

STATS 305B: Final Exam

Write your name here:

Instructions:

- Your grade will come from your **best 5 out of 7 questions**. Each question is worth 20 points so the max score is 100.
- Write on the exam. We will scan it. If you use the extra pages at the back, label clearly.
- You can bring handwritten notes on **two sides** of an 8.5x11" piece of paper.
- Unless otherwise specified, you can write your answers using the "named distribution PDF shorthand," e.g. write the pdf of a Gaussian distribution with mean μ and variance σ^2 as $\mathcal{N}(x; \mu, \sigma^2)$.

Some tips:

1. It's usually a good idea to look through the whole exam before taking it to make sure there aren't missing pages; and so that you roughly know what you are up against.
2. Most questions have a relatively simple answer. If you find yourself doing integration by parts, rethink your approach.

Stanford Honor Code

1. The Honor Code is an undertaking of the students, individually and collectively:
 - that they will not give or receive aid in examinations; that they will not give or receive unpermitted aid in class work, in the preparation of reports, or in any other work that is to be used by the instructor as the basis of grading;
 - that they will do their share and take an active part in seeing to it that others as well as themselves uphold the spirit and letter of the Honor Code.
2. The faculty on its part manifests its confidence in the honor of its students by refraining from proctoring examinations and from taking unusual and unreasonable precautions to prevent the forms of dishonesty mentioned above. The faculty will also avoid, as far as practicable, academic procedures that create temptations to violate the Honor Code.
3. While the faculty alone has the right and obligation to set academic requirements, the students and faculty will work together to establish optimal conditions for honorable academic work.

Problem 1: *Exponential Family Distributions (20 pts)*

Consider the negative binomial distribution for a random variable $X \in \mathbb{N}$ parameterized by $r \in \mathbb{R}_+$ and $\rho \in [0, 1]$. Its pmf is,

$$\text{NB}(x; r, \rho) = \frac{\Gamma(r+x)}{x! \Gamma(r)} \rho^x (1-\rho)^r,$$

where $\Gamma(\cdot)$ is the gamma function.

- (a) (4 pts) Assume r is fixed. Write the negative binomial pmf in exponential family form. What is its natural parameter η , sufficient statistic $t(x)$, and log normalizer $A(\eta)$?

- (b) (4 pts) Using the log normalizer, compute the expected value $\mathbb{E}[X]$ where $X \sim \text{NB}(r, \rho)$.

- (c) (4 pts) Using the log normalizer, compute the variance $\text{Var}[X]$ where $X \sim \text{NB}(r, \rho)$.

- (d) (4 pts) Suppose $\lambda \sim \text{Ga}(r, \frac{1-\rho}{\rho})$ and $X | \lambda \sim \text{Po}(\lambda)$, where Ga denotes the gamma distribution with density $\text{Ga}(\lambda; a, b) = \frac{b^a}{\Gamma(a)} \lambda^{a-1} e^{-b\lambda}$, and where Po denotes the Poisson distribution with pmf $\text{Po}(x; \lambda) = \frac{1}{x!} e^{-\lambda} \lambda^x$. Compute the pmf of the marginal distribution,

$$p(x; r, \rho) = \int \text{Po}(x; \lambda) \text{Ga}(\lambda; r, \frac{1-\rho}{\rho}) d\lambda$$

- (e) (4 pts) The ratio of the variance to the mean is a common index of a distribution's *dispersion*. Consider a negative binomial distribution and a Poisson distribution with the same mean. Which has higher dispersion? Explain why your answer makes sense in light of Problem 1d.

Problem 2: Logistic Regression (20 pts)

Consider a logistic regression model,

$$p(y_i | \mathbf{x}_i) = \text{Bern}(y_i | \sigma(\boldsymbol{\beta}^\top \mathbf{x}_i))$$

where $\boldsymbol{\beta} \in \mathbb{R}^2$ and $\mathbf{x}_i \in \mathbb{R}^2$ are two-dimensional weights and covariates, respectively, and where $\sigma(a) = 1/(1 + e^{-a})$ is the logistic function. Suppose you observe the following set of independent observations,

$x_{i,1}$	$x_{i,2}$	y_i
1	0	0
0	1	1
-1	0	0
0	-1	1

Table 1: Four data points for Problem 2

(a) (1 pt) Sketch the data with an "x" where $y_i = 1$ and an "o" where $y_i = 0$.

