

Readings:

K&F: 16.2

Jordan Chapter 20

K&F: 4.5, 12.2, 12.3

EM for BNs Kalman Filters Gaussian BNs

Graphical Models – 10708

Carlos Guestrin

Carnegie Mellon University

November 17th, 2006

1

Thus far, fully supervised learning

- We have assumed fully supervised learning:
- Many real problems have missing data:

The general learning problem with missing data

- Marginal likelihood – $\mathbf{x}$ is observed, $\mathbf{z}$ is missing:

$$\begin{aligned}\ell(\theta : \mathcal{D}) &= \log \prod_{j=1}^m P(\mathbf{x}_j | \theta) \\ &= \sum_{j=1}^m \log P(\mathbf{x}_j | \theta) \\ &= \sum_{j=1}^m \log \sum_{\mathbf{z}} P(\mathbf{x}_j, \mathbf{z} | \theta)\end{aligned}$$

10-708 – ©Carlos Guestrin 2006

3

E-step

- $\mathbf{x}$ is observed, $\mathbf{z}$ is missing
- Compute probability of missing data given current choice of θ
 - $Q(\mathbf{z} | \mathbf{x}_j)$ for each $\mathbf{x}_j$
 - e.g., probability computed during classification step
 - corresponds to “classification step” in K-means

$$Q^{(t+1)}(\mathbf{z} | \mathbf{x}_j) = P(\mathbf{z} | \mathbf{x}_j, \theta^{(t)})$$

10-708 – ©Carlos Guestrin 2006

4

Jensen's inequality

$$\ell(\theta : \mathcal{D}) = \sum_{j=1}^m \log \sum_{\mathbf{z}} P(\mathbf{z} | \mathbf{x}_j) P(\mathbf{x}_j | \theta)$$

- **Theorem:** $\log \sum_{\mathbf{z}} P(\mathbf{z}) f(\mathbf{z}) \geq \sum_{\mathbf{z}} P(\mathbf{z}) \log f(\mathbf{z})$

10-708 – ©Carlos Guestrin 2006

5

Applying Jensen's inequality

- Use: $\log \sum_{\mathbf{z}} P(\mathbf{z}) f(\mathbf{z}) \geq \sum_{\mathbf{z}} P(\mathbf{z}) \log f(\mathbf{z})$

$$\ell(\theta^{(t)} : \mathcal{D}) = \sum_{j=1}^m \log \sum_{\mathbf{z}} Q^{(t+1)}(\mathbf{z} | \mathbf{x}_j) \frac{P(\mathbf{z}, \mathbf{x}_j | \theta^{(t)})}{Q^{(t+1)}(\mathbf{z} | \mathbf{x}_j)}$$

10-708 – ©Carlos Guestrin 2006

6

The M-step maximizes lower bound on weighted data

- Lower bound from Jensen's:

$$\ell(\theta^{(t)} : \mathcal{D}) \geq \sum_{j=1}^m \sum_{\mathbf{z}} Q^{(t+1)}(\mathbf{z} | \mathbf{x}_j) \log P(\mathbf{z}, \mathbf{x}_j | \theta^{(t)}) + H(Q^{(t+1)})$$

- Corresponds to weighted dataset:

- ☐ $\langle \mathbf{x}_1, \mathbf{z}=1 \rangle$ with weight $Q^{(t+1)}(\mathbf{z}=1 | \mathbf{x}_1)$
- ☐ $\langle \mathbf{x}_1, \mathbf{z}=2 \rangle$ with weight $Q^{(t+1)}(\mathbf{z}=2 | \mathbf{x}_1)$
- ☐ $\langle \mathbf{x}_1, \mathbf{z}=3 \rangle$ with weight $Q^{(t+1)}(\mathbf{z}=3 | \mathbf{x}_1)$
- ☐ $\langle \mathbf{x}_2, \mathbf{z}=1 \rangle$ with weight $Q^{(t+1)}(\mathbf{z}=1 | \mathbf{x}_2)$
- ☐ $\langle \mathbf{x}_2, \mathbf{z}=2 \rangle$ with weight $Q^{(t+1)}(\mathbf{z}=2 | \mathbf{x}_2)$
- ☐ $\langle \mathbf{x}_2, \mathbf{z}=3 \rangle$ with weight $Q^{(t+1)}(\mathbf{z}=3 | \mathbf{x}_2)$
- ☐ ...

