

4. The Scaling Problem

INDEX

1. Limitations of the basic EKF SLAM algorithm

- 1. The linearization problem**
- 2. The scaling problem**

2. Alternatives

- 1. Independent local maps**
- 2. Map joining**
- 3. Hierarchical SLAM**
- 4. Divide and Conquer SLAM**

3. Conclusion

1. Limitations of the basic EKF SLAM algorithm

The scaling problem

The EKF SLAM algorithm

Algorithm 1 SLAM:

$\mathbf{x}_0^B = \mathbf{0}; \mathbf{P}_0^B = \mathbf{0} \{Map\ initialization\}$

$[\mathbf{z}_0, \mathbf{R}_0] = \text{get_measurements}$

$[\mathbf{x}_0^B, \mathbf{P}_0^B] = \text{add_new_features}(\mathbf{x}_0^B, \mathbf{P}_0^B, \mathbf{z}_0, \mathbf{R}_0)$

for $k = 1$ to steps do

$[\mathbf{x}_{R_k}^{R_{k-1}}, \mathbf{Q}_k] = \text{get_odometry}$

→ $[\mathbf{x}_{k|k-1}^B, \mathbf{P}_{k|k-1}^B] = \text{EKF_prediction}(\mathbf{x}_{k-1}^B, \mathbf{P}_{k-1}^B, \mathbf{x}_{R_k}^{R_{k-1}}, \mathbf{Q}_k)$

$[\mathbf{z}_k, \mathbf{R}_k] = \text{get_measurements}$

→ $\mathcal{H}_k = \text{data_association}(\mathbf{x}_{k|k-1}^B, \mathbf{P}_{k|k-1}^B, \mathbf{z}_k, \mathbf{R}_k)$

→ $[\mathbf{x}_k^B, \mathbf{P}_k^B] = \text{EKF_update}(\mathbf{x}_{k|k-1}^B, \mathbf{P}_{k|k-1}^B, \mathbf{z}_k, \mathbf{R}_k, \mathcal{H}_k)$

→ $[\mathbf{x}_k^B, \mathbf{P}_k^B] = \text{add_new_features}(\mathbf{x}_k^B, \mathbf{P}_k^B, \mathbf{z}_k, \mathbf{R}_k, \mathcal{H}_k)$

end for

The prediction step

EKF SLAM prediction

$$\hat{\mathbf{x}}_{k|k-1}^B = \begin{bmatrix} \hat{\mathbf{x}}_{R_{k-1}}^B \oplus \hat{\mathbf{x}}_{R_k}^{R_{k-1}} \\ \hat{\mathbf{x}}_{F_{1,k-1}}^B \\ \vdots \\ \hat{\mathbf{x}}_{F_{m,k-1}}^B \end{bmatrix}$$

$$\mathbf{P}_{k|k-1}^B \simeq \mathbf{F}_k \mathbf{P}_{k-1}^B \mathbf{F}_k^T + \mathbf{G}_k \mathbf{Q}_k \mathbf{G}_k^T$$

$$\mathbf{F}_k = \begin{bmatrix} \mathbf{J}_{1 \oplus} \left\{ \hat{\mathbf{x}}_{R_{k-1}}^B, \hat{\mathbf{x}}_{R_k}^{R_{k-1}} \right\} & \mathbf{0} & \cdots & \mathbf{0} \\ \mathbf{0} & \mathbf{I} & & \vdots \\ \vdots & & \ddots & \\ \mathbf{0} & \cdots & & \mathbf{I} \end{bmatrix} ; \quad \mathbf{G}_k = \begin{bmatrix} \mathbf{J}_{2 \oplus} \left\{ \hat{\mathbf{x}}_{R_{k-1}}^B, \hat{\mathbf{x}}_{R_k}^{R_{k-1}} \right\} \\ \mathbf{0} \\ \vdots \\ \mathbf{0} \end{bmatrix}$$

EKF SLAM prediction

$$\mathbf{P}_{k-1|k-1}^B = \begin{pmatrix} \mathbf{P}_R & \mathbf{P}_{RF_1} & \dots & \mathbf{P}_{RF_n} \\ \mathbf{P}_{RF_1}^T & \mathbf{P}_{F_1} & \dots & \mathbf{P}_{F_1 F_n} \\ \vdots & \vdots & \ddots & \vdots \\ \mathbf{P}_{RF_n}^T & \mathbf{P}_{F_1 F_n}^T & \dots & \mathbf{P}_{F_n} \end{pmatrix}$$

$$\mathbf{P}_{k|k-1}^B = \begin{pmatrix} \boxed{\mathbf{J}_{1\oplus} \mathbf{P}_R \mathbf{J}_{1\oplus}^T + \mathbf{J}_{2\oplus} \mathbf{Q}_k \mathbf{J}_{2\oplus}^T} & \boxed{\mathbf{J}_{1\oplus} \mathbf{P}_{RF_1}} & \dots & \boxed{\mathbf{J}_{1\oplus} \mathbf{P}_{RF_n}} \\ \mathbf{J}_{1\oplus}^T \mathbf{P}_{RF_1}^T & \mathbf{P}_{F_1} & \dots & \mathbf{P}_{F_1 F_n} \\ \vdots & \vdots & \ddots & \vdots \\ \mathbf{J}_{1\oplus}^T \mathbf{P}_{RF_n}^T & \mathbf{P}_{F_1 F_n}^T & \dots & \mathbf{P}_{F_n} \end{pmatrix}$$

EKF prediction is $O(n)$

Adding new features

EKF SLAM: add new features

$$\mathbf{x}_k^B = \begin{pmatrix} \mathbf{x}_{R_k}^B \\ \mathbf{x}_{F_{1,k}}^B \\ \vdots \\ \mathbf{x}_{F_{n,k}}^B \end{pmatrix} \Rightarrow \mathbf{x}_{k+}^B = \begin{pmatrix} \mathbf{x}_{R_k}^B \\ \mathbf{x}_{F_{1,k}}^B \\ \vdots \\ \mathbf{x}_{F_{n,k}}^B \\ \mathbf{x}_{F_{n+1,k}}^B \end{pmatrix} = \begin{pmatrix} \mathbf{x}_{R_k}^B \\ \mathbf{x}_{F_{1,k}}^B \\ \vdots \\ \mathbf{x}_{F_{n,k}}^B \\ \mathbf{x}_{R_k}^B \oplus \mathbf{z}_i \end{pmatrix}$$

Linearization:

$$\mathbf{x}_{k+}^B \simeq \hat{\mathbf{x}}_{k+}^B + \mathbf{F}_k(\mathbf{x}_k^B - \hat{\mathbf{x}}_k^B) + \mathbf{G}_k(\mathbf{z}_i - \hat{\mathbf{z}}_i)$$

$$\mathbf{P}_{k+}^B = \mathbf{F}_k \mathbf{P}_k^B \mathbf{F}_k^T + \mathbf{G}_k \mathbf{R}_k \mathbf{G}_k^T$$

Where:

$$\mathbf{F}_k = \frac{\partial \mathbf{x}_{k+}^B}{\partial \mathbf{x}_k^B} = \begin{pmatrix} \mathbf{I} & 0 & \dots & 0 \\ 0 & \ddots & \dots & 0 \\ 0 & 0 & \dots & \mathbf{I} \\ \mathbf{J}_{1 \oplus \{\hat{\mathbf{x}}_{R_k}^B, \hat{\mathbf{z}}_i\}} & 0 & \dots & 0 \end{pmatrix}; \mathbf{G}_k = \frac{\partial \mathbf{x}_{k+}^B}{\partial \mathbf{z}_i} = \begin{pmatrix} 0 \\ \vdots \\ 0 \\ \mathbf{J}_{2 \oplus \{\hat{\mathbf{x}}_{R_k}^B, \hat{\mathbf{z}}_i\}} \end{pmatrix}$$

EKF SLAM: add new features

$$\mathbf{P}_k^W = \begin{pmatrix} \mathbf{P}_R & \mathbf{P}_{RF_1} & \cdots & \mathbf{P}_{RF_n} \\ \mathbf{P}_{RF_1}^T & \mathbf{P}_{F_1} & \cdots & \mathbf{P}_{F_1 F_n} \\ \vdots & \vdots & \ddots & \vdots \\ \mathbf{P}_{RF_n}^T & \mathbf{P}_{F_1 F_n}^T & \cdots & \mathbf{P}_{F_n} \end{pmatrix}$$

$$\mathbf{P}_{k+}^W = \begin{pmatrix} \mathbf{P}_R & \mathbf{P}_{RF_1} & \cdots & \mathbf{P}_{RF_n} & \mathbf{P}_R \mathbf{J}_{1\oplus}^T \\ \mathbf{P}_{RF_1}^T & \mathbf{P}_{F_1} & \cdots & \mathbf{P}_{F_1 F_n} & \mathbf{P}_{RF_1}^T \mathbf{J}_{1\oplus}^T \\ \vdots & \vdots & \ddots & \vdots & \vdots \\ \mathbf{P}_{RF_n}^T & \mathbf{P}_{F_1 F_n}^T & \cdots & \mathbf{P}_{F_n} & \mathbf{P}_{RF_n}^T \mathbf{J}_{1\oplus}^T \\ \mathbf{J}_{1\oplus} \mathbf{P}_R & \mathbf{J}_{1\oplus} \mathbf{P}_{RF_1} & \cdots & \mathbf{J}_{1\oplus} \mathbf{P}_{RF_n} & \mathbf{J}_{1\oplus} \mathbf{P}_R \mathbf{J}_{1\oplus}^T + \mathbf{J}_{2\oplus} \mathbf{R}_k \mathbf{J}_{2\oplus}^T \end{pmatrix}$$

Adding new features is $O(n)$

The update step

EKF SLAM: map update

m observations:

$$\mathbf{z}_k = \mathbf{h}_k(\mathbf{x}_{\mathcal{F}_k}^B) + \mathbf{w}_k$$

$$\mathbf{z}_k \simeq \mathbf{h}_k(\hat{\mathbf{x}}_{\mathcal{F}_{k|k-1}}^B) + \mathbf{H}_k(\mathbf{x}_{\mathcal{F}_k}^B - \hat{\mathbf{x}}_{\mathcal{F}_{k|k-1}}^B)$$

$$\mathbf{H}_k = \left. \frac{\partial \mathbf{h}_k}{\partial \mathbf{x}_{\mathcal{F}_k}^B} \right|_{(\hat{\mathbf{x}}_{\mathcal{F}_{k|k-1}}^B)}$$

Filter gain: $\mathbf{K}_k = \mathbf{P}_{k|k-1}^B \mathbf{H}_k^T (\mathbf{H}_k \mathbf{P}_{k|k-1}^B \mathbf{H}_k^T + \mathbf{R}_k)^{-1}$

State update: $\hat{\mathbf{x}}_k^B = \hat{\mathbf{x}}_{k|k-1}^B + \mathbf{K}_k(\mathbf{z}_k - \mathbf{h}_k(\hat{\mathbf{x}}_{k|k-1}^B))$

Covariance update: $\mathbf{P}_k^B = (\mathbf{I} - \mathbf{K}_k \mathbf{H}_k) \mathbf{P}_{k|k-1}^B$

The innovation matrix

$$\begin{array}{ccccc} \mathbf{S}_k = & \mathbf{H}_k & \mathbf{P}_{k|k-1} & \mathbf{H}_k^T + \mathbf{R}_k & \\ r \times r & r \times n & n \times n & n \times r & r \times r \end{array}$$

$O(rn^2)$ operations?

The innovation matrix

$$\mathbf{S}_k = \mathbf{H}_k \mathbf{P}_{k|k-1} \mathbf{H}_k^T + \mathbf{R}_k$$

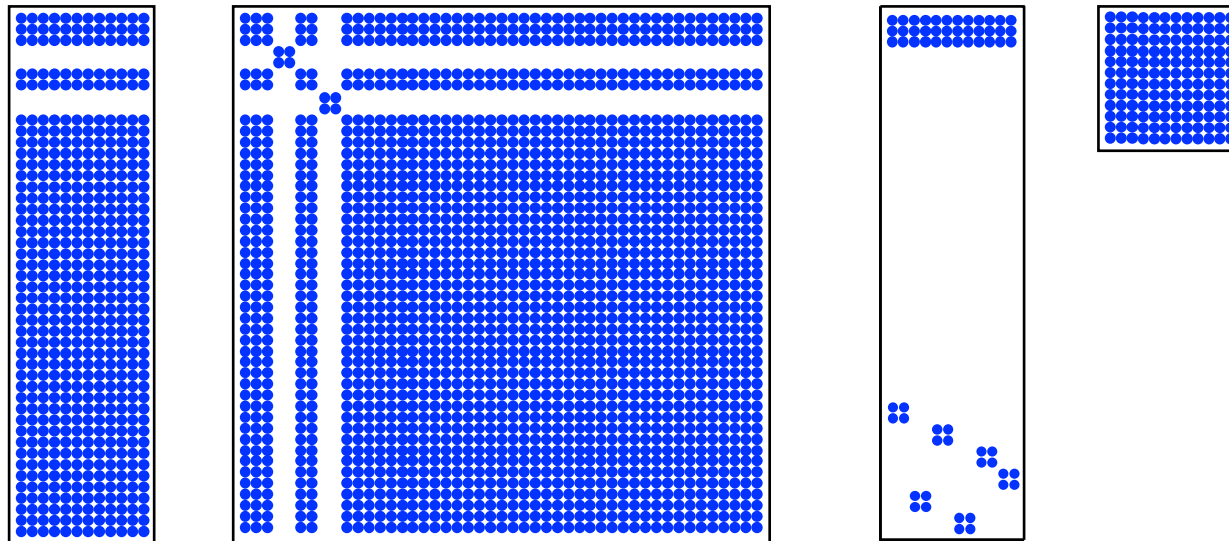
rxr rxn nxn txr rxr

$$O(rn) = O(n) \text{ operations}$$

The Kalman gain matrix

$$\mathbf{K}_k = \mathbf{P}_{k|k-1} \mathbf{H}_k^T (\mathbf{S}_k)^{-1}$$

$n \times r$ $n \times n$ $t \times r$ $r \times r$

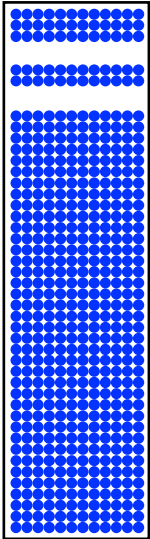


$$O(r^2n) = O(n) \text{ operations}$$

The covariance matrix

$$\mathbf{P}_k = (\mathbf{I} - \mathbf{K}_k \mathbf{H}_k) \mathbf{P}_{k|k-1}$$

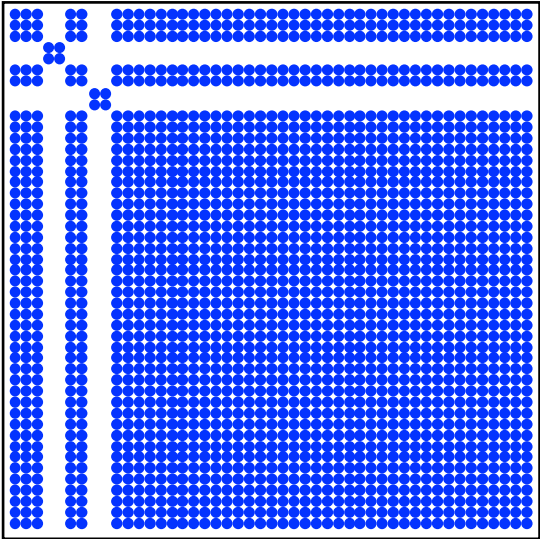
$$\dots = \dots \mathbf{K}_k \mathbf{H}_k \mathbf{P}_{k|k-1}$$



$n \times r$



$r \times t$



$n \times n$

$$O(rn^2) = O(n^2) \text{ operations}$$

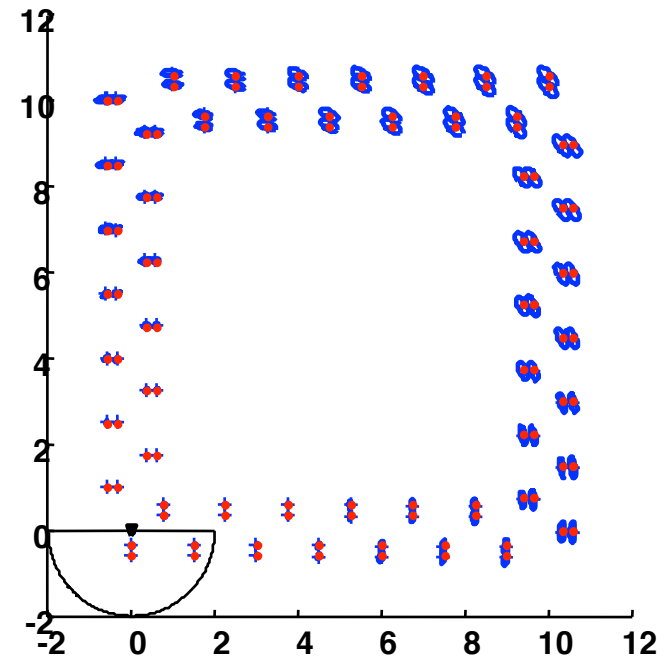
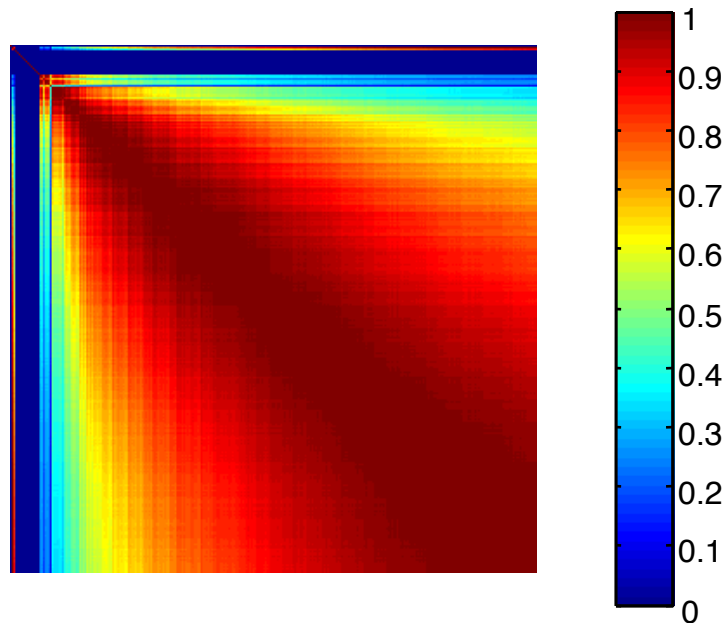
Efforts to reduce complexity

- Decoupled Stochastic Mapping (Leonard and Feder, 2000) (Jensfelt 2001) **$O(1)$**
- Local Mapping Algorithm (Chong and Kleeman 1999) **$O(1)$**
- Suboptimal SLAM (Guivant and Nebot 2001) **$O(n)$**
- Sparse Weight Filter (Julier 2001) **$O(n)$**
- Sparse Extended Information Filter (Thrun et al 2003) **$O(1)$**
- Postponement (Davidson 1998, Knight, Davidson and Reed 2001)
- Compressed Filter (Guivant and Nebot 2001)
- Constrained Local Submap Filter (Williams 2001)
- Map Joining (Tardós et. al, 2002)

Approximate, or pessimistic solutions

**Exact solutions that delay global map updating, and strongly reduce cost.
But still $O(n^2)$**

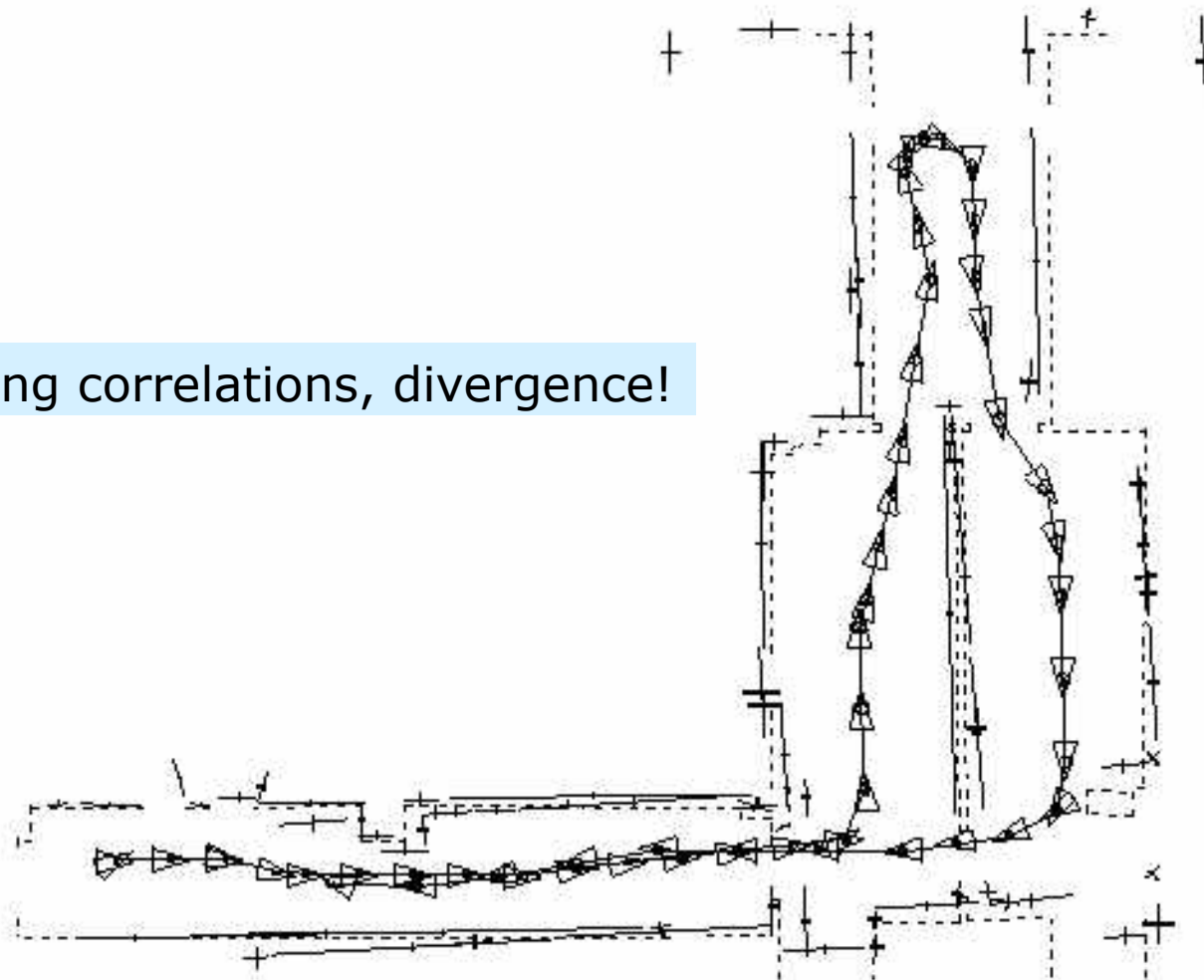
Are correlations necessary?



```
sigmas = sqrt(diag(Cov))' ;  
Corr=diag(1./sigmas)*Cov*diag(1./sigmas) ;
```

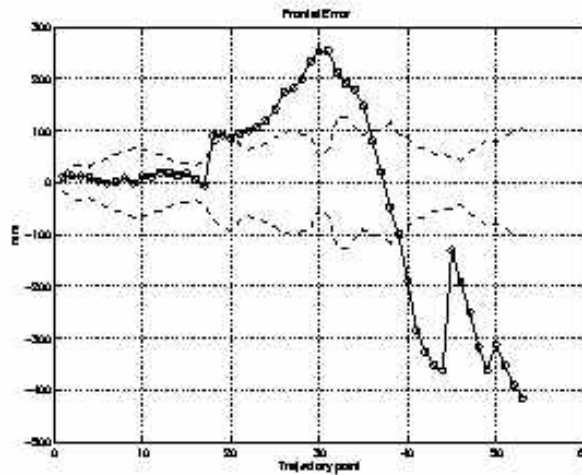
The importance of Correlations

Ignoring correlations, divergence!

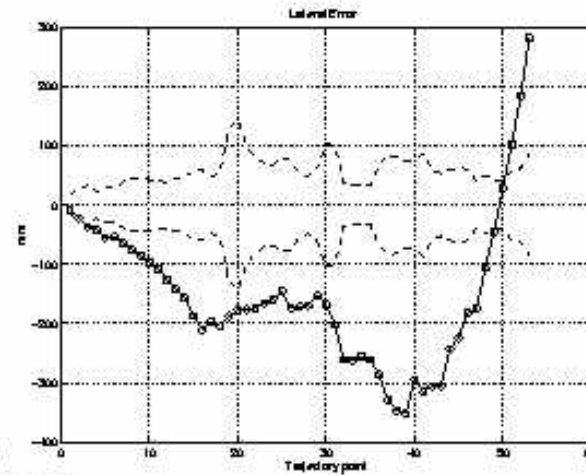


The importance of Correlations

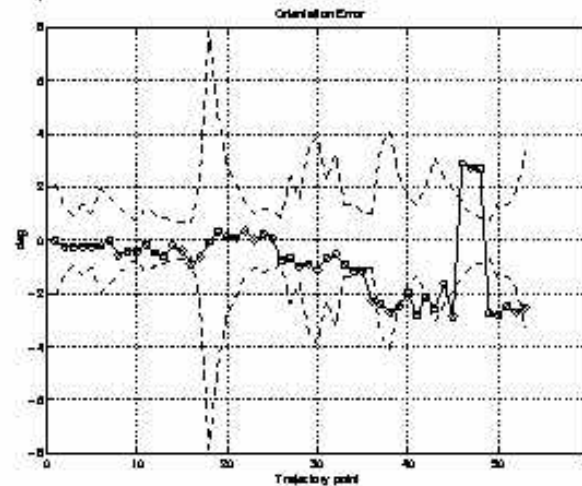
Frontal



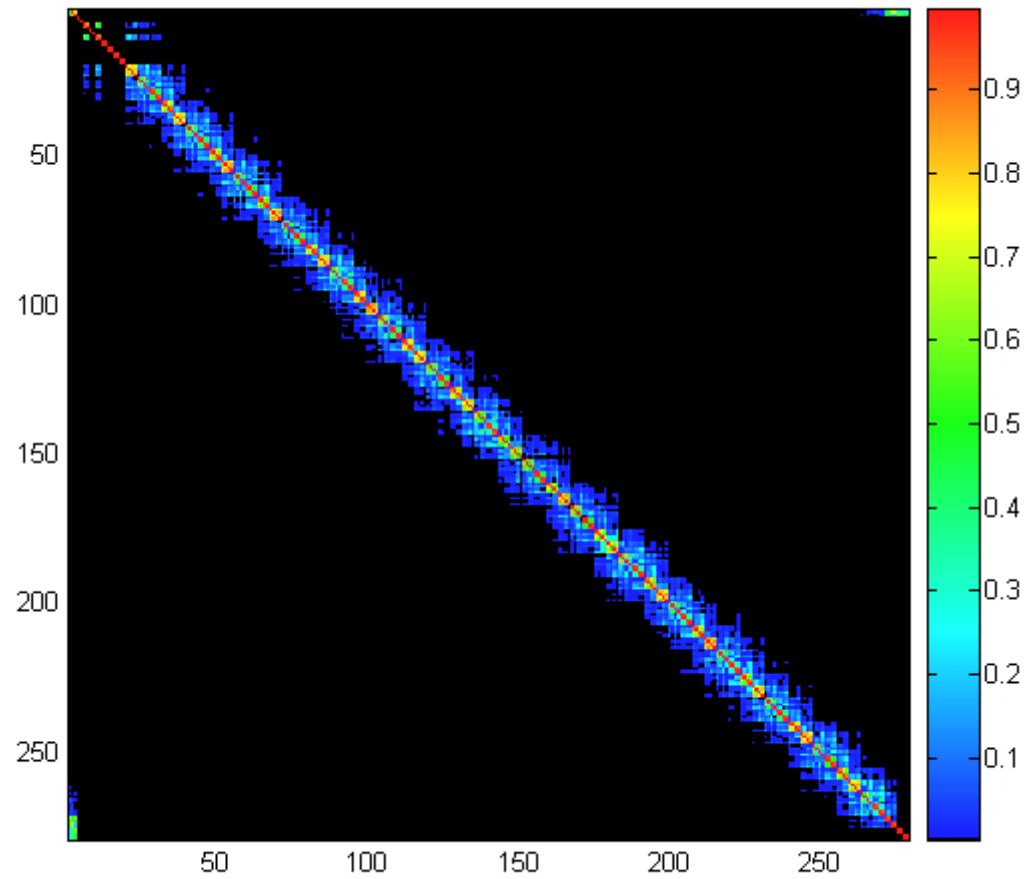
Lateral



Angular



Inverse correlations



This observation is the basis of SEIFs

The linearization problem

Consistency of EKF-SLAM

- Nice “**convergence**” properties of P_k^W (Dissanayake et al. 2001):
 - Landmark covariance decreases monotonically
 - In the limit, landmarks become fully correlated
 - In the limit, landmark covariance reaches a lower bound related to the initial vehicle covariance
- But SLAM is a non-linear problem
 - The inherent approximations due to linearizations can lead to **divergence (inconsistency)** of the EKF
 - » see for example (Jazwinski, 1970)

EKF-SLAM: Robot Motion

$$\mathbf{x}_{R_k}^B = \mathbf{x}_{R_{k-1}}^B \oplus \mathbf{x}_{R_k}^{R_{k-1}}$$

Odometry model (white noise):

$$\begin{aligned}\mathbf{x}_{R_k}^{R_{k-1}} &= \hat{\mathbf{x}}_{R_k}^{R_{k-1}} + \mathbf{v}_k \\ E[\mathbf{v}_k] &= \mathbf{0} \\ E[\mathbf{v}_k \mathbf{v}_j^T] &= \delta_{kj} \mathbf{Q}_k\end{aligned}$$

EKF prediction:

$$\begin{aligned}\hat{\mathbf{x}}_{k|k-1}^B &= \begin{bmatrix} \hat{\mathbf{x}}_{R_{k-1}}^B \oplus \hat{\mathbf{x}}_{R_k}^{R_{k-1}} \\ \hat{\mathbf{x}}_{F_{1,k-1}}^B \\ \vdots \\ \hat{\mathbf{x}}_{F_{m,k-1}}^B \end{bmatrix} & \mathbf{F}_k &= \begin{bmatrix} \mathbf{J}_{1 \oplus \left\{ \hat{\mathbf{x}}_{R_{k-1}}^B, \hat{\mathbf{x}}_{R_k}^{R_{k-1}} \right\}} & \mathbf{0} & \cdots & \mathbf{0} \\ \mathbf{0} & \mathbf{I} & & \vdots \\ \vdots & & \ddots & \\ \mathbf{0} & \cdots & & \mathbf{I} \end{bmatrix} \\ \mathbf{P}_{\mathcal{F}_{k|k-1}}^B &= \mathbf{F}_k \mathbf{P}_{k-1}^B \mathbf{F}_k^T + \mathbf{G}_k \mathbf{Q}_k \mathbf{G}_k^T & \mathbf{G}_k &= \begin{bmatrix} \mathbf{J}_{2 \oplus \left\{ \hat{\mathbf{x}}_{R_{k-1}}^B, \hat{\mathbf{x}}_{R_k}^{R_{k-1}} \right\}} \\ \mathbf{0} \\ \vdots \\ \mathbf{0} \end{bmatrix}\end{aligned}$$

