

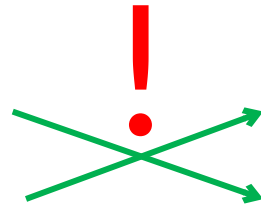
TU Dresden, Microsoft

1. RANSAC
2. Camera localization problem
3. Learning camera localization
4. Our end-to-end learning approach, DSAC
5. Experiments
6. Conclusions and Outlook

How much do you understand out of it?
Is this outline helpful for you?

Typical talk outline

1. Title
2. **Outline**
3. **Problem description**
4. Method/Solution
5. Experiments
6. Conclusions
7. Future work

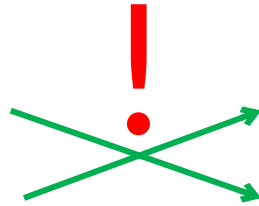


1. Title
2. **Problem description**
3. **Outline**
4. Method/Solution
5. Experiments
6. Conclusions
7. Future work

Problem description explains the outline.

Typical talk outline

1. Title
2. **Outline**
3. **Problem description**
4. Method/Solution
5. Experiments
6. Conclusions
7. Future work



1. Title
2. **Problem description**
3. ~~Outline~~
4. Method/Solution
5. Experiments
6. Conclusions
7. Future work

Problem description explains the outline.
You do not need an outline for most short (< 30 mins) talks.

Video: Provide your timing

1. Title
2. Problem description
3. Method/Solution
4. Experiments
5. Conclusions
6. Future work

Example of a talk



YouTube Video: [DSAC – Differentiable RANSAC for Camera Localization](#)

Typical talk timing

Time from the beginning of the talk:

- | | |
|------------------------|----------------------------|
| 1. Title | < 1 min: Thanks, coauthors |
| 2. Problem description | < 5 min |
| 3. Method/Solution | > 50 % of the talk |
| 4. Experiments | } 20-30 % of the talk |
| 5. Conclusions | |
| 6. Future work | |

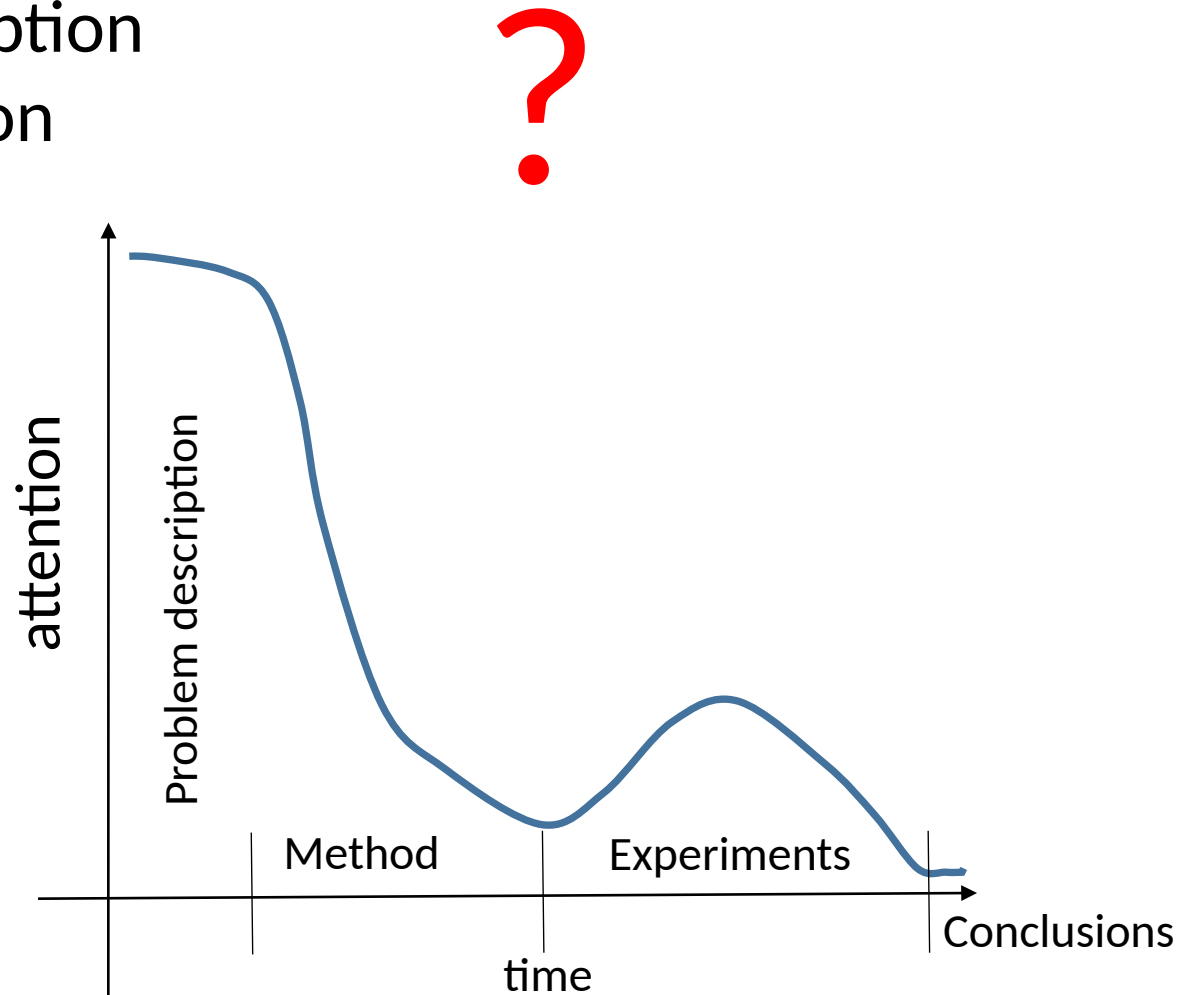
The most important part of the talk

1. Title
2. Problem description
3. Method/Solution
4. Experiments
5. Conclusions
6. Future work



The most important part of the talk

1. Title
2. Problem description
3. Method/Solution
4. Experiments
5. Conclusions
6. Future work



Video: The most important part of the talk

Problem description:

- simplified description (top view, clear for broad audience)
- more details (necessary to pose your actual problem)
- existing solutions
- deficiencies you address (problem statement)
- advertisement (short description of your results)

Example of a talk

Scene Coordinate Regression [Sho13]

Image I → Scene Coordinate Regression → RANSAC (Robust Scene (CR) Scene Plane Estimation) → 3D Pose $\hat{t}$

$p_i \in \mathbb{R}^2$ → $y_i \in \mathbb{R}^3$

[Sho13] "Scene Coordinate Regression Forests for Camera Relocalization in RGB-D Images", Shotton et al., CVPR'13

CVPR 2017 July 23-28 2017

Computer Vision and Learning Lab

19

1. RANSAC
2. Camera localization problem
- 3. Learning camera localization**
4. Our end-to-end learning approach, DSAC
5. Experiments
6. Conclusions and Outlook

When:

- long talk
- involved method

1. **RANSAC**
2. Camera localization problem
3. Learning camera localization
4. Our end-to-end learning approach, DSAC
5. Experiments
6. Conclusions and Outlook

...

1. RANSAC
- 2. Camera localization problem**
3. Learning camera localization
4. Our end-to-end learning approach, DSAC
5. Experiments
6. Conclusions and Outlook

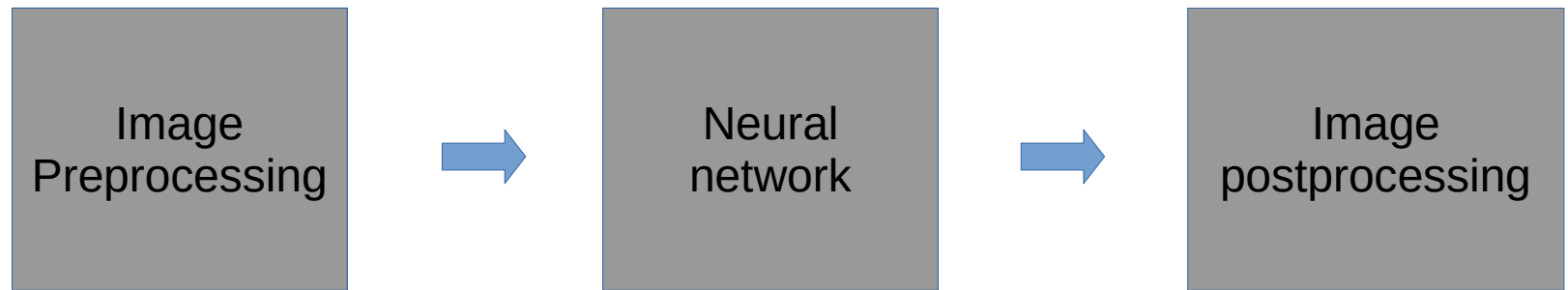
Camera Localization Problem

...

