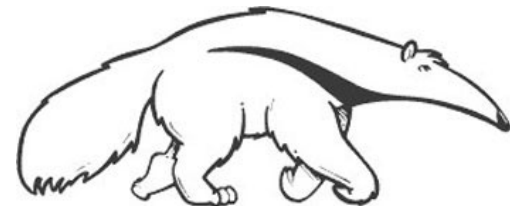


Algorithms for Causal Probabilistic Graphical Models

Class 2: **Decomposition & Variational Methods**

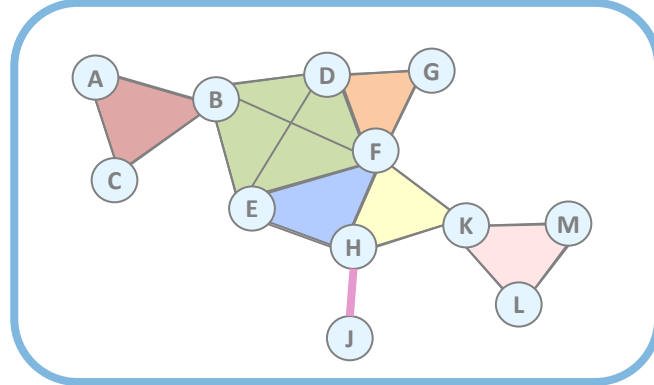
Athens Summer School on AI
July 2024

Prof. Rina Dechter
Prof. Alexander Ihler

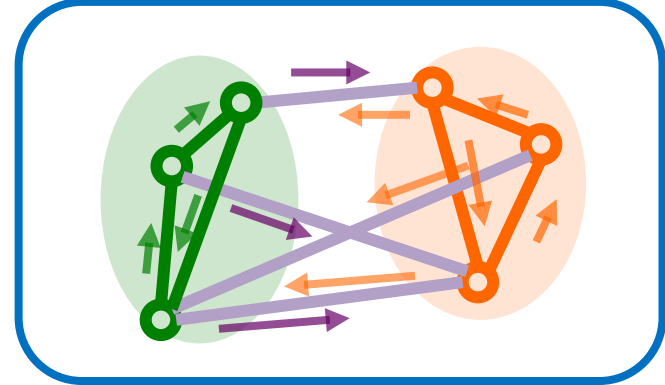


Outline of Lectures

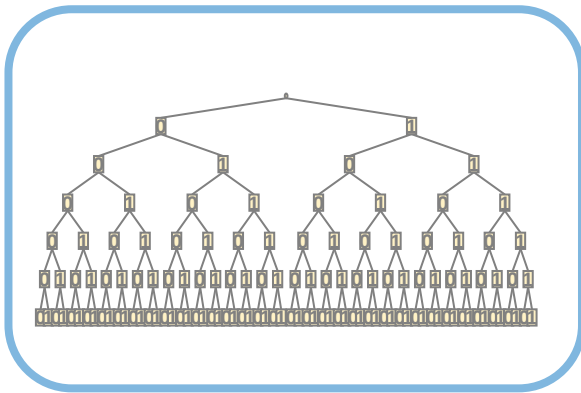
Class 1: Introduction & Inference



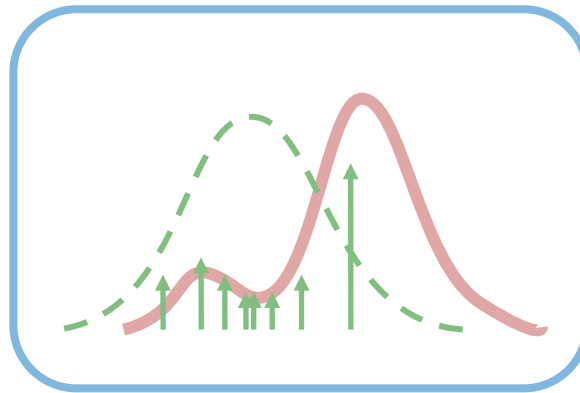
Class 2: Bounds & Variational Methods



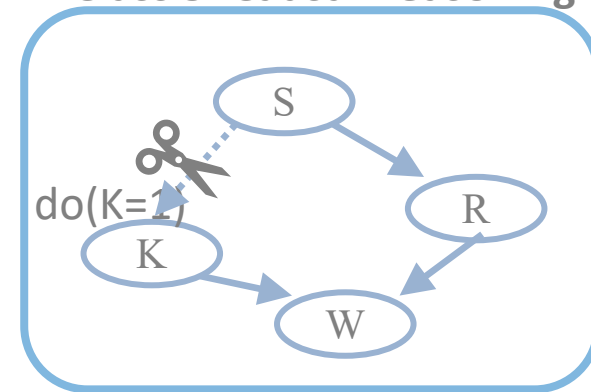
Class 3: Search Methods



Class 4: Monte Carlo Methods



Class 5: Causal Reasoning

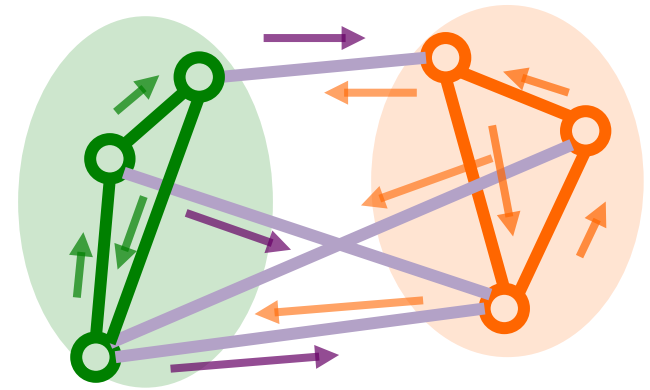


Approximate Inference

- Two main schools of approximate inference

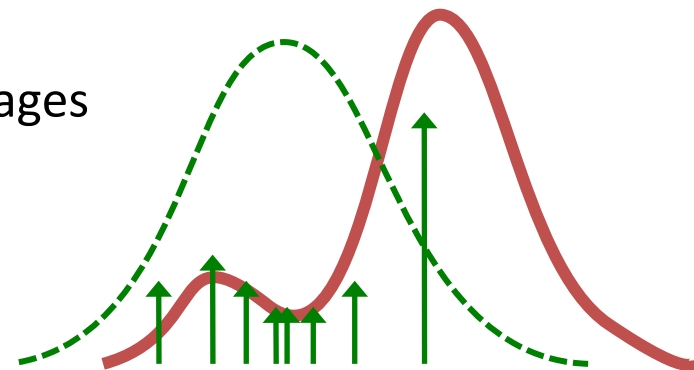
- **Variational methods** [Class 2]

- Frame “inference” as convex optimization & approximate (constraints, objectives)
- Reason about “beliefs”; pass messages
- Fast approximations & bounds
- Quality often limited by memory



- **Monte Carlo sampling** [Class 4]

- Approximate expectations with sample averages
- Estimates are asymptotically correct
- Can be hard to gauge finite sample quality



Outline

Review: Graphical Models

Decomposition Bounds

Variational Optimization

Convexity & Duality

Regions & Higher-order Approximations

Graphical models

A *graphical model* consists of:

$X = \{X_1, \dots, X_n\}$ -- variables

$D = \{D_1, \dots, D_n\}$ -- domains (we'll assume discrete)

$F = \{f_{\alpha_1}, \dots, f_{\alpha_m}\}$ -- functions or “factors”

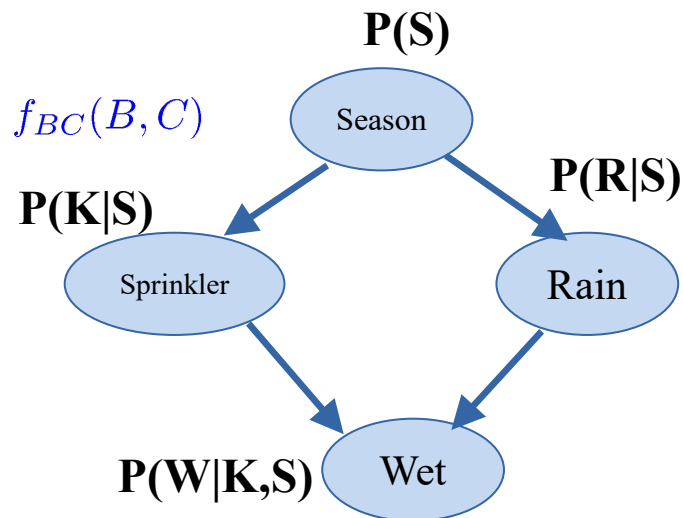
and a *combination operator*

$$A \in \{0, 1\}$$

$$B \in \{0, 1\}$$

$$C \in \{0, 1\}$$

$$f_{AB}(A, B), \quad f_{BC}(B, C)$$



The *combination operator* defines an overall function from the individual factors,

$$\text{e.g., “*”} : P(S, K, R, W) = P(S) \cdot P(K|S) \cdot P(R|S) \cdot P(W|K, S)$$

Notation:

Discrete X_i values called “states”

“Tuple” or “configuration”: states taken by a set of variables

“Scope” of f : set of variables that are arguments to a factor f

often index factors by their scope, e.g., $f_{\alpha}(X_{\alpha})$, $X_{\alpha} \subseteq X$

Canonical forms

A *graphical model* consists of:

$X = \{X_1, \dots, X_n\}$ -- variables

$D = \{D_1, \dots, D_n\}$ -- domains

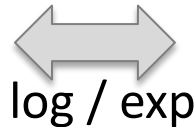
$F = \{f_{\alpha_1}, \dots, f_{\alpha_m}\}$ -- functions or “factors”

and a *combination operator*

Typically either multiplication or summation; mostly equivalent:

$$f_{\alpha}(X_{\alpha}) \geq 0$$

$$F(X) = \prod_{\alpha} f_{\alpha}(X_{\alpha})$$



$$\theta_{\alpha}(X_{\alpha}) = \log f_{\alpha}(X_{\alpha}) \in \mathbb{R}$$

$$\theta(X) = \log F(x) = \sum_{\alpha} \theta_{\alpha}(X_{\alpha})$$

Product of nonnegative factors
(probabilities, 0/1, etc.)

Sum of factors
(costs, utilities, etc.)

Probabilistic Reasoning Problems

- Exact Inference by elimination or search
- Complexity:

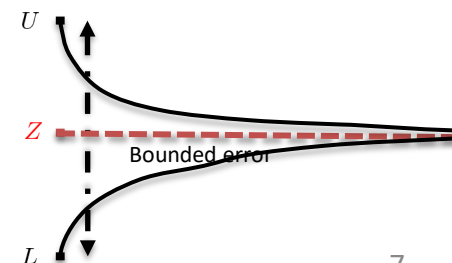
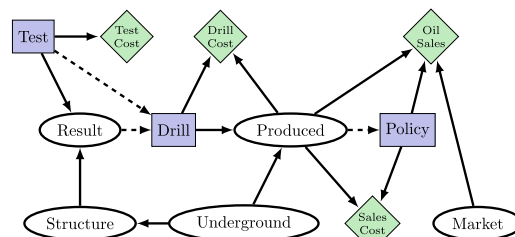
Causal effects

Max-Inference:	$f(x^*) = \max_x \prod_{\alpha} f_{\alpha}(x_{\alpha})$
Sum-Inference:	$Z = \sum_x \prod_{\alpha} f_{\alpha}(x_{\alpha})$
Mixed-Inference (MMAP):	$f_M(x_M^*) = \max_{x_M} \sum_{x_S} \prod_{\alpha} f_{\alpha}(x_{\alpha})$
Mixed-Inference (MEU):	$\text{MEU} = \max_{D_1, \dots, D_m} \sum_{X_1, \dots, X_n} \left(\prod_{P_i \in P} P_i \right) \times \left(\sum_{r_i \in R} r_i \right)$

$e^{\text{tree-width}}$

Harder

Influence
diagrams &
planning



Outline

Review: Graphical Models

Decomposition Bounds

Variational Optimization

Convexity & Duality

Regions & Higher-order Approximations

Decomposition bounds

- Upper & lower bounds via approximate problem decomposition
- Example: MAP inference $F(x) = f_1(x) + f_2(x)$

X	F(X)
0	2.0
1	4.0
2	5.0
3	4.0

 $=$

X	$f_1(X)$
0	1.0
1	2.0
2	3.0
3	4.0

 $+$

X	$f_2(X)$
0	1.0
1	2.0
2	2.0
3	0.0

$$\begin{aligned}\max_x F(x) &= \max_x [f_1(x) + f_2(x)] \\ 5.0 &\leq \left[\max_x f_1(x) + \max_x f_2(x) \right] = 4.0 + 2.0 = 6.0\end{aligned}$$

- Relaxation: two “copies” of x , no longer required to be equal
- Bound is tight (equality) if f_1, f_2 agree on maximizing value x

Mini-Bucket Approximation

Split a bucket into mini-buckets $\rightarrow$ bound complexity

bucket (X) =

$$\left\{ f_1, f_2, \dots, f_r, f_{r+1}, \dots, f_n \right\}$$

$$\lambda_X(\cdot) = \max_x \prod_{i=1}^n f_i(x, \dots)$$

$$\left\{ f_1, \dots, f_r \right\}$$

$$\lambda_{X,1}(\cdot) = \max_x \prod_{i=1}^r f_i(x, \dots)$$

$$\left\{ f_{r+1}, \dots, f_n \right\}$$

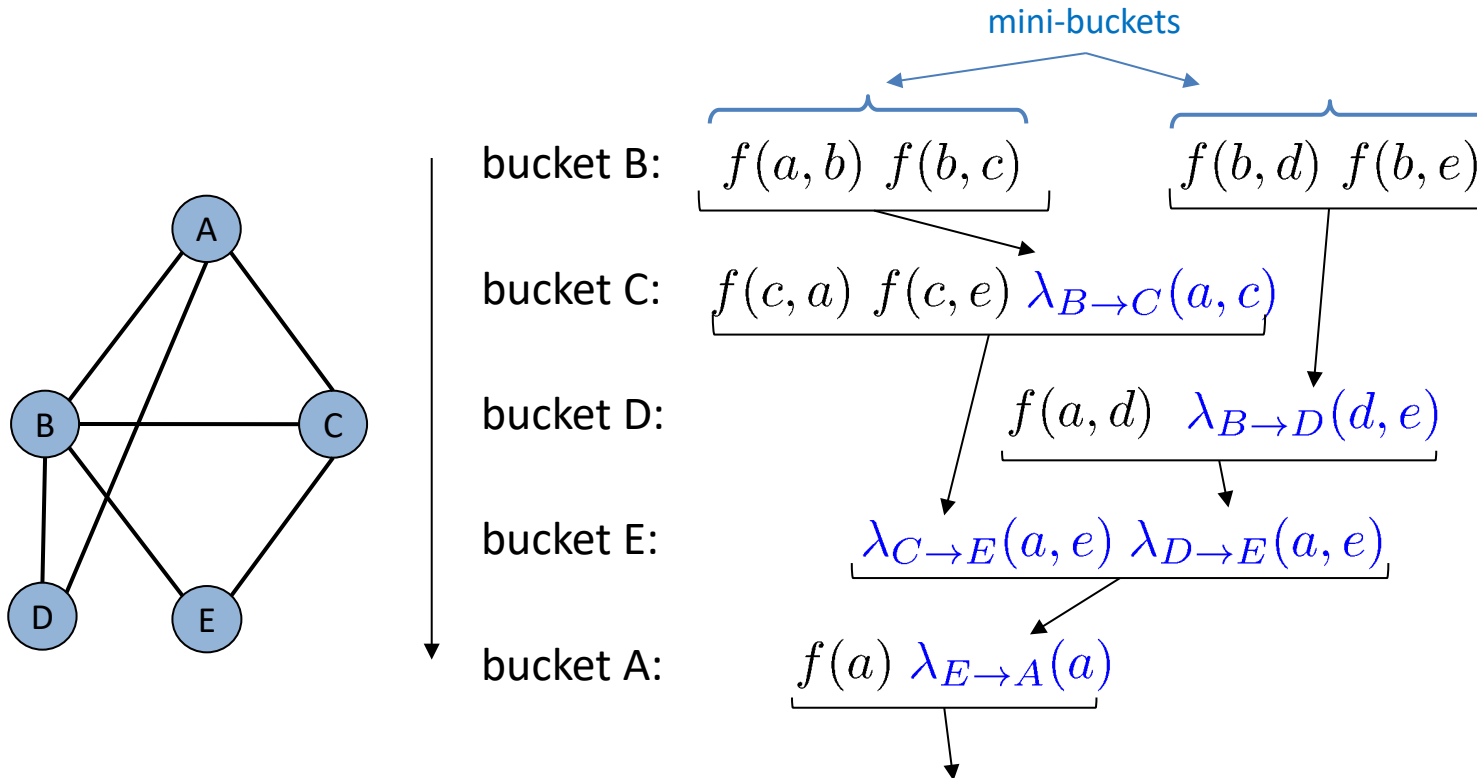
$$\lambda_{X,2}(\cdot) = \max_x \prod_{i=r+1}^n f_i(x, \dots)$$

$$\lambda_X(\cdot) \leq \lambda_{X,1}(\cdot) \lambda_{X,2}(\cdot)$$

Exponential complexity decrease: $O(e^n) \rightarrow O(e^r) + O(e^{n-r})$

Mini-Bucket Elimination

[Dechter & Rish 2003]



