

Assignment 1: Probabilistic Reasoning, Maximum Likelihood, Classification

For due date see <https://courses.cs.sfu.ca>

This assignment is to be done individually.

Important Note: The university policy on academic dishonesty (cheating) will be taken very seriously in this course. You may not provide or use any solution, in whole or in part, to or by another student. Any mark in this assignment may be changed on the basis of an oral exam where you need to explain your solution.

You are encouraged to discuss the concepts involved in the questions with other students. If you are in doubt as to what constitutes acceptable discussion, please ask! Further, please take advantage of office hours offered by the instructor and the TA if you are having difficulties with this assignment.

DO NOT:

- Give/receive code or proofs to/from other students
- Use Google to find solutions for assignment

DO:

- Meet with other students to discuss assignment (it is best not to take any notes during such meetings, and to re-work assignment on your own)
 - Use online resources (e.g. Wikipedia) to understand the concepts needed to solve the assignment
-

Practice: Joint and Conditional Probabilities (10 marks)

Go to www.aispace.org and start the “belief and decision network” tool. Load the sample file “File/Load Sample Problem/Simple Diagnostic Example”. *If you have difficulty with the UBC tool (e.g. Java doesn’t run, Simple Diagnostic isn’t found), I suggest you download the jar file from the course website.*

We will use this to test some of the basic probability laws. The AISpace tool can do many of these calculations for you, but the purpose of the exercise is to learn the principles behind the calculations. You can use the tool to check your answers, but you should compute them yourself using the proba-

bility calculus together with the conditional probabilities from the Bayesian network. *Compute the following joint probabilities up to 6 significant digits.*

1. Use the product formula of Bayes nets and the conditional probability parameters specified by AIspace to compute the probability that: all nodes are true.

$$P(\text{all nodes are True}) = 0.00128$$

2. Use the product formula of Bayes nets and the conditional probability parameters specified by AIspace to compute the probability that: all nodes are true except for Sore Throat, and that Sore Throat is false.

$$P(\text{SoreThroat=False, all other nodes True}) = 0.00299$$

3. Show how can you use these two joint probabilities to compute the probability that: all nodes other than Sore Throat are true. (Where the value of Sore Throat is unspecified.)

$$\begin{aligned} P(\text{all nodes other than SoreThroat True}) \\ &= P(\text{SoreThroat=True, all other nodes True}) + P(\text{SoreThroat=False, all other nodes True}) \\ &= 0.00128 + 0.00299 \\ &= 0.00427 \end{aligned}$$

4. Verify the product formula:

$$P(\text{all nodes are true}) = P(\text{Sore Throat} = \text{true} \mid \text{all other nodes are true}) \times P(\text{all other nodes are true}).$$

You may get the first conditional probability by executing a query with the tool.

$$\begin{aligned} P(\text{all nodes are true}) \\ &= P(\text{SoreThroat True} \mid \text{all other nodes are True}) \times P(\text{all other nodes are True}) \\ &= 0.3 \times 0.00427 \\ &= 0.00128 \end{aligned}$$

5. Compute the probability that Sore Throat is true and that Fever is true. (Hint: If you use the right formula, you need only 4 conditional probabilities.)

$$\begin{aligned} P(\text{SoreThroat} = \text{True, Fever} = \text{True}) &= P(\text{SoreThroat} = \text{True, Fever} \\ &= \text{True} \mid \text{Influenza} = \text{True}) \times P(\text{Influenza} = \text{True}) \end{aligned}$$

$$\begin{aligned}
& + P(\text{SoreThroat} = \text{True}, \text{Fever} = \text{True} \mid \text{Influenza} = \text{False}) \times P(\text{Influenza} = \text{False}) \\
& = P(\text{SoreThroat} = \text{True} \mid \text{Influenza} = \text{True}) \times P(\text{Fever} = \text{True} \mid \text{Influenza} = \text{True}) \times P(\text{Influenza} = \text{True}) \\
& + P(\text{SoreThroat} = \text{True} \mid \text{Influenza} = \text{False}) \times P(\text{Fever} = \text{True} \mid \text{Influenza} = \text{False}) \times P(\text{Influenza} = \text{False}) \\
& = 0.3 \times 0.9 \times 0.05 + 0.001 \times 0.05 \times 0.95 \\
& = 0.0135
\end{aligned}$$

You can enter the computed probabilities in the table below.