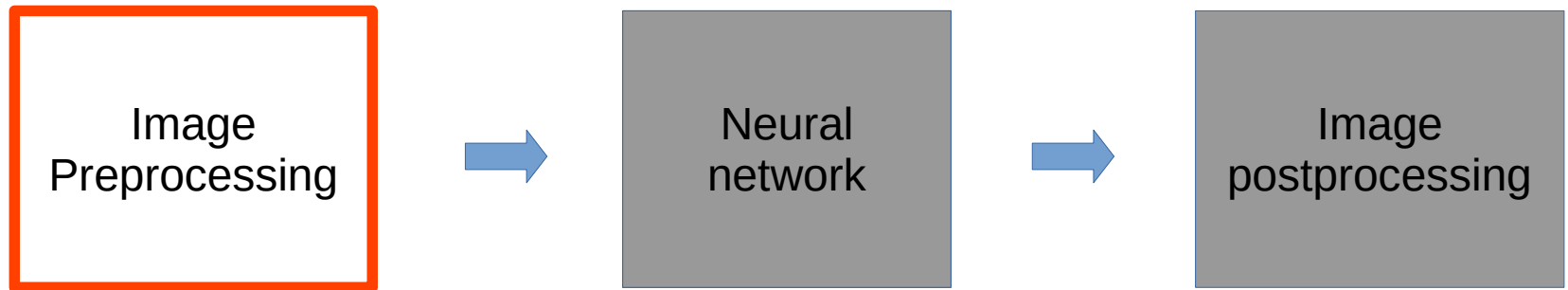
Hint: For complex talks consider conclusions to each part

1. RANSAC
2. Camera localization problem
- 3. Learning camera localization**
4. Our end-to-end learning approach, DSAC
5. Experiments
6. Conclusions and Outlook

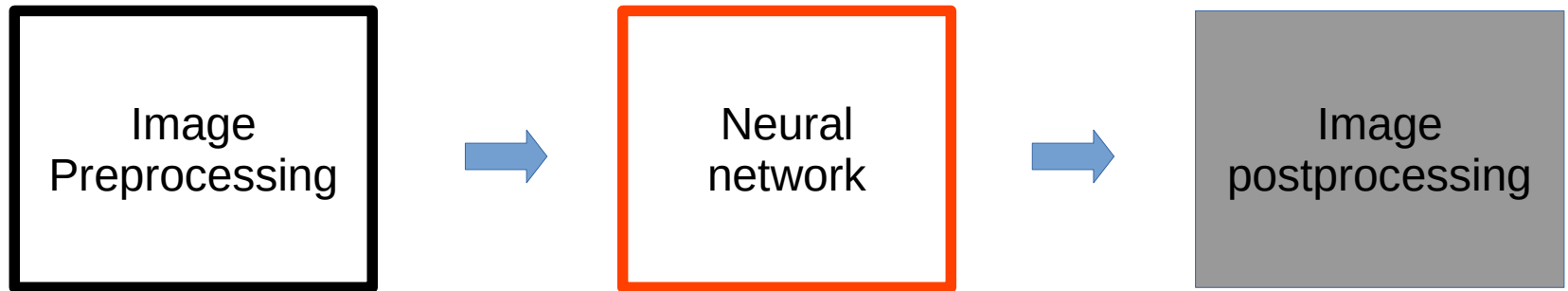
Pipeline-Based Outline



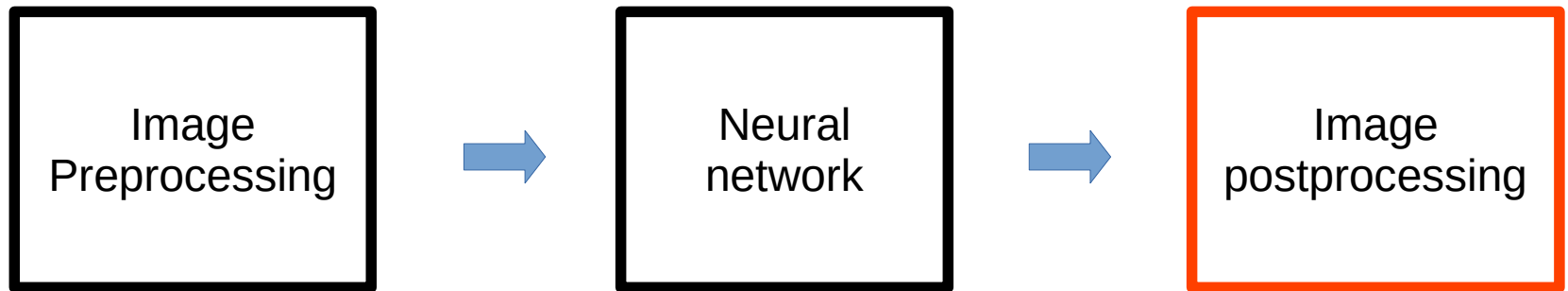
Pipeline-Based Outline



Pipeline-Based Outline



Pipeline-Based Outline



Principles of method/experiment description

1. Problem formulation



Find: Point correspondences, camera motion

Principles of method/experiment description

1. Problem formulation



Find: Point correspondences, camera motion

2. Problem formalization:

$$\min_{R,t} \sum_{i=1}^n \|I_l(i) - I_r[R,t](i)\|^2$$

Principles of method/experiment description

1. Problem formulation



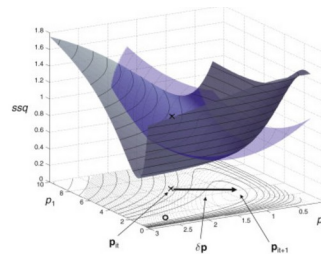
Find: Point correspondences, camera motion

2. Problem formalization:

$$\min_{R,t} \sum_{i=1}^n \|I_l(i) - I_r[R,t](i)\|^2$$

3. Problem solution:

Gauss-Newton method



Principles of method/experiment description

1. Problem formulation



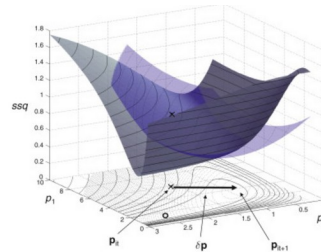
Find: Point correspondences, camera motion

2. Problem formalization:

$$\min_{R,t} \sum_{i=1}^n \|I_l(i) - I_r[R,t](i)\|^2$$

3. Problem solution:

Gauss-Newton method



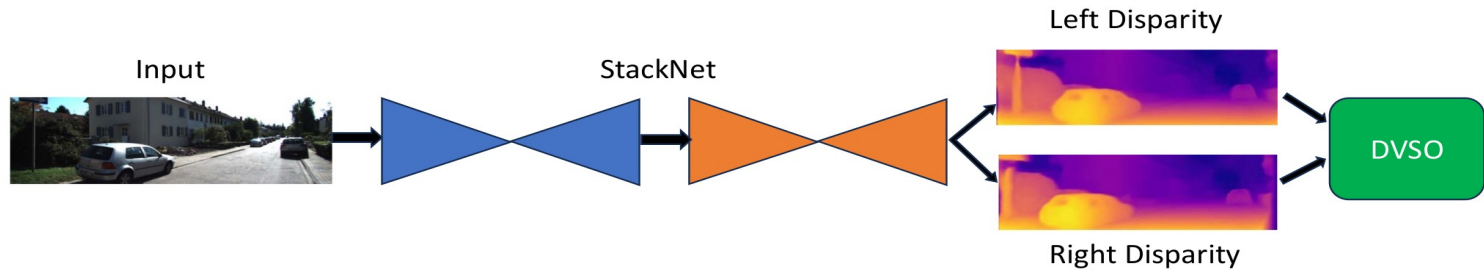
1. What to solve?