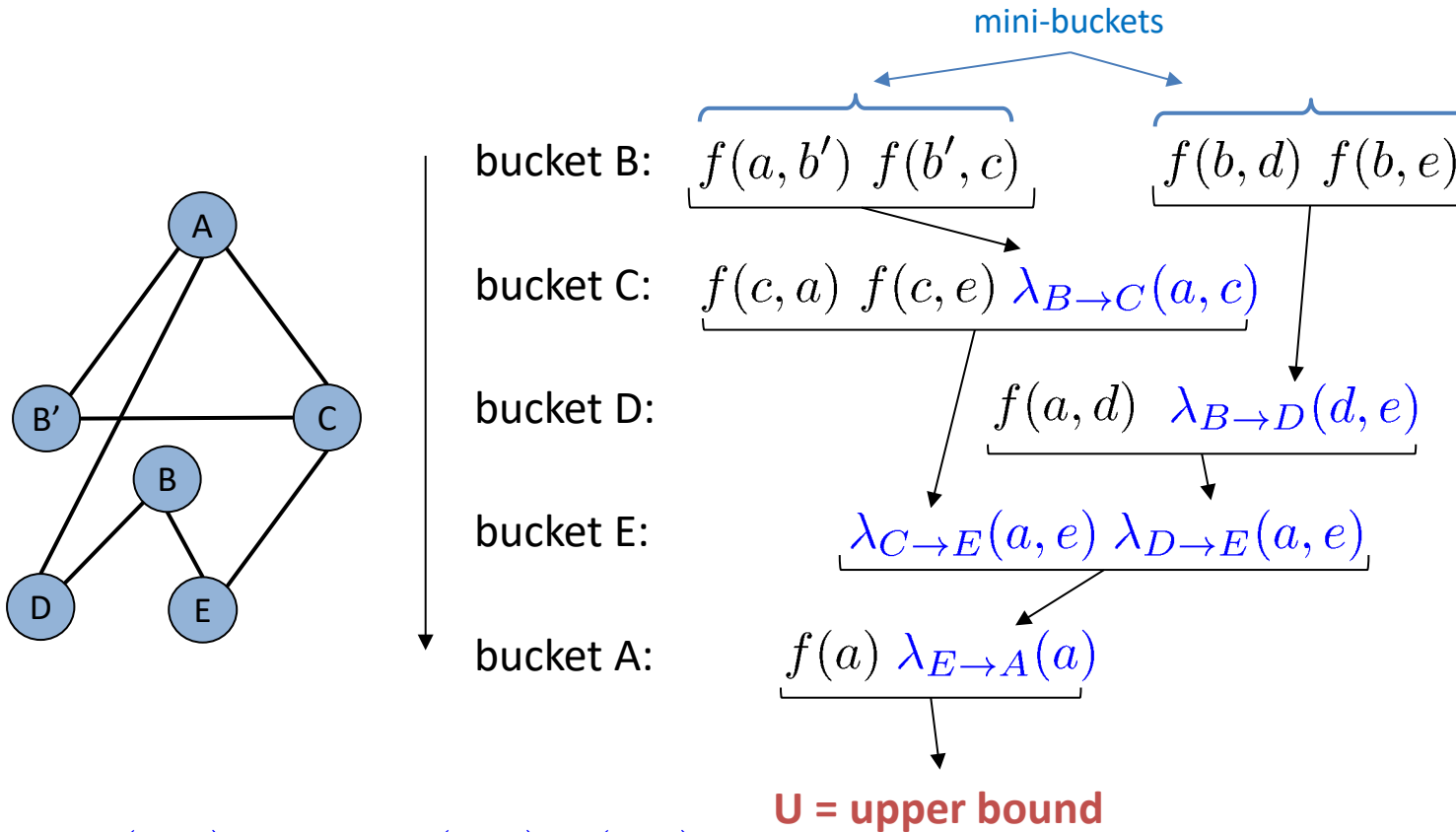
$$\lambda_{B \rightarrow C}(a, c) = \max_b f(a, b) f(b, c)$$

$$\lambda_{B \rightarrow D}(d, e) = \max_b f(b, d) f(b, e)$$

$$\lambda_{C \rightarrow E}(a, e) = \max_c \dots$$

Mini-Bucket Elimination

[Dechter & Rish 2003]



$$\lambda_{B \rightarrow C}(a, c) = \max_b f(a, b) f(b, c)$$

$$\lambda_{B \rightarrow D}(d, e) = \max_b f(b, d) f(b, e)$$

$$\lambda_{C \rightarrow E}(a, e) = \max_c \dots$$

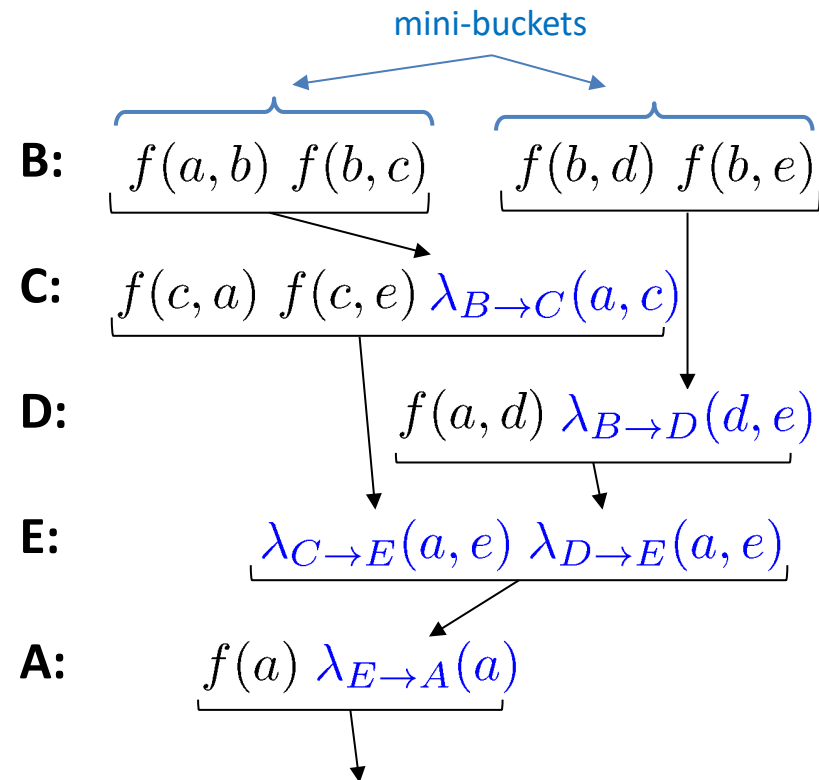
Can interpret process as “duplicating” B
[Kask et al. 2001, Geffner et al. 2007,
Choi et al. 2007, Johnson et al. 2007]

Mini-Bucket Decoding

- Assign values in reverse order using approximate messages

$$\begin{aligned}
 \mathbf{b}^* &= \arg \max_b f(a^*, b) \cdot f(b, c^*) \\
 &\quad \cdot f(b, d^*) \cdot f(b, e^*) \\
 \mathbf{c}^* &= \arg \max_c f(c, a^*) \cdot f(c, e^*) \cdot \lambda_{B \rightarrow C}(a^*, c) \\
 \mathbf{d}^* &= \arg \max_d f(a^*, d) \cdot \lambda_{B \rightarrow D}(d, e^*) \\
 \mathbf{e}^* &= \arg \max_e \lambda_{C \rightarrow E}(a^*, e) \cdot \lambda_{D \rightarrow E}(a^*, e) \\
 \mathbf{a}^* &= \arg \max_a f(a) \cdot \lambda_{E \rightarrow A}(a)
 \end{aligned}$$

Greedy configuration = lower bound



U = upper bound

Properties of MBE(i)

- **Complexity:** $O(r \exp(i))$ time and $O(\exp(i))$ space
- Yields a lower bound and an upper bound
- **Accuracy:** determined by upper/lower (U/L) bound
- Possible use of mini-bucket approximations
 - As **anytime algorithms**
 - As **heuristics** in search
- Other tasks (similar mini-bucket approximations)
 - Belief updating, Marginal MAP, MEU, WCSP, Max-CSP
[Dechter and Rish, 1997], [Liu and Ihler, 2011], [Liu and Ihler, 2013]

Tightening the bound

- Reparameterization (or, “cost shifting”)
 - Decrease bound without changing overall function

$$f_{AB}(a, b) + f_{BC}(b, c)$$

A	B	$f_1(A, B)$
0	0	2.0
1	0	3.5
0	1	1.0
1	1	3.0

+

B	C	$f_2(B, C)$
0	0	1.0
0	1	0.0
1	0	1.0
1	1	3.0

=

A	B	C	$F(A, B, C)$
0	0	0	3.0
0	0	1	2.0
0	1	0	2.0
0	1	1	4.0
1	0	0	4.5
1	0	1	3.5
1	1	0	4.0
1	1	1	6.0

$$\max_{a,b} f_1(a, b) + \lambda_{B \rightarrow AB}(b)$$

$$+ \max_{b,c} f_2(b, c) + \lambda_{B \rightarrow BC}(b)$$

$$\lambda_{B \rightarrow AB}(b) + \lambda_{B \rightarrow BC}(b) = 0$$

A	B	$f_1(A, B)$	$\lambda(B)$
0	0	2.0	0
1	0	3.5	
0	1	1.0	+1
1	1	3.0	

+

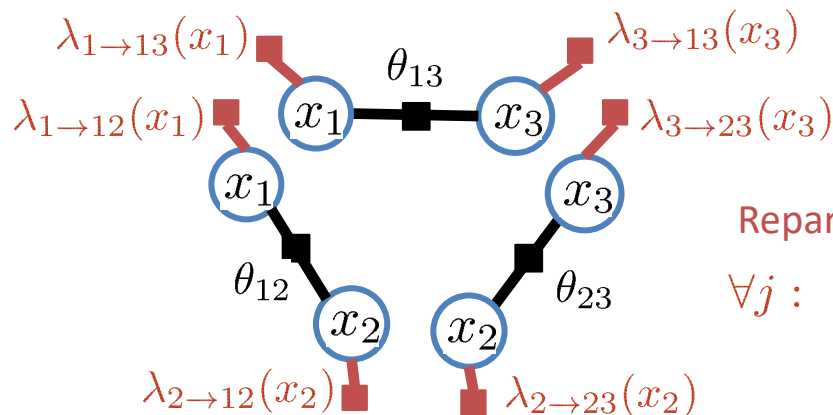
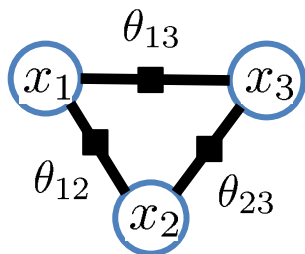
B	C	$f_2(B, C)$	$-\lambda(B)$
0	0	1.0	0
0	1	0.0	
1	0	1.0	-1
1	1	3.0	

(Adjusting functions
cancel each other)

(Decomposition bound is exact)

Decomposition for MAP

Add factors that “adjust”
each local term, but
cancel out in total



Reparameterization:

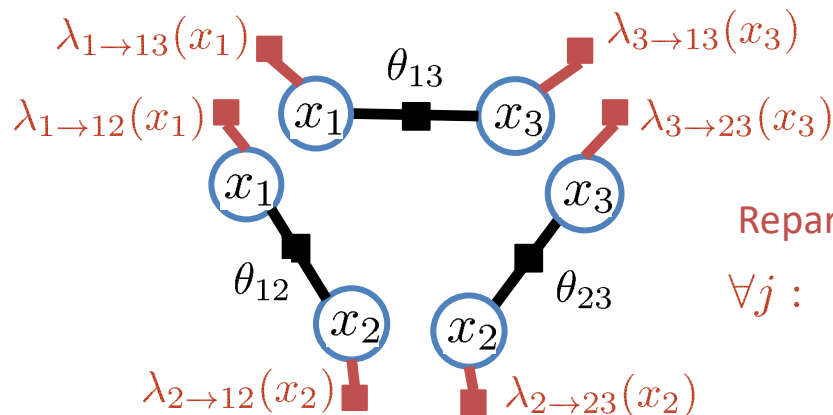
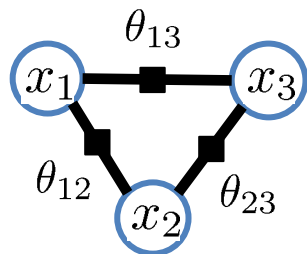
$$\forall j : \sum_{\alpha \ni j} \lambda_{j \rightarrow \alpha}(x_j) = 0$$

$$\log f(\mathbf{x}^*) = \max_{\mathbf{x}} \sum_{\alpha} \theta_{\alpha}(\mathbf{x}_{\alpha}) \leq \min_{\{\lambda_{i \rightarrow \alpha}\}} \sum_{\alpha} \max_{\mathbf{x}_{\alpha}} \left[\theta_{\alpha}(\mathbf{x}_{\alpha}) + \sum_{i \in \alpha} \lambda_{i \rightarrow \alpha}(x_i) \right]$$

- Bound solution using decomposed optimization
- Solve independently: optimistic bound
- Tighten the bound by reparameterization
 - Enforces lost equality constraints using Lagrange multipliers

Decomposition for MAP

Add factors that “adjust”
each local term, but
cancel out in total



Reparameterization:

$$\forall j : \sum_{\alpha \ni j} \lambda_{j \rightarrow \alpha}(x_j) = 0$$

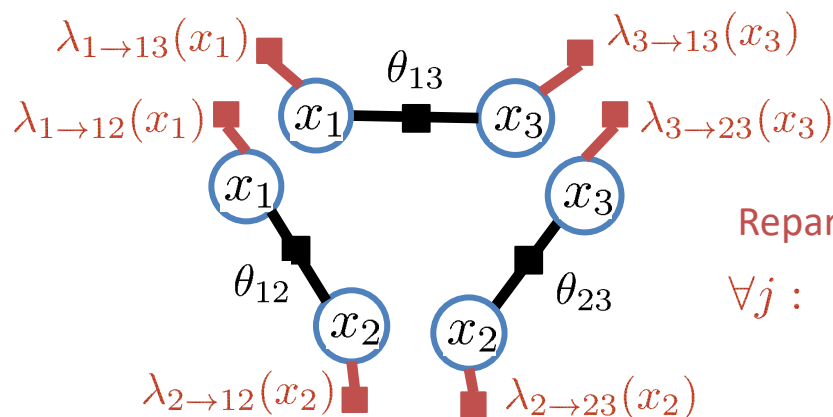
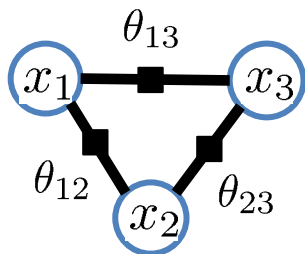
$$\log f(\mathbf{x}^*) = \max_{\mathbf{x}} \sum_{\alpha} \theta_{\alpha}(\mathbf{x}_{\alpha}) \leq \min_{\{\lambda_{i \rightarrow \alpha}\}} \sum_{\alpha} \max_{\mathbf{x}_{\alpha}} \left[\theta_{\alpha}(\mathbf{x}_{\alpha}) + \sum_{i \in \alpha} \lambda_{i \rightarrow \alpha}(x_i) \right]$$

- Many names for the same class of bounds

- Dual decomposition [Komodakis et al. 2007]
- TRW, MPLP [Wainwright et al. 2005; Globerson & Jaakkola 2007]
- Soft arc consistency [Cooper & Schiex 2004]
- Max-sum diffusion [Warner 2007]

Decomposition for MAP

Add factors that “adjust”
each local term, but
cancel out in total



Reparameterization:

$$\forall j : \sum_{\alpha \ni j} \lambda_{j \rightarrow \alpha}(x_j) = 0$$

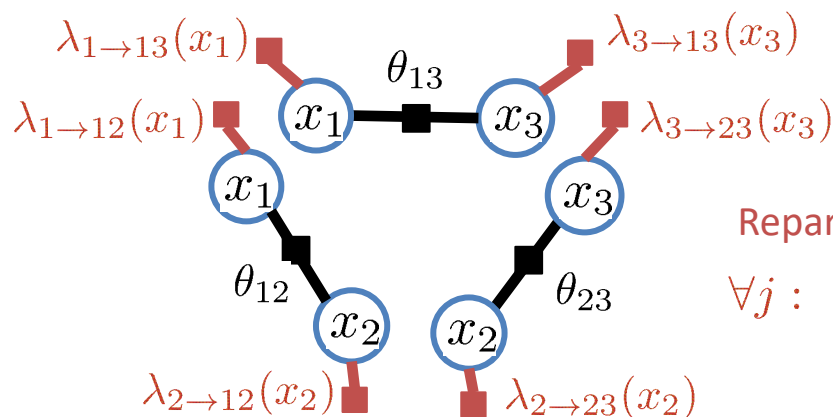
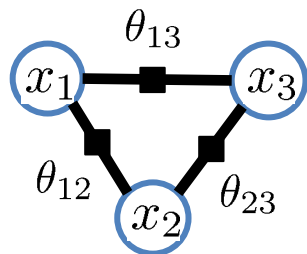
$$\log f(\mathbf{x}^*) = \max_{\mathbf{x}} \sum_{\alpha} \theta_{\alpha}(\mathbf{x}_{\alpha}) \leq \min_{\{\lambda_{i \rightarrow \alpha}\}} \sum_{\alpha} \max_{\mathbf{x}_{\alpha}} \left[\theta_{\alpha}(\mathbf{x}_{\alpha}) + \sum_{i \in \alpha} \lambda_{i \rightarrow \alpha}(x_i) \right]$$

- Many ways to optimize the bound:

- Sub-gradient descent [Komodakis et al. 2007; Jojic et al. 2010]
- Coordinate descent [Warner 2007; Globerson & Jaakkola 2007; Sontag 2009; Ihler et al. 2012]
- Proximal optimization [Ravikumar et al. 2010]
- ADMM [Meshi & Globerson 2011; Martins et al. 2011; Forouzan & Ihler 2013]