(b) (5 pts) Write the log likelihood function, $\mathcal{L}(\boldsymbol{\beta})$, and simplify.

(c) (5 pts) Sketch the function $\log(1 + e^a) + \log(1 + e^{-a})$ for $a \in \mathbb{R}$. Label your axes and key features of the graph, like the y -intercept.

(d) (5 pts) Sketch the contours of the log likelihood $\mathcal{L}(\boldsymbol{\beta})$ for $\boldsymbol{\beta} \in \mathbb{R}^2$.

(e) (4 pts) Solve for the MLE $\hat{\boldsymbol{\beta}}_{\text{MLE}}$.

Problem 3: Mixture Models (20 pts)

Let $\mathbf{y}_i \in \mathbb{N}^D$ be a vector of counts. For example, it could correspond to the number of spikes measured in each of D neurons in the i -th time bin. Let $z_i \in \{1, \dots, K\}$ denote a discrete class assignment for that corresponding measurement. Consider the following mixture model,

$$p(\{\mathbf{y}_i, z_i\}_{i=1}^n \mid \boldsymbol{\pi}, \{\boldsymbol{\lambda}_k\}_{k=1}^K) = \prod_{i=1}^n \text{Cat}(z_i \mid \boldsymbol{\pi}) \prod_{d=1}^D \text{Po}(y_{i,d} \mid \lambda_{z_i,d})$$

where $\boldsymbol{\pi} \in \Delta_{K-1}$ is a distribution over the K classes and $\boldsymbol{\lambda}_k = (\lambda_{k,1}, \dots, \lambda_{k,D})^\top$ is a vector of rates corresponding to class k . The goal is to obtain a maximum likelihood estimate of $\{\boldsymbol{\lambda}_k\}_{k=1}^K$, assuming $\boldsymbol{\pi}$ is known.

(a) (10 pts) **E step:** Compute the *responsibilities*, $\omega_{i,k} = \Pr(z_i = k \mid \mathbf{y}_i, \boldsymbol{\pi}, \{\boldsymbol{\lambda}_j\}_{j=1}^K)$.

(b) (10 pts) **M step:** solve for the parameters $\boldsymbol{\lambda}_k$ that maximize the expected log likelihood, in terms of the responsibilities from part (a).

Problem 4: Variational Autoencoders (20 pts)

Consider the following generative model for observations $\mathbf{x} \in \mathbb{R}^D$ and latent variables $z \in \mathbb{R}$,

$$p(z, \mathbf{x}) = N(z \mid 0, 1)N(\mathbf{x} \mid \boldsymbol{\theta} \cdot z, \sigma^2 \mathbf{I}),$$

where the $\boldsymbol{\theta} \in \mathbb{R}^D$ is a parameter to be learned, and σ^2 is a fixed hyperparameter. Assume the amortized posterior (aka encoder) is of the form,

$$q(z \mid \mathbf{x}) = N(z \mid \boldsymbol{\phi}^\top \mathbf{x}, \nu^2),$$

where $\boldsymbol{\phi} \in \mathbb{R}^D$ and $\nu^2 \in \mathbb{R}_+$ are variational parameters to be learned. This is a variational autoencoder (VAE) with a **linear** encoder and decoder.

- (a) (7 pts) Write the local ELBO $\mathcal{L}(\boldsymbol{\theta}, \boldsymbol{\phi}, \nu) \leq \log p(\mathbf{x}; \boldsymbol{\theta})$ for a single observation $\mathbf{x}$. Leave your answer in terms of the Gaussian density function.

- (b) (2 pts) When you maximize the ELBO, what quantity is minimized?

(c) (9 pts) Fixing θ , solve for the variational parameters ϕ^* and ν^* that maximize the ELBO.

