

CS221 Problem Session

Final Review

1) Problem 1: Inferencia

You are the president of the small nation of Inferencia, and you have been charged with choosing which of your country's two rival soccer teams - the Bayesians or the Markovians - should represent Inferencia at the upcoming Olympics. You'd like to send whichever team is more popular, so you decide to model the monthly evolution of the two teams' fanbases during the months leading up to the Olympics using a dynamic Bayesian network.

Let B_t denote the number of fans that the Bayesians have in month t , and let M_t denote the number of fans that the Markovians have in month t . You have no way of observing these quantities directly, but you can observe two other quantities which they influence: let J_t denote the number of jerseys sold by the Bayesians in month t , and let A_t denote the attendance of the monthly exhibition game between the Bayesians and the Markovians in month t .

The fanbases of the two teams evolve according to the following model, where each month a fan is either gained or lost with equal probability:

$$\Pr(M_{t+1}|M_t) = \begin{cases} \frac{1}{2} & \text{if } M_{t+1} = M_t - 1 \\ \frac{1}{2} & \text{if } M_{t+1} = M_t + 1 \\ 0 & \text{otherwise} \end{cases} \quad \Pr(B_{t+1}|B_t) = \begin{cases} \frac{1}{2} & \text{if } B_{t+1} = B_t - 1 \\ \frac{1}{2} & \text{if } B_{t+1} = B_t + 1 \\ 0 & \text{otherwise} \end{cases}$$

The Bayesian fans are big spenders - almost every fan buys a jersey each month! We model the fanbase size's influence on jersey sales by:

$$\Pr(J_t|B_t) = \begin{cases} 0.3 & \text{if } J_t = B_t \\ 0.25 & \text{if } J_t = B_t - 1 \\ 0.2 & \text{if } J_t = B_t - 2 \\ 0.15 & \text{if } J_t = B_t - 3 \\ 0.1 & \text{if } J_t = B_t - 4 \\ 0 & \text{otherwise} \end{cases}$$

Lastly, because most fans attend each monthly exhibition (although sometimes more, and sometimes fewer), we model the influence of the fanbase sizes on the exhibition

attendance by:

$$\Pr(A_t|B_t, M_t) = \begin{cases} 0.14 & \text{if } A_t = B_t + M_t \\ 0.13 & \text{if } |A_t - (B_t + M_t)| = 1 \\ 0.11 & \text{if } |A_t - (B_t + M_t)| = 2 \\ 0.09 & \text{if } |A_t - (B_t + M_t)| = 3 \\ 0.06 & \text{if } |A_t - (B_t + M_t)| = 4 \\ 0.04 & \text{if } |A_t - (B_t + M_t)| = 5 \\ 0 & \text{otherwise} \end{cases}$$

Note that the assumptions and inferences made in individual parts (i.e. **(a)**, **(b)**, etc.) of this problem do *not* carry over from one to the next; the only assumptions you may make in a given part are those which are explicitly stated in that part's description.

- (a) Model the changing fanbases as a Bayesian network. You should create 8 nodes: B_t , B_{t+1} , M_t , M_{t+1} , A_t , A_{t+1} , J_t , and J_{t+1} . Indicate which nodes correspond to latent/hidden fanbase counts and which correspond to the observable emissions.



(b) **Domain Consistencies**

As a first step, we will not concern ourselves with which fanbase counts are *probable*, but instead which counts are even *possible*. Suppose that we observe, in our first month of collecting data, that $J_1 = 75$ and $A_1 = 100$. Give the domains for M_1 and B_1 that are consistent with these observations. You need only give the consistent domains (using either set notation or inequality notation).

(c) **Inference**

Suppose the Bayesian's manager took a nationwide poll in month t that concluded they had exactly 75 fans. Suppose additionally that in month $t + 2$, the Bayesians sell 73 jerseys. What is the probability that in month $t + 2$ the Bayesians have 77 fans?

- i. What is the probability that in month $t + 1$ the Bayesians sell 72 jerseys?

$$\Pr(J_{t+1} = 72 | B_t = 75) =$$

- ii. What is the probability that in month $t + 2$ that the Bayesians have 77 fans given that that they had 75 in month t and sold 73 jerseys in month $t + 2$?

$$\Pr(B_{t+2} = 77 | B_t = 75, J_{t+2} = 73) =$$

(d) **Gibbs Sampling**

Inference is exhausting; you decide that you'd be satisfied with simply being able to draw samples from distributions rather than specifying them exactly. In particular, you want to sample joint assignments to the variables $\{B_t, M_t, A_t, J_t\}_{t=1}^T$ for some time horizon T . You decide to implement Gibbs sampling for this purpose, but something's not right! What additional information, beyond what we've given you, would allow you to perform Gibbs sampling? Briefly explain.