Probability to be Computed	Your Result
P(all nodes true)	0.00128
P(Sore Throat = False, all other nodes true)	0.00299
P(all nodes other than Sore Throat true)	0.00427
P(all nodes are true) = P(Sore Throat = true all other nodes are true) x P(all other nodes are true).	0.00128304 = 0.3 x 0.0042768
P(Sore Throat = true, Fever = True)	0.0135475

Theory: Conditional Probabilities (9 marks)

Exercise 13.3 in Russell and Norvig AMAI.

1. True.

$$\begin{aligned}
P(a \mid b, c) &= P(b \mid a, c) \\
\implies \frac{P(a,b,c)}{P(b,c)} &= \frac{P(b,a,c)}{P(a,c)} \\
\implies P(a,c) &= P(b,c) \\
\implies \frac{P(a,c)}{P(c)} &= \frac{P(b,c)}{P(c)} \\
\therefore P(a \mid b, c) &= P(b \mid a, c) \implies P(a \mid c) = P(b \mid c)
\end{aligned}$$

2. False.

Counter-Example: Consider two independent coin flips a and b, such that $c = b$.

$$P(a = \text{Head} \mid b = \text{Tail}, c = \text{Tail}) = P(a = \text{Head}) = 0.5$$

$$\begin{aligned}
P(b = \text{Tail} \mid c = \text{Tail}) &= 1 \\
P(b = \text{Tail}) &= 0.5 \\
P(b = \text{Tail} \mid c = \text{Tail}) &\neq P(b = \text{Tail})
\end{aligned}$$

3. False.

Counter-Example: Consider two independent coin flips a and b, such that $c = a \text{ xor } b$.

$$P(a = \text{Tail} \mid b = \text{Tail}) = P(a = \text{Tail}) = 0.5$$

$$P(a = \text{Tail} \mid b = \text{Tail}, c = \text{Tail}) = 1$$

$$P(a = \text{Tail} \mid c = \text{Tail}) = 0.5$$

$$P(a = \text{Tail} \mid b = \text{Tail}, c = \text{Tail}) \neq P(a = \text{Tail} \mid c = \text{Tail})$$

Theory: Expectations (5 marks)

Any function $f(X)$ of a discrete random variable X defines a random variable. (Define $p(f(X) = y) = \sum_{x:f(x)=y} p(x)$.) Similarly, given a joint probability $p(X_1, \dots, X_n)$, any function $f(X_1, \dots, X_n)$ is also a random variable. Show that the expected value of the sum of two random variables is the sum of the expectations. In symbols, show that

Let denote S the sample space underlying a random experiment with elements $s \in S$. Let's define two random variables X_1 and X_2 whose domains are S .

$$\begin{aligned}
&E[X_1 + X_2] \\
&= \sum_{x_1, x_2 \in S} (x_1 + x_2) \times P(X_1 = x_1, X_2 = x_2) \\
&= \sum_{x_1 \in S} \sum_{x_2 \in S} x_1 \times P(X_1 = x_1, X_2 = x_2) + \sum_{x_1 \in S} \sum_{x_2 \in S} x_2 \times P(X_1 = x_1, X_2 = x_2) \\
&= \sum_{x_1 \in S} x_1 \sum_{x_2 \in S} P(X_1 = x_1, X_2 = x_2) + \sum_{x_2 \in S} x_2 \sum_{x_1 \in S} P(X_1 = x_1, X_2 = x_2) \\
&= \sum_{x_1 \in S} x_1 \times P(X_1 = x_1) + \sum_{x_2 \in S} x_2 \times P(X_2 = x_2) \\
&= E[X_1] + E[X_2]
\end{aligned}$$

Practice: Decision Tree Learning with ID3 (10 marks)

Figure 1 provides data about whether a customer will wait for a table in a restaurant or not. Assume that ID3 splits first on the Pat attribute (for Patrons). Show the following for the branch $Pat = Full$.