Decomposition for MAP

Add factors that “adjust”
each local term, but
cancel out in total



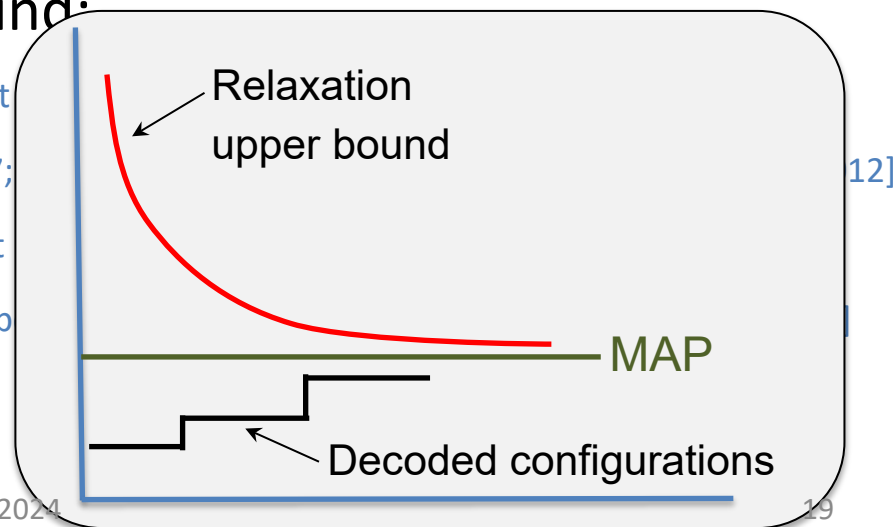
Reparameterization:

$$\forall j : \sum_{\alpha \ni j} \lambda_{j \rightarrow \alpha}(x_j) = 0$$

$$\log f(\mathbf{x}^*) = \max_{\mathbf{x}} \sum_{\alpha} \theta_{\alpha}(\mathbf{x}_{\alpha}) \leq \min_{\{\lambda_{i \rightarrow \alpha}\}} \sum_{\alpha} \max_{\mathbf{x}_{\alpha}} \left[\theta_{\alpha}(\mathbf{x}_{\alpha}) + \sum_{i \in \alpha} \lambda_{i \rightarrow \alpha}(x_i) \right]$$

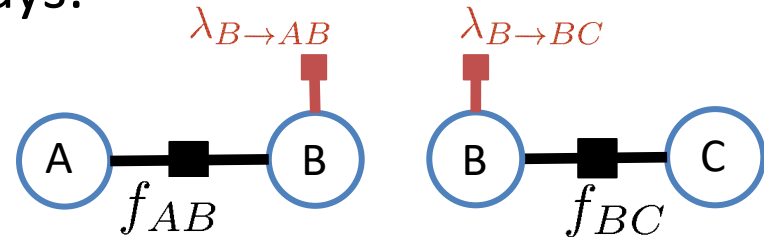
- Many ways to optimize the bound:

- Sub-gradient descent [Komodakis et al.]
- Coordinate descent [Warner 2007;]
- Proximal optimization [Ravikumar et al.]
- ADMM [Meshi & Globerson]



Optimizing the bound

- Can optimize the bound in various ways:
 - (Sub-)gradient descent



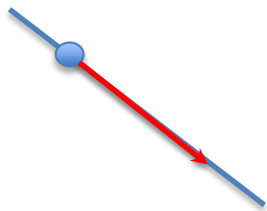
$$=$$

A	B	$f_1(A,B)$	$\lambda(B)$
0	0	1.0	0
1	0	0.0	
0	1	0.0	0
1	1	2.5	
0	2	1.0	0
1	2	3.0	

$$+$$

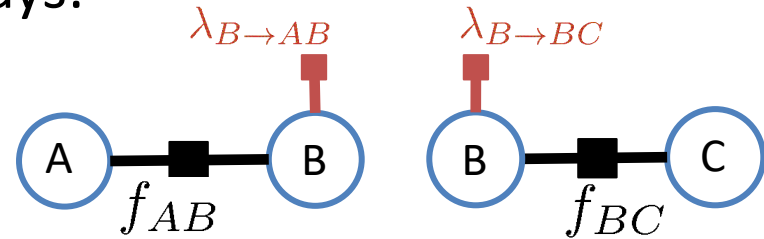
B	C	$f_2(B,C)$	$-\lambda(B)$
0	0	5.0	0
0	1	2.0	
1	0	1.0	0
1	1	1.5	
2	0	0.2	0
2	1	0.0	

$$\max_x f_1(a, b) + \max_x f_2(b, c) + \lambda_{B \rightarrow AB}(b) + \lambda_{B \rightarrow BC}(b)$$



Optimizing the bound

- Can optimize the bound in various ways:
 - (Sub-)gradient descent



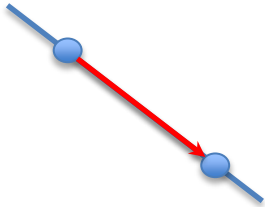
$$=$$

A	B	$f_1(A,B)$	$\lambda(B)$
0	0	1.0	+1
1	0	0.0	
0	1	0.0	0
1	1	2.5	
0	2	1.0	-1
1	2	3.0	

$$+$$

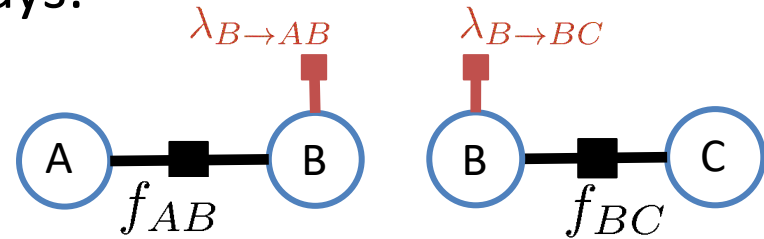
B	C	$f_2(B,C)$	$-\lambda(B)$
0	0	5.0	-1
0	1	2.0	
1	0	1.0	0
1	1	1.5	
2	0	0.2	+1
2	1	0.0	

$$\max_x f_1(a, b) + \max_x f_2(b, c) + \lambda_{B \rightarrow AB}(b) + \lambda_{B \rightarrow BC}(b)$$



Optimizing the bound

- Can optimize the bound in various ways:
 - (Sub-)gradient descent



$$=$$

A	B	$f_1(A,B)$	$\lambda(B)$
0	0	1.0	+1
1	0	0.0	
0	1	0.0	0
1	1	2.5	
0	2	1.0	-1
1	2	3.0	

$$+$$

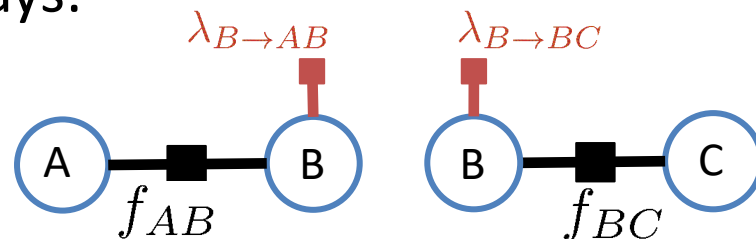
B	C	$f_2(B,C)$	$-\lambda(B)$
0	0	5.0	-1
0	1	2.0	
1	0	1.0	0
1	1	1.5	
2	0	0.2	+1
2	1	0.0	

$$\max_x f_1(a, b) + \max_x f_2(b, c) + \lambda_{B \rightarrow AB}(b) + \lambda_{B \rightarrow BC}(b)$$



Optimizing the bound

- Can optimize the bound in various ways:
 - (Sub-)gradient descent



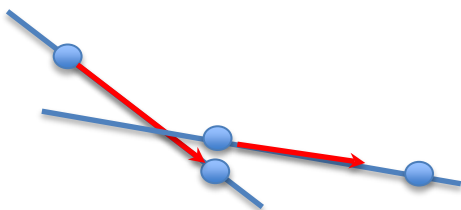
$$=$$

A	B	$f_1(A,B)$	$\lambda(B)$
0	0	1.0	+2
1	0	0.0	
0	1	0.0	-1
1	1	2.5	
0	2	1.0	-1
1	2	3.0	

$$+$$

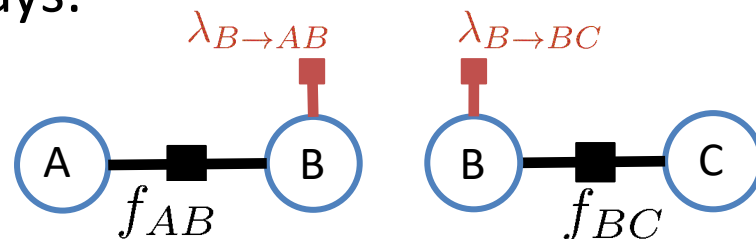
B	C	$f_2(B,C)$	$-\lambda(B)$
0	0	5.0	-2
0	1	2.0	
1	0	1.0	+1
1	1	1.5	
2	0	0.2	+1
2	1	0.0	

$$\max_x f_1(a, b) + \max_x f_2(b, c) + \lambda_{B \rightarrow AB}(b) + \lambda_{B \rightarrow BC}(b)$$



Optimizing the bound

- Can optimize the bound in various ways:
 - (Sub-)gradient descent



Both parts agree on the optima value(s): zero subgradient

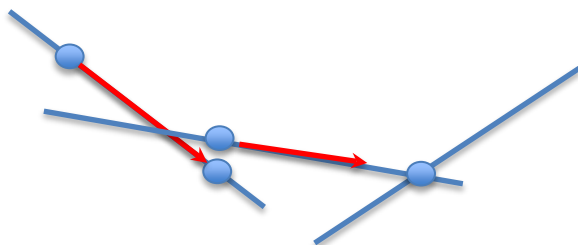
=

A	B	$f_1(A,B)$	$\lambda(B)$
0	0	1.0	+2
1	0	0.0	
0	1	0.0	-1
1	1	2.5	
0	2	1.0	-1
1	2	3.0	

+

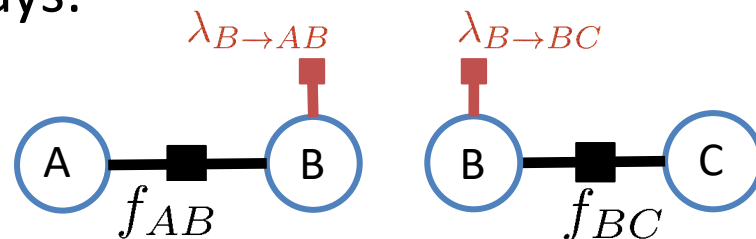
B	C	$f_2(B,C)$	$-\lambda(B)$
0	0	5.0	-2
0	1	2.0	
1	0	1.0	+1
1	1	1.5	
2	0	0.2	+1
2	1	0.0	

$$\max_x f_1(a, b) + \max_x f_2(b, c) + \lambda_{B \rightarrow AB}(b) + \lambda_{B \rightarrow BC}(b)$$



Optimizing the bound

- Can optimize the bound in various ways:
 - (Sub-)gradient descent
 - Coordinate descent



Easy to minimize over a single variable, e.g. B:

Find maxima for each B
Match values between f 's

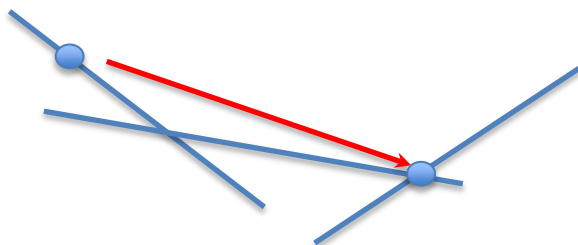
=

A	B	$f_1(A,B)$	$\lambda(B)$
0	0	1.0	
1	0	0.0	
0	1	0.0	
1	1	2.5	
0	2	1.0	
1	2	3.0	

+

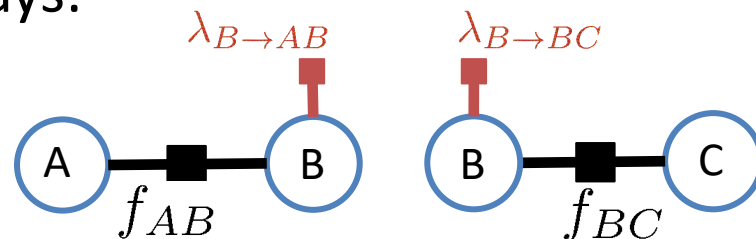
B	C	$f_2(B,C)$	$-\lambda(B)$
0	0	5.0	
0	1	2.0	
1	0	1.0	
1	1	1.5	
2	0	0.2	
2	1	0.0	

$$\max_x f_1(a, b) + \max_x f_2(b, c) + \lambda_{B \rightarrow AB}(b) + \lambda_{B \rightarrow BC}(b)$$



Optimizing the bound

- Can optimize the bound in various ways:
 - (Sub-)gradient descent
 - Coordinate descent



Easy to minimize over a single variable, e.g. B:

Find maxima for each B
Match values between f 's

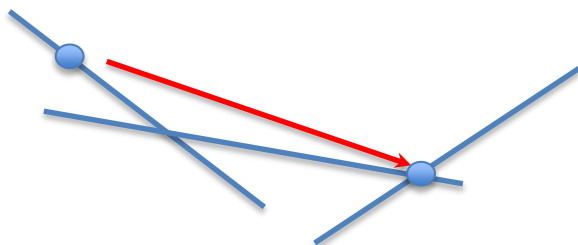
=

A	B	$f_1(A,B)$	$\lambda(B)$
0	0	1.0	- 0.5
1	0	0.0	+2.5
0	1	0.0	- 1.25
1	1	2.5	+0.75
0	2	1.0	- 1.5
1	2	3.0	+0.1

+

B	C	$f_2(B,C)$	$-\lambda(B)$
0	0	5.0	+0.5
0	1	2.0	- 2.5
1	0	1.0	+1.25
1	1	1.5	- 0.75
2	0	0.2	+1.5
2	1	0.0	- 0.1

$$\max_x f_1(a, b) + \lambda_{B \rightarrow AB}(b) + \max_x f_2(b, c) + \lambda_{B \rightarrow BC}(b)$$



Mini-Bucket as Decomposition

[Ihler et al. 2012]

$$\max_{a,c,b} \log \left[f(a,b) \cdot f(b,c) / \lambda_{B \rightarrow C}(a,c) \right] = 0$$

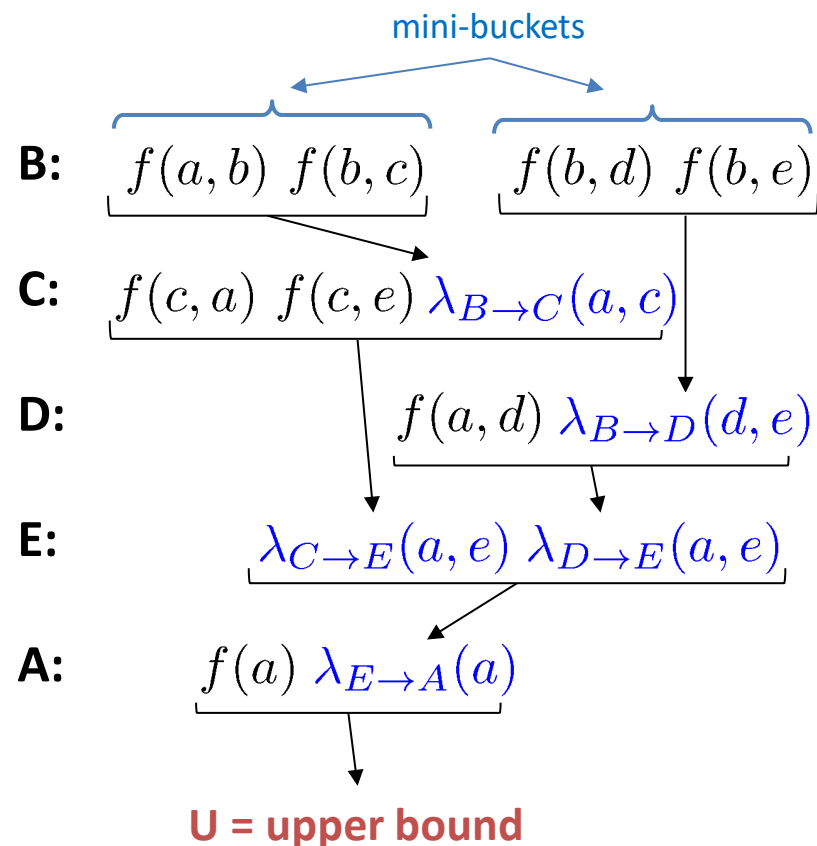
$$\max_{b,d,e} \log \left[f(b,d) \cdot f(b,e) / \lambda_{B \rightarrow D}(d,e) \right] = 0$$

$$\max_{a,e,c} \log \left[f(c,a) f(c,e) \lambda_{B \rightarrow C} / \lambda_{C \rightarrow E} \right] = 0$$

$$\max_{a,d,e} \log \left[f(a,d) \lambda_{B \rightarrow D} / \lambda_{D \rightarrow E} \right] = 0$$

$$\max_{a,d} \log \left[\lambda_{C \rightarrow E} \lambda_{D \rightarrow E} / \lambda_{E \rightarrow A} \right] = 0$$