Linearization

EKF-SLAM: Map Update

Feature observations:

$$\mathbf{z}_k = \mathbf{h}_k(\mathbf{x}_k^B) + \mathbf{w}_k$$

$$\mathbf{z}_k \simeq \mathbf{h}_k(\hat{\mathbf{x}}_{k|k-1}^B) + \mathbf{H}_k(\mathbf{x}_k^B - \hat{\mathbf{x}}_{k|k-1}^B)$$

$$\mathbf{H}_k = \left. \frac{\partial \mathbf{h}_k}{\partial \mathbf{x}_k^B} \right|_{(\hat{\mathbf{x}}_{k|k-1}^B)}$$

Linearization

EKF map update:

$$\hat{\mathbf{x}}_k^B = \hat{\mathbf{x}}_{k|k-1}^B + \mathbf{K}_k(\mathbf{z}_k - \mathbf{h}_k(\hat{\mathbf{x}}_{k|k-1}^B))$$

$$\mathbf{P}_k^B = (\mathbf{I} - \mathbf{K}_k \mathbf{H}_k) \mathbf{P}_{k|k-1}^B$$

$$\mathbf{K}_k = \mathbf{P}_{k|k-1}^B \mathbf{H}_k^T (\mathbf{H}_k \mathbf{P}_{k|k-1}^B \mathbf{H}_k^T + \mathbf{R}_k)^{-1}$$

The Consistency Problem

True map value

$$\mathbf{x}_k^W$$

EKF-SLAM
estimation

$$\hat{\mathbf{x}}_k^W$$

$$\mathbf{P}_k^W$$

- An estimator is **consistent** if:

$$E \left[\mathbf{x}_k^W - \hat{\mathbf{x}}_k^W \right] = 0$$

$$E \left[\left(\mathbf{x}_k^W - \hat{\mathbf{x}}_k^W \right) \left(\mathbf{x}_k^W - \hat{\mathbf{x}}_k^W \right)^T \right] = \mathbf{P}_k^W$$

Unbiased

The Mean Square Error matches the filter computed Covariance

- Pessimistic covariance is OK (but not too pessimistic)
- Optimistic covariance = Inconsistency = Filter divergence

Consistency Testing

1. Normalized Estimation Error Squared NEES

$$D^2 = \left(\mathbf{x}_k^W - \hat{\mathbf{x}}_k^W \right)^T \left(\mathbf{P}_k^W \right)^{-1} \left(\mathbf{x}_k^W - \hat{\mathbf{x}}_k^W \right)$$

$$D^2 \leq \chi_{r,1-\alpha}^2$$

**True map required
→ Simulations**

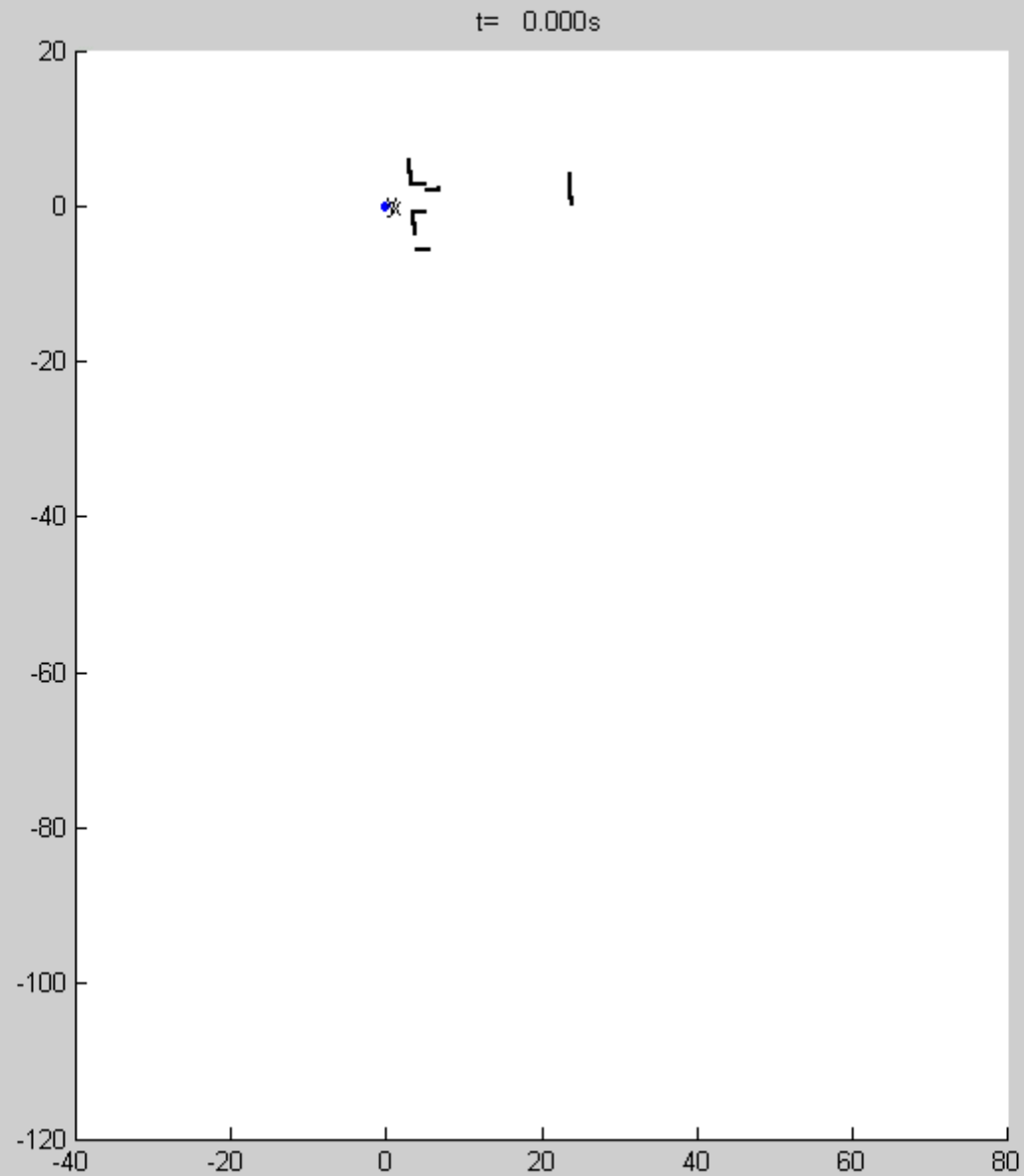
2. Innovation test (observation $i \rightarrow$ map feature j)

$$D_{ij}^2 = \left(\mathbf{z}_i - \mathbf{h}_j(\hat{\mathbf{x}}_k^W) \right)^T \left(\mathbf{H}_j \mathbf{P}_k^W \mathbf{H}_j^T + \mathbf{R}_i \right)^{-1} \left(\mathbf{z}_i - \mathbf{h}_j(\hat{\mathbf{x}}_k^W) \right)$$

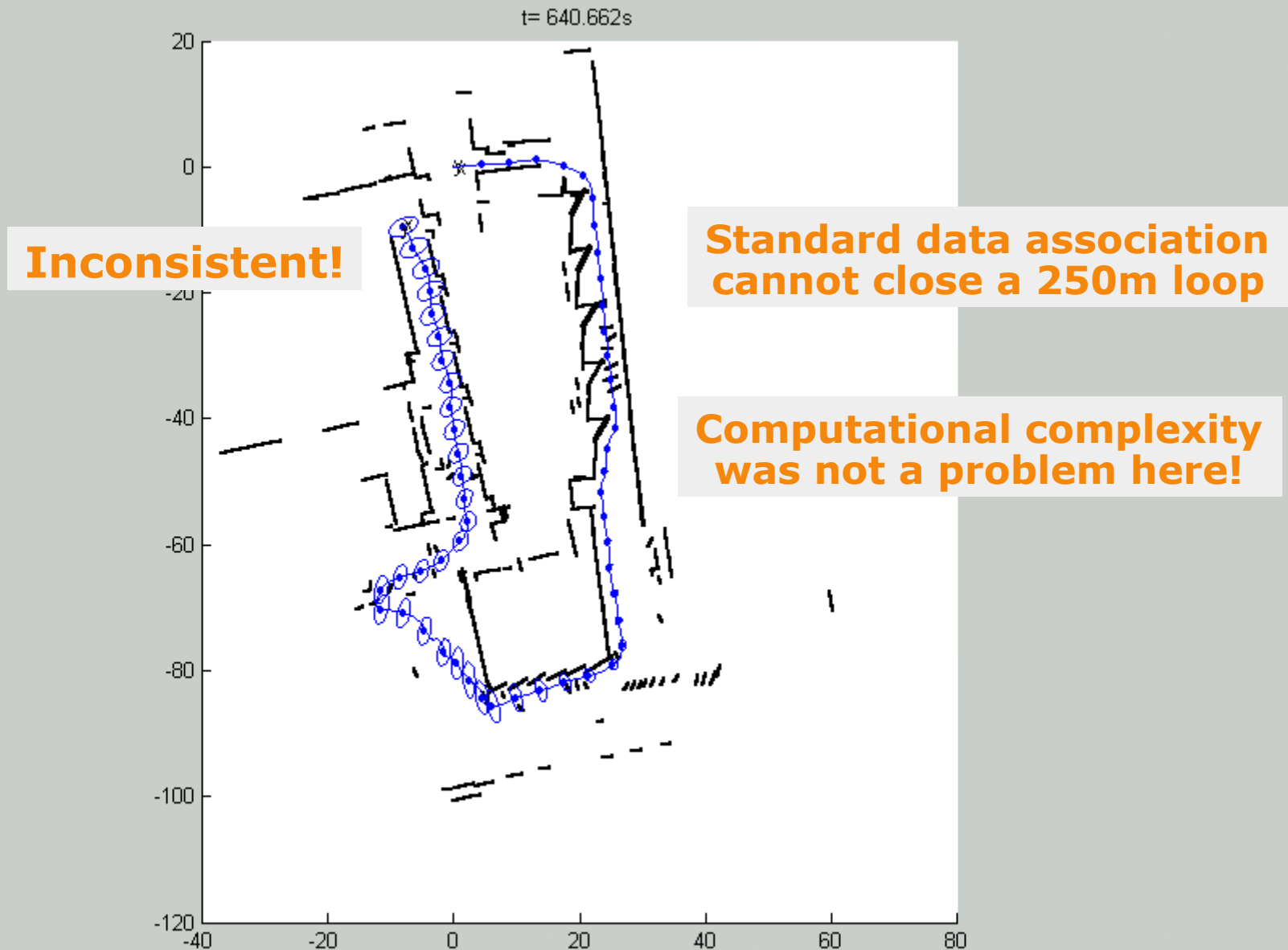
$$D_{ij}^2 \leq \chi_{d,1-\alpha}^2$$

**Critical when
closing big
loops**

EKF-SLAM: Real Example



EKF-SLAM: Real Example

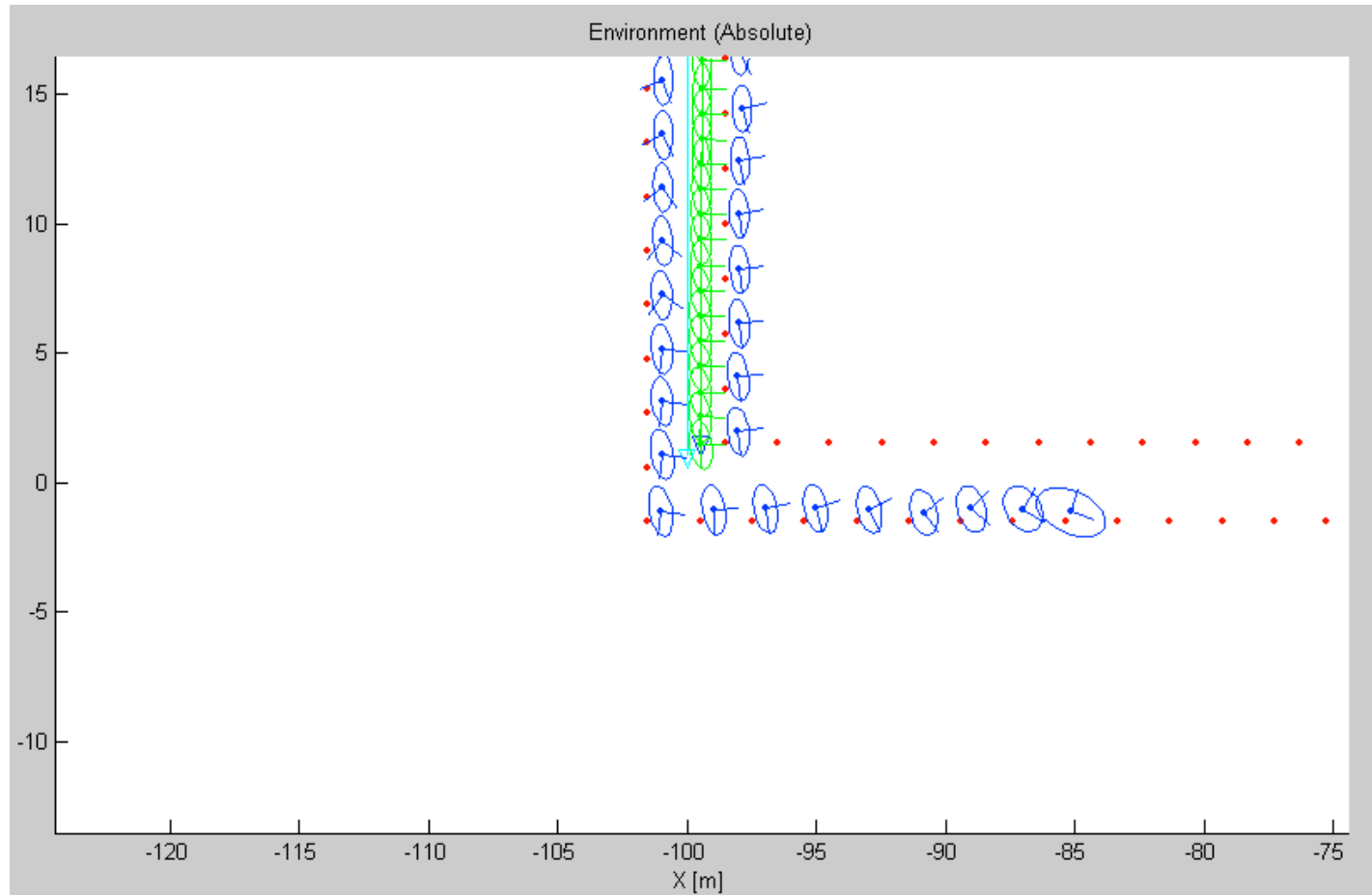


EKF-SLAM: Simulation

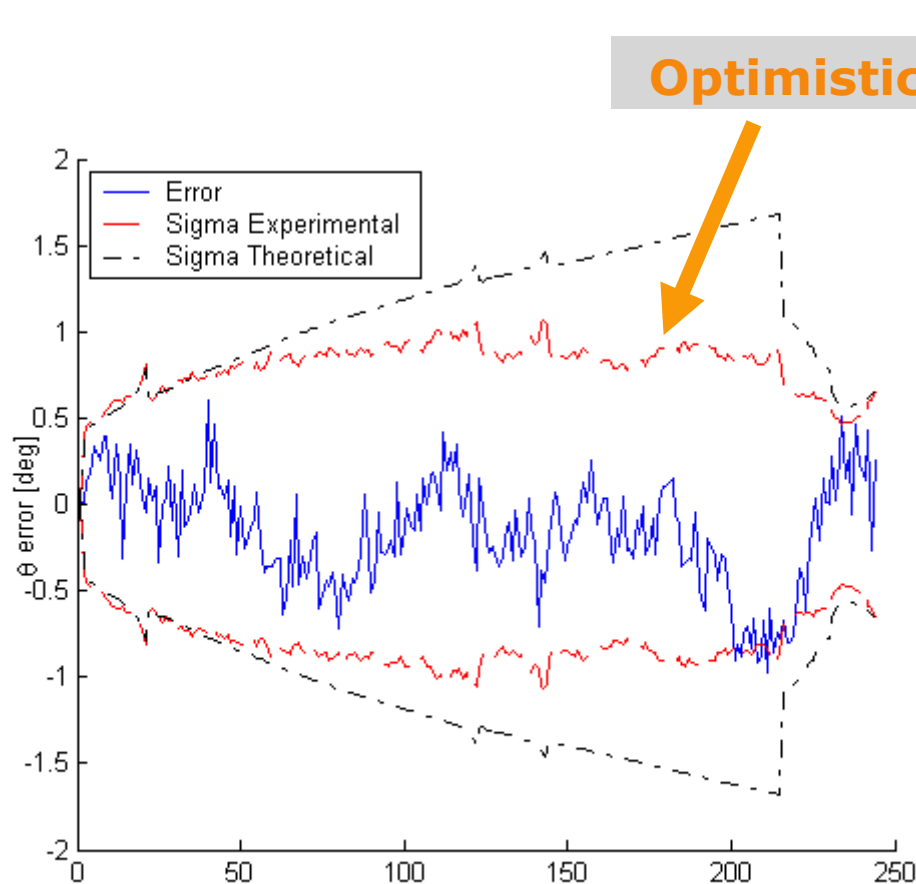
- Simulation conditions
 - Perfect data association
 - Ideal odometry and measurement noise
 - » white, Gaussian, known covariance
- Advantages of simulation:
 - Consistency can be tested against the true map
 - A simulation with noise=0 gives the theoretical map covariance (without linearization errors)

EKF-SLAM: Simulat

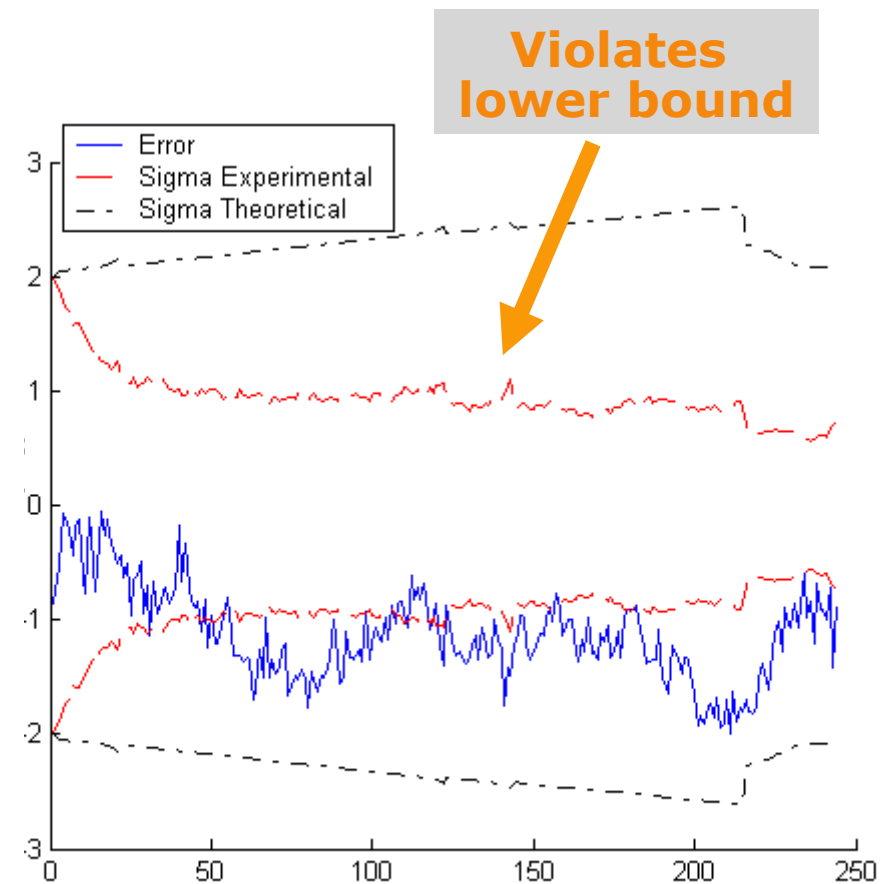
Perfect data association
and noise model



EKF-SLAM: Covariance



Initial uncertainty = 0

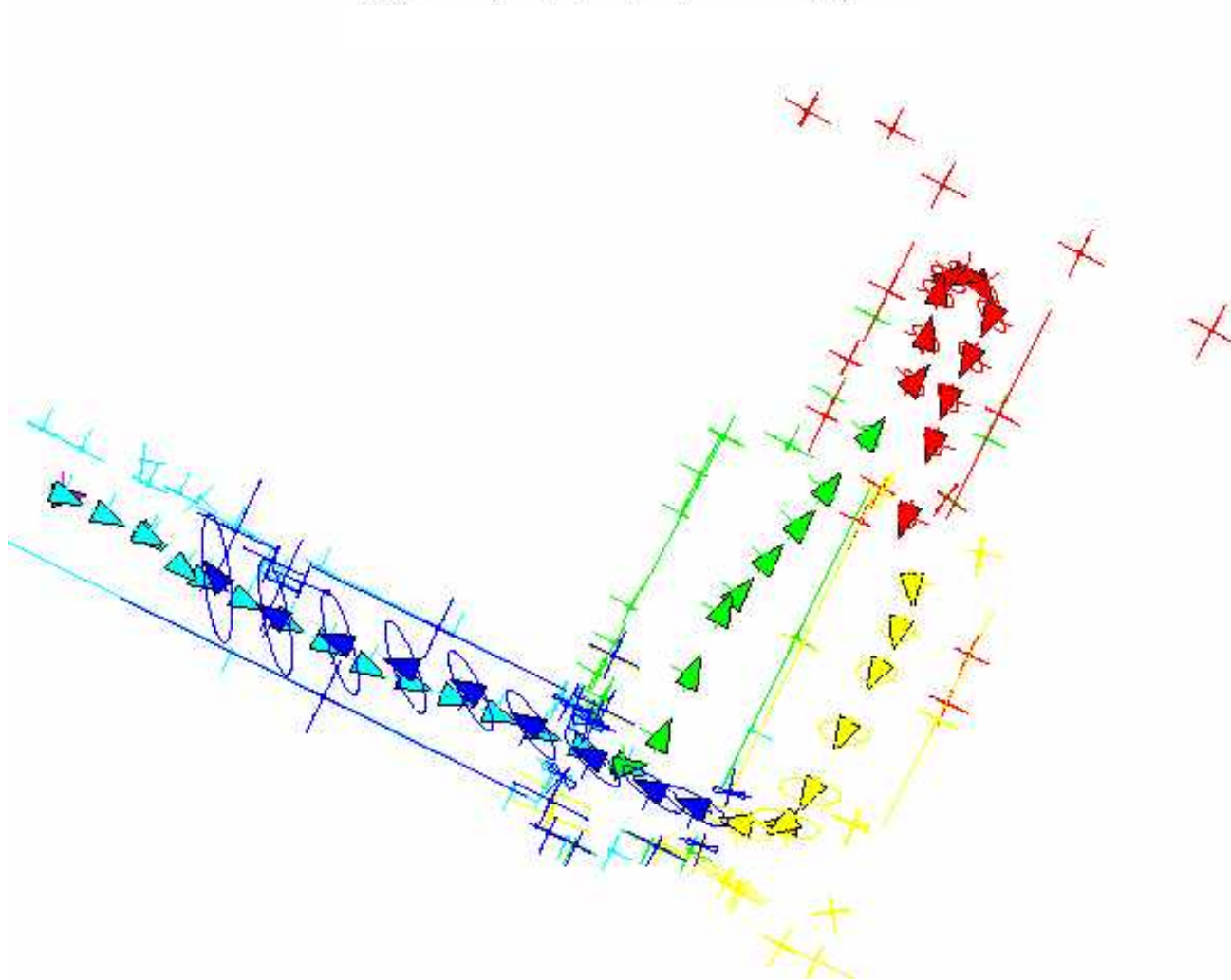


Initial uncertainty > 0

J.A. Castellanos, J. Neira, J.D. Tardós, **Limits to the Consistency of EKF-based SLAM**, 5th IFAC Symposium on Intelligent Autonomous Vehicles, Lisbon, July 2004

Overcoming these problems: Local maps

Local mapping



Local map building

- Periodically, the robot starts a new map, relative to its current location:

$$\hat{\mathbf{x}}_{R_0}^B = \mathbf{0}$$

$$\mathbf{P}_{R_0}^B = \mathbf{0}$$

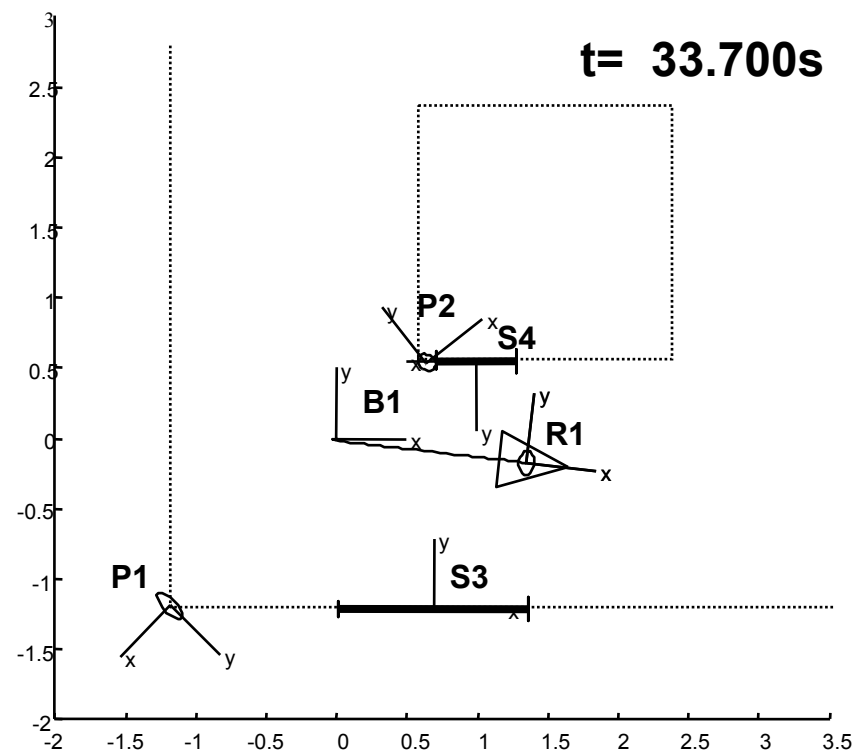
- Given measurements:

$$D^{1\dots k_1} = \left\{ \mathbf{u}_1 \mathbf{z}_1 \dots \mathbf{u}_{k_1} \mathbf{z}_{k_1} \right\}$$

$$\mathbf{u}_k = \hat{\mathbf{x}}_{R_k}^{R_{k-1}}$$

- EKF approximates the conditional mean:

$$\hat{\mathbf{x}}_{\mathcal{F}_1}^{B_1} \simeq E \left[\mathbf{x}_{\mathcal{F}_1}^{B_1} \mid D^{1\dots k_1}, \mathcal{H}^{1\dots k_1} \right]$$



Local map building

- Second map: $D^{k_1+1 \dots k_2} = \left\{ \mathbf{u}_{k_1+1} \mathbf{z}_{k_1+1} \dots \mathbf{u}_{k_2} \mathbf{z}_{k_2} \right\}$

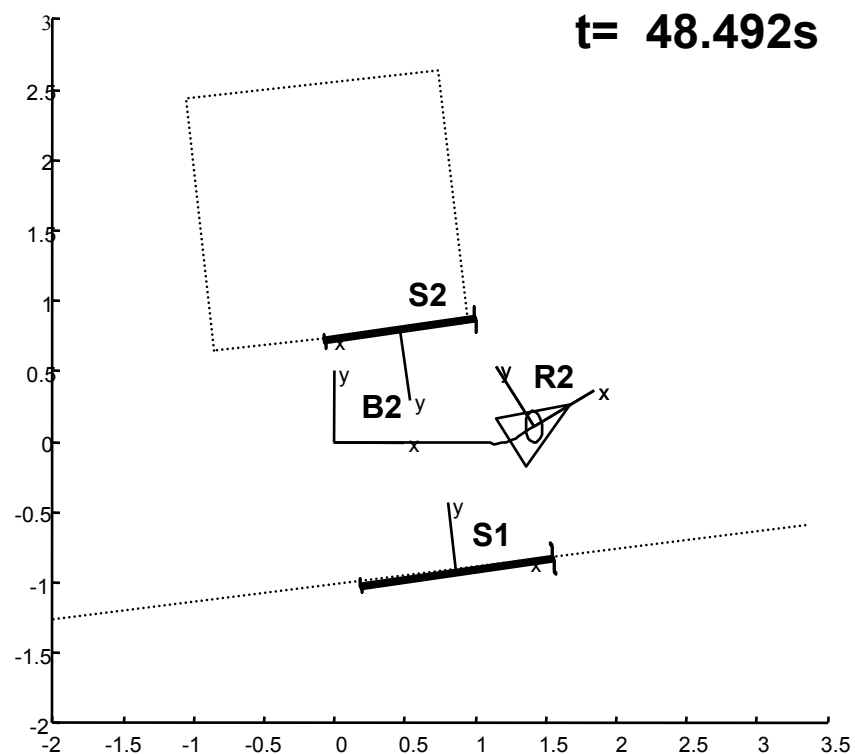
$$\hat{\mathbf{x}}_{\mathcal{F}_2}^{B_2} \simeq E \left[\mathbf{x}_{\mathcal{F}_2}^{B_2} \mid D^{k_1+1 \dots k_2}, \mathcal{H}^{k_1+1 \dots k_2} \right]$$

