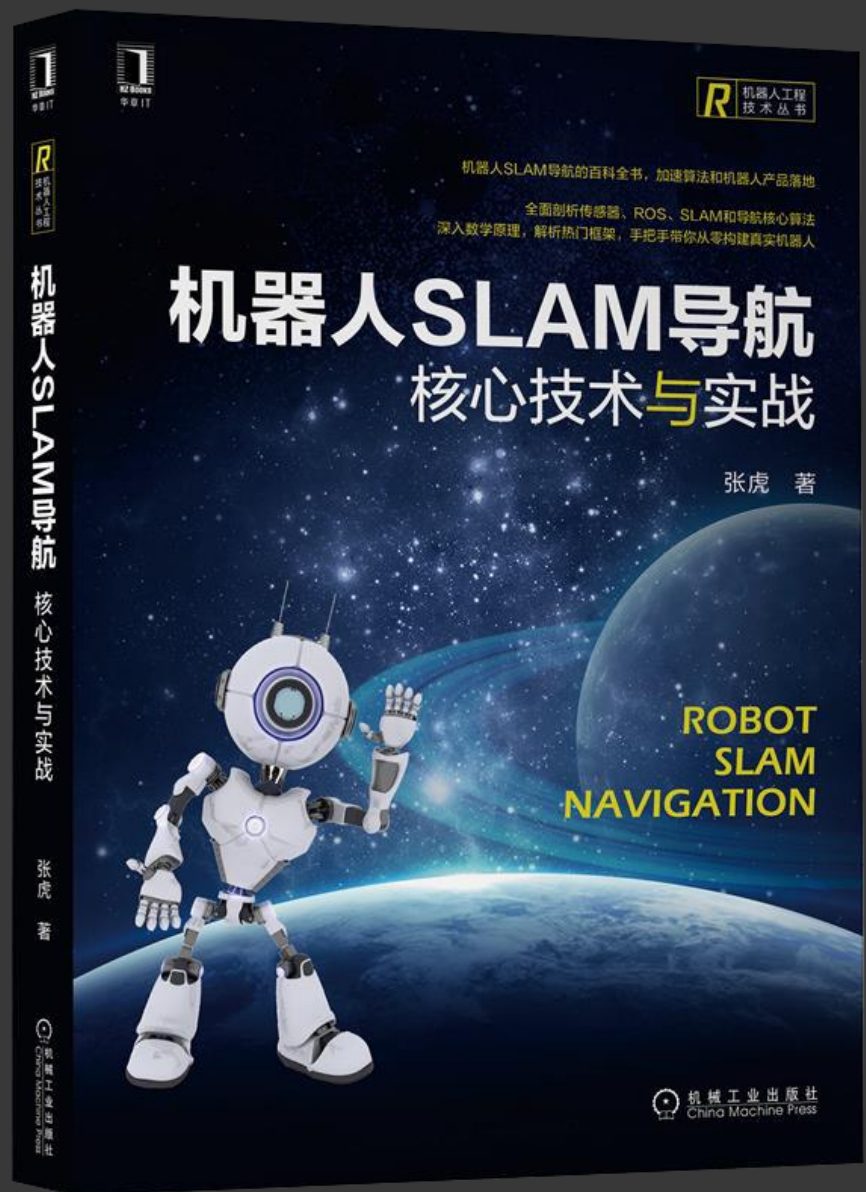




第1季

第7章：SLAM中的数学基础



主讲人：张虎
(小虎哥哥爱学习)

- 先导课
- 第1季：快速梳理知识要点与学习方法 ✓
- 第2季：详细推导数学公式与代码解析
- 第3季：代码实操以及真实机器人调试
- 答疑课

----- (永久免费 • 系列课程 • 长期更新) -----

本书内容安排

一、编程基础篇

第1章：ROS入门必备知识

第2章：C++编程范式

第3章：OpenCV图像处理

二、硬件基础篇

第4章：机器人传感器

第5章：机器人主机

第6章：机器人底盘

三、SLAM篇

第7章：SLAM中的数学基础

第8章：激光SLAM系统

第9章：视觉SLAM系统

第10章：其他SLAM系统

四、自主导航篇

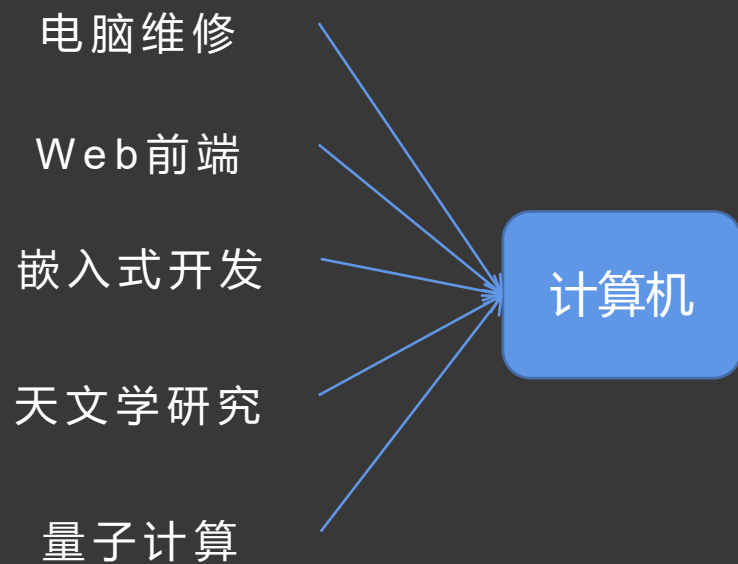
第11章：自主导航中的数学基础

第12章：典型自主导航系统

第13章：机器人SLAM导航综合实战

搞SLAM为什么还要学编程知识和硬件知识？

SLAM，你真的了解吗？



- SLAM并不仅仅指代某种特定算法（求解方法）
- SLAM是一种问题的总称
- 你可能是一个SLAMer，也可能是码农、数学家、焊板子的、机械狗、嵌入式、AI调包侠、...

SLAM，你真的了解吗？

问题	求解方法（也叫算法）
1+1=?	• 计算器 • 算盘 • 笔算 • 心算 • 扳手指 • ...
SLAM (同时定位与建图)	• Gmapping • Cartographer • LOAM • ORB-SLAM • Rtabmap • ...

注意：不要以偏概全，某种算法（比如ORB-SLAM）并不能代表SLAM的全貌，你所做的是是一个很细分的方向而已。

SLAM的价值是什么？



萌新：SLAM是不是不行了，搞SLAM还有前途吗？



大佬：

- ① SLAM是人类生产生活的客观需求，SLAM是一种问题。
- ② 只要SLAM问题一天不解决，那么搞SLAM就一定有前途。
- ③ 至少目前来看，SLAM问题还没有出现完美的解决方案。

?

SLAM的价值是什么？

躯体使智能得以延展

AI的最终归宿是机器人

机器人的完全自主化

自主移动技术

SLAM

- 如果人工智能技术仅仅停留在虚拟的网络和数据之中，那么其挖掘并利用新知识的能力将很难扩展开来。
- 换句话说，人工智能需要真实的**躯体**。

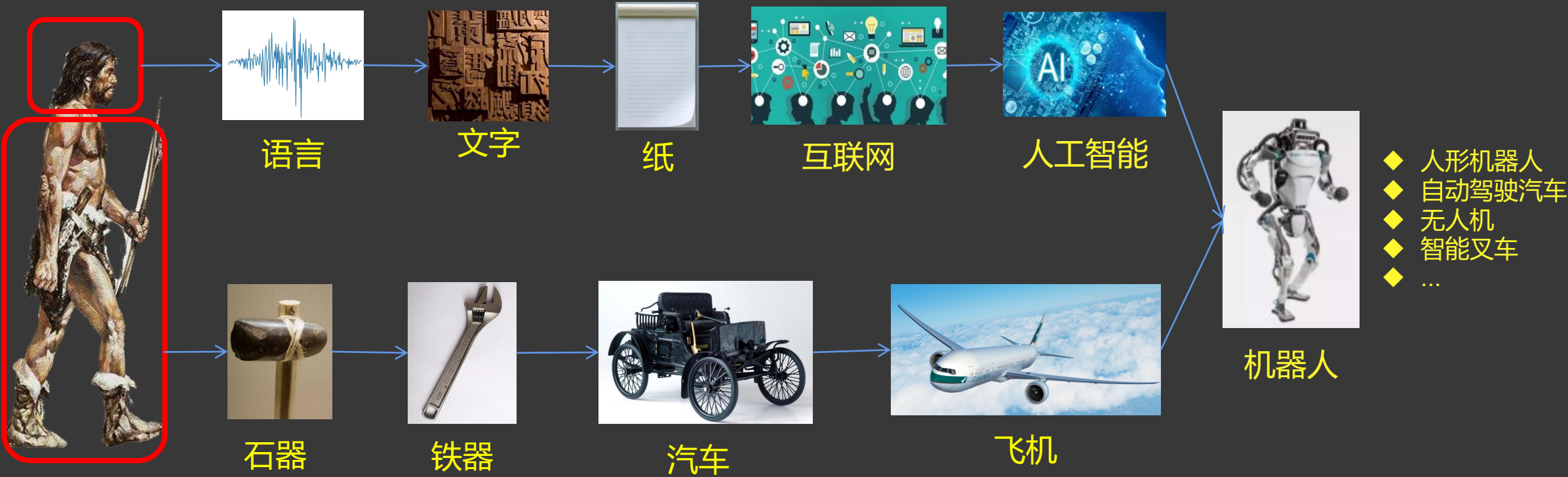


目前的人工智能

互联网的虚拟世界

真实世界

SLAM的价值是什么？



SLAM的价值是什么？



- ✓ 让机器人实现完全自主化也就成了人类一直以来的梦想。
- ✓ 没有外界指令的干预
- ✓ 机器人能通过传感器和执行机构与环境自动发生交互
- ✓ 完成特定的任务

机器人自主化
底层技术基石



实现机器人自主化的关键

SLAM的价值是什么？



为什么需要自主移动？



机械臂

VS



机械臂+移动底盘



智能音响

VS



智能音响+移动底盘

SLAM的价值是什么？



早期计算机



智能手机



自主移动机器人

- 轮式底盘
- 双足人形机器人
- 四足机器狗
- 无人机
- 无人船
- ...

*** 机器人是人工智能技术应用能力的有效延展，而能自主移动的机器人更是极大地拓展了人工智能技术的应用范围。**

SLAM的价值是什么？

躯体使智能得以延展

AI的最终归宿是机器人

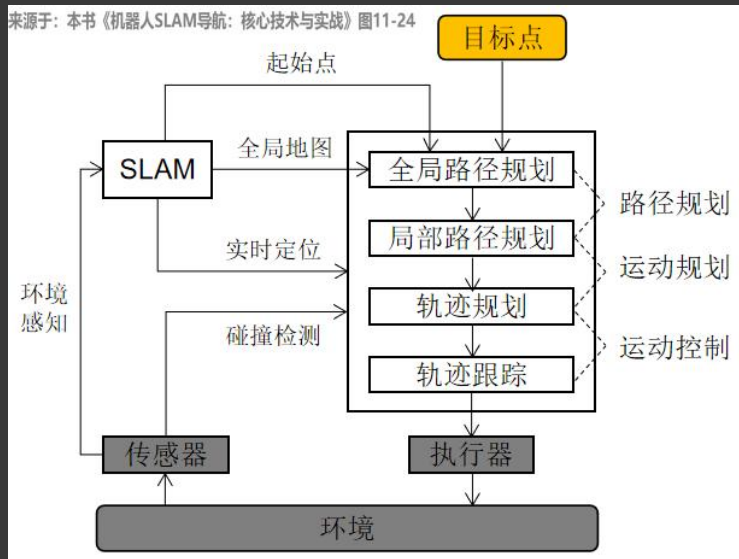
机器人的完全自主化

自主移动技术

SLAM

从工程上看，自主移动实质上就是解决从地点A到地点B的问题

- ① 我在哪？
- ② 我将到何处去？
- ③ 我该如何去



目前SLAM的概念已经泛化了，对定位、导航、自主移动等应用的笼统习惯指代。SLAM问题的求解方法也早已脱离滤波或优化等传统方法，现在的新求解方法与传统的SLAM理论框架可能有较大区别。

如果不用SLAM这个名词，你也可以用ABC、123之类的名词来指代。所以，大家不用再纠结于SLAM这个名词本身的定义，名词的定义往往只是约定俗成而已。

SLAM的应用领域？

- 机器人
- 自动驾驶
- 其他

自主移动是一种根本性需求



火星探测



农业采摘



服务机器人



深海勘探



电力巡检



扫地机器人



工厂物料运送AGV



送餐



岩洞勘测机器人



割草

SLAM的应用领域？

- 机器人
- 自动驾驶
- 其他

自主移动是一种根本性需求



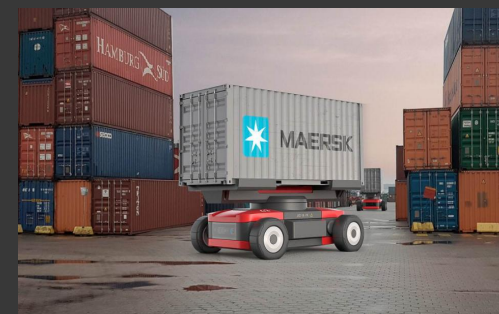
无人驾驶



物流配送



无人清扫车



港口集装箱搬运



无人货运卡车



景区无人驾驶观光车



无人驾驶公交车

SLAM的应用领域?

- 机器人
- 自动驾驶
- 其他

自主移动是一种根本性需求



测绘无人机



建筑物手持SLAM建图仪



AR/VR中的应用

内容概要

7.1 SLAM发展简史

7.2 SLAM中的概率理论

7.3 估计理论

7.4 基于贝叶斯网络的状态估计

7.5 基于因子图的状态估计

7.6 典型SLAM算法

7.7 SFM、BA和SLAM比较

7.1 SLAM发展简史

- SLAM的起源
- 数据关联、收敛和一致性
- SLAM的基本理论
- SLAM的学习方法



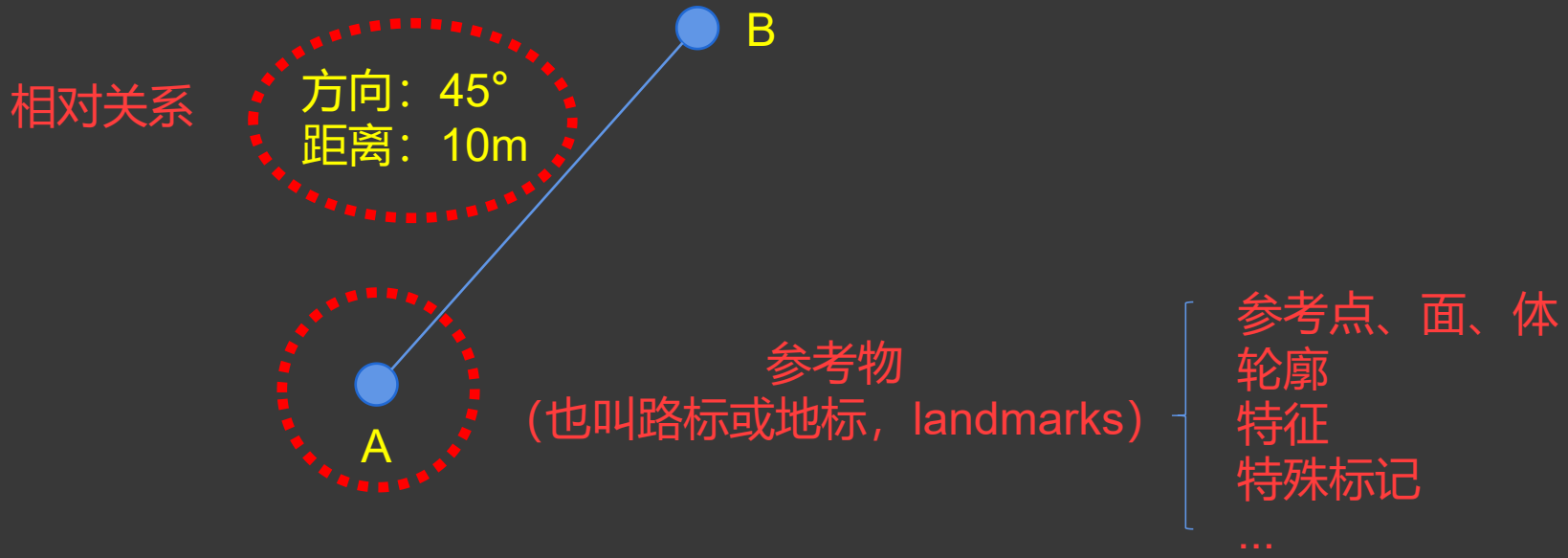
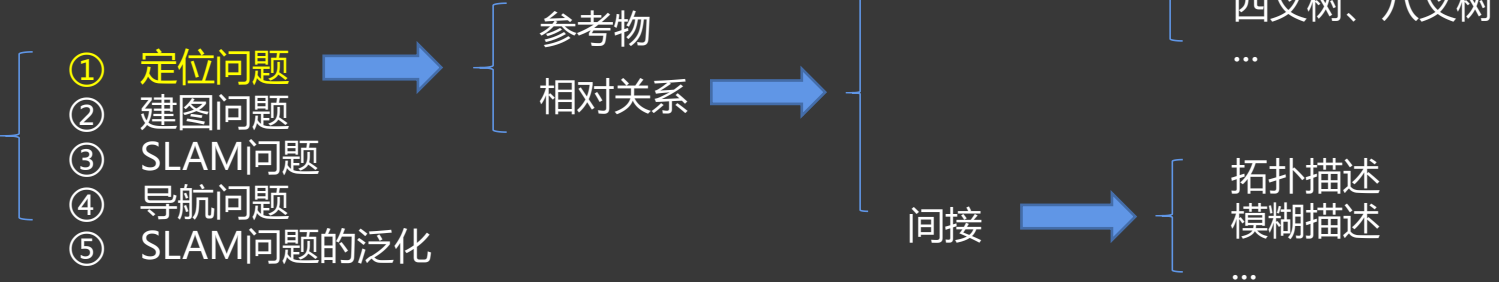
SLAM并不是某种具体的算法，而是一类问题的总称。

问题	求解方法（也叫算法）
1+1=?	• 计算器 • 算盘 • 笔算 • 心算 • 扳手指 • ...
SLAM (同时定位与建图)	• Gmapping • Cartographer • LOAM • ORB-SLAM • Rtabmap • ...

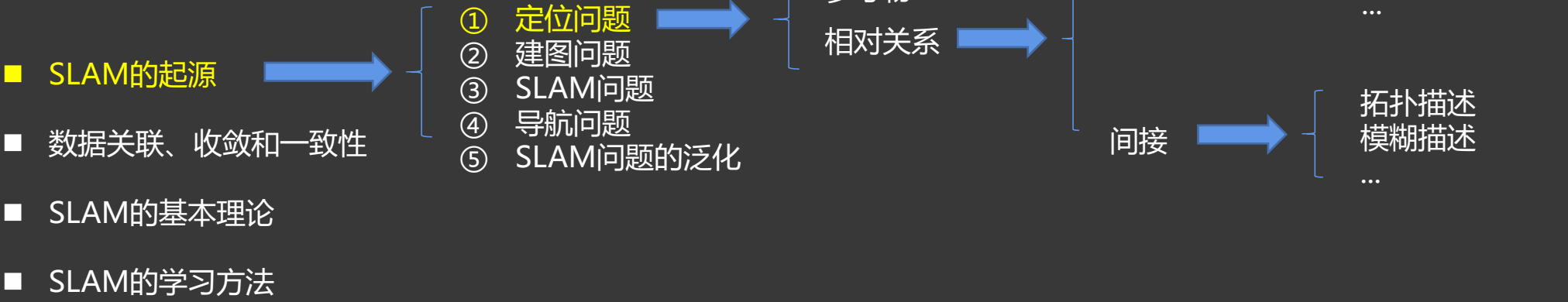
7.1 SLAM发展简史

■ SLAM的起源

- 数据关联、收敛和一致性
- SLAM的基本理论
- SLAM的学习方法



7.1 SLAM发展简史



欧氏坐标

$A = (1,2)$
 $\overrightarrow{AB} = (2,3)$
 $B = A + \overrightarrow{AB} = (3,5)$

极坐标

$A = (0,0^\circ)$
 $B = (10,45^\circ)$

球坐标

$A = (0,0^\circ,0^\circ)$
 $B = (10,30^\circ,45^\circ)$

四叉树

$A = (1,1,1,1,...)$
 $B = (1,4,2,3,...)$

八叉树

$A = (1,1,1,1,...)$
 $B = (1,8,5,7,...)$

7.1 SLAM发展简史

■ SLAM的起源

■ 数据关联、收敛和一致性

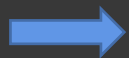
■ SLAM的基本理论

■ SLAM的学习方法

- ① 定位问题
- ② 建图问题
- ③ SLAM问题
- ④ 导航问题
- ⑤ SLAM问题的泛化

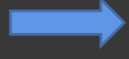
参考物
相对关系

直接



欧氏坐标
极坐标
球坐标
二叉树、八叉树
...

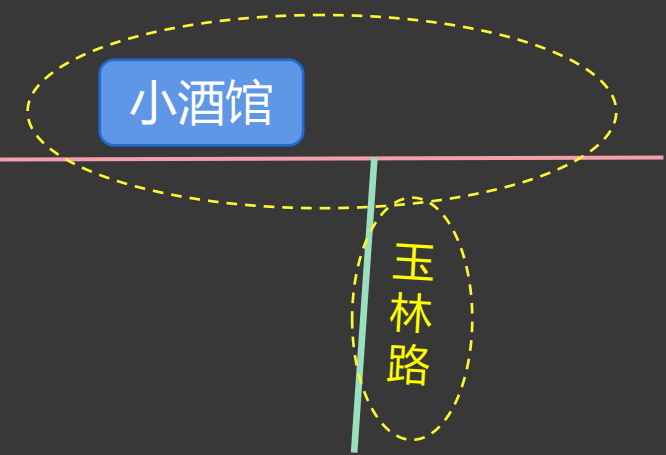
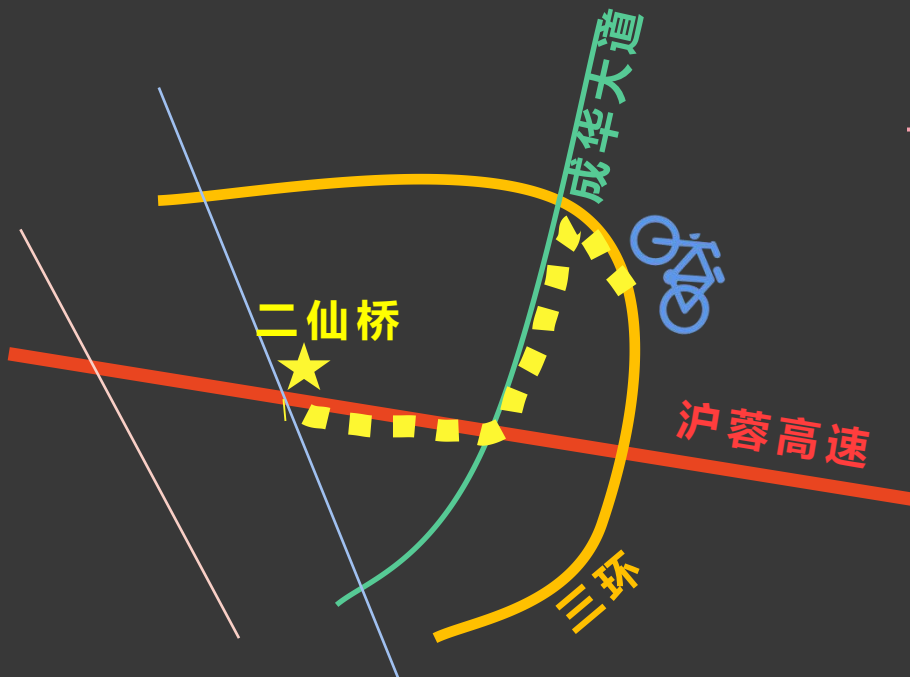
间接



拓扑描述
模糊描述
...

拓扑描述

模糊描述



【成都旅游打卡】

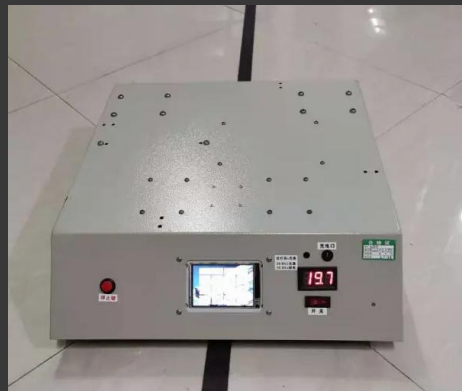
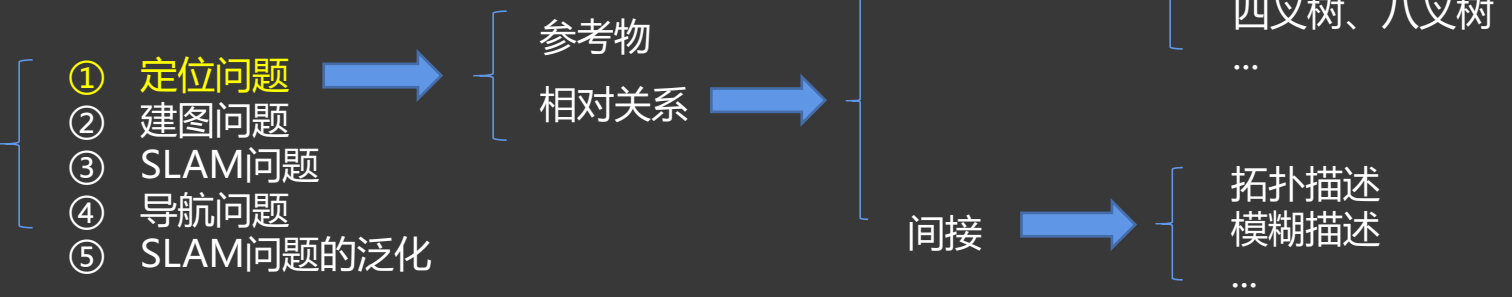
走到玉林路的尽头
坐在小酒馆的门口

成都天府
国际机场

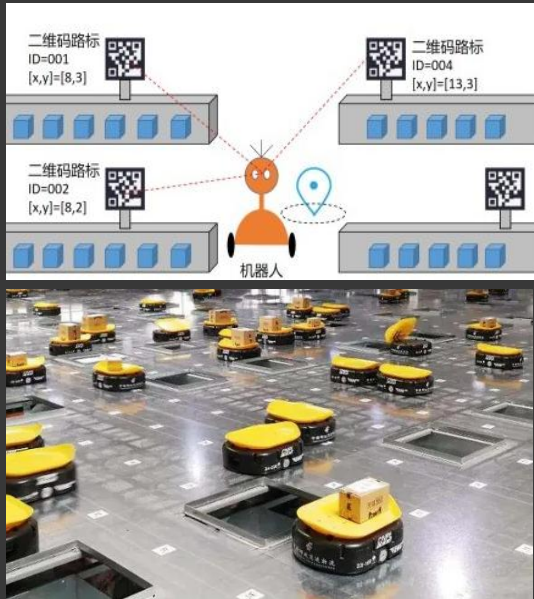


7.1 SLAM发展简史

- SLAM的起源
- 数据关联、收敛和一致性
- SLAM的基本理论
- SLAM的学习方法



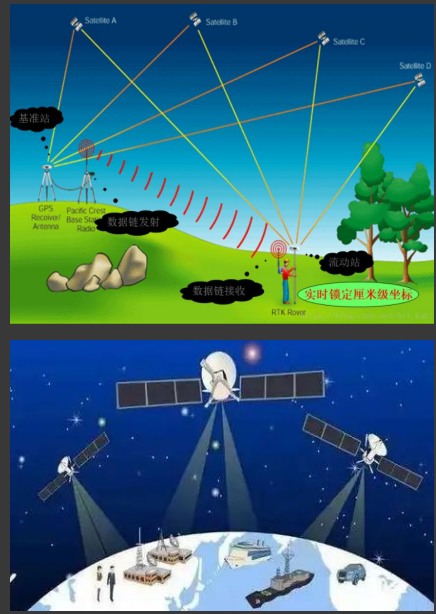
磁条参考物



二维码参考物



无线电参考物



卫星参考物



7.1 SLAM发展简史

■ SLAM的起源



- 数据关联、收敛和一致性
- SLAM的基本理论
- SLAM的学习方法

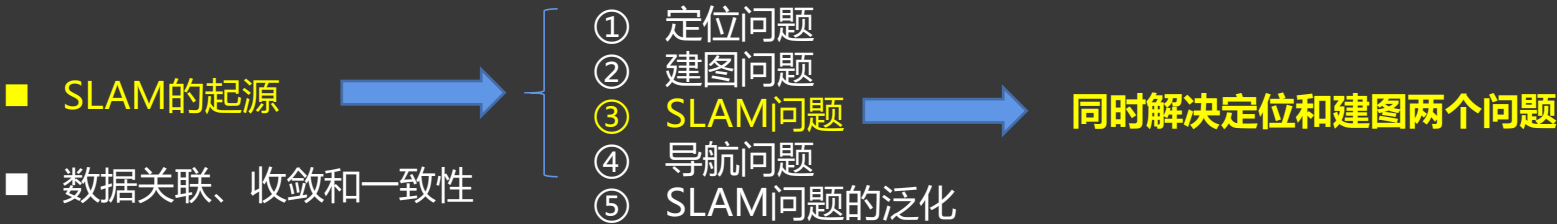
- ① 定位问题
- ② 建图问题
- ③ SLAM问题
- ④ 导航问题
- ⑤ SLAM问题的泛化



人体三维结构=切片1+切片2+切片3+...