$$\max_a \log \left[f(a) \lambda_{E \rightarrow A}(a) \right] = \log U$$

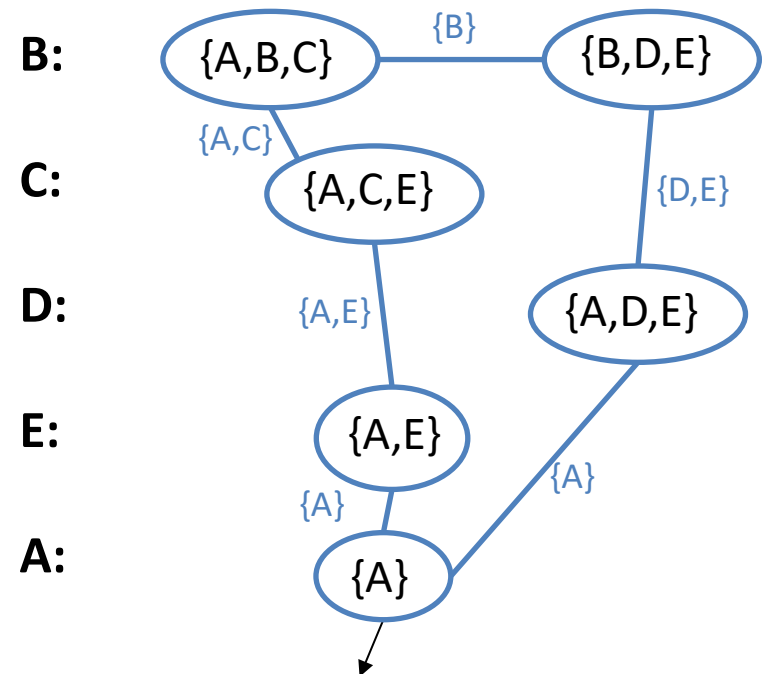


Mini-Bucket as Decomposition

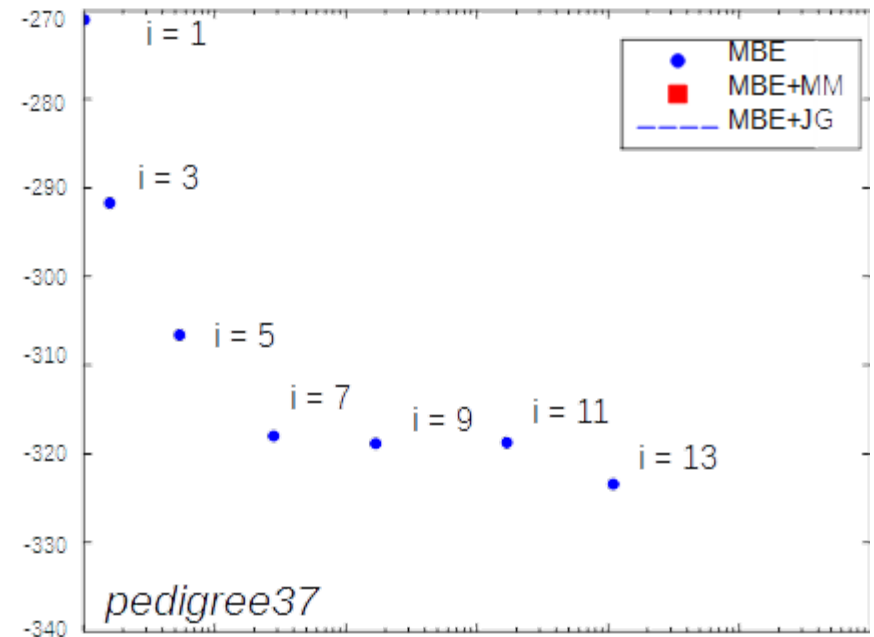
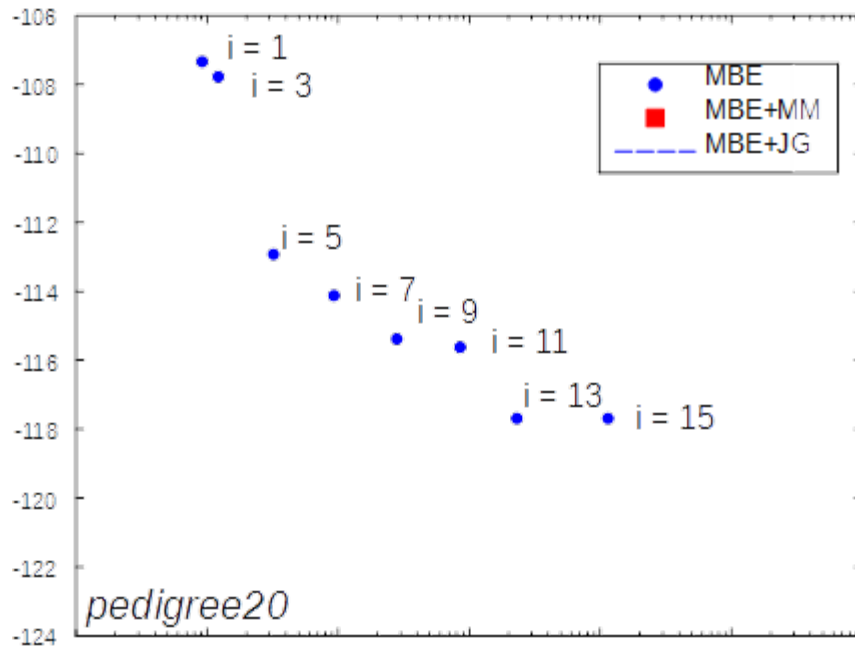
[Ihler et al. 2012]

- Downward pass as cost shifting
- Can also do cost shifting within mini-buckets:
“Join graph” message passing
- “Moment-matching” version:
One message exchange within each bucket, during downward sweep
- Optimal bound defined by cliques (“regions”) and cost-shifting f’n scopes (“coordinates”)

Join graph:

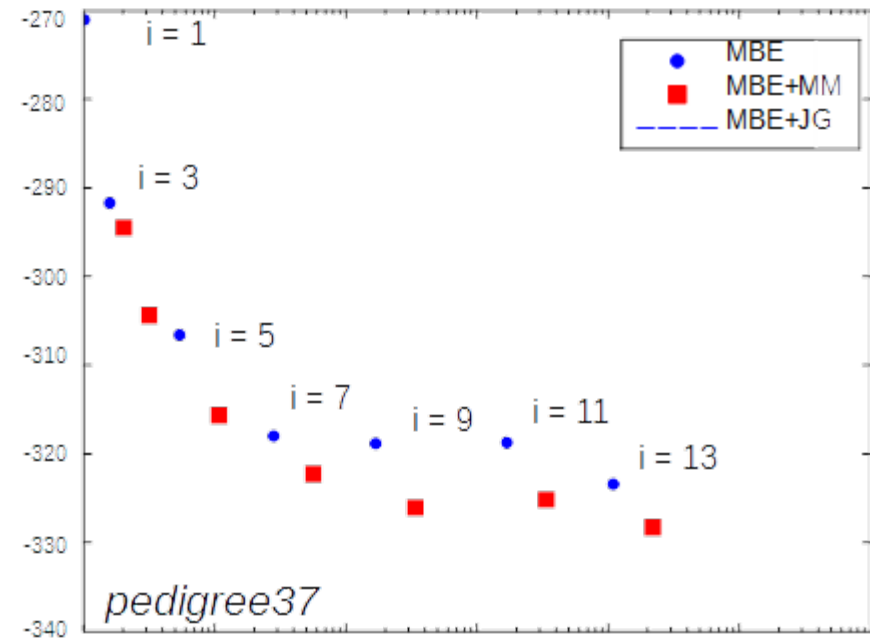
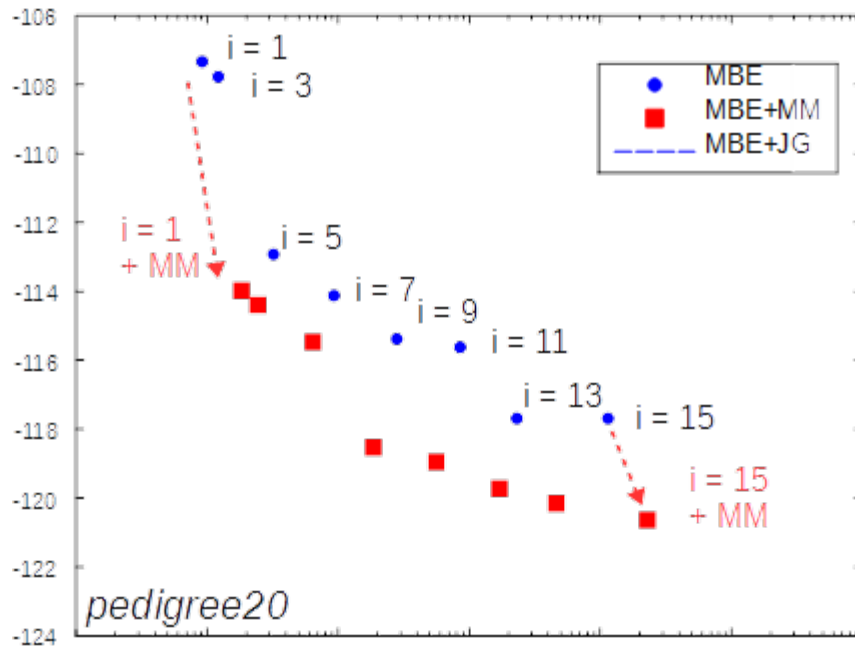


Anytime Approximation



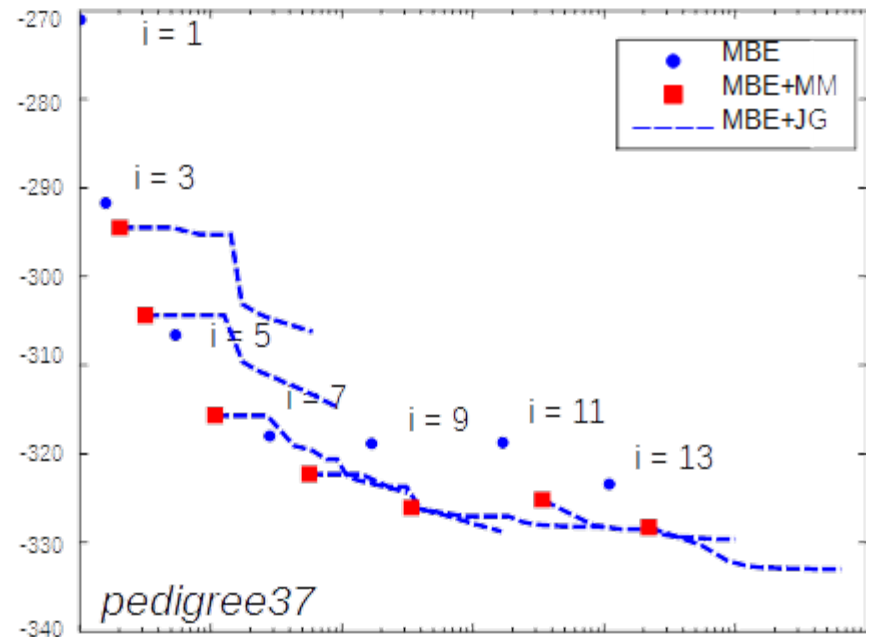
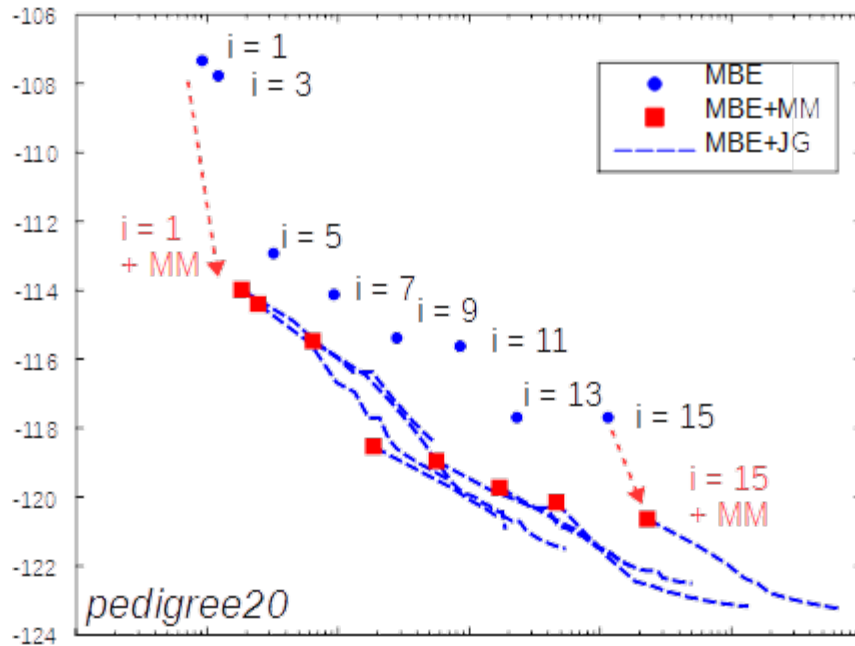
- Can tighten the bound in various ways
 - Cost-shifting (improve consistency between cliques)
 - Increase i -bound (higher order consistency)
- Simple moment-matching step improves bound significantly

Anytime Approximation



- Can tighten the bound in various ways
 - Cost-shifting (improve consistency between cliques)
 - Increase i -bound (higher order consistency)
- Simple moment-matching step improves bound significantly

Anytime Approximation



- Can tighten the bound in various ways
 - Cost-shifting (improve consistency between cliques)
 - Increase i -bound (higher order consistency)
- Simple moment-matching step improves bound significantly

Decomposition for Sum

$$F(x) = f_1(x) \cdot f_2(x)$$

- Generalize technique to sum via Holder's inequality:

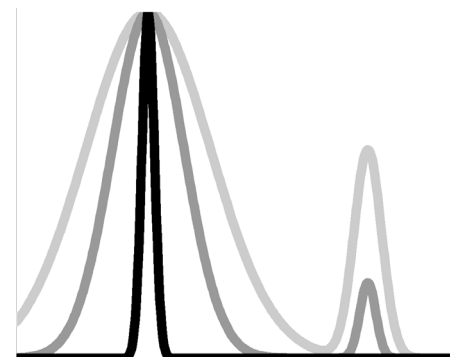
$$\sum_x f_1(x) \cdot f_2(x) \leq \left[\sum_x f_1(x)^{\frac{1}{w_1}} \right]^{w_1} \cdot \left[\sum_x f_2(x)^{\frac{1}{w_2}} \right]^{w_2} \quad w_1 + w_2 = 1$$

- Define the weighted (or powered) sum:

$$\sum_{x_1}^{w_1} f(x_1) = \left[\sum_{x_1} f(x_1)^{\frac{1}{w_1}} \right]^{w_1}$$

- “Temperature” interpolates between sum & max:
- Different weights do not commute:

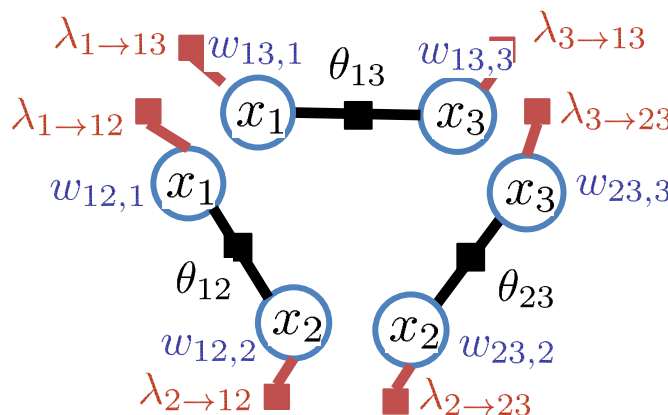
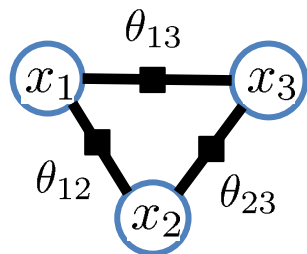
$$\sum_{x_1}^{w_1} \sum_{x_2}^{w_2} f(x_1, x_2) \neq \sum_{x_2}^{w_2} \sum_{x_1}^{w_1} f(x_1, x_2)$$



$$\lim_{w \rightarrow 0^+} \sum_x^w f(x) = \max_x f(x)$$

Decomposition for Sum

[Peng, Liu, Ihler 2015]



Reparameterization:

$$\forall j : \sum_{\alpha \ni j} \lambda_{j \rightarrow \alpha}(x_j) = 0$$

$$\log Z = \log \sum_{\mathbf{x}} \exp \left[\sum_{\alpha} \theta_{\alpha}(\mathbf{x}_{\alpha}) \right] \leq \min_{\substack{\{\lambda_{i \rightarrow \alpha}\} \\ \{w_{\alpha, i}\}}} \sum_{\alpha} \log \sum_{\mathbf{x}_{\alpha}}^{\mathbf{w}_{\alpha}} \exp \left[\theta_{\alpha}(\mathbf{x}_{\alpha} + \sum_{i \in \alpha} \lambda_{i \rightarrow \alpha}(x_i)) \right]$$

- Fixed elimination order
- Assign weight per clique & variable
- Again, tighten bound by reparameterization
 - Can also optimize over weights

Weights:

$$\forall j : \sum_{\alpha \ni j} \mathbf{w}_{\alpha, j} = 1$$

Ex: $\mathbf{w}_{12} = \begin{bmatrix} 0.5 & 0.3 & - \end{bmatrix}$
 $\mathbf{w}_{13} = \begin{bmatrix} 0.5 & - & 0.6 \end{bmatrix}$
 $\mathbf{w}_{23} = \begin{bmatrix} - & 0.7 & 0.4 \end{bmatrix}$

Weighted Mini-bucket

[Liu & Ihler 2011]

$$\lambda_{B \rightarrow C} = \sum_b^{w_{B1}} f(a, b) \cdot f(b, c) \quad w_{B1} + w_{B2} = 1$$