10-708 – ©Carlos Guestrin 2006

7

The M-step

$$\ell(\theta^{(t)} : \mathcal{D}) \geq \sum_{j=1}^m \sum_{\mathbf{z}} Q^{(t+1)}(\mathbf{z} | \mathbf{x}_j) \log P(\mathbf{z}, \mathbf{x}_j | \theta^{(t)}) + H(Q^{(t+1)})$$

- Maximization step:

$$\theta^{(t+1)} \leftarrow \arg \max_{\theta} \sum_{j=1}^m \sum_{\mathbf{z}} Q^{(t+1)}(\mathbf{z} | \mathbf{x}_j) \log P(\mathbf{z}, \mathbf{x}_j | \theta)$$

- Use expected counts instead of counts:

- ☐ If learning requires $\text{Count}(\mathbf{x}, \mathbf{z})$
- ☐ Use $E_{Q^{(t+1)}}[\text{Count}(\mathbf{x}, \mathbf{z})]$

10-708 – ©Carlos Guestrin 2006

8

Convergence of EM

- Define potential function $F(\theta, Q)$:

$$\ell(\theta : \mathcal{D}) \geq F(\theta, Q) = \sum_{j=1}^m \sum_{\mathbf{z}} Q(\mathbf{z} | \mathbf{x}_j) \log \frac{P(\mathbf{z}, \mathbf{x}_j | \theta)}{Q(\mathbf{z} | \mathbf{x}_j)}$$

- EM corresponds to coordinate ascent on F
 - Thus, maximizes lower bound on marginal log likelihood
 - As seen in Machine Learning class last semester

10-708 – ©Carlos Guestrin 2006

9

Announcements

- Lectures the rest of the semester:
 - **Special time: Monday Nov 20 - 5:30-7pm, Wean 4615A:**
Gaussian GMs & Kalman Filters
 - **Special time: Monday Nov 27 - 5:30-7pm, Wean 4615A:**
Dynamic BNs
 - Wed. 11/30, regular class time: Causality (Richard Scheines)
 - Friday 12/1, regular class time: Finish Dynamic BNs & Overview of Advanced Topics
- Deadlines & Presentations:
 - Project Poster Presentations: Dec. 1st 3-6pm (NSH Atrium)
 - Project write up: Dec. 8th by 2pm by email
 - 8 pages – limit will be strictly enforced
 - Final: Out Dec. 1st, Due Dec. 15th by 2pm (strict deadline)

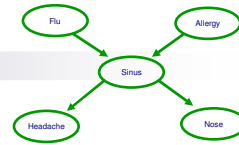
10-708 – ©Carlos Guestrin 2006

10

Data likelihood for BNs

- Given structure, log likelihood of fully observed data:

$$\log P(\mathcal{D} \mid \theta_{\mathcal{G}}, \mathcal{G})$$



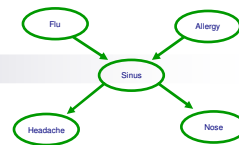
10.708 – ©Carlos Guestrin 2006

11

Marginal likelihood

- What if S is hidden?

$$\log P(\mathcal{D} \mid \theta_{\mathcal{G}}, \mathcal{G})$$



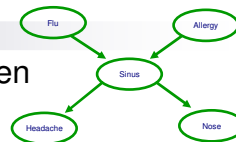
10.708 – ©Carlos Guestrin 2006

12

Log likelihood for BNs with hidden data

- Marginal likelihood – $\mathbf{O}$ is observed, $\mathbf{H}$ is hidden

$$\begin{aligned}\ell(\theta : \mathcal{D}) &= \sum_{j=1}^m \log P(\mathbf{o}^{(j)} | \theta) \\ &= \sum_{j=1}^m \log \sum_{\mathbf{h}} P(\mathbf{h}, \mathbf{o}^{(j)} | \theta)\end{aligned}$$