(d) (2 pts) With the optimal variational parameters found above, is the ELBO tight? I.e., does $\mathcal{L}(\theta, \phi^*, \nu^*)$ equal $\log p(\mathbf{x}; \theta)$?

Problem 5: Hidden Markov Models (HMMs) (20 pts)

A *factorial* HMM is a hidden Markov model with $L \geq 2$ levels of latent states. The latent states evolve independently at each level, but they all jointly determine the likelihood for the observations. Let $\mathbf{z}^{(\ell)} = (z_1^{(\ell)}, \dots, z_T^{(\ell)})$ denote the sequence of latent states at level ℓ . Assume that for all levels and time steps, the discrete latent states take values $z_t^{(\ell)} \in \{1, \dots, K\}$. Let $\mathbf{x} = (x_1, \dots, x_T)$ denote the sequence of observations. The joint distribution factors as,

$$p(\mathbf{x}, \{\mathbf{z}^{(\ell)}\}_{\ell=1}^L) = \left[\prod_{\ell=1}^L p(z_1^{(\ell)}) \prod_{t=2}^T p(z_t^{(\ell)} | z_{t-1}^{(\ell)}) \right] \times \prod_{t=1}^T p(x_t | z_t^{(1)}, \dots, z_t^{(L)}).$$

- (a) (7 pts) Write the factorial HMM above as a standard HMM on an extended state space. (Define your extended state space and the corresponding transition probabilities.)

- (b) (6 pts) What is the time complexity of the forward-backward for a factorial HMM? Your answer should be in terms of T , K , and L .

- (c) (7 pts) Consider a factorial HMM with binary latent states, $z_t^{(\ell)} \in \{0, 1\}$. Let $\mathbf{z}_t = (z_t^{(1)}, \dots, z_t^{(L)}) \in \{0, 1\}^L$ denote the vector of latent states at time t . Assume a Gaussian likelihood,

$$p(x_t | z_t^{(1)}, \dots, z_t^{(L)}; \boldsymbol{\theta}) = \mathcal{N}(x_t | \boldsymbol{\theta}^\top \mathbf{z}_t, 1).$$

where $\boldsymbol{\theta} = (\theta_1, \dots, \theta_L) \in \mathbb{R}^L$. Intuitively, the mean of x_t is a sum of contributions θ_l for each level that is "on" at time t .

With a small number of levels (say, $L < 10$), we can use the forward-backward algorithm to compute posterior expectations $\boldsymbol{\omega}_t = \mathbb{E}[\mathbf{z}_t | \mathbf{x}]$ and $\boldsymbol{\Omega}_t = \mathbb{E}[\mathbf{z}_t \mathbf{z}_t^\top | \mathbf{x}]$ for all $t = 1, \dots, T$. Find a closed-form solution for the M-step,

$$\boldsymbol{\theta}^* = \arg \max_{\boldsymbol{\theta}} \mathbb{E}_{p(\mathbf{z}|\mathbf{x})} \left[\sum_{t=1}^T \log p(x_t | z_t^{(1)}, \dots, z_t^{(L)}; \boldsymbol{\theta}) \right].$$

Your answer should be in terms of $\boldsymbol{\omega}_t$ and/or $\boldsymbol{\Omega}_t$.

Problem 6: Recurrent Neural Networks (RNNs) (20 pts)

Consider a *linear* RNN that consumes a sequence of inputs $\mathbf{x} = (x_1, \dots, x_L)$ with $x_\ell \in \mathbb{R}$ and produces a single output at the end, $y \in \mathbb{R}$. The log likelihood is,

$$p(y | \mathbf{x}) = \mathcal{N}(y | h_{L+1}, 1)$$

where the hidden states $h_\ell \in \mathbb{R}$ obey deterministic, linear dynamics for $\ell = 0, \dots, L$,

$$h_{\ell+1} = wh_\ell + bx_\ell.$$

Assume $h_0 = 0$. The model has three real-valued parameters w , b , and x_0 .

(a) (5 pts) Write the model as a linear regression,

$$p(y | \mathbf{x}) = \mathcal{N}(y | \theta_0 + \boldsymbol{\theta}^\top \mathbf{x}, 1)$$

where $\theta_0 \in \mathbb{R}$ is a scalar intercept and $\boldsymbol{\theta} \in \mathbb{R}^L$ is a vector of weights. Give an expression for θ_0 and the entries of $\boldsymbol{\theta}$.