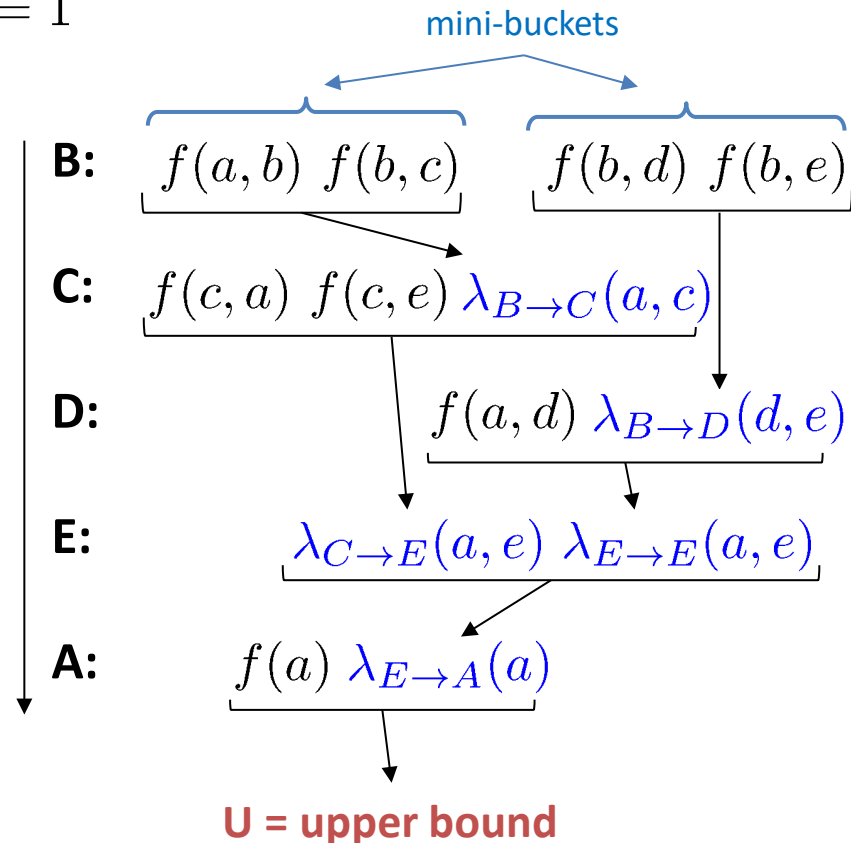
$$\lambda_{B \rightarrow D} = \sum_b^{w_{B2}} f(b, d) \cdot f(b, e)$$

$$\lambda_{C \rightarrow E} = \sum_c f(c, a) \cdot f(c, e) \cdot \lambda_{B \rightarrow C}$$

$$\vdots$$

Compute downward messages
using weighted sum

Upper bound if all weights positive
(corresponding lower bound if only one positive, rest negative)



Outline

Review: Graphical Models

Decomposition Bounds

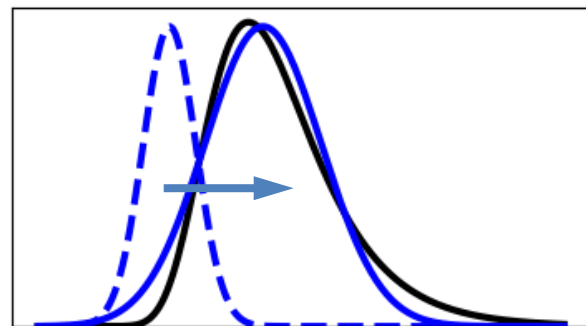
Variational Optimization

Convexity & Duality

Regions & Higher-order Approximations

Variational methods

- “Variational” = calculus of variations
 - Optimization of a “functional” (function of a function)
- Idea:
 - frame “inference” (maximization or marginals, partition f’n) as a (continuous) optimization problem
- Ex: fit a surrogate model $q(x)$; use to answer questions about $p(x)$
- Why?
 - We’re really good at continuous optimization: (stochastic) gradient descent, etc.
- Problem?
 - How can we optimize $q(x)$ without inference about $p(x)$?



Ex: BN with Evidence

- Suppose we have a Bayesian network with some evidence $E=e$

$$p(x) = \tilde{p}(x|E = e)$$

“target” distribution we’re interested in

$$\propto f(x) = \tilde{p}(x, E = e)$$

but, only able to evaluate up to a constant

$$p(x) = f(x)/Z \quad Z = p(E = e)$$

“probability of evidence”

- The KL-divergence between q & p works out very conveniently:

$$D(q||p) = -H(x; q) - \mathbb{E}_q[\log f(x)] + \log Z \geq 0$$

$$\Rightarrow \quad \underbrace{\log Z}_{\text{probability of evidence}} \geq \underbrace{H(x; q) + \mathbb{E}_q[\log f(x)]}_{\text{Evaluate or estimate from } q(x)}$$

probability
of evidence

Evaluate or estimate from $q(x)$
We can maximize this over $q(x)$!

Sometimes called the ELBO = “Evidence Lower Bound”

Stochastic Variational Inference (in Pyro)

(1) Define our target, unnormalized model (may have evidence, etc.)

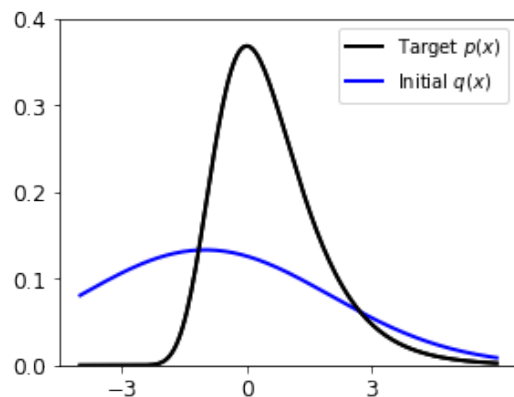
```
def model():  
    X = pyro.sample('X', dist.Gumbel(torch.tensor([0.0]), torch.tensor([1.0])))
```

(2) Define our variational approximation, $q(x)$ and initialize its parameters:

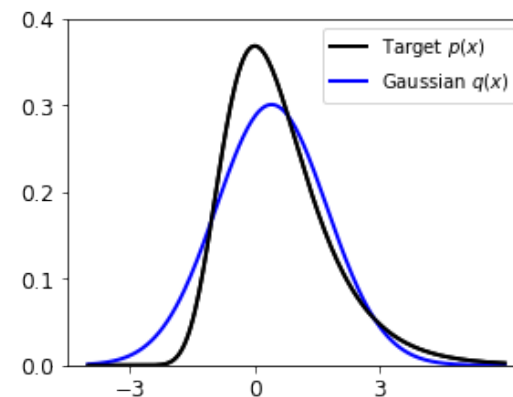
```
def guide():  
    mu = pyro.param("mu", torch.tensor(-1.0) )  
    var = pyro.param("var", torch.tensor(3.0), constraint=constraints.positive)  
    X = pyro.sample("X", dist.Normal(mu, var))
```

(3) Optimize the bound using gradient descent

```
optimizer = pyro.optim.Adam({"lr": 0.01})  
svi = pyro.infer.SVI(model, guide, optimizer, loss=pyro.infer.Trace_ELBO())  
for step in range(3000): svi.step()
```



(gradient descent)



Variational methods

- Answer queries by fitting a simpler “proxy” model
 - Optimize the KL divergence (proxy to target)

$$\begin{aligned} D(q||p) &= \sum_x q(x) \log \left[\frac{q(x)}{\frac{1}{Z} f(x)} \right] \\ &= \underbrace{-H(x; q) - \mathbb{E}_q[\log f(x)]}_{\text{Can evaluate or estimate from } q(x)} + \underbrace{\log Z}_{\text{Constant – depends only on } f(x)!} \end{aligned}$$

- Needs proxy $q(x)$ to be “easy” or “nice”!
 - What kinds of $q(x)$ are nice?
 - Need to be able to evaluate expectations & evaluate/estimate entropy
 - Continuous-valued x ? $q(x)$ Gaussian, etc.
 - Discrete x ? High-dimensional x ? Make $q(x)$ simple in terms of its graph!

Mean Field

- We can design lower bounds by restricting $q(x)$
 - Naïve mean field: $q(x)$ is fully independent
 - Entropy $H(q)$ is then easy:

$$q(x) = \prod_i q_i(x_i)$$

$$) \quad H(q) = \sum_i H(q_i)$$

- Optimizing the bound via coordinate ascent:

$$\begin{aligned}\mathbb{E}_q[\theta(x)] + H(q) &= \mathbb{E}_q \left[\sum_{\alpha \ni i} \theta_\alpha(x_\alpha) \right] + H(q_i) + \text{const} \\ &= \mathbb{E}_{q_i} [\log g(x_i)] + H(q_i) \\ &= D(q_i \parallel g_i)\end{aligned}$$

$$\log g_i(x_i) = \mathbb{E}_{q_{\neg i}} \left[\sum_{\alpha \ni i} \theta_\alpha(x_\alpha) \right]$$

$$q_{\neg i}(x) = \prod_{j \neq i} q_j(x_j)$$

Coordinate update:

$$\Rightarrow q_i(x_i) \propto \exp \left[\mathbb{E}_{q_{\neg i}} \left[\sum_{\alpha \ni i} \theta_\alpha(x_\alpha) \right] \right]$$

Mean Field

- We can design lower bounds by restricting $q(x)$

- Naïve mean field: $q(x)$ is fully independent

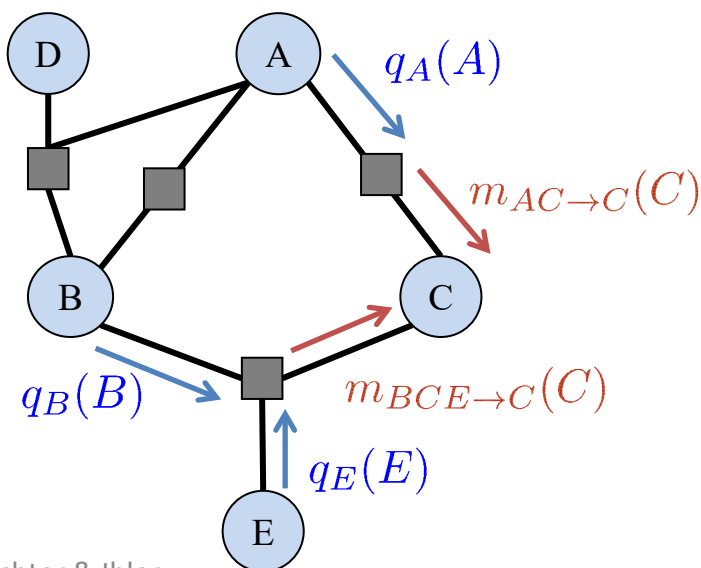
- Entropy $H(q)$ is then easy:

$$q(x) = \prod_i q_i(x_i)$$

$$\Rightarrow H(q) = \sum_i H(q_i)$$

- Optimizing the bound via coordinate ascent:

$$q_i(x_i) \propto \exp \left[\mathbb{E}_{q_{-i}} \left[\sum_{\alpha \ni i} \theta_{\alpha}(x_{\alpha}) \right] \right]$$



“Message passing” interpretation:

Updates depend only on X_i ’s Markov blanket

Naïve Mean Field

- 1: Initialize $\{q_i(X_i)\}$
 - 2: **while** not converged **do**
 - 3: **for** $i = 1 \dots n$ **do**
 - 4: $m_{\alpha \rightarrow i}(x_i) = \exp \left[\sum_{x_{\alpha \setminus i}} \theta_{\alpha}(x_{\alpha}) \prod_{j \in \alpha \setminus i} q_j(x_j) \right]$
 - 5: $q_i(x_i) \propto \prod_{\alpha \ni i} m_{\alpha \rightarrow i}(x_i)$
-

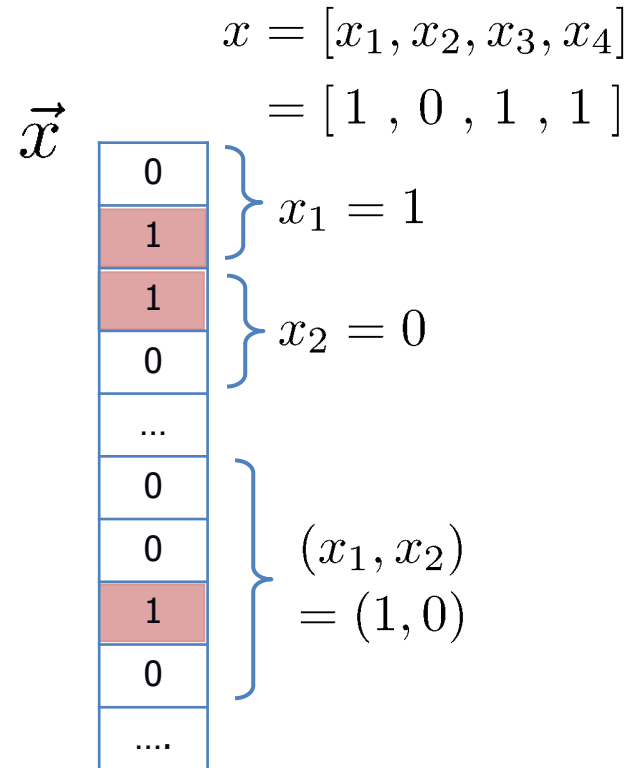
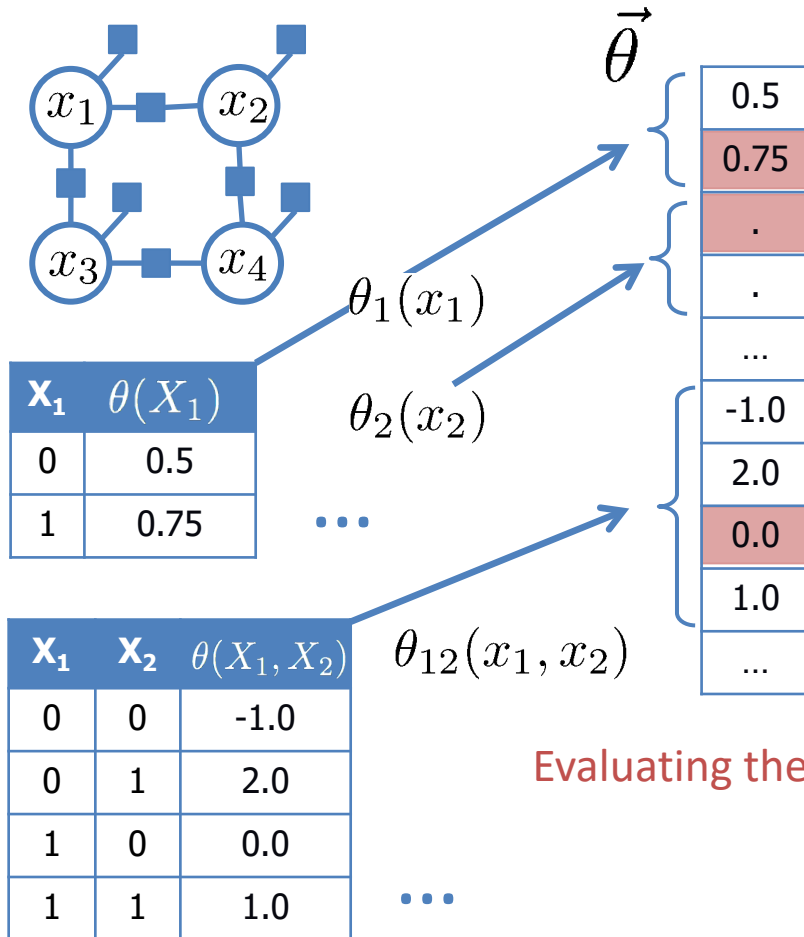
Optimization Perspective

- “Variational” = calculus of variations
 - Optimization of a “functional” (function of a function)
- **Exponential family distributions**
 - Inference tasks are **convex** in the model natural parameters!
- Very elegant perspective based on convex optimization (discussed here for background / perspective)

Vector space representation

- Represent the (log) model and state in a vector space

$$\theta(x) = \theta_1(x_1) + \theta_2(x_2) + \dots + \theta_{12}(x_1, x_2) + \dots$$



Evaluating the function is a dot product in the vector space:

$$\theta(x) = \vec{\theta} \cdot \vec{x}$$

Inference Tasks & Convexity

- Distribution is log-linear (exponential family):

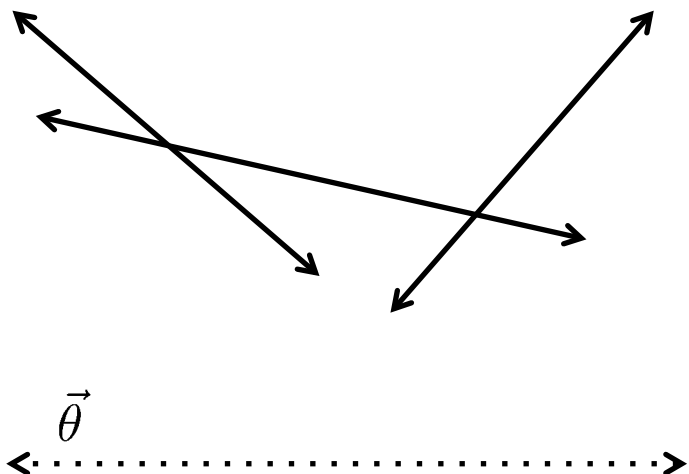
$$p(x) = \frac{1}{Z} f(x) \propto \exp [\vec{\theta} \cdot u(x)]$$

$\vec{\theta}$ “natural parameters”

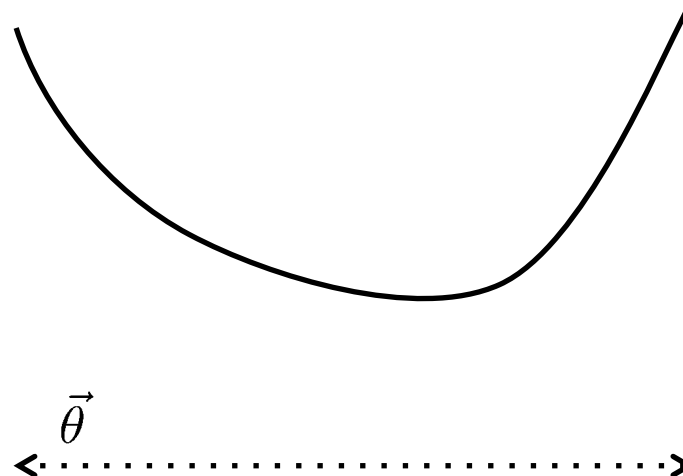
$u(x) = \vec{x}$ “features”

- Tasks of interest are convex functions of the model:

$$\log f(\mathbf{x}^*) = \max_{\vec{x} \in \mathcal{X}} \vec{\theta} \cdot \vec{x}$$