1. How to solve?

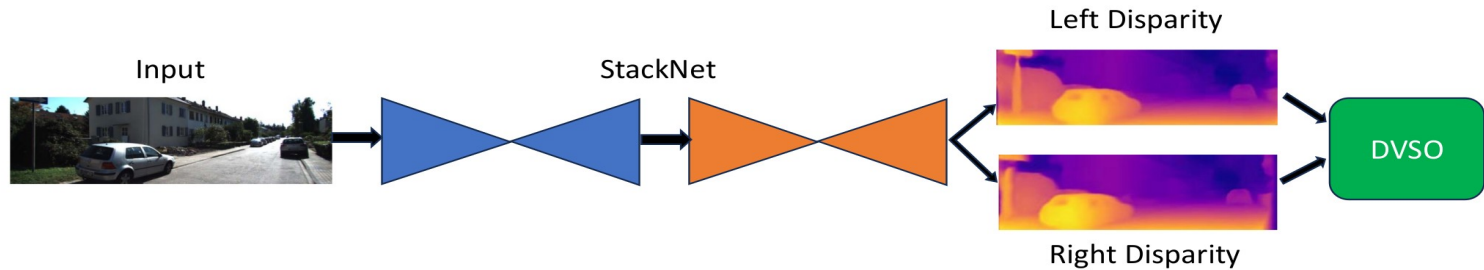
Principles of method/experiment description

1. From overview:

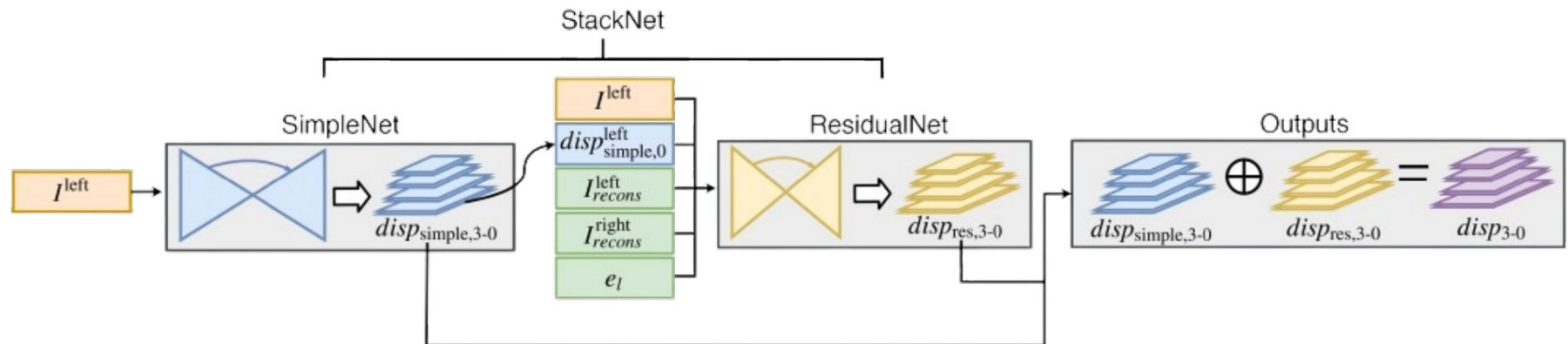


Principles of method/experiment description

1. From overview:

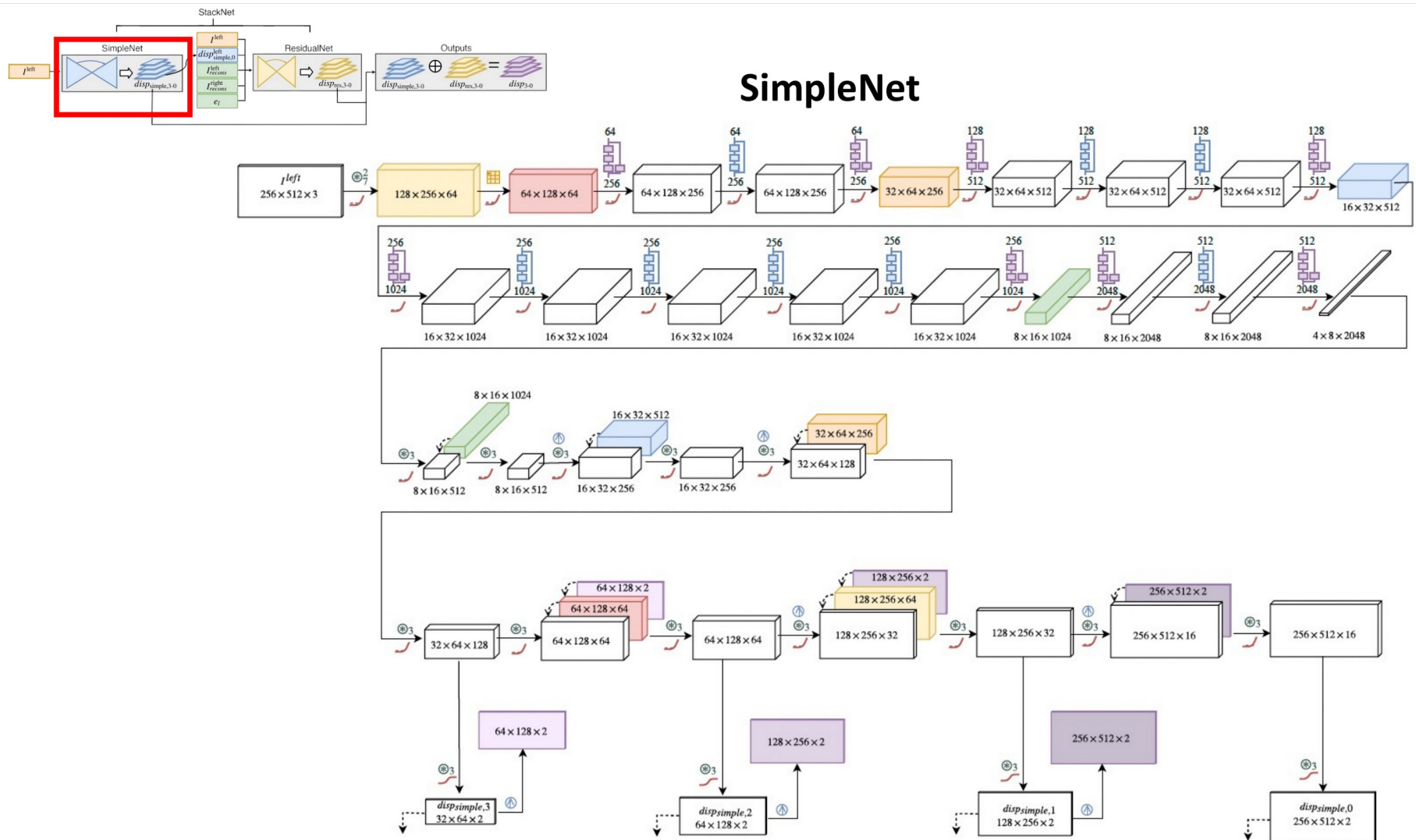


2. To details:



Principles of method/experiment description

But not too much details!