- No information is shared:

$$D^{1 \dots k_1} \cap D^{k_1+1 \dots k_2} = \emptyset$$

Maps are uncorrelated

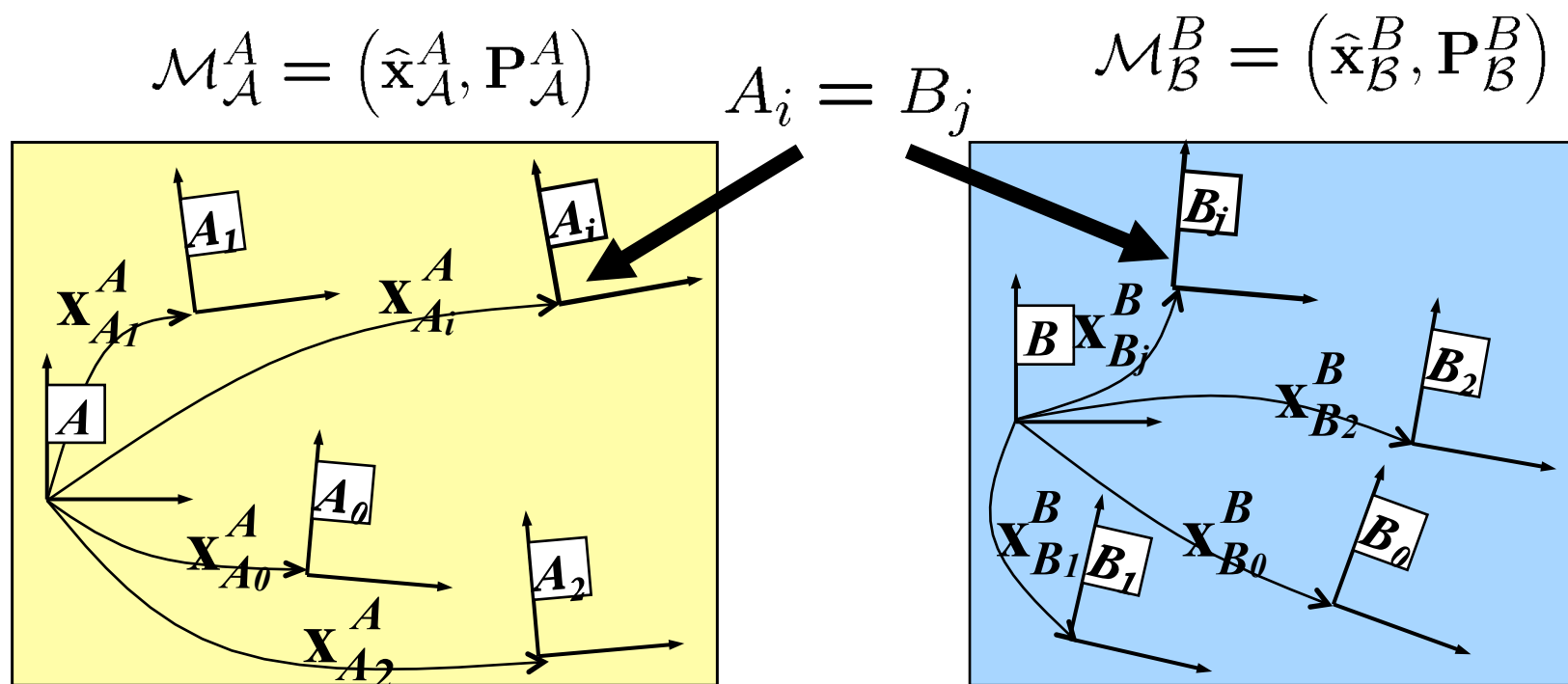
- Common reference:
 $B_2 = R_1$



Map Joining

Map Joining

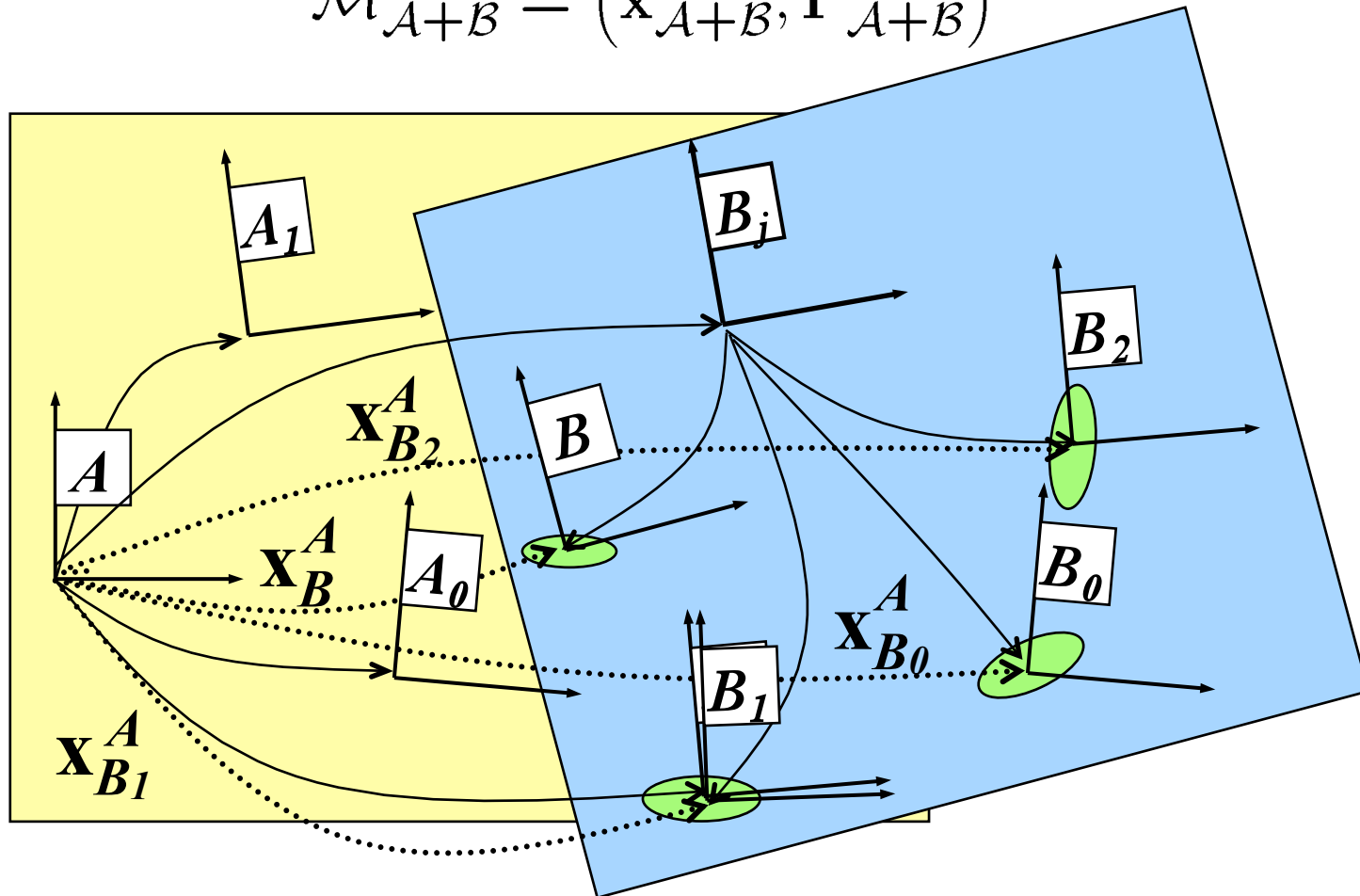
- Given:
 - Two statistically independent stochastic maps
 - A common reference



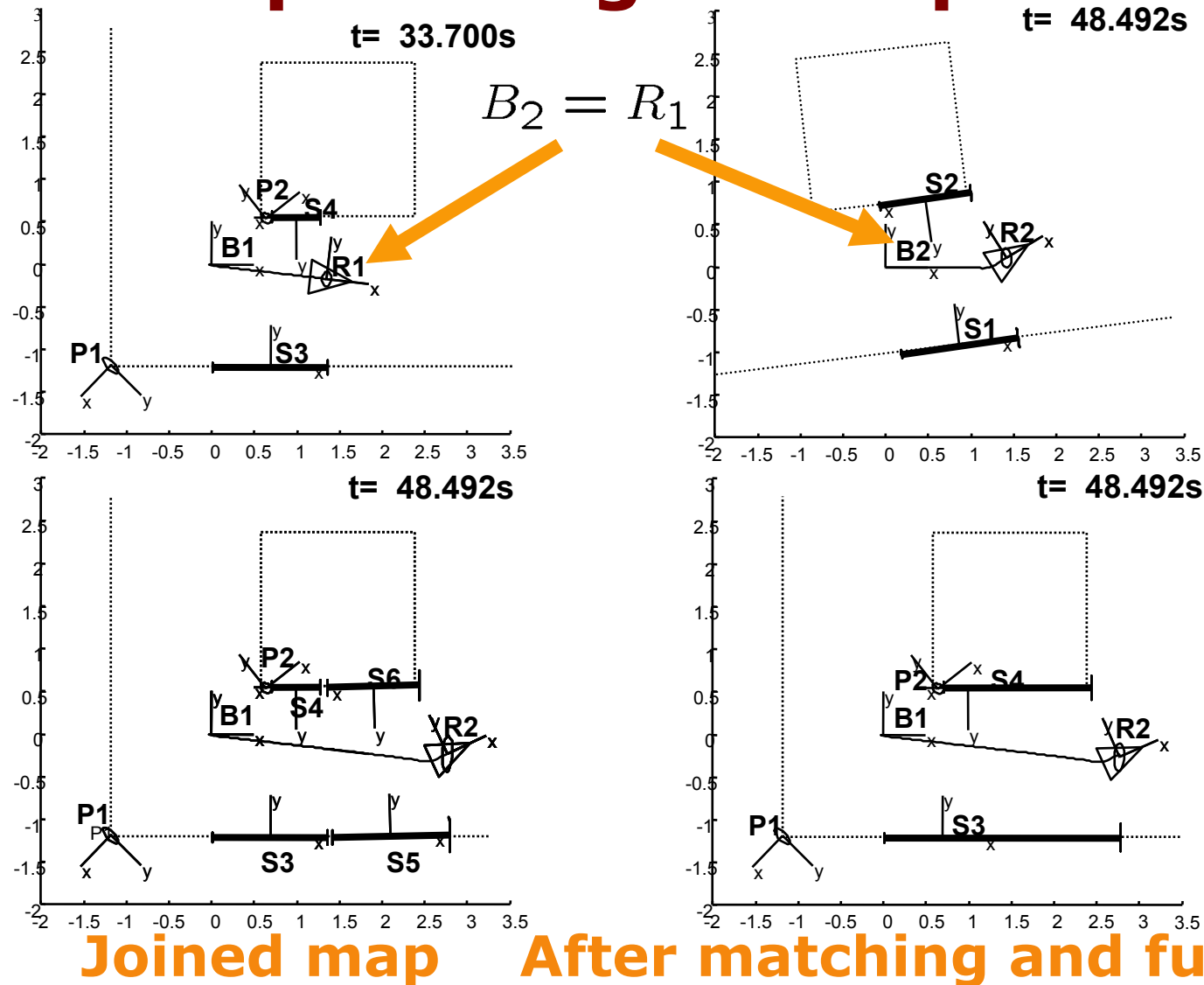
Map Joining

- Conveys the information of the two maps into a single **fully consistent** stochastic map:

$$\mathcal{M}_{A+B}^A = (\hat{\mathbf{x}}_{A+B}^A, \mathbf{P}_{A+B}^A)$$



Map Joining: Example



Joined map After matching and fusion

Map Joining

- New state vector: $\hat{\mathbf{x}}_{\mathcal{A}+\mathcal{B}}^A = \begin{bmatrix} \hat{\mathbf{x}}_{\mathcal{A}}^A \\ \hat{\mathbf{x}}_{\mathcal{B}}^A \end{bmatrix} = \begin{bmatrix} \hat{\mathbf{x}}_{A_i}^A \oplus \hat{\mathbf{x}}_{B_0}^{B_j} \\ \vdots \\ \hat{\mathbf{x}}_{A_i}^A \oplus \hat{\mathbf{x}}_{B_m}^{B_j} \end{bmatrix}$
- New covariance matrix:

$$\begin{aligned} \mathbf{P}_{\mathcal{A}+\mathcal{B}}^A &= \mathbf{J}_{\mathcal{A}}^{A+\mathcal{B}} \mathbf{P}_{\mathcal{A}}^A \mathbf{J}_{\mathcal{A}}^{A+\mathcal{B}T} + \mathbf{J}_{\mathcal{B}}^{A+\mathcal{B}} \mathbf{P}_{\mathcal{B}}^{B_j} \mathbf{J}_{\mathcal{B}}^{A+\mathcal{B}T} \\ &= \begin{bmatrix} \mathbf{P}_{\mathcal{A}}^A & \mathbf{P}_{\mathcal{A}}^A \mathbf{J}_1^T \\ \mathbf{J}_1 \mathbf{P}_{\mathcal{A}}^A & \mathbf{J}_1 \mathbf{P}_{\mathcal{A}}^A \mathbf{J}_1^T \end{bmatrix} + \begin{bmatrix} \mathbf{0} & \mathbf{0} \\ \mathbf{0} & \mathbf{J}_2 \mathbf{P}_{\mathcal{B}}^{B_j} \mathbf{J}_2^T \end{bmatrix} \end{aligned}$$

$$\begin{aligned} \mathbf{J}_{\mathcal{A}}^{A+\mathcal{B}} &= \frac{\partial \hat{\mathbf{x}}_{\mathcal{A}+\mathcal{B}}^A}{\partial \hat{\mathbf{x}}_{\mathcal{A}}^A} = \begin{bmatrix} \mathbf{I} \\ \mathbf{J}_1 \end{bmatrix} & \mathbf{J}_1 &= \begin{bmatrix} \mathbf{0} & \cdots & \mathbf{J}_{1\oplus} \left\{ \hat{\mathbf{x}}_{A_i}^A, \hat{\mathbf{x}}_{B_0}^{B_j} \right\} & \cdots & \mathbf{0} \\ \vdots & & \vdots & & \vdots \\ \mathbf{0} & \cdots & \mathbf{J}_{1\oplus} \left\{ \hat{\mathbf{x}}_{A_i}^A, \hat{\mathbf{x}}_{B_m}^{B_j} \right\} & \cdots & \mathbf{0} \end{bmatrix} \\ \mathbf{J}_{\mathcal{B}}^{A+\mathcal{B}} &= \frac{\partial \hat{\mathbf{x}}_{\mathcal{A}+\mathcal{B}}^A}{\partial \hat{\mathbf{x}}_{\mathcal{B}}^{B_j}} = \begin{bmatrix} \mathbf{0} \\ \mathbf{J}_2 \end{bmatrix} & \mathbf{J}_2 &= \begin{bmatrix} \mathbf{J}_{2\oplus} \left\{ \hat{\mathbf{x}}_{A_i}^A, \hat{\mathbf{x}}_{B_0}^{B_j} \right\} & \cdots & \mathbf{0} \\ \vdots & \ddots & \vdots \\ \mathbf{0} & \cdots & \mathbf{J}_{2\oplus} \left\{ \hat{\mathbf{x}}_{A_i}^A, \hat{\mathbf{x}}_{B_m}^{B_j} \right\} \end{bmatrix} \end{aligned}$$

Matching and Fusion

- Matching function:

$$\mathbf{f}_{ij_i}(\mathbf{x}) = 0$$

- Joint matching function for the hypothesis:

$$\mathbf{f}_{\mathcal{H}}(\mathbf{x}) = \begin{bmatrix} \mathbf{f}_{1j_1}(\mathbf{x}) \\ \vdots \\ \mathbf{f}_{mj_m}(\mathbf{x}) \end{bmatrix} \simeq \mathbf{h}_{\mathcal{H}} + \mathbf{H}_{\mathcal{H}}(\mathbf{x} - \hat{\mathbf{x}}) = 0$$

- Joint innovation test:

$$D_{\mathcal{H}}^2 = \mathbf{h}_{\mathcal{H}}^T (\mathbf{H}_{\mathcal{H}} \mathbf{P} \mathbf{H}_{\mathcal{H}}^T)^{-1} \mathbf{h}_{\mathcal{H}} < \chi_{d,\alpha}^2$$

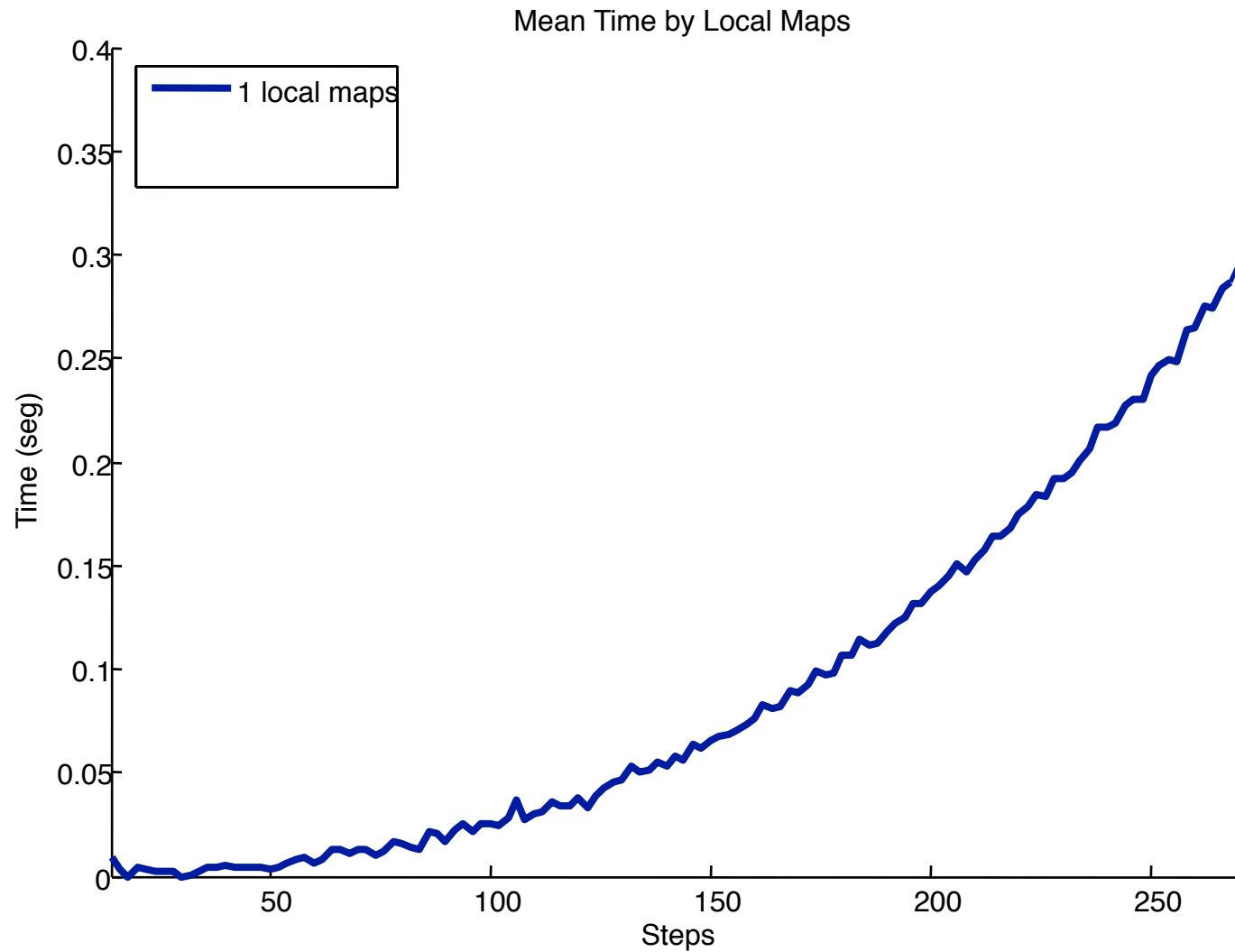
- Map update using EKF:

$$\hat{\mathbf{x}}_k = \hat{\mathbf{x}}_{k-1} - \mathbf{K}_k \mathbf{h}_{\mathcal{H}}$$

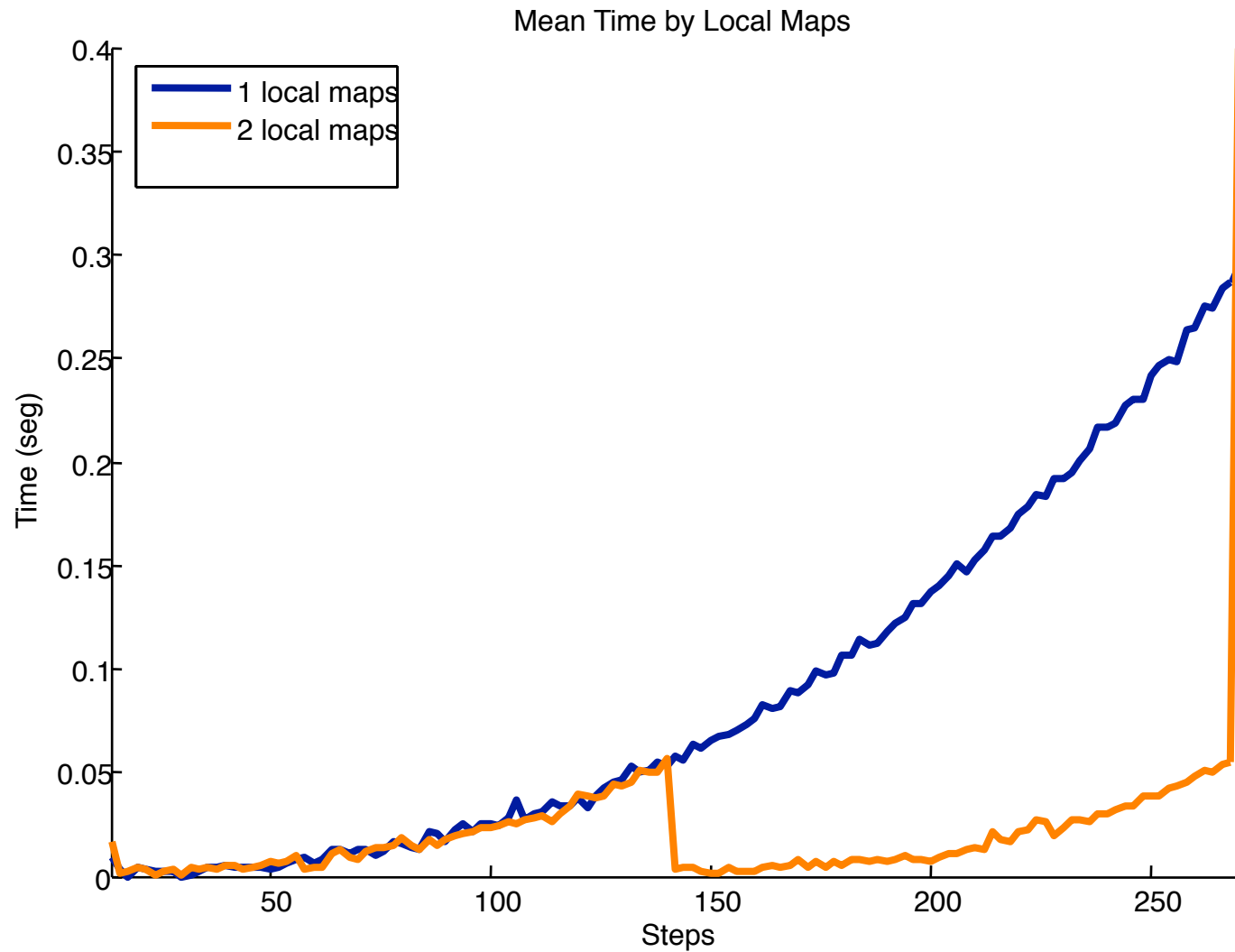
$$\mathbf{P}_k = (\mathbf{I} - \mathbf{K}_k \mathbf{H}_{\mathcal{H}}) \mathbf{P}_{k-1}$$

$$\mathbf{K}_k = \mathbf{P}_{k-1} \mathbf{H}_{\mathcal{H}}^T (\mathbf{H}_{\mathcal{H}} \mathbf{P}_{k-1} \mathbf{H}_{\mathcal{H}}^T)^{-1}$$

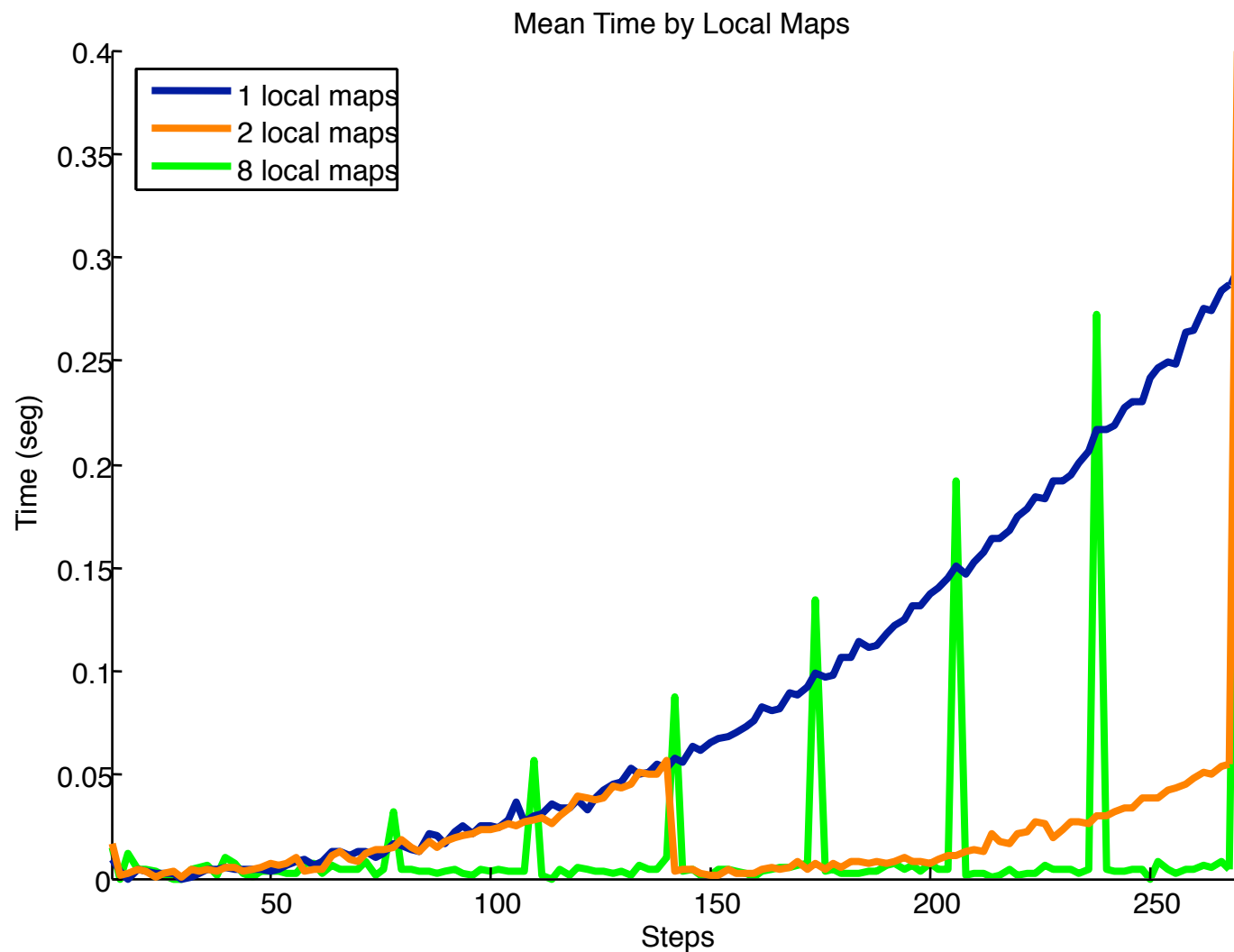
EKF updates



Map Joining



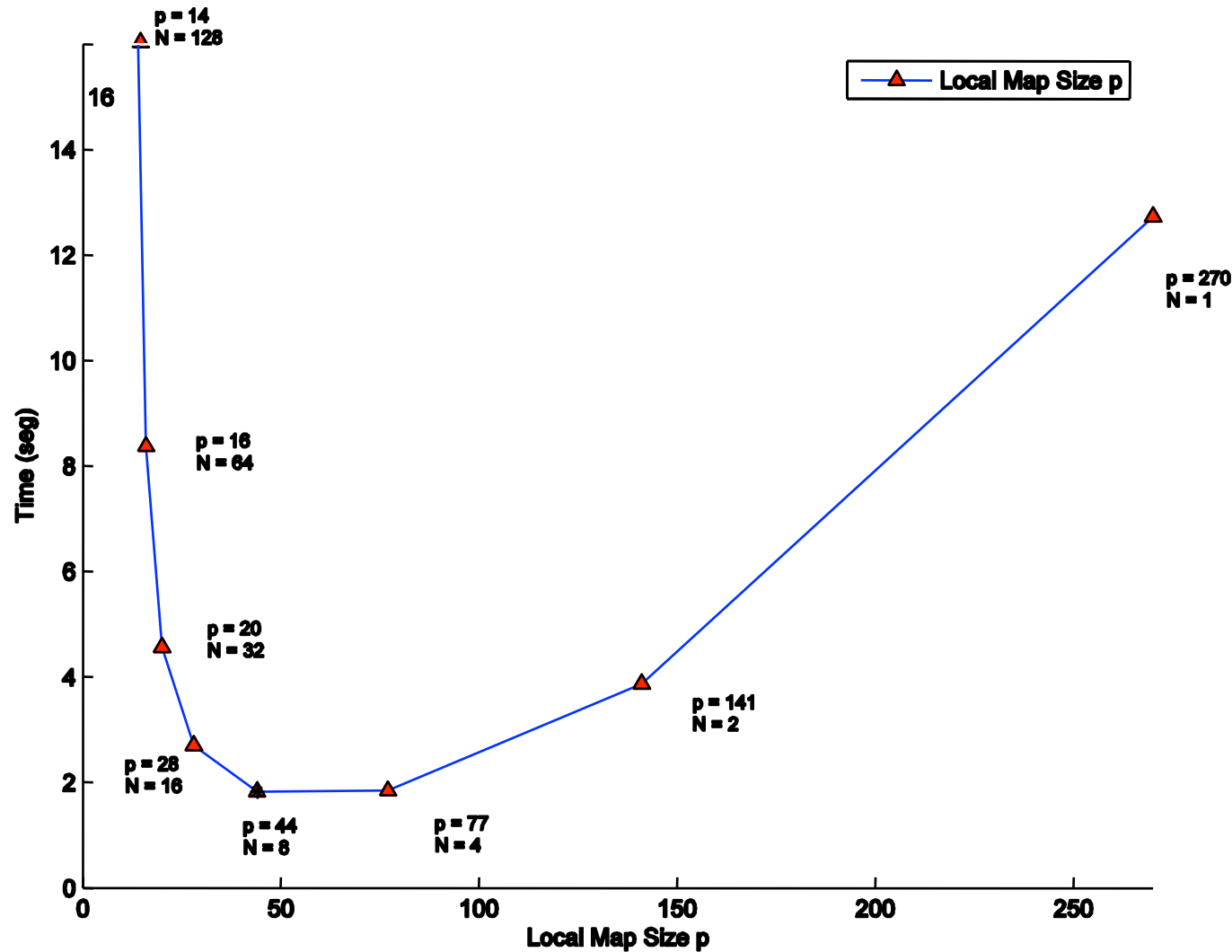
Map Joining



Map joining is $O(n^2)$

Local map size

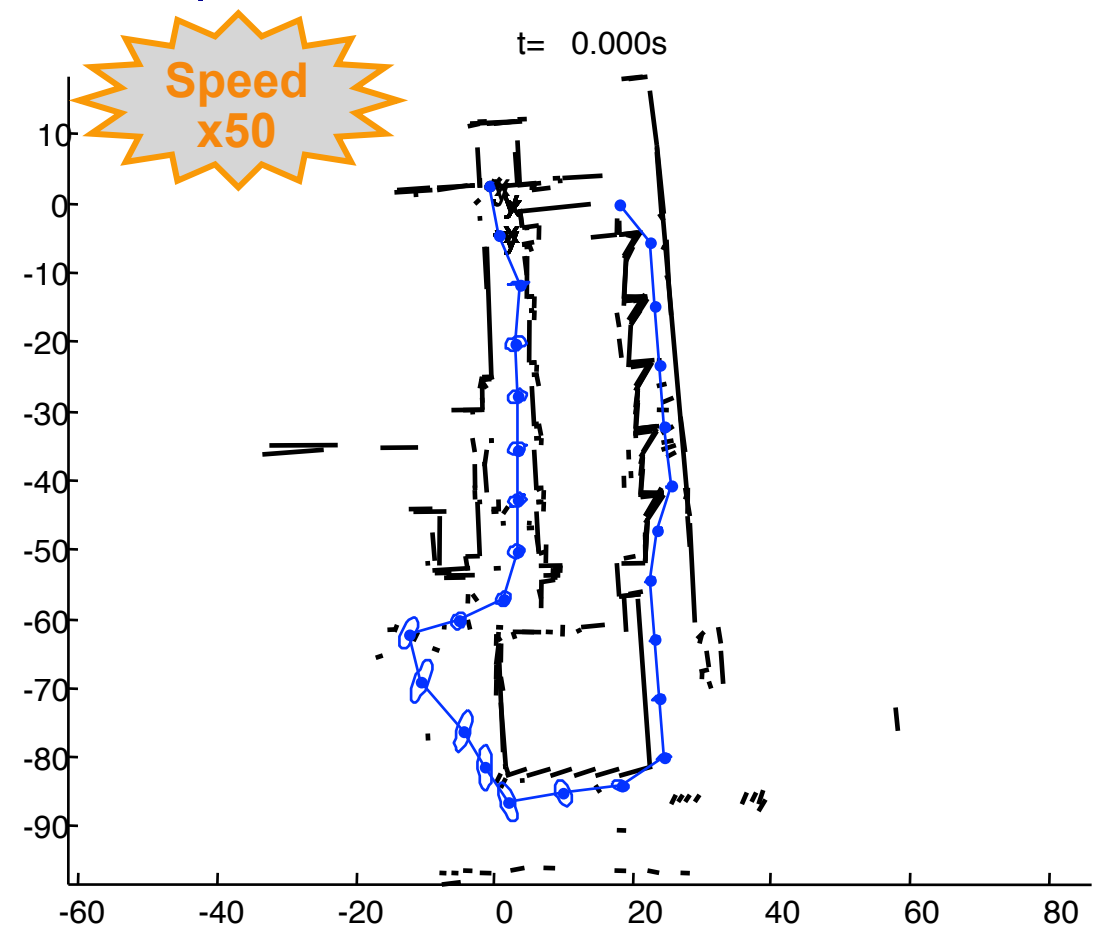
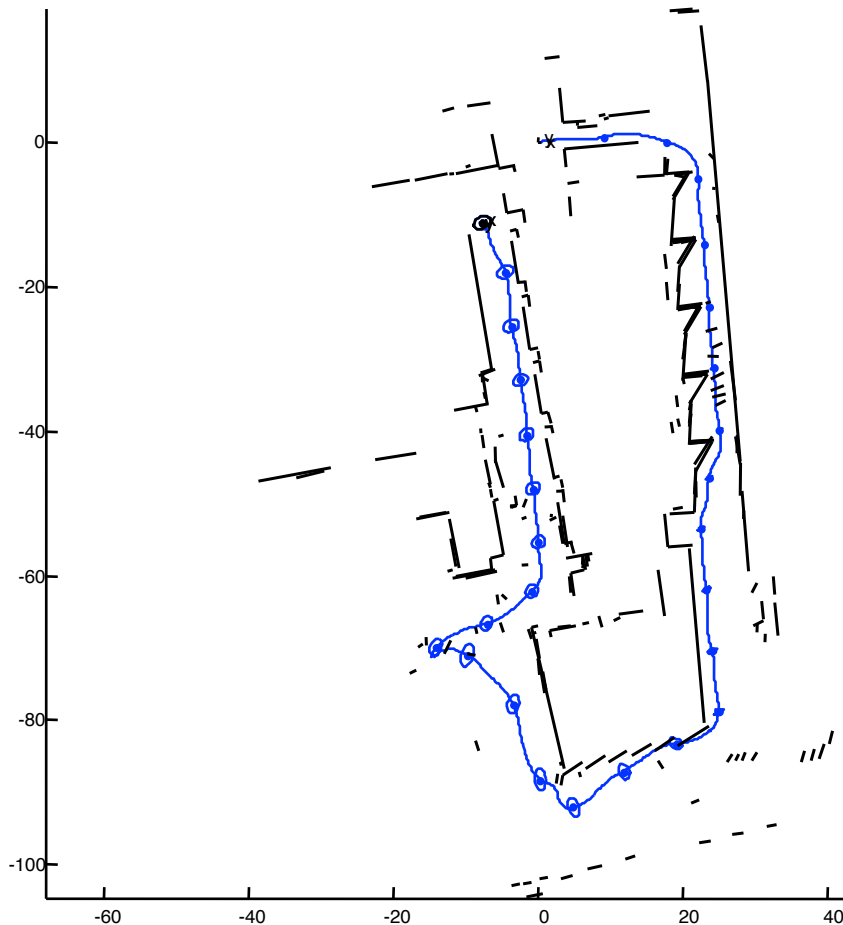
- Size matters!