1. The next attribute chosen by ID3. There may be a tie among several attributes; you can list all or just one of them.
2. The expected information gain associated with the next attribute. Compare this with the expected information gain for Hungry.
3. How you calculated the expected information gain.

Consider:

$$H(S) = -p_+ \times \log_2(p_+) - p_- \times \log_2(p_-)$$

$$\text{Gain}(S, A) = H(S) - \sum_{v \in \text{values}(A)} \frac{|S_v|}{|S|} H(S_v)$$

$$H(\text{Pat}=\text{Full})=0.91829 \quad [2T, 4F]$$

Alt:

$$H(\text{Alt}=\text{True})=0.97095 \quad [2T, 3F]$$

$$H(\text{Alt}=\text{False})=0 \quad [0T, 1F]$$

$$\begin{aligned} \text{Gain}(\text{Pat}=\text{Full}, \text{Alt}) &= H(\text{Pat}=\text{Full}) - \frac{5}{6}H(\text{Alt}=\text{True}) - \frac{1}{6}H(\text{Alt}=\text{False}) \\ &= 0.10917 \end{aligned}$$

Bar:

$$H(\text{Bar}=\text{True})=0.91829 \quad [1T, 2F]$$

$$H(\text{Bar}=\text{False})=0.91829 \quad [1T, 2F]$$

$$\begin{aligned} \text{Gain}(\text{Pat}=\text{Full}, \text{Bar}) &= H(\text{Pat}=\text{Full}) - \frac{1}{2}H(\text{Bar}=\text{True}) - \frac{1}{2}H(\text{Bar}=\text{False}) \\ &= 0 \end{aligned}$$

Fri:

$$H(\text{Fri}=\text{True})=1 \quad [2T, 3F]$$

$$H(\text{Fri}=\text{False})=0 \quad [0T, 1F]$$

$$\begin{aligned} \text{Gain}(\text{Pat}=\text{Full}, \text{Fri}) &= H(\text{Pat}=\text{Full}) - \frac{5}{6}H(\text{Fri}=\text{True}) - \frac{1}{6}H(\text{Fri}=\text{False}) \\ &= 0.1091 \end{aligned}$$

Hun:

$$H(\text{Hun}=\text{True})=1 \quad [2T, 2F]$$

$$H(\text{Hun}=\text{False})=0 \quad [0T, 2F]$$

$$\begin{aligned} \text{Gain}(\text{Pat}=\text{Full}, \text{Hun}) &= H(\text{Pat}=\text{Full}) - \frac{2}{3}H(\text{Hun}=\text{True}) - \frac{1}{3}H(\text{Hun}=\text{False}) \\ &= 0.2516 \end{aligned}$$

Price:

$$H(\text{Price}=\$)=1 \text{ [2T, 2F]}$$

$$H(\text{Price}=\$\$\$)=0 \text{ [0T, 2F]}$$

$$\begin{aligned} \text{Gain}(\text{Pat}=\text{Full}, \text{Price}) &= H(\text{Pat}=\text{Full}) - \frac{2}{3}H(\text{Price}=\$) - \frac{1}{3}H(\text{Price}=\$\$\$) \\ &= 0.2516 \end{aligned}$$

Rain:

$$H(\text{Rain}=\text{True})=0 \text{ [0T, 1F]}$$

$$H(\text{Rain}=\text{False})=0.97095 \text{ [2T, 3F]}$$

$$\begin{aligned} \text{Gain}(\text{Pat}=\text{Full}, \text{Rain}) &= H(\text{Pat}=\text{Full}) - \frac{1}{6}H(\text{Rain}=\text{True}) - \frac{5}{6}H(\text{Rain}=\text{False}) \\ &= 0.1091 \end{aligned}$$

