

PRACTICE FINAL EXAM

CSC412 WINTER 2025
PROBABILISTIC MACHINE LEARNING

University of Toronto
Faculty of Arts & Science

Duration - 3 hours

Aids allowed: Two double-sided handwritten $8.5'' \times 11''$ or A4 aid sheets.

Exam reminders:

- Fill out your name and student number on the top of this page.
- Do not begin writing the actual exam until the announcements have ended and the Exam Facilitator has started the exam.
- Write all answers in the provided answer booklets.
- Blank scrap paper is provided at the back of the exam.
- If you possess an unauthorized aid during an exam, you may be charged with an academic offence.
- Turn off and place all cell phones, smart watches, electronic devices, and unauthorized study materials in your bag under your desk. If it is left in your pocket, it may be an academic offence.
- When you are done your exam, raise your hand for someone to come and collect your exam. Do not collect your bag and jacket before your exam is handed in.
- If you are feeling ill and unable to finish your exam, please bring it to the attention of an Exam Facilitator so it can be recorded before leaving the exam hall.
- In the event of a fire alarm, do not check your cell phone when escorted outside.

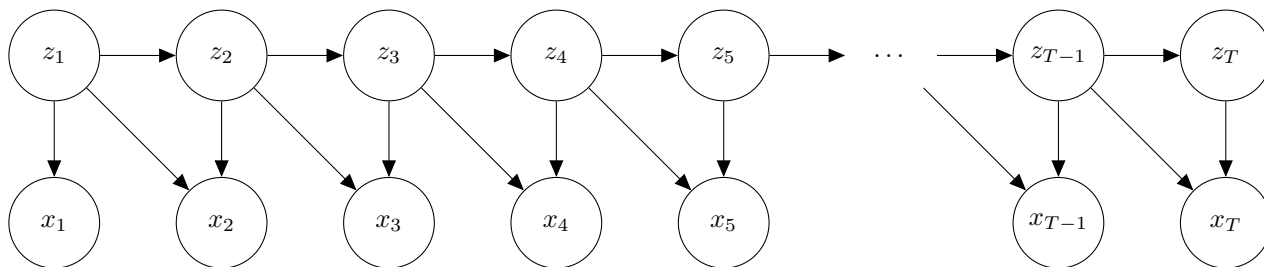
Hand in all examination materials at the end

1. Decision theory (10 points). Imagine you are writing a quiz that has a true or false section. To discourage random guessing, the quiz awards x points for a correct answer, y points for a false answer, and z points for no answer.

- (8 points) You think you know the correct answer with probability θ . How high must θ be, as a function of x , y , and z , before the expected number of points is higher for choosing the most likely answer, versus leaving the question blank?
- (2 points) How high must θ be, before the expected number of points is higher for guessing the correct answer, when $x = 2$, $y = -2$, and $z = 0$?

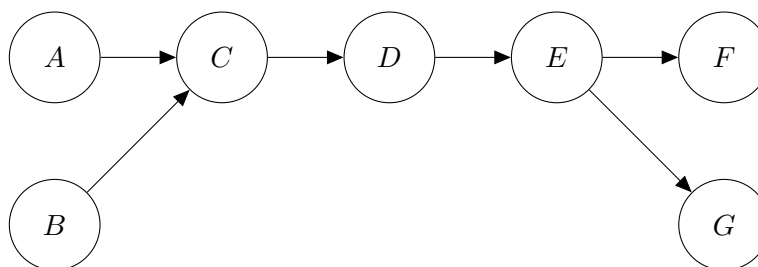
2. Graphical model analysis (20 points).

- (5 points) Consider the graphical model shown below, a 2nd-order hidden Markov model:



Write the factorization of the joint distribution over $p(z_1, z_2, \dots, z_T, x_1, x_2, \dots, x_T)$ implied by this model.

- (10 points) Consider another graphical model:



Answer true or false, no need to show your work:

- $A \perp\!\!\!\perp B$
- $B \perp\!\!\!\perp G$
- $F \perp\!\!\!\perp G$
- $A \perp\!\!\!\perp B|C$
- $A \perp\!\!\!\perp B|D$
- $A \perp\!\!\!\perp B|G$
- $F \perp\!\!\!\perp G|E$
- $F \perp\!\!\!\perp G|A$

3. (5 points) Draw the graphical model for

$$p(x_1, x_2, \dots, x_N, y_1, y_2, \dots, y_N, z_1, z_2, \dots, z_N, \theta, \pi) = p(\theta)p(\pi) \prod_{i=1}^N p(y_i|x_i, z_i, \theta)p(x_i|z_i)p(z_i|\pi)$$

3. Variational Inference (10 points). Hint for this section: Jensen's inequality states that when f is concave, $f(\mathbb{E}[z]) \geq \mathbb{E}[f(z)]$.