Do not start **Experimental section** with details:

Experiments



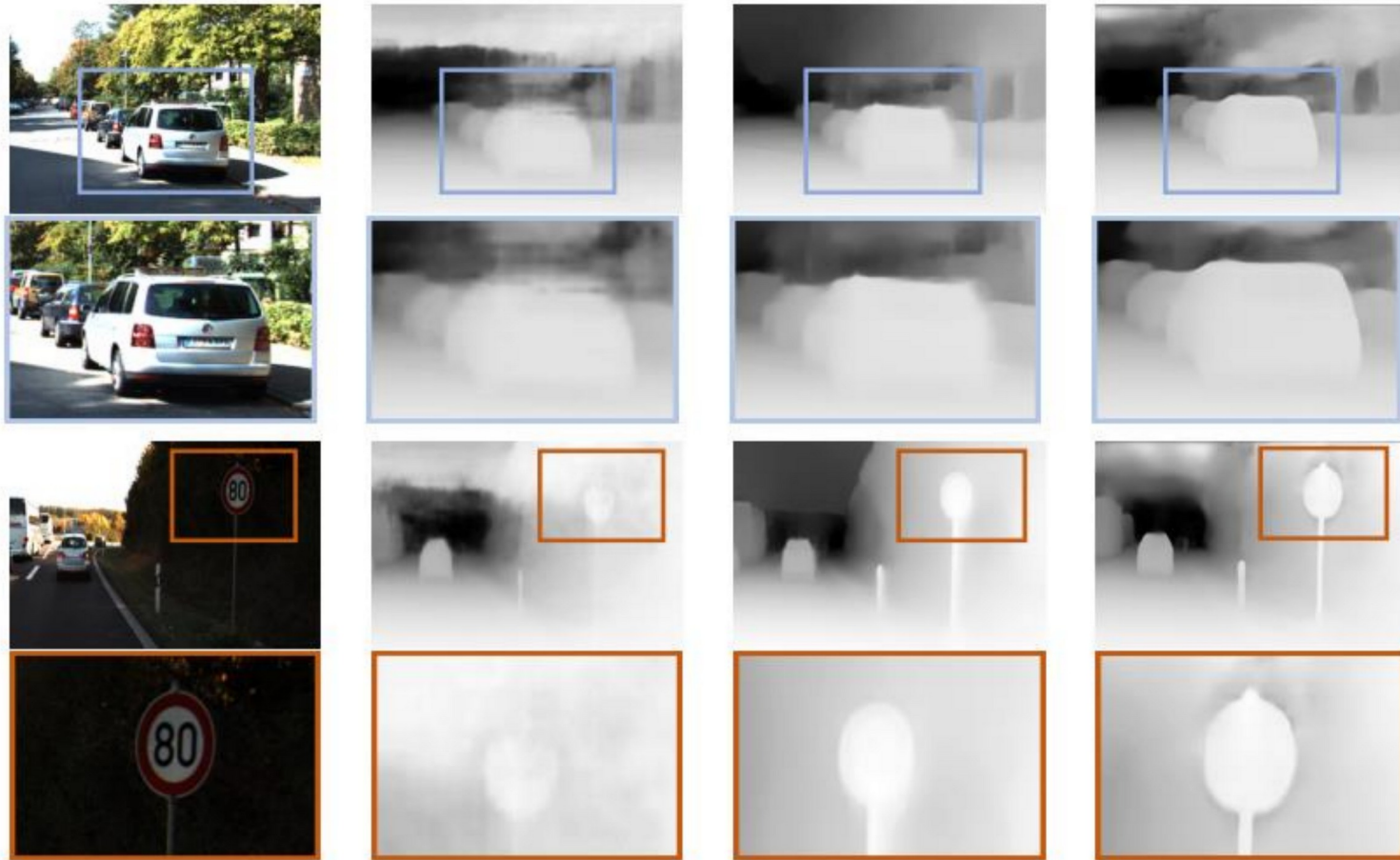
For training: Adam Optimizer

Backbone network: ResNet-101, ResNext-101 and DenseNet-161

Pretrained weights for ImageNet Large Scale Visual Recognition Challenge Dataset

Principles of method/experiment description

Better start with results:



Input

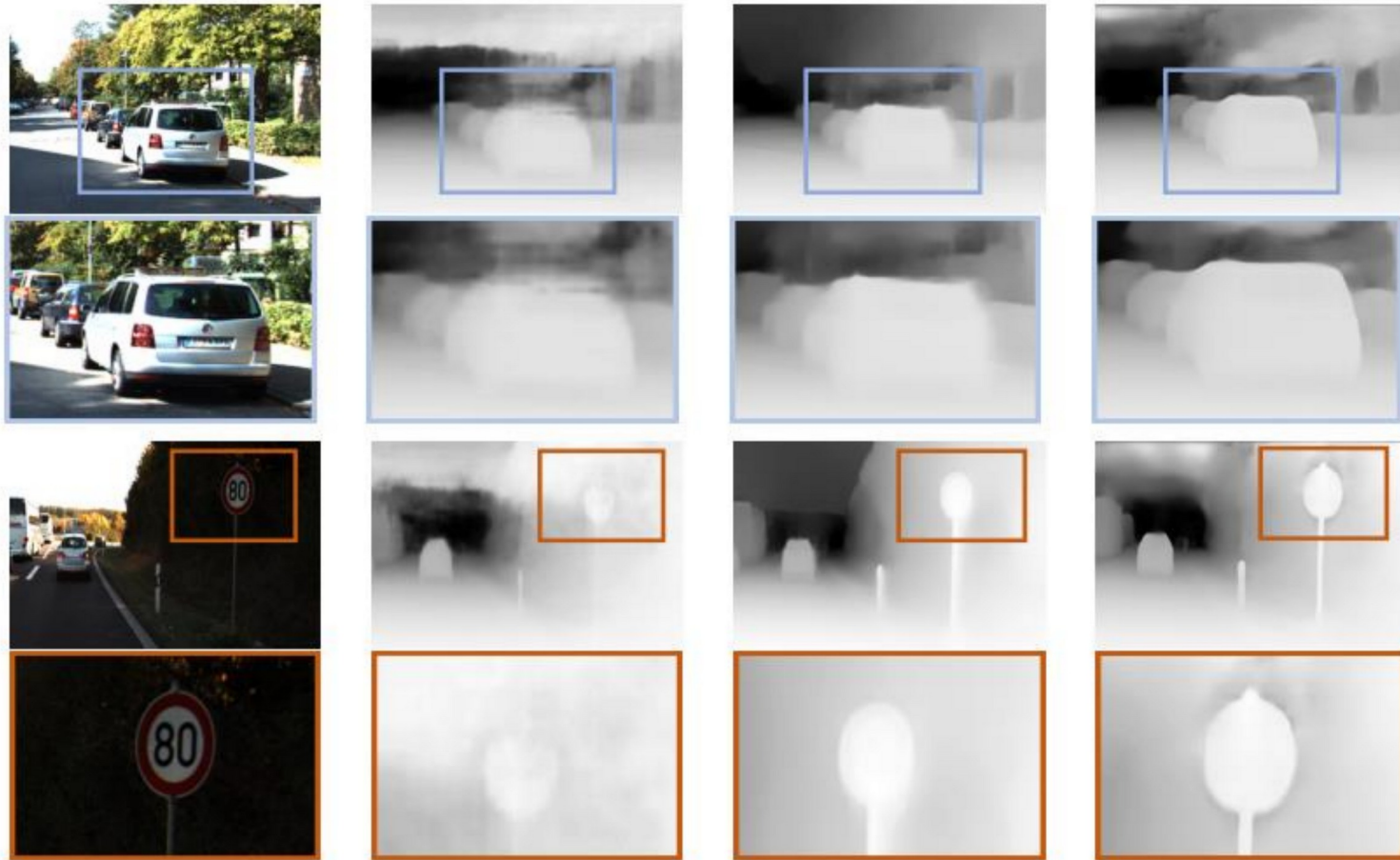
Fu et al.
ResNet-101

Yin et al.
ResNext-101

Big To Small
ResNet-101

Principles of method/experiment description

Better start with results:



Input

Fu et al.
ResNet-101

Yin et al.
ResNext-101

Big To Small
ResNet-101

... and then give details, if required

Slides Content

What is allowed on the slides and what not

Information Representation

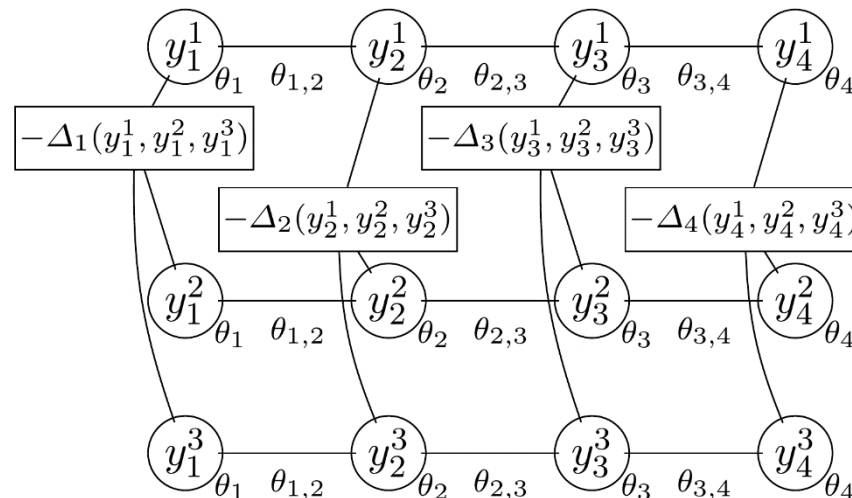
Text

A combinatorial optimization problem can be characterized by a mapping $f : R^m \rightarrow \{0,1\}^n$, where the problem description is mapped to a binary vector representing an optimal problem solution. In case of ILPs the description consists of the coefficients of the objective function and the constraint matrix. Computation of the mapping f has in general a complexity that grows at least as an exponent of m .