Res:

$$H(\text{Res}=\text{True})=0 \text{ [0T, 2F]}$$

$$H(\text{Res}=\text{False})=1 \text{ [2T, 2F]}$$

$$\begin{aligned} \text{Gain}(\text{Pat}=\text{Full}, \text{Res}) &= H(\text{Pat}=\text{Full}) - \frac{2}{6}H(\text{Res}=\text{True}) - \frac{4}{6}H(\text{Res}=\text{False}) \\ &= 0.2516 \end{aligned}$$

Type:

$$H(\text{Type}=\text{Thai})=1 \text{ [1T, 1F]}$$

$$H(\text{Type}=\text{French})=0 \text{ [0T, 1F]}$$

$$H(\text{Type}=\text{Burger})=1 \text{ [1T, 1F]}$$

$$H(\text{Type}=\text{Italian})=0 \text{ [0T, 1F]}$$

$$\begin{aligned} \text{Gain}(\text{Pat}=\text{Full}, \text{Type}) &= H(\text{Pat}=\text{Full}) - \frac{2}{6}H(\text{Type}=\text{Thai}) - \frac{1}{6}H(\text{Type}=\text{French}) \\ &\quad - \frac{2}{6}H(\text{Type}=\text{Burger}) - \frac{1}{6}H(\text{Type}=\text{Italian}) = 0.2516 \end{aligned}$$

Est:

$$H(\text{Est}=10-30)=0 \text{ [1T, 1F]}$$

$$H(\text{Est}=30-60)=1 \text{ [1T, 1F]}$$

$$H(\text{Est}=>60)=1 \text{ [0T, 2F]}$$

$$\begin{aligned} \text{Gain}(\text{Pat}=\text{Full}, \text{Est}) &= H(\text{Pat}=\text{Full}) - \frac{2}{6}H(\text{Est}=10-30) - \frac{2}{6}H(\text{Est}=30-60) \\ &\quad - \frac{2}{6}H(\text{Est}=>60) = 0.2516 \end{aligned}$$

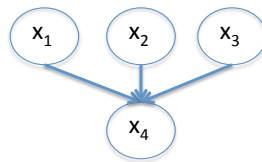
The highest information gain possible is 0.2516 and the tie is between Hungry, Price, Reservation, Type and Estimated Wait Time.

Example	Attributes										Target <i>WillWait</i>
	<i>Alt</i>	<i>Bar</i>	<i>Fri</i>	<i>Hun</i>	<i>Pat</i>	<i>Price</i>	<i>Rain</i>	<i>Res</i>	<i>Type</i>	<i>Est</i>	
X_1	<i>T</i>	<i>F</i>	<i>F</i>	<i>T</i>	<i>Some</i>	<i>\$\$\$</i>	<i>F</i>	<i>T</i>	<i>French</i>	<i>0-10</i>	<i>T</i>
X_2	<i>T</i>	<i>F</i>	<i>F</i>	<i>T</i>	<i>Full</i>	<i>\$</i>	<i>F</i>	<i>F</i>	<i>Thai</i>	<i>30-60</i>	<i>F</i>
X_3	<i>F</i>	<i>T</i>	<i>F</i>	<i>F</i>	<i>Some</i>	<i>\$</i>	<i>F</i>	<i>F</i>	<i>Burger</i>	<i>0-10</i>	<i>T</i>
X_4	<i>T</i>	<i>F</i>	<i>T</i>	<i>T</i>	<i>Full</i>	<i>\$</i>	<i>F</i>	<i>F</i>	<i>Thai</i>	<i>10-30</i>	<i>T</i>
X_5	<i>T</i>	<i>F</i>	<i>T</i>	<i>F</i>	<i>Full</i>	<i>\$\$\$</i>	<i>F</i>	<i>T</i>	<i>French</i>	<i>>60</i>	<i>F</i>
X_6	<i>F</i>	<i>T</i>	<i>F</i>	<i>T</i>	<i>Some</i>	<i>\$\$</i>	<i>T</i>	<i>T</i>	<i>Italian</i>	<i>0-10</i>	<i>T</i>
X_7	<i>F</i>	<i>T</i>	<i>F</i>	<i>F</i>	<i>None</i>	<i>\$</i>	<i>T</i>	<i>F</i>	<i>Burger</i>	<i>0-10</i>	<i>F</i>
X_8	<i>F</i>	<i>F</i>	<i>F</i>	<i>T</i>	<i>Some</i>	<i>\$\$</i>	<i>T</i>	<i>T</i>	<i>Thai</i>	<i>0-10</i>	<i>T</i>
X_9	<i>F</i>	<i>T</i>	<i>T</i>	<i>F</i>	<i>Full</i>	<i>\$</i>	<i>T</i>	<i>F</i>	<i>Burger</i>	<i>>60</i>	<i>F</i>
X_{10}	<i>T</i>	<i>T</i>	<i>T</i>	<i>T</i>	<i>Full</i>	<i>\$\$\$</i>	<i>F</i>	<i>T</i>	<i>Italian</i>	<i>10-30</i>	<i>F</i>
X_{11}	<i>F</i>	<i>F</i>	<i>F</i>	<i>F</i>	<i>None</i>	<i>\$</i>	<i>F</i>	<i>F</i>	<i>Thai</i>	<i>0-10</i>	<i>F</i>
X_{12}	<i>T</i>	<i>T</i>	<i>T</i>	<i>T</i>	<i>Full</i>	<i>\$</i>	<i>F</i>	<i>F</i>	<i>Burger</i>	<i>30-60</i>	<i>T</i>