Computational cost for
different local map sizes

Map Joining closes the loop!

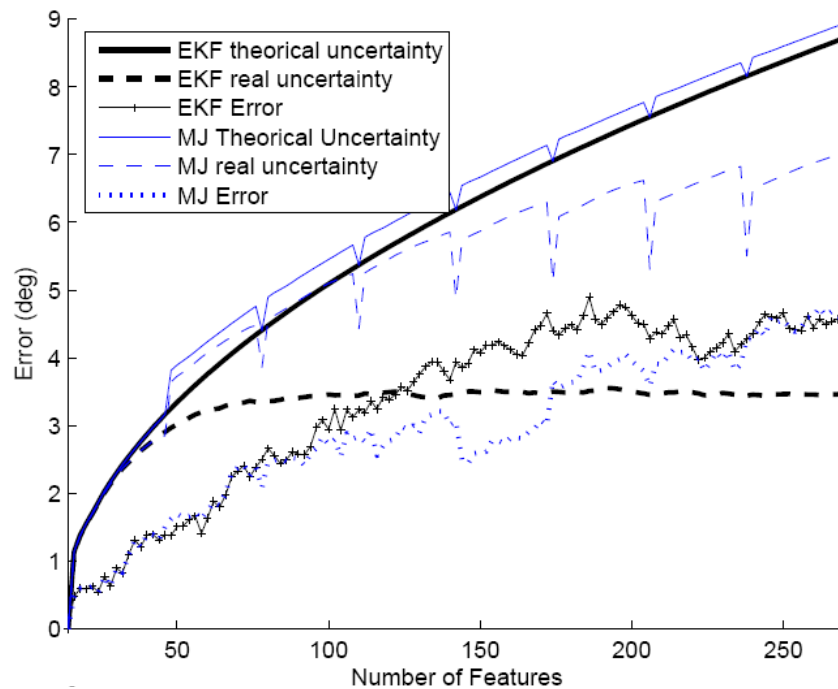
- One full SLAM run
- Map joining of 28 local maps



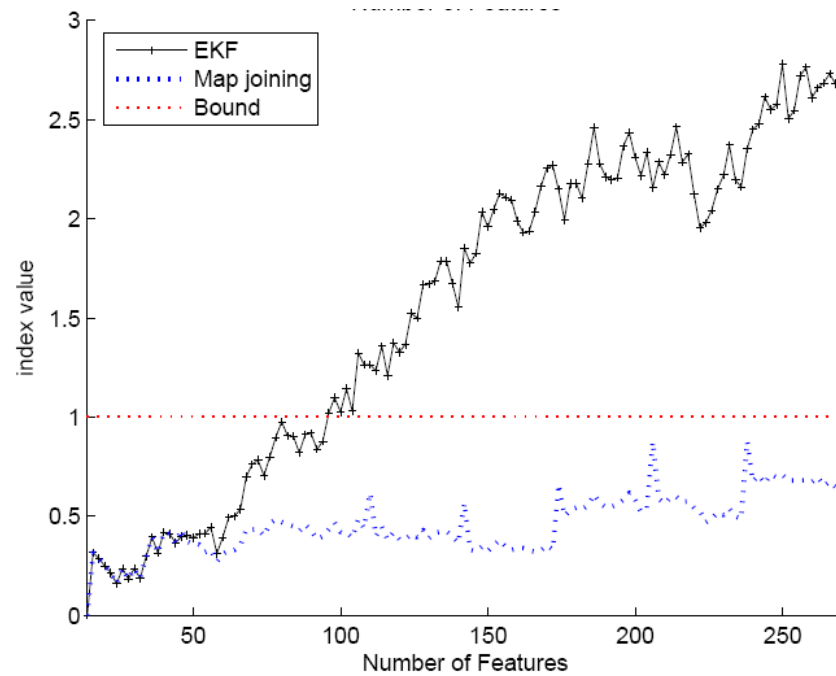
Local maps bound
linearization error effects

Map Joining

- Mean vehicle orientation error for full EKF and Map Joining 8 local maps

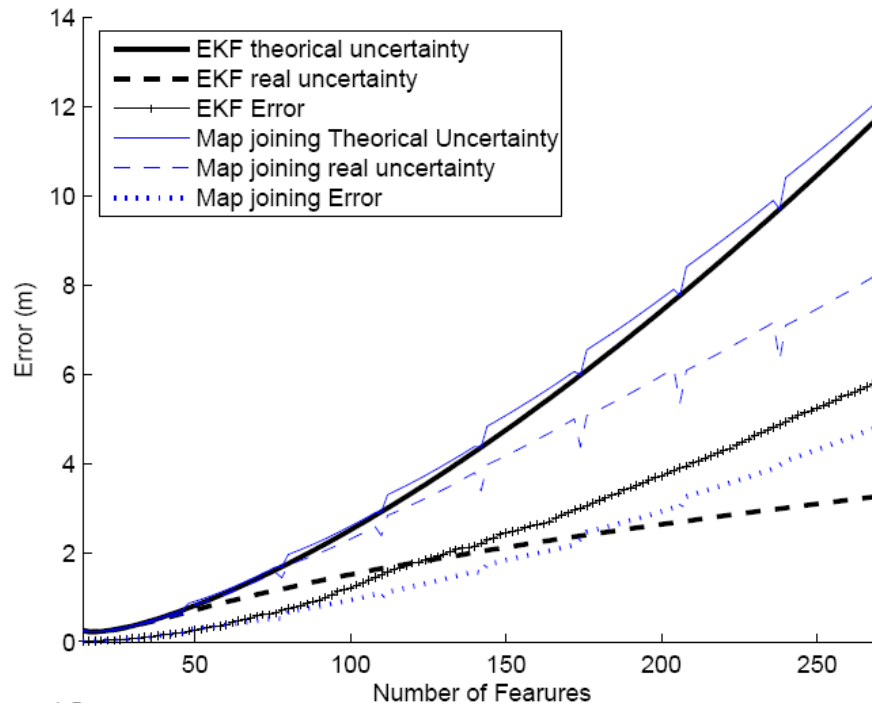


Vehicle orientation
error

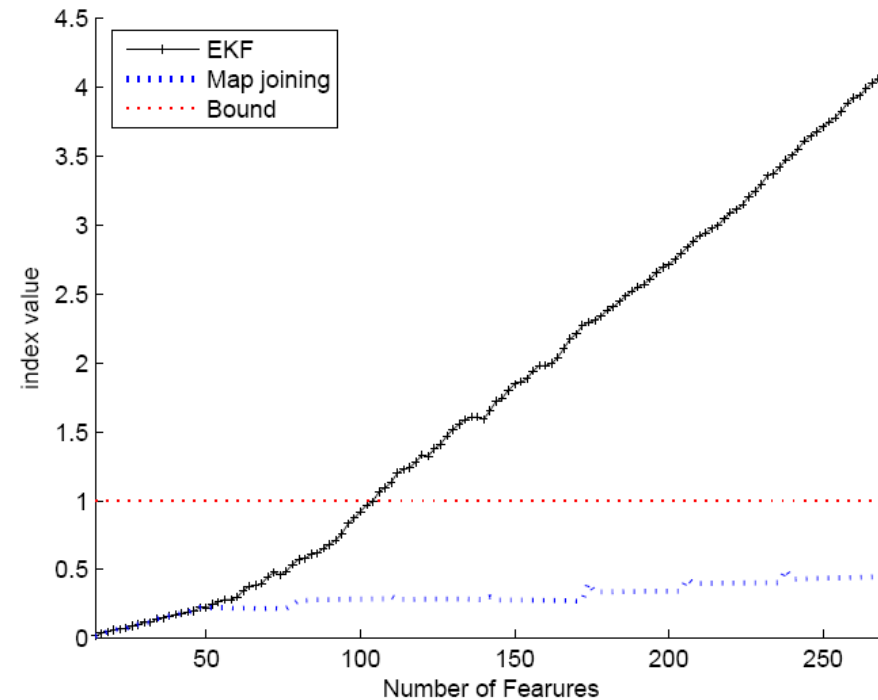


Vehicle orientation
consistency index

Map Joining



Mean feature
position error

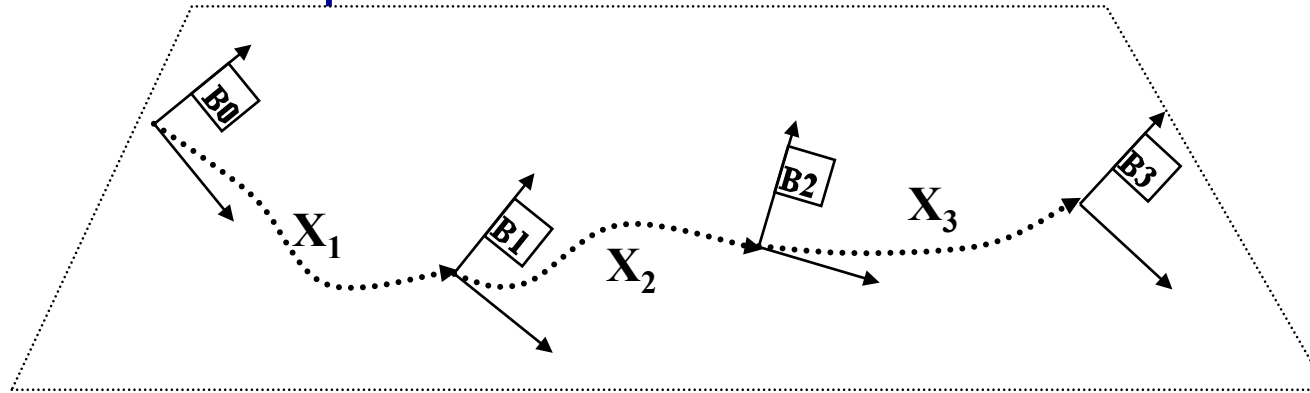


Mean feature
consistency index

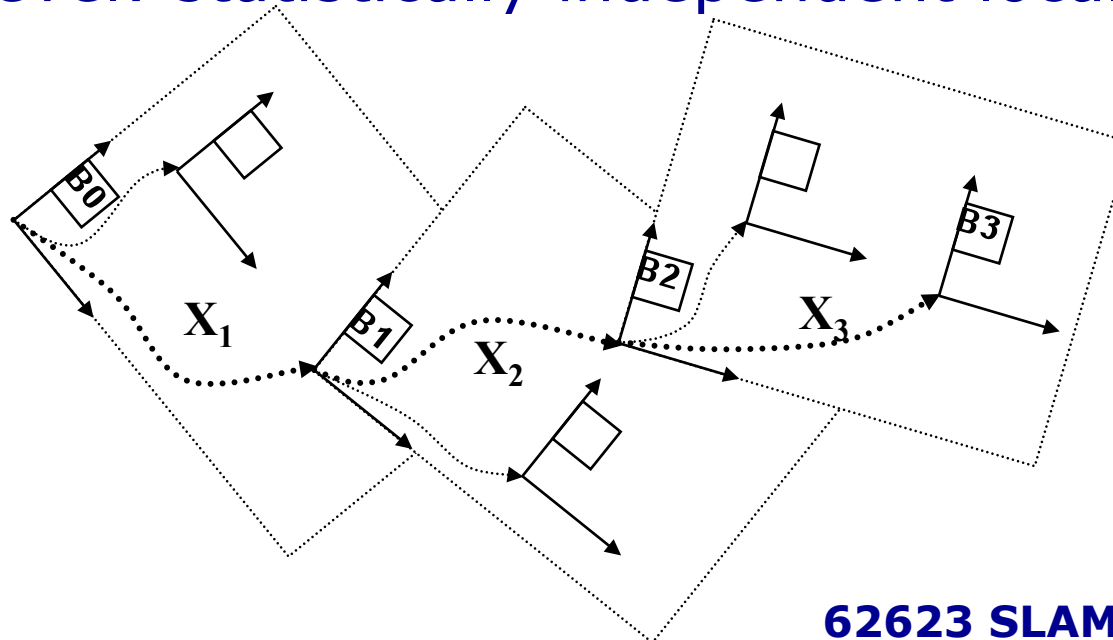
Hierarchical SLAM

Hierarchical SLAM

- Global level: adjacency graph and relative stochastic map



- Local level: statistically independent local maps



Hierarchical SLAM

- Local maps:

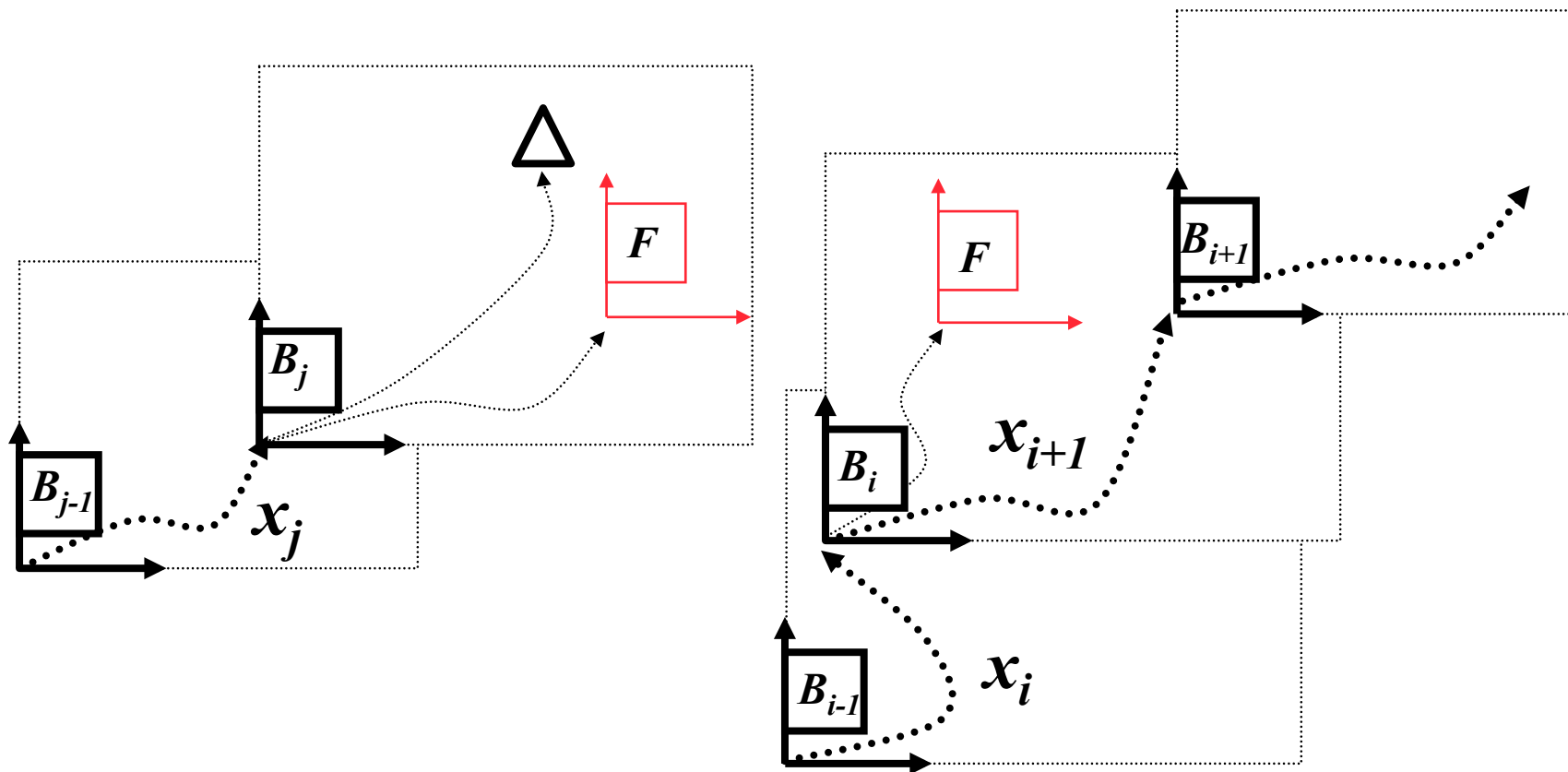
$$\hat{\mathbf{x}}_{\mathcal{F}}^B = \begin{bmatrix} \hat{\mathbf{x}}_R^B \\ \vdots \\ \hat{\mathbf{x}}_{F_n}^B \end{bmatrix}; \mathbf{P}_{\mathcal{F}} = \begin{bmatrix} \mathbf{P}_R^B & \cdots & \mathbf{P}_{RF_n}^B \\ \vdots & \ddots & \vdots \\ \mathbf{P}_{F_n R}^B & \cdots & \mathbf{P}_{F_n F_n}^B \end{bmatrix}$$

- Global relative map:

$$\hat{\mathbf{x}} = \begin{bmatrix} \vdots \\ \hat{\mathbf{x}}_i \\ \vdots \end{bmatrix}; \mathbf{P} = \begin{bmatrix} . & 0 & 0 \\ 0 & \mathbf{P}_i & 0 \\ 0 & 0 & . \end{bmatrix}$$

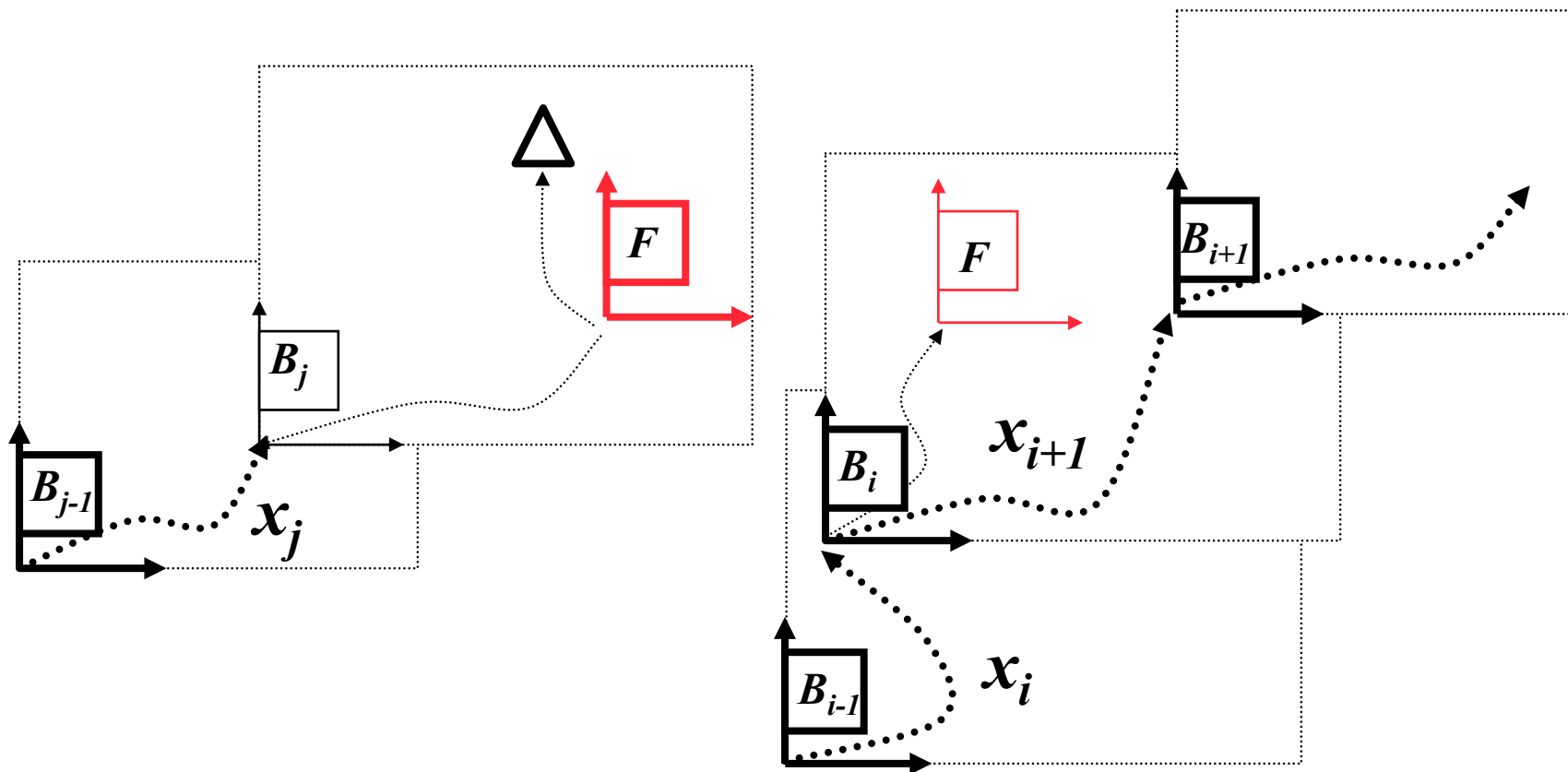
Block diagonal

Loop closing



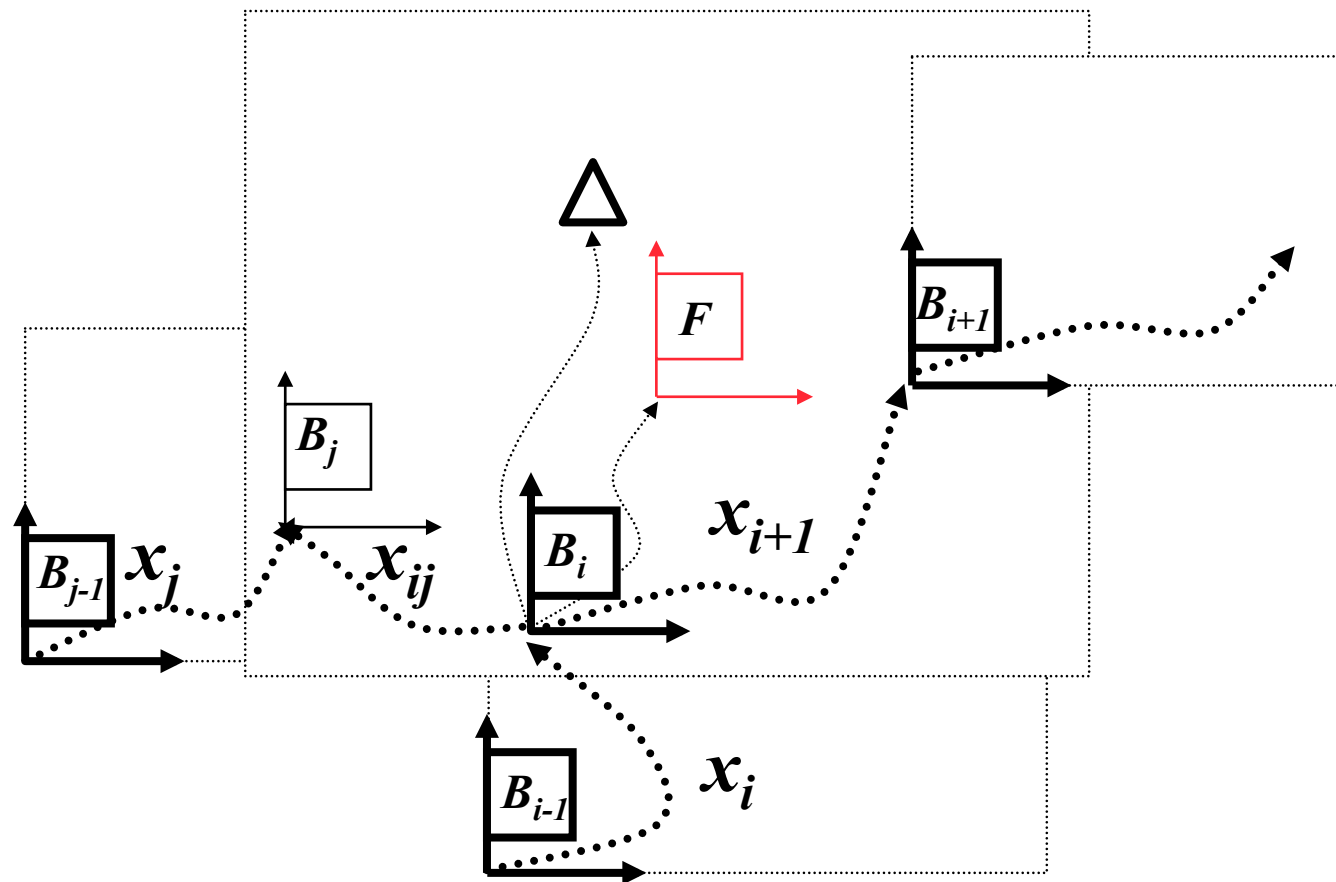
Find a common reference
between maps i and j

Loop closing



Change base reference of
map i to F .

Loop closing



Join maps i and j .

