

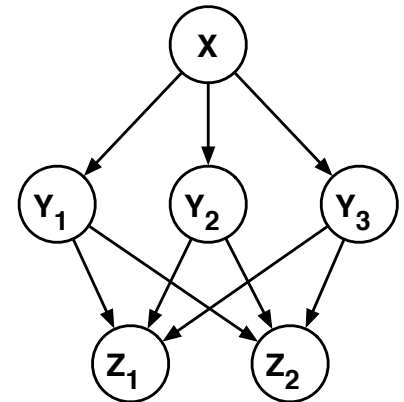
Out: Thu Jul 7**Due:** Tue Jul 12**Recommended:** The Unreasonable Effectiveness of Data<https://www.youtube.com/watch?v=yvDCzhbjYWs>

4.1 Node clustering

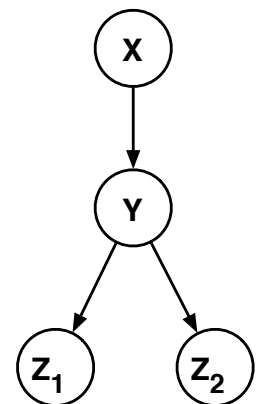
Consider the belief network shown below over binary variables X , Y_1 , Y_2 , Y_3 , Z_1 , and Z_2 . The network can be transformed into a polytree by clustering the nodes Y_1 , Y_2 , and Y_3 into a single node Y . From the CPTs in the original belief network, fill in the missing elements of the CPTs for the polytree.

X	$P(Y_1 = 1 X)$	$P(Y_2 = 1 X)$	$P(Y_3 = 1 X)$
0	0.8	0.7	0.2
1	0.4	0.5	0.9

Y_1	Y_2	Y_3	$P(Z_1 = 1 Y_1, Y_2, Y_3)$	$P(Z_2 = 1 Y_1, Y_2, Y_3)$
0	0	0	0.1	0.8
1	0	0	0.2	0.7
0	1	0	0.3	0.6
0	0	1	0.4	0.5
1	1	0	0.5	0.4
1	0	1	0.6	0.3
0	1	1	0.7	0.2
1	1	1	0.8	0.1

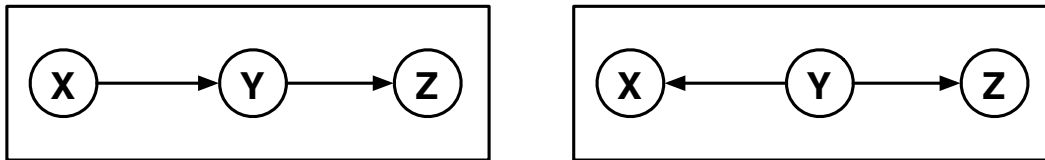


Y_1	Y_2	Y_3	Y	$P(Y X=0)$	$P(Y X=1)$	$P(Z_1=1 Y)$	$P(Z_2=1 Y)$
0	0	0	1				
1	0	0	2				
0	1	0	3				
0	0	1	4				
1	1	0	5				
1	0	1	6				
0	1	1	7				
1	1	1	8				



4.2 Maximum likelihood estimation

Consider the two DAGs shown below, which are defined over the same nodes X , Y , and Z but differ in the directionality of their edges.



For these DAGs, consider the maximum likelihood CPTs obtained from “fully observed” data $\{(x_t, y_t, z_t)\}_{t=1}^T$ in which each example provides a complete instantiation of the nodes X , Y , Z . Also, let $C(x)$ count the number of examples in which $X = x$, let $C(y)$ count the number of examples in which $Y = y$, let $C(x, y)$ count the number of examples in which $X = x$ and $Y = y$, and let $C(y, z)$ count the number of examples in which $Y = y$ and $Z = z$.