(b) (5 pts) Plot the weights θ_ℓ as a function of ℓ for $\ell = 0, 1, \dots, L = 20$ assuming $w = e^{-1/20} \approx 0.95$, $b = 1$, and $x_0 = 1$. Label your axes and key points like the y-intercept.

(c) (5 pts) Compute the derivative of the log likelihood with respect to the initial condition x_0 .

(d) (5 pts) Why is the solution to part (c) problematic for learning x_0 with simple gradient descent?

Problem 7: Transformers (20 pts)

Instead of the standard softmax attention, consider a more general form of *kernel* attention,

$$A_{t,s} = \frac{\text{sim}(\mathbf{U}_q \mathbf{x}_t, \mathbf{U}_k \mathbf{x}_s)}{\sum_{s'=1}^T \text{sim}(\mathbf{U}_q \mathbf{x}_t, \mathbf{U}_k \mathbf{x}_{s'})},$$

where $\mathbf{U}_q \mathbf{x}_t \in \mathbb{R}^K$ are the *queries* and $\mathbf{U}_k \mathbf{x}_s \in \mathbb{R}^K$ are the *keys*. Assume $K < D$. The function $\text{sim}(\cdot, \cdot)$ is a *kernel* that measures similarity between two vectors. (Softmax attention corresponds to the exponential kernel, $\text{sim}(\mathbf{a}, \mathbf{b}) = \exp\{\mathbf{a}^\top \mathbf{b}\}$.) Here, we will consider the quadratic kernel instead,

$$\text{sim}(\mathbf{a}, \mathbf{b}) = (\mathbf{a}^\top \mathbf{b})^2.$$

- (a) (4 pts) Show that with the quadratic kernel, $\text{sim}(\mathbf{a}, \mathbf{b})$ can be written as an inner product $\phi(\mathbf{a})^\top \phi(\mathbf{b})$ where $\phi : \mathbb{R}^K \mapsto \mathbb{R}^P$ maps the vectors into a feature space of possibly higher dimension. *Hint: note that $\text{Tr}(\mathbf{A}^\top \mathbf{B}) = \text{vec}(\mathbf{A})^\top \text{vec}(\mathbf{B})$ is an inner product.*

- (b) (6 pts) In class, we wrote the output of a single-headed self-attention step as $\mathbf{Y} = \mathbf{A}\mathbf{X}$ where $\mathbf{X}, \mathbf{Y} \in \mathbb{R}^{T \times D}$ are matrices with rows $\mathbf{x}_t^\top$ and $\mathbf{y}_t^\top$, respectively, and $\mathbf{A} \in \mathbb{R}^{T \times T}$ is the matrix of attention weights. Show that with the quadratic kernel, the outputs matrix can be written as,

$$\mathbf{y}_t = \frac{\mathbf{V} \phi(\mathbf{U}_q \mathbf{x}_t)}{\mathbf{w}^\top \phi(\mathbf{U}_q \mathbf{x}_t)}$$

for some matrix $\mathbf{V} \in \mathbb{R}^{D \times P}$ and vector $\mathbf{w} \in \mathbb{R}^P$ that are the same for all outputs t . (Here, $\mathbf{x}_t$ and $\mathbf{y}_t$ are column vectors.)

(c) (4 pts) What is the computational complexity of the single headed self-attention step using the quadratic kernel? Your answer should be in terms of T , D , and/or K , but not P .

(d) (4 pts) What is the computation complexity of the single headed self-attention step using standard softmax attention? Your answer should be in terms of T , D , and/or K .

(e) (2 pts) In what regimes is quadratic attention more efficient? Do you expect Transformers to often be in this regime?

Blank page for extra work (please label clearly for GradeScope scanning).

Blank page for extra work (please label clearly for GradeScope scanning).

Blank page for extra work (please label clearly for GradeScope scanning).

Blank page for extra work (please label clearly for GradeScope scanning).