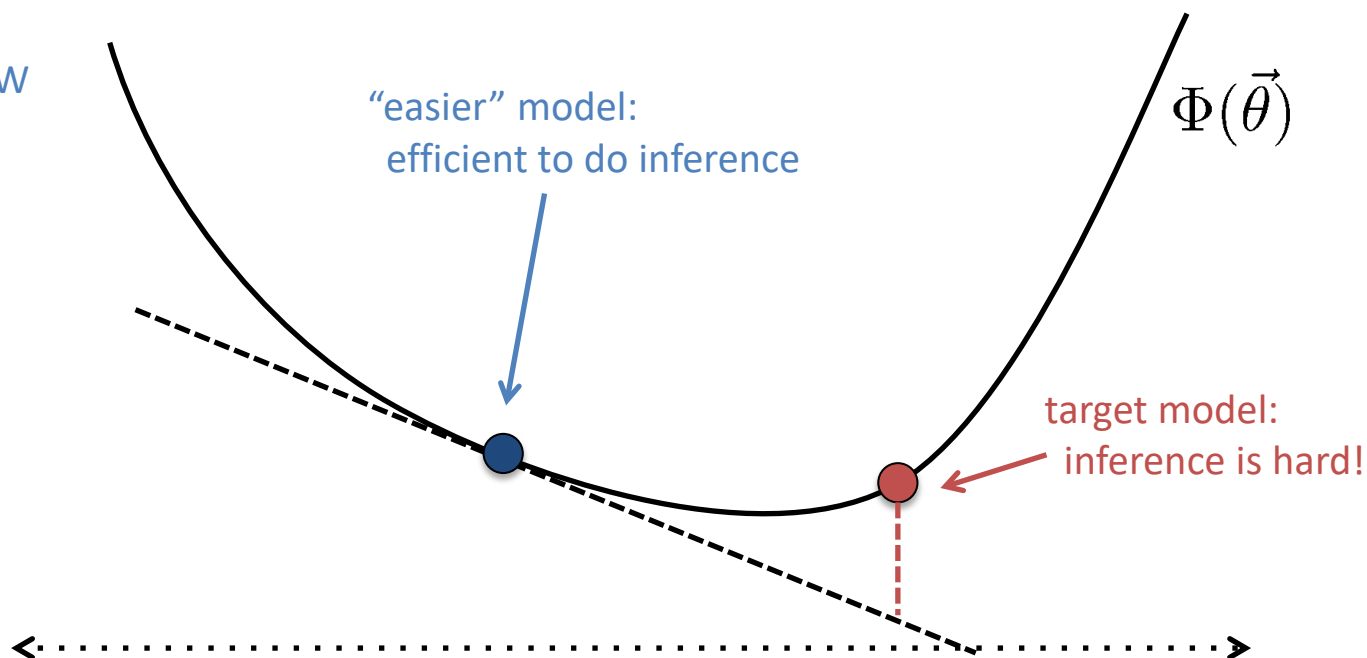
$$\log Z = \log \sum_{\vec{x} \in \mathcal{X}} \exp [\vec{\theta} \cdot \vec{x}]$$



Bounds via Convexity

- Convexity relates target to “nearby” models
 - Some of these models are easy to solve! (trees, etc.)
 - Inference at easy models + convexity tells us something about our model!
- Lower bounds:

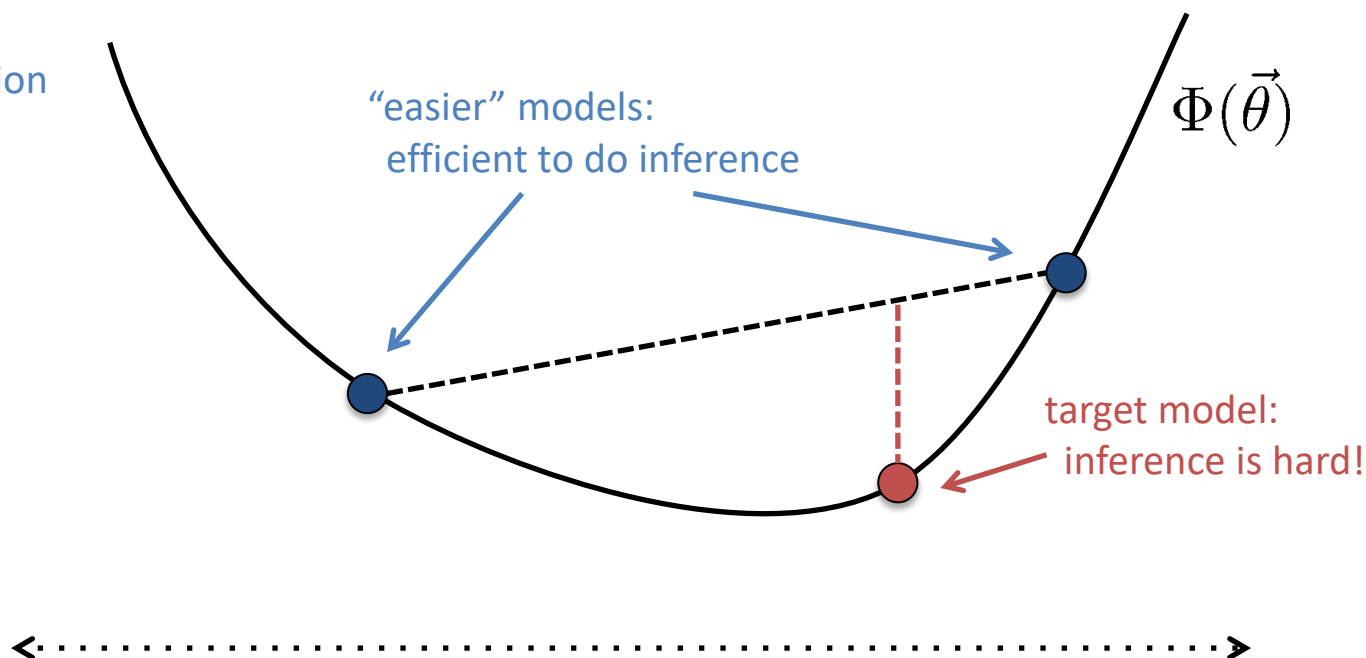
Mean field
Negative TRW
...



Bounds via Convexity

- Convexity relates target to “nearby” models
 - Some of these models are easy to solve! (trees, etc.)
 - Inference at easy models + convexity tells us something about our model!
- Upper bounds:

TRW
Decomposition
...



Tree-reweighted MAP

$$\Phi_0(\vec{\theta}) = \max_{\vec{x} \in \mathcal{X}} \vec{\theta} \cdot \vec{x}$$

- Let T_1, T_2 be two (or more) tree-structured models, with

$$\vec{\theta} = w_1 \vec{\theta}^{(1)} + w_2 \vec{\theta}^{(2)}$$

- Each T_i is easy to solve:

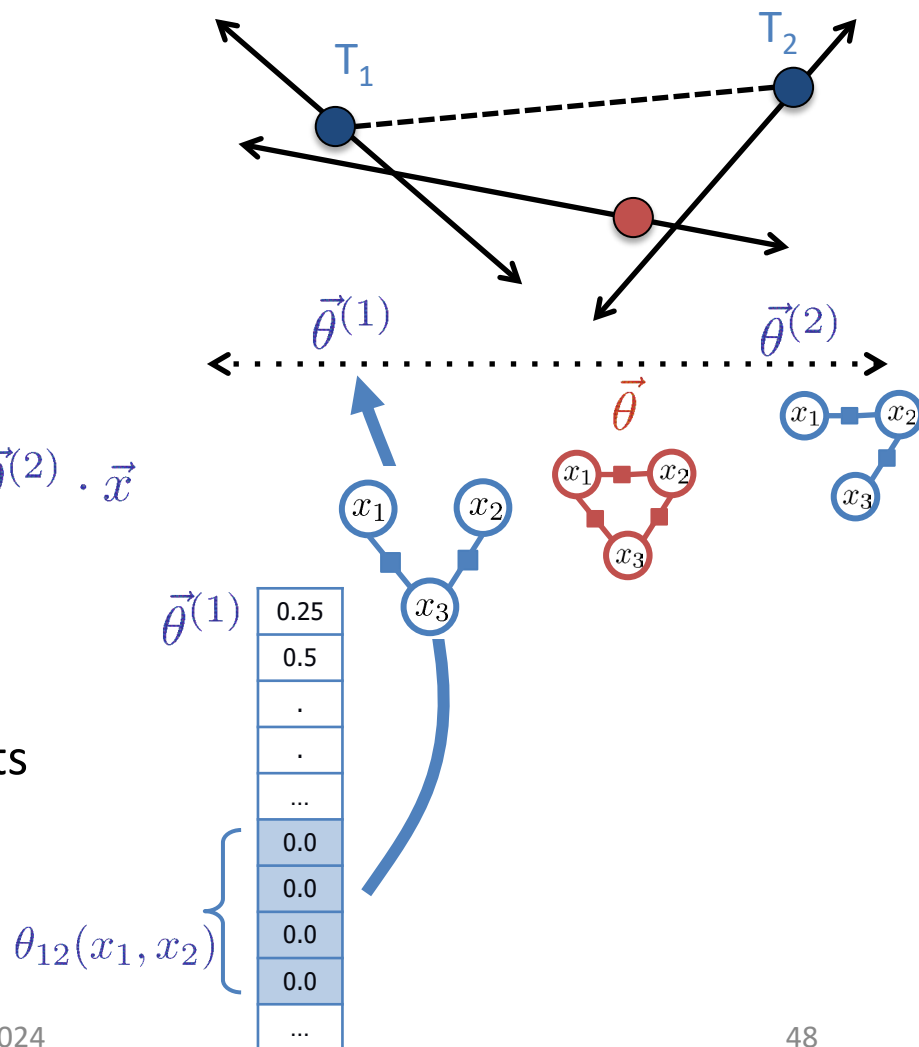
$$\vec{x}^{*(1)} = \max_{\vec{x}} \vec{\theta}^{(1)} \cdot \vec{x}$$

- And by convexity,

$$\max_{\vec{x}} \vec{\theta} \cdot \vec{x} \leq w_1 \max_{\vec{x}} \vec{\theta}^{(1)} \cdot \vec{x} + w_2 \max_{\vec{x}} \vec{\theta}^{(2)} \cdot \vec{x}$$

- Minimize bound?

- Convex objective, linear constraints



Decomposition Bounds

- TRW MAP is equivalent to MAP decomposition

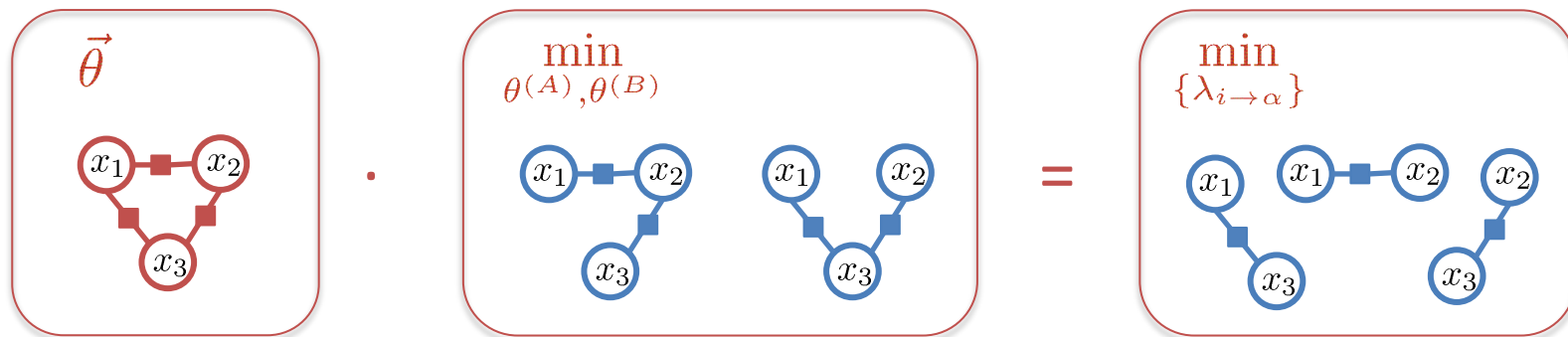
$$\begin{aligned}
 \max_{\vec{x}} [\vec{\theta} \cdot \vec{x}] &\leq \min_{\theta^{(1)}, \theta^{(2)}} \max_{\vec{x}} [w_1 \vec{\theta}^{(1)} \cdot \vec{x}] + \max_{\vec{x}} [w_2 \vec{\theta}^{(2)} \cdot \vec{x}] \\
 &= \min_{\theta^{(A)}, \theta^{(B)}} \max_{\vec{x}} [\vec{\theta}^{(A)} \cdot \vec{x}] + \max_{\vec{x}} [\vec{\theta}^{(B)} \cdot \vec{x}] \\
 &= \min_{\{\lambda_{i \rightarrow \alpha}\}} \sum_{\alpha} \max_{\vec{x}_{\alpha}} [(\vec{\theta}_{\alpha} + \sum_i \vec{\lambda}_{i \rightarrow \alpha}) \cdot \vec{x}_{\alpha}]
 \end{aligned}$$

$$\vec{\theta} = w_1 \vec{\theta}^{(1)} + w_2 \vec{\theta}^{(2)}$$

$$\vec{\theta} = \vec{\theta}^{(A)} + \vec{\theta}^{(B)}$$

$$\vec{\theta} = \sum_{\alpha \ni i} \vec{\lambda}_{i \rightarrow \alpha}$$

(on trees, decomposition bound = exact inference)



More compact
Faster optimization
Reparameterization “messages”

Tree-reweighted Sum

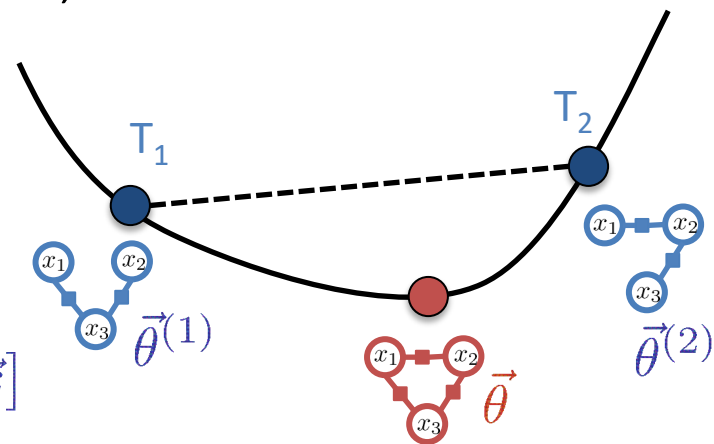
$$\Phi_1(\vec{\theta}) = \log \sum_{\vec{x} \in \mathcal{X}} \exp [\vec{\theta} \cdot \vec{x}]$$

- Let T_1, T_2 be two (or more) tree-structured models, with

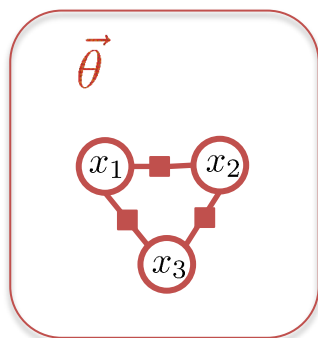
$$\vec{\theta} = w_1 \vec{\theta}^{(1)} + w_2 \vec{\theta}^{(2)} = \vec{\theta}^{(A)} + \vec{\theta}^{(B)}$$

- Again, we have

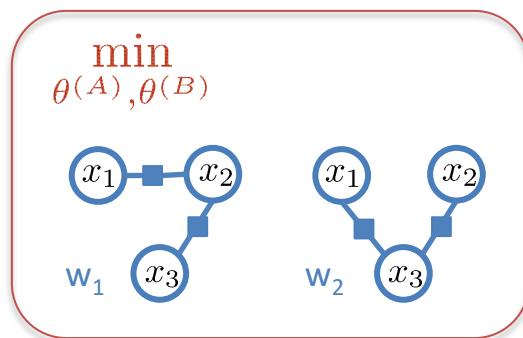
$$\begin{aligned} \Phi_1(\vec{\theta}) &\leq w_1 \Phi_1(\vec{\theta}^{(1)}) + w_2 \Phi_1(\vec{\theta}^{(2)}) \\ &= \log \sum_{\vec{x}}^{w_1} \exp [\vec{\theta}^{(A)} \vec{x}] + \log \sum_{\vec{x}}^{w_2} \exp [\vec{\theta}^{(B)} \vec{x}] \end{aligned}$$



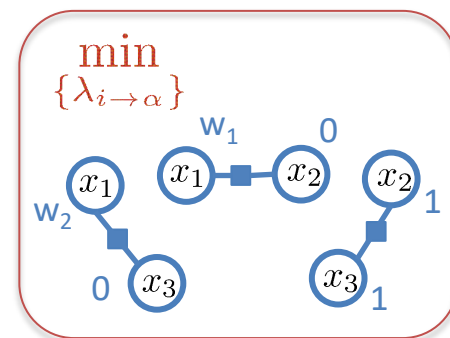
$$\sum_x^w f(x) = \left[\sum_x f(x)^{\frac{1}{w}} \right]^w$$



.



=



(if T_1, T_2 share an elimination order)

Outline

Review: Graphical Models

Decomposition Bounds

Variational Optimization

Convexity & Duality

Regions & Higher-order Approximations

Variational forms

- Reframe inference task as an optimization over distributions $q(x)$
- Ex: MAP inference $\max_x \log f(x) = \log f(x^*) = \max_{q \in \mathbb{P}} \mathbb{E}_q[\log f(x)]$

Optimal $q(x)$ puts all mass on optimal value(s) of x : $q^*(x) = \mathbb{1}[x = x^*]$
(mass on any other values of x reduces the average)

- Sum inference: $\log Z = \log \sum_x f(x) = \max_{q \in \mathbb{P}} \mathbb{E}_q[\log f(x)] + H(x; q)$

Proof:

$$D(q||p) = \sum_x q(x) \log \left[\frac{q(x)}{\frac{1}{Z} f(x)} \right] \quad (\text{Kullback-Leibler divergence})$$

$$= -H(x; q) - \mathbb{E}_q[\log f(x)] + \log Z$$

$$\Rightarrow \log Z \geq \mathbb{E}_q[\log f(x)] + H(x; q)$$

Equal iff

$$q(x) = p(x) = \frac{1}{Z} f(x)$$

- How to optimize over distributions q ?