- (a) Express the maximum likelihood estimates for $P(X)$, $P(Y|X)$, and $P(Z|Y)$ in terms of the counts of x , y , and z . Note that these are the CPTs of the left DAG.
- (b) Express the maximum likelihood estimates for $P(Y)$, $P(X|Y)$, and $P(Z|Y)$ in terms of the counts of x , y , and z . Note that these are the CPTs of the right DAG.
- (c) From your answers in parts (a) and (b), show that the maximum likelihood CPTs in these different DAGs give rise to the same joint distribution over X , Y , and Z .
- (d) Are there any conditional independence relations implied by one DAG that are not implied by the other? Briefly justify your answer.

4.3 Statistical language modeling

In this problem, you will explore some simple statistical models of English text. Download the data files on the course website for this assignment. These files contain unigram and bigram counts for 500 frequently occurring tokens in English text. These tokens include actual words as well as punctuation symbols and other textual markers. In addition, an “unknown” token is used to represent all words that occur outside this basic vocabulary.

- (a) Compute the maximum likelihood estimate of the unigram distribution $P_u(w)$ over words w . Print out a table of all the words w that start with the letter “A”, along with their unigram probabilities $P_u(w)$. (You do not need to print out the full unigram distribution over all 500 words.)

- (b) Compute the maximum likelihood estimate of the bigram distribution $P_b(w'|w)$. Print out a table of the ten most likely words w' to follow the word “THE”, along with their bigram probabilities $P_b(w'|w = \text{THE})$. (You do not need need to print out the full bigram matrix.)

- (c) Consider the sentence “**The stock market fell by one hundred points last week.**” Ignoring punctuation, compute and compare the log-likelihoods (using the natural logarithm) of this sentence under the unigram and bigram models:

$$\begin{aligned}\mathcal{L}_u &= \log \left[P_u(\text{the}) P_u(\text{stock}) P_u(\text{market}) \dots P_u(\text{points}) P_u(\text{last}) P_u(\text{week}) \right] \\ \mathcal{L}_b &= \log \left[P_b(\text{the}|\langle s \rangle) P_b(\text{stock}|\text{the}) P_b(\text{market}|\text{stock}) \dots P_b(\text{last}|\text{points}) P_b(\text{week}|\text{last}) \right]\end{aligned}$$

In the equation for the bigram log-likelihood, the token $\langle s \rangle$ is used to mark the beginning of a sentence. Which model yields the highest log-likelihood?

- (d) Consider the sentence “**The sixteen officials sold fire insurance.**” Ignoring punctuation, compute and compare the log-likelihoods (using the natural logarithm) of this sentence under the unigram and bigram models:

$$\begin{aligned}\mathcal{L}_u &= \log \left[P_u(\text{the}) P_u(\text{sixteen}) P_u(\text{officials}) \dots P_u(\text{sold}) P_u(\text{fire}) P_u(\text{insurance}) \right] \\ \mathcal{L}_b &= \log \left[P_b(\text{the}|\langle s \rangle) P_b(\text{sixteen}|\text{the}) P_b(\text{officials}|\text{sixteen}) \dots P_b(\text{fire}|\text{sold}) P_b(\text{insurance}|\text{fire}) \right]\end{aligned}$$

Which pairs of adjacent words in this sentence are not observed in the training corpus? What effect does this have on the log-likelihood from the bigram model?

- (e) Consider the so-called *mixture* model that predicts words from a weighted interpolation of the unigram and bigram models:

$$P_m(w'|w) = (1 - \lambda)P_u(w') + \lambda P_b(w'|w),$$

where $\lambda \in [0, 1]$ determines how much weight is attached to each prediction. Under this mixture model, the log-likelihood of the sentence from part (d) is given by:

$$\mathcal{L}_m = \log \left[P_m(\text{the}|\langle s \rangle) P_m(\text{sixteen}|\text{the}) P_m(\text{officials}|\text{sixteen}) \dots P_m(\text{fire}|\text{sold}) P_m(\text{insurance}|\text{fire}) \right].$$

Compute and plot the value of this log-likelihood $\mathcal{L}_m$ (using the natural logarithm) as a function of the parameter $\lambda \in [0, 1]$. From your results, deduce the optimal value of λ to two significant digits.

- (f) Along with your answers, turn in a printed copy of all your source code for this problem. As usual, you may program in the language of your choice.