10-708 – ©Carlos Guestrin 2006

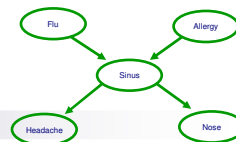
13

E-step for BNs

- E-step computes probability of hidden vars $\mathbf{h}$ given $\mathbf{o}$

$$Q^{(t+1)}(\mathbf{x} | \mathbf{o}) = P(\mathbf{x} | \mathbf{o}, \theta^{(t)})$$

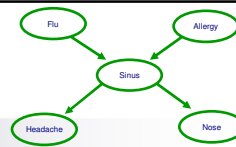
- Corresponds to inference in BN



10-708 – ©Carlos Guestrin 2006

14

The M-step for BNs



■ Maximization step:

$$\theta^{(t+1)} \leftarrow \arg \max_{\theta} \sum_{\mathbf{x}} Q^{(t+1)}(\mathbf{h} \mid \mathbf{o}) \log P(\mathbf{h}, \mathbf{o} \mid \theta)$$

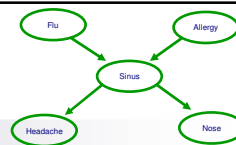
■ Use expected counts instead of counts:

- If learning requires Count(**h**,**o**)
- Use $E_{Q^{(t+1)}}[\text{Count}(\mathbf{h}, \mathbf{o})]$

10-708 – ©Carlos Guestrin 2006

15

M-step for each CPT



■ M-step decomposes per CPT

- Standard MLE:

$$P(X_i = x_i \mid \mathbf{Pa}_{X_i} = \mathbf{z}) = \frac{\text{Count}(X_i = x_i, \mathbf{Pa}_{X_i} = \mathbf{z})}{\text{Count}(\mathbf{Pa}_{X_i} = \mathbf{z})}$$

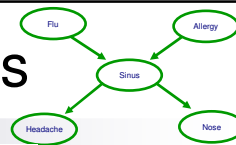
- M-step uses expected counts:

$$P(X_i = x_i \mid \mathbf{Pa}_{X_i} = \mathbf{z}) = \frac{\text{ExCount}(X_i = x_i, \mathbf{Pa}_{X_i} = \mathbf{z})}{\text{ExCount}(\mathbf{Pa}_{X_i} = \mathbf{z})}$$

10-708 – ©Carlos Guestrin 2006

16

Computing expected counts



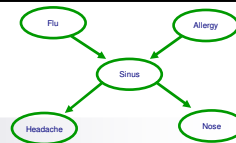
$$P(X_i = x_i \mid \text{Pa}_{X_i} = \mathbf{z}) = \frac{\text{ExCount}(X_i = x_i, \text{Pa}_{X_i} = \mathbf{z})}{\text{ExCount}(\text{Pa}_{X_i} = \mathbf{z})}$$

- M-step requires expected counts:
 - For a set of vars **A**, must compute $\text{ExCount}(\mathbf{A}=\mathbf{a})$
 - Some of **A** in example j will be observed
 - denote by $\mathbf{A}_O = \mathbf{a}_O^{(j)}$
 - Some of **A** will be hidden
 - denote by $\mathbf{A}_H$
- Use inference (E-step computes expected counts):
 - $\text{ExCount}^{(t+1)}(\mathbf{A}_O = \mathbf{a}_O^{(j)}, \mathbf{A}_H = \mathbf{a}_H) \leftarrow P(\mathbf{A}_H = \mathbf{a}_H \mid \mathbf{A}_O = \mathbf{a}_O^{(j)}, \boldsymbol{\theta}^{(t)})$