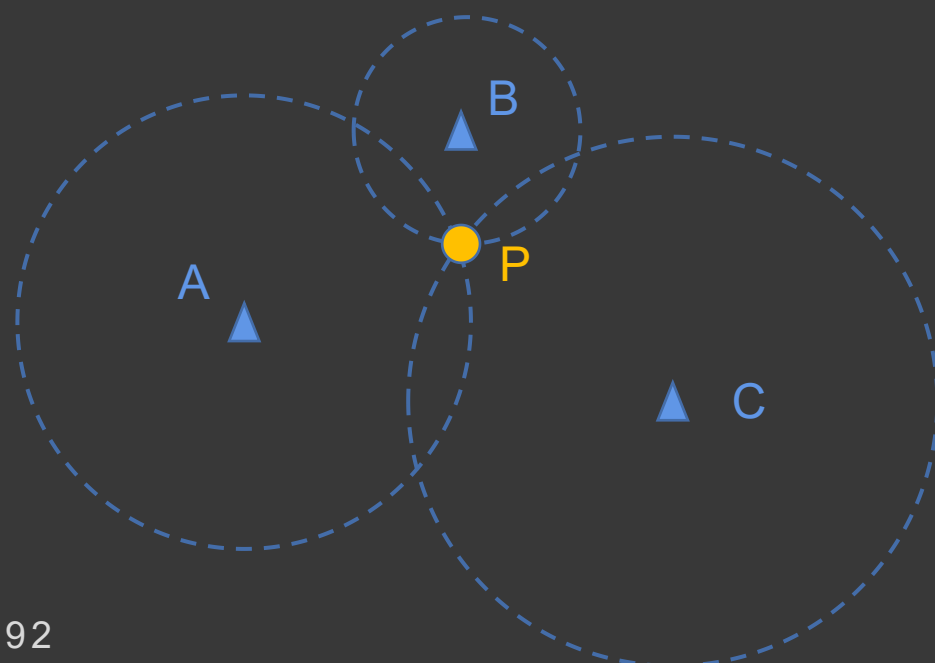
建图的本质都是基于已知位置切片的拼接

7.1 SLAM发展简史



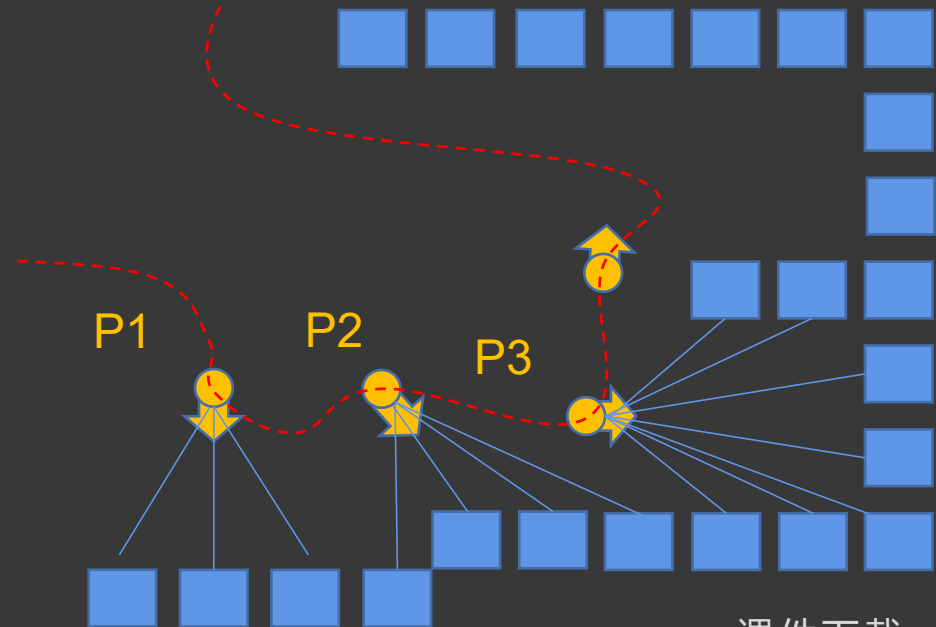
定位问题：

- 已知全局地图坐标信息
- 利用观测求解机器人在地图中的坐标



建图问题：

- 已知机器人的实时坐标
- 利用观测求解被观测物体的坐标



7.1 SLAM发展简史

■ SLAM的起源

■ 数据关联、收敛和一致性

■ SLAM的基本理论

■ SLAM的学习方法

- ① 定位问题
- ② 建图问题
- ③ SLAM问题
- ④ 导航问题
- ⑤ SLAM问题的泛化

同时解决定位和建图两个问题

定位问题

VS

建图问题

地图
信息

推测

定位
信息

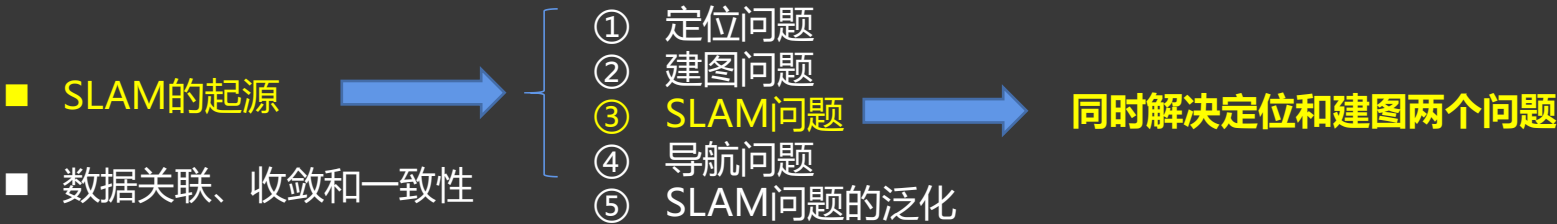
定位
信息

推测

地图
信息

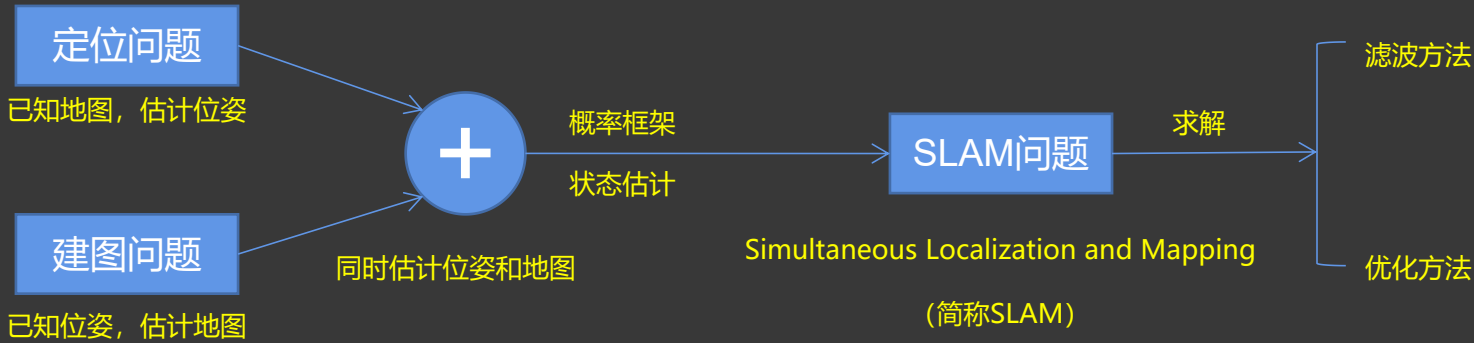


7.1 SLAM发展简史



谁都会提出几个问题，关键是要能给出解决问题的思路才行

(PS:要么解决问题，要么解决提出问题的人)



7.1 SLAM发展简史

■ SLAM的起源

■ 数据关联、收敛和一致性

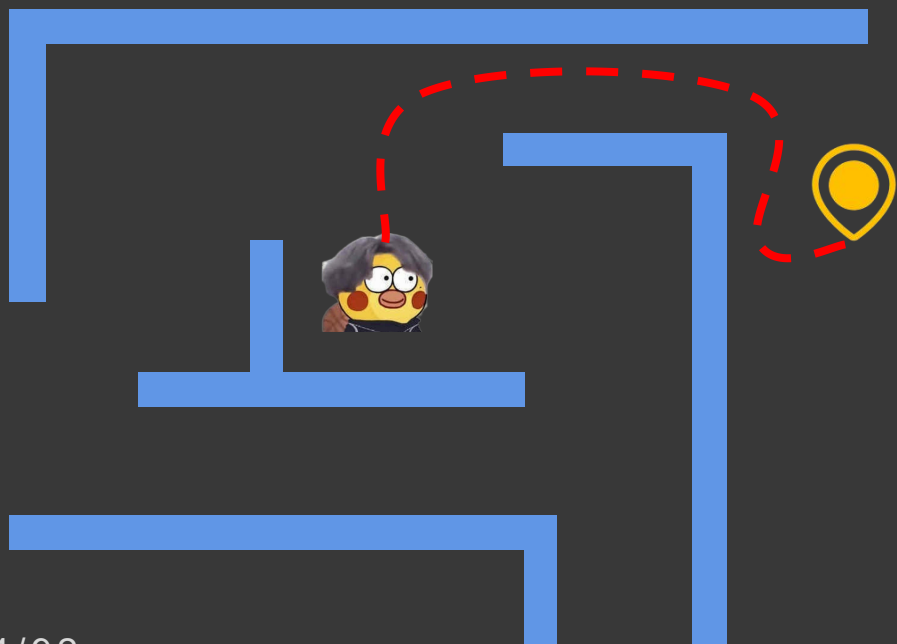
■ SLAM的基本理论

■ SLAM的学习方法

- ① 定位问题
- ② 建图问题
- ③ SLAM问题
- ④ 导航问题
- ⑤ SLAM问题的泛化

从工程上看，自主移动实质上就是解决从地点A到地点B的问题

- ① 我在哪？
- ② 我将到何处去？
- ③ 我该如何去



我在哪？ -> 取决于何种参考系

我将到何处去？ -> 取决于目的地在参考系下采用的描述方式

我该如何去？ -> 取决于路线和运动方式

7.1 SLAM发展简史

■ SLAM的起源

■ 数据关联、收敛和一致性

■ SLAM的基本理论

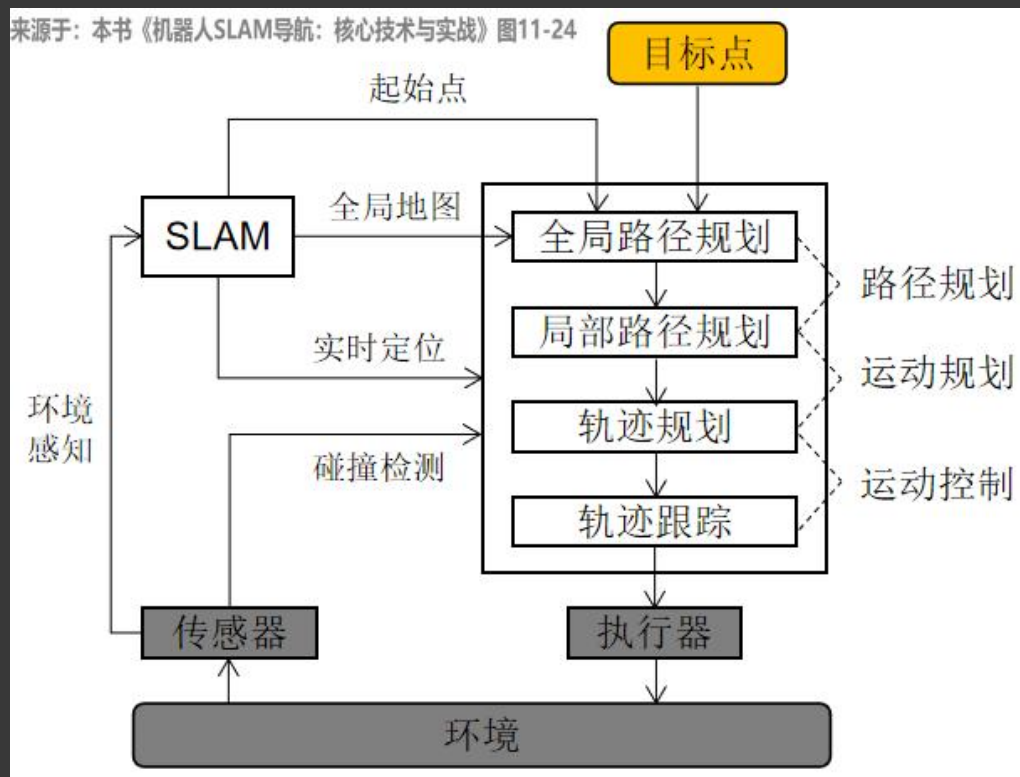
■ SLAM的学习方法

- ① 定位问题
- ② 建图问题
- ③ SLAM问题
- ④ 导航问题
- ⑤ SLAM问题的泛化

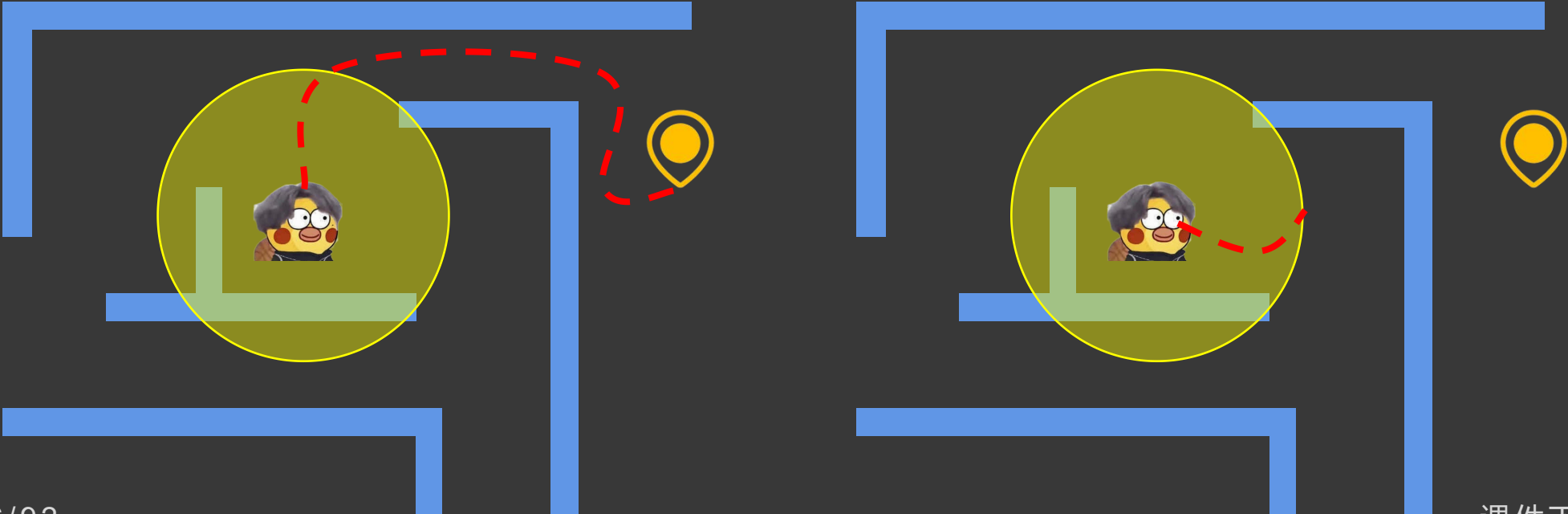
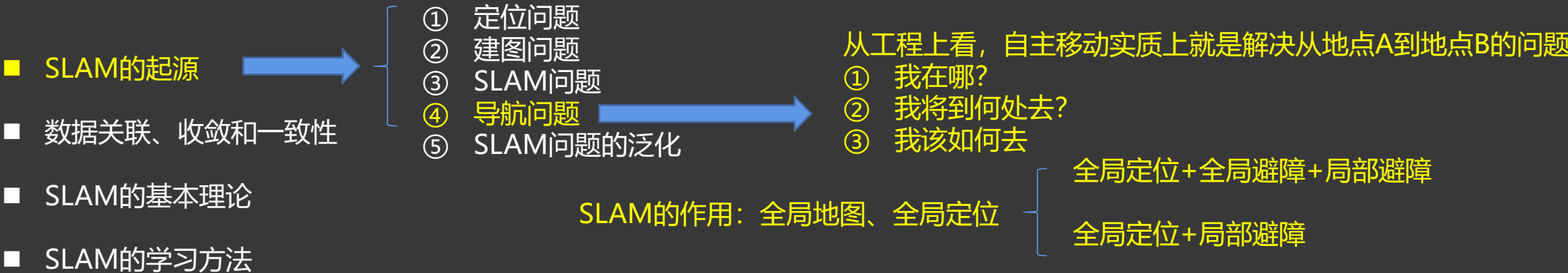
从工程上看，自主移动实质上就是解决从地点A到地点B的问题

- ① 我在哪？
- ② 我将到何处去？
- ③ 我该如何去

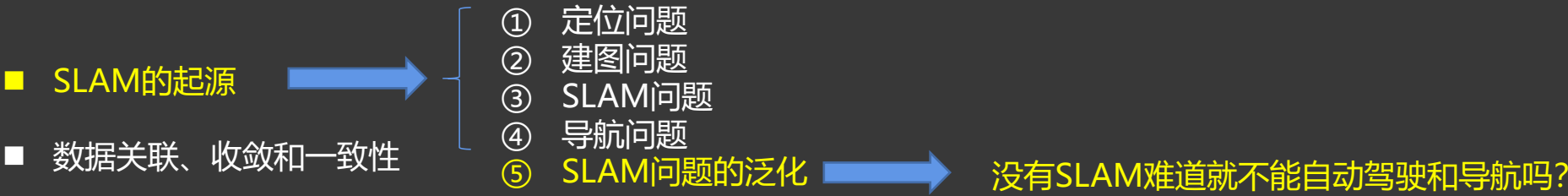
虽然其它领域（比如VR/AR/测绘/三维建模）也讨论SLAM问题，但自动驾驶/机器人导航对SLAM问题更迫切



7.1 SLAM发展简史



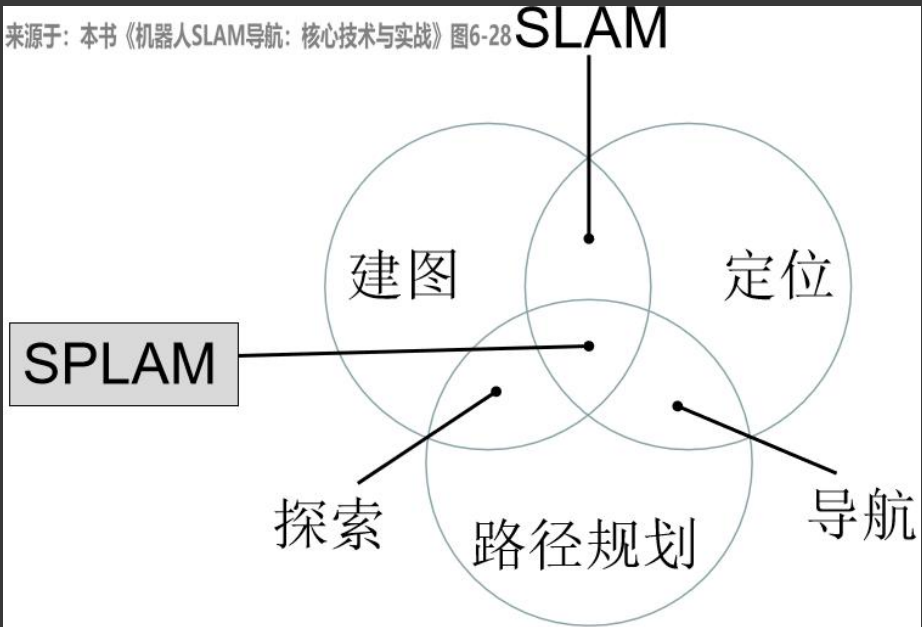
7.1 SLAM发展简史



- GPS信号覆盖盲区
- GPS定位延迟
- 高精度定位
- 全局决策
- 未知环境探索

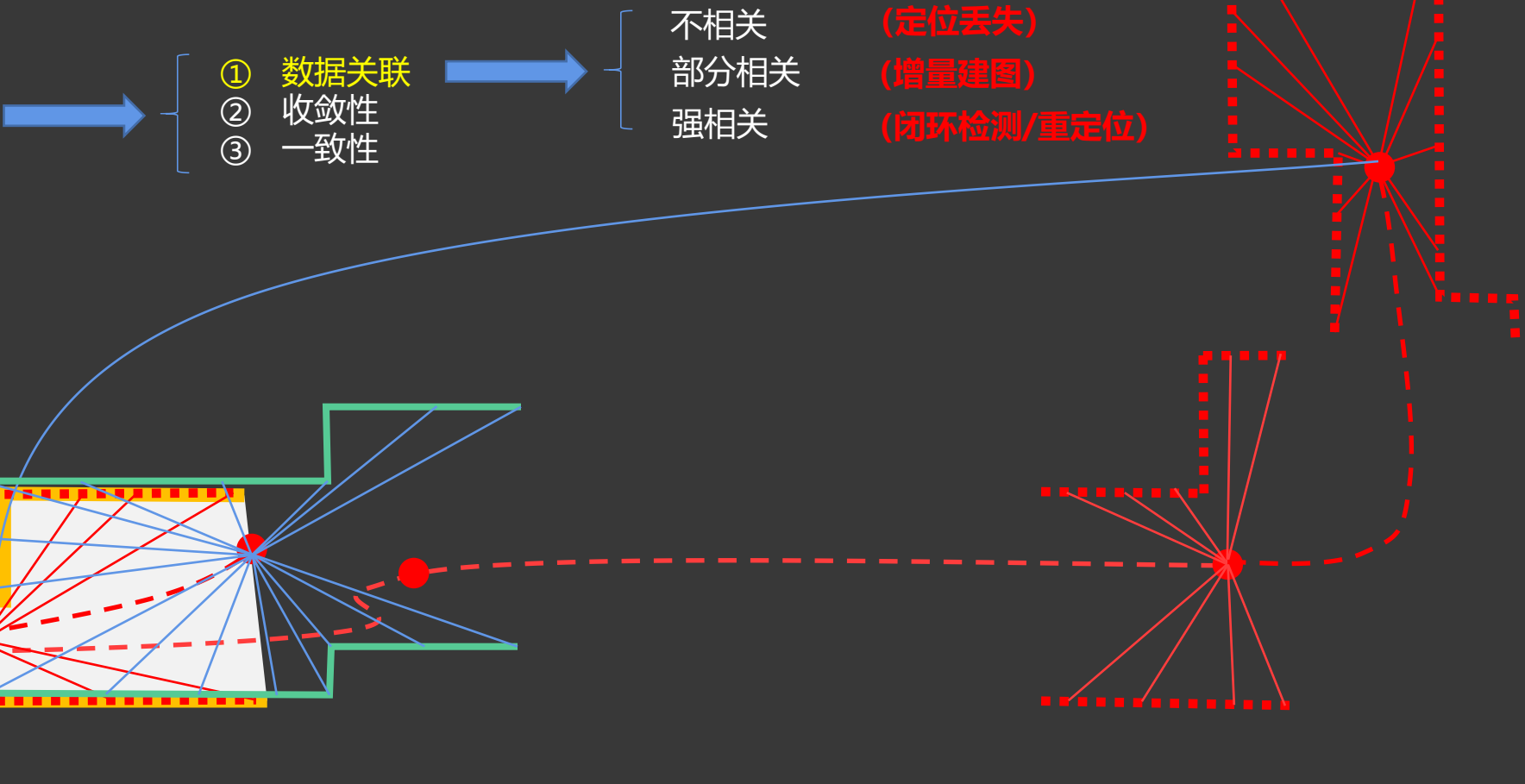
目前SLAM的概念已经泛化了，对定位、导航、自主移动等应用的笼统习惯指代。SLAM问题的求解方法也早已脱离滤波或优化等传统方法，现在的新求解方法与传统的SLAM理论框架可能有较大区别。

如果不用SLAM这个名词，你也可以用ABC、123之类的名词来指代。所以，大家不用再纠结于SLAM这个名词本身的定义，名词的定义往往只是约定俗成而已。



7.1 SLAM发展简史

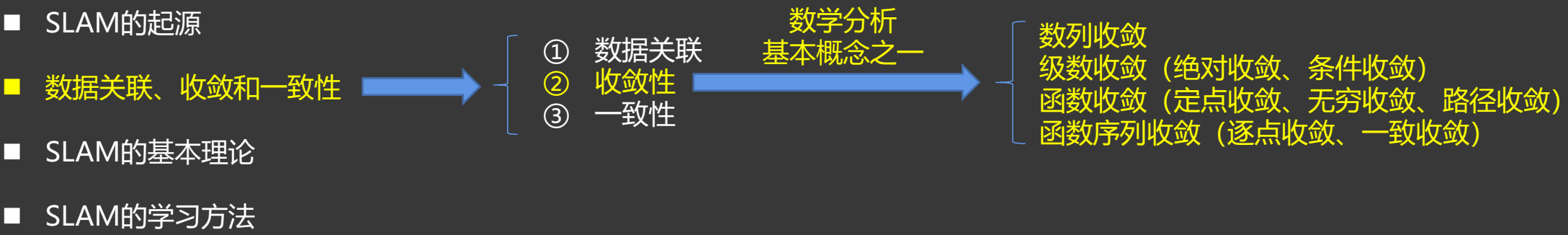
- SLAM的起源
- 数据关联、收敛和一致性
- SLAM的基本理论
- SLAM的学习方法



数据关联的数据形式是多样的，除了基于激光轮廓，还可以基于激光特征点、图像特征、机器学习特征等。

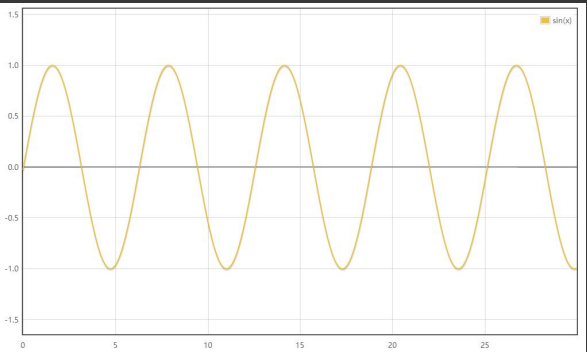
(其实目前的人工智能算法，也就相当于一个数据关联模型)