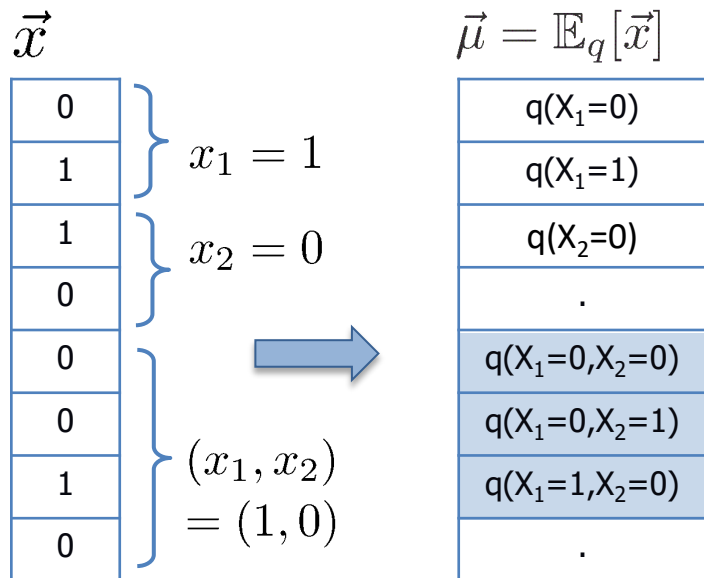
The marginal polytope

- Rewrite $\log f(x^*) = \max_{q \in \mathbb{P}} \mathbb{E}_q[\log f(x)] = \max_{q \in \mathbb{P}} \mathbb{E}_q[\vec{\theta} \cdot \vec{x}] = \max_{\vec{\mu} \in \mathcal{M}} \vec{\theta} \cdot \vec{\mu}$

and similarly, $\log Z = \max_{\vec{\mu} \in \mathcal{M}} \vec{\theta} \cdot \vec{\mu} + H(\vec{\mu})$
 (max entropy given ¹)

$$\vec{\mu} = \mathbb{E}_q[\vec{x}]$$

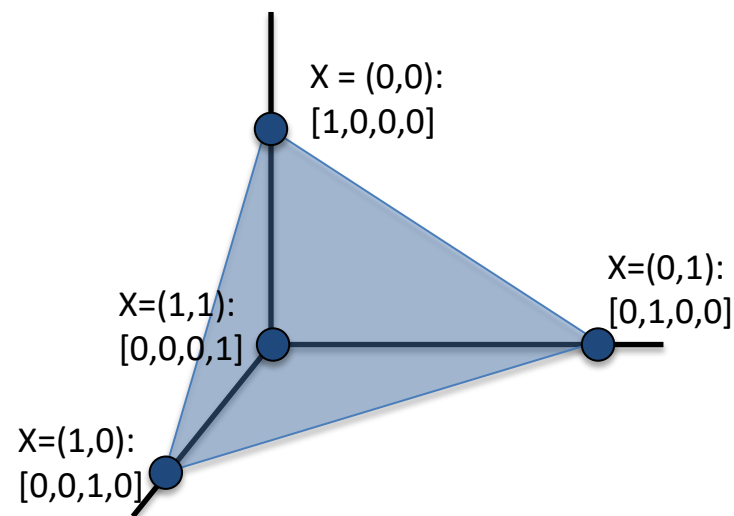
(the marginal probabilities of q)



$$\mathcal{M} = \{ \vec{\mu} : \exists q : \vec{\mu} = \mathbb{E}_q[\vec{x}] \}$$

(set of all valid marginal probabilities of q)

“marginal polytope”



Variational perspectives

- Replace $q \in \mathbb{P}$ and $H(q)$ with simpler approximations

$$\log p(x^*) = \max_{q \in \mathbb{P}} \mathbb{E}_q[\log f(x)]$$

$$\log Z = \max_{q \in \mathbb{P}} \mathbb{E}_q[\log f(x)] + H(x; q)$$

- Algorithms and their properties:

	Method	distributions	entropy	value
Max:	Linear programming	$q \in \mathbb{L} \supseteq \mathbb{P}$	n/a	$\hat{p}_{lp} \geq p(x^*)$
Sum:	Mean field	$\{q = \prod q_i(x_i)\} \subseteq \mathbb{P}$	exact	$Z_{mf} \leq Z$
	Belief propagation	$q \in \mathbb{L} \supseteq \mathbb{P}$	$H_\beta \approx H(q)$	$Z_\beta \approx Z$
	Tree-reweighted	$q \in \mathbb{L} \supseteq \mathbb{P}$	$H_{tr} \geq H(q)$	$Z_{tr} \geq Z$

Variational perspectives

- Replace $q \in \mathbb{P}$ and $H(q)$ with simpler approximations

$$\log p(x^*) = \max_{q \in \mathbb{P}} \mathbb{E}_q[\log f(x)]$$

$$\log Z = \max_{q \in \mathbb{P}} \mathbb{E}_q[\log f(x)] + H(x; q)$$

- Algorithms and their properties:

	Method	distributions	entropy	value
Max:	Linear programming	$q \in \mathbb{L} \supseteq \mathbb{P}$	n/a	$\hat{p}_{lp} \geq p(x^*)$
Sum:	Mean field	$\{q = \prod q_i(x_i)\} \subseteq \mathbb{P}$	exact	$Z_{mf} \leq Z$
	Belief propagation	$q \in \mathbb{L} \supseteq \mathbb{P}$	$H_\beta \approx H(q)$	$Z_\beta \approx Z$
	Tree-reweighted	$q \in \mathbb{L} \supseteq \mathbb{P}$	$H_{tr} \geq H(q)$	$Z_{tr} \geq Z$

The local polytope

- Unfortunately, M has a large number of constraints
 - Enforce only a few, easy to check constraints?
 - Equivalent to a linear programming relaxation of original ILP

$\mu \in \mathbb{L}$: “local consistency” polytope

$$\vec{\mu} = \mathbb{E}_q[\vec{x}]$$

$q(X_1=0)$
$q(X_1=1)$
$q(X_2=0)$
$q(X_2=1)$
...
$q(X_1=0, X_2=0)$
$q(X_1=0, X_2=1)$
$q(X_1=1, X_2=0)$
$q(X_1=1, X_2=1)$
...

$$\mu_{i;k} \in [0, 1]$$

$$\mu_{ij;kl} \in [0, 1]$$

All probabilities
are within $[0,1]$

The local polytope

- Unfortunately, M has a large number of constraints
 - Enforce only a few, easy to check constraints?
 - Equivalent to a linear programming relaxation of original ILP

$\mu \in \mathbb{L}$: “local consistency” polytope

$$\vec{\mu} = \mathbb{E}_q[\vec{x}]$$

$q(X_1=0)$
$q(X_1=1)$
$q(X_2=0)$
$q(X_2=1)$
...
$q(X_1=0, X_2=0)$
$q(X_1=0, X_2=1)$
$q(X_1=1, X_2=0)$
$q(X_1=1, X_2=1)$
...

$$\mu_{i;k} \in [0, 1]$$

$$\mu_{ij;kl} \in [0, 1]$$

All probabilities
are within $[0,1]$

$$\sum_k \mu_{i;k} = 1$$

$$\sum_{k,l} \mu_{ij;kl} = 1$$

Each marginal probability
is normalized to sum to one

The local polytope

- Unfortunately, M has a large number of constraints
 - Enforce only a few, easy to check constraints?
 - Equivalent to a linear programming relaxation of original ILP

$\mu \in \mathbb{L}$: “local consistency” polytope

$$\vec{\mu} = \mathbb{E}_q[\vec{x}]$$

$q(X_1=0)$
$q(X_1=1)$
$q(X_2=0)$
$q(X_2=1)$
...
$q(X_1=0, X_2=0)$
$q(X_1=0, X_2=1)$
$q(X_1=1, X_2=0)$
$q(X_1=1, X_2=1)$
...

$$\mu_{i;k} \in [0, 1]$$

$$\mu_{ij;kl} \in [0, 1]$$

All probabilities
are within $[0,1]$

$$\sum_k \mu_{i;k} = 1$$

$$\sum_{k,l} \mu_{ij;kl} = 1$$

Each marginal probability
is normalized to sum to one

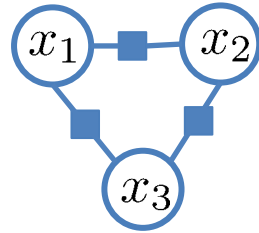
$$\sum_l \mu_{ij;kl} = \mu_{i;k}$$

$$\sum_k \mu_{ij;kl} = \mu_{j;l}$$

Marginal of (x_i, x_j)
is consistent with marginal of x_i

(& similarly, consistent with x_j)

The local polytope



- Local polytope does not enforce all the constraints of $\mathbb{M}$:
 - Ex: all pairwise probabilities locally consistent, but no joint $q(x)$ exists:

$$\mu_1 = \mu_2 = \mu_3$$

$$\begin{pmatrix} 0.5 \\ 0.5 \end{pmatrix}$$

$$\mu_{12} \quad x_2$$

$$x_1 \begin{pmatrix} 0.5 & 0 \\ 0 & 0.5 \end{pmatrix}$$

$$(x_1 = x_2)$$

$$\mu_{13} \quad x_3$$

$$x_1 \begin{pmatrix} 0.5 & 0 \\ 0 & 0.5 \end{pmatrix}$$

$$(x_1 = x_3)$$

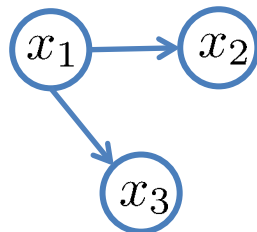
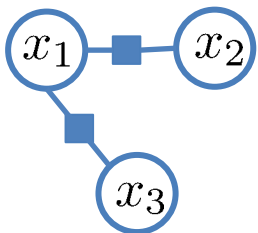
$$\mu_{23} \quad x_3$$

$$x_2 \begin{pmatrix} 0 & 0.5 \\ 0.5 & 0 \end{pmatrix}$$

$$(x_2 \neq x_3)$$

(also illustrates connection to arc consistency in CSPs, etc.)

- But, trees remain easy
 - If we only specify the marginals on a tree, we can construct $q(x)$



$$q(x) = q(x_1) \cdot q(x_2|x_1) \cdot q(x_3|x_1)$$

$$= \mu_1 \cdot \frac{\mu_{12}}{\mu_1} \cdot \frac{\mu_{13}}{\mu_1}$$

$\mathbb{L} = \mathbb{M}$ on tree-structured distributions

Duality relationship

- Local polytope LP & MAP decomposition are Lagrangian duals:

$$\log f(x^*) \leq \max_{\mu} \left[\sum_{i,k} \theta_{i;k} \mu_{i;k} + \sum_{i,j,k,l} \theta_{ij;kl} \mu_{ij;kl} \right]$$

subject to (a) normalization constraints (enforce explicitly)

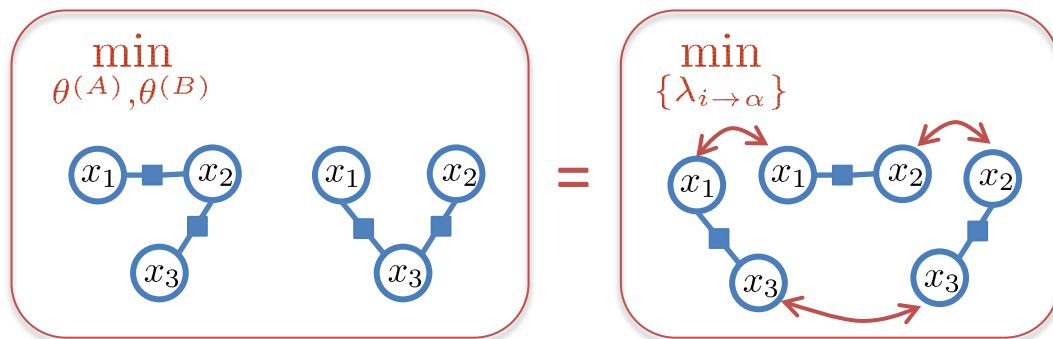
(b) consistency: $\sum_l \mu_{ij;kl} = \mu_{i;k}$, $\sum_k \mu_{ij;kl} = \mu_{j;l}$ (use Lagrange)

$$\begin{aligned} L &= \max_{\mu} \min_{\lambda} \sum_{i,k} \theta_{i;k} \mu_{i;k} + \sum_{i,j,k,l} \theta_{ij;kl} \mu_{ij;kl} + \sum_{i,j,k} \lambda_{i \rightarrow ij;k} \left(\sum_l \mu_{ij;kl} - \mu_{i;k} \right) \\ &\leq \min_{\lambda} \max_{\mu} \sum_{i,k} \theta_{i;k} \mu_{i;k} + \sum_{i,j,k,l} \theta_{ij;kl} \mu_{ij;kl} + \sum_{i,j,k} \lambda_{i \rightarrow ij;k} \left(\sum_l \mu_{ij;kl} - \mu_{i;k} \right) \\ &= \min_{\lambda} \max_{\mu} \sum_{i,k} (\theta_{i;k} - \sum_j \lambda_{i \rightarrow ij;k}) \mu_{i;k} + \sum_{i,j,k,l} (\theta_{ij;kl} + \lambda_{i \rightarrow ij;k} + \lambda_{j \rightarrow ij;l}) \mu_{ij;kl} \\ &= \min_{\lambda} \sum_{i,k} \max_k (\theta_{i;k} - \sum_j \lambda_{i \rightarrow ij;k}) + \sum_{i,j,k,l} \max_{k,l} (\theta_{ij;kl} + \lambda_{i \rightarrow ij;k} + \lambda_{j \rightarrow ij;l}) \end{aligned}$$

Duality: MAP

Primal

$$\min_{\{\lambda_{i \rightarrow \alpha}\}} \sum_{\alpha} \max_{\mathbf{x}_{\alpha}} \left[\theta_{\alpha}(\mathbf{x}_{\alpha}) + \sum_{i \in \alpha} \lambda_{i \rightarrow \alpha}(x_i) \right]$$

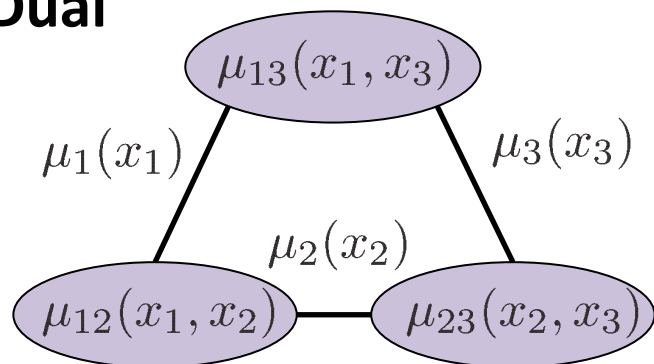


Reason about subproblems

“Messages” adjust overlapping subproblems

Reparameterize subproblems to decrease upper bound

Dual



$$\max_{\vec{\mu} \in \mathbb{L}} \vec{\theta} \cdot \vec{\mu}$$

Reason about “beliefs” (marginals)

Constraints enforce overlapping beliefs are consistent

Optimum over beliefs gives upper bound

Outline

Review: Graphical Models

Decomposition Bounds

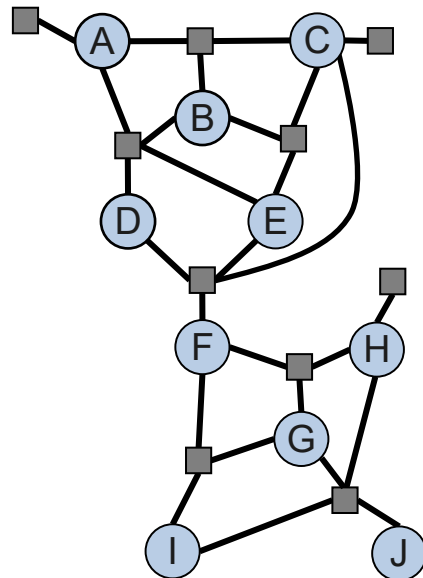
Variational Optimization

Convexity & Duality

Regions & Higher-order Approximations

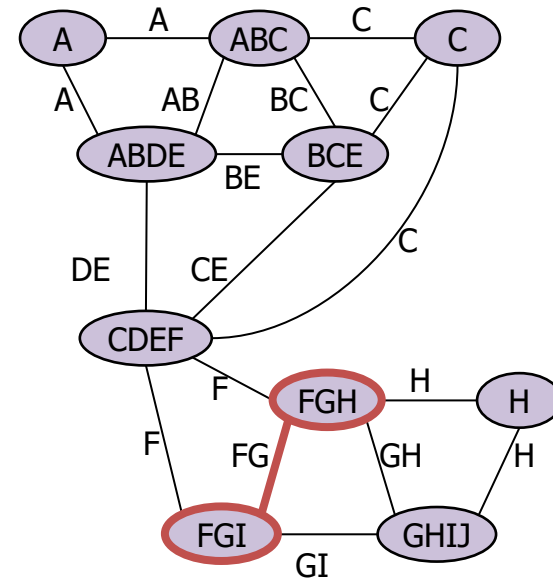
Regions

- Generalize local consistency enforcement



Factor graph

Separators = coordinates
of bound optimization (,)