10-708 – ©Carlos Guestrin 2006

17

Data need not be hidden in the same way



- When data is fully observed
 - A data point is
- When data is partially observed
 - A data point is
- But unobserved variables can be different for different data points
 - e.g.,
- Same framework, just change definition of expected counts
 - $\text{ExCount}^{(t+1)}(\mathbf{A}_O = \mathbf{a}_O^{(j)}, \mathbf{A}_H = \mathbf{a}_H) \leftarrow P(\mathbf{A}_H = \mathbf{a}_H \mid \mathbf{A}_O = \mathbf{a}_O^{(j)}, \boldsymbol{\theta}^{(t)})$

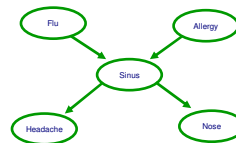
10-708 – ©Carlos Guestrin 2006

18

Learning structure with missing data

[K&F 16.6]

- Known BN structure: Use expected counts, learning algorithm doesn't change
- Unknown BN structure:
 - Can use expected counts and score model as when we talked about structure learning
 - But, very slow...
 - e.g., greedy algorithm would need to redo inference for every edge we test...
- (Much Faster) **Structure-EM**: Iterate:
 - compute expected counts
 - do a some structure search (e.g., many greedy steps)
 - repeat
- **Theorem**: Converges to local optima of marginal log-likelihood
 - details in the book



10-708 - ©Carlos Guestrin 2006

19

What you need to know about learning with missing data

- EM for Bayes Nets
- E-step: inference computes expected counts
 - Only need expected counts over X_i and $\mathbf{Pa}_{x_i}$
- M-step: expected counts used to estimate parameters
- Which variables are hidden can change per datapoint
 - Also, use labeled and unlabeled data → some data points are complete, some include hidden variables
- Structure-EM:
 - iterate between computing expected counts & many structure search steps

10-708 - ©Carlos Guestrin 2006

20

Adventures of our BN hero

- Compact representation for probability distributions
 - Fast inference
 - Fast learning
 - Approximate inference
 - But... Who are the most popular kids?
- 1. Naïve Bayes**
- 2 and 3.**
Hidden Markov models (HMMs)
Kalman Filters

21

The Kalman Filter

- An HMM with Gaussian distributions
- Has been around for at least 50 years
- Possibly the most used graphical model ever
- It's what
 - does your cruise control
 - tracks missiles
 - controls robots
 - ...
- And it's so simple...
 - Possibly explaining why it's so used
- Many interesting models build on it...
 - An example of a Gaussian BN (more on this later)

22

Example of KF – SLAT

Simultaneous Localization and Tracking

[Funiak, Guestrin, Paskin, Sukthankar '06]

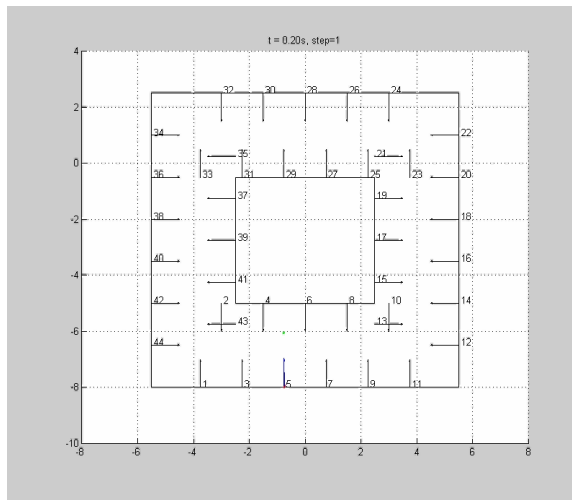
- Place some cameras around an environment, don't know where they are
- Could measure all locations, but requires lots of grad. student (Stano) time
- Intuition:
 - A person walks around
 - If camera 1 sees person, then camera 2 sees person, learn about relative positions of cameras

23

Example of KF – SLAT

Simultaneous Localization and Tracking

[Funiak, Guestrin, Paskin, Sukthankar '06]