7.1 SLAM发展简史



$\{x_n\} = \{0.1 \quad 0.01 \quad 0.001 \quad 0.0001 \quad \dots\}$

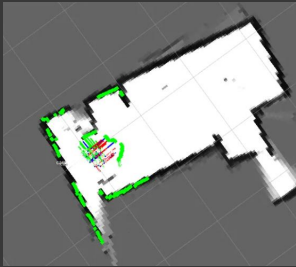
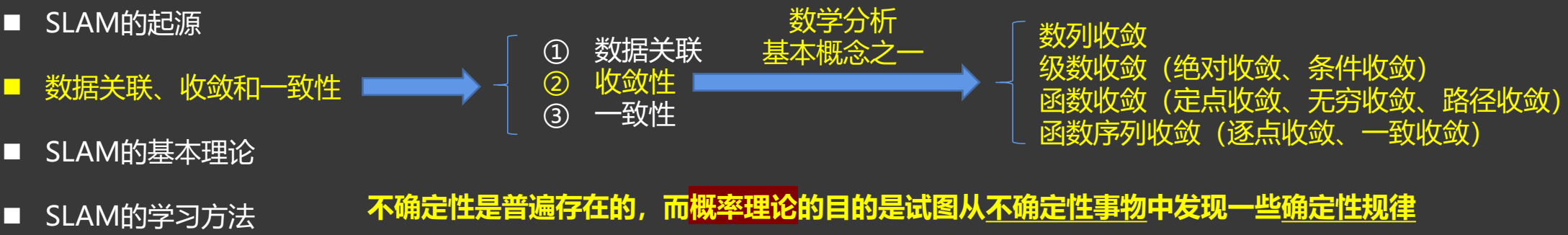
$\lim_{n \rightarrow \infty} x_n = 0$



$\lim_{x \rightarrow \infty} \sin(x) = ?$

本章讨论的SLAM问题主要是基于概率框架建模，所以这里重点讨论随机变量的收敛性。

7.1 SLAM发展简史



$$\lim_{n \rightarrow \infty} \left(\frac{T_1 + T_2 + \dots + T_n}{n} \right) \rightarrow T$$

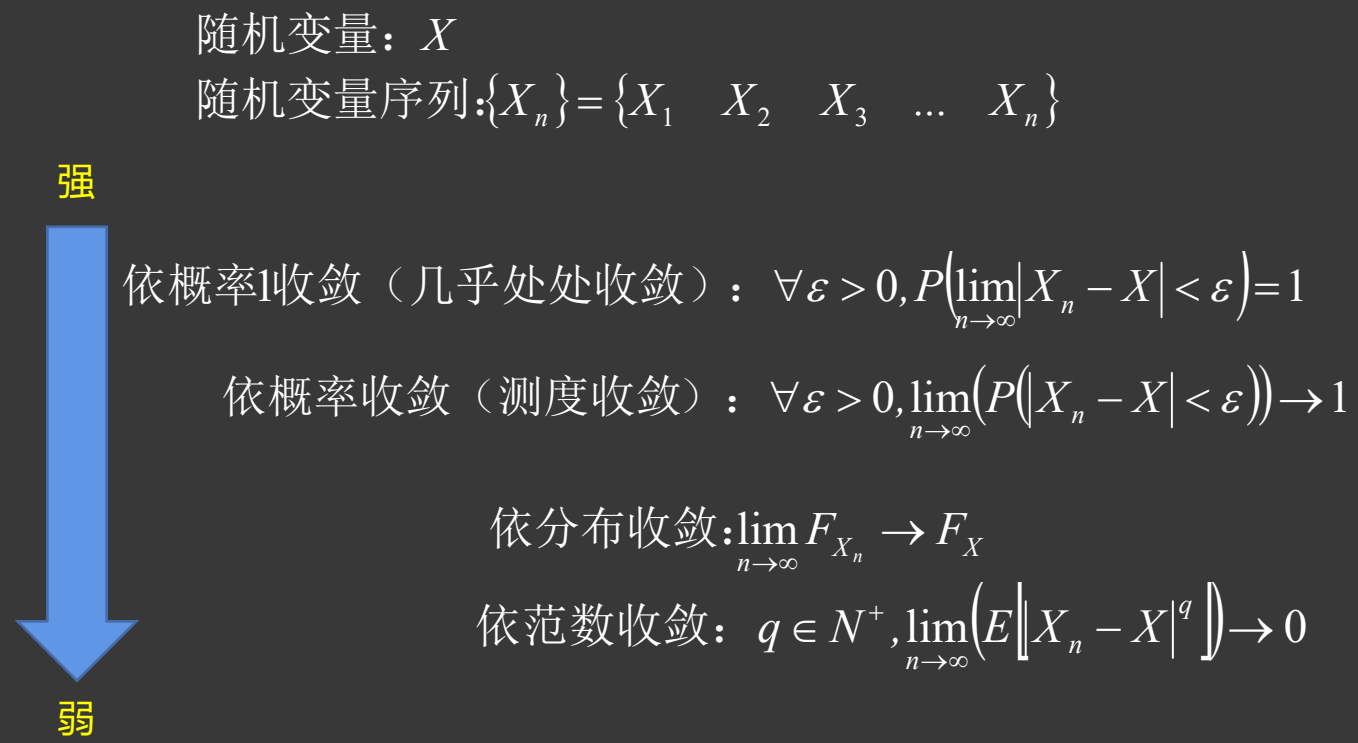
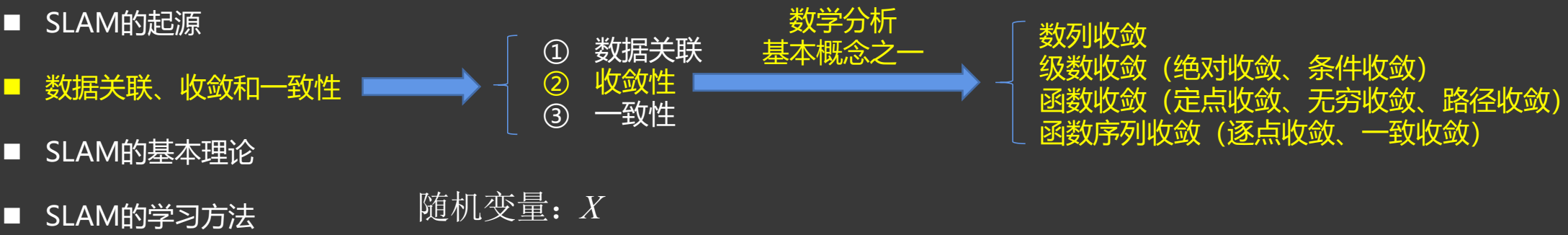
$$\lim_{n \rightarrow \infty} \left(\frac{S_1 + S_2 + \dots + S_n}{n} \right) \overset{?}{\rightarrow} S_{n+1}$$

SLAM数学模型
(数据关联过程)

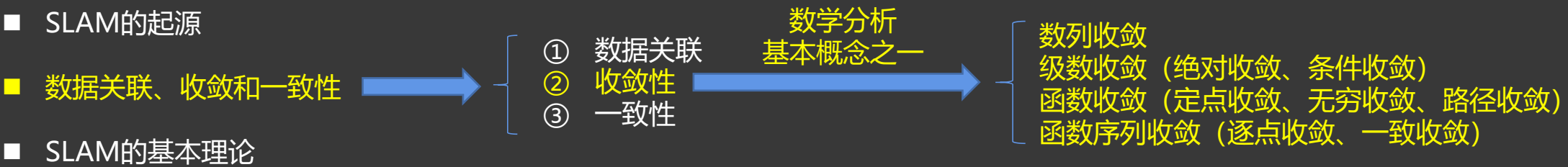
$$\lim_{n \rightarrow \infty} f(z_1, z_2, z_3, \dots, z_n) \overset{?}{\rightarrow} \langle pose, map \rangle$$

切入问题的角度（数学建模方式）对能否发现问题中的有用规律至关重要，收敛性一定程度可以作为评价指标，但也不绝对。

7.1 SLAM发展简史



7.1 SLAM发展简史



随机变量： X
随机变量序列： $\{X_n\} = \{X_1 \ X_2 \ X_3 \ \dots \ X_n\}$ (独立同分布)

弱大数定律： $E(X) \neq \infty, \forall \varepsilon > 0, \lim_{n \rightarrow \infty} P\left(\left|\frac{X_1 + X_2 + \dots + X_n}{n} - E(X)\right| < \varepsilon\right) \rightarrow 1$ (依概率收敛)

强大数定律： $E(X) \neq \infty, D(X) \neq \infty, \forall \varepsilon > 0, P\left(\lim_{n \rightarrow \infty} \left|\frac{X_1 + X_2 + \dots + X_n}{n} - E(X)\right| < \varepsilon\right) = 1$ (依概率1收敛)

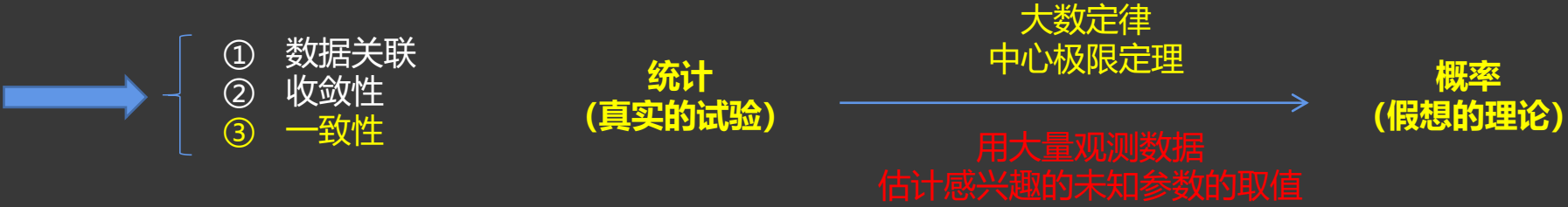
中心极限定理： $\lim_{n \rightarrow \infty} (X_1 + X_2 + \dots + X_n) \sim N(n \cdot E(X), n \cdot D(X))$

大数定律是将统计平均用来表示数学期望的依据

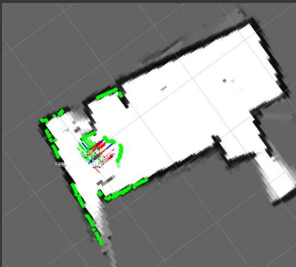
中心极限定理可以解释正态分布为何普遍存在

7.1 SLAM发展简史

- SLAM的起源
- 数据关联、收敛和一致性
- SLAM的基本理论
- SLAM的学习方法



SLAM数学模型
(数据关联过程)



$$\lim_{n \rightarrow \infty} f(z_1, z_2, z_3, \dots, z_n) \xrightarrow{?} \langle pose, map \rangle$$

- 收敛性如何?
- 收敛 or 发散
 - 一致性 (依概率1收敛、依概率收敛、依分布收敛、依范数收敛)
 - 收敛速度

7.1 SLAM发展简史

- SLAM的起源

■ 数据关联、收敛和一致性

■ SLAM的基本理论

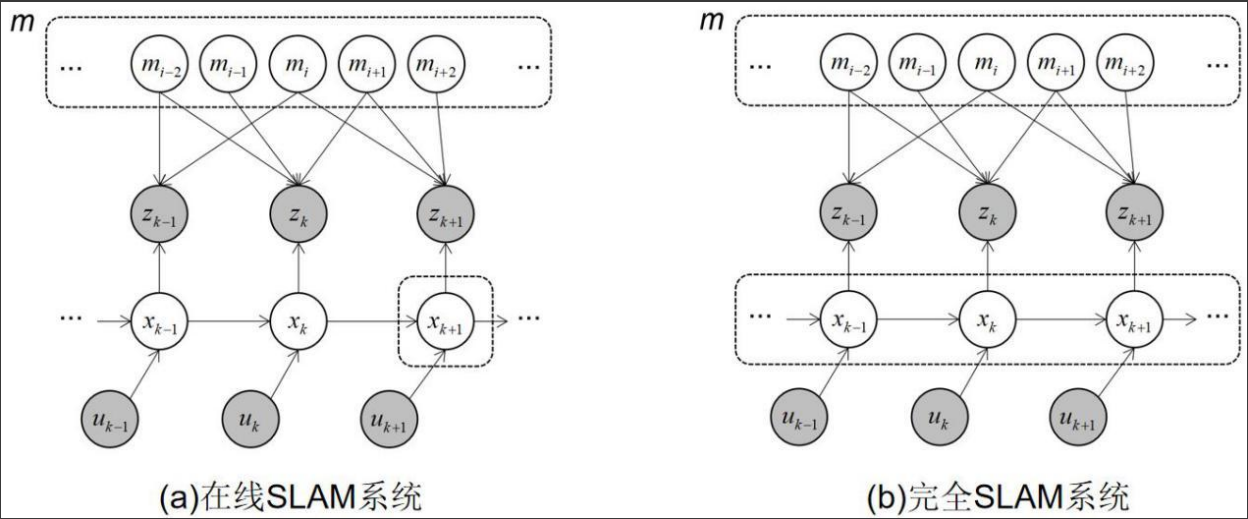
■ SLAM的学习方法
- 求解SLAM问题的方法，称为SLAM算法

① SLAM求解方法 (算法)

② SLAM求解方法 (算法) 的分类

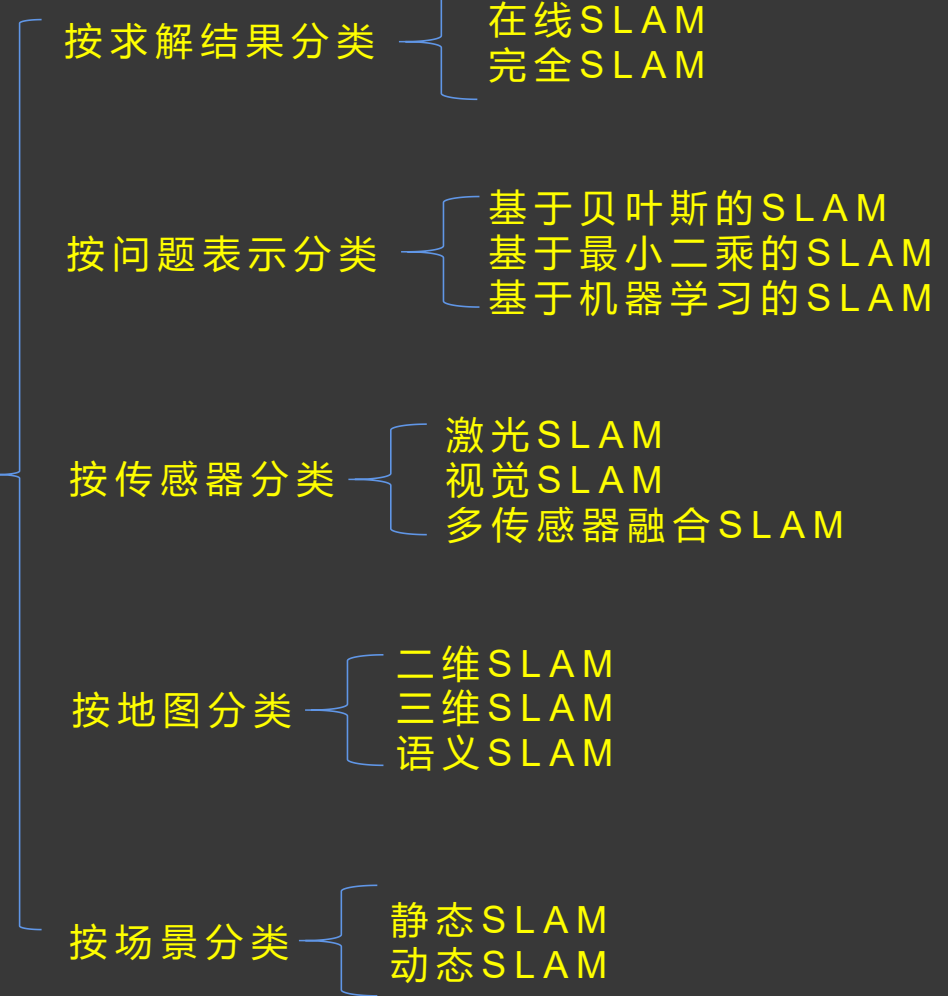
③ SLAM学习路线

		在线SLAM系统	完全SLAM系统
贝叶斯网络表示	EKF滤波	EKF-SLAM	
	粒子滤波		Fast-SLAM Gmapping
最小二乘表示	直接法		Graph-SLAM
	优化法		Cartographer LOAM ORB-SLAM
机器学习表示	端到端法	DeepVO	
	改进模块法	NeRF-SLAM NICE-SLAM	CNN-SLAM
	生物启发法	RatSLAM NeuroSLAM	



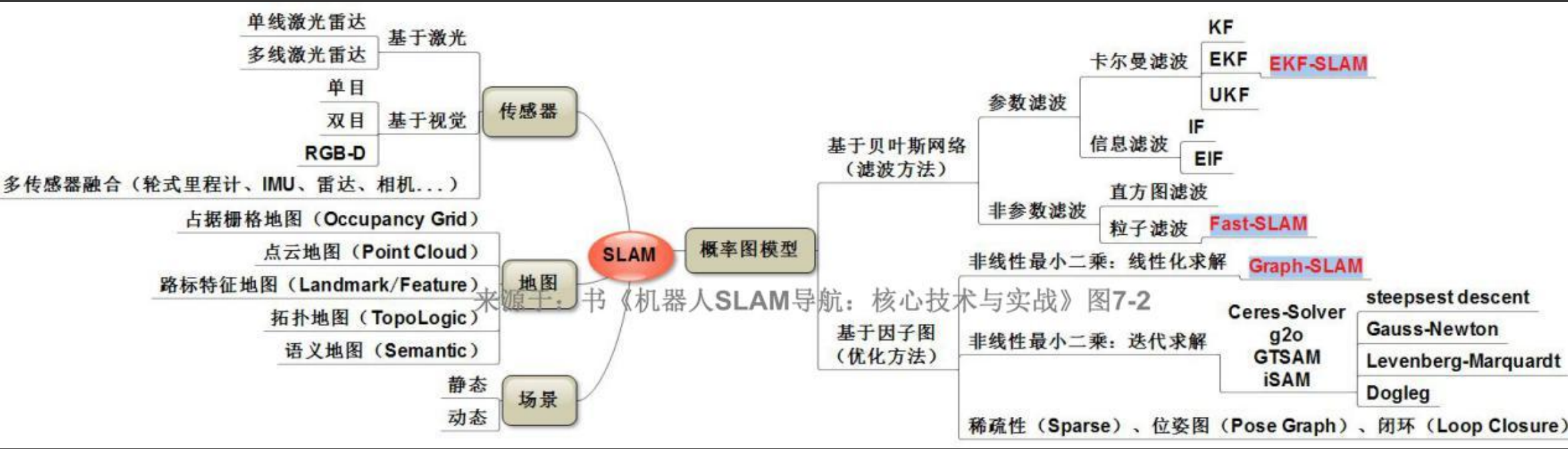
7.1 SLAM发展简史

- SLAM的起源
- 数据关联、收敛和一致性
- SLAM的基本理论
- SLAM的学习方法



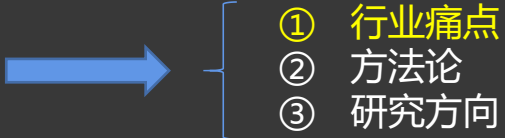
7.1 SLAM发展简史

- SLAM的起源
- 数据关联、收敛和一致性
- SLAM的基本理论
- SLAM的学习方法



来源于：书《机器人SLAM导航：核心技术与实战》图7-2

7.1 SLAM发展简史

- SLAM的起源
 - 数据关联、收敛和一致性
 - SLAM的基本理论
 - SLAM的学习方法
- 
- ① 行业痛点
② 方法论
③ 研究方向

SLAM是一个软硬件相结合，理论加实战的浩大工程性问题

企业的目的是要将**SLAM**技术真正落地到产品，
而不是等你慢慢研究数学理论或简单调几个参数，
每个从业者都应该建立起大局观以避免重复造轮子，
换句话说我们缺乏**SLAM**全栈人才，这也正是我写这本书的初衷。

7.1 SLAM发展简史

- SLAM的起源
- 数据关联、收敛和一致性
- SLAM的基本理论
- SLAM的学习方法



- ① 行业痛点
- ② 方法论
- ③ 研究方向

螺旋上升

将SLAM开源代码部署到真实机器人

针对性深入学习书籍中的理论

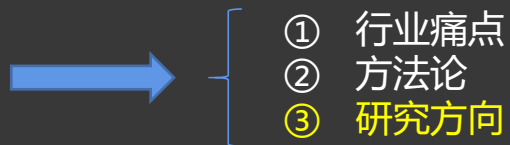
在数据集上跑跑SLAM开源代码

涉猎书籍、博客、论文

7.1 SLAM发展简史

- SLAM的起源
- 数据关联、收敛和一致性
- SLAM的基本理论

- SLAM的学习方法



巨大的信息量是获取研究思路的前提

- ① 发现新应用领域：水下SLAM应用、手动编辑SLAM地图、SLAM云计算化
- ② 软件工程的优化：去ROS化、裸机级SLAM、跨平台兼容
- ③ 改进SLAM某些功能模块：特征提取新方法、闭环检测可靠性、持久化建图机制
- ④ 多传感器融合：传感器标定、时间戳同步、去干扰数据、数据关联
- ⑤ AI+SLAM：端到端SLAM、语义SLAM、特征工程
- ⑥ 克服异常场景：玻璃障碍物、传感器盲区、光照变化、上下坡问题

内容概要

7.1 SLAM发展简史

7.2 SLAM中的概率理论

7.3 估计理论

7.4 基于贝叶斯网络的状态估计

7.5 基于因子图的状态估计

7.6 典型SLAM算法

7.7 SFM、BA和SLAM比较

7.2 SLAM中的概率理论

概率理论，是什么？

不确定性是普遍存在的，而概率理论的目的是试图从不确定性事物中发现一些确定性规律

概率理论不一定具有明确物理意义，很多时候只是借用这套理论体系以便于问题的研究

概率理论是归纳演绎的方法，第1步：统计试验，第2步：概率推理，第3步：统计预测

真实

统计试验： $P(A,B)$ 、 $P(A)$ 、 $P(B)$



虚构

概率推理： $P(A|B) = P(A,B)/P(B)$



真实

统计预测： $P(A|B) = ?$

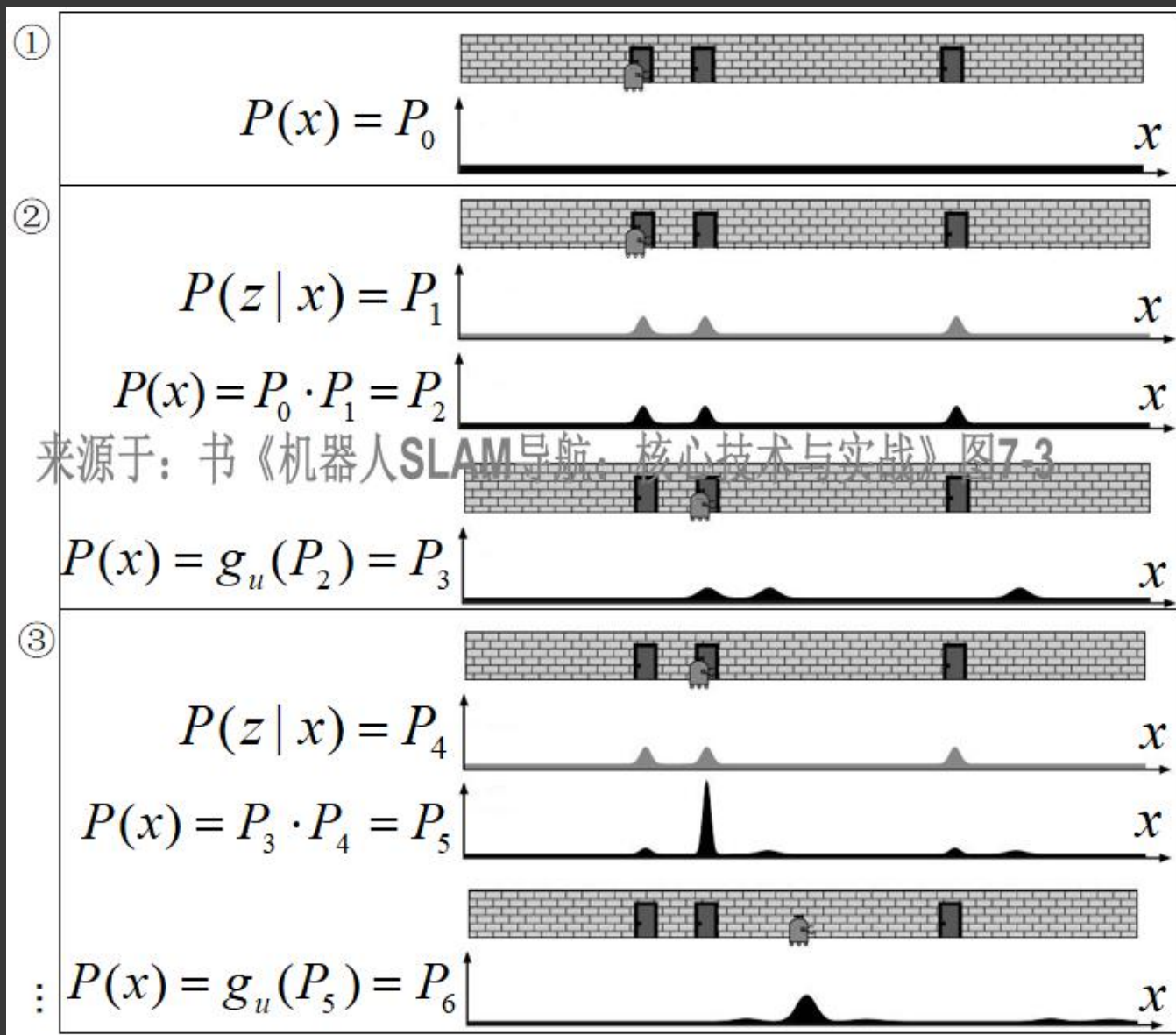
研究SLAM问题必须要用概率理论吗？

- 不一定，只是目前来看概率理论对解决SLAM问题有显著效果。
- 现在用神经网络、机器学习大模型、仿生等理论也取得了很多效果。

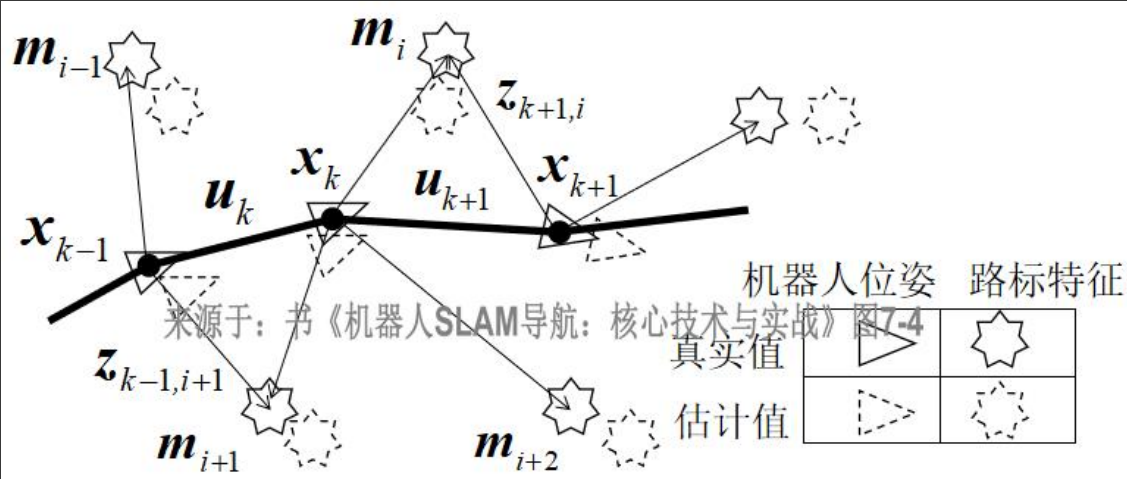
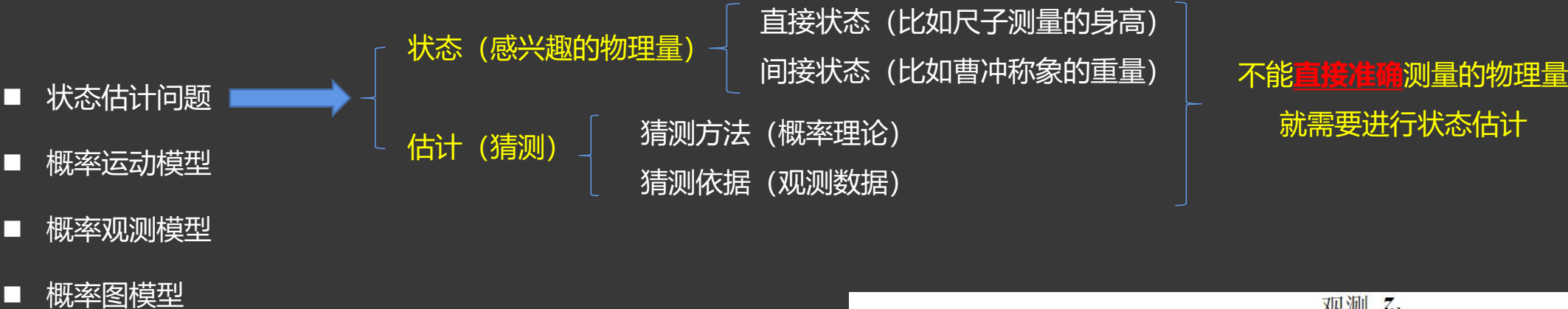
7.2 SLAM中的概率理论