(e) **Exact Filtering**

You now want to begin making inferences as to the sizes of the teams' fanbases given only observations of attendances and jersey sales. Recall that exact inference of this kind in dynamic Bayesian networks can be achieved using a dynamic programming approach - for example, in the context of Hidden Markov Models, we used the forward-backward algorithm to do filtering and smoothing.

Give recursive expressions for the following filtering queries. Leave your expressions in terms of known probabilities.

- i. Let's start by making inferences based only on observed jersey sales. Denote $F_t(b_t) = \Pr(B_t = b_t | J_1 = j_1, \dots, J_t = j_t)$. Give a recursive expression for $F_t(b_t)$ assuming that you've already computed $F_{t-1}(b_{t-1})$ for all b_{t-1} .

- ii. Let's bring in the observed attendances as well! Now, denote $F_t(b_t, m_t) = \Pr(B_t = b_t, M_t = m_t | J_1 = j_1, \dots, J_t = j_t, A_1 = a_1, \dots, A_t = a_t)$. Give a recursive expression for $F_t(b_t, m_t)$ assuming that you've already computed $F_{t-1}(b_{t-1}, m_{t-1})$ for all b_{t-1} and all m_{t-1} .

(f) **Particle Filtering**

Throughout this problem, you are free to leave quantities in terms of unevaluated expressions (i.e. you may write $0.75 \cdot 0.5$ instead of 0.375).

Computing all of those terms exactly seems tedious, so you instead decide to employ particle filtering to quickly and painlessly provide you with approximate solutions. You're fine with a (very) crude approximation, so you only use two particles.

- i. Suppose you begin with the two particles $(B_1 = 80, M_1 = 75)$ and $(B_1 = 82, M_1 = 74)$. You then observe that $J_1 = 79$ and $A_1 = 154$. Compute the weights that you should assign to the two particles based on this evidence.

- ii. Using these weights, we now resample two new particles. Provide this sampling distribution.

Probability of sampling a new particle to be $(B_1 = 80, M_1 = 75) =$

Probability of sampling a new particle to be $(B_1 = 82, M_1 = 74) =$

- iii. Suppose both of our new particles are sampled to be $(B_1 = 80, M_1 = 75)$. We now extend these particles using our dynamics models. What is the probability that a particular one of these two particles is extended to: $(B_1 = 80, M_1 = 75, B_2 = 78, M_2 = 76)$?

$(B_1 = 80, M_1 = 76, B_2 = 79, M_2 = 75)$?

$(B_1 = 80, M_1 = 75, B_2 = 79, M_2 = 76)$?

- iv. Suppose now that you have access to a large number of particles which are approximating the distribution over $(B_1, \dots, B_n, M_1, \dots, M_n)$. The Olympics are happening in 6 months, but you have to decide now which team to send so that they can start preparing! You decide to make predictions of B_{n+6} and M_{n+6} in order to send whichever team you predict to be more popular during the month in which the Olympics will be held. Explain in a few sentences how you would use your particles for making this decision.

2) PS9 Problem 4: Knowledge Base

Imagine we are building a knowledge base of propositions in first order logic and want to make inferences based on what we know. We will deal with a simple setting, where we only have three objects in the world: Alice, Carol, and Bob. Our predicates are as follows:

- $\text{Employee}(x)$: x is an employee.
- $\text{Boss}(x)$: x is a boss.
- $\text{Works}(x)$: x works.
- $\text{Paid}(x)$: x gets paid.