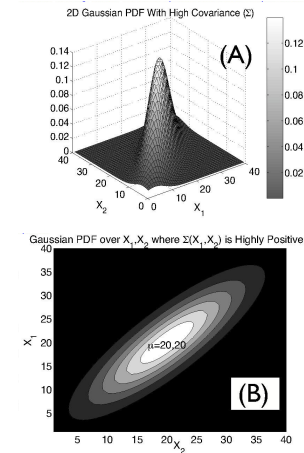
24

Multivariate Gaussian

$$p(X_1, \dots, X_n) = \frac{1}{(2\pi)^{n/2} |\Sigma|^{1/2}} \exp \left\{ -\frac{1}{2} (\mathbf{x} - \mu)^T \Sigma^{-1} (\mathbf{x} - \mu) \right\}$$

Mean vector:

Covariance matrix:



Conditioning a Gaussian

■ Joint Gaussian:

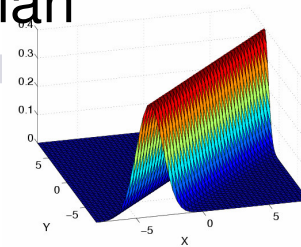
□ $p(X, Y) \sim N(\mu; \Sigma)$

■ Conditional linear Gaussian:

□ $p(Y|X) \sim N(\mu_{Y|X}; \sigma^2)$

$$\mu_{Y|X} = \mu_Y + \frac{\sigma_{YX}}{\sigma_X^2} (x - \mu_x)$$

$$\sigma_{Y|X}^2 = \sigma_Y^2 - \frac{\sigma_{YX}^2}{\sigma_X^2}$$



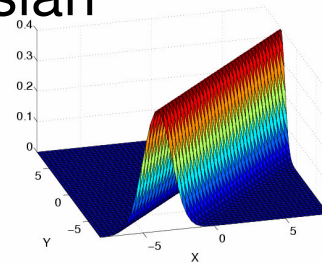
Gaussian is a “Linear Model”

- $\mu_{Y|X} = \mu_Y + \frac{\sigma_{YX}}{\sigma_X^2}(x - \mu_x)$
- Conditional linear Gaussian:
 - $p(Y|X) \sim N(\beta_0 + \beta X; \sigma^2)$
 - $\sigma_{Y|X}^2 = \sigma_Y^2 - \frac{\sigma_{YX}^2}{\sigma_X^2}$

27

Conditioning a Gaussian

- Joint Gaussian:
 - $p(X, Y) \sim N(\mu; \Sigma)$
- Conditional linear Gaussian:
 - $p(Y|X) \sim N(\mu_{Y|X}; \Sigma_{YY|X})$



$$\mu_{Y|X} = \mu_Y + \Sigma_{YX} \Sigma_{XX}^{-1} (x - \mu_x)$$

$$\Sigma_{YY|X} = \Sigma_{YY} - \Sigma_{YX} \Sigma_{XX}^{-1} \Sigma_{XY}$$

28

Conditional Linear Gaussian (CLG) – general case

■ Conditional linear Gaussian:

□ $p(Y|X) \sim \mathcal{N}(\beta_0 + BX; \Sigma_{YY|X})$

$$\mu_{Y|X} = \mu_Y + \Sigma_{YX} \Sigma_{XX}^{-1} (x - \mu_x)$$

$$\Sigma_{YY|X} = \Sigma_{YY} - \Sigma_{YX} \Sigma_{XX}^{-1} \Sigma_{XY}$$

29

Understanding a linear Gaussian – the 2d case

- Variance increases over time (motion noise adds up)
- Object doesn't necessarily move in a straight line

30

Tracking with a Gaussian 1

- $p(X_0) \sim \mathcal{N}(\mu_0, \Sigma_0)$
- $p(X_{i+1}|X_i) \sim \mathcal{N}(B X_i + \beta; \Sigma_{X_{i+1}|X_i})$

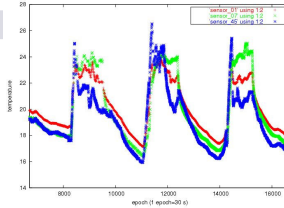
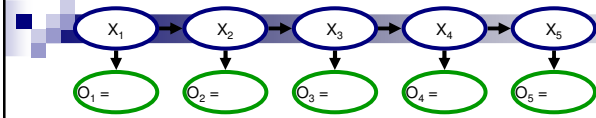
31