机器人问题中为什么存在不确定性，原因包括传感器测量噪声、电机控制偏差、计算机软件计算精度近似等

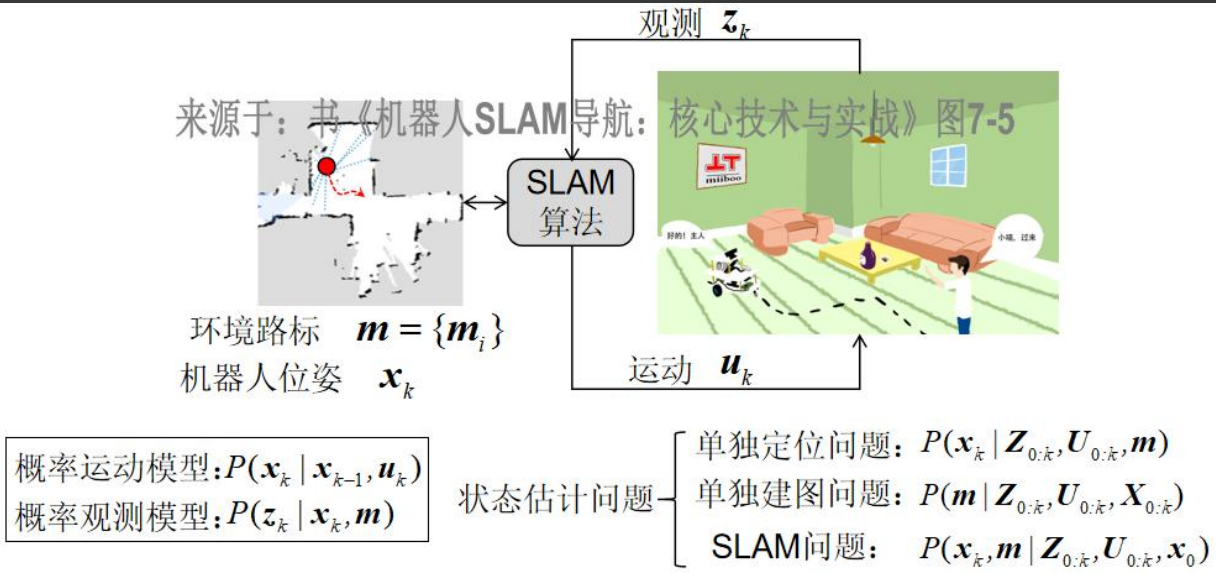
利用概率描述机器人中的不确定性，这样机器人中的不确定性就可以在概率理论框架下被计算和推演，即所谓的概率机器人学



7.2 SLAM中的概率理论



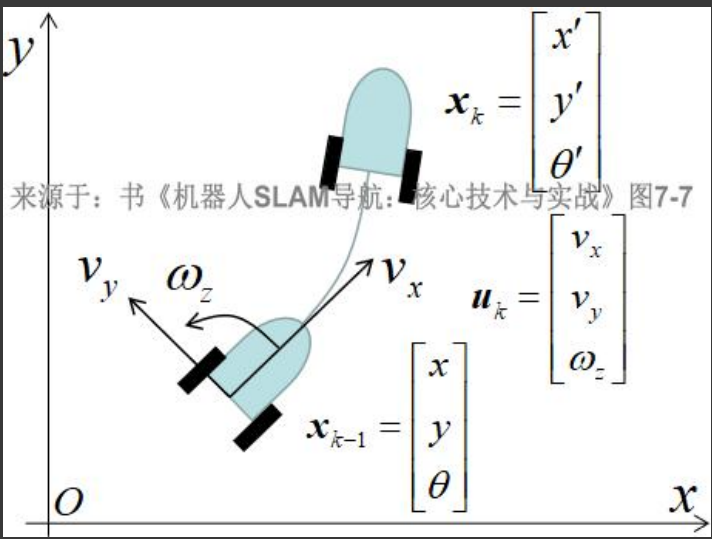
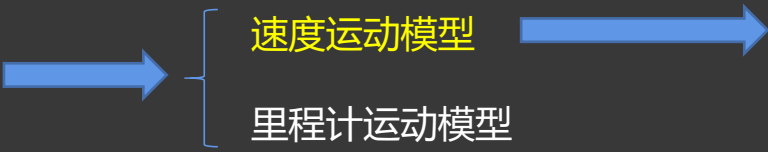
机器人SLAM中的状态估计问题



机器人SLAM中的状态估计问题的概率描述

7.2 SLAM中的概率理论

- 状态估计问题
- 概率运动模型
- 概率观测模型
- 概率图模型



来源于：书《机器人SLAM导航：核心技术与实战》图7-7

状态转移方程：

$$\mathbf{x}_k = \begin{bmatrix} x' \\ y' \\ \theta' \end{bmatrix} = g(\mathbf{x}_{k-1}, \mathbf{u}_k) = \begin{bmatrix} x \\ y \\ \theta \end{bmatrix} + \begin{bmatrix} \cos \theta & -\sin \theta & 0 \\ \sin \theta & \cos \theta & 0 \\ 0 & 0 & 1 \end{bmatrix} \begin{bmatrix} v_x \\ v_y \\ \omega_z \end{bmatrix} \Delta t$$

误差描述：

$$\begin{aligned} \mathbf{u}_k &\sim N(\bar{\mathbf{u}}_k, \Sigma_{\mathbf{u}_k}) \\ \mathbf{x}_{k-1} &\sim N(\bar{\mathbf{x}}_{k-1}, \Sigma_{\mathbf{x}_{k-1}}) \\ \mathbf{x}_k &\sim N(\bar{\mathbf{x}}_k, \Sigma_{\mathbf{x}_k}) \end{aligned}$$

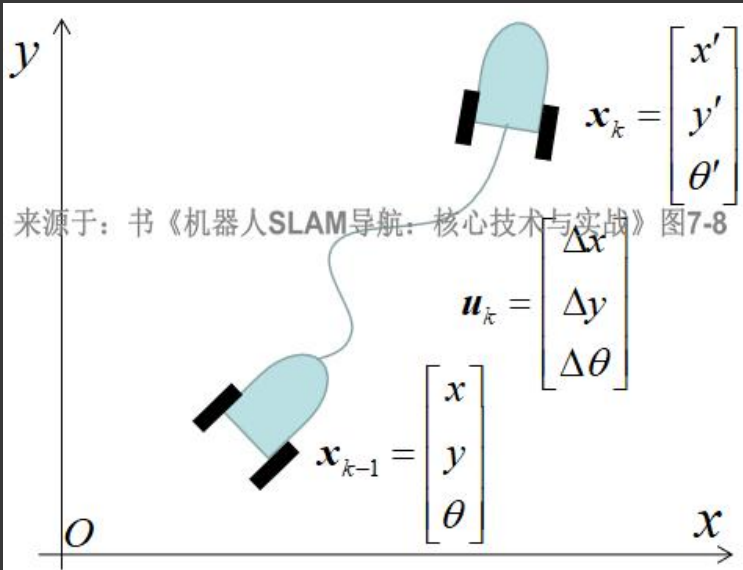
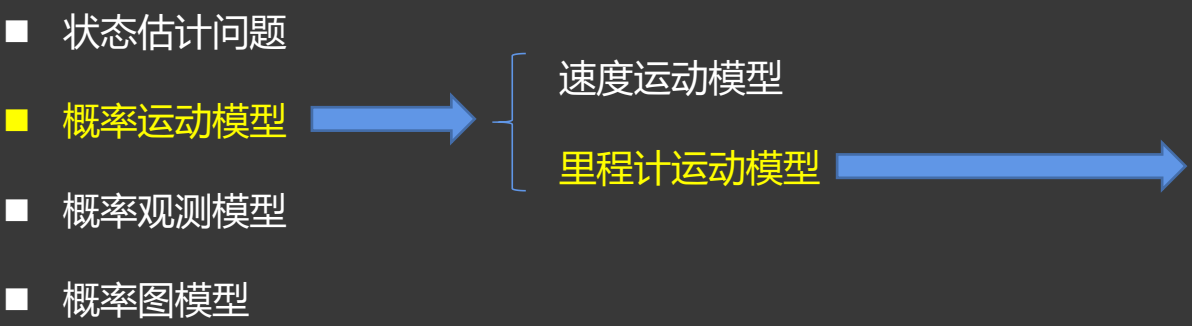
在没有先验信息的情况，通常
将误差假设为高斯分布

$$\Sigma_{\mathbf{u}_k} = \begin{bmatrix} \sigma_{v_x}^2 & 0 & 0 \\ 0 & \sigma_{v_y}^2 & 0 \\ 0 & 0 & \sigma_{\omega_z}^2 \end{bmatrix}$$

$$\Sigma_{\mathbf{x}_k} = \begin{bmatrix} \frac{\partial g}{\partial \mathbf{x}_{k-1}} & \frac{\partial g}{\partial \mathbf{u}_k} \end{bmatrix} \begin{bmatrix} \Sigma_{\mathbf{x}_{k-1}} & \mathbf{0}_{3 \times 3} \\ \mathbf{0}_{3 \times 3} & \Sigma_{\mathbf{u}_k} \end{bmatrix} \begin{bmatrix} \frac{\partial g}{\partial \mathbf{x}_{k-1}} & \frac{\partial g}{\partial \mathbf{u}_k} \end{bmatrix}^T$$

(其中, $\begin{bmatrix} \frac{\partial g}{\partial \mathbf{x}_{k-1}} & \frac{\partial g}{\partial \mathbf{u}_k} \end{bmatrix}$ 为函数g的雅克比矩阵)

7.2 SLAM中的概率理论



状态转移方程：

$$\mathbf{x}_k = \begin{bmatrix} x' \\ y' \\ \theta' \end{bmatrix} = g(\mathbf{x}_{k-1}, \mathbf{u}_k) = \begin{bmatrix} x \\ y \\ \theta \end{bmatrix} + \begin{bmatrix} \Delta x \\ \Delta y \\ \Delta \theta \end{bmatrix}$$

误差描述：

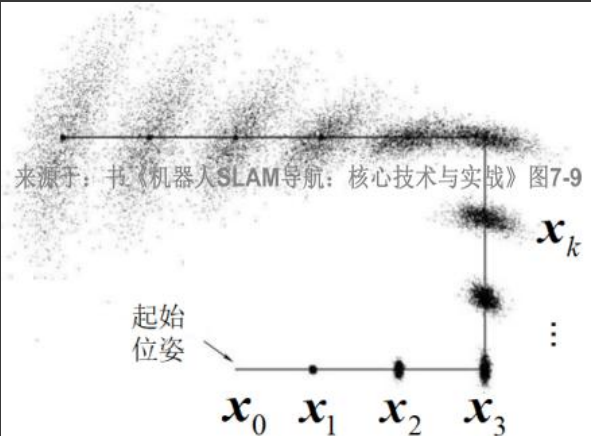
$$\mathbf{u}_k \sim N(\bar{\mathbf{u}}_k, \Sigma_{\mathbf{u}_k})$$
$$\mathbf{x}_{k-1} \sim N(\bar{\mathbf{x}}_{k-1}, \Sigma_{\mathbf{x}_{k-1}})$$
$$\mathbf{x}_k \sim N(\bar{\mathbf{x}}_k, \Sigma_{\mathbf{x}_k})$$

在没有先验信息的情况，通常
将误差假设为高斯分布

$$\Sigma_{\mathbf{u}_k} = \begin{bmatrix} \sigma_{\Delta x}^2 & 0 & 0 \\ 0 & \sigma_{\Delta y}^2 & 0 \\ 0 & 0 & \sigma_{\Delta \theta}^2 \end{bmatrix}$$
$$\begin{aligned} \sigma_{\Delta x}^2 &= \xi_{\Delta xy} + \alpha_1 \sqrt{(\Delta x)^2 + (\Delta y)^2} + \alpha_2 |\Delta \theta| \\ \sigma_{\Delta y}^2 &= \xi_{\Delta xy} + \alpha_1 \sqrt{(\Delta x)^2 + (\Delta y)^2} + \alpha_2 |\Delta \theta| \\ \sigma_{\Delta \theta}^2 &= \xi_{\Delta \theta} + \alpha_3 \sqrt{(\Delta x)^2 + (\Delta y)^2} + \alpha_4 |\Delta \theta| \end{aligned}$$

$$\Sigma_{\mathbf{x}_k} = \begin{bmatrix} \frac{\partial g}{\partial \mathbf{x}_{k-1}} & \frac{\partial g}{\partial \mathbf{u}_k} \end{bmatrix} \begin{bmatrix} \Sigma_{\mathbf{x}_{k-1}} & \mathbf{0}_{3 \times 3} \\ \mathbf{0}_{3 \times 3} & \Sigma_{\mathbf{u}_k} \end{bmatrix} \begin{bmatrix} \frac{\partial g}{\partial \mathbf{x}_{k-1}} \\ \frac{\partial g}{\partial \mathbf{u}_k} \end{bmatrix}^T$$

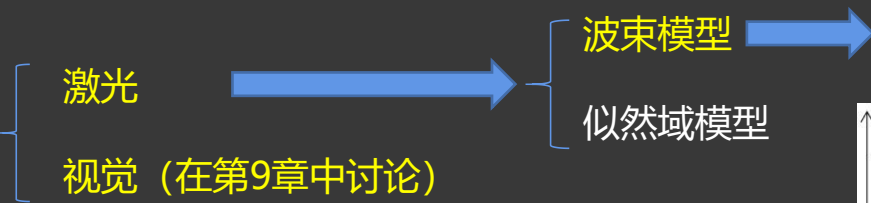
(其中, $\begin{bmatrix} \frac{\partial g}{\partial \mathbf{x}_{k-1}} & \frac{\partial g}{\partial \mathbf{u}_k} \end{bmatrix}$ 为函数g的雅克比矩阵)



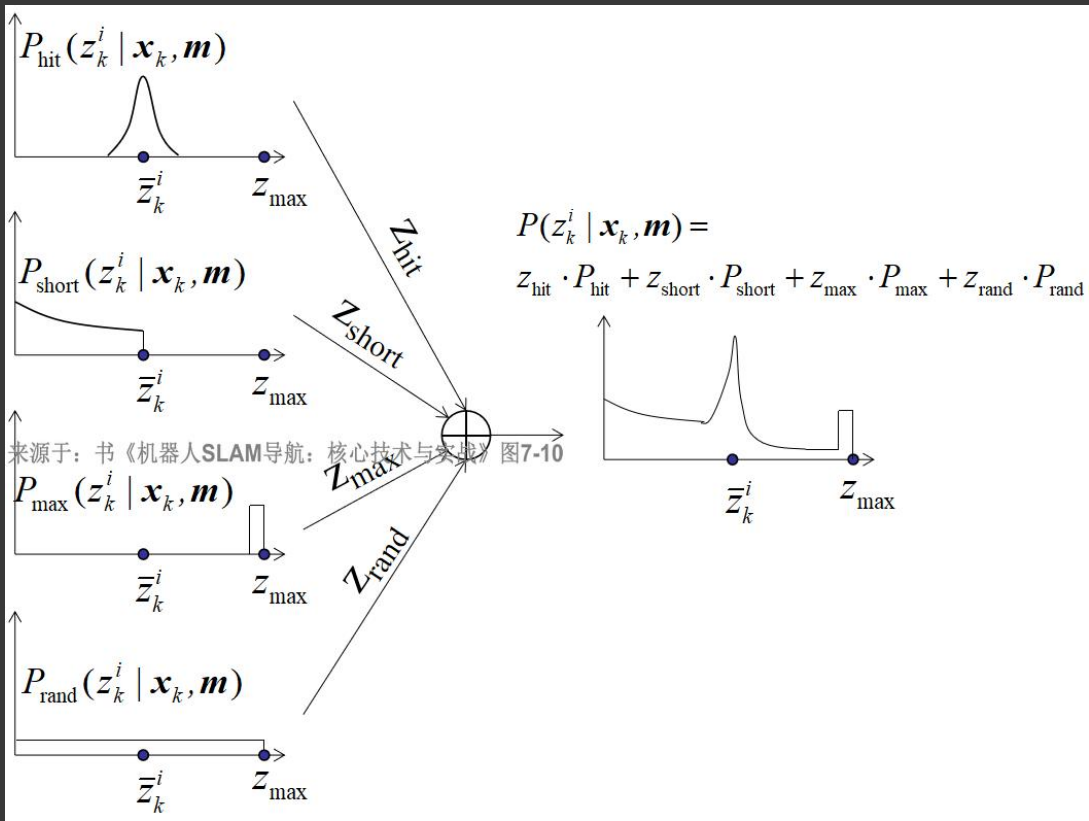
机器人位姿不确定性随运动的关系

7.2 SLAM中的概率理论

- 状态估计问题
- 概率运动模型
- 概率观测模型
- 概率图模型

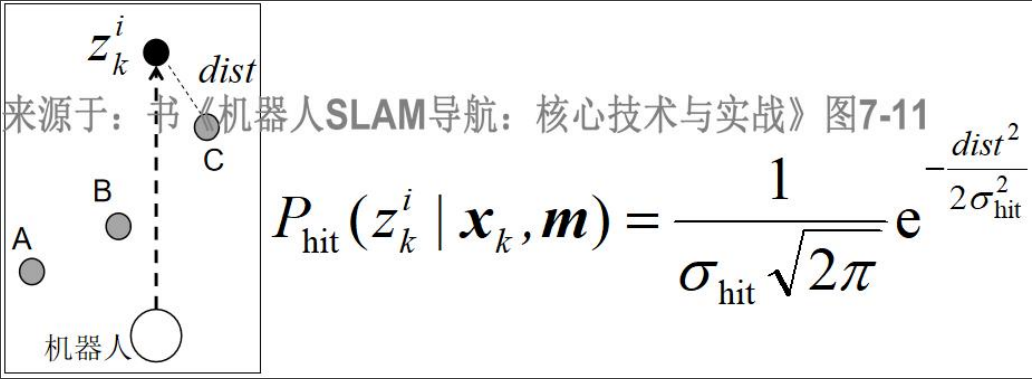
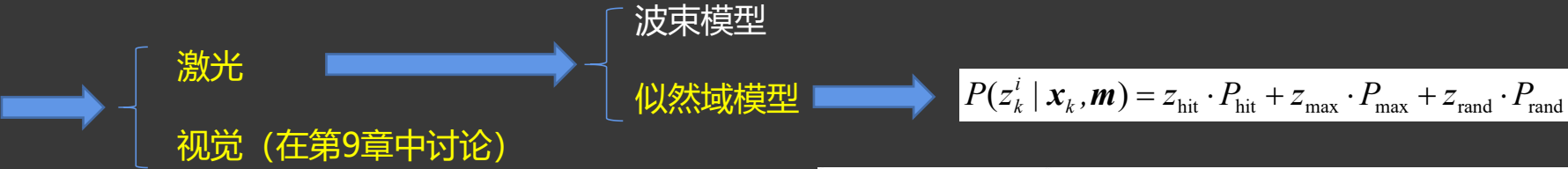


$$P(z_k^i | \mathbf{x}_k, \mathbf{m}) = z_{\text{hit}} \cdot P_{\text{hit}} + z_{\text{short}} \cdot P_{\text{short}} + z_{\text{max}} \cdot P_{\text{max}} + z_{\text{rand}} \cdot P_{\text{rand}}$$



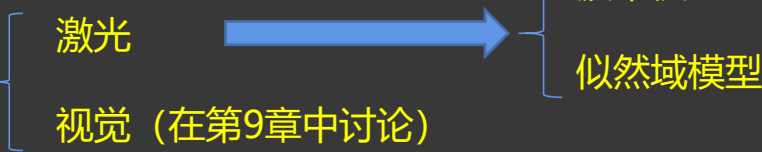
7.2 SLAM中的概率理论

- 状态估计问题
- 概率运动模型
- 概率观测模型
- 概率图模型



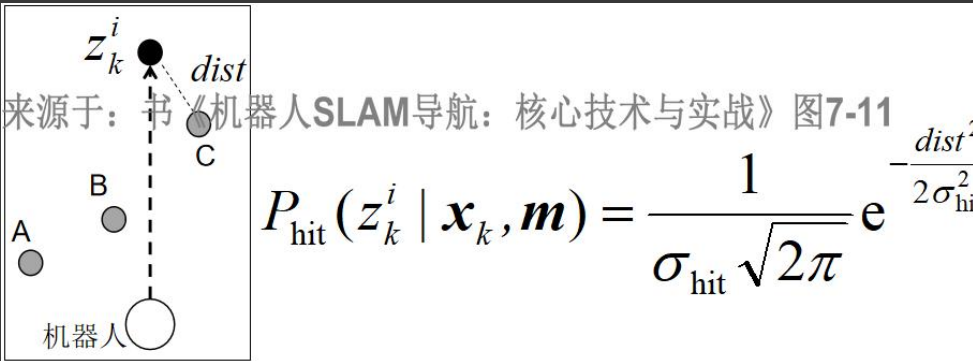
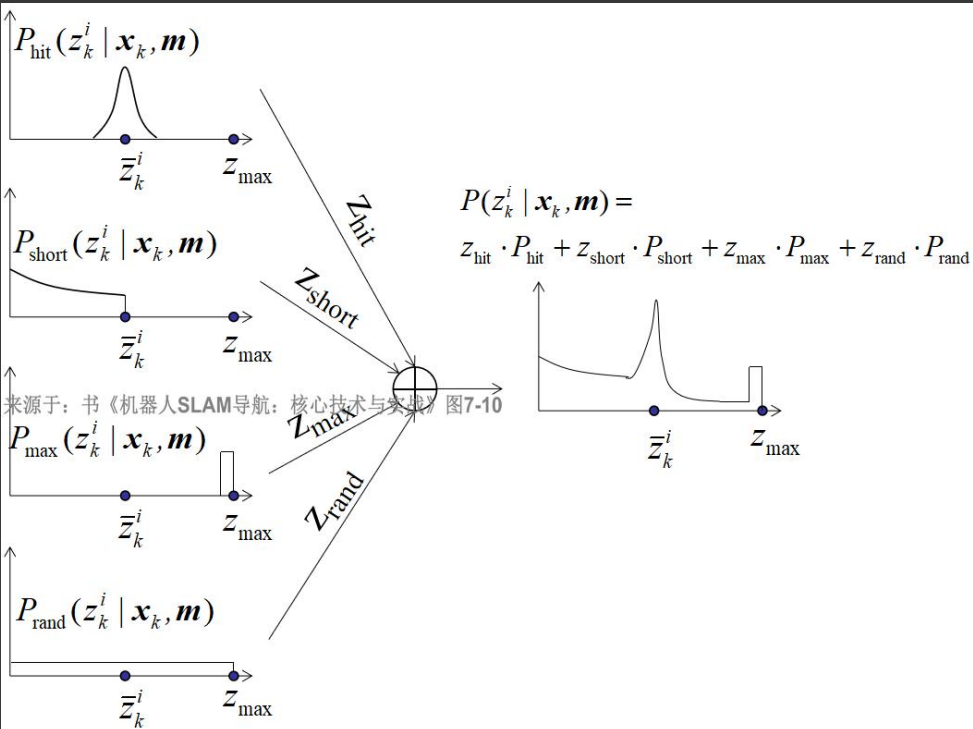
7.2 SLAM中的概率理论

- 状态估计问题
- 概率运动模型
- 概率观测模型
- 概率图模型



波束模型缺乏光滑性，导致构建出来的地图具有很多毛刺。

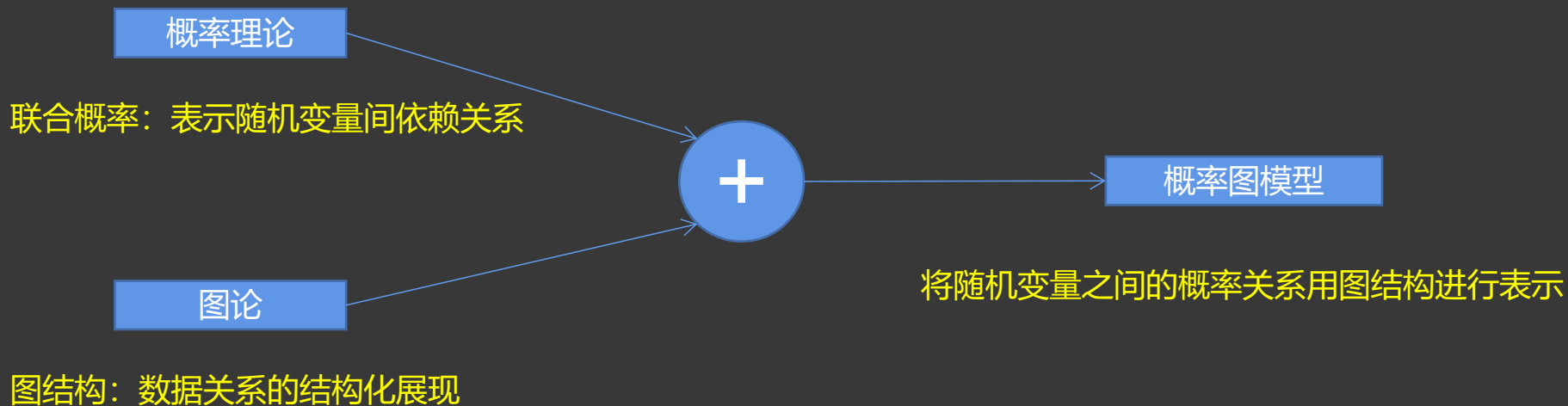
似然域模型较为平滑，构建出来的地图不易出现毛刺。



7.2 SLAM中的概率理论

- 状态估计问题
- 概率运动模型
- 概率观测模型

■ 概率图模型 → 当有了概率运动模型和概率观测模型后，怎么描述运动模型与观测模型之间的概率关系呢？



7.2 SLAM中的概率理论

- 状态估计问题
- 概率运动模型
- 概率观测模型
- 概率图模型

连续随机变量 VS 离散随机变量
(比如观测温度) (比如掷骰子)