Hierarchical SLAM

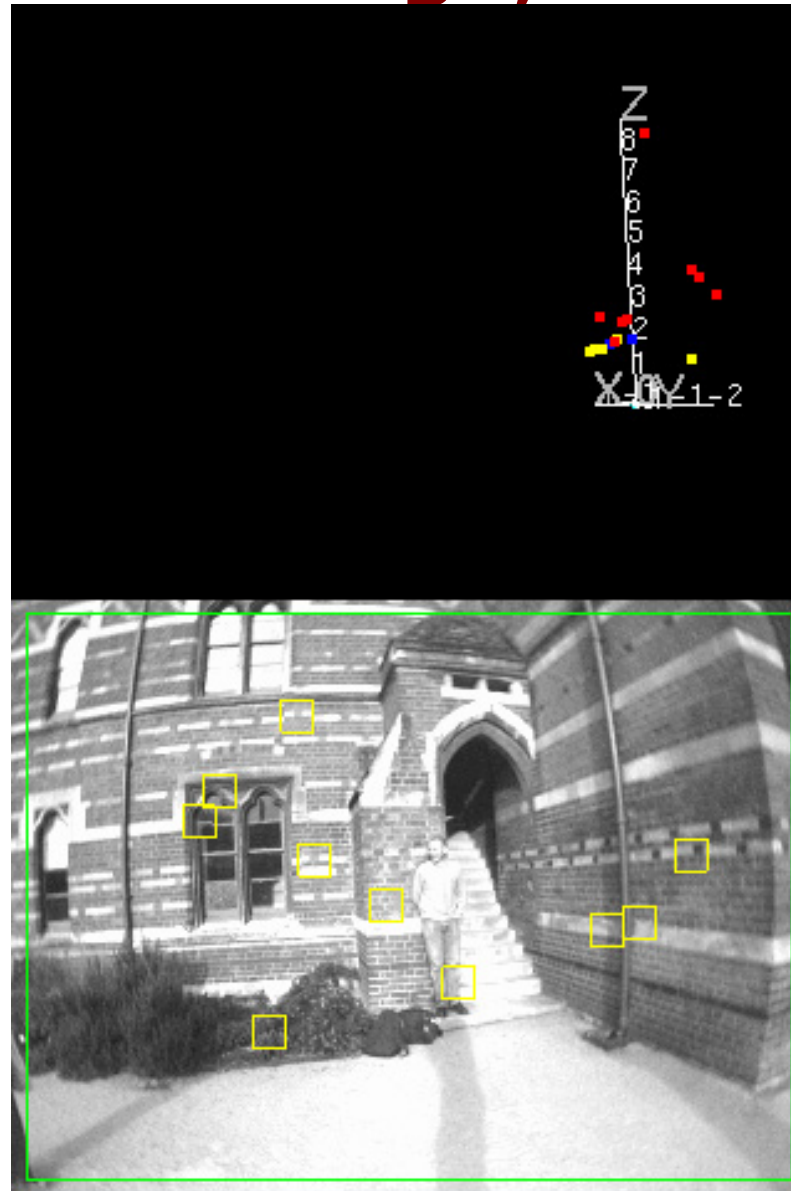
- Global relative map before loop closing:

$$\hat{\mathbf{x}} = \begin{bmatrix} \vdots \\ \hat{\mathbf{x}}_{i+1} \\ \vdots \\ \hat{\mathbf{x}}_j \end{bmatrix}; \mathbf{P} = \begin{bmatrix} . & 0 & 0 & 0 \\ 0 & \mathbf{P}_{i+1} & 0 & 0 \\ 0 & 0 & . & 0 \\ 0 & 0 & 0 & \mathbf{P}_j \end{bmatrix}$$

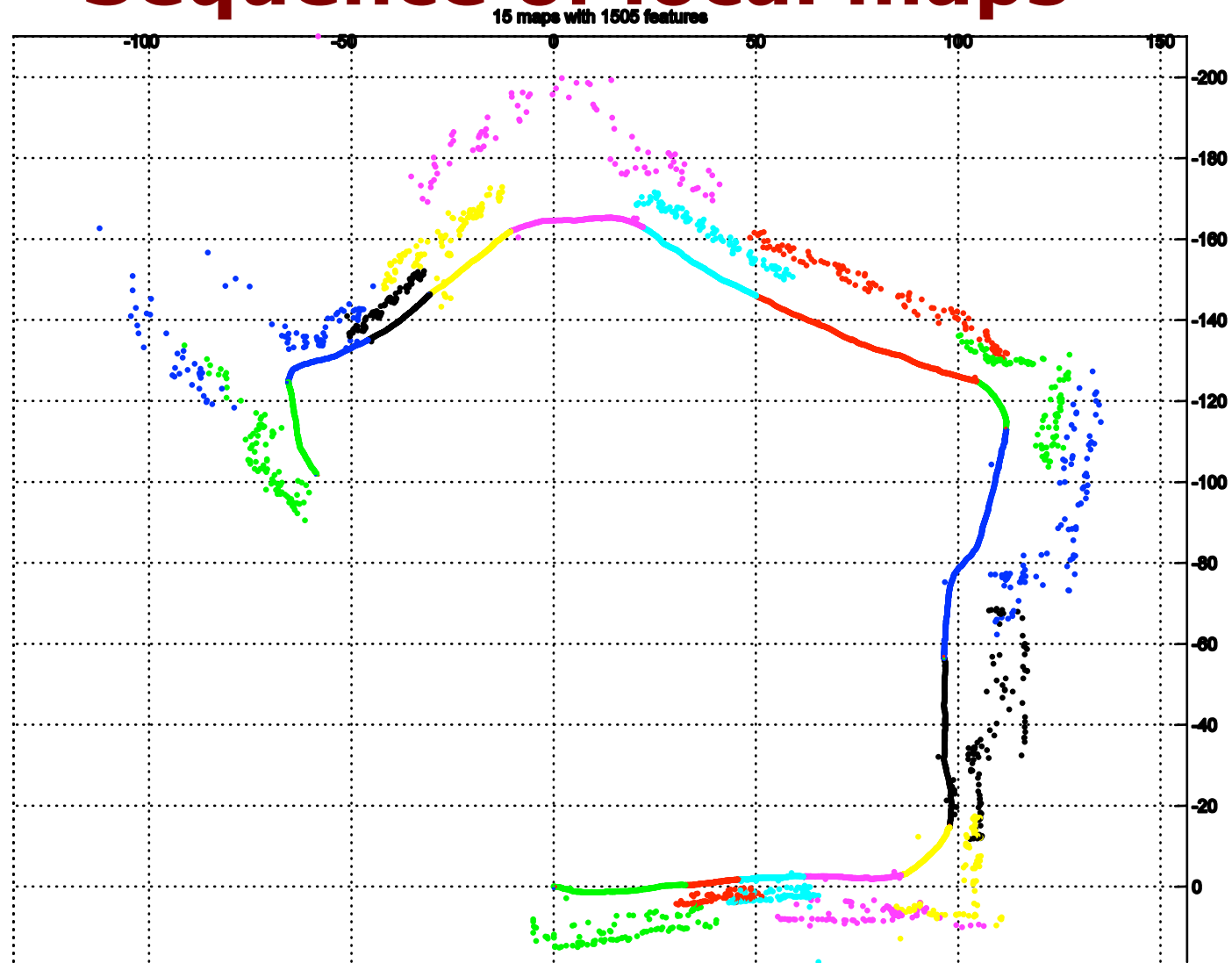
- After loop closing:

$$\hat{\mathbf{x}} = \begin{bmatrix} \vdots \\ \hat{\mathbf{x}}_{i+1} \\ \hat{\mathbf{x}}_{ij} \\ \vdots \\ \hat{\mathbf{x}}_j \end{bmatrix}; \mathbf{P} = \begin{bmatrix} . & 0 & 0 & 0 & 0 \\ 0 & \mathbf{P}_{i+1} & \mathbf{P}_{i+1,ij} & 0 & 0 \\ 0 & \mathbf{P}_{i+1,ij}^T & \mathbf{P}_{ij} & 0 & 0 \\ 0 & 0 & 0 & . & 0 \\ 0 & 0 & 0 & 0 & \mathbf{P}_j \end{bmatrix}$$

Keble College, Oxford

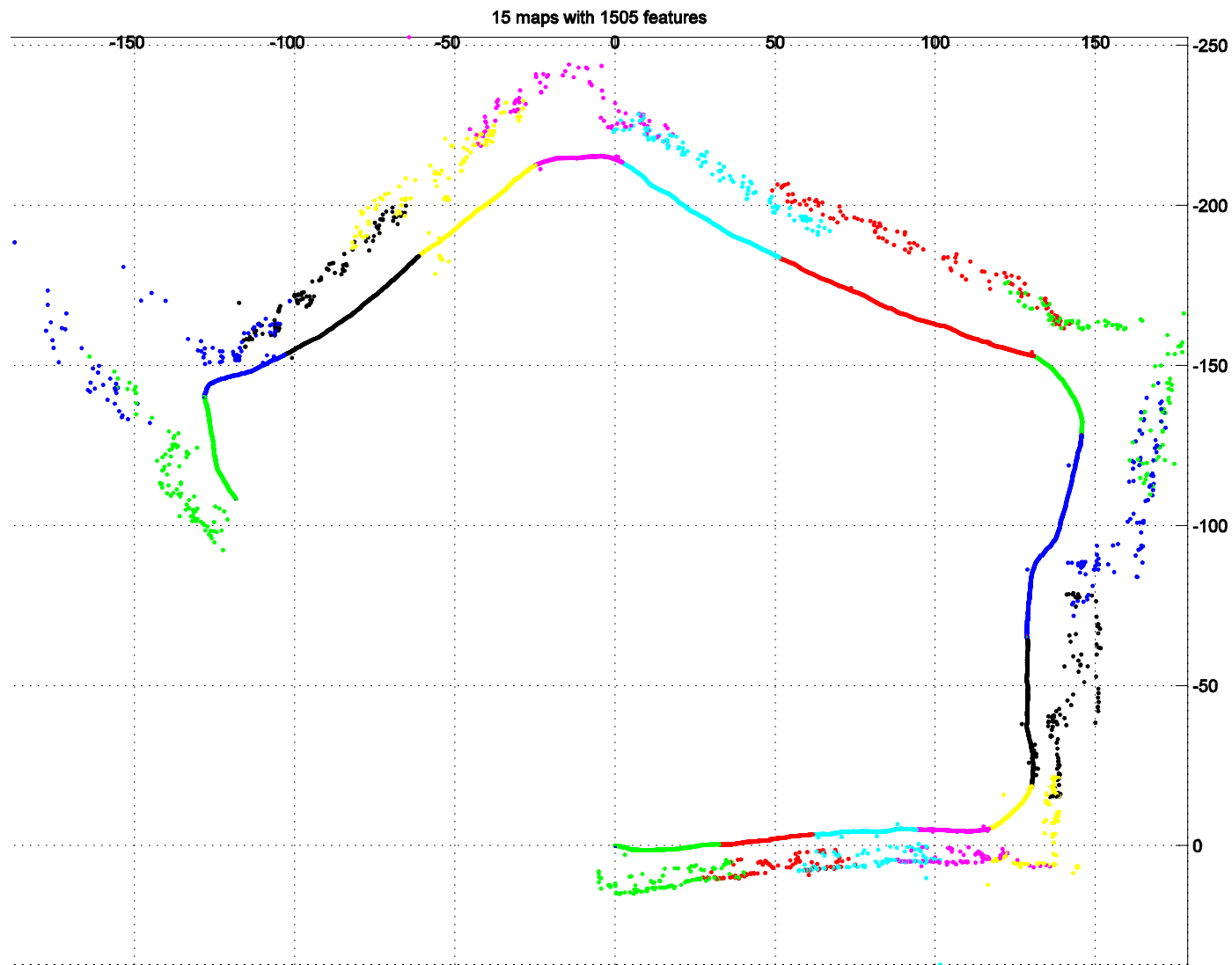


Sequence of local maps



The scale is arbitrary (not observable)

With scale compensation



Map Matching

Unary Constraints

Cloud 1

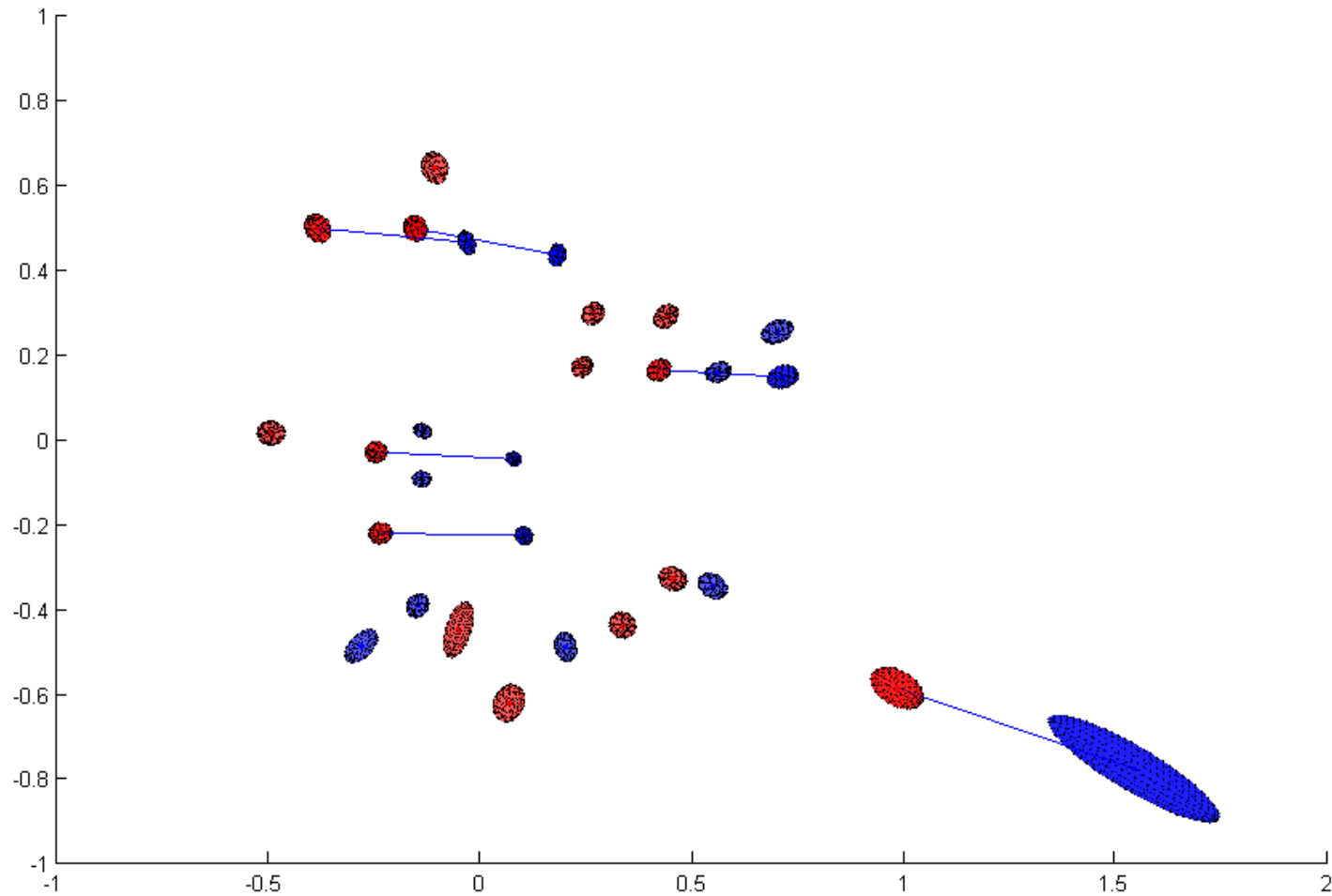


Cloud 2

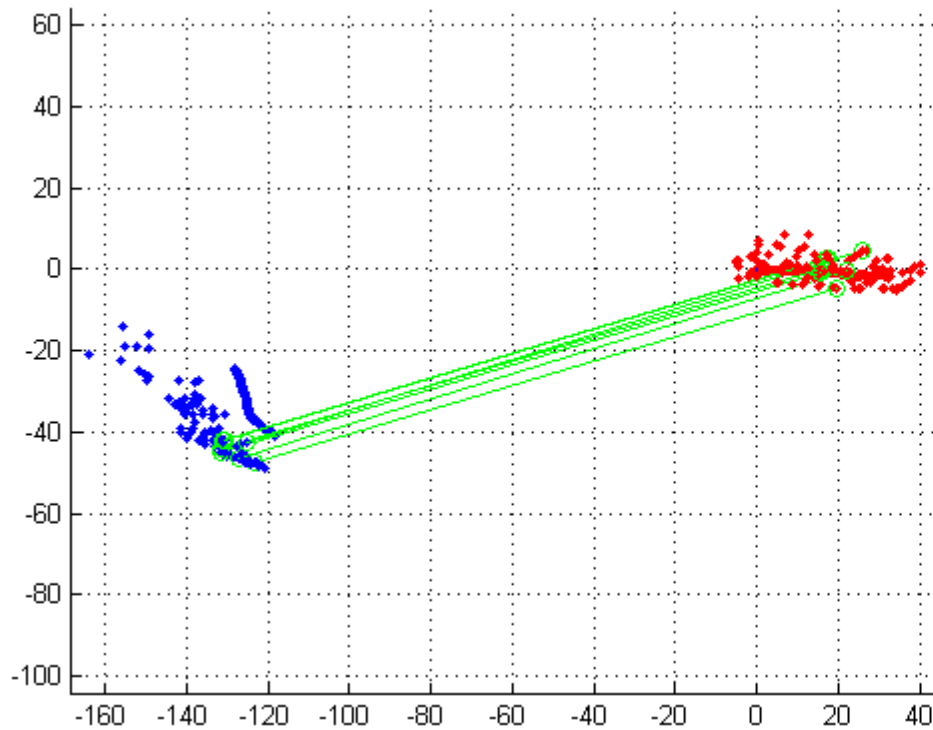


Map Matching

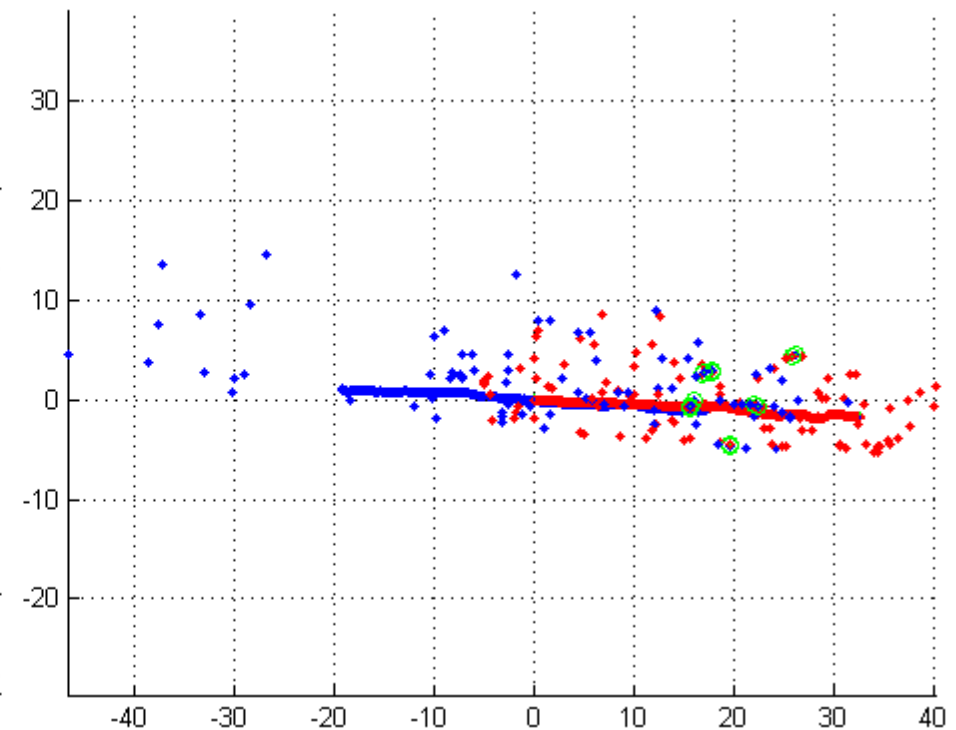
Binary Constraints



Map-to-Map Matching



Matchings found

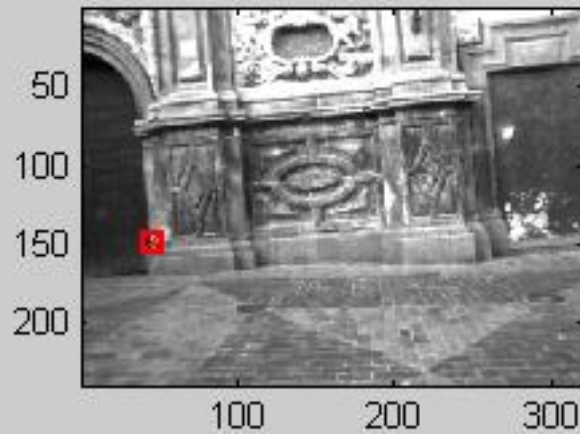


Aligned Submaps

Map Matching

LOOP CLOSING

FIRST MAP



LAST MAP



Nonlinear constrained optimization

- Minimize corrections to the global map, subject to the loop constraint:

$$\min_{\mathbf{x}} \frac{1}{2}(\mathbf{x} - \hat{\mathbf{x}})^T \mathbf{P}^{-1}(\mathbf{x} - \hat{\mathbf{x}})$$
$$\mathbf{h}(\mathbf{x}) = 0$$

- Sequential Quadratic Programming (SQP) :

$$\mathbf{H}_i = \left[\left. \frac{\partial \mathbf{h}}{\partial \mathbf{x}_1} \right|_{\hat{\mathbf{x}}_i} \left. \frac{\partial \mathbf{h}}{\partial \mathbf{x}_2} \right|_{\hat{\mathbf{x}}_i} \cdots \left. \frac{\partial \mathbf{h}}{\partial \mathbf{x}_{n-1}} \right|_{\hat{\mathbf{x}}_i} \left. \frac{\partial \mathbf{h}}{\partial \mathbf{x}_n} \right|_{\hat{\mathbf{x}}_i} \right]$$

$$\mathbf{P}_i = \mathbf{P}_0 - \mathbf{P}_0 \mathbf{H}_i^T \left(\mathbf{H}_i \mathbf{P}_0 \mathbf{H}_i^T \right)^{-1} \mathbf{H}_i \mathbf{P}_0$$

$$\hat{\mathbf{x}}_{i+1} = \hat{\mathbf{x}}_i - \mathbf{P}_i \mathbf{P}_0^{-1} (\hat{\mathbf{x}}_i - \hat{\mathbf{x}}_0) - \mathbf{P}_0 \mathbf{H}_i^T \left(\mathbf{H}_i \mathbf{P}_0 \mathbf{H}_i^T \right)^{-1} \hat{\mathbf{h}}_i$$

» Iterate until convergence

Nonlinear constrained optimization

- A more efficient version:

$$\hat{\mathbf{x}}_{i+1} = \hat{\mathbf{x}}_0 + \mathbf{P}_0 \mathbf{H}_i^T \left(\mathbf{H}_i \mathbf{P}_0 \mathbf{H}_i^T \right)^{-1} \left(\mathbf{H}_i (\hat{\mathbf{x}}_i - \hat{\mathbf{x}}_0) - \hat{\mathbf{h}}_i \right)$$

» Iterate until convergence

- Complexity:
 - $\mathbf{P}_0$ is block diagonal
 - $\mathbf{H}_i$ is sparse with nonzeros only for the maps in the loop
 - The iteration is linear with the number of maps in the loop
- Convergence:
 - Converges in 2 or 3 iterations (for loops around 300m)
 - For bigger errors, may it converge to a local minimum ??

Solution 2 (for EKF fans)

- We could impose the loop constraints using an imprecise measurement function:

$$\mathbf{z} = \mathbf{h}(\mathbf{x}) + \mathbf{w} = \mathbf{0}$$

- With Covariance:

$$\mathbf{P}_z = Cov(\mathbf{w}) = \begin{bmatrix} \mathbf{P}_{z_1} & \mathbf{0} & \cdots & \mathbf{0} \\ \mathbf{0} & \mathbf{P}_{z_2} & \cdots & \mathbf{0} \\ \vdots & \vdots & \ddots & \vdots \\ \mathbf{0} & \mathbf{0} & \cdots & \mathbf{P}_{z_l} \end{bmatrix}$$

- Estimated errors in closing the loops:

$$\mathbf{h}(\hat{\mathbf{x}}) = \begin{bmatrix} h_1(\hat{\mathbf{x}}) \\ h_2(\hat{\mathbf{x}}) \\ \vdots \\ h_l(\hat{\mathbf{x}}) \end{bmatrix}$$

Iterated Extended Kalman Filter

- Jacobian of the measurement function:

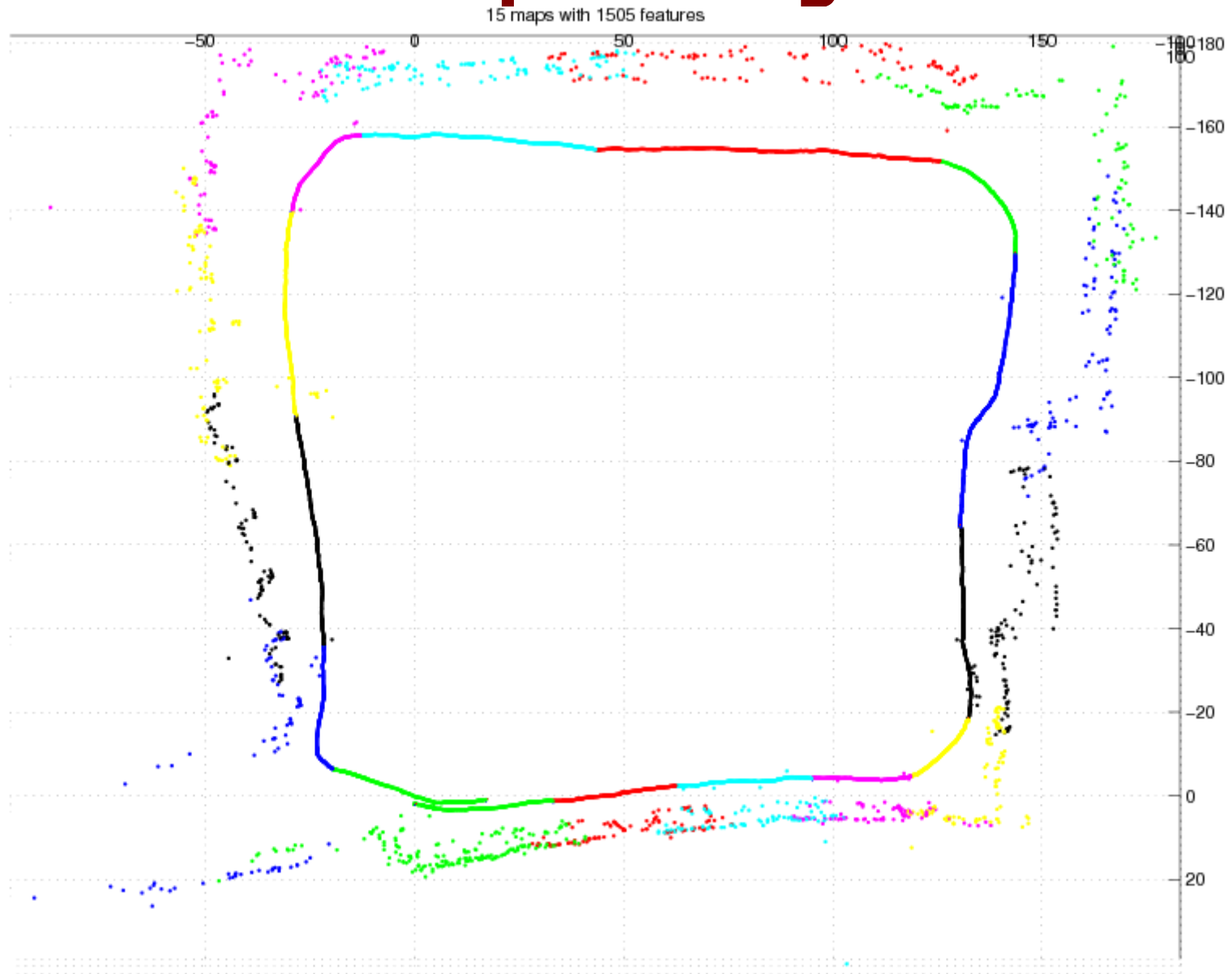
$$\mathbf{H}_i = \left[\left. \frac{\partial \mathbf{h}}{\partial \mathbf{x}_1} \right|_{\hat{\mathbf{x}}_i} \left. \frac{\partial \mathbf{h}}{\partial \mathbf{x}_2} \right|_{\hat{\mathbf{x}}_i} \cdots \left. \frac{\partial \mathbf{h}}{\partial \mathbf{x}_{n-1}} \right|_{\hat{\mathbf{x}}_i} \left. \frac{\partial \mathbf{h}}{\partial \mathbf{x}_n} \right|_{\hat{\mathbf{x}}_i} \right]$$

- Iterated EKF equations:

$$\begin{aligned} \mathbf{P}_i &= \mathbf{P}_0 - \mathbf{P}_0 \mathbf{H}_i^T \left[\mathbf{H}_i \mathbf{P}_0 \mathbf{H}_i^T + \mathbf{P}_z \right]^{-1} \mathbf{H}_i \mathbf{P}_0 \\ \hat{\mathbf{x}}_{i+1} &= \hat{\mathbf{x}}_i - \mathbf{P}_i \mathbf{P}_0^{-1} (\hat{\mathbf{x}}_i - \hat{\mathbf{x}}_0) + \mathbf{P}_0 \mathbf{H}_i^T \left(\mathbf{H}_i \mathbf{P}_0 \mathbf{H}_i^T + \mathbf{P}_z \right)^{-1} (\mathbf{z} - \hat{\mathbf{h}}_i) \end{aligned}$$

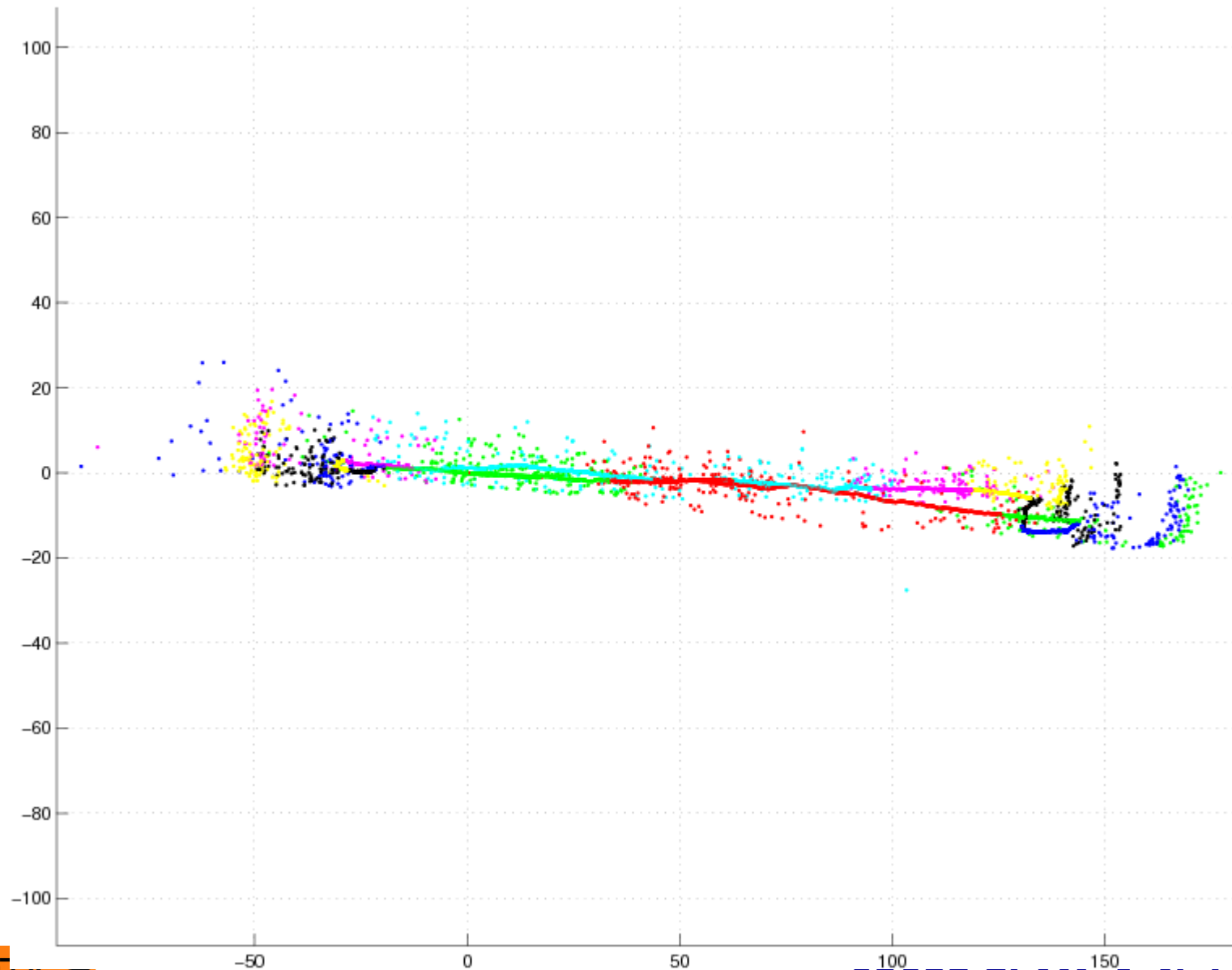
- With exact loop constraint, $\mathbf{z} = 0$ and $\mathbf{P}_z = 0$, IEKF is equivalent to nonlinear optimization with SQP

Loop closing



Loop closing (lateral view)

15 maps with 1505 features



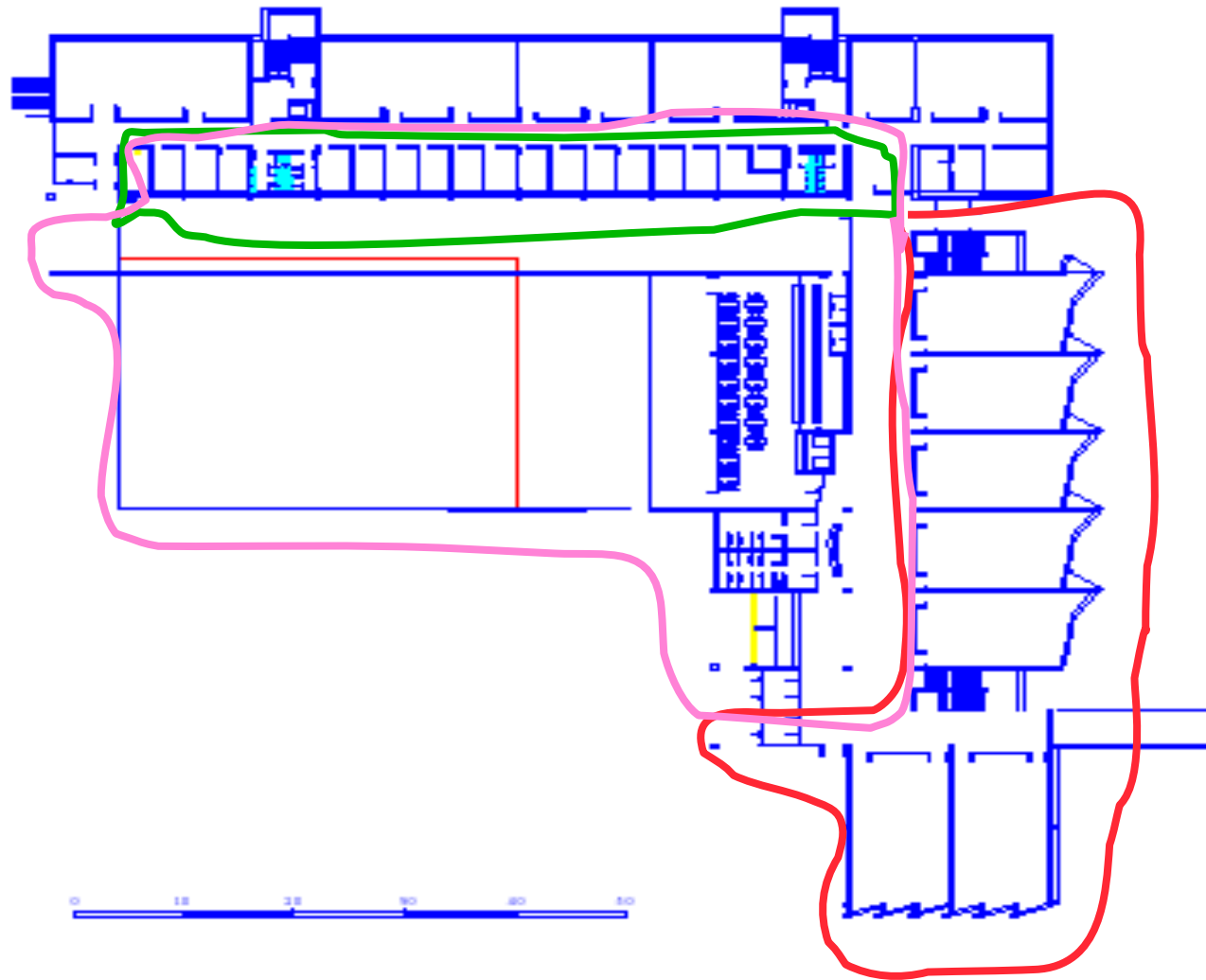
Keble College, Oxford (290m)



L. Clemente, A. Davison, I. Reid, J. Neira and J.D. Tardós **Mapping Large Loops with a Single Hand-Held Camera.** Robotics: Science and Systems, 2007.

