



Temporal Models

$X_i^{(t)}$ represents the instantiation of variable X_i at time t . X_i is a template variable

each "possible world" is a trajectory.

our goal is to represent a joint distribution over such trajectories

A dynamic system over the template variable X satisfies the Markov assumption if $\forall t \geq 0$

$$(X^{(t+1)} \perp X^{(0:t-1)} | X^{(t)})$$

$$P(X^{(0:T)}) = P(X^{(0)}) \cdot \prod_{t=1}^T P(X^{(t)} | X^{(t-1)})$$

A Markovian dynamic system is stationary (invariant/homogeneous) if

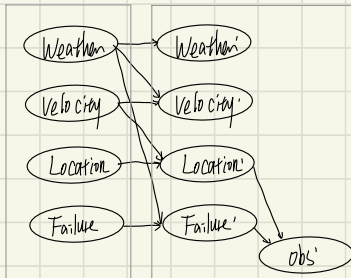
$P(X^{(t+1)} | X^{(t)})$ is same for all t . In this case, we can represent the process

using a transition model $P(x' | x)$

$$P(X^{(t+1)} = x' | X^{(t)} = x) = P(X = x' | X = x) \quad \forall t \geq 0$$

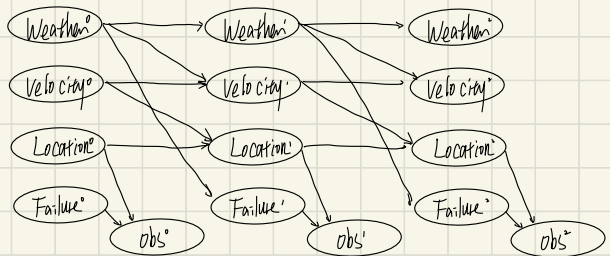
Dynamic Bayesian Network

A 2-time-slice Bayesian network (2-tbn) for a process over X is a conditional Bayesian network over X' given X_0 , where $X_0 \subseteq X$ is a set of interface variables



time slice t

time slice $t+1$



unrolled DBN

In a Conditional Bayesian Network, only X' have parents or CPDs

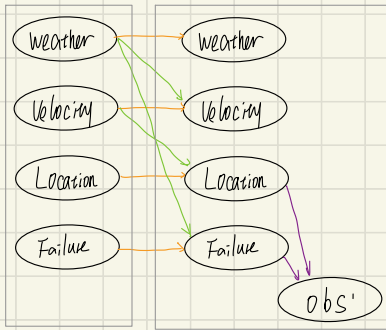
the interface variables X_0 are those variables whose values at time t have a direct effect on variables at $t+1$

$\therefore$ only variables in X_0 can be parents of variables in X'

A 2-time-slice Bayesian Network represents the conditional distribution

$$P(X | X_0) = P(X' | X_0) = \prod_{i=1}^T P(X_i' | Pa(X_i'))$$

For each template variable X_i , the CPD $P(X_i | Pa(X_i))$ is a template factor
it will be instantiated multiple times within the model, for multiple $X_i^{(t)}$ and their parents



$\Rightarrow$: inter-time-slice edges

$\rightarrow$: edges of form $X \rightarrow X'$ are persistence edge
 X is persistence variable

$\rightarrow$: intra-time-slice

if the effect of one variable on the other is immediate
(much shorter than time granularity in the model)
there's a intra-time-slice

A dynamic Bayesian Network is a pair $\langle B_0, B_{\rightarrow} \rangle$, where B_0 is a Bayesian network over $X^{(0)}$ representing the initial distribution over states. And $B_{\rightarrow}$ is a 2-TBN for the process
For any desired time span $T \geq 0$, the distribution over $X^{(0:T)}$ is defined as a unrolled Bayesian network
(B_0 is the initial state, $B_{\rightarrow}$ is the transition)

State - Observation Models

Partition variables into 2 groups: X which is always hidden
 O which is always observed

A state-observation model utilizes two independence assumptions:

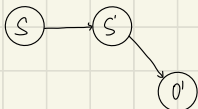
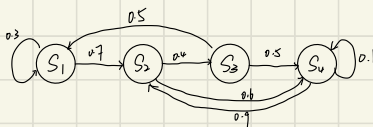
Markov assumption: $(X^{(t+1)} \perp X^{(0:t-1)} | X^{(t)})$

observation only depend on current state: $(O^{(t)} \perp X^{(0:t-1)}, X^{(t+1:t+\infty)} | X^{(t)})$

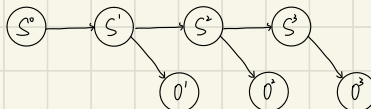
The probabilistic model consists 2 components:

transition model $p(X'|X)$ and observation model $p(O|X)$

State - Observation Models : Hidden Markov Models



2-TBN for a HMM



The unrolled DBN

A Hidden Markov Model (HMM) is a special of a simple DBN which has a sparse transition model

State - Observation Models : Linear Dynamical System

A linear dynamical model represents a system of one or more real-valued variables that evolve linearly over time, with some Gaussian noise