随机变量集合: $X = \{A, B, C, D\}$

边缘分布: $P(A) = \sum_D \sum_C \sum_B P(A, B, C, D)$
(概率加法)

联合分布: $P(A, B, C, D) = P(A, B, C)P(D | A, B, C)$
(概率乘法)
 $= P(A, B)P(C | A, B)P(D | A, B, C)$
 $= P(A)P(B | A)P(C | A, B)P(D | A, B, C)$

贝叶斯准则: $P(B | A) = \frac{P(A | B)P(B)}{P(A)}$

概率分布的展开形式，可以简化吗？

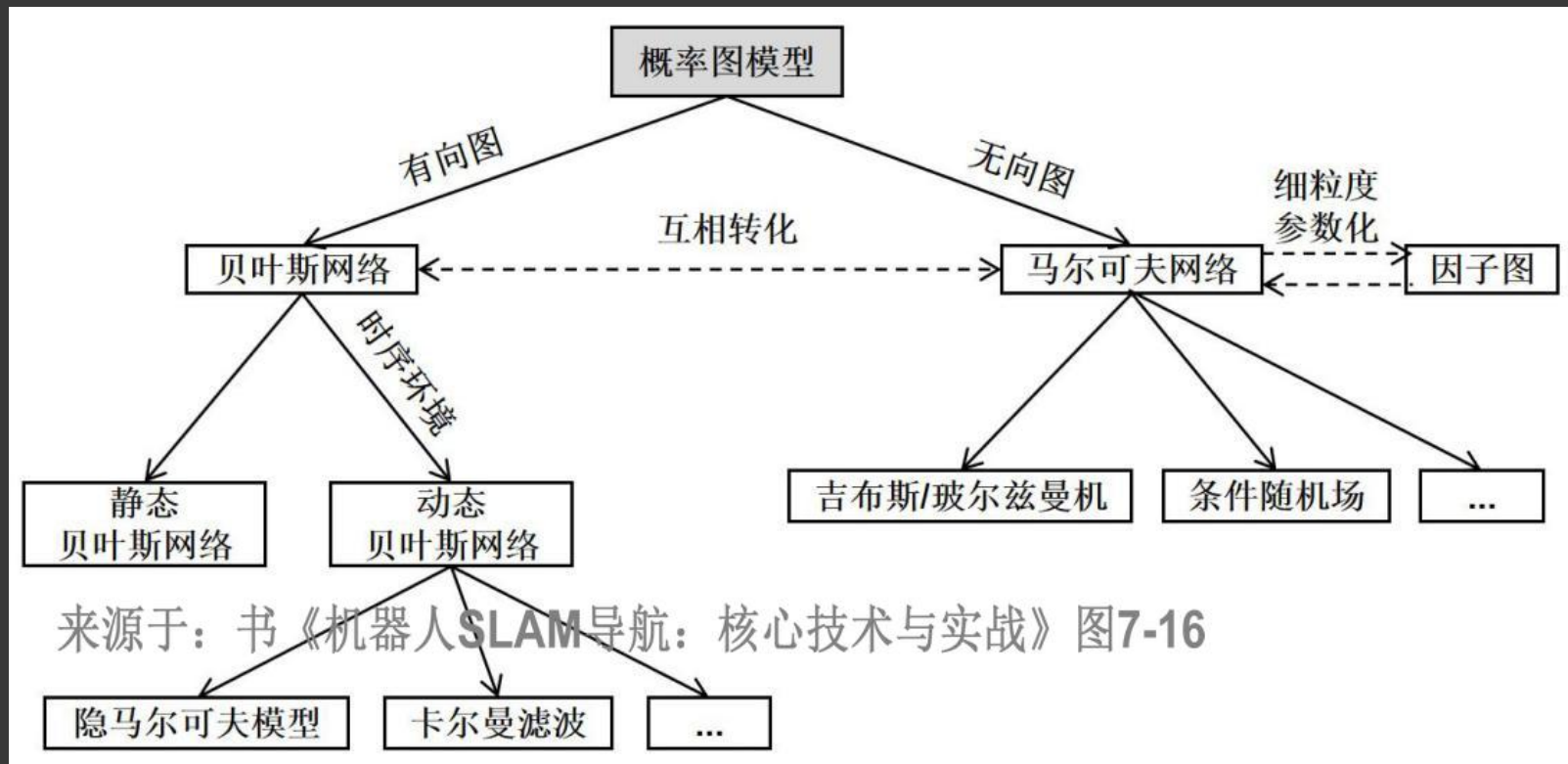
利用图论结构化的数据表示方式，可以让随机变量之间的依赖关系变得更为直观。

一般形式

$$\begin{aligned} P(B | A, C) &= \frac{P(B, A, C)}{P(A, C)} \\ &= \frac{P(A | B, C)P(B, C)}{P(A, C)} = \frac{P(A | B, C)P(B | C)P(C)}{P(A | C)P(C)} = \frac{P(A | B, C)P(B | C)}{P(A | C)} \end{aligned}$$

7.2 SLAM中的概率理论

- 状态估计问题
- 概率运动模型
- 概率观测模型
- 概率图模型



7.2 SLAM中的概率理论

- 状态估计问题
- 概率运动模型
- 概率观测模型

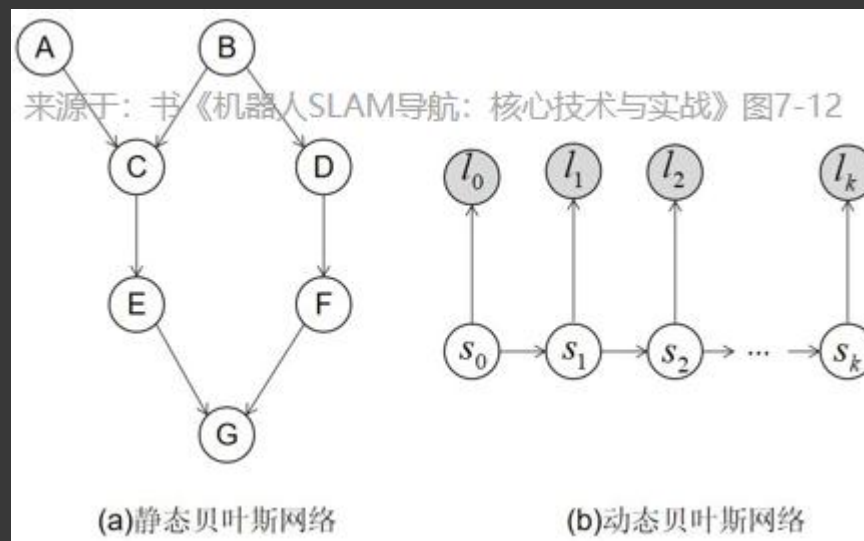
- 概率图模型



贝叶斯网络

马尔可夫网络

因子图



箭头表示随机变量之间依赖的因果关系

条件独立性

联合分布化简形式:

$$P(A, B, C, D, E, F, G) = P(A)P(B)P(C | A, B)P(D | B)P(E | C)P(F | D)P(G | E, F)$$

联合分布化简一般形式:

$$P(X_1, X_2, \dots, X_n) = \prod_{i=1}^n P(X_i | Pa_{X_i})$$

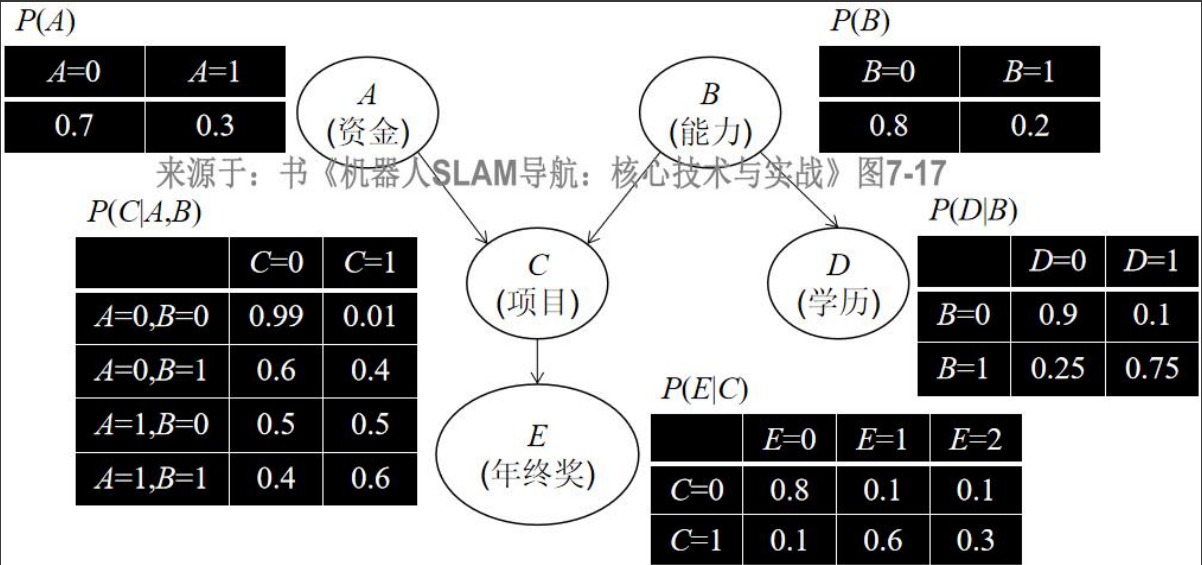
7.2 SLAM中的概率理论

利用贝叶斯网络进行SLAM问题的表示及推理过程

- 状态估计问题
 - 概率运动模型
 - 概率观测模型
 - 概率图模型
- 贝叶斯网络

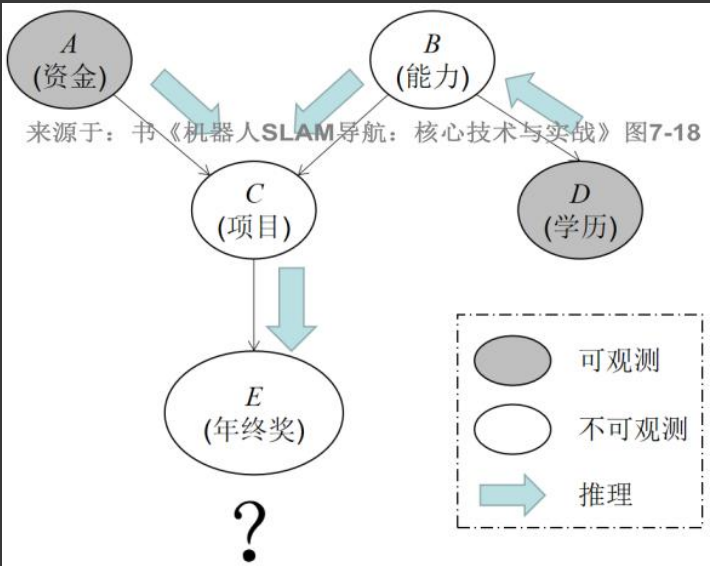
马尔可夫网络

因子图



静态贝叶斯网络【表示】

- 【学习】
- 模型结构
 - 模型参数



静态贝叶斯网络【推理】

- 给定观测
- 先验信息 (模型结构、模型参数)
- 预测结果

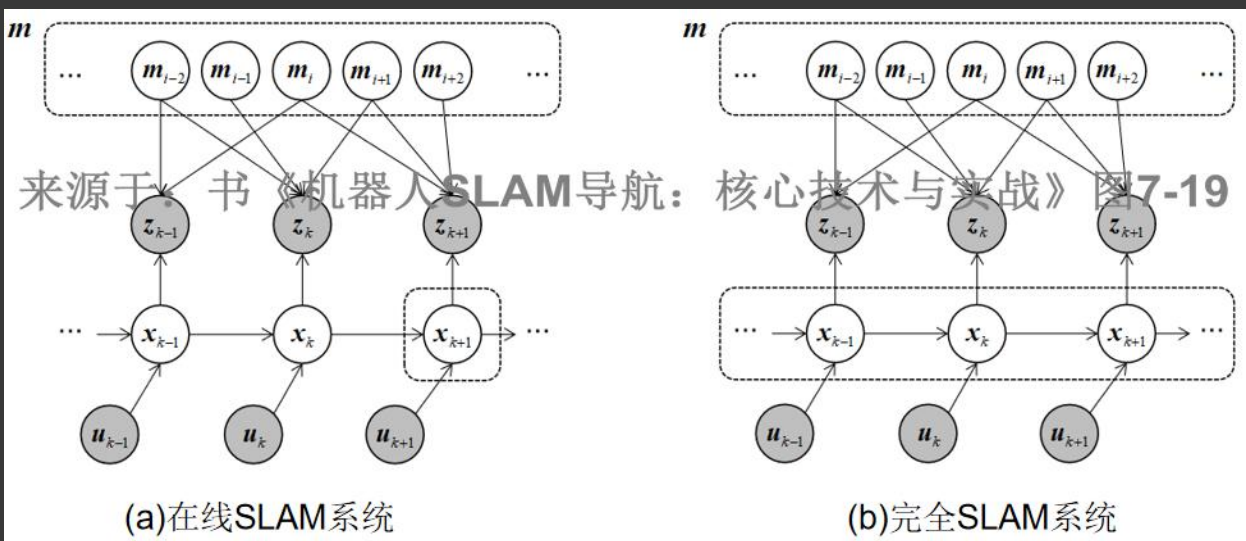
7.2 SLAM中的概率理论

利用贝叶斯网络进行SLAM问题的表示及推理过程

- 状态估计问题
- 概率运动模型
- 概率观测模型
- 概率图模型

贝叶斯网络
马尔可夫网络
因子图

动态贝叶斯网络【对SLAM问题的表示】



- 【学习】
- 模型结构
 - 模型参数

动态贝叶斯网络【对SLAM问题的推理】

在线SLAM: $P(x_k, m | Z_{0:k}, U_{0:k}, x_0)$

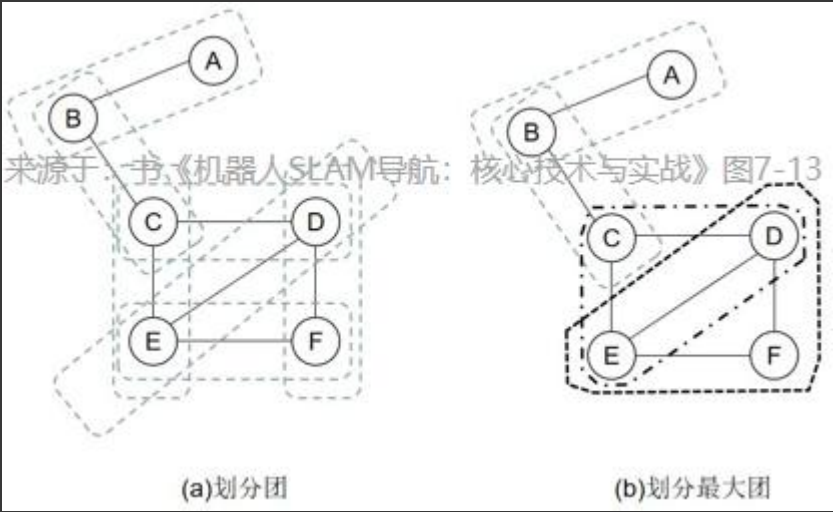
完全SLAM: $P(X_{0:k}, m | Z_{0:k}, U_{0:k}, x_0)$

- 给定观测
- 先验信息 (模型结构、模型参数)
- 预测结果

7.2 SLAM中的概率理论

- 状态估计问题
 - 概率运动模型
 - 概率观测模型
 - 概率图模型
- ➡

贝叶斯网络
马尔可夫网络
因子图



连线表示随机变量之间隐式相关性，可表示循环依赖

团中的节点必须两两之间都有连接，团的势能函数是一个抽象函数，需要在实际问题中被具体化为特定形式

不能再往节点中加入节点的团，就是最大团

条件独立性

联合分布化简形式：

$$P(A,B,C,D,E,F) = \frac{1}{Z} \psi_{Q_1}(A,B) \psi_{Q_2}(B,C) \psi_{Q_3}(C,D) \psi_{Q_4}(C,E) \psi_{Q_5}(D,E) \psi_{Q_6}(D,F) \psi_{Q_7}(E,F)$$

联合分布最大团化简形式：

$$P(A,B,C,D,E,F) = \frac{1}{Z} \psi_{Q_1}(A,B) \psi_{Q_2}(B,C) \psi_{Q_8}(C,D,E) \psi_{Q_9}(D,E,F)$$

联合分布最大团化简一般形式：

$$P(X_1,X_2,...,X_n) = \frac{1}{Z} \prod_{i=1}^m \psi_{Q_i}(Q_i)$$

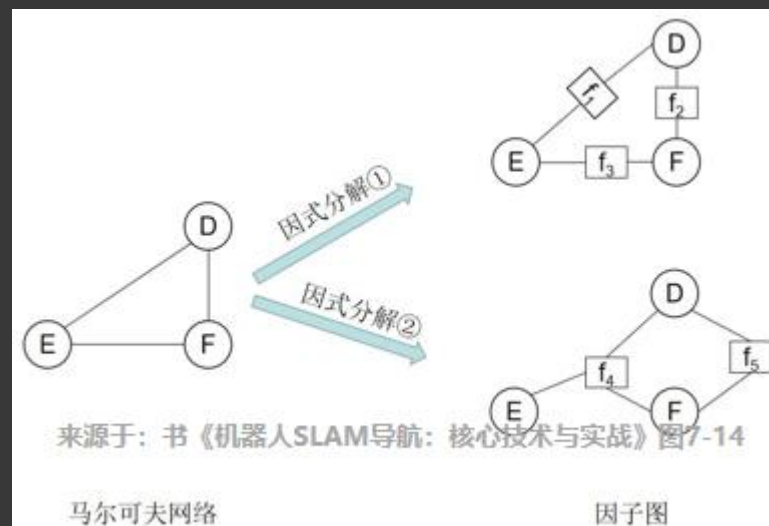
7.2 SLAM中的概率理论

- 状态估计问题
- 概率运动模型
- 概率观测模型

- 概率图模型



贝叶斯网络
马尔可夫网络
因子图



马尔可夫网络中团的势能函数只能笼统表示团内节点之间的相关性

将团的势能函数进行因式分解，分解出来的因子项对随机变量之间的关系描述更具体

因式分解不唯一，所以马尔可夫网络到因子图的转换也不唯一

联合分布因子项化简形式1：

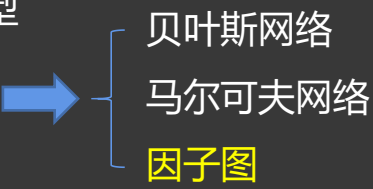
$$\psi_{Q_9}(D, E, F) = f_1(D, E) \cdot f_2(D, F) \cdot f_3(E, F)$$

联合分布因子项化简形式2：

$$\psi_{Q_9}(D, E, F) = f_4(D, E, F) \cdot f_5(D, F)$$

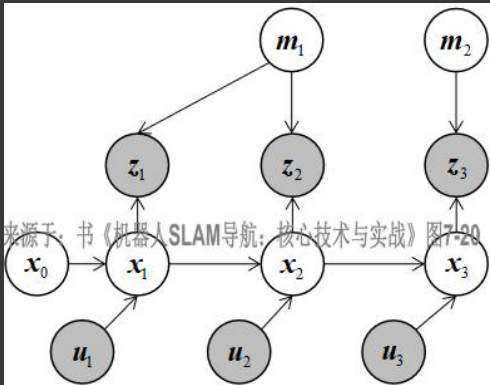
7.2 SLAM中的概率理论

- 状态估计问题
- 概率运动模型
- 概率观测模型
- 概率图模型

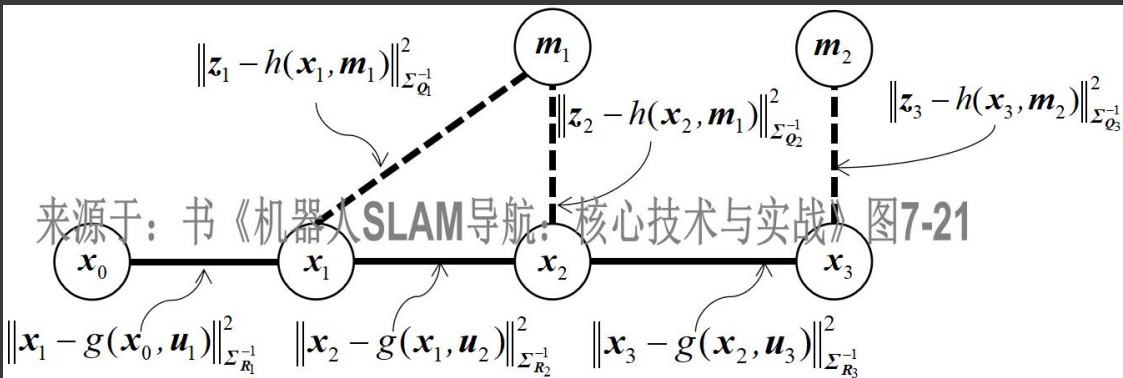


将贝叶斯网络箭头的条件概率
替换成因子图的最小二乘约束边

利用因子图进行SLAM问题的表示及推理过程



贝叶斯网络表示的SLAM问题



来源于：书《机器人SLAM导航：核心技术与实战》图7-21

最小二乘约束形式因子图表示的SLAM问题

设：

机器人位姿状态：	$X = \{x_0, x_1, x_2, x_3\}$
路标：	$m = \{m_1, m_2\}$
运动量：	$U = \{u_1, u_2, u_3\}$
观测量：	$Z = \{z_1, z_2, z_3\}$

SLAM问题表示（联合分布）：

$$P(X, m, Z, U) = P(m_1)P(m_2)P(u_0)P(u_1)P(u_2) \} \text{先验, 可忽略} \\ \times P(x_0)P(x_1 | x_0, u_1)P(x_2 | x_1, u_2)P(x_3 | x_2, u_3) \} \text{运动} \\ \times P(z_1 | x_1, m_1)P(z_2 | x_2, m_1)P(z_3 | x_3, m_2) \} \text{观测}$$

SLAM问题推理（状态估计）：

（利用给定观测状态估计未知状态，比如最大后验估计）

$$S_{\text{MAP}} = \arg \max_S P(S | V)$$

$$P(S | V) \\ \propto P(x_0)P(x_1 | x_0, u_1)P(x_2 | x_1, u_2)P(x_3 | x_2, u_3) \} \text{运动} \\ \times P(z_1 | x_1, m_1)P(z_2 | x_2, m_1)P(z_3 | x_3, m_2) \} \text{观测}$$

可直接量测信息：	$V = \{Z, U\}$
不可直接量测信息：	$S = \{X, m\}$
	（待估计量）

在获得直接量测信息U和Z后,利用条件概率分布关系对不可量测信息X和m进行推理，这就是SLAM问题的状态估计过程。

$$P(S | V) \propto \psi(S) = \prod_i \psi_i(X, m, Z, U)$$

由于估计过程关心的是X和m，所以可以用运动和观测这两个约束来构建因子函数，并将不关心的U和Z隐藏到因子函数约束之中。

最小二乘估计

$$S_{\text{MAP}} = \arg \max_S \psi(S) \\ \propto \arg \max_S (\sum_{i1} \ln \psi_{i1}(x_k, x_{k-1}) + \sum_{i2} \ln \psi_{i2}(x_k, m_j)) \\ \propto \arg \max_S (-\frac{1}{2} \sum_{i1} (x_k - g(x_{k-1}, u_k))^T \Sigma_{R_k}^{-1} (x_k - g(x_{k-1}, u_k)) \\ - \frac{1}{2} \sum_{i2} (z_k - h(x_k, m_j))^T \Sigma_{Q_k}^{-1} (z_k - h(x_k, m_j))) \\ \propto \arg \min_S (\sum_{i1} \|x_k - g(x_{k-1}, u_k)\|_{\Sigma_{R_k}^{-1}}^2 + \sum_{i2} \|z_k - h(x_k, m_j)\|_{\Sigma_{Q_k}^{-1}}^2)$$

- 假设运动噪声和观测噪声都是0均值高斯分布
- 用运动误差和观测误差分别来构造运动约束和观测约束这2种因子项
- 误差项采用分布的协方差矩阵加权求和

内容概要

7.1 SLAM发展简史

7.2 SLAM中的概率理论

7.3 估计理论

7.4 基于贝叶斯网络的状态估计

7.5 基于因子图的状态估计

7.6 典型SLAM算法

7.7 SFM、BA和SLAM比较

7.3 估计理论

不管是用贝叶斯网络还是因子图，一旦SLAM问题用概率图模型得到表示后，接下来就是利用可观测量推理不可观测量，也就是说SLAM问题的求解过程是一个状态估计。

7.3 估计理论

所谓估计，就是研究某问题时感兴趣的参数不能够通过精确测量得知，只能通过一组观测样本值来猜测参数的可能取值。

有时候估计量比较直观，比如估计室内的温度，感兴趣的参数 θ 就是温度；

有时候估计量并不直观，比如估计机械零件的精度，感兴趣的参数 θ 是误差分布函数 f 的特征量。

- 估计量的性质
- 估计量的构建
- 各估计量比较

估计既然是一种猜测行为，可以毫无根据的瞎猜，也可以依据严密的逻辑策略进行科学的猜测。

这就涉及到如何评价估计量的好坏程度，借助估计量的性质可以对估计好坏程度进行评价。

①一致性：

强

↓

弱

$\hat{\theta} \xrightarrow[Z=\{z_k\}|k \rightarrow \infty]{} \theta$

$\hat{\theta} \xrightarrow[Z=\{z_k\}|k \rightarrow \infty]{P} \theta$

$\hat{\theta} \xrightarrow[Z=\{z_k\}|k \rightarrow \infty]{d} \theta$

$\hat{\theta} \xrightarrow[Z=\{z_k\}|k \rightarrow \infty]{q} \theta$

依概率1收敛（几乎处处收敛）

依概率收敛（测度收敛）

依分布收敛

依范数收敛

②偏差性：

$E[\hat{\theta}^k]$

$E[(\hat{\theta} - E[\hat{\theta}])^k]$

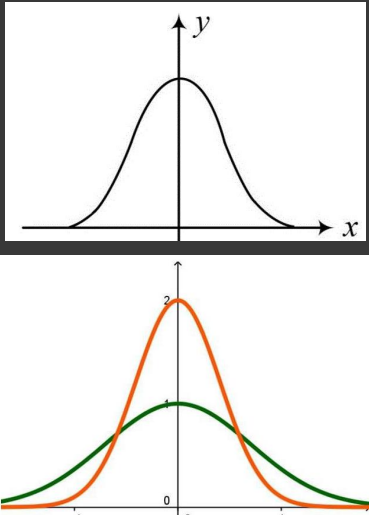
k=1，一阶原点矩（期望）

k=2，二阶中心矩（方差）

无偏估计？
 $E[\hat{\theta}] = \theta$

最小方差估计？

无偏估计只是保证估计量在期望上是正确的，估计量本身还是有不确定性，方差描述了估计量的不确定性，方差越小估计不确定性越小



7.3 估计理论

好的估计量应该是**最小方差无偏估计** (Minimum Variance Unbiased Estimation, MVUE) ,

然而这会非常困难, 因此需要寻找**近似方法**

- 估计量的性质
- 估计量的构建
- 各估计量比较

最大似然估计 (MLE) :

$$\begin{aligned}\hat{\theta}_{\text{MLE}} &= \arg \max_{\theta} l(\theta | Z) \\ &= \arg \max_{\theta} \sum_k \ln(P(z_k | \theta))\end{aligned}$$

最小二乘估计 (LSE) :

$$\begin{aligned}\hat{\theta}_{\text{LSE}} &= \arg \min_{\theta} J(\theta) \\ &= \arg \min_{\theta} \sum_k (z_k - h(x_k, \theta))^2\end{aligned}$$

贝叶斯估计 (Bayes) :

$$\begin{aligned}\hat{\theta}_{\text{Bayes}} &= \arg \min_{\theta} \gamma \\ &= \arg \min_{\theta} E_{P(Z, \theta)}[C(e)] \\ &= \arg \min_{\theta} \iint C(\theta - \hat{\theta}) P(Z, \theta) dZ d\theta \\ &= \arg \min_{\theta} \int \left[\int C(\theta - \hat{\theta}) P(\theta | Z) d\theta \right] P(Z) dZ\end{aligned}$$

最小均方误差估计 (MMSE) :

$$\begin{aligned}\hat{\theta}_{\text{MMSE}} &= \arg \min_{\theta} E_{P(Z, \theta)}[C(e)] \\ &= \arg \min_{\theta} \int \left[\int (\theta - \hat{\theta})^2 P(\theta | Z) d\theta \right] P(Z) dZ\end{aligned}$$

$$\hat{\theta}_{\text{MMSE}} = \int \theta P(\theta | Z) d\theta = E_{P(\theta | Z)}[\theta]$$

化简

最大后验估计 (MAP) :

$$\begin{aligned}\hat{\theta}_{\text{MAP}} &= \arg \min_{\theta} E_{P(Z, \theta)}[C(e)] \\ &= \arg \min_{\theta} \int \left[\int_{-\infty}^{\hat{\theta}-\delta} 1 \cdot P(\theta | Z) d\theta + \int_{\hat{\theta}+\delta}^{\infty} 1 \cdot P(\theta | Z) d\theta \right] P(Z) dZ \\ &= \arg \min_{\theta} \int \left[1 - \int_{\hat{\theta}-\delta}^{\hat{\theta}+\delta} P(\theta | Z) d\theta \right] P(Z) dZ \\ &\propto \arg \max_{\theta} \int_{\hat{\theta}-\delta}^{\hat{\theta}+\delta} P(\theta | Z) d\theta\end{aligned}$$

$$\begin{aligned}\hat{\theta}_{\text{MAP}} &= \arg \max_{\theta} P(\theta | Z) \\ &= \arg \max_{\theta} \frac{P(Z | \theta) P(\theta)}{P(Z)} \\ &\propto \arg \max_{\theta} \underbrace{P(Z | \theta)}_{\text{似然}} \cdot \underbrace{P(\theta)}_{\text{先验}}\end{aligned}$$

化简

均方误差与方差的区别和联系?