Tracking with Gaussians 2 – Making observations

- We have $p(X_i)$
- Detector observes $O_i=o_i$
- Want to compute $p(X_i|O_i=o_i)$
- Use Bayes rule:
- Require a CLG observation model
 - $p(O_i|X_i) \sim \mathcal{N}(W X_i + v; \Sigma_{O_i|X_i})$

32

Operations in Kalman filter



- Compute $p(X_t | O_{1:t} = o_{1:t})$
- Start with $p(X_0)$
- At each time step t :
 - **Condition** on observation
 $p(X_t | o_{1:t}) \propto p(X_t | o_{1:t-1})p(o_t | X_t)$
 - **Prediction** (Multiply transition model)
 $p(X_{t+1}, X_t | o_{1:t}) = p(X_{t+1} | X_t)p(X_t | o_{1:t})$
 - **Roll-up** (marginalize previous time step)
 $p(X_{t+1} | o_{1:t}) = \int_{X_t} p(X_{t+1}, x_t | o_{1:t}) dx_t$
- I'll describe one implementation of KF, there are others
 - Information filter

33

Exponential family representation of Gaussian: Canonical Form

$$p(X_1, \dots, X_n) = \frac{1}{(2\pi)^{n/2} |\Sigma|^{1/2}} \exp \left\{ -\frac{1}{2} (\mathbf{x} - \mu)^T \Sigma^{-1} (\mathbf{x} - \mu) \right\}$$

34

Canonical form

$$\begin{aligned}
 p(X_1, \dots, X_n) &= \frac{1}{(2\pi)^{n/2} |\Sigma|^{1/2}} \exp \left\{ -\frac{1}{2} (\mathbf{x} - \mu)^T \Sigma^{-1} (\mathbf{x} - \mu) \right\} \\
 &= K \exp \left\{ \eta^T \mathbf{x} - \frac{1}{2} \mathbf{x}^T \Lambda \mathbf{x} \right\}
 \end{aligned}$$

- Standard form and canonical forms are related:

$$\mu = \Lambda^{-1} \eta$$

$$\Sigma = \Lambda^{-1}$$

- Conditioning is easy in canonical form
- Marginalization easy in standard form

35

Conditioning in canonical form

$$p(X_t | o_{1:t}) \propto p(X_t | o_{1:t-1}) p(o_t | X_t)$$

- First multiply: $p(A, B) = p(A) p(B | A)$

$$p(A) : \quad \eta_1, \quad \Lambda_1$$

$$p(B | A) : \quad \eta_2, \quad \Lambda_2$$

$$p(A, B) : \quad \eta_3 = \eta_1 + \eta_2, \quad \Lambda_3 = \Lambda_1 + \Lambda_2$$

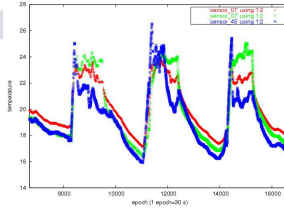
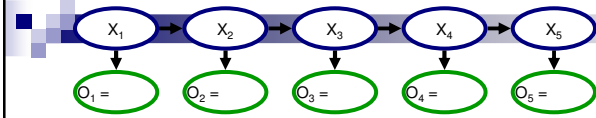
- Then, condition on value $B = y$ $p(A | B = y)$

$$\eta_{A|B=y} = \eta_A - \Lambda_{AB} \cdot y$$

$$\Lambda_{AA|B=y} = \Lambda_{AA}$$

36

Operations in Kalman filter



- Compute $p(X_t | O_{1:t} = o_{1:t})$
- Start with $p(X_0)$
- At each time step t :
 - **Condition** on observation

$$p(X_t | o_{1:t}) \propto p(X_t | o_{1:t-1})p(o_t | X_t)$$
 - **Prediction** (Multiply transition model)