A linear dynamical system (often Kalman filter) can be viewed as a dynamical Bayesian network where all variables are all continuous and all of the dependencies are linear Gaussian

$$\begin{aligned} P(X^{(t)} | X^{(t-1)}) &= \mathcal{N}(A X^{(t-1)}; Q) & A \in \mathbb{R}^{n \times n} & Q \in \mathbb{R}^{n \times n} \\ P(O^{(t)} | X^{(t)}) &= \mathcal{N}(H X^{(t)}; R) & H \in \mathbb{R}^{m \times n} & R \in \mathbb{R}^{m \times m} \end{aligned}$$

A nonlinear variant of linear dynamical system (often called an extended Kalman filter) have nonlinear transition models and observation models

$$P(X^{(t)} | X^{(t-1)}) = f(X^{(t-1)}, u^{(t-1)})$$

$$P(O^{(t)} | X^{(t)}) = g(X^{(t)}, w^{(t)})$$

where f and g are deterministic nonlinear functions, and $u^{(t)}, w^{(t)}$ are Gaussian r.v.s.

Plate Models

In the plate formalism, object types are called types are called "plates"

The fact that multiple objects in the class share the same set of attributes and same probabilistic model is the basis for the use of the term "plate"

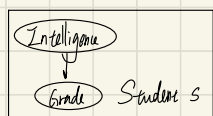
more explicit that all variables $X^{(d)}$ are sampled from the CPD

have a set of random variables $X^{(d)}$ ($d \in D$) that all have the same domain $\text{val}(X)$ and are sampled from the same distribution

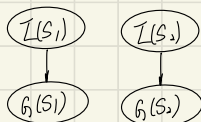


eg X is the template coin toss, $X^{(i)}$ instantiate X as the i th toss
 $X^{(1)}, X^{(2)} \dots \sim \text{iid Bernoulli}(B_x)$

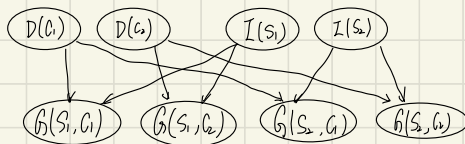
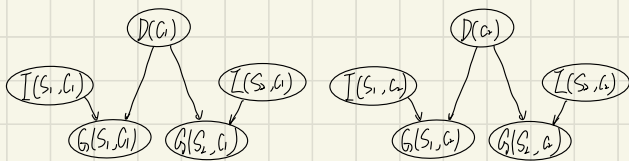
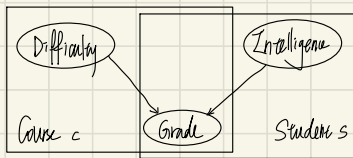
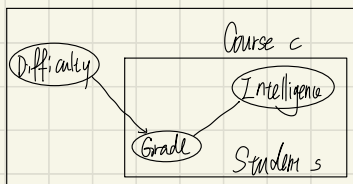
all $X^{(d)}$ are generated from the same template



instantiation



implicitly assume that $P(I(s))$ and $P(G(s) | I(s))$ is same for all s



∴ in blood type example, Genotype(child) depends on Genotype(father), plate models fail to capture this dependencies

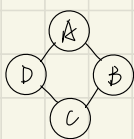
// Probabilistic Relational Model

• Undirected Templated Model

A relational Markov network M_{λ} is defined in terms of a set of template features λ where each template feature $\lambda \in \Lambda$ comprises

- A real-valued template feature f_{λ} whose arguments are $\mathcal{N}(\lambda) = \{A_1(u_1), \dots, A_n(u_n)\}$
- A weight $w_{\lambda} \in \mathbb{R}$

$$\alpha(\lambda) = \{u_i\}$$



define $\begin{cases} \text{a binary predicate (relation) } \text{Study-Pair}(S_1, S_2) \\ \text{a predicate (attribute) } \text{Misconception}(S) \end{cases}$

define a template feature f_{λ} over pairs $(\text{Misconception}(S_1), \text{Misconception}(S_2))$

$$f_{\lambda}(\text{Misconception}(S_1), \text{Misconception}(S_2), \text{Study-Pair}(S_1, S_2)) = \begin{cases} 1 & \text{if } (\text{Study-Pair}(S_1, S_2) = \text{true}) \&\& (\text{Mis}(S_1) = \text{Mis}(S_2)) \\ 0 & \text{otherwise} \end{cases}$$

$$\mathcal{N}(\lambda_{\lambda}) = \{\text{Study-Pair}(S_1, S_2), \text{Miscon}(S_1), \text{Miscon}(S_2)\}$$

$$\alpha(\lambda_{\lambda}) = \{S_1, S_2\}$$

Given a Markov network M and an object skeleton K , define a ground Gibbs distribution P_K^M

the variables in the network

P_K^M contains a term $\exp(w_{\lambda} \cdot f_{\lambda}(K))$ for each template $\lambda \in \Lambda$ and each assignment