Formulas

$$\mathbf{K} = \text{diag}(\text{vec}(\mathbf{K}_p)) + (\mathbf{G}_2 \otimes \mathbf{G}_1) \text{diag}(\text{vec}(\mathbf{K}_q))(\mathbf{H}_2 \otimes \mathbf{H}_1)^\top$$

Images



Text

A combinatorial optimization problem can be characterized by a mapping $f : R^m \rightarrow \{0,1\}^n$, where the problem description is mapped to a binary vector representing an optimal problem solution. In case of ILPs the description consists of the coefficients of the objective function and the constraint matrix. Computation of the mapping f has in general a complexity that grows at least as an exponent of m .

Information Representation: Text vs. Formulas

Text

A combinatorial optimization problem can be characterized by a mapping $f : R^m \rightarrow \{0,1\}^n$, where the problem description is mapped to a binary vector representing an optimal problem solution. In case of ILPs the description consists of the coefficients of the objective function and the constraint matrix. Computation of the mapping f has in general a complexity that grows at least as an exponent of m .

Formulas

$f : R^m \rightarrow \{0,1\}^n$ - optimization problem, NP-hard

Information Representation: Text vs. Formulas

Text

A combinatorial optimization problem can be characterized by a mapping $f : R^m \rightarrow \{0,1\}^n$, where the problem description is mapped to a binary vector representing an optimal problem solution. In case of ILPs the description consists of the coefficients of the objective function and the constraint matrix. Computation of the mapping f has in general a complexity that grows at least as an exponent of m .

Formulas

$f : R^m \rightarrow \{0,1\}^n$ - optimization problem, NP-hard

Example:
$$f(A, c) = \arg \min_{\substack{Ax \leq b \\ x \in \{0,1\}^m}} \langle c, x \rangle$$

Information Representation: Text vs. Formulas

Text

~~A combinatorial optimization problem can be characterized by a mapping $f : R^m \rightarrow \{0,1\}^n$, where the problem description is mapped to a binary vector representing an optimal problem solution. In case of ILPs the description consists of the coefficients of the objective function and the constraint matrix. Computation of the mapping f has in general a complexity that grows at least as an exponent of m .~~

Formulas

$f : R^m \rightarrow \{0,1\}^n$ - optimization problem, NP-hard

Example: $f(A, c) = \arg \min_{\substack{Ax \leq b \\ x \in \{0,1\}^m}} \langle c, x \rangle$

Allowed: keywords, names, terminology

Avoid non-standard acronyms! SVM, VAE, ICA, CINN, LAP, PCA, CNN, MRF, CRF, MAP?

Information Representation: Text vs. Images

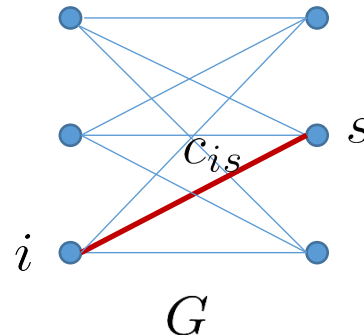
Text: c_{ij} is the similarity between node i and j in graph G

Information Representation: Text vs. Images

Text: c_{ij} is the similarity between node i and j in graph G

Image:

c_{ij} - similarity

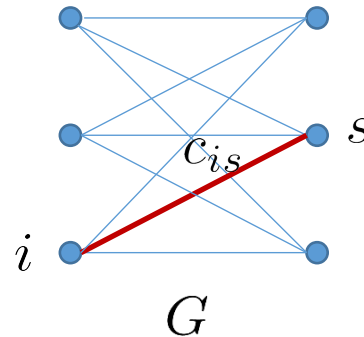


Information Representation: Text vs. Images

~~Text. c_{ij} is the similarity between node i and j in graph G~~

Image:

c_{ij} - similarity



Information Representation: Text in Tables

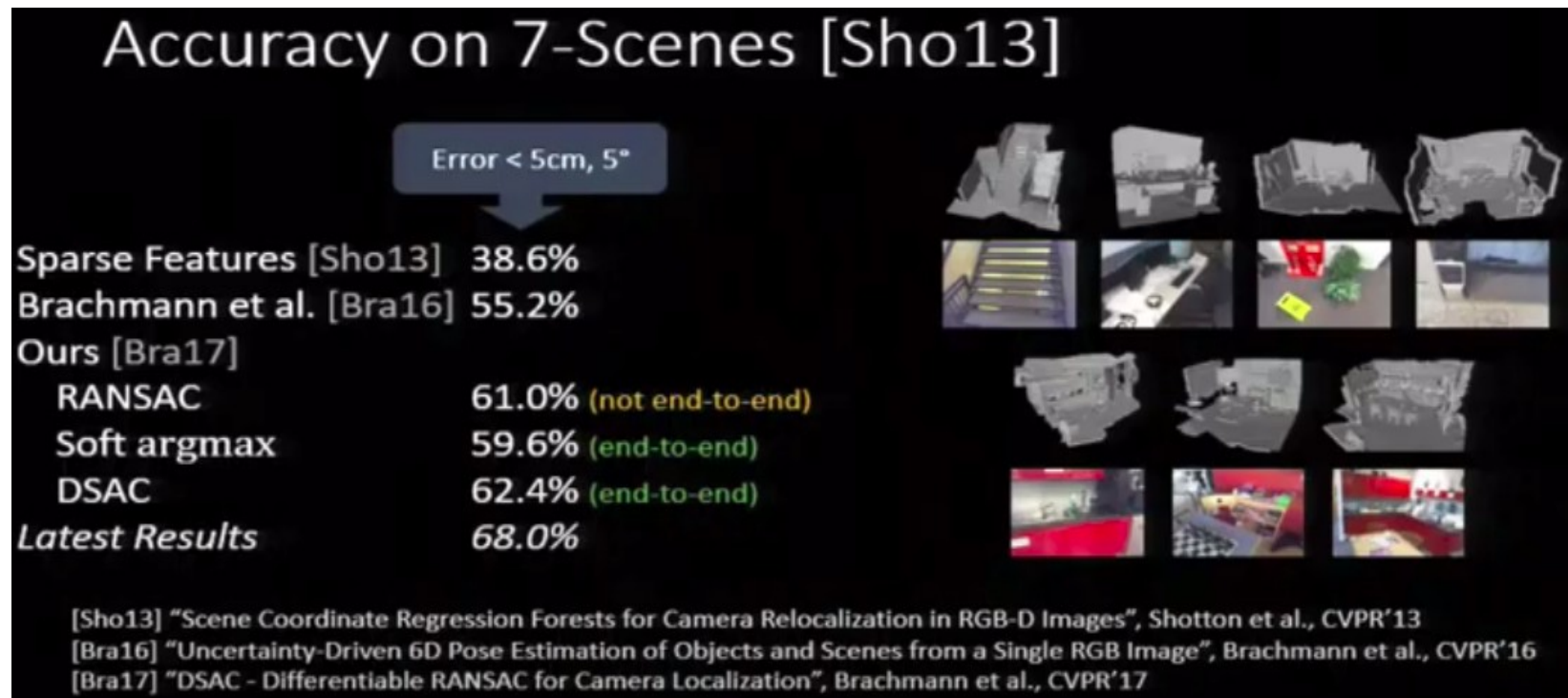
Table 1. Accuracy measured as the percentage of test images where the pose error is below 5cm and 5°. *Complete* denotes the combined set of frames (17000) of all scenes. Numbers in **green** denote improved accuracy after end-to-end training for SoftAM resp. DSAC compared to componentwise training. Similarly, **red** numbers denote decreased accuracy. **Bold** numbers indicate the best result for each scene.