$$p(X_{t+1}, X_t | o_{1:t}) = p(X_{t+1} | X_t)p(X_t | o_{1:t})$$
 - **Roll-up** (marginalize previous time step)

$$p(X_{t+1} | o_{1:t}) = \int_{X_t} p(X_{t+1}, x_t | o_{1:t}) dx_t$$

37

Prediction & roll-up in canonical form

$$p(X_{t+1} | o_{1:t}) = \int_{X_t} p(X_{t+1} | x_t) p(x_t | o_{1:t}) dx_t$$

- First multiply: $p(A, B) = p(A)p(B | A)$

- Then, marginalize X_t : $p(A) = \int_B p(A, b) db$

$$\begin{aligned} \eta_A^m &= \eta_A - \Lambda_{AB} \Lambda_{BB}^{-1} \eta_B \\ \Lambda_{AA}^m &= \Lambda_{AA} - \Lambda_{AB} \Lambda_{BB}^{-1} \Lambda_{BA} \end{aligned}$$

38

What if observations are not CLG?

- Often observations are not CLG
 - CLG if $O_i = B X_i + \beta_o + \varepsilon$
- Consider a motion detector
 - $O_i = 1$ if person is likely to be in the region
- Posterior is not Gaussian

39

Linearization: incorporating non-linear evidence

- $p(O_i|X_i)$ not CLG, but...
- Find a Gaussian approximation of $p(X_i, O_i) = p(X_i) p(O_i|X_i)$
- Instantiate evidence $O_i = o_i$ and obtain a Gaussian for $p(X_i|O_i = o_i)$
- Why do we hope this would be any good?
 - Locally, Gaussian may be OK

40

Linearization as integration

- Gaussian approximation of $p(X_i, O_i) = p(X_i) p(O_i|X_i)$
- Need to compute moments
 - $E[O_i]$
 - $E[O_i^2]$
 - $E[O_i X_i]$
- Note: Integral is product of a Gaussian with an arbitrary function

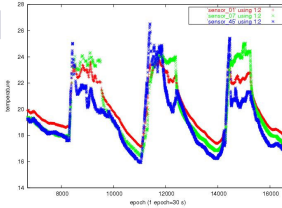
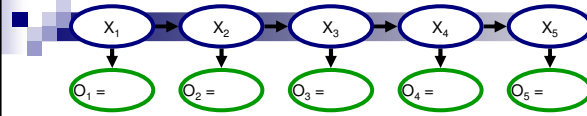
41

Linearization as numerical integration

- **Product of a Gaussian with arbitrary function**
- Effective numerical integration with **Gaussian quadrature** method
 - Approximate integral as **weighted sum over integration points**
 - Gaussian quadrature defines location of points and weights
- Exact if arbitrary function is **polynomial of bounded degree**
- **Number of integration points exponential** in number of dimensions d
- **Exact monomials** requires exponentially fewer points
 - For **$2d+1$ points**, this method is equivalent to effective **Unscented Kalman filter**
 - **Generalizes to many more points**

42

Operations in non-linear Kalman filter



- Compute $p(X_t | O_{1:t} = o_{1:t})$
- Start with $p(X_0)$
- At each time step t :
 - **Condition** on observation (use **numerical integration**)

$$p(X_t | o_{1:t}) \propto p(X_t | o_{1:t-1})p(o_t | X_t)$$
 - **Prediction** (Multiply transition model, use **numerical integration**)

$$p(X_{t+1}, X_t | o_{1:t}) = p(X_{t+1} | X_t)p(X_t | o_{1:t})$$
 - **Roll-up** (marginalize previous time step)

$$p(X_{t+1} | o_{1:t}) = \int_{X_t} p(X_{t+1}, x_t | o_{1:t}) dx_t$$

43

What you need to know about Kalman Filters

- **Kalman filter**
 - Probably most used BN
 - Assumes Gaussian distributions
 - Equivalent to linear system
 - Simple matrix operations for computations
- **Non-linear Kalman filter**
 - Usually, observation or motion model not CLG
 - Use numerical integration to find Gaussian approximation

44