Multivehicle SLAM

- Experiment



Multivehicle SLAM

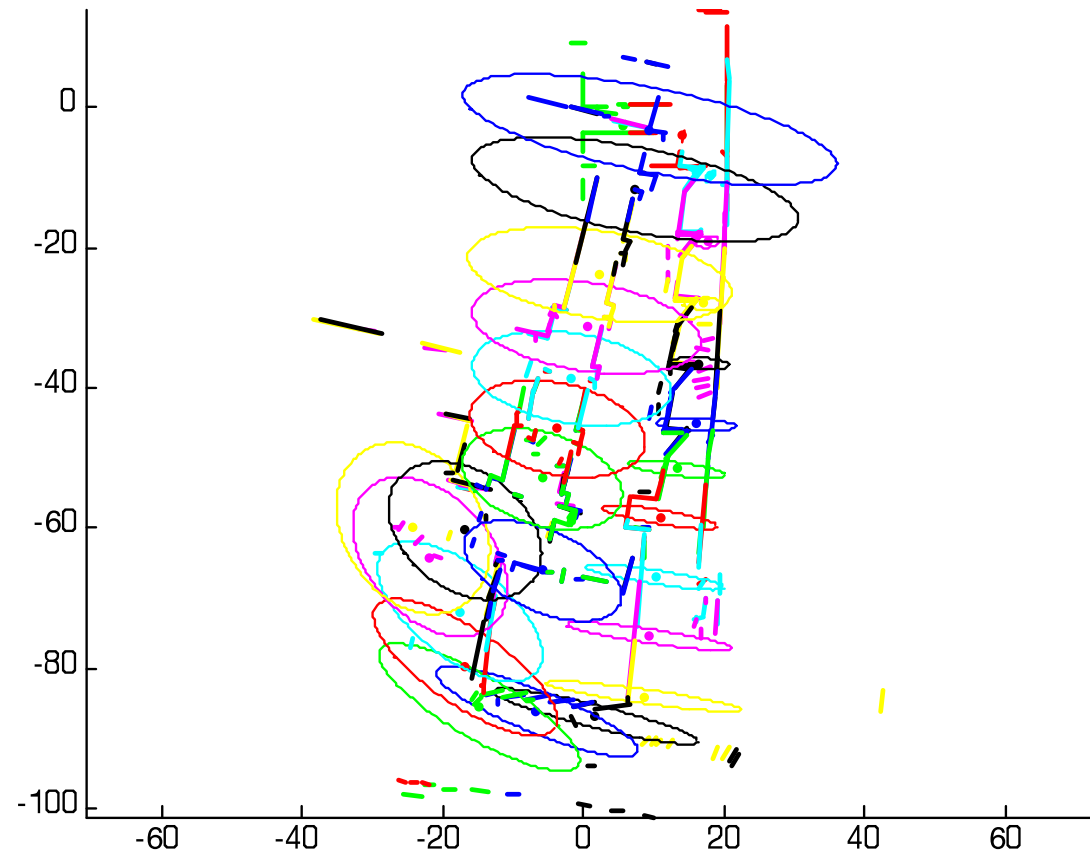
- Generalization to closing several loops simultaneously:

$$\mathbf{h}_j(\mathbf{x}) = \mathbf{x}_{j_1} \oplus \mathbf{x}_{j_2} \oplus \cdots \oplus \mathbf{x}_{j_{n_j-1}} \oplus \mathbf{x}_{j_{n_j}} = 0$$

$$\mathbf{h} = \begin{bmatrix} \mathbf{h}_1 \\ \mathbf{h}_2 \\ \vdots \\ \mathbf{h}_l \end{bmatrix}$$

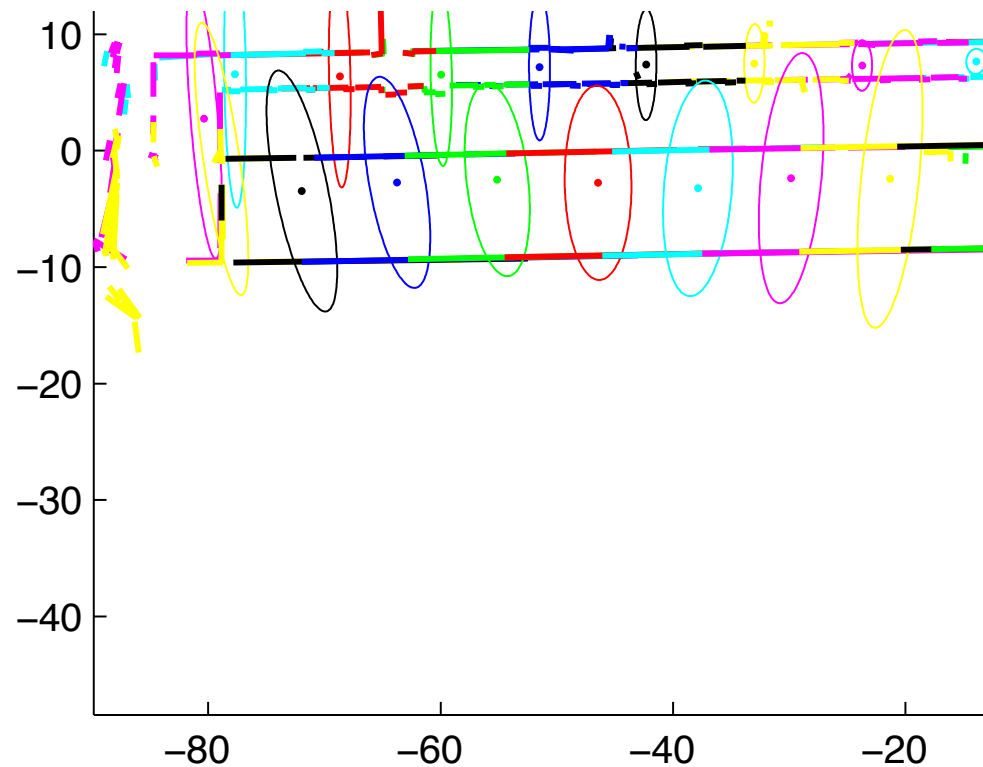
Experiments

- Local maps, first robot



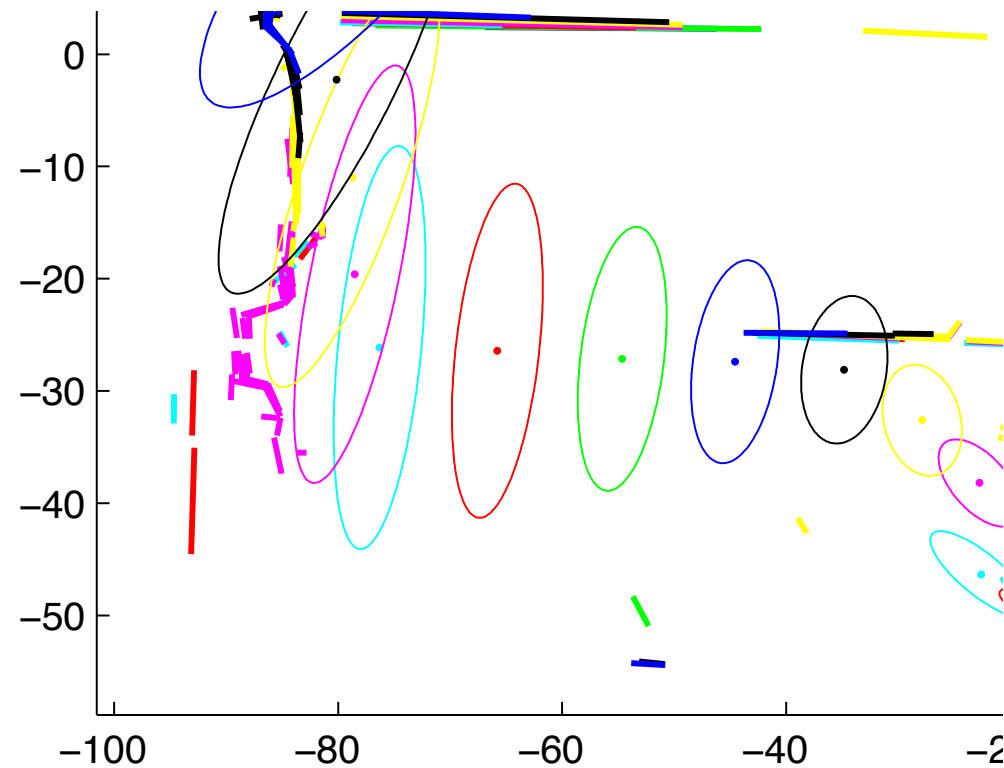
Experiments

- Local maps, second robot



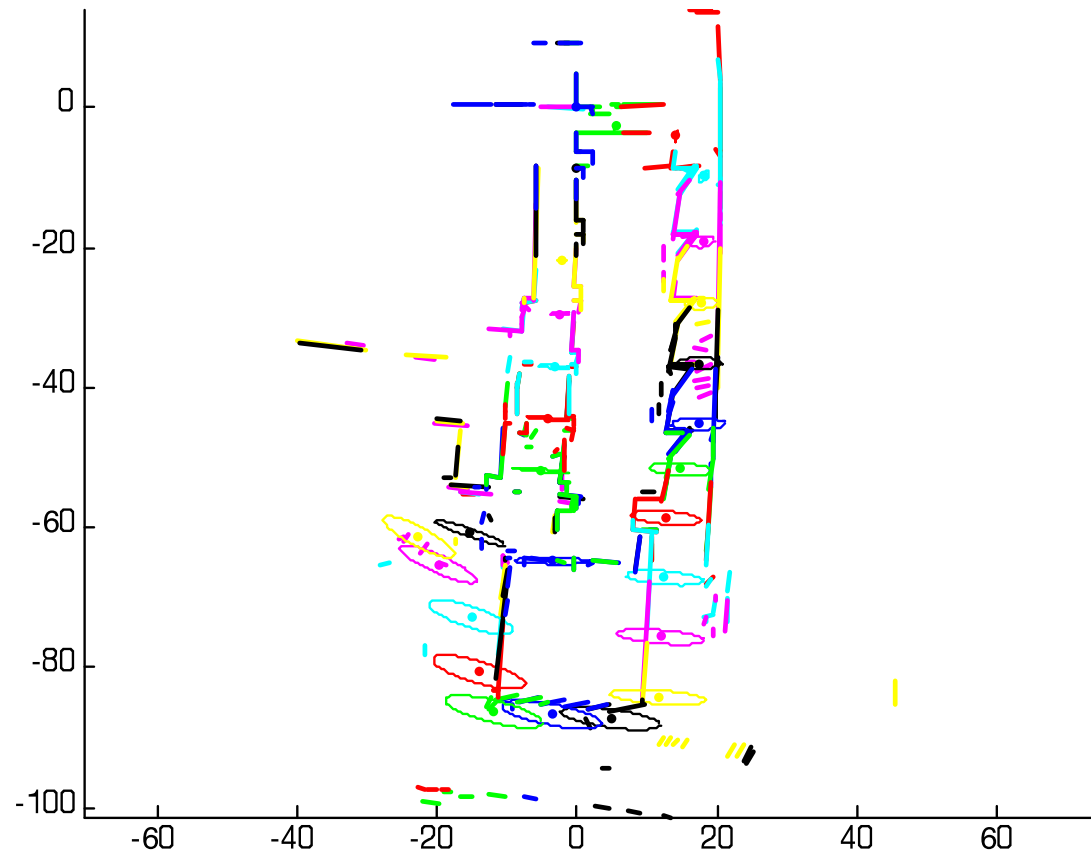
Experiments

- Local maps, third robot



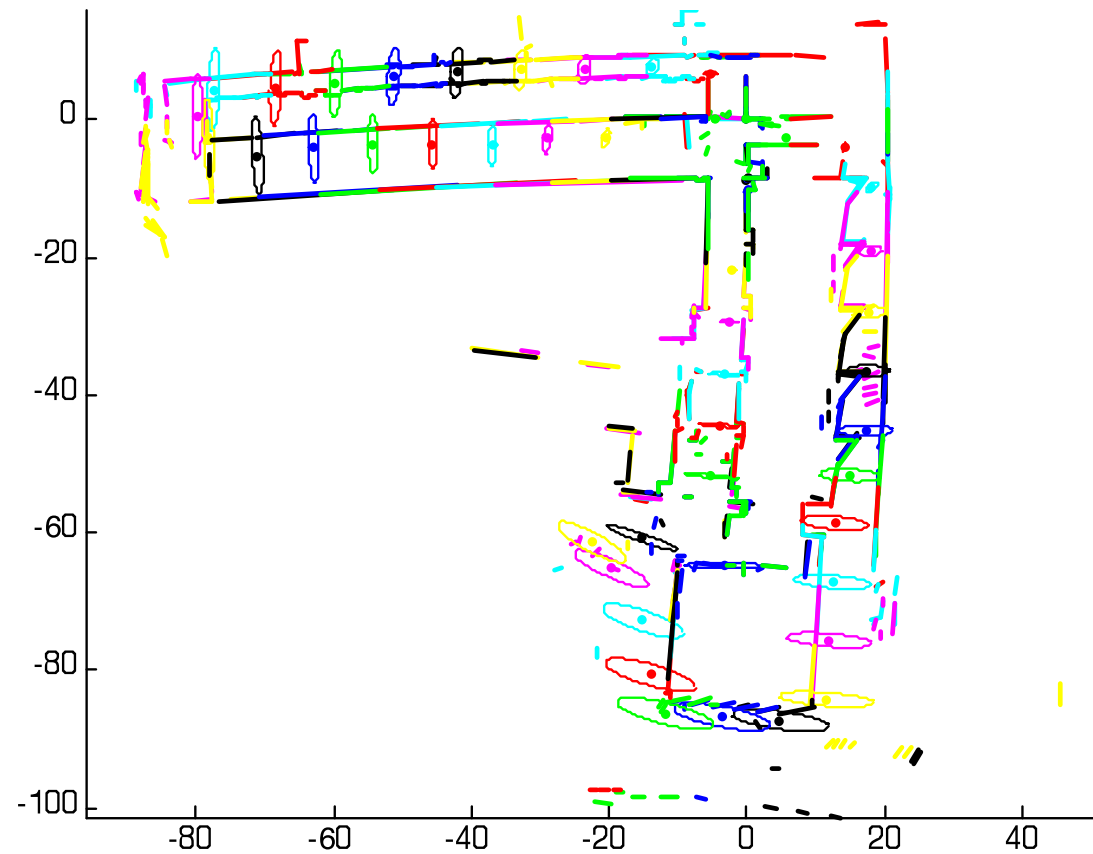
Hierarchical SLAM

- Imposing loop constraint:



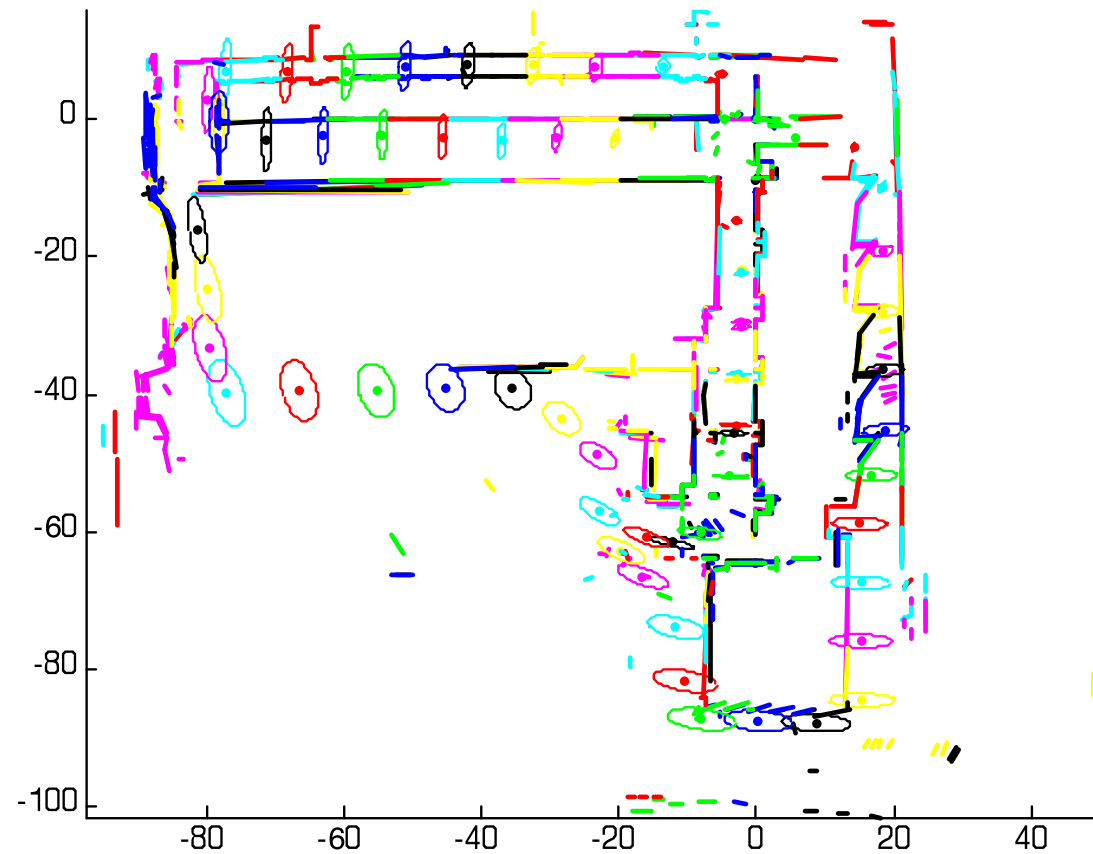
Hierarchical SLAM

- Second robot



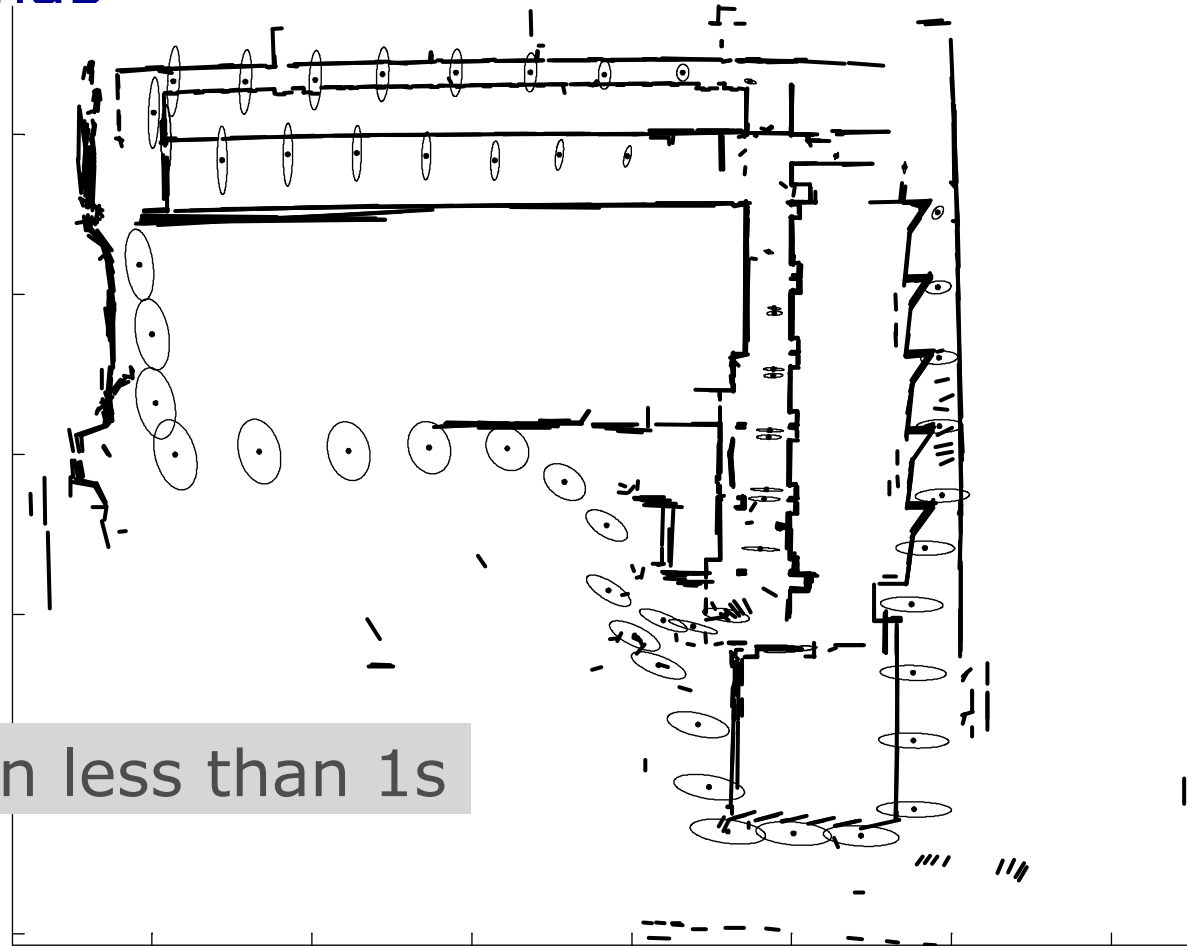
Hierarchical SLAM

- Third robot



Hierarchical SLAM

- Final map



Results in less than 1s

Divide and Conquer SLAM

- Experimental setup



**A bumblebee, a laptop
and a firewire cable**

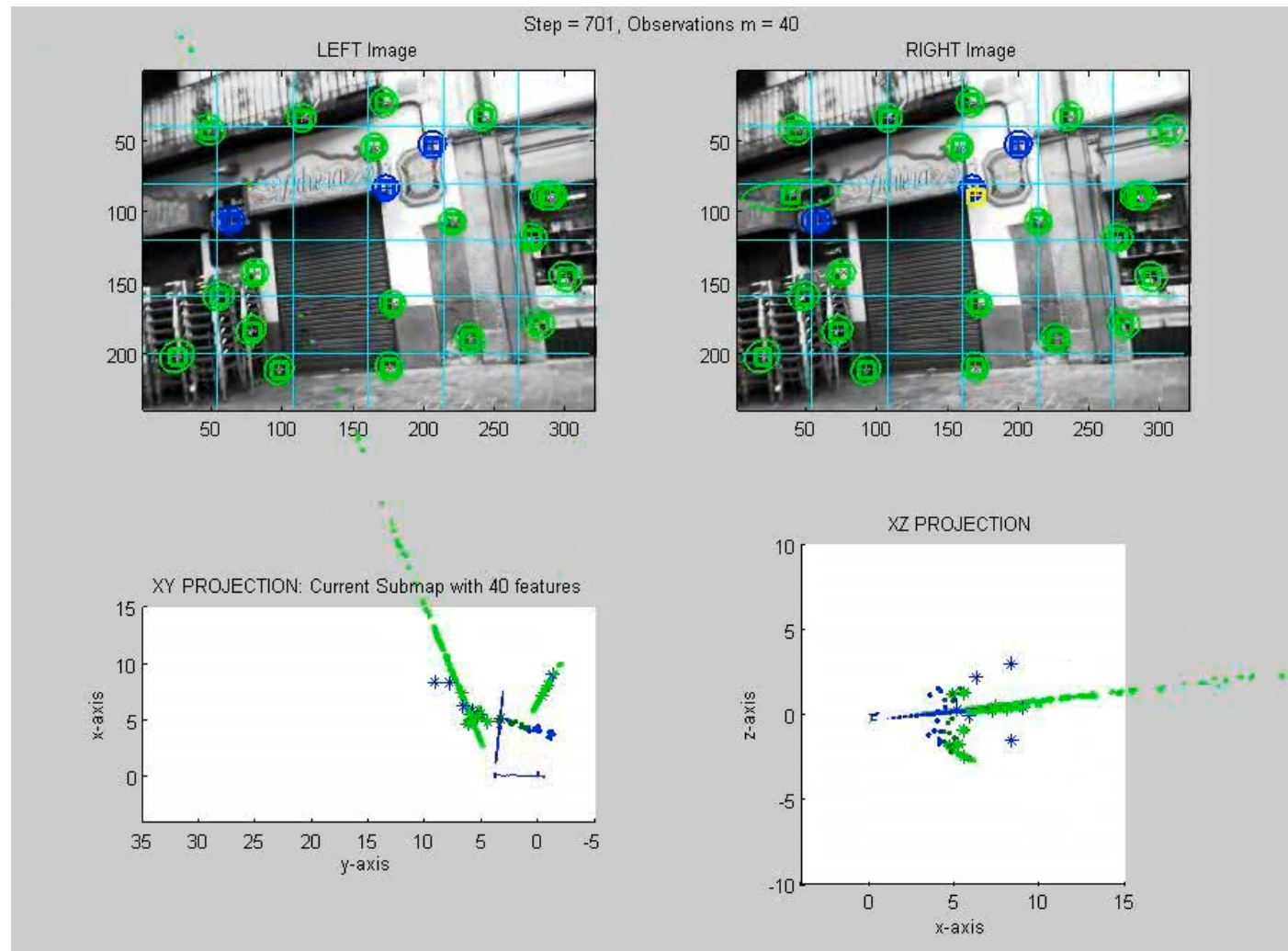
Pure Stereo SLAM



Pure Stereo SLAM

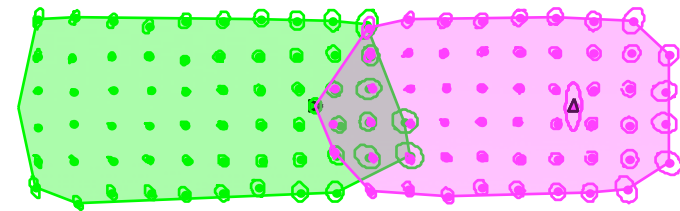
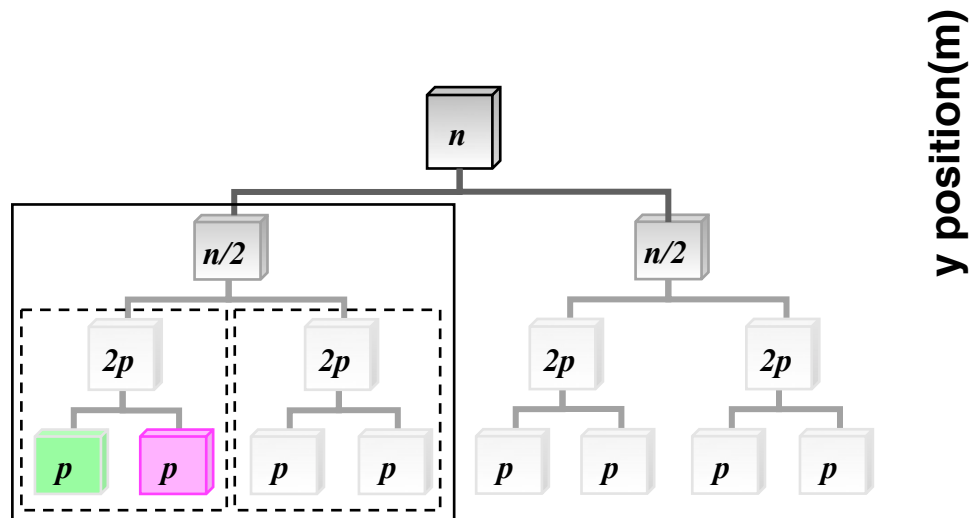


