

CARNEGIE MELLON UNIVERSITY 07-280

HOMEWORK 9

DUE: Thursday, Mar. 19, 2026

<https://www.cs.cmu.edu/~07280>

INSTRUCTIONS

- **Format:** Use the provided LaTeX template to write your answers in the appropriate locations within the *.tex files and then compile a pdf for submission. We try to mark these areas with STUDENT SOLUTION HERE comments. Make sure that you don't change the size or location of any of the answer boxes and that your answers are within the dedicated regions for each question/part. If you do not follow this format, we may deduct points.

You may also digitally annotate the pdf. Illegible handwriting will lead to lost points. However, we suggest that try to do at least some of your work directly in LaTeX.

- **How to submit written component:** Submit to Gradescope a pdf with your answers. Again, make sure your answer boxes are aligned with the original pdf template.
- **How to submit programming component:** See Programming component for details on how to submit to the Gradescope autograder.
- **Policy:** See the course website for homework policies, including late policy, and academic integrity policies.
- **Proofs and Derivations:** For full credit, you must clearly show all nontrivial steps. Explicit reasons for each step are not required but can be helpful, especially if the step is not obvious.

Name	
Andrew ID	
Hours to complete all components (nearest hour)	

1 [18 pts] MLE

Consider the following distribution with parameters k and α :

$$p(x \mid k, \alpha) = \mathbb{I}(x \geq k) \frac{\alpha k^\alpha}{x^{\alpha+1}} = \begin{cases} \frac{\alpha k^\alpha}{x^{\alpha+1}} & x \geq k \\ 0 & \text{otherwise} \end{cases}$$

We also have that $k \in (0, \infty)$ and $\alpha \in (0, \infty)$.

This distribution is often used for modeling the distribution of wealth in society, fitting the trend that a large portion of wealth is held by a small fraction of the population. This is due to its nature as a skewed, heavy-tailed distribution (notice that the probability decays polynomially in value of the random variable x , unlike Gaussian, exponential, Laplace or Poisson distributions where the probability decays exponentially in the value of the random variable).

Suppose you have a dataset $\mathcal{D}$ which contains N i.i.d samples $x^{(1)}, x^{(2)}, \dots, x^{(N)}$ drawn from the above distribution.

Note: For convenience in this optimization problem, let us define $\log 0$ to be $-\infty$; this won't affect the answers in the end; it essentially just acts as a notational convenience.

1. [4 pts] Write the likelihood $\mathcal{L}(k, \alpha; \mathcal{D})$ and log-likelihood $\ell(k, \alpha; \mathcal{D})$. You may keep the indicator function in your answers (rather than trying to write cases or otherwise avoid the indicator function).

Your Answer

2. [2 pts] What is the numerical value of the likelihood (not the log-likelihood) if we have $\alpha = 2$, $k = 5$, and $N = 4$ data points $\mathcal{D} = \{2, 3, 7, 4\}$?

3. **[8 pts]** Give the MLE for the parameter α , assuming the parameter k is fixed and $k \leq x^{(i)} \forall i$. *Hint:* There is a closed-form solution.

 $\hat{\alpha}_{MLE}$

4. **[4 pts]** Next, give the MLE for the parameter k . Here α may be any fixed value or set to its MLE.

Hint: You may be tempted to conclude that MLE of k is infinity, but when $k = \infty$, what happens to $p(x | k, \alpha)$?

 $\hat{k}_{MLE}$

2 [20 pts] MLE for Categorical Variables

The categorical distribution is a discrete distribution with K possible values (often corresponding to K discrete events). We'll represent this distribution as a one-hot vector of K random binary random variables, $Y = [Y_1, Y_2, \dots, Y_K]^\top$, where Y_k takes on the value 1 if the outcome is in class k and 0 otherwise. It is described by a parameter vector $\phi = (\phi_1, \phi_2, \dots, \phi_K)$ where each ϕ_k determines the probability that the random variable is in category k . Also note that $\phi_k \geq 0$ and $\sum_{k=1}^K \phi_k = 1$. Note that the Bernoulli distribution is a special case of the categorical distribution where $K = 2$.

For a dataset $\mathcal{D}$ of N i.i.d. samples drawn from the categorical distribution, the likelihood function can be written as:

$$\mathcal{L}(\phi; \mathcal{D}) = \prod_{i=1}^N p(\mathbf{y}^{(i)} | \phi) = \prod_{i=1}^N \prod_{k=1}^K \phi_k^{\mathbb{I}(y_k^{(i)}=1)}$$

To help us to remove the indicator function notation, we can define N_k as the number of points in $\mathcal{D}$ that belong to the k -th category.

1. [4 pts] Write the likelihood and log-likelihood functions $p(\mathcal{D} | \phi)$ and $\log p(\mathcal{D} | \phi)$ in terms of N_k .