Dual graph

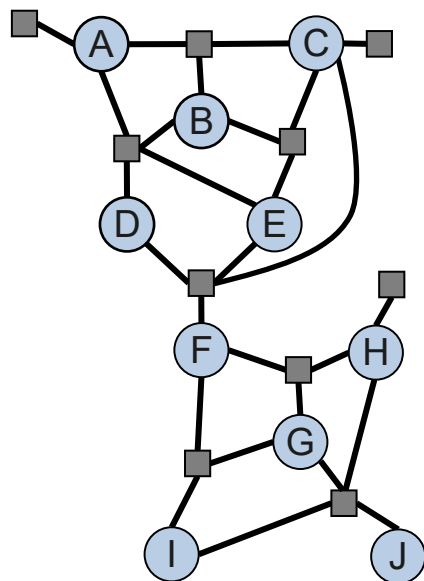
Beliefs: $\mu_{FGH}, \mu_{FGI}, \dots$

Consistency:

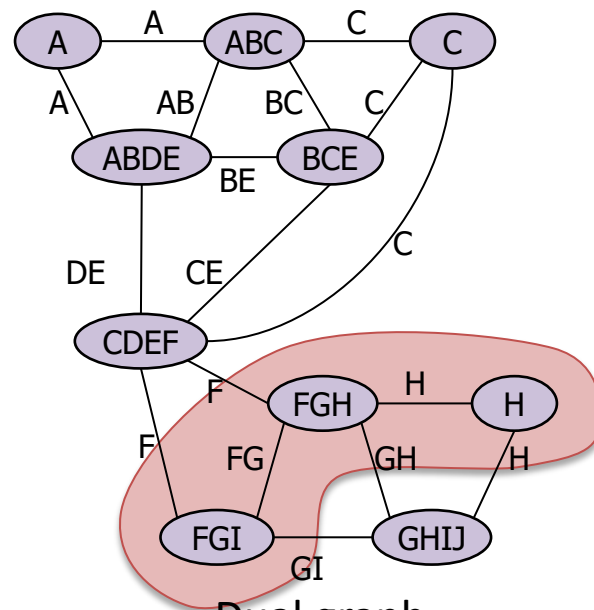
$$\sum_a \mu_{FGH}(f, g, h) = \mu_{FG}(f, g) = \sum_i \mu_{FGI}(f, g, i)$$

Regions

- Generalize local consistency enforcement
- Larger regions: more consistent; more costly to represent



Factor graph



Dual graph

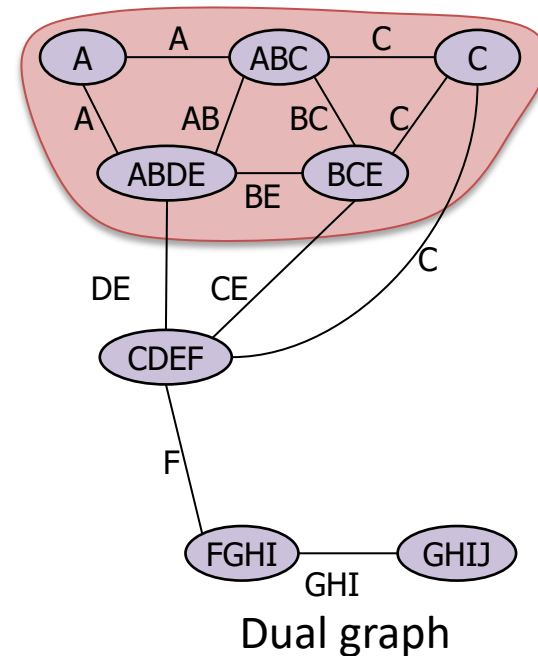
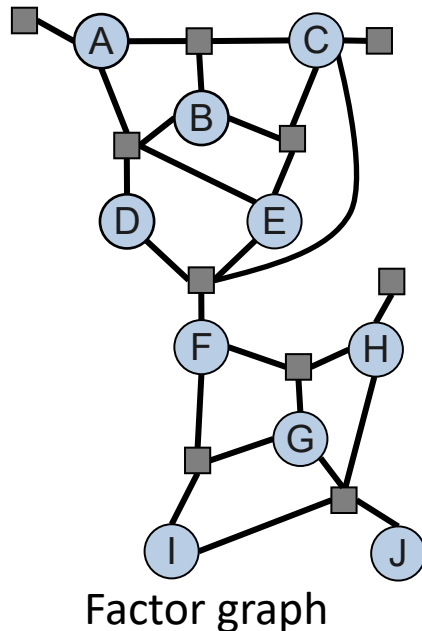
Beliefs: $\mu_{FGH}, \mu_{FGI}, \dots$

Consistency:

$$\sum_a \mu_{FGH}(f, g, h) = \mu_{FG}(f, g) = \sum_i \mu_{FGI}(f, g, i)$$

Regions

- Generalize local consistency enforcement
- Larger regions: more consistent; more costly to represent



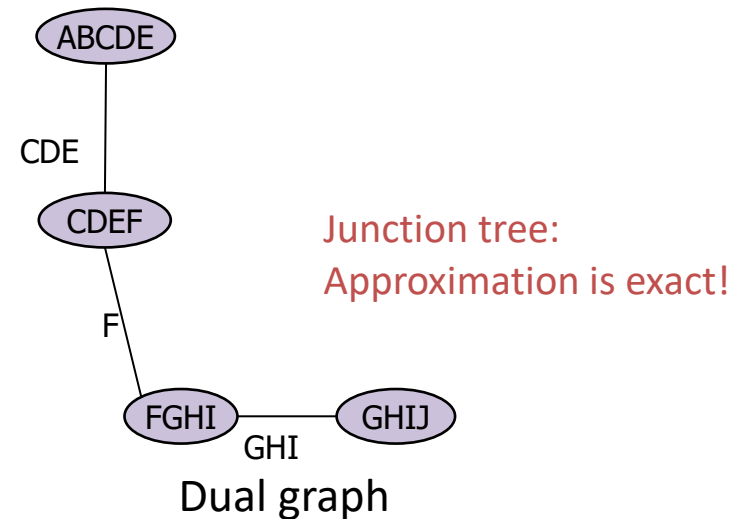
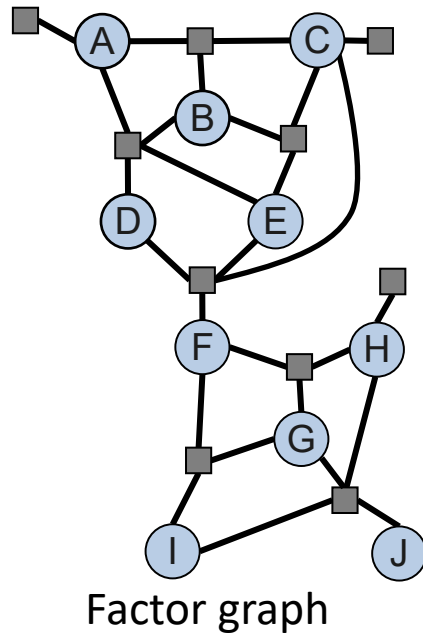
Beliefs: $\mu_{FGHI}, \dots$

Consistency:

$$\sum_a \mu_{FGHI}(f, g, h, i) = \mu_{GHI}(g, h, i) = \dots$$

Regions

- Generalize local consistency enforcement
- Larger regions: more consistent; more costly to represent

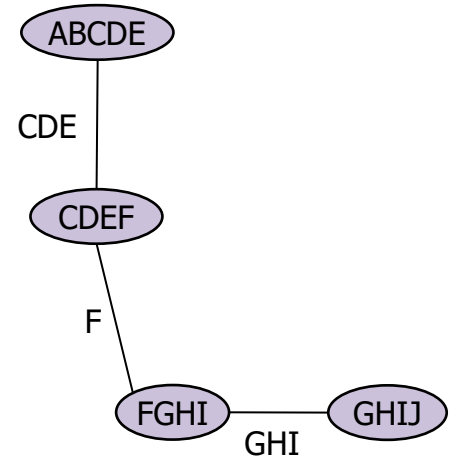
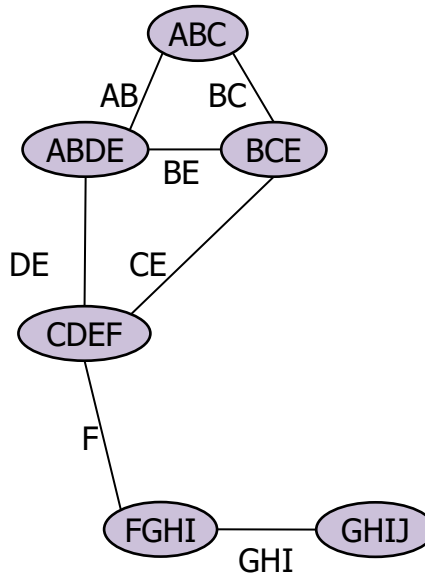
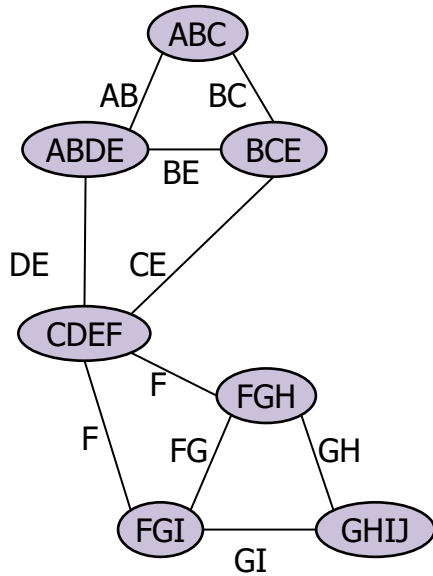


Beliefs: $\mu_{FGHI}, \dots$

Consistency:

$$\sum_a \mu_{FGHI}(f, g, h, i) = \mu_{GHI}(g, h, i) = \dots$$

Regions

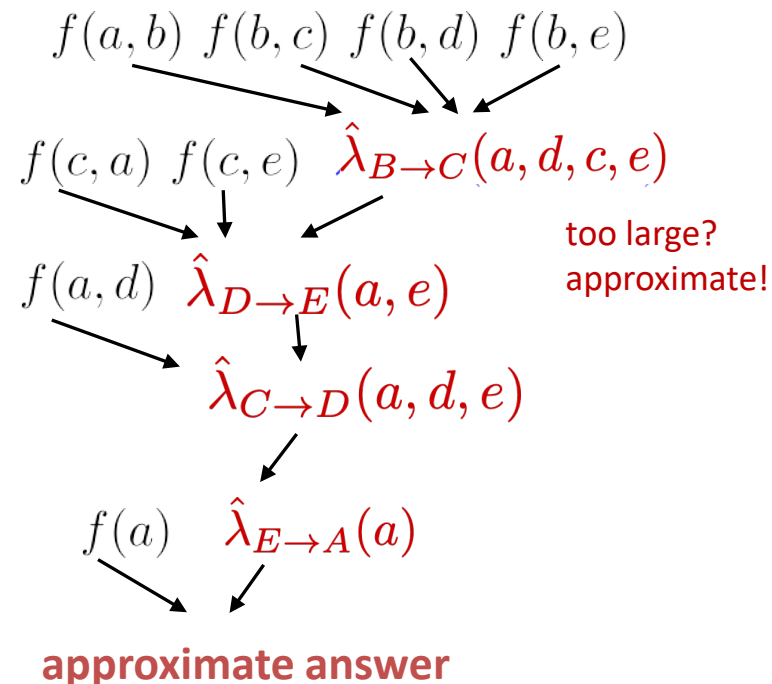
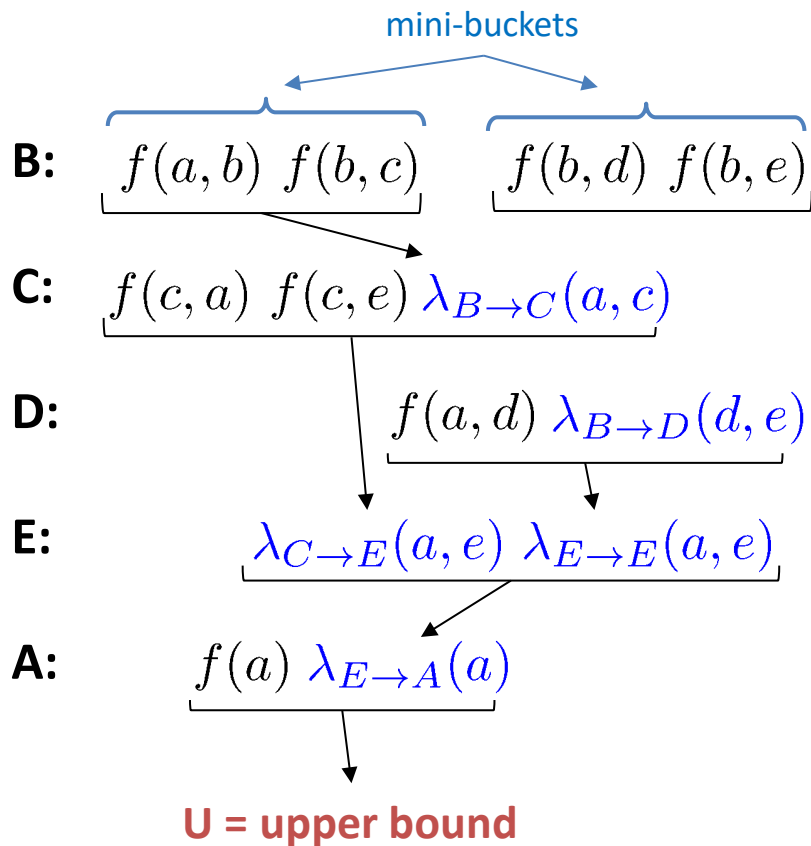


more accuracy

less complexity

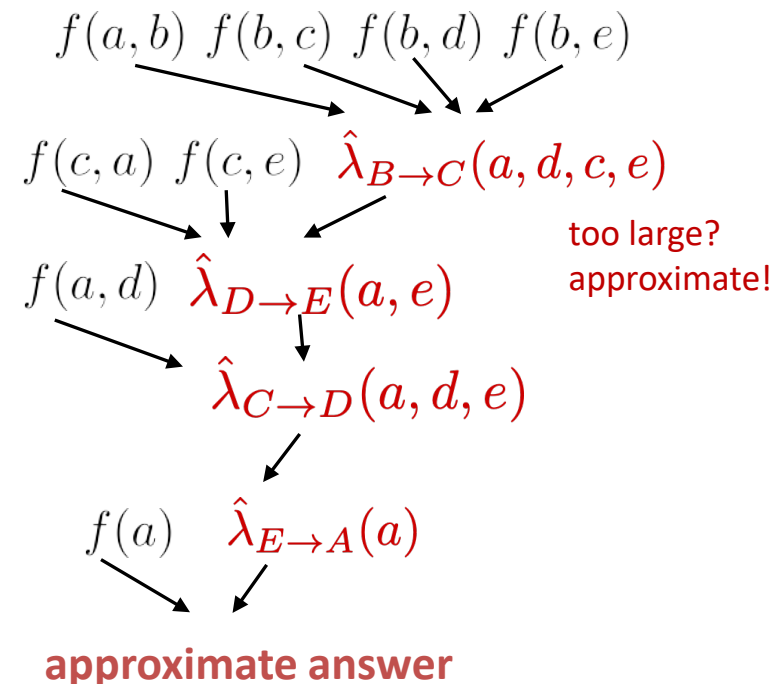
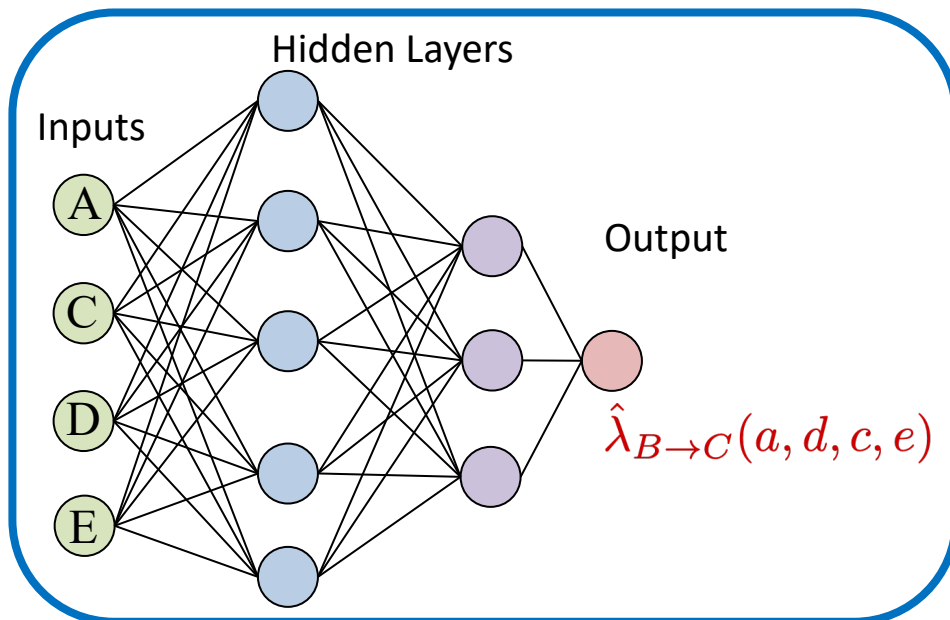
Approximate Elimination

- Mini-bucket: “structural” approximation with bounded complexity
- Can also approximate functions parametrically



Approximate Elimination

- Mini-bucket: “structural” approximation with bounded complexity
- Can also approximate functions parametrically
 - Ex: “Neuro Bucket Elimination”
 - Bounded complexity neural network for each message
 - Train to fit a subsample of table



Summary: Variational methods

- Build approximations via an optimization perspective
 - **Primal** form: decomposition into simpler problems
 - **Dual** form: optimization over local “beliefs”
- Deterministic bounds and approximations
 - Convex upper bounds
 - Non-convex lower bounds
 - Bethe approximation & belief propagation
- Scalable, “local approximation” viewpoint
 - Optimization as local message passing
- Can improve quality through increasing region size
 - But, requires exponentially increasing memory & time, or approximation