Basic EKF SLAM



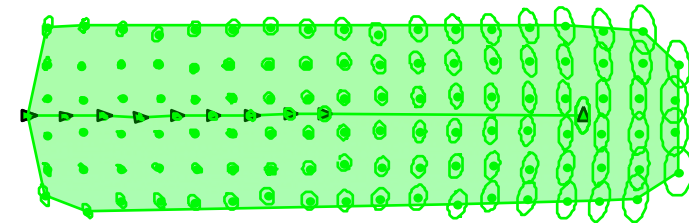
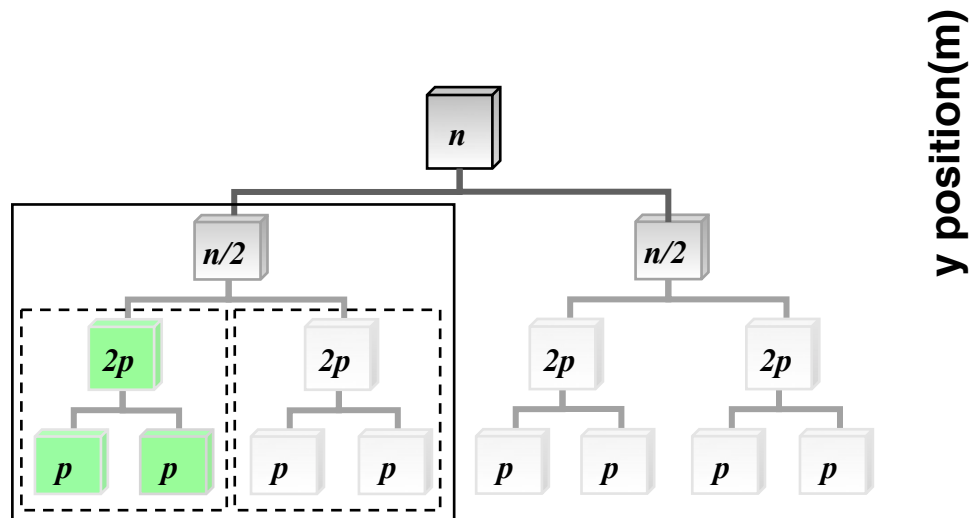
Divide & Conquer SLAM

Number of Maps : 2



Divide & Conquer SLAM

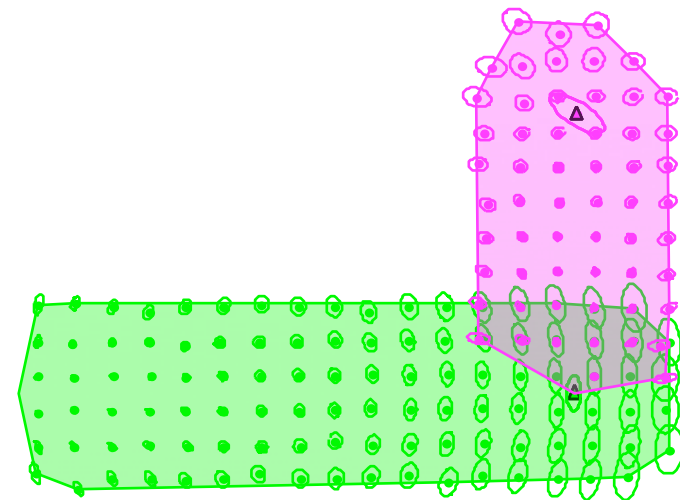
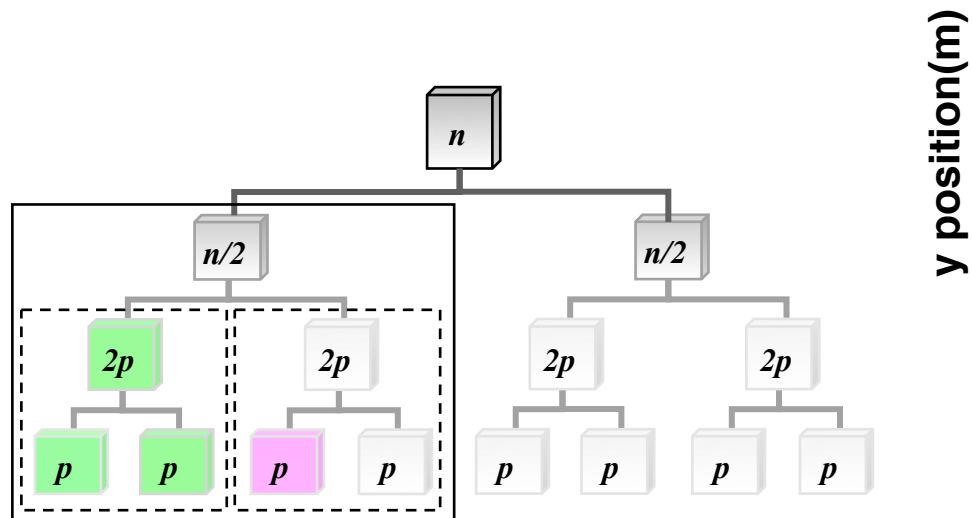
Number of Maps : 1



x position(m)

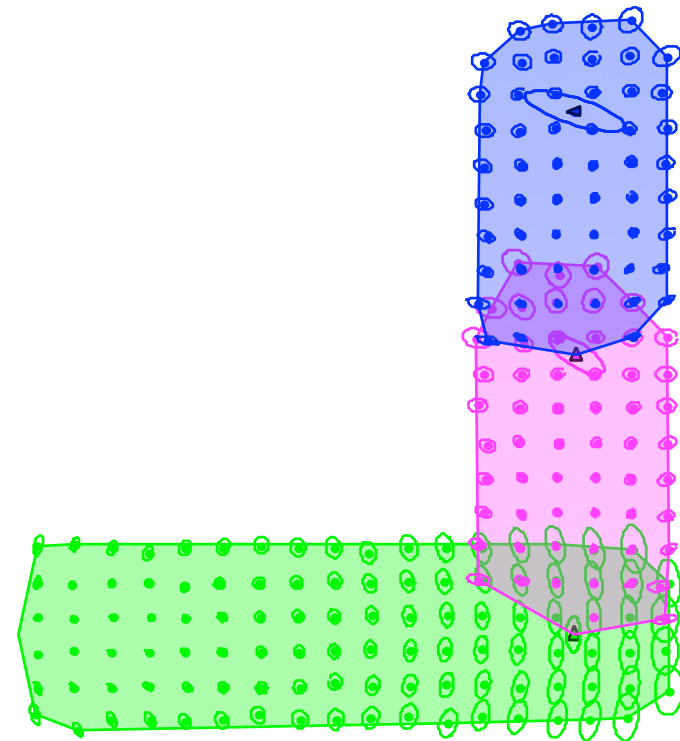
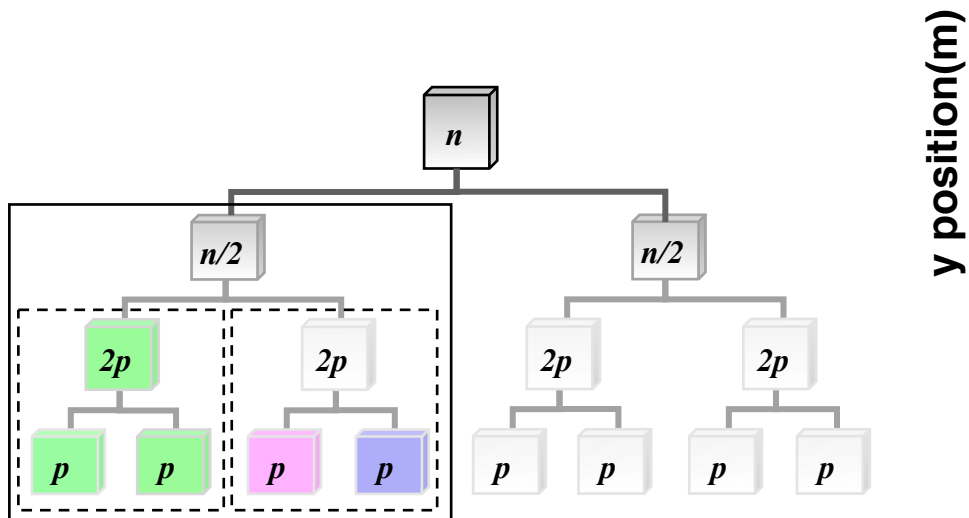
Divide & Conquer SLAM

Number of Maps : 2



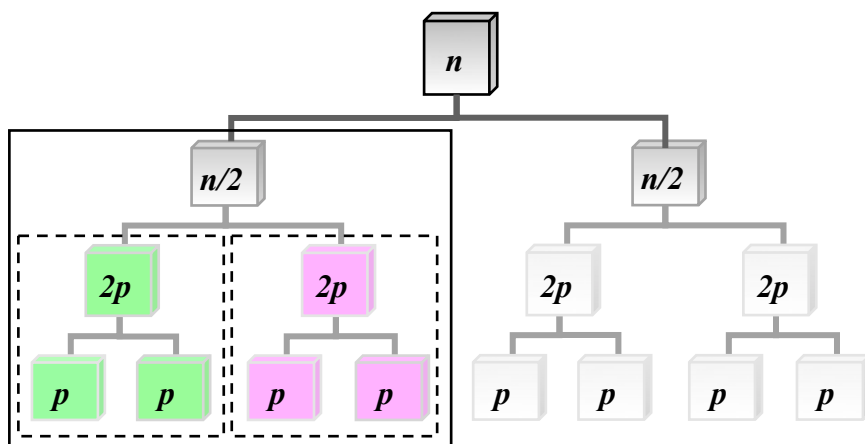
Divide & Conquer SLAM

Number of Maps : 3

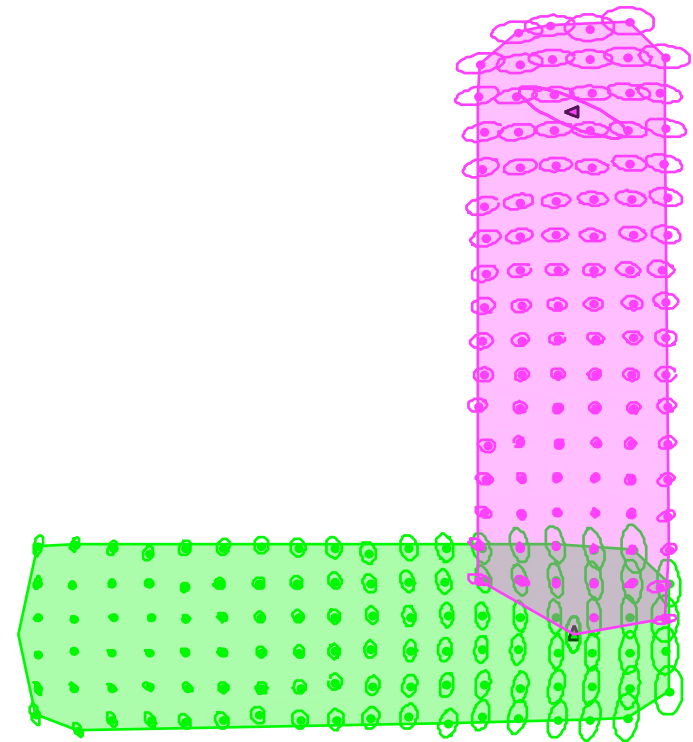


Divide & Conquer SLAM

Number of Maps : 2

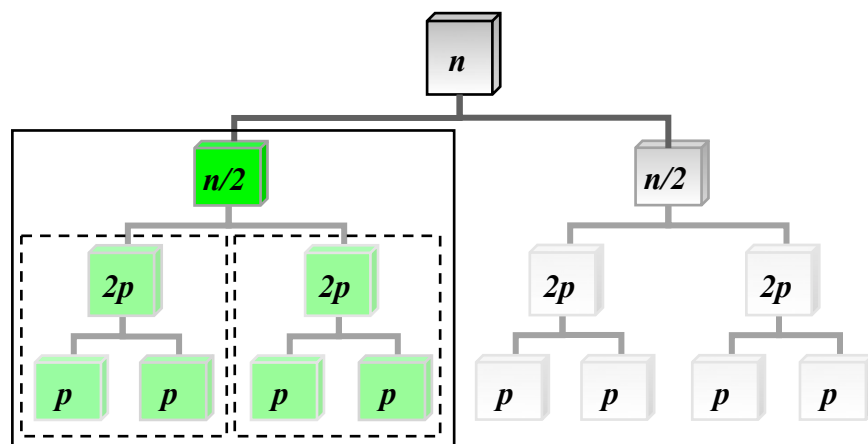


y position(m)



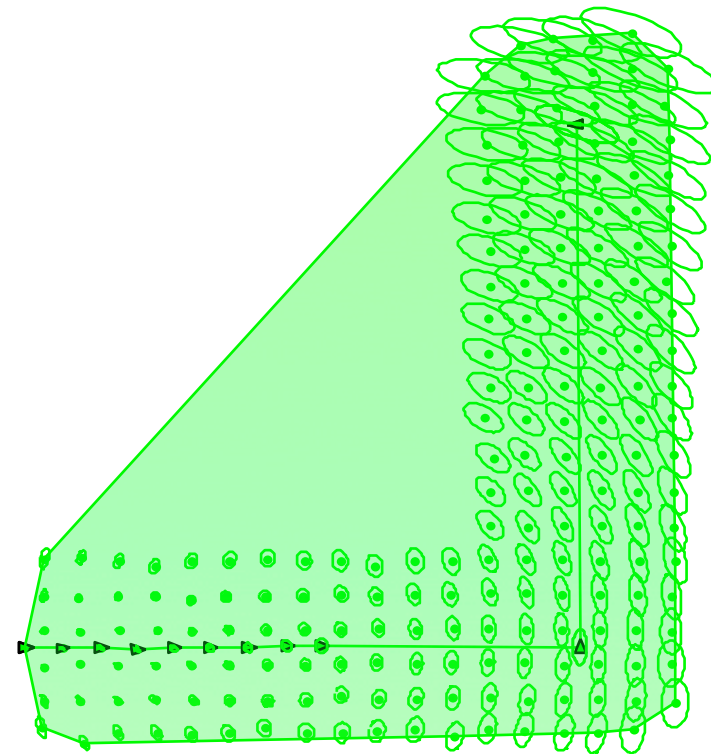
x position(m)

Divide & Conquer SLAM



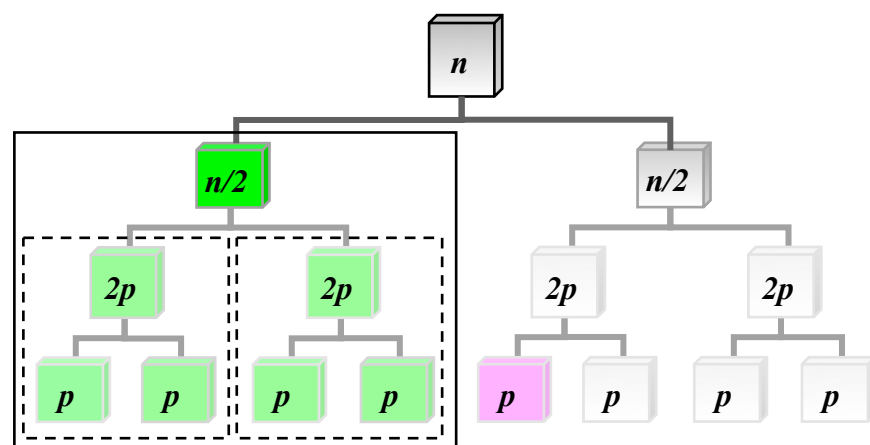
y position(m)

Number of Maps : 1



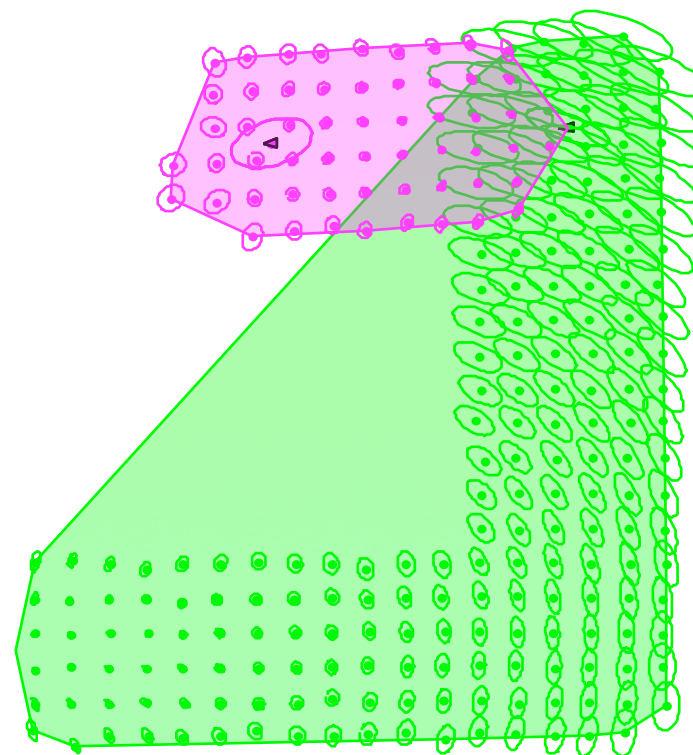
x position(m)

Divide & Conquer SLAM



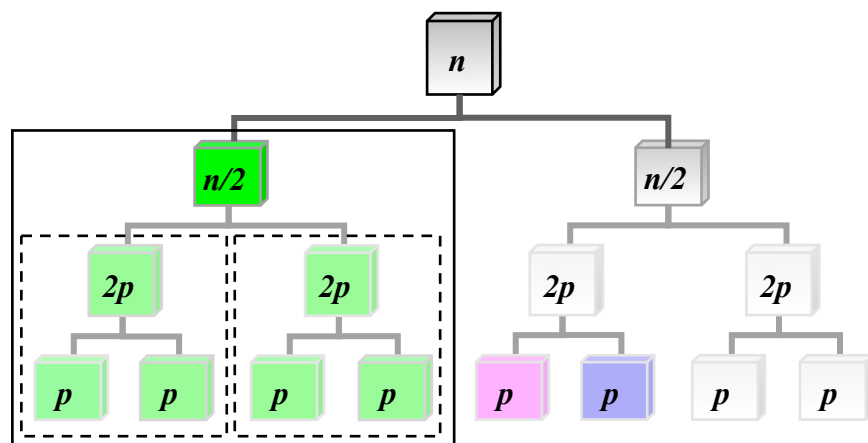
y position(m)

Number of Maps : 2



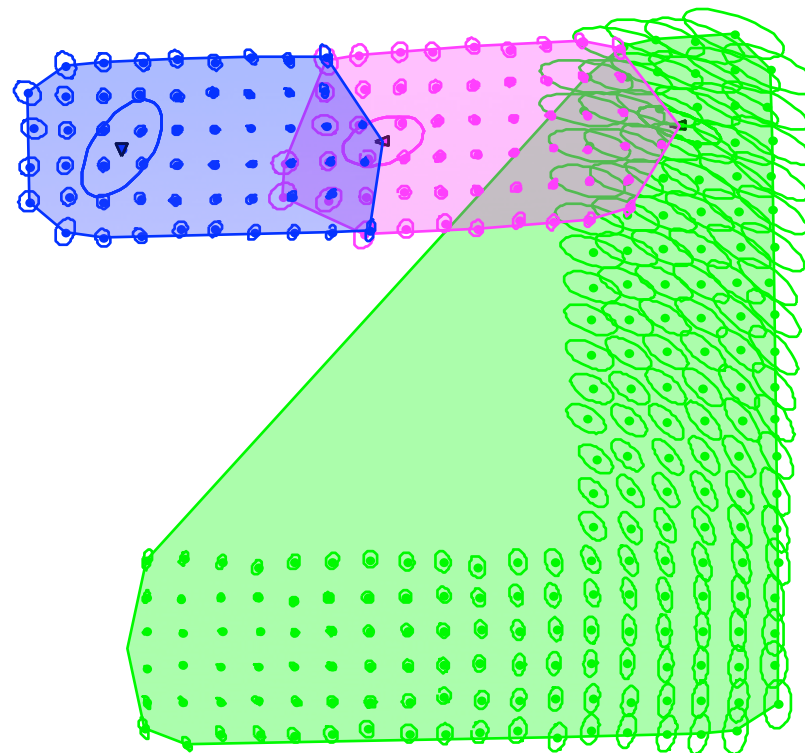
x position(m)

Divide & Conquer SLAM



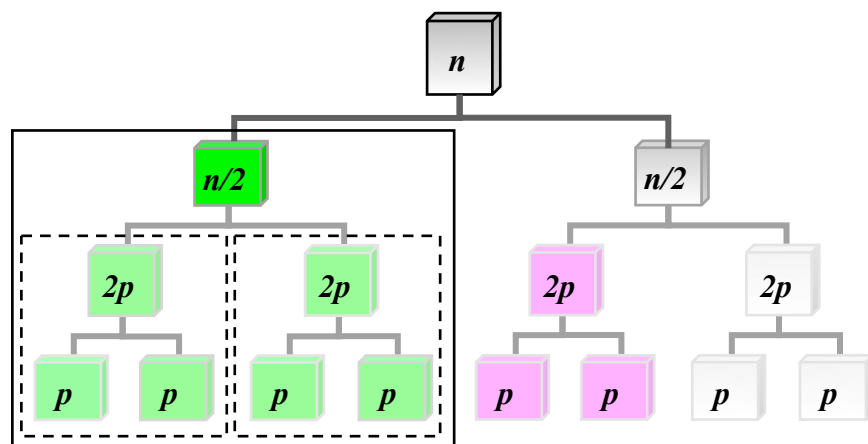
y position(m)

Number of Maps : 3



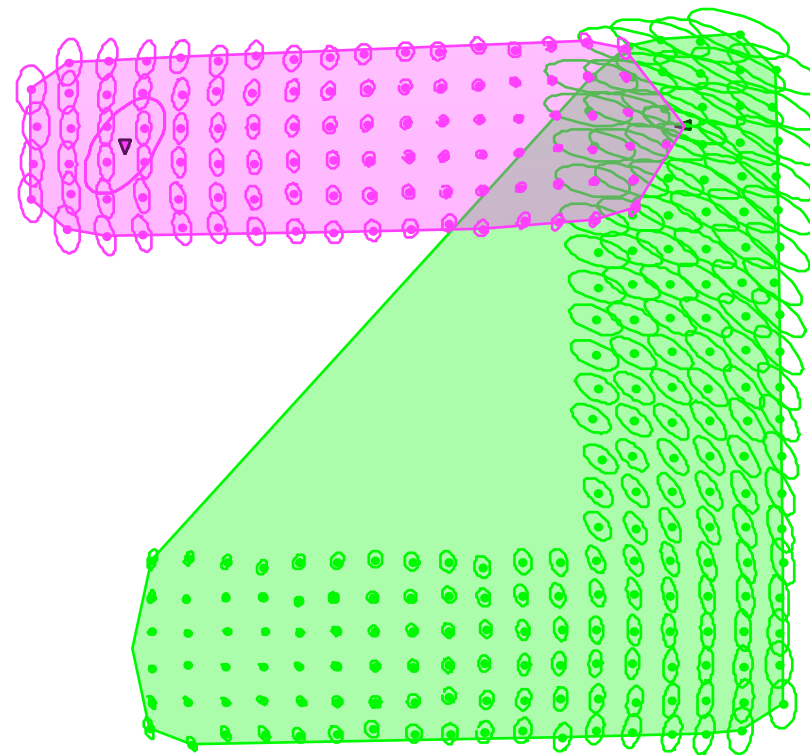
x position(m)

Divide & Conquer SLAM



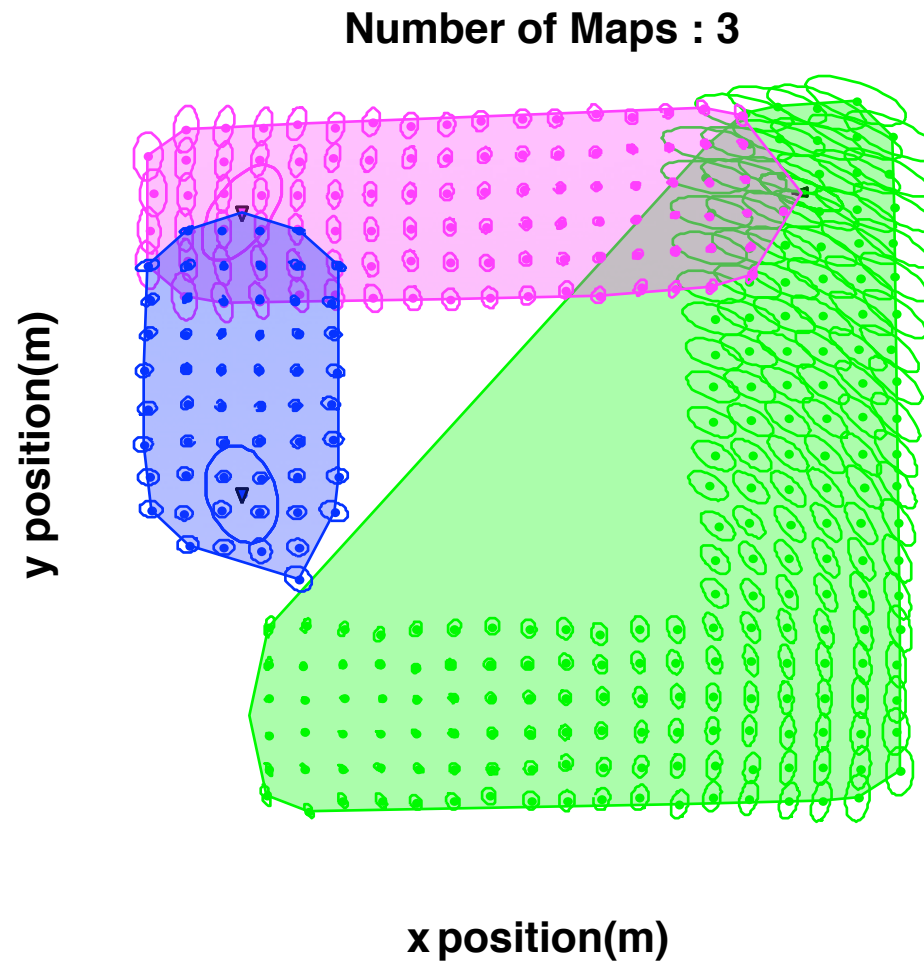
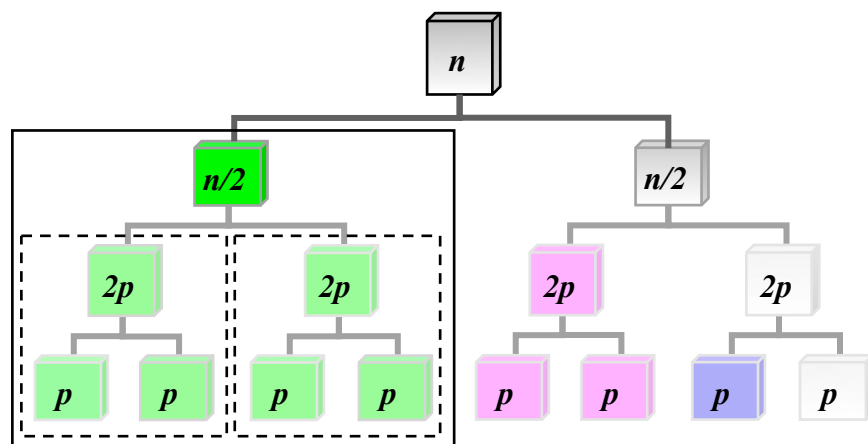
y position(m)

Number of Maps : 2

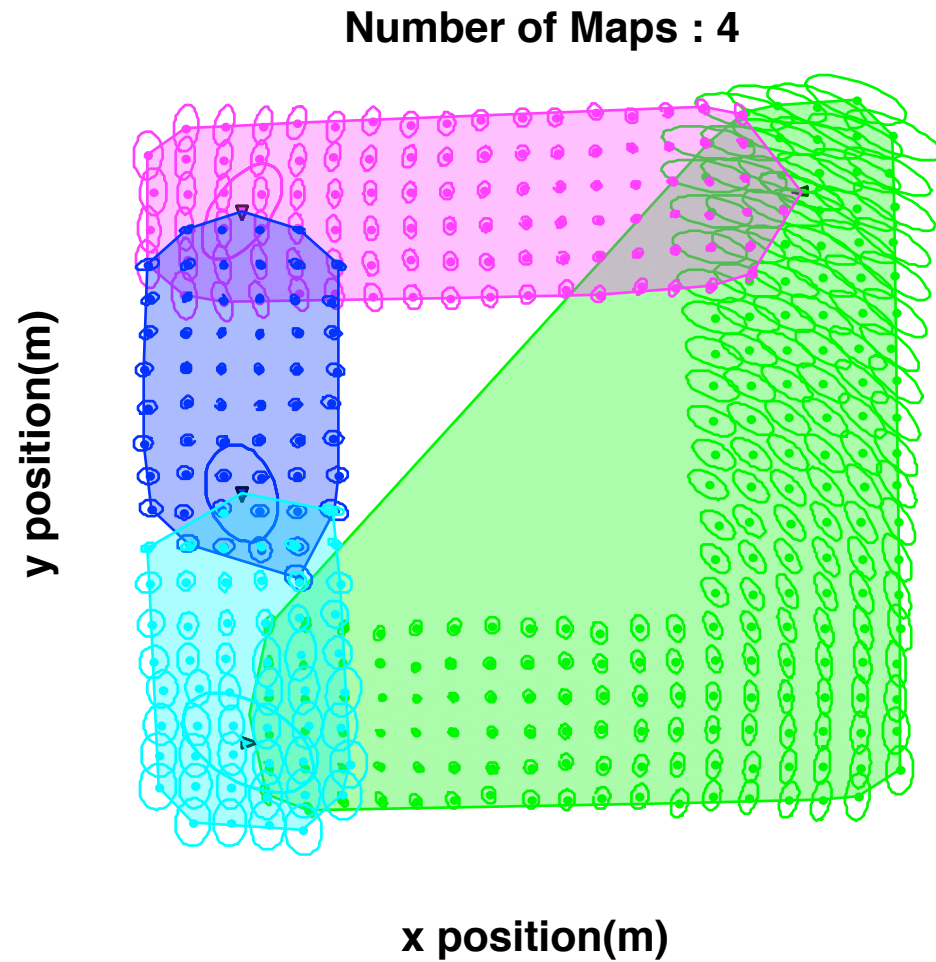
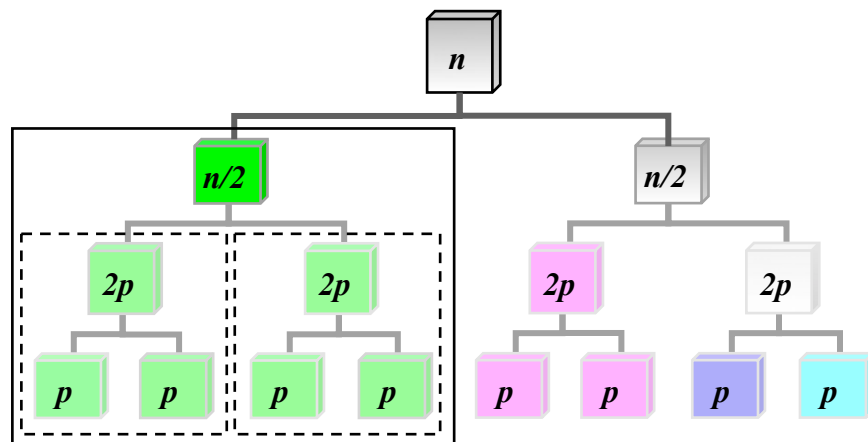


x position(m)