The knowledge base we have constructed consists of the following propositions:

- (a) $\text{Boss}(\text{Carol})$
 - (b) $\text{Employee}(\text{Bob})$
 - (c) $\text{Paid}(\text{Carol}) \wedge \text{Works}(\text{Carol})$
 - (d) $\text{Paid}(\text{Alice})$
 - (e) $\forall x (\text{Employee}(x) \leftrightarrow \neg \text{Boss}(x))$
 - (f) $\forall x (\text{Employee}(x) \rightarrow \text{Works}(x))$
 - (g) $\forall x ((\text{Paid}(x) \wedge \neg \text{Works}(x)) \rightarrow \text{Boss}(x))$
- (a) We know from class that one technique we can use to perform inference with our knowledge base is to propositionalize the statements of first-order logic into statements of propositional logic. Practice this by propositionalizing statement (6) from our knowledge base.
- (b) If we translated the statement "Anyone who is not a boss either works or does not get paid" into first-order logic and added it to our knowledge base, how would the size of the set of valid models representing our knowledge base change, and why?

- (c) Using only our original knowledge base (not including the statement from part (b)), we want to answer the question "Does everyone work?" We first translate the sentence "everyone works" into first order logic as statement f . Determine the answer to our query by considering the following questions of satisfiability:

- ① Is $KB \cup \neg f$ satisfiable? Answer yes/no. If yes, fill in the following table with T for true and F for false to show that there is a satisfying model.

x	Employee(x)	Boss(x)	Works(x)	Paid(x)
Alice				
Bob				
Carol				

- ② Is $KB \cup f$ satisfiable? Answer yes/no. If yes, fill in the following table with T for true and F for false to show that there is a satisfying model.

x	Employee(x)	Boss(x)	Works(x)	Paid(x)
Alice				
Bob				
Carol				

- ③ Based on your answers to the previous two parts, does our knowledge base entail f , contradict f , or is f contingent? And what should the answer to our original question "Does everyone work?" be?

3) Problem 3: CA Assignment (Winter 21, Problem 1)

Every quarter, the Stanford computer science department assigns graduate students as course assistants (CAs). Students who wish to serve as CAs fill out an application in which they can list the classes they'd like to CA for. After the application due date, the department matches applicants to courses, taking into account student preferences as well as how many course assistants each class needs. Here's the formal CA-assignment problem setup:

- There are n students $S_1, \dots, S_n$ who apply for CAships.
- There are m courses $C_1, \dots, C_m$ that have CA openings.
- Each student S_i specifies arbitrary non-negative preferences $P_1^{(i)}, \dots, P_m^{(i)} \geq 0$ for each of the m classes. A large preference value $P_j^{(i)}$ means student S_i really wants to CA for class C_j , and a preference value of 0 for $P_j^{(i)}$ means student S_i does not want to CA for class C_j .

The CA-matching process must adhere to the following requirements:

- Each course C_i can have a maximum of M_i course assistants.
- Every student must be matched to exactly one class for which they have specified a positive preference (assume each student has at least one such preference).

Model the CA-matching process with a CSP with n variables, one for each student $S_1, \dots, S_n$. Our CSP should find the *maximum weight assignment*, where the weights are determined by student preferences.

(a) What is the domain of each variable and what is the cardinality?

(b) What are the factors? State the arity of each.

- (c) We imagine a small setting of this problem for 3 students S_1, S_2, S_3 and 3 courses C_1, C_2, C_3 . The student preferences are given by the following table:

	C_1	C_2	C_3
S_1	3	0	0
S_2	2	1	3
S_3	5	3	0

Additionally, classes C_1 and C_2 can have a maximum of 1 CA each, and class C_3 can have at most 2 CAs.

Apply the CSP you designed to this small setting and enforce arc-consistency amongst its variables. In particular, write out each variable and its domain after arc-consistency has been enforced. For example, if you have a variable X_i with a domain $\{a, b, c\}$ after enforcing arc-consistency, you should write

$$X_i : \{a, b, c\}$$

- (d) True or False, with justification.
- The least constrained value (LCV) heuristic would be a useful optimization for our CA-assignment CSP.
 - The most constrained variable (MCV) heuristic would be a useful optimization for our CA-assignment CSP.

- iii. If we use the ICM algorithm to solve our CA-assignment CSP, everytime we modify a single variable assignment our factor recomputation will be on the order of n (recall that n is the number of students applying for a CA assignment).

- iv. If we use beam search with different beam sizes k to solve our CA-assignment CSP, our solution's assignment weight will always increase as we increase the beam size k .

- (e) Explain how you would modify your CSP from part a. to allow for the possibility that some students aren't matched to a course. You should encode the (realistic) assumption that not receiving a CAship is the least-preferable assignment for the student. (Note that students can still give a preference of 0 for a class if they do not want to be a CA for that class, and your modification should not prohibit this.)