(在无偏估计下, 均方误差等于方差)

$$\begin{aligned}MSE[\hat{\theta}] &= E[(\hat{\theta} - \theta)^2] \\ &= E[\hat{\theta}^2 - 2\hat{\theta}\theta + \theta^2] \\ &= E[\hat{\theta}^2] - 2\theta E[\hat{\theta}] + \theta^2 \\ &= (E[\hat{\theta}^2] - (E[\hat{\theta}])^2) + (E[\hat{\theta}] - \theta)^2 \\ &= D[\hat{\theta}] + (E[\hat{\theta}] - \theta)^2\end{aligned}$$

7.3 估计理论

实际情况，观测到的样本数量不可能无穷多，因此一致性只是理论上需要满足的条件。在样本数量有限时，讨论估计值与实际值之间的偏差将更有意义，也就是偏差性。所以估计策略的目标是让估计量的偏差量尽量小，最小方差无偏估计无疑是一个理想的估计量。然而，最小方差无偏估计实现非常困难，因此需要寻找近似策略

- 估计量的性质
- 估计量的构建
- 各估计量比较
 - 从策略角度对比
 - 从模型角度对比
 - 等价转换关系

最大似然估计假设概率分布模型是知道的，只是概率分布模型中的参数 θ 未知。尝试 θ 的不同取值，看模型是否能输出与观测结果一样的数据。当某个 θ 取值能让模型输出与观测结果一样的数据的概率最大，那么这个 θ 取值就是最适合的估计参数。

$$\begin{aligned}\hat{\theta}_{MLE} &= \arg \max_{\theta} l(\theta | Z) \\ &= \arg \max_{\theta} \sum_k \ln(P(z_k | \theta))\end{aligned}$$

最小二乘估计是计算观测样本点与模型实际点之间的平方误差，求使得该平方误差最小的参数值 θ ，求出来的这个参数值 θ 就是估计参数

$$\begin{aligned}\hat{\theta}_{LSE} &= \arg \min_{\theta} J(\theta) \\ &= \arg \min_{\theta} \sum_k (z_k - h(x_k, \theta))^2\end{aligned}$$

贝叶斯估计首先也是构建关于误差的代价函数，然后最小化代价函数在概率密度下的期望

$$\begin{aligned}\hat{\theta}_{Bayes} &= \arg \min_{\theta} \gamma \\ &= \arg \min_{\theta} E_{P(Z, \theta)}[C(e)] \\ &= \arg \min_{\theta} \iint C(\theta - \hat{\theta}) P(Z, \theta) dZ d\theta \\ &= \arg \min_{\theta} \int \left[\int C(\theta - \hat{\theta}) P(\theta | Z) d\theta \right] P(Z) dZ\end{aligned}$$

- 可以看出，贝叶斯估计中的 $C(e)$ 是关于误差 e 的抽象函数，也就是说贝叶斯估计是一个通用的形式，具体的估计量形式与函数 $C(e)$ 的具体实现有关。最小均方误差估计和最大后验估计，就是贝叶斯估计的常见形式。
- **最小均方误差估计**就是用 θ 在后验分布上的概率均值作为估计取值。
- **最大后验估计**就是用 θ 在后验分布上最大值（也叫的众数）处对应的值作为估计取值。

7.3 估计理论

- 估计量的性质

■ 估计量的构建

■ 各估计量比较
- 从策略角度对比

从模型角度对比

等价转换关系

	(概率) 模型	模型待估参数
最大似然估计	已知	未知确定量
最小二乘估计	未知	未知确定量
(贝叶斯估计)	已知	未知随机量
最小均方误差估计	已知	未知随机量
最大后验估计	已知	未知随机量

来源于：书《机器人SLAM导航：核心技术与实战》表7-1

(概率) 模型已知 VS (概率) 模型未知

模型待估参数确定 VS 模型待估参数随机

经典估计 VS 贝叶斯估计

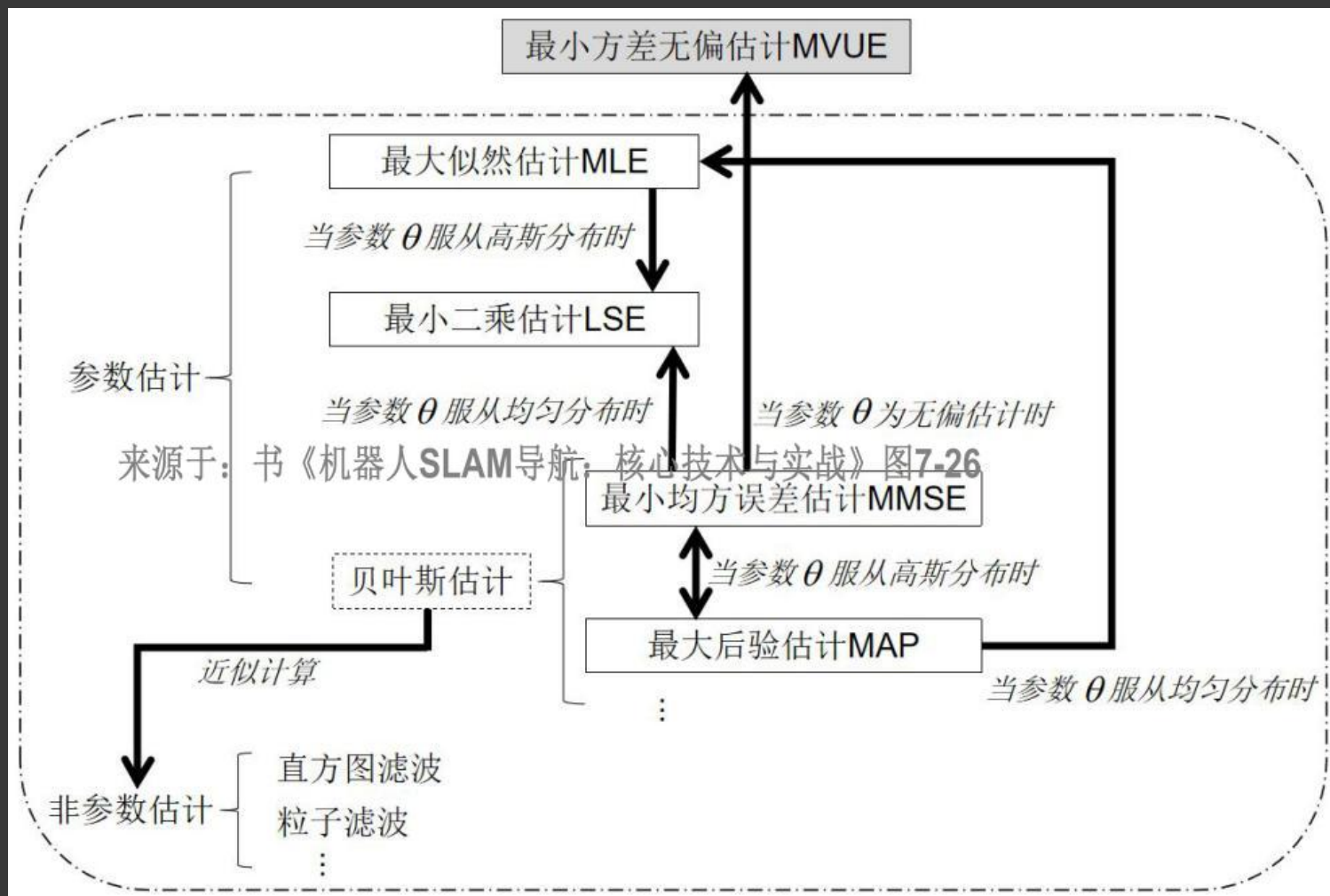
7.3 估计理论

■ 估计量的性质

■ 估计量的构建

■ 各估计量比较

从策略角度对比
从模型角度对比
等价转换关系



内容概要

7.1 SLAM发展简史

7.2 SLAM中的概率理论

7.3 估计理论

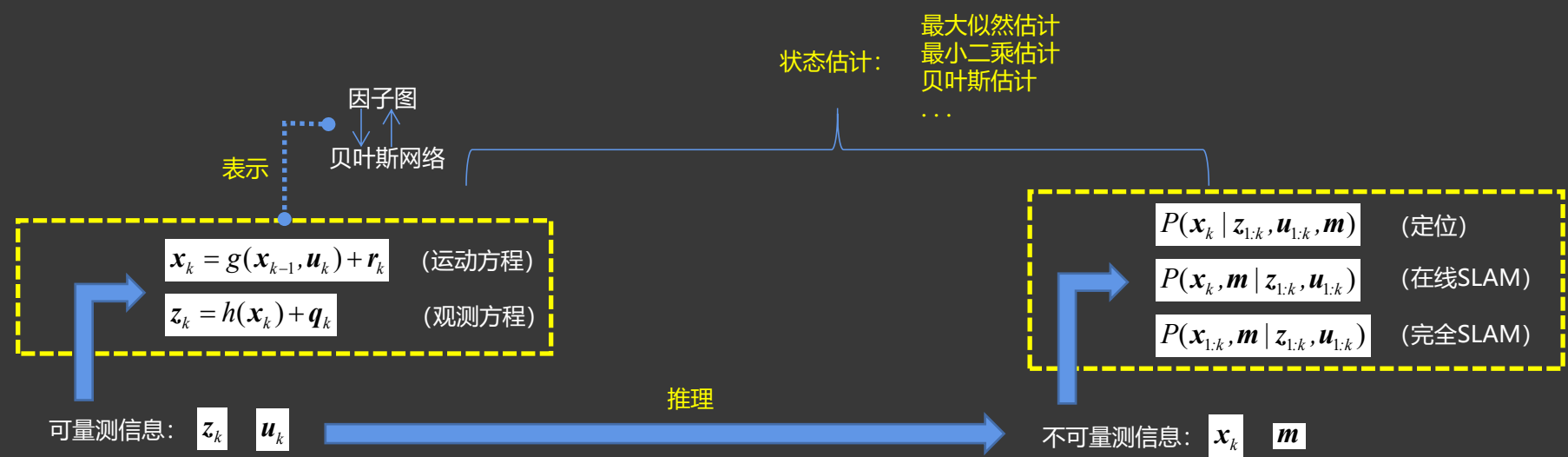
7.4 基于贝叶斯网络的状态估计

7.5 基于因子图的状态估计

7.6 典型SLAM算法

7.7 SFM、BA和SLAM比较

7.4 基于贝叶斯网络的状态估计



7.4 基于贝叶斯网络的状态估计

■ 贝叶斯估计  以贝叶斯网络（表示）、定位(待估问题)、贝叶斯估计（求解方法）这种情况为例来讨论：

■ 参数化实现

■ 非参数化实现

置信度，其实是对**后验概率分布**的符号化简写：

$$\begin{aligned} bel(\mathbf{x}_k) &= P(\mathbf{x}_k | \mathbf{z}_{1:k}, \mathbf{u}_{1:k}) \\ &= \frac{P(\mathbf{z}_k | \mathbf{x}_k, \mathbf{z}_{1:k-1}, \mathbf{u}_{1:k}) P(\mathbf{x}_k | \mathbf{z}_{1:k-1}, \mathbf{u}_{1:k})}{P(\mathbf{z}_k | \mathbf{z}_{1:k-1}, \mathbf{u}_{1:k})} \\ &= \eta \cdot P(\mathbf{z}_k | \mathbf{x}_k, \mathbf{z}_{1:k-1}, \mathbf{u}_{1:k}) P(\mathbf{x}_k | \mathbf{z}_{1:k-1}, \mathbf{u}_{1:k}) \end{aligned}$$

贝叶斯准则

贝叶斯网络中的条件独立性：

$$P(\mathbf{z}_k | \mathbf{x}_k, \mathbf{z}_{1:k-1}, \mathbf{u}_{1:k}) = P(\mathbf{z}_k | \mathbf{x}_k)$$

全概率展开：

$$\begin{aligned} \overline{bel(\mathbf{x}_k)} &= P(\mathbf{x}_k | \mathbf{z}_{1:k-1}, \mathbf{u}_{1:k}) \\ &= \int P(\mathbf{x}_k | \mathbf{x}_{k-1}, \mathbf{z}_{1:k-1}, \mathbf{u}_{1:k}) P(\mathbf{x}_{k-1} | \mathbf{z}_{1:k-1}, \mathbf{u}_{1:k}) d\mathbf{x}_{k-1} \end{aligned}$$

贝叶斯网络中的条件独立性：

$$P(\mathbf{x}_k | \mathbf{x}_{k-1}, \mathbf{z}_{1:k-1}, \mathbf{u}_{1:k}) = P(\mathbf{x}_k | \mathbf{x}_{k-1}, \mathbf{u}_k)$$

常识：

$$P(\mathbf{x}_{k-1} | \mathbf{z}_{1:k-1}, \mathbf{u}_{1:k}) = P(\mathbf{x}_{k-1} | \mathbf{z}_{1:k-1}, \mathbf{u}_{1:k-1})$$

$$bel(\mathbf{x}_k) = \eta \cdot P(\mathbf{z}_k | \mathbf{x}_k) \cdot \overline{bel(\mathbf{x}_k)}$$

$$\begin{aligned} \overline{bel(\mathbf{x}_k)} &= P(\mathbf{x}_k | \mathbf{z}_{1:k-1}, \mathbf{u}_{1:k}) \\ &= \int P(\mathbf{x}_k | \mathbf{x}_{k-1}, \mathbf{z}_{1:k-1}, \mathbf{u}_{1:k}) P(\mathbf{x}_{k-1} | \mathbf{z}_{1:k-1}, \mathbf{u}_{1:k}) d\mathbf{x}_{k-1} \\ &= \int P(\mathbf{x}_k | \mathbf{x}_{k-1}, \mathbf{u}_k) P(\mathbf{x}_{k-1} | \mathbf{z}_{1:k-1}, \mathbf{u}_{1:k-1}) d\mathbf{x}_{k-1} \\ &= \int P(\mathbf{x}_k | \mathbf{x}_{k-1}, \mathbf{u}_k) \cdot \overline{bel(\mathbf{x}_{k-1})} d\mathbf{x}_{k-1} \end{aligned}$$

递归贝叶斯滤波：

$$\begin{cases} \overline{bel(\mathbf{x}_k)} = \int P(\mathbf{x}_k | \mathbf{x}_{k-1}, \mathbf{u}_k) \cdot \overline{bel(\mathbf{x}_{k-1})} d\mathbf{x}_{k-1} \} \text{运动预测} \\ bel(\mathbf{x}_k) = \eta \cdot P(\mathbf{z}_k | \mathbf{x}_k) \overline{bel(\mathbf{x}_k)} \} \text{观测更新} \end{cases}$$

思考：

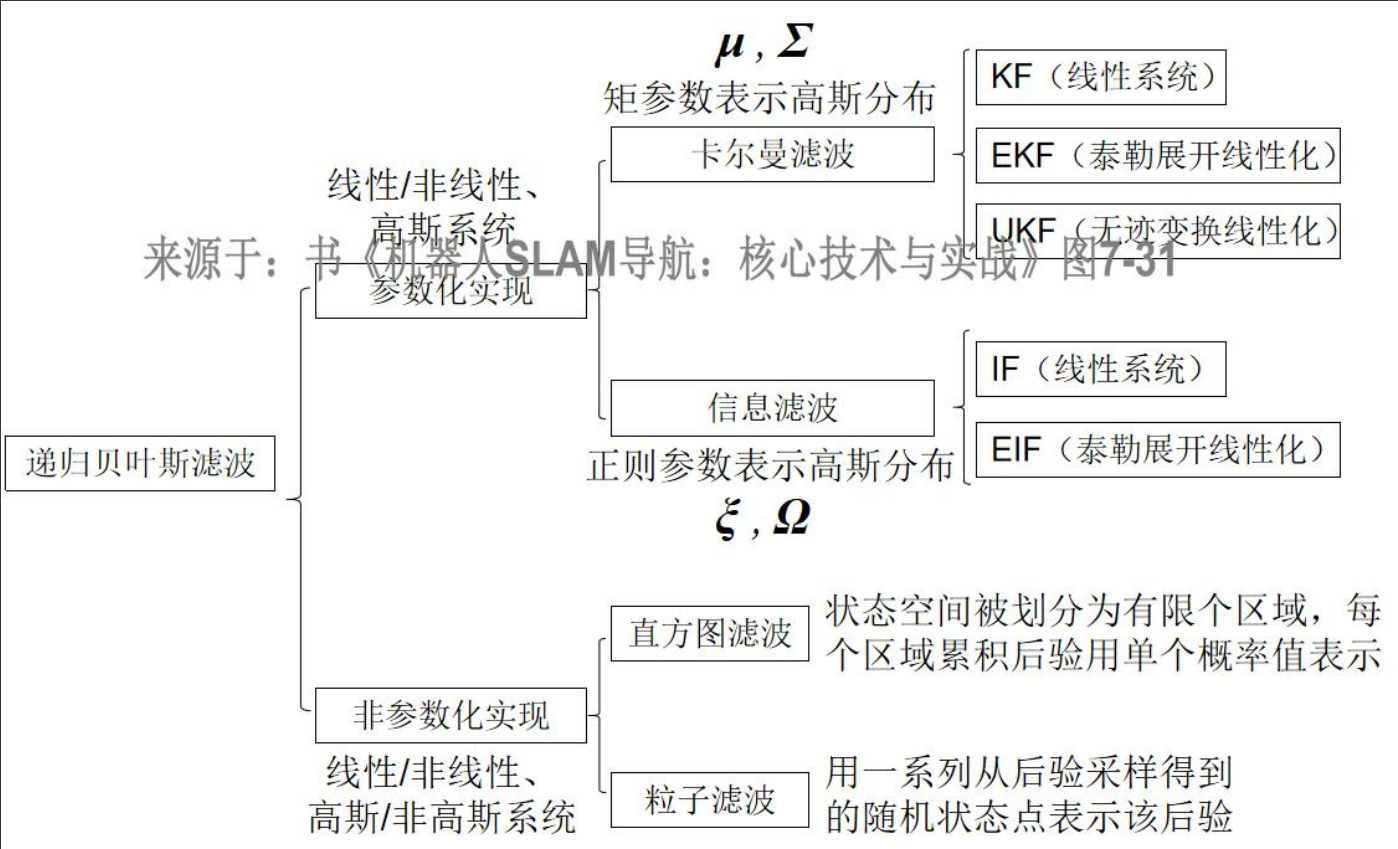
- 递归过程为何是求待估量的概率分布（即置信度）？
- 为何不直接求出待估量的具体取值呢？

7.4 基于贝叶斯网络的状态估计

- 贝叶斯估计 ➡ 以贝叶斯网络（表示）、定位(待估问题)、贝叶斯估计（求解方法）这种情况为例来讨论：
- 参数化实现
- 非参数化实现

由于后验概率分布 $P(\mathbf{x}_k | \mathbf{z}_{1:k}, \mathbf{u}_{1:k})$ 没有给定确切的形式，也就是说递归贝叶斯滤波是一种通用框架。在给定不同形式的 $P(\mathbf{x}_k | \mathbf{z}_{1:k}, \mathbf{u}_{1:k})$ 分布后，递归贝叶斯滤波也就对应不同形式的算法实现。

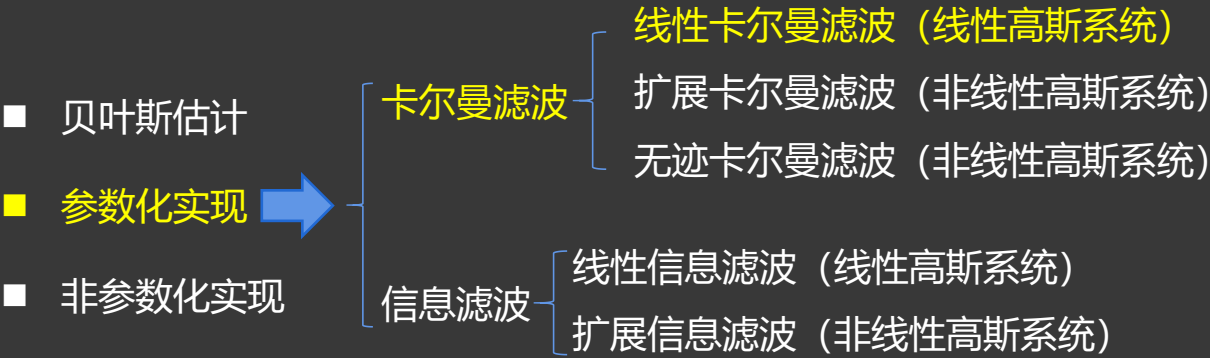
	高斯分布	非高斯分布
线性系统	线性高斯系统：KF、IF	线性非高斯系统
非线性系统	非线性高斯系统：EKF、UKF、EIF	非线性非高斯系统：HF、PF



7.4 基于贝叶斯网络的状态估计

高斯分布可以用矩参数(均值和方差)进行表示,那么多维高斯分布为:

$$P(\mathbf{x}) = \frac{1}{\sqrt{(2\pi)^n \det(\Sigma)}} \exp\left(-\frac{1}{2}(\mathbf{x} - \boldsymbol{\mu})^T \Sigma^{-1}(\mathbf{x} - \boldsymbol{\mu})\right)$$



运动方程: $\mathbf{x}_k = \mathbf{A}_k \mathbf{x}_{k-1} + \mathbf{B}_k \mathbf{u}_k + \mathbf{r}_k$ → 概率运动模型:

$$P(\mathbf{x}_k | \mathbf{x}_{k-1}, \mathbf{u}_k) = \frac{1}{\sqrt{(2\pi)^n \det(\mathbf{R}_k)}} \exp\left(-\frac{1}{2}(\mathbf{x}_k - (\mathbf{A}_k \mathbf{x}_{k-1} + \mathbf{B}_k \mathbf{u}_k))^T \mathbf{R}_k^{-1}(\mathbf{x}_k - (\mathbf{A}_k \mathbf{x}_{k-1} + \mathbf{B}_k \mathbf{u}_k))\right)$$

观测方程: $\mathbf{z}_k = \mathbf{C}_k \mathbf{x}_k + \mathbf{q}_k$ → 概率观测模型:

$$P(\mathbf{z}_k | \mathbf{x}_k) = \frac{1}{\sqrt{(2\pi)^n \det(\mathbf{Q}_k)}} \exp\left(-\frac{1}{2}(\mathbf{z}_k - \mathbf{C}_k \mathbf{x}_k)^T \mathbf{Q}_k^{-1}(\mathbf{z}_k - \mathbf{C}_k \mathbf{x}_k)\right)$$

待估状态量置信度的高斯分布形式:

$$bel(\mathbf{x}_k) = \frac{1}{\sqrt{(2\pi)^n \det(\Sigma_k)}} \exp\left(-\frac{1}{2}(\mathbf{x}_k - \boldsymbol{\mu}_k)^T \Sigma_k^{-1}(\mathbf{x}_k - \boldsymbol{\mu}_k)\right)$$

* 卡尔曼滤波其实就是利用下面的5个核心公式,递归计算状态 $\mathbf{x}_k$ 置信度的参数 $\boldsymbol{\mu}_k$ 和 Σ_k

① $\overline{\boldsymbol{\mu}}_k = \mathbf{A}_k \boldsymbol{\mu}_{k-1} + \mathbf{B}_k \mathbf{u}_k$

② $\overline{\Sigma}_k = \mathbf{A}_k \Sigma_{k-1} \mathbf{A}_k^T + \mathbf{R}_k$

③ $\mathbf{K}_k = \overline{\Sigma}_k \mathbf{C}_k^T (\mathbf{C}_k \overline{\Sigma}_k \mathbf{C}_k^T + \mathbf{Q}_k)^{-1}$ 卡尔曼增益

④ $\boldsymbol{\mu}_k = \overline{\boldsymbol{\mu}}_k + \mathbf{K}_k (\mathbf{z}_k - \mathbf{C}_k \overline{\boldsymbol{\mu}}_k)$

⑤ $\Sigma_k = (\mathbf{I} - \mathbf{K}_k \mathbf{C}_k) \overline{\Sigma}_k$

运动预测

观测更新

高斯分布矩参数化

高斯分布矩参数化

递归贝叶斯滤波

$\overline{bel}(\mathbf{x}_k) = \int P(\mathbf{x}_k | \mathbf{x}_{k-1}, \mathbf{u}_k) \cdot bel(\mathbf{x}_{k-1}) d\mathbf{x}_{k-1}$

$bel(\mathbf{x}_k) = \eta \cdot P(\mathbf{z}_k | \mathbf{x}_k) \overline{bel}(\mathbf{x}_k)$

运动预测

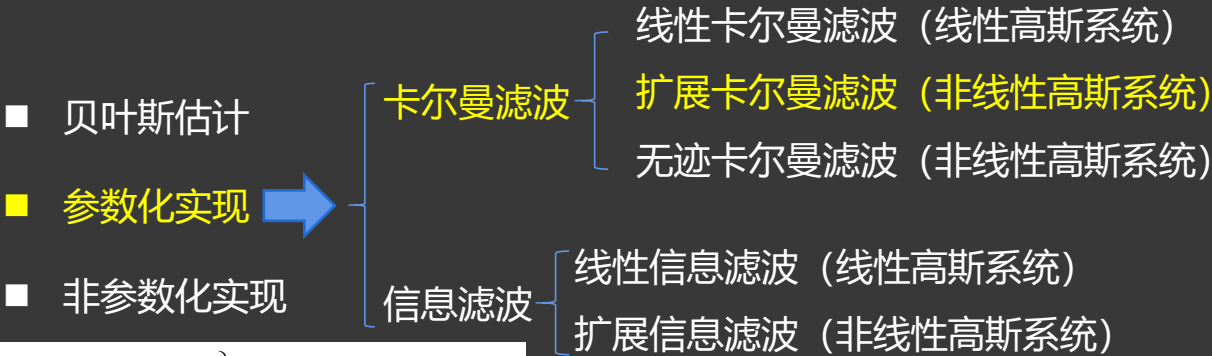
观测更新

7.4 基于贝叶斯网络的状态估计

高斯分布可以用矩参数（均值和方差）进行表示，那么多维高斯分布为：

$$P(\mathbf{x}) = \frac{1}{\sqrt{(2\pi)^n \det(\Sigma)}} \exp\left(-\frac{1}{2}(\mathbf{x} - \boldsymbol{\mu})^T \Sigma^{-1}(\mathbf{x} - \boldsymbol{\mu})\right)$$