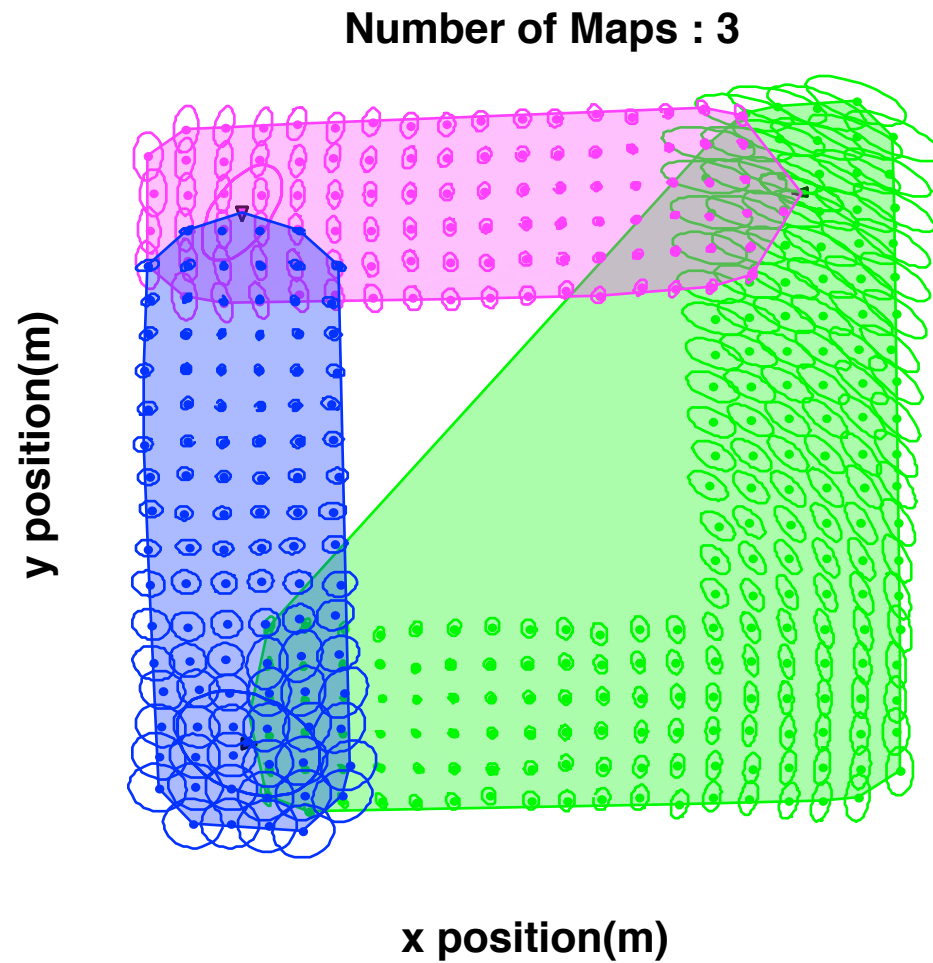
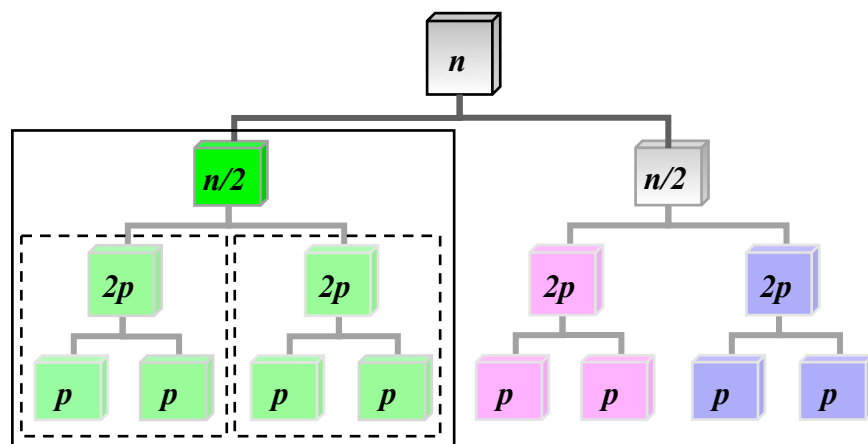
Divide & Conquer SLAM



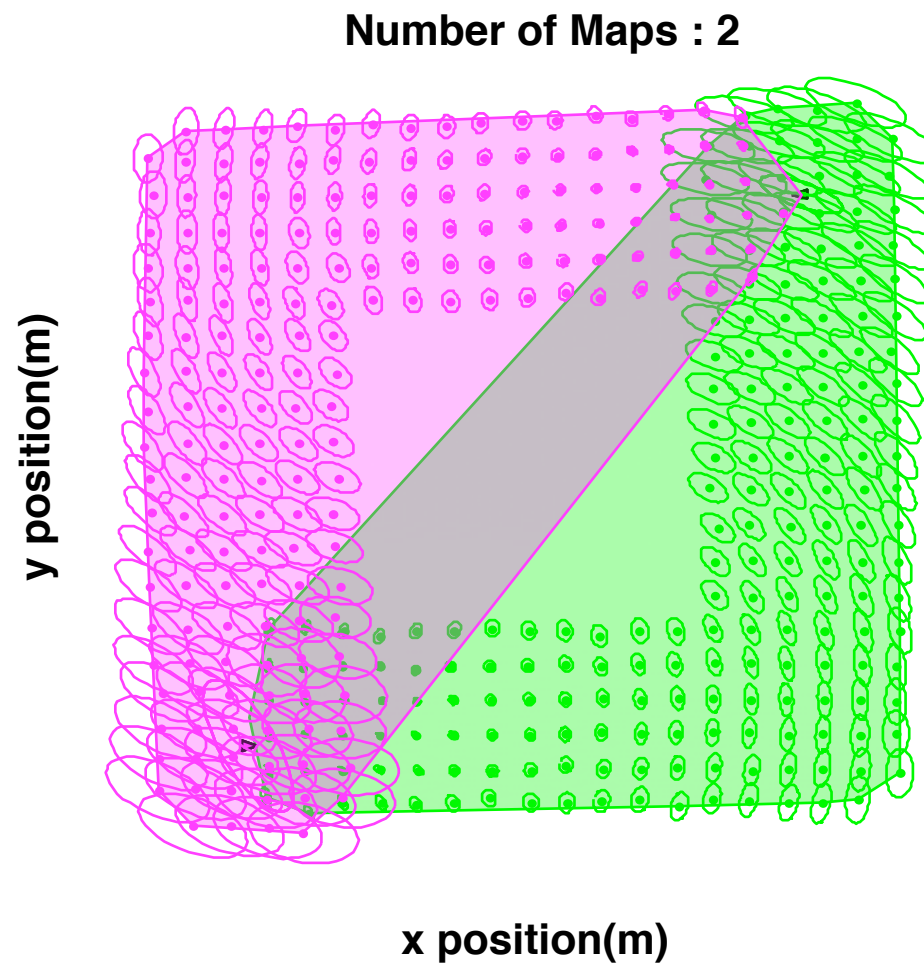
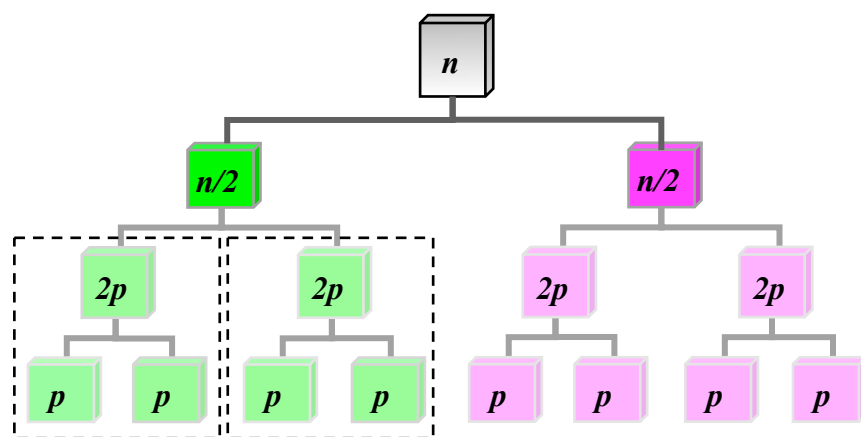
Divide & Conquer SLAM



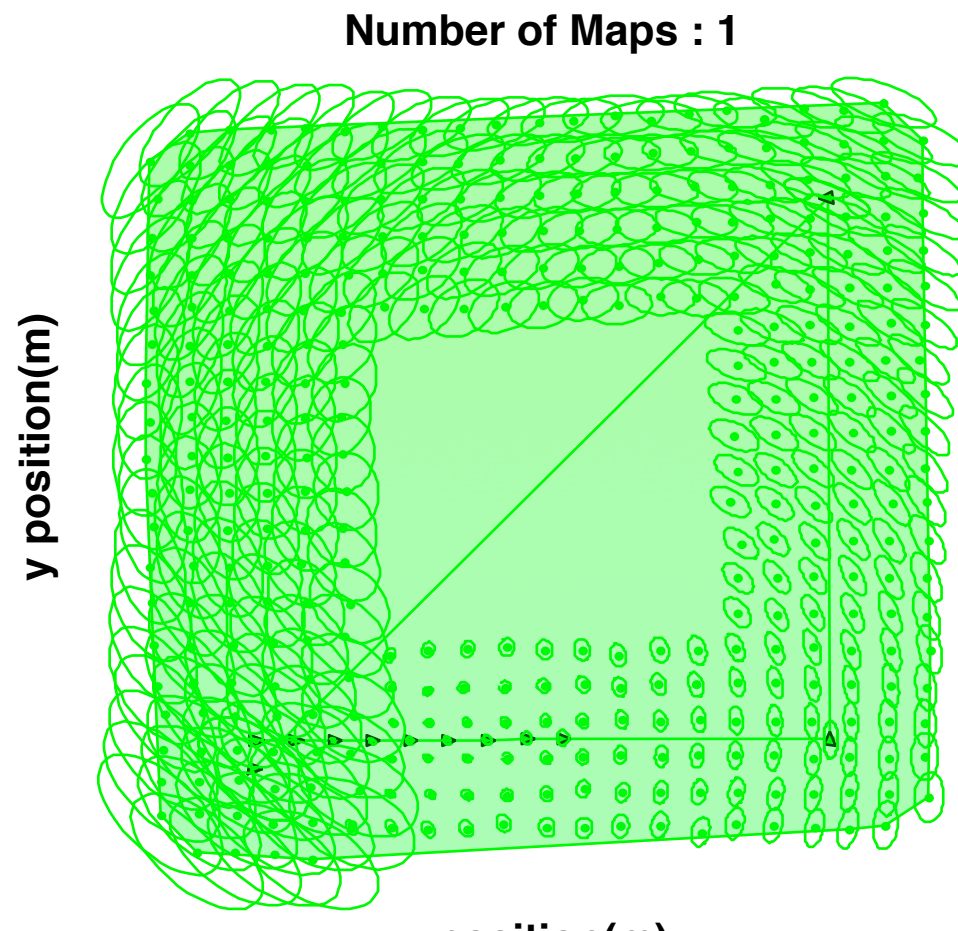
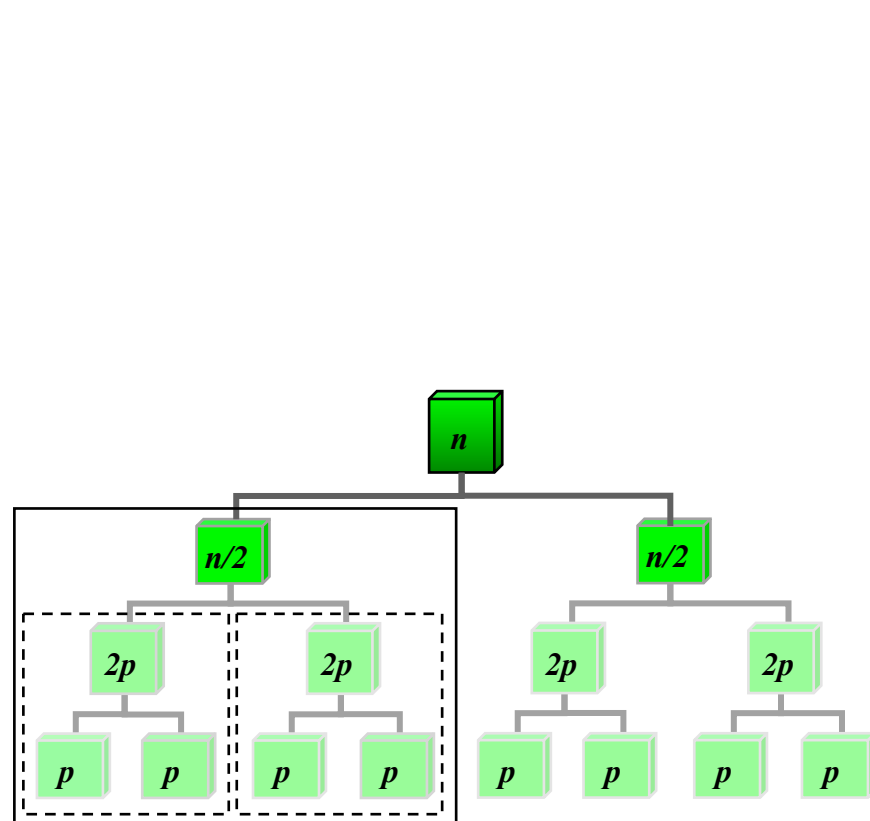
Divide & Conquer SLAM



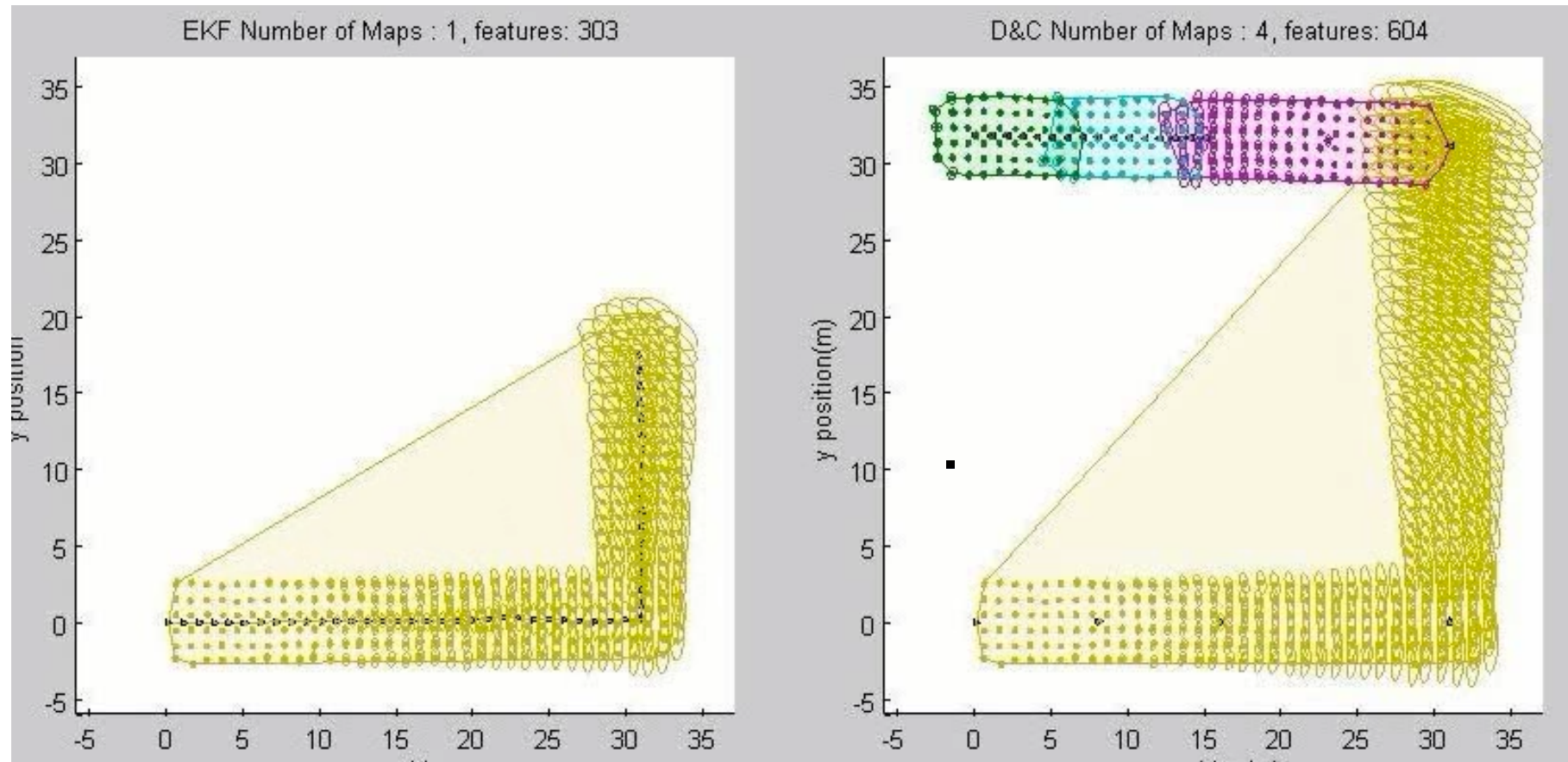
Divide & Conquer SLAM



Divide & Conquer SLAM

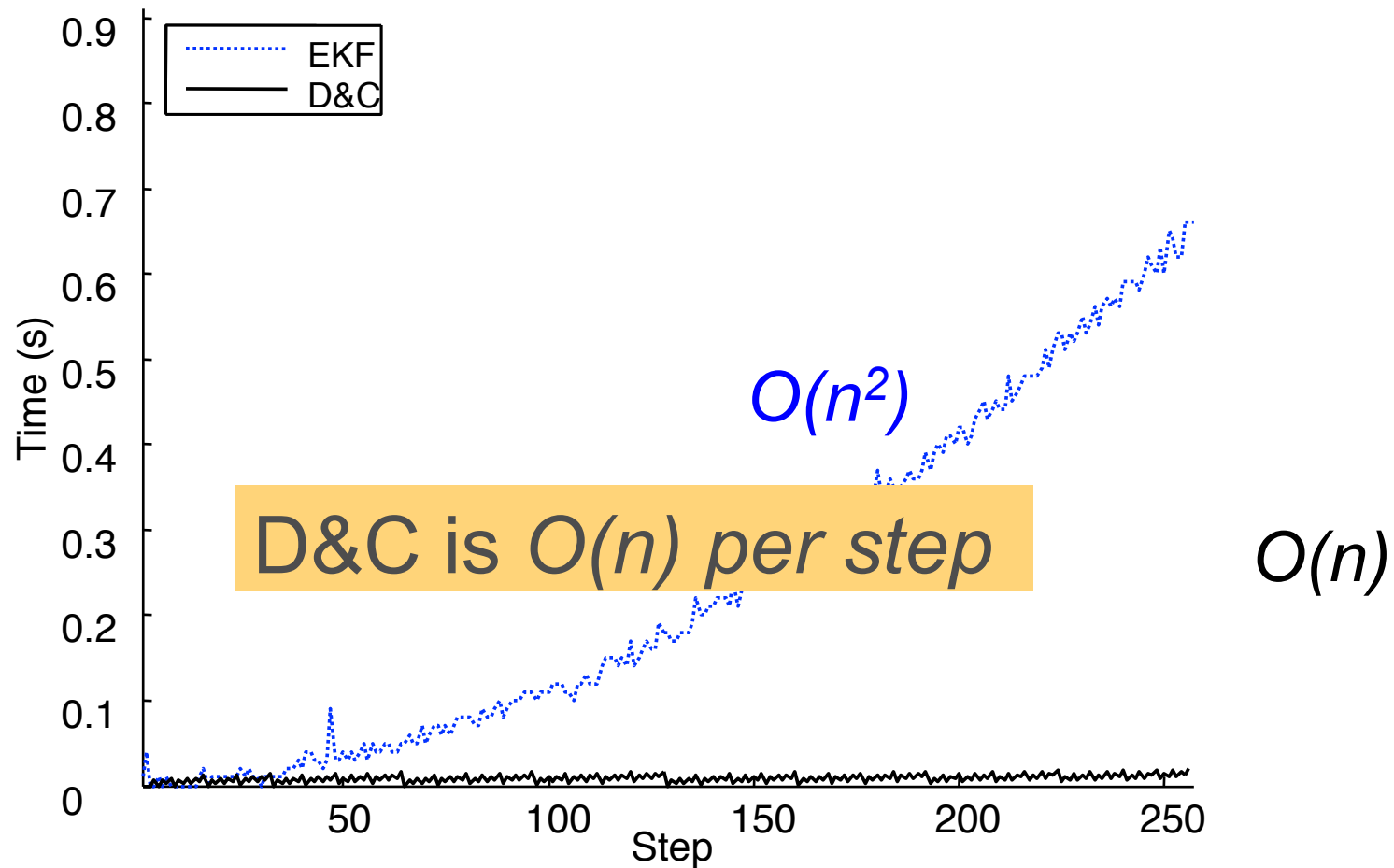


Loop Trajectory

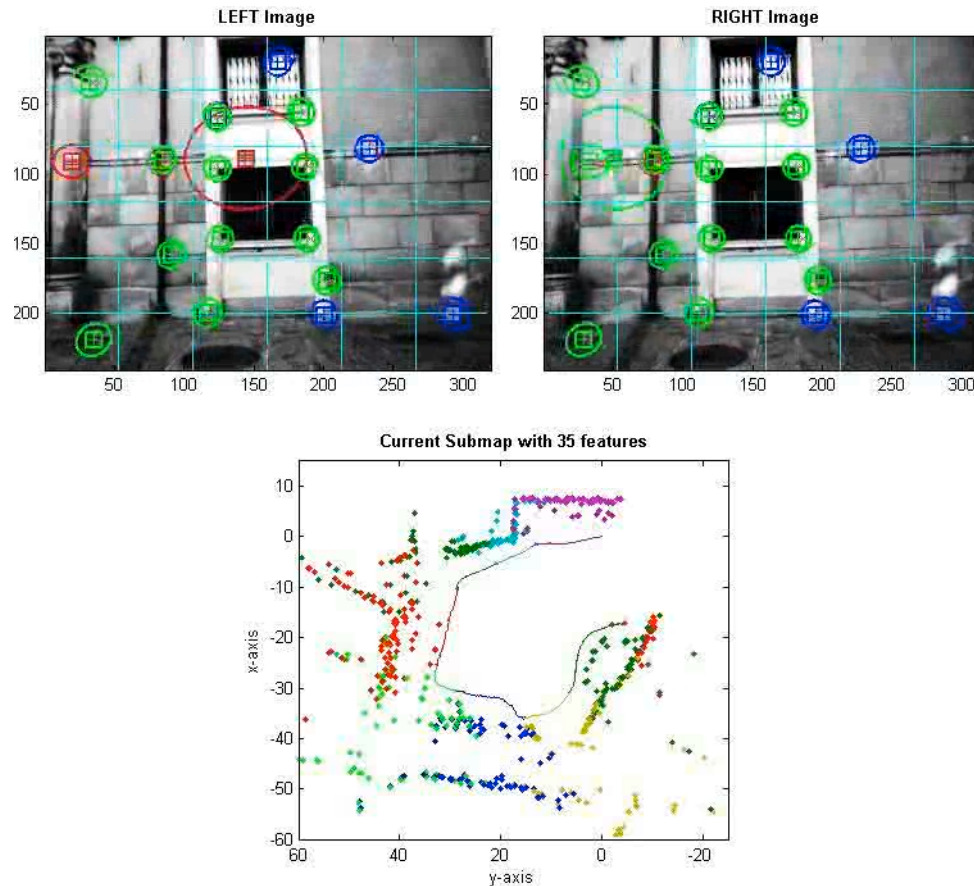


L. Paz, J. Neira and J.D. Tardós **Divide and Conquer: EKF SLAM in $O(n)$** . IEEE Transactions on Robotics, October 2008.

Amortized cost per step



6DOF SLAM with stereo

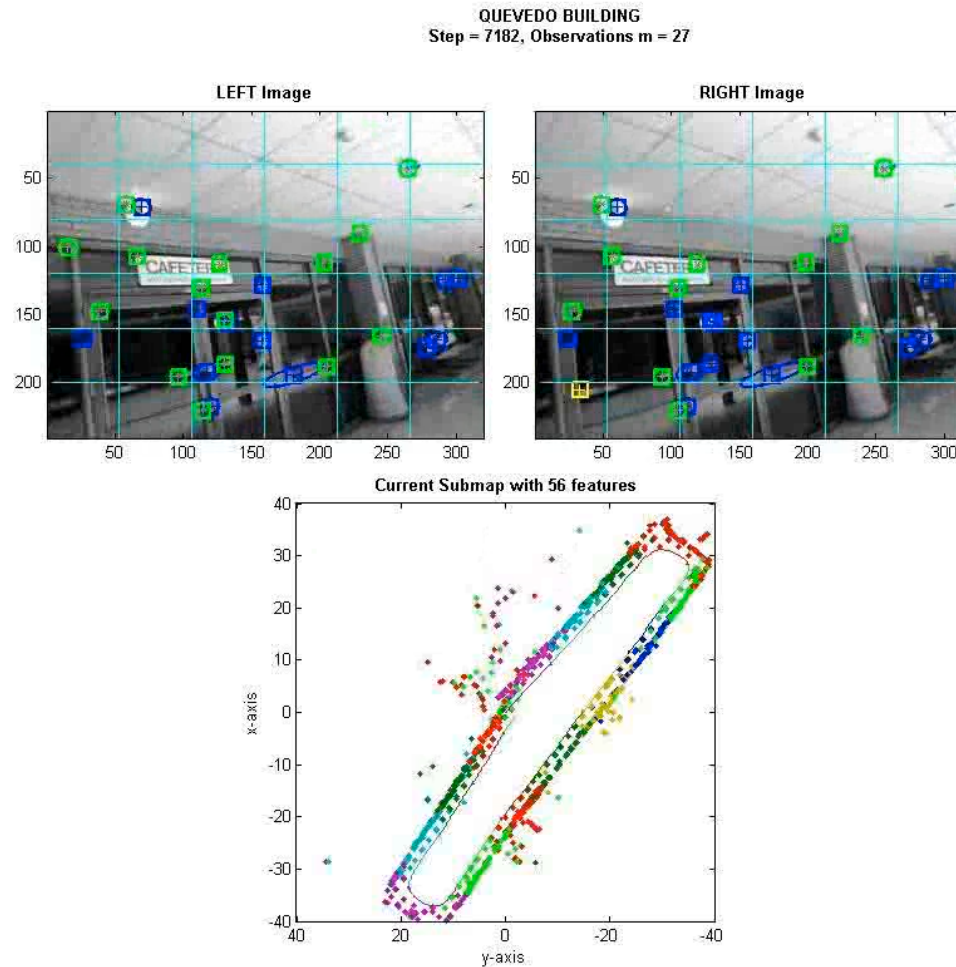


L. Paz, P. Pinies, J. Neira and J.D. Tardós **Large Scale 6DOF SLAM with Stereo-in-Hand.** IEEE Transactions on Robotics, 2008.

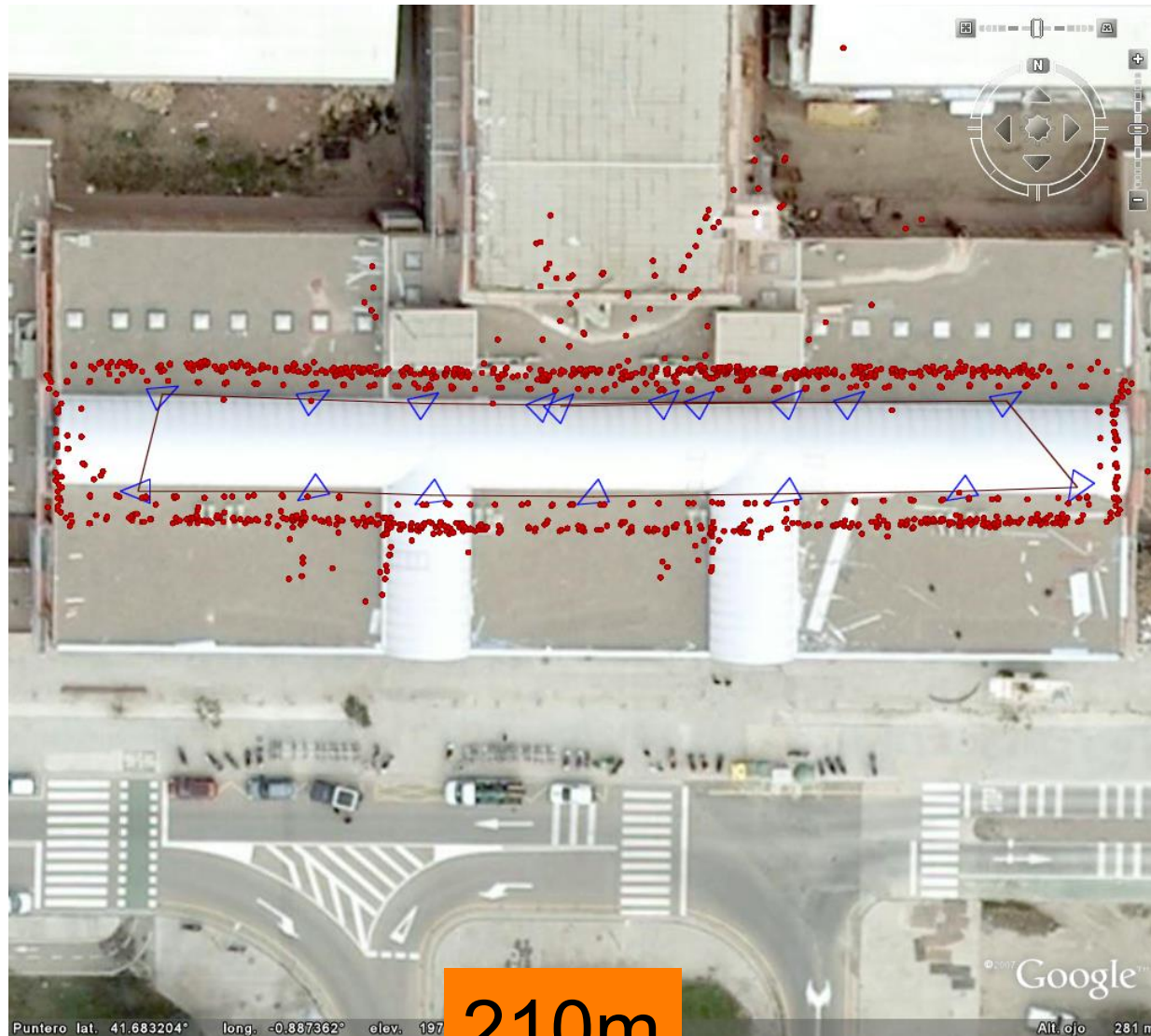
6Dof Stereo SLAM, outdoors



6Dof Stereo SLAM, indoors



6Dof Stereo SLAM, indoors



210m

Conclusions

- EKF-SLAM is only consistent for:
 - The linear case (1D robot)
 - Small scale maps ($< 100\text{m}$)
- Inconsistency only becomes evident if:
 - Ground truth is available
 - Trying to close a big loop ($> 100\text{m}$)
- The consistency problem appears before the computational complexity problem !