	Sparse Features [36]	Brachmann <i>et al.</i> [5]	Ours: Trained Componentwise			Ours: Trained End-To-End	
			RANSAC	SoftAM	DSAC	SoftAM	DSAC
Chess	70.7%	94.9%	94.9%	94.8%	94.7%	94.2% -0.6%	94.6% -0.1%
Fire	49.9%	73.5%	75.1%	75.6%	75.3%	76.9% +1.3%	74.3% -1.0%
Heads	67.6%	48.1%	72.5%	74.5%	71.9%	74.0% -0.5%	71.7% -0.2%
Office	36.6%	53.2%	70.4%	71.3%	69.2%	56.6% -14.7%	71.2% +2.0%
Pumpkin	21.3%	54.5%	50.7%	50.6%	50.3%	51.9% +1.3%	53.6% +3.3%
Kitchen	29.8%	42.2%	47.1%	47.8%	46.2%	46.2% -1.6%	51.2% +5.0%
Stairs	9.2%	20.1%	6.2%	6.5%	5.3%	5.5% -1.0%	4.5% -0.8%
Average	40.7%	55.2%	59.5%	60.1%	59.0%	57.9% -2.2%	60.1% +1.1%
Complete	38.6%	55.2%	61.0%	61.6%	60.3%	57.8% -3.8%	62.5% +2.2%

Paper-style table.

Make it short enough to explain everything!

Information Representation: Text in Tables



Presentation-style table.

$$\mathbf{K} = \text{diag}(\text{vec}(\mathbf{K}_p)) + (\mathbf{G}_2 \otimes \mathbf{G}_1) \text{diag}(\text{vec}(\mathbf{K}_q))(\mathbf{H}_2 \otimes \mathbf{H}_1)^\top$$

Information Representation: Formulas

Factorization I

$$\mathbf{K} = \text{diag}(\text{vec}(\mathbf{K}_p)) + (\mathbf{G}_2 \otimes \mathbf{G}_1) \text{diag}(\text{vec}(\mathbf{K}_q)) (\mathbf{H}_2 \otimes \mathbf{H}_1)^T.$$

$$\mathbf{K}_p = \begin{pmatrix} 1 & 1 & 4 \\ 10 & 3 & 2 \\ 2 & 9 & 3 \\ 1 & 2 & 8 \end{pmatrix}$$

$$\dots = \begin{pmatrix} 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 10 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 2 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ \hline 0 & 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 3 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 9 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 2 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ \hline 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 4 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 2 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 3 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 8 & 0 & 0 & 0 \end{pmatrix}$$

Factorization II

$$\mathbf{K} = \text{diag}(\text{vec}(\mathbf{K}_p)) + (\mathbf{G}_2 \otimes \mathbf{G}_1) \text{diag}(\text{vec}(\mathbf{K}_q)) (\mathbf{H}_2 \otimes \mathbf{H}_1)^T.$$

$$\mathbf{G}_2 \otimes \mathbf{G}_1 = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 1 & 0 \\ 0 & 0 & 0 & 1 \end{pmatrix} \otimes \mathbf{G}_1 = \begin{pmatrix} \mathbf{G}_1 & \mathbf{0} & \mathbf{0} & \mathbf{0} \\ \mathbf{0} & \mathbf{G}_1 & \mathbf{G}_1 & \mathbf{0} \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \mathbf{G}_1 \end{pmatrix}$$

Information Representation: Formulas

Factorization III

$$\mathbf{K} = \text{diag}(\text{vec}(\mathbf{K}_p)) + (\mathbf{G}_2 \otimes \mathbf{G}_1) \text{diag}(\text{vec}(\mathbf{K}_q)) (\mathbf{H}_2 \otimes \mathbf{H}_1)^T.$$

$$\begin{bmatrix} \mathbf{G}_1 & 0 & 0 & 0 \\ 0 & \mathbf{G}_1 & \mathbf{G}_1 & 0 \\ 0 & 0 & 0 & \mathbf{G}_1 \end{bmatrix} \cdot \begin{bmatrix} d(\mathbf{k}_{(a,b)}^q) & 0 & 0 & 0 \\ 0 & d(\mathbf{k}_{(b,c)}^q) & 0 & 0 \\ 0 & 0 & d(\mathbf{k}_{(b,a)}^q) & 0 \\ 0 & 0 & 0 & d(\mathbf{k}_{(c,b)}^q) \end{bmatrix} \cdot \begin{bmatrix} 0 & \mathbf{H}_1 & 0 \\ 0 & 0 & \mathbf{H}_1 \\ \mathbf{H}_1 & 0 & 0 \\ 0 & \mathbf{H}_1 & 0 \end{bmatrix}$$

$$= \begin{array}{c|ccc} & a & b & c \\ \hline a & 0 & \mathbf{G}_1 \text{diag}(\mathbf{k}_{(a,b)}^q) \mathbf{H}_1^T & 0 \\ b & \mathbf{G}_1 \text{diag}(\mathbf{k}_{(b,a)}^q) \mathbf{H}_1^T & 0 & \mathbf{G}_1 \text{diag}(\mathbf{k}_{(b,c)}^q) \mathbf{H}_1^T \\ c & 0 & \mathbf{G}_1 \text{diag}(\mathbf{k}_{(c,b)}^q) \mathbf{H}_1^T & 0 \end{array}$$

Avoid complex notations: $\text{vec}(X)$, $\text{diag}(x)$

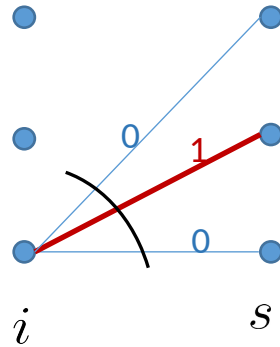
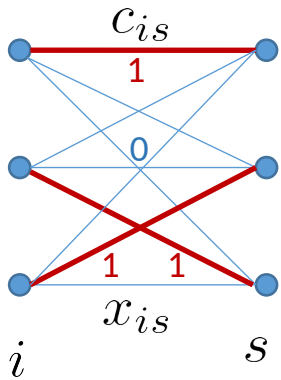
$\text{vec}(X) \Rightarrow X$ - matrix, x - resp. vector

$$\text{diag}(x) \Rightarrow \mathbf{D} = \begin{bmatrix} x_1 & 0 & 0 \\ 0 & x_2 & 0 \\ 0 & 0 & x_3 \end{bmatrix}$$

Formulas vs. images

$$\begin{aligned} \min_{x \in \{0,1\}^{n \times n}} \quad & \sum_{i=1}^n \sum_{s=1}^n c_{is} x_{is} \\ \text{s.t.:} \quad & \sum_{s=1}^n x_{is} = 1 \quad \forall i \\ & \sum_{i=1}^n x_{is} = 1 \quad \forall s \end{aligned}$$