实际机器人运动方程和观测方程往往是非线性的，也就是说高斯随机变量经过非线性函数g和h变换后都将变成非高斯随机变量，因此计算过程需要先线性化近似处理



① $\overline{\boldsymbol{\mu}}_k = g(\boldsymbol{\mu}_{k-1}, \mathbf{u}_k)$

② $\overline{\boldsymbol{\Sigma}}_k = \mathbf{G}_k \boldsymbol{\Sigma}_{k-1} \mathbf{G}_k^T + \mathbf{R}_k$

运动预测

③ $\mathbf{K}_k = \overline{\boldsymbol{\Sigma}}_k \mathbf{H}_k^T (\mathbf{H}_k \overline{\boldsymbol{\Sigma}}_k \mathbf{H}_k^T + \mathbf{Q}_k)^{-1}$

卡尔曼增益

④ $\boldsymbol{\mu}_k = \overline{\boldsymbol{\mu}}_k + \mathbf{K}_k (z_k - h(\overline{\boldsymbol{\mu}}_k))$

⑤ $\boldsymbol{\Sigma}_k = (\mathbf{I} - \mathbf{K}_k \mathbf{H}_k) \overline{\boldsymbol{\Sigma}}_k$

观测更新

扩展卡尔曼滤波 (EKF)

$$\begin{aligned} g(\mathbf{x}_{k-1}, \mathbf{u}_k) &\approx g(\mathbf{x}_{k-1}, \mathbf{u}_k)|_{\mathbf{x}_{k-1}=\overline{\boldsymbol{\mu}}_{k-1}} + g'(\mathbf{x}_{k-1}, \mathbf{u}_k)|_{\mathbf{x}_{k-1}=\overline{\boldsymbol{\mu}}_{k-1}} \bullet (\mathbf{x}_{k-1} - \overline{\boldsymbol{\mu}}_{k-1}) \\ &= g(\overline{\boldsymbol{\mu}}_{k-1}, \mathbf{u}_k) + g'(\overline{\boldsymbol{\mu}}_{k-1}, \mathbf{u}_k) \bullet (\mathbf{x}_{k-1} - \overline{\boldsymbol{\mu}}_{k-1}) \\ &= g(\overline{\boldsymbol{\mu}}_{k-1}, \mathbf{u}_k) + \mathbf{G}_k \bullet (\mathbf{x}_{k-1} - \overline{\boldsymbol{\mu}}_{k-1}) \end{aligned}$$

运动方程： $\mathbf{x}_k = g(\mathbf{x}_{k-1}, \mathbf{u}_k) + \mathbf{r}_k$

一阶泰勒展开进行线性化

→ 概率运动模型：

$$P(\mathbf{x}_k | \mathbf{x}_{k-1}, \mathbf{u}_k) = \frac{1}{\sqrt{(2\pi)^n \det(\mathbf{R}_k)}} \exp\left(-\frac{1}{2}(\mathbf{x}_k - (g(\overline{\boldsymbol{\mu}}_{k-1}, \mathbf{u}_k) + \mathbf{G}_k \bullet (\mathbf{x}_{k-1} - \overline{\boldsymbol{\mu}}_{k-1})))^T \bullet \mathbf{R}_k^{-1}(\mathbf{x}_k - (g(\overline{\boldsymbol{\mu}}_{k-1}, \mathbf{u}_k) + \mathbf{G}_k \bullet (\mathbf{x}_{k-1} - \overline{\boldsymbol{\mu}}_{k-1})))\right)$$

观测方程： $z_k = h(\mathbf{x}_k) + q_k$

一阶泰勒展开进行线性化

→ 概率观测模型：

$$P(z_k | \mathbf{x}_k) = \frac{1}{\sqrt{(2\pi)^n \det(\mathbf{Q}_k)}} \exp\left(-\frac{1}{2}(z_k - (h(\overline{\boldsymbol{\mu}}_k) + \mathbf{H}_k \bullet (\mathbf{x}_k - \overline{\boldsymbol{\mu}}_k)))^T \bullet \mathbf{Q}_k^{-1}(z_k - (h(\overline{\boldsymbol{\mu}}_k) + \mathbf{H}_k \bullet (\mathbf{x}_k - \overline{\boldsymbol{\mu}}_k)))\right)$$

$$\begin{aligned} h(\mathbf{x}_k) &\approx h(\mathbf{x}_k)|_{\mathbf{x}_k=\overline{\boldsymbol{\mu}}_k} + h'(\mathbf{x}_k)|_{\mathbf{x}_k=\overline{\boldsymbol{\mu}}_k} \bullet (\mathbf{x}_k - \overline{\boldsymbol{\mu}}_k) \\ &= h(\overline{\boldsymbol{\mu}}_k) + h'(\overline{\boldsymbol{\mu}}_k) \bullet (\mathbf{x}_k - \overline{\boldsymbol{\mu}}_k) \\ &= h(\overline{\boldsymbol{\mu}}_k) + \mathbf{H}_k \bullet (\mathbf{x}_k - \overline{\boldsymbol{\mu}}_k) \end{aligned}$$

来源于：书《机器人SLAM导航：核心技术与实战》图7-27

一阶泰勒展开进行线性化

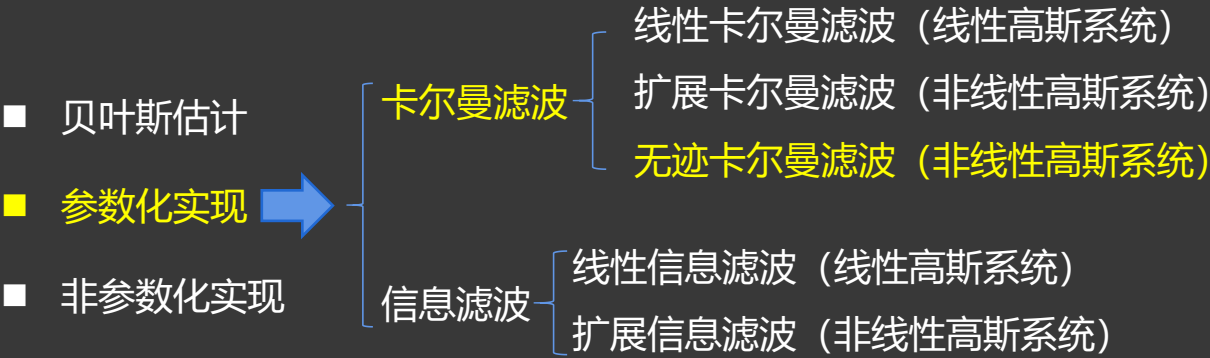
70/92

课件下载：www.xiihoo.com

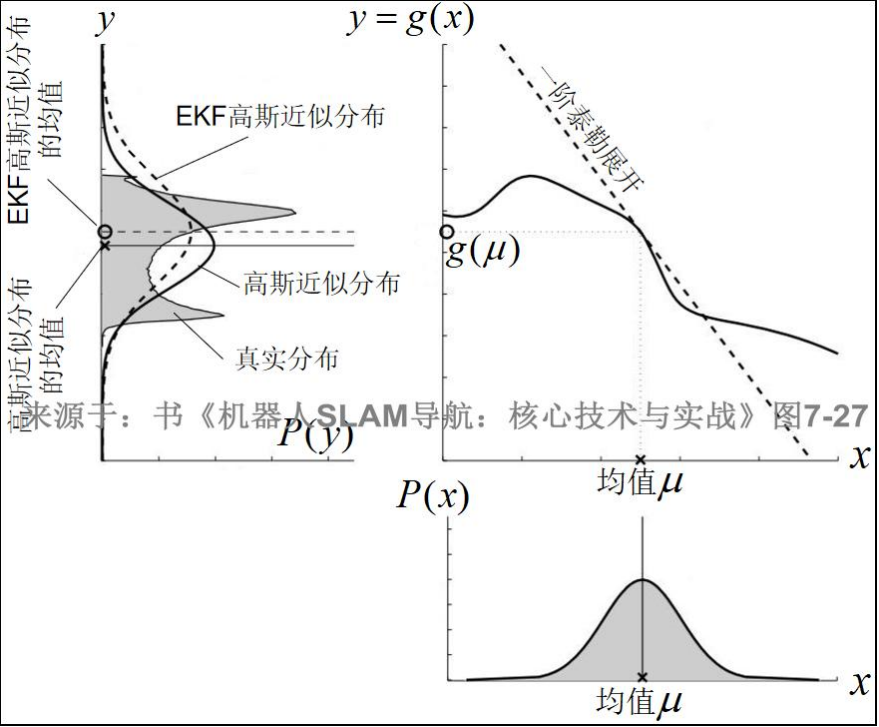
7.4 基于贝叶斯网络的状态估计

高斯分布可以用矩参数（均值和方差）进行表示，那么多维高斯分布为：

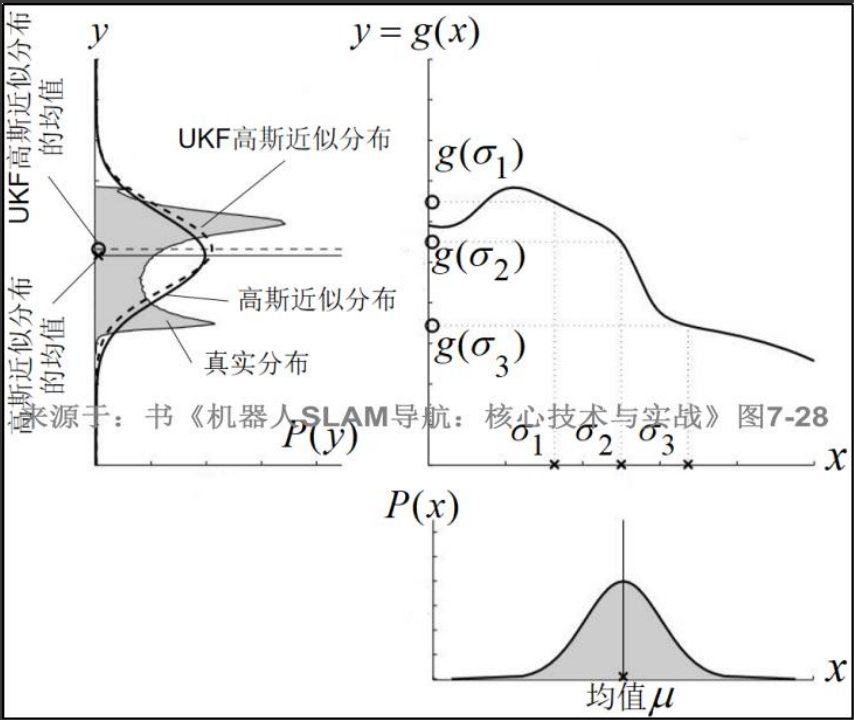
$$P(x) = \frac{1}{\sqrt{(2\pi)^n \det(\Sigma)}} \exp\left(-\frac{1}{2}(x-\mu)^T \Sigma^{-1}(x-\mu)\right)$$



- 如果x的不确定性和函数g的非线性度都小，那么扩展卡尔曼滤波（EKF）中线性化近似效果是好的。
- 但是，当x的不确定性和函数g的非线性度较大时，扩展卡尔曼滤波（EKF）的近似效果就不好了，此时可以用无迹卡尔曼滤波（UKF）来处理。
- 除了EKF和UKF方法，还有很多其他处理非线性系统线性近似的方法，这里就不再过大幅展开了



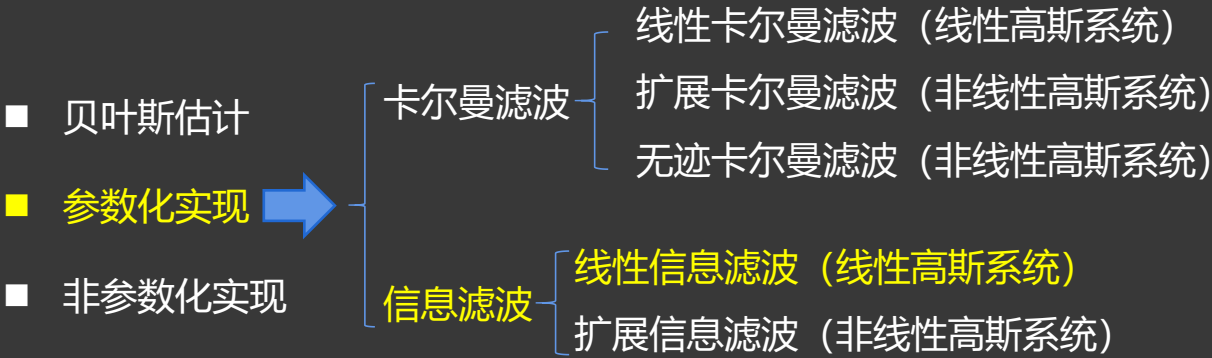
VS



7.4 基于贝叶斯网络的状态估计

高斯分布可以用矩参数（均值和方差）进行表示，那么多维高斯分布为：

$$P(\mathbf{x}) = \frac{1}{\sqrt{(2\pi)^n \det(\Sigma)}} \exp\left(-\frac{1}{2}(\mathbf{x} - \boldsymbol{\mu})^T \Sigma^{-1}(\mathbf{x} - \boldsymbol{\mu})\right)$$



矩参数（均值和方差）可以转换成正则参数（信息矩阵和信息向量）：

$$\begin{cases} \boldsymbol{\Omega} = \Sigma^{-1} \\ \boldsymbol{\xi} = \Sigma^{-1} \boldsymbol{\mu} \end{cases}$$

那么正则参数形式的多维高斯分布为：

$$\begin{aligned} P(\mathbf{x}) &= \frac{1}{\sqrt{(2\pi)^n \det(\Sigma)}} \exp\left(-\frac{1}{2}(\mathbf{x} - \boldsymbol{\mu})^T \Sigma^{-1}(\mathbf{x} - \boldsymbol{\mu})\right) \\ &= \frac{1}{\sqrt{(2\pi)^n \det(\Sigma)}} \exp\left(-\frac{1}{2} \mathbf{x}^T \Sigma^{-1} \mathbf{x} + \mathbf{x}^T \Sigma^{-1} \boldsymbol{\mu} - \frac{1}{2} \boldsymbol{\mu}^T \Sigma^{-1} \boldsymbol{\mu}\right) \\ &= \underbrace{\frac{1}{\sqrt{(2\pi)^n \det(\Sigma)}} \exp\left(-\frac{1}{2} \boldsymbol{\mu}^T \Sigma^{-1} \boldsymbol{\mu}\right)}_{\text{常数项}} \cdot \exp\left(-\frac{1}{2} \mathbf{x}^T \underbrace{\Sigma^{-1}}_{\boldsymbol{\Omega}} \mathbf{x} + \mathbf{x}^T \underbrace{\Sigma^{-1} \boldsymbol{\mu}}_{\boldsymbol{\xi}}\right) \\ &= \eta \cdot \exp\left(-\frac{1}{2} \mathbf{x}^T \boldsymbol{\Omega} \mathbf{x} + \mathbf{x}^T \boldsymbol{\xi}\right) \end{aligned}$$

矩参数和正则参数只是表示高斯分布的2种不同方式，高斯分布的本质是一样的。
因此，线性信息滤波的推导与线性卡尔曼滤波是类似的

$$\left\{ \begin{array}{l} \textcircled{1} \overline{\boldsymbol{\Omega}}_k = (\mathbf{A}_k \boldsymbol{\Omega}_{k-1}^{-1} \mathbf{A}_k^T + \mathbf{R}_k)^{-1} \\ \textcircled{2} \overline{\boldsymbol{\xi}}_k = \overline{\boldsymbol{\Omega}}_k (\mathbf{A}_k \boldsymbol{\Omega}_{k-1}^{-1} \boldsymbol{\xi}_{k-1} + \mathbf{B}_k \mathbf{u}_k) \end{array} \right\} \text{运动预测}$$
$$\left\{ \begin{array}{l} \textcircled{3} \boldsymbol{\Omega}_k = \overline{\boldsymbol{\Omega}}_k + \mathbf{C}_k^T \boldsymbol{\Omega}_k^{-1} \mathbf{C}_k \\ \textcircled{4} \boldsymbol{\xi}_k = \overline{\boldsymbol{\xi}}_k + \mathbf{C}_k^T \boldsymbol{\Omega}_k^{-1} \mathbf{z}_k \end{array} \right\} \text{观测更新}$$

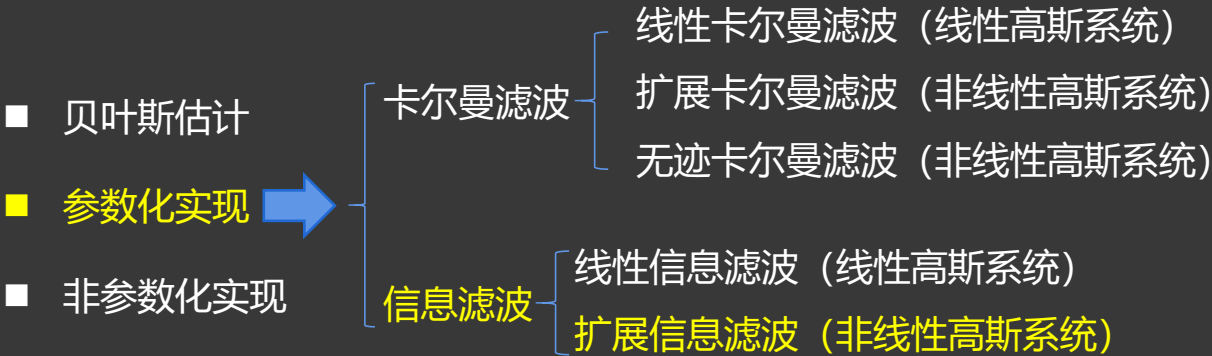
线性信息滤波
(IF)

- 可以发现线性卡尔曼滤波和线性信息滤波的性能是对偶的。
- 线性卡尔曼滤波中的运动预测是增量的，而线性信息滤波中的运动预测要求矩阵逆运算，所以线性卡尔曼滤波在运动预测上计算效率更高。
- 不过，线性卡尔曼滤波中的观测更新要求矩阵逆运算，而线性信息滤波中的观测更新是增量的，所以线性信息滤波在观测更新上计算效率更高。

7.4 基于贝叶斯网络的状态估计

高斯分布可以用矩参数（均值和方差）进行表示，那么多维高斯分布为：

$$P(\mathbf{x}) = \frac{1}{\sqrt{(2\pi)^n \det(\Sigma)}} \exp\left(-\frac{1}{2}(\mathbf{x} - \boldsymbol{\mu})^T \Sigma^{-1}(\mathbf{x} - \boldsymbol{\mu})\right)$$



矩参数（均值和方差）可以转换成正则参数（信息矩阵和信息向量）：

$$\begin{cases} \boldsymbol{\Omega} = \Sigma^{-1} \\ \boldsymbol{\xi} = \Sigma^{-1} \boldsymbol{\mu} \end{cases}$$

那么正则参数形式的多维高斯分布为：

$$\begin{aligned} P(\mathbf{x}) &= \frac{1}{\sqrt{(2\pi)^n \det(\Sigma)}} \exp\left(-\frac{1}{2}(\mathbf{x} - \boldsymbol{\mu})^T \Sigma^{-1}(\mathbf{x} - \boldsymbol{\mu})\right) \\ &= \frac{1}{\sqrt{(2\pi)^n \det(\Sigma)}} \exp\left(-\frac{1}{2} \mathbf{x}^T \Sigma^{-1} \mathbf{x} + \mathbf{x}^T \Sigma^{-1} \boldsymbol{\mu} - \frac{1}{2} \boldsymbol{\mu}^T \Sigma^{-1} \boldsymbol{\mu}\right) \\ &= \underbrace{\frac{1}{\sqrt{(2\pi)^n \det(\Sigma)}} \exp\left(-\frac{1}{2} \boldsymbol{\mu}^T \Sigma^{-1} \boldsymbol{\mu}\right)}_{\text{常数项}} \cdot \exp\left(-\frac{1}{2} \mathbf{x}^T \underbrace{\Sigma^{-1}}_{\boldsymbol{\Omega}} \mathbf{x} + \mathbf{x}^T \underbrace{\Sigma^{-1} \boldsymbol{\mu}}_{\boldsymbol{\xi}}\right) \\ &= \eta \cdot \exp\left(-\frac{1}{2} \mathbf{x}^T \boldsymbol{\Omega} \mathbf{x} + \mathbf{x}^T \boldsymbol{\xi}\right) \end{aligned}$$

矩参数和正则参数只是表示高斯分布的2种不同方式，高斯分布的本质是一样的。
对于非线性高斯系统，扩展信息滤波与扩展卡尔曼滤波是类似的

$$\left\{ \begin{aligned} \overline{\boldsymbol{\mu}}_{k-1} &= \boldsymbol{\Omega}_{k-1}^{-1} \boldsymbol{\xi}_{k-1} \\ \overline{\boldsymbol{\mu}}_k &= g(\overline{\boldsymbol{\mu}}_{k-1}, \mathbf{u}_k) \\ \text{① } \overline{\boldsymbol{\Omega}}_k &= (\mathbf{G}_k \boldsymbol{\Omega}_{k-1}^{-1} \mathbf{G}_k^T + \mathbf{R}_k)^{-1} \\ \text{② } \overline{\boldsymbol{\xi}}_k &= \overline{\boldsymbol{\Omega}}_k \bullet \overline{\boldsymbol{\mu}}_k \end{aligned} \right\}$$

运动预测

$$\left\{ \begin{aligned} \text{③ } \boldsymbol{\Omega}_k &= \overline{\boldsymbol{\Omega}}_k + \mathbf{H}_k^T \mathbf{Q}_k^{-1} \mathbf{H}_k \\ \text{④ } \boldsymbol{\xi}_k &= \overline{\boldsymbol{\xi}}_k + \mathbf{H}_k^T \mathbf{Q}_k^{-1} (\mathbf{z}_k - h(\overline{\boldsymbol{\mu}}_k) + \mathbf{H}_k \overline{\boldsymbol{\mu}}_k) \end{aligned} \right\}$$

观测更新

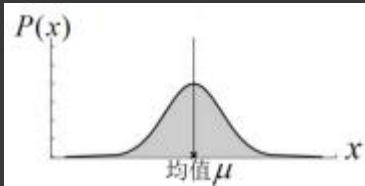
扩展信息滤波
(EIF)

对实际非线性系统进行
线性化近似处理

7.4 基于贝叶斯网络的状态估计

- 贝叶斯估计
- 参数化实现
- 非参数化实现

实际问题往往除了是**非线性系统**外，通常还是**非高斯系统**。



高斯分布

$$P(x) = \frac{1}{\sqrt{(2\pi)^n \det(\Sigma)}} \exp\left(-\frac{1}{2}(x - \mu)^T \Sigma^{-1}(x - \mu)\right)$$

(概率密度函数可解析表示)

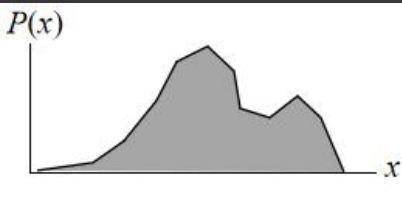
概率运动模型: $P(x_k | x_{k-1}, u_k)$ = 可解析

概率观测模型: $P(z_k | x_k)$ = 可解析

待估计量的置信度: $bel(x_k) = P(x_k | z_{1:k}, u_{1:k})$ = 可解析

参数化实现

(对概率分布进行参数化表示)



非高斯分布

$$P(x) = \alpha_0 + \alpha_1 x + \alpha_2 x^2 + \dots + \alpha_n x^n + \dots$$

(概率密度函数不可解析表示，用无限级数逼近)

$P(x_k | x_{k-1}, u_k)$ = 不可解析

$P(z_k | x_k)$ = 不可解析

$bel(x_k) = P(x_k | z_{1:k}, u_{1:k})$ = 不可解析

?