$\mathcal{L}(\phi; \mathcal{D})$

$\ell(\phi; \mathcal{D})$

If we just took the derivative of the log-likelihood and set it equal to zero, we would get an estimate saying that each ϕ_k should be equal to infinity, regardless of the value of N_k . So, something is a little broken here... What we're missing is to account for the constraint that $\sum_{k=1}^K \phi_k = 1$. Here is the constrained optimization problem:

$$\begin{aligned} \underset{\phi}{\operatorname{argmin}} \quad & -\ell(\phi; \mathcal{D}) \\ \text{s.t.} \quad & \sum_{k=1}^K \phi_k = 1 \end{aligned}$$

To solve this constrained optimization problem, we're going to use a really cool trick called Lagrange multipliers. We'll give you just enough information on Lagrange multipliers in this write-up, but if you are curious and want more information, you could start with Lagrange multiplier tutorials in [Khan Academy](#) or [Paul's Notes](#).

- **Super brief instructions for using Lagrange multipliers to solve a generic constrained optimization problem**

Goal: Solve the following optimization:

$$\begin{aligned} \underset{\phi}{\operatorname{argmin}} \quad & f(\phi) \\ \text{s.t.} \quad & g(\phi) = 0 \end{aligned}$$

Step 1: Construct the “Lagrangian”, $L(\phi, \lambda)$, with the added scalar variable λ . Note that we're going to use the letter L for the Lagrangian rather than the typical $\mathcal{L}$ to avoid confusion with the notation for the likelihood.

$$L(\phi, \lambda) = f(\phi) + \lambda g(\phi)$$

Step 2: Find the “saddle point” of the Lagrangian, specifically the ϕ and λ where:

$$\nabla L(\phi, \lambda) = \mathbf{0}$$

For our categorical MLE problem, we can formulate the Lagrangian (Step 1) with $f(\phi)$ as the negative log-likelihood and $g(x) = (\sum_{k=1}^K \phi_k) - 1$.

2. **[4 pts]** (Lagrange multipliers Step 2a) Write the partial derivatives of the Lagrangian with respect to parameter ϕ_k and with respect to λ :

$$\partial L / \partial \phi_k$$

$\partial L / \partial \lambda$

3. **[8 pts]** (Lagrange multipliers Step 2b) Set your expressions for $\partial L / \partial \phi_k$ and $\partial L / \partial \lambda$ both equal to zero and use them as a system of equations to solve for the MLE estimate for ϕ_k .

 ϕ_k^{MLE}

4. **[4 pts]** Finally, let's plug in some numbers to see if our results above match our intuition. For a dataset with $N = 100$ points $K = 4$ categories and $N_1 = 10$, $N_2 = 20$, $N_3 = 30$, and $N_4 = 40$, what is the MLE of the vector ϕ . Write your answer as a vector of numerical values.

ϕ_{MLE}

3 [12 pts] Programming Experiments

Complete this written section **after** finishing the programming notebook `hw9_student.ipynb` (especially Sections 3.5–3.9 and 4.1). Use your own generated outputs and printed statistics from those sections.

1. [6 pts] Sampling strategy analysis (Sections 3.5–3.9)

(a) **Temperature effects:** Using your Section 3.9 experiments, describe what happens when:

- Temperature is low (e.g., 0.5)
- Temperature is high (e.g., 5.0)

When might each setting be useful?

Temperature Effects

(b) **Top-k effects:** How does reducing k affect output diversity and coherence? Explain using your observations from Section 3.9.

Top-k Effects

(c) **Combined sampling:** From your Section 3.9 experiments, does combining temperature and top-k produce better results than using either method alone? What combination worked best for Hamilton-style text, and why?

Combined Sampling

2. [6 pts] Data sparsity in n-gram models (Section 4.1)

Use the coverage table/results you computed in Section 4.1.

- (a) **Interpreting coverage:** Describe the trend in n-gram coverage from unigrams to 4-grams. Why does coverage drop so quickly as n increases?

Coverage Trend

- (b) **Why sparsity matters:** Explain why unseen n-grams receive probability 0 and how this affects:

- Generation of novel sequences
- Evaluation of realistic text
- Higher-order n-gram models in particular

Impact of Sparsity

4 Collaboration Policy

After you have completed all other components of this assignment, report your answers to the following collaboration questions.

1. Did you receive any help whatsoever from anyone in solving this assignment? If so, include full details including names of people who helped you and the exact nature of help you received.

Your Answer

2. Did you give any help whatsoever to anyone in solving this assignment? If so, include full details including names of people you helped and the exact nature of help you offered.

Your Answer

3. Did you find or come across code that implements any part of this assignment? If so, include full details including the source of the code and how you used it in the assignment.

Your Answer