Formulas vs. images



$$i : \sum_{s=1}^n x_{is} = 1$$

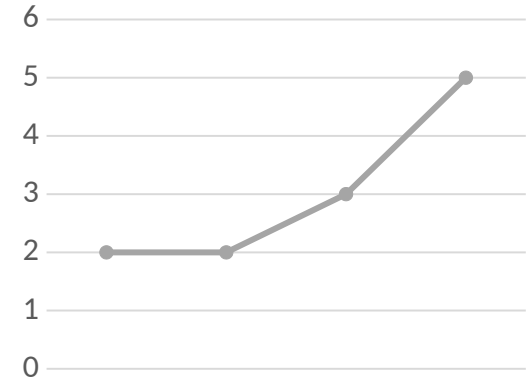
$$\begin{aligned} \min_{x \in \{0,1\}^{n \times n}} \quad & \sum_{i=1}^n \sum_{s=1}^n c_{is} x_{is} \\ \text{s.t.:} \quad & \sum_{s=1}^n x_{is} = 1 \quad \forall i \\ & \sum_{i=1}^n x_{is} = 1 \quad \forall s \end{aligned}$$

c_{is} - cost of matching $i \leftrightarrow s$

$x_{is} \in \{0, 1\}$ - edge selectors

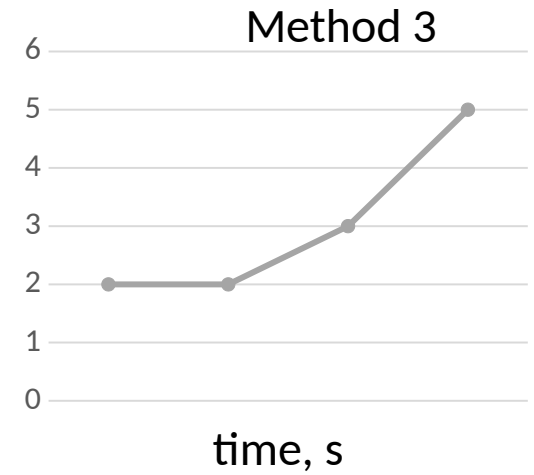
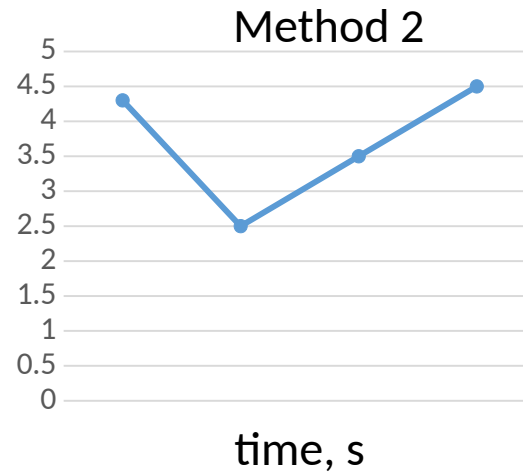
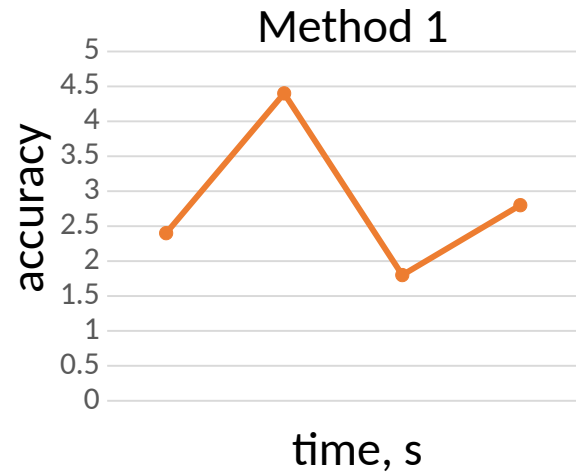
$\sum_{is} c_{is} x_{is}$ - total matching cost

Text Usage



(left, orange) – performance of Method 1
(center, blue) – performance of Method 2
(right, gray) – performance of Method 3

Text Usage



Talk vs. paper or lecture notes

Table 1: **Characteristics of datasets.** For all datasets used for evaluation we state number of instances (*inst.*), number of nodes (n), number of labels ($|\mathcal{L}|$), and graph density in percent (*dens.*).

	<i>inst.</i>	n	$ \mathcal{L} $	<i>dens.</i> (%)
hotel	105	30	$= n$	100
house	105	30	$= n$	100
car [†]	30	19 - 49	$= n$	11 - 27
motor [†]	20	15 - 52	$= n$	10 - 32
flow	6	48 - 126	$\approx n$	45 - 98
opengm	4	19 / 20	$= n$	66 / 100
worms	30	558	$\approx 2.4 n$	≈ 1.5
pairs	16	511 - 565	$\approx n$	≈ 20

[†] Zero edges were removed. Prior to this, graph density was 100 %.

What is wrong with this slide?

Algorithm 2: Iterative Pruning Arc Consistency

Input: Cost vector $f \in \mathbb{R}^{\mathcal{I}}$, test labeling $y \in \mathcal{X}$;

Output: Strictly improving substitution p ;

```
1  $(\forall v \in \mathcal{V}) \mathcal{Y}_v := \mathcal{X}_v \setminus \{y_v\}$ ;  
2 while true  
3   Construct verification problem  $g := (I - [p])^\top f$  with  $p$   
   defined by (7);  
4   Use dual solver to find  $\varphi$  such that  $g^\varphi$  is arc consistent;  
5    $\mathcal{O}_v(\varphi) := \{i \in \mathcal{X}_v \mid g_v^\varphi(i) = 0\}$ ;  
6   if  $(\forall v \in \mathcal{V}) \mathcal{O}_v(\varphi) \cap \mathcal{Y}_v = \emptyset$  then return  $p$ ;  
7   for  $v \in \mathcal{V}$  do  
8     Pruning of substitutions:  $\mathcal{Y}_v := \mathcal{Y}_v \setminus \mathcal{O}_v(\varphi)$ ;
```

What is wrong with this slide?

Talk vs. paper or lecture notes

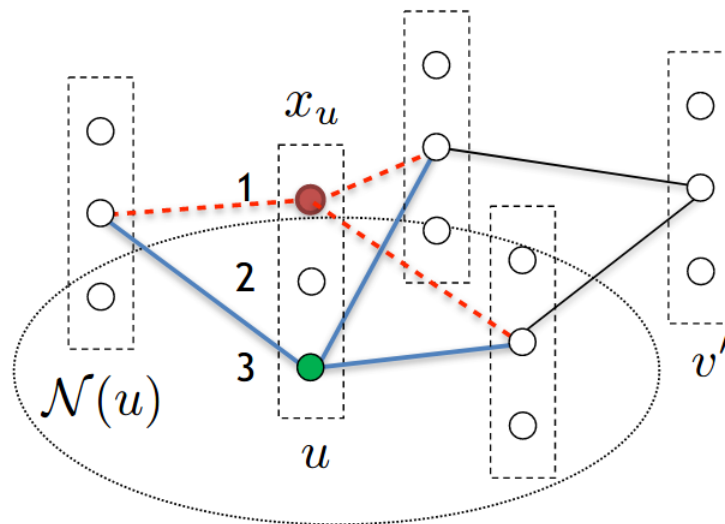


Fig. 2: Dead end elimination / dominance. Variables are shown as boxes and their possible labels as circles. Label $x_u = 1$ is substituted with label $x_u = 3$. If for any configuration of neighbors $x_{\mathcal{N}(u)}$ the energy does not increase (only the terms inside $\{u\} \cup \mathcal{N}(u)$ contribute to the difference), label $x_u = 1$ can be eliminated without loss of optimality.

What is wrong with this slide?

The Best Method to Solve the Universal Problem

Max Mustermann

Seminar on

„Optimization in Machine Learning and Computer Vision“

Heidelberg 2023

What is missing here?

The Best Method to Solve the Universal Problem

Y. Zhang, J. Schmidt, A. Zimmersinger
University of Toronto, DeepResearch Ltd.
2022

Presenter: Max Mustermann

Seminar on

„Optimization in Machine Learning and Computer Vision“

Heidelberg 2023

Authors, affiliations and the publication year!

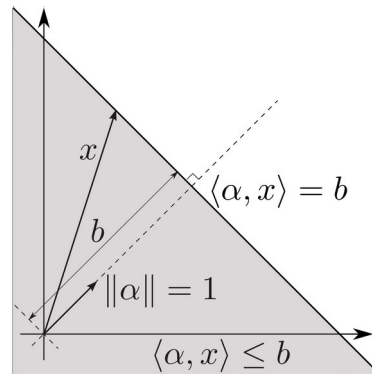
Other recommendations and requirements

- Add **page numbers**: Easier to reference
- Using PowerPoint use **IguanaTex**: Nicely looking formulas
- Use PowerPoint tools or **InkScape**: To draw your own illustrations
- Avoid messy templates: They eat up too much slide space

Slide

15

$$\max_{\substack{i \in K \\ i < y}} \sum_{t=1}^n f^t(x_i)$$



I eat up your slide space

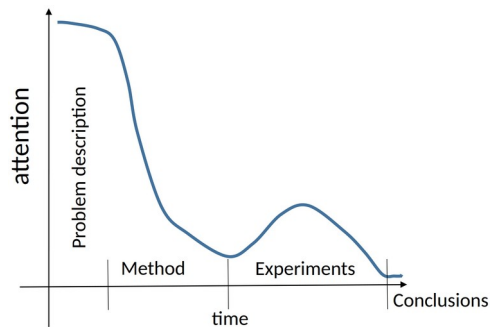
Yes, I do



Thank you for your attention!