非参数化实现

(对概率分布进行非参数化样本模拟)

7.4 基于贝叶斯网络的状态估计

- 贝叶斯估计
 - 参数化实现
 - 非参数化实现
- 直方图滤波

粒子滤波

利用区域划分将待估状态量从连续空间离散化到离散状态空间

①离散化的概率运动模型

$$P(\mathbf{x}_k^i | \mathbf{x}_{k-1}^j, \mathbf{u}_k) \approx \eta_1 \cdot P(c(\mathbf{x}_k^i) | c(\mathbf{x}_{k-1}^j), \mathbf{u}_k)$$

②离散化的概率观测模型

$$P(z_k | \mathbf{x}_k^i) \approx \eta_2 \cdot P(z_k | c(\mathbf{x}_k^i))$$

递归贝叶斯滤波

$$\begin{cases} \overline{bel(\mathbf{x}_k)} = \int P(\mathbf{x}_k | \mathbf{x}_{k-1}, \mathbf{u}_k) \cdot \overline{bel(\mathbf{x}_{k-1})} d\mathbf{x}_{k-1} \end{cases} \text{运动预测}$$
$$\begin{cases} bel(\mathbf{x}_k) = \eta \cdot P(z_k | \mathbf{x}_k) \overline{bel(\mathbf{x}_k)} \end{cases} \text{观测更新}$$

离散化

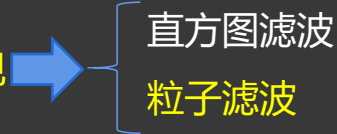
③离散递归贝叶斯滤波

(遍历待估状态的所有离散区域)

$$\begin{cases} \overline{bel(\mathbf{x}_k^i)} = \sum_j P(\mathbf{x}_k^i | \mathbf{x}_{k-1}^j, \mathbf{u}_k) \cdot \overline{bel(\mathbf{x}_{k-1}^j)} \end{cases} \text{运动预测}$$
$$\begin{cases} bel(\mathbf{x}_k^i) = \eta \cdot P(z_k | \mathbf{x}_k^i) \cdot \overline{bel(\mathbf{x}_k^i)} \end{cases} \text{观测更新}$$

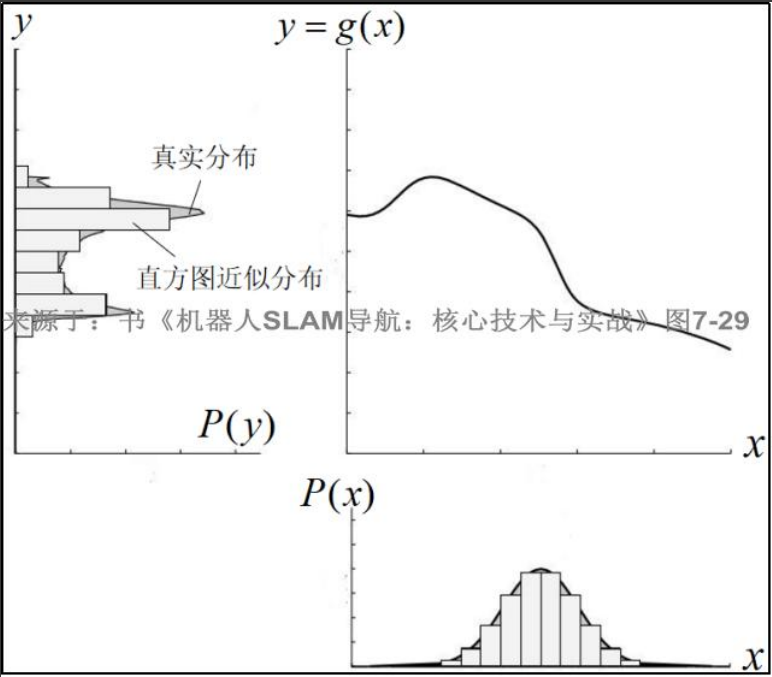
7.4 基于贝叶斯网络的状态估计

- 贝叶斯估计
- 参数化实现
- 非参数化实现

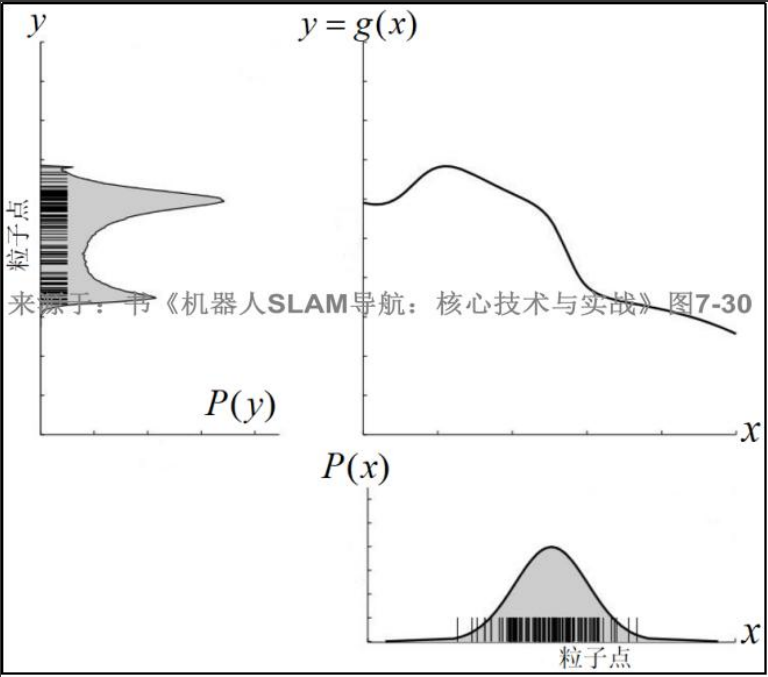


* 直方图滤波，可以认为是用区域划分方法对待估量的状态空间进行离散化，用每个划分区域上质心的概率代表整个区域的概率大小

* 粒子滤波，可以认为是用随机粒子点对待估量的状态空间进行离散化，用粒子点的稠密度代表该取值空间处的概率大小



直方图近似表示概率分布



粒子近似表示概率分布

内容概要

7.1 SLAM发展简史

7.2 SLAM中的概率理论

7.3 估计理论

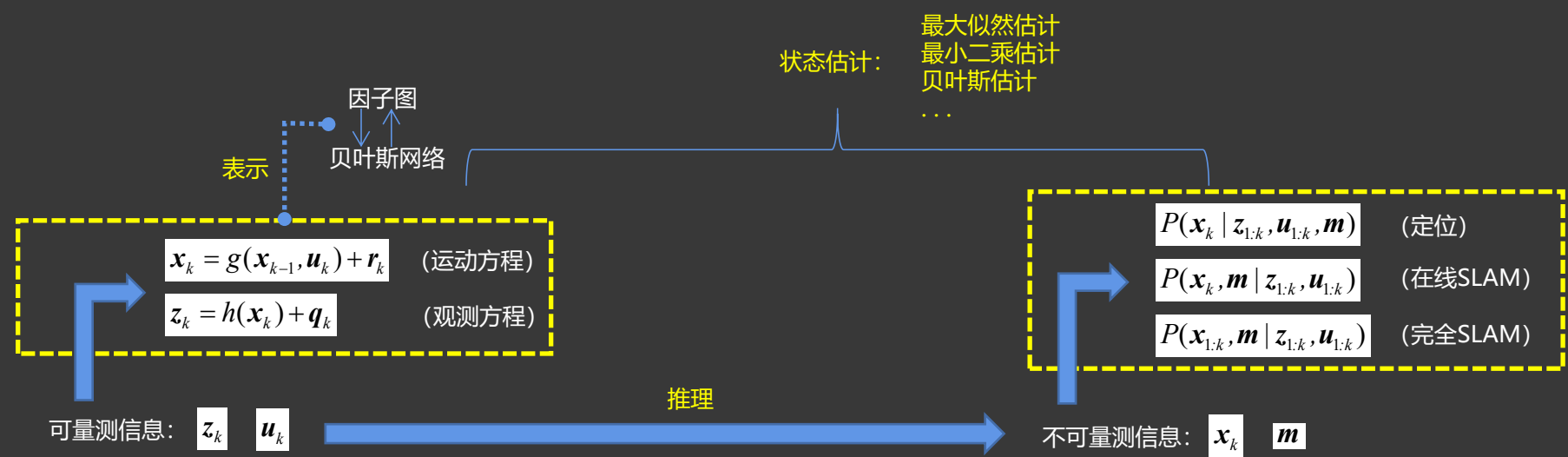
7.4 基于贝叶斯网络的状态估计

7.5 基于因子图的状态估计

7.6 典型SLAM算法

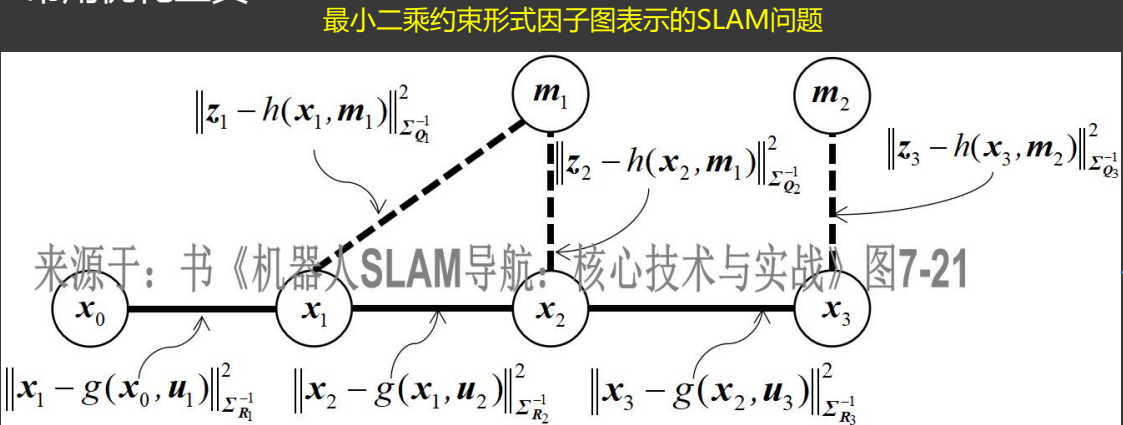
7.7 SFM、BA和SLAM比较

7.5 基于因子图的状态估计



7.5 基于因子图的状态估计

- 非线性最小二乘估计 ➡ 以因子图（表示）、完全SLAM(待估问题)、最小二乘估计（求解方法）这种情况为例来讨论：
- 直接求解方法
- 优化方法
- 各种优化方法对比
- 常用优化工具



最小二乘估计
(实际一般为非线性最小二乘估计)

$$\begin{aligned} \mathcal{S}_{\text{MAP}} &= \arg \max_{\mathcal{S}} \psi(\mathcal{S}) \\ &\propto \arg \max_{\mathcal{S}} \left(\sum_{i1} \ln \psi_{i1}(\mathbf{x}_k, \mathbf{x}_{k-1}) + \sum_{i2} \ln \psi_{i2}(\mathbf{x}_k, \mathbf{m}_j) \right) \\ &\propto \arg \max_{\mathcal{S}} \left(-\frac{1}{2} \sum_{i1} (\mathbf{x}_k - g(\mathbf{x}_{k-1}, \mathbf{u}_k))^T \Sigma_{R_k}^{-1} (\mathbf{x}_k - g(\mathbf{x}_{k-1}, \mathbf{u}_k)) \right. \\ &\quad \left. - \frac{1}{2} \sum_{i2} (z_k - h(\mathbf{x}_k, \mathbf{m}_j))^T \Sigma_{Q_k}^{-1} (z_k - h(\mathbf{x}_k, \mathbf{m}_j)) \right) \\ &\propto \arg \min_{\mathcal{S}} \left(\sum_{i1} \|\mathbf{x}_k - g(\mathbf{x}_{k-1}, \mathbf{u}_k)\|_{\Sigma_{R_k}^{-1}}^2 + \sum_{i2} \|z_k - h(\mathbf{x}_k, \mathbf{m}_j)\|_{\Sigma_{Q_k}^{-1}}^2 \right) \end{aligned}$$

思考：
这里的非线性是指什么？

非线性最小二乘估计的通用形式
(严格说是，非线性加权最小二乘)

$$\min_{\mathbf{x}} \sum_i \|\mathbf{y}_i - \mathbf{f}(\mathbf{x}_i)\|_{\Sigma_i^{-1}}^2$$

7.5 基于因子图的状态估计

- 非线性最小二乘估计
- 直接求解方法
- 优化方法
- 各种优化方法对比
- 常用优化工具

非线性最小二乘问题

$$\min_x \sum_i \|y_i - f(x_i)\|_{\Sigma_i^{-1}}^2$$

泰勒展开进行线性化近似

$$f(x) = \frac{f(x_0)}{0!} + \frac{f'(x_0)}{1!}(x - x_0) + \frac{f''(x_0)}{2!}(x - x_0)^2 + \dots + \frac{f^{(n)}(x_0)}{n!}(x - x_0)^n + R_n(x)$$

$$\begin{aligned} f(x_i) &\approx f(x_i)|_{x_i=x\theta_i} + f'(x_i)|_{x_i=x\theta_i} \bullet (x_i - x\theta_i) \\ &= f(x\theta_i) + f'(x\theta_i) \bullet (x_i - x\theta_i) \\ &= f(x\theta_i) + F_i \bullet (x_i - x\theta_i) \end{aligned}$$

(一般取一阶泰勒展开)

$$\begin{aligned} \min_x \sum_i \|y_i - f(x_i)\|_{\Sigma_i^{-1}}^2 &\propto \min_x \sum_i \|y_i - (f(x\theta_i) + F_i \bullet (x_i - x\theta_i))\|_{\Sigma_i^{-1}}^2 \\ &= \min_x \sum_i \|(y_i - f(x\theta_i) + F_i \bullet x\theta_i) - F_i \bullet x_i\|_{\Sigma_i^{-1}}^2 \\ &= \min_x \sum_i \left\| \underbrace{\Sigma_i^{-1/2}(y_i - f(x\theta_i) + F_i \bullet x\theta_i)}_{b_i} - \underbrace{\Sigma_i^{-1/2}F_i}_{A_i} \bullet x_i \right\|^2 \\ &= \min_x \sum_i \|b_i - A_i \bullet x_i\|^2 \end{aligned}$$

Cholesky分解

$$\begin{aligned} Ax = b &\Leftrightarrow (A^T A)x = A^T b \\ &\Leftrightarrow (R^T R)x = A^T b \\ &\Leftrightarrow R^T (Rx) = A^T b \\ &\Leftrightarrow R^T y = A^T b \end{aligned} \quad \left. \begin{array}{l} \\ \\ \\ \end{array} \right\} \text{其中: } A^T A = R^T R, Rx = y$$

QR分解

$$\begin{aligned} Ax = b &\Leftrightarrow (QR)x = b \\ &\Leftrightarrow Q^{-1}QRx = Q^{-1}b \\ &\Leftrightarrow Rx = Q^{-1}b \\ &\Leftrightarrow Rx = Q^T b \end{aligned} \quad \left. \begin{array}{l} \\ \\ \\ \end{array} \right\} \text{其中: } \begin{array}{l} A = QR \\ Q \text{ 是正交矩阵 } (Q^{-1} = Q^T) \\ R \text{ 是上三角矩阵} \end{array}$$

解线性方程

理想情况
最小值可取到0

$$\begin{cases} A_1 \bullet x_1 = b_1 \\ A_2 \bullet x_2 = b_2 \\ \dots \\ A_i \bullet x_i = b_i \\ \dots \end{cases}, \text{令 } A = \begin{bmatrix} A_1 & 0 & \dots & 0 & \dots \\ 0 & A_2 & \dots & 0 & \dots \\ \dots & \dots & \dots & \dots & \dots \\ 0 & 0 & \dots & A_i & \dots \\ \dots & \dots & \dots & \dots & \dots \end{bmatrix}, x = \begin{bmatrix} x_1 \\ x_2 \\ \dots \\ x_i \\ \dots \end{bmatrix}, b = \begin{bmatrix} b_1 \\ b_2 \\ \dots \\ b_i \\ \dots \end{bmatrix} \Rightarrow Ax = b$$

7.5 基于因子图的状态估计

由于实际SLAM问题的非线性最小二乘中，很难找到合适的线性化方法，初值也比较难确定，并且代价函数的误差往往不能最小化到0值，所以上面介绍的直接求解方法很难利用在实际问题中。实际工程中，通常采用优化方法进行迭代求解。

- 非线性最小二乘估计
- 直接求解方法
- 优化方法
- 各种优化方法对比
- 常用优化工具

非线性最小二乘问题

$$\psi(x) = \sum_i \|y_i - f(x_i)\|_{\Sigma_i^{-1}}^2$$

实际情况
最小值取不到0
初值难确定
线性化困难

迭代策略

梯度下降算法

$$\Delta x_{\text{grad}} = -\alpha \cdot \nabla \psi(x^{(k)})$$
$$x^{(k+1)} = x^{(k)} + \Delta x_{\text{grad}}$$

最速下降算法

$$\alpha_k = \arg \min_{\alpha \geq 0} \psi(x^{(k)} - \alpha \cdot \nabla \psi(x^{(k)}))$$
$$\Delta x_{\text{SD}} = -\alpha_k \cdot \nabla \psi(x^{(k)})$$
$$x^{(k+1)} = x^{(k)} + \Delta x_{\text{SD}}$$

高斯-牛顿算法

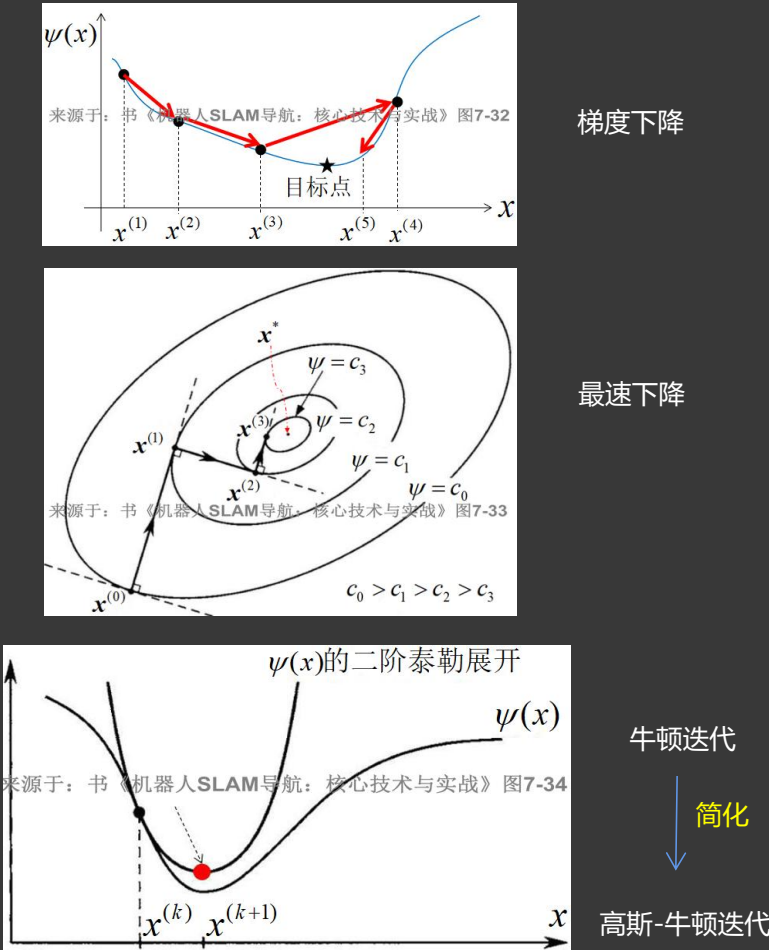
$$\Delta x_{\text{GN}} = -(J_e^T J_e)^{-1} \cdot J_e^T \cdot e$$
$$x^{(k+1)} = x^{(k)} + \Delta x_{\text{GN}}$$

列文伯格-马夸尔特算法

$$\Delta x_{\text{LM}}, \mu_{k+1} = LM(x^{(k)}, \mu_k)$$
$$x^{(k+1)} = x^{(k)} + \Delta x_{\text{LM}}$$

狗腿算法

$$\Delta x_{\text{DL}}, d_{k+1} = Dogleg(x^{(k)}, d_k)$$
$$x^{(k+1)} = x^{(k)} + \Delta x_{\text{DL}}$$



7.5 基于因子图的状态估计

- 非线性最小二乘估计
 - 直接求解方法
 - 优化方法
 - 各种优化方法对比
 - 常用优化工具
- 可以发现，每一种方法都是前一种的改进，下降速度越来越好，但是每一步迭代的计算代价越来越高。
 - 因此，在实际应用中要结合实际问题的合理选择。比如Dogleg算法虽然每步迭代下降效果是最好的，但是每步迭代付出的计算代价也最昂贵，站在整个迭代过程看，迭代效率并不一定比最普通的梯度下降法好。

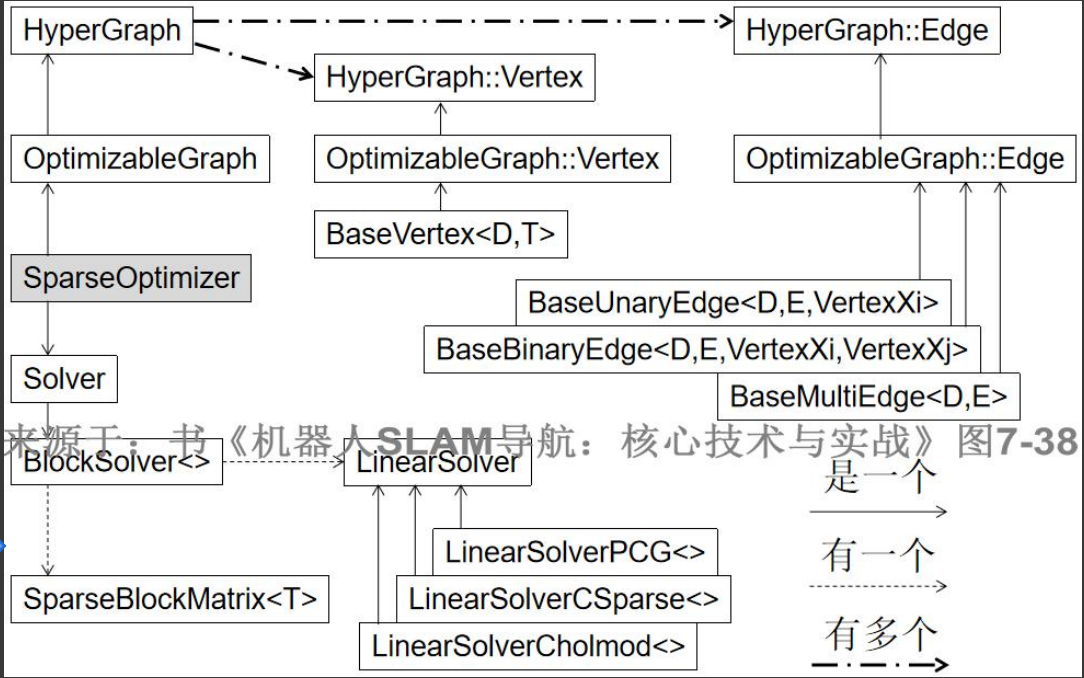
7.5 基于因子图的状态估计

- 非线性最小二乘估计
- 直接求解方法
- 优化方法
- 各种优化方法对比
- 常用优化工具

图优化工具g2o (轻量级)

非线性优化工具Ceres-Solver (复杂但性能强)

增量优化工具GTSAM (精度稍低但支持增量优化)



g2o框架的C++类

来源于：书《机器人SLAM导航：核心技术与实战》图7-38

是一个

有一个

有多个

内容概要

7.1 SLAM发展简史

7.2 SLAM中的概率理论

7.3 估计理论

7.4 基于贝叶斯网络的状态估计

7.5 基于因子图的状态估计

7.6 典型SLAM算法

7.7 SFM、BA和SLAM比较

7.6 典型SLAM算法

		在线SLAM系统	完全SLAM系统
贝叶斯网络表示	EKF滤波	EKF-SLAM	
	粒子滤波		Fast-SLAM Gmapping
最小二乘表示	直接法		Graph-SLAM
	优化法		Cartographer LOAM ORB-SLAM
机器学习表示	端到端法	DeepVO	
	改进模块法	NeRF-SLAM NICE-SLAM	CNN-SLAM
	生物启发法	RatSLAM NeuroSLAM	

7.6 典型SLAM算法

- EKF-SLAM
- Fast-SLAM
- Graph-SLAM
- 现今主流SLAM算法

$$P(\underbrace{x_k, m_{1:L}}_{s_k} | z_{1:k}, u_{1:k}) \sim N(\mu_k, \Sigma_k)$$

其中,

$$s_k = [x_k, m_1, m_2, \dots, m_L]^T$$
$$\mu_k = [\mu_{x_k}, \mu_{m_1}, \mu_{m_2}, \dots, \mu_{m_L}]^T$$
$$\Sigma_k = \begin{bmatrix} \Sigma_{x_k} & \Sigma_{x_k m_1} & \dots & \Sigma_{x_k m_L} \\ \Sigma_{m_1 x_k} & \Sigma_{m_1} & \dots & \Sigma_{m_1 m_L} \\ \dots & \dots & \dots & \dots \\ \Sigma_{m_L x_k} & \Sigma_{m_L m_1} & \dots & \Sigma_{m_L} \end{bmatrix} = \begin{bmatrix} \Sigma_{xx} & \Sigma_{xm} \\ \Sigma_{mx} & \Sigma_{mm} \end{bmatrix}$$

$$s_k = g(s_{k-1}, u_k) + r_k \text{ 或 } \begin{bmatrix} x_k \\ m_1 \\ \dots \\ m_L \end{bmatrix} = \begin{bmatrix} g'(x_{k-1}, u_k) + r'_k \\ m_1 \\ \dots \\ m_L \end{bmatrix}$$

$$z_k = h(s_k) + q_k$$

其中 r_k 为运动过程携带的高斯噪声，协方差矩阵记为 $R_k = \begin{bmatrix} R'_k & 0 \\ 0 & 0 \end{bmatrix}$ ； q_k 为观测过程携带的高斯噪声，协方差矩阵记为 Q_k 。

在线SLAM系统求解

$$\left. \begin{aligned} & \textcircled{1} \overline{\mu}_k = g(\mu_{k-1}, u_k) \\ & \textcircled{2} \overline{\Sigma}_k = G_k \Sigma_{k-1} G_k^T + R_k \end{aligned} \right\} \text{运动预测}$$
$$\left. \begin{aligned} & \textcircled{3} K_k = \overline{\Sigma}_k H_k^T (H_k \overline{\Sigma}_k H_k^T + Q_k)^{-1} \\ & \textcircled{4} \mu_k = \overline{\mu}_k + K_k (z_k - h(\overline{\mu}_k)) \\ & \textcircled{5} \Sigma_k = (I - K_k H_k) \overline{\Sigma}_k \end{aligned} \right\} \text{卡尔曼增益} \quad \text{观测更新}$$

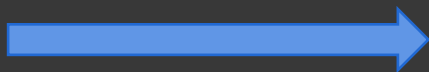
卡尔曼滤波

以EKF-SLAM为代表的**在线SLAM系统**，其EKF滤波方法是一种增量求解方法，只需要输入新量测信息就能完成对SLAM状态量的求解，这种增量求解方法具有很大优势。但是EKF-SLAM也有两个明显的缺点，一方面，其将所估计后验分布假设成高斯分布，非线性高斯系统必须被线性化近似成线性高斯系统，这种简单粗暴的**线性化近似**容易引起较大误差；另一方面，状态量的协方差矩阵的复杂度随路标数量是二次方增长，也就是说其**无法建立大规模地图**。为了解决线性化和构建大规模地图的问题，大家开始放弃EKF-SLAM这种在线SLAM增量求解的好处，转向**完全SLAM系统**的研究。

7.6 典型SLAM算法

■ EKF-SLAM

■ Fast-SLAM



完全SLAM系统求解

$$P(x_{1:k}, m_{1:L} | z_{1:k}, u_{1:k})$$

■ Graph-SLAM

■ 现今主流SLAM算法

$$P(x_{1:k}, m_{1:L} | z_{1:k}, u_{1:k}) \propto \underbrace{P(x_{1:k} | z_{1:k}, u_{1:k})}_{\text{定位问题 (采用PF实现)}} \bullet \underbrace{P(m_{1:L} | x_{1:k}, z_{1:k})}_{\text{建图问题 (采用EKF实现)}}$$

7.6 典型SLAM算法

- EKF-SLAM
- Fast-SLAM
- Graph-SLAM
- 现今主流SLAM算法

完全SLAM系统求解

$$P(x_{1:k}, m_{1:L} | z_{1:k}, u_{1:k})$$

$$P(x_{1:k}, m_{1:L} | z_{1:k}, u_{1:k}) \propto \min \sum_{i=1}^m \|e_i\|^2$$

最小二乘问题

直接法解线性方程

7.6 典型SLAM算法

- EKF-SLAM
- Fast-SLAM
- Graph-SLAM
- 现今主流SLAM算法



- 第8章：激光SLAM系统
- 第9章：视觉SLAM系统
- 第10章：其他SLAM系统

内容概要

7.1 SLAM发展简史

7.2 SLAM中的概率理论

7.3 估计理论

7.4 基于贝叶斯网络的状态估计

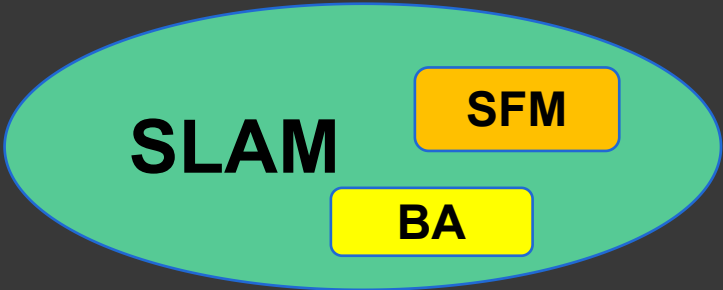
7.5 基于因子图的状态估计

7.6 典型SLAM算法

7.7 SFM、BA和SLAM比较

7.7 SFM、BA和SLAM比较

名称	描述
SFM	SFM (structure from motion, 运动重建) 是计算机视觉领域的概念, 通过收集到的无序图片进行离线三维重建, 可以理解为多个图片的拼接。SFM一般使用在比较大的范围, 比如构建一个城市的百度地图。SFM具有两个鲜明的特点, 一个特点是输入算法的图片是 无序 的, 另一个特点是重建过程 离线 进行对实时性没有要求。
BA	BA (bundle adjustment, 光束平差法) 是摄影测量领域的概念, 通过估计运动相机的位姿来反推场景中路标点的位姿, 场景路标点又可以反过来作用于相机位姿估计, 这其实就是一个局部优化问题。相机位姿与路标的通过相机观测模型进行约束关联, 其实就是用到了重投影误差这个约束。与SFM相比, 输入BA算法的图片是需要 有顺序 的。
SLAM	SLAM是移动机器人领域的概念, 利用机器人的运动模型和观测模型数据实时估计机器人位姿和路标。与BA相比, SLAM除了有观测数据还有运动数据, SLAM有回环检测能通过全局优化消除累积误差, 并且SLAM计算要求尽可能的 实时性 。从这里可以看出, SLAM其实涵盖了SFM和BA所研究的问题, 是一个更为复杂和困难的问题。



- 例程源码下载： https://github.com/xiihoo/Books_Robot_SLAM_Navigation
- 课件PPT下载： www.xiihoo.com

敬请关注,长期更新...

下集预告