1. (5 points) For the joint distribution $p(x, z)$, suppose we are trying to approximate a conditional distribution $p(z|x)$ using distribution $q(z|x)$. Show that for any distribution q , the "evidence lower bound"

$$\mathcal{L}(\phi) = \mathbb{E}_{q(z|x)}[\log p(x, z) - \log q(z|x)]$$

will be less than or equal to the log marginal likelihood $\log p(x)$. You can assume p and q are positive everywhere.

2. (5 points) If a training set $x_1, x_2, \dots, x_N$ are drawn i.i.d. from $p(x|\theta)$ and the parameter $\hat{\theta}$ is estimated from the data, show that the expected log-probability of the data under $\hat{\theta}$ will be smaller in expectation on a validation set of data drawn from the same distribution $p(x|\theta)$ than it will be on the training set. That is, show that, for all $\hat{\theta}$,

$$\mathbb{E}_{p(x|\theta)}[\log p(x|\hat{\theta})] \leq \mathbb{E}_{p(x|\theta)}[\log p(x|\theta)].$$

You can assume p and q are positive everywhere.

- 4. Monte Carlo Estimators (10 points).** Recall the Simple Monte Carlo estimator:

$$\hat{e}(x_1, x_2, \dots, x_S) = \frac{1}{S} \sum_{i=1}^S f(x^{(i)}), \quad \text{where each } x^{(i)} \sim p(x) \text{ independently.}$$

- (2 points) Show that this is an unbiased estimator of $\mathbb{E}_{p(x)}[f(x)]$.
- (4 points) Find the variance of this estimator as a function of S .
- (4 points) Imagine you have a distribution $p(x)$ whose normalized density you can evaluate, but which it is difficult to sample from. You also have another distribution $q(x)$, that you can sample from, and also evaluate its density. Using these two distributions, write an unbiased estimator of $\mathbb{E}_{p(x)}[f(x)]$ that can be computed without access to samples from $p(x)$.

5. Mixture of Experts (20 points).

- (10 points) Draw a three layered MoE based on the switch transformer architecture.
- (10 points) Based on the switch transformer architecture, define the weight matrices and show how they would be multiplied together for a one layered MoE . You may choose a small layer width for simplicity

6. Diffusion (10 points).

- (5 points) Derive the loss function for a basic diffusion model based on the forward and backwards process.
- (5 points) Explain how your loss function would change if you were training a prompt-able image diffusion model

7. Word2vec (15 points). You are working with a dataset of M molecules built from some combination of any number of 35 atoms. You are interested in creating vector representations of the atoms to be used in downstream tasks. The data is represented as graphs with atoms being nodes, and edges corresponding to there being a bond between the two atoms. Describe how you could train a model to produce embeddings for atoms using this dataset, incorporating the idea that "atoms A and B are similar if they often bond to the same atoms". In your answer include the following:

1. (5 points) What is your model?
2. (4 points) What is the loss function?
3. (4 points) How is the data sampled in the training process?
4. (2 points) Is negative sampling necessary in this case?

8. Decision theory - 15 pts. Recall the density of the normal distribution $\mathcal{N}(\mu, \sigma^2)$

$$p(x|\mu, \sigma) = \frac{1}{\sqrt{2\pi\sigma^2}} \exp \left\{ -\frac{1}{2\sigma^2}(x - \mu)^2 \right\}$$

Suppose we have a classification problem with two classes $t \in \{0, 1\}$ and input x is 1-dimensional satisfying

$$x|t = 0 \sim \mathcal{N}(\mu_0, \sigma_0^2)$$

$$x|t = 1 \sim \mathcal{N}(\mu_1, \sigma_1^2)$$

We assume that, a priori, both classes are equally likely. In each of the below scenarios, mathematically derive

1. the optimal decision rule that minimizes the misclassification rate,
2. the resulting value of the misclassification rate.

Decision rule will be specified by two disjoint regions $\mathcal{R}_0$ and $\mathcal{R}_1$ with $\mathcal{R}_0 \cup \mathcal{R}_1 = \mathbb{R}$. If $x \in \mathcal{R}_0$ we classify x as class 0, otherwise class 1. The misclassification rate is given by

$$p(x \in \mathcal{R}_0, t = 1) + p(x \in \mathcal{R}_1, t = 0).$$

- (a) (5 pts) Suppose $\mu_0 \neq \mu_1$ and $\sigma_0 = \sigma_1$.
- (b) (5 pts) Suppose $\mu_0 = \mu_1$ and $\sigma_0 = \sigma_1$.
- (c) (5 pts) Suppose $\mu_0 = \mu_1$ and $\sigma_0 \neq \sigma_1$.

End of exam