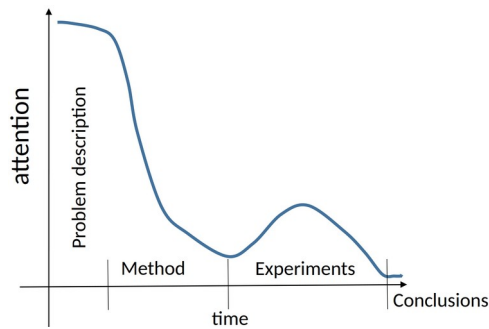
Last Slide = Take-Home Message

Last Slide = Take-Home Message



Problem description < 5 min, the most important

Last Slide = Take-Home Message

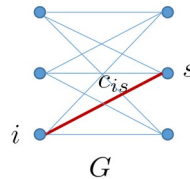


Problem description < 5 min, the most important

~~A combinatorial optimization problem can be characterized by a mapping $f: R^m \rightarrow \{0,1\}^n$, where the problem description is mapped to a binary vector representing an optimal problem solution. In case of ILPs the description consists of the coefficients of the objective function and the constraint matrix. Computation of the mapping f has in general a complexity that grows at least as an exponent of m .~~

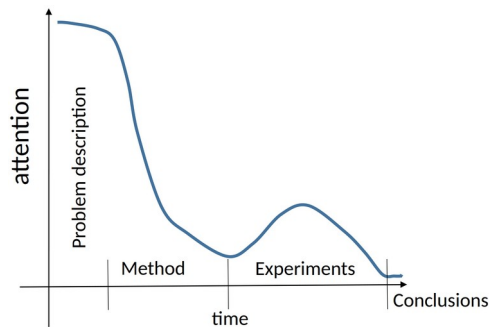


$$f: R^m \rightarrow \{0,1\}^n$$



Illustrations and formulas instead of text

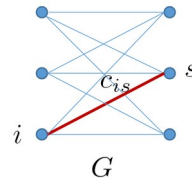
Last Slide = Take-Home Message



Problem description < 5 min, the most important

~~A combinatorial optimization problem can be characterized by a mapping $f: R^m \rightarrow \{0,1\}^n$, where the problem description is mapped to a binary vector representing an optimal problem solution. In case of ILPs the description consists of the coefficients of the objective function and the constraint matrix. Computation of the mapping f has in general a complexity that grows at least as an exponent of m .~~

$$f: R^m \rightarrow \{0,1\}^n$$



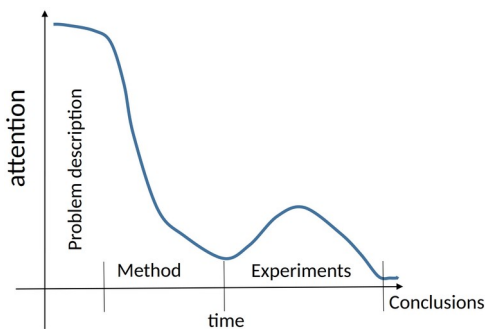
Illustrations and formulas instead of text

~~$$\mathbf{K} = \text{diag}(\text{vec}(\mathbf{K}_p)) + (\mathbf{G}_2 \otimes \mathbf{G}_1) \text{diag}(\text{vec}(\mathbf{K}_a))(\mathbf{H}_2 \otimes \mathbf{H}_1)^\top$$~~

$$\sum_{ij} c_{ij} x_{ij} + \sum_{ij,kl} c_{ij,kl} x_{ij,kl}$$

Avoid and simplify complex formulas

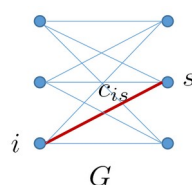
Last Slide = Take-Home Message



Problem description < 5 min, the most important

~~A combinatorial optimization problem can be characterized by a mapping $f: R^m \rightarrow \{0,1\}^n$, where the problem description is mapped to a binary vector representing an optimal problem solution. In case of ILPs the description consists of the coefficients of the objective function and the constraint matrix. Computation of the mapping f has in general a complexity that grows at least as an exponent of m .~~

$$f: R^m \rightarrow \{0,1\}^n$$



Illustrations and formulas instead of text

$$\mathbf{K} = \text{diag}(\text{vec}(\mathbf{K}_p)) + (\mathbf{G}_2 \otimes \mathbf{G}_1) \text{diag}(\text{vec}(\mathbf{K}_a))(\mathbf{H}_2 \otimes \mathbf{H}_1)^T$$
$$\sum_{ij} c_{ij} x_{ij} + \sum_{ij,kl} c_{ij,kl} x_{ij,kl}$$

Avoid and simplify complex formulas

~~Table 1. Accuracy measured as the percentage of test images where the pose error is below 5cm and 5°. Complete denotes the combined set of frames (17000) of all scenes. Numbers in green denote improved accuracy after end-to-end training for SoftAM resp. DSAC compared to componentwise training. Similarly, red numbers denote decreased accuracy. Bold numbers indicate the best result for each scene.~~

	Sparse Features [36]	Brachmann et al. [5]	Ours: Trained Componentwise			Ours: Trained End-To-End		
			RANSAC	SoftAM	DSAC	SoftAM	DSAC	
Chess	70.7%	94.9%	94.9%	94.8%	94.7%	94.2% -0.6%	94.6% -0.1%	
Fire	49.9%	73.5%	75.1%	75.6%	75.3%	76.9% +1.3%	74.3% -1.0%	
Heads	67.6%	48.1%	72.5%	74.5%	71.9%	74.0% -0.5%	71.7% -0.2%	
Office	36.6%	53.2%	70.4%	71.3%	69.2%	56.6% -14.7%	71.2% +2.0%	
Pumpkin	21.3%	54.5%	50.7%	50.6%	50.8%	51.9% +1.3%	53.6% +3.3%	
Kitchen	29.8%	42.2%	47.1%	47.8%	46.2%	46.2% -1.6%	51.2% +5.0%	
Stairs	9.2%	20.1%	6.2%	6.5%	5.3%	5.5% -1.0%	4.5% -0.8%	
Average	40.7%	55.2%	59.5%	60.1%	59.0%	57.9% -2.2%	60.1% +1.1%	
Complete	38.6%	55.2%	61.0%	61.6%	60.3%	57.8% -3.8%	62.5% +4.7%	



Error < 5cm, 5°	
Sparse Features [Sho13]	38.6%
Brachmann et al. [Bra16]	55.2%
Ours [Bra17]	
RANSAC	61.0% (not end-to-end)
Soft argmax	59.6% (end-to-end)
DSAC	62.4% (end-to-end)
Latest Results	68.0%

Avoid paper-style elements in presentation

Do's of an entertaining talk

- **Joke** (if you feel comfortable with that)
- **Advertise** (e.g. after problem formulation)
- **Simplify** (give examples)
- **Intrigue** (e.g. seemingly correct conclusions)

Checklist/Feedback criteria

1. Talk structure:

- Is the general structure of the presentation correct, as in Sl. 14?
- Is the **Problem description** composed as in Sl. 19? Was the timing ok? (<5 min)
- Would you change anything in the structure of the **Method/Solution** part?
- Useless or missing/insufficient **Outline**?

2. Slide content:

- Too much text?
- Too few explanatory images?
- Difficult formulas with insufficient explanation (did you understand them)?
- Long image/table captions?
- Enumeration of formulas, descriptions starting with **Table, Figure, Algorithm** etc?
- Are there paper-style slides, tables, algorithms, theorems etc. ?
- Title page ok?

3. Presentation and question answering:

- Missing/incorrect/uncertain answers to the questions?
- Intonation/voice modulation: Emphasis on important things?
- Acoustical problems: Too fast/not loud enough/unclear pronunciation etc.?

A good practice applied math presentation

YouTube Video: [An explicit analysis of the entropic penalty in linear programming](#)

Thanks you for attention!

**This is just a kind reminder to submit
your slides 2 weeks in advance**

Thanks you for attention!

Send your final slides before your presentation