Figure 1: Data Set for Question 5, Decision Tree Learning

Theory: Maximum Likelihood Parameter Estimation for Bayesian Networks (8 marks)

Consider learning the parameters for the Bayesian network shown in the figure.



Suppose we have a training set $(\mathbf{x}^1, \mathbf{x}^2, \dots, \mathbf{x}^N)$, where each $\mathbf{x}^i = (x_1^i, \dots, x_4^i)$ is a vector containing values for all 4 random variables in the network.

- Write down the likelihood and the log-likelihood of the training data given a parameter setting of the Bayes net. Please use the following notation.
 - θ_{ijk} = the conditional probability of node i taking on value k given that the parents of i are in state j . In our example, $i = 1, \dots, 4$, and $k = 1, \dots, L$. If $i = 1, 2, 3$, then $j = 0$, that is j is just a dummy index since the first three nodes have no parents. If $i = 4$,

then $j = 1, \dots, L^3$.

- (b) n_{ijk} = the number of training cases where node i takes on value k and the parents of i are in state j .

$$\begin{aligned}
 P(X | \theta) &= \prod_{n=1:N} P(x_1^n, x_2^n, x_3^n, x_4^n | \theta) \\
 &= \prod_{n=1:N} P(x_4^n | x_1^n, x_2^n, x_3^n, \theta) \times P(x_3^n | \theta) \times P(x_2^n | \theta) \times P(x_1^n | \theta) \\
 &= \left(\prod_{k=1:L} \theta_{10k}^{n_{10k}} \times \theta_{20k}^{n_{20k}} \times \theta_{30k}^{n_{30k}} \right) \times \left(\prod_{k=1:L} \prod_{j=1:L^3} \theta_{4jk}^{n_{4jk}} \right) \\
 &= \left(\prod_{i=1:3} \prod_{k=1:L} \theta_{i0k}^{n_{i0k}} \right) \times \left(\prod_{k=1:L} \prod_{j=1:L^3} \theta_{4jk}^{n_{4jk}} \right)
 \end{aligned}$$

$$\begin{aligned}
 \log(P(X | \theta)) &= \log\left(\left(\prod_{i=1:3} \prod_{k=1:L} \theta_{i0k}^{n_{i0k}}\right) \times \left(\prod_{k=1:L} \prod_{j=1:L^3} \theta_{4jk}^{n_{4jk}}\right)\right) \\
 &= \left(\sum_{i=1:3} \sum_{k=1:L} \log(\theta_{i0k}^{n_{i0k}})\right) + \left(\sum_{k=1:L} \sum_{j=1:L^3} \log(\theta_{4jk}^{n_{4jk}})\right) \\
 &= \left(\sum_{i=1:3} \sum_{k=1:L} n_{i0k} \times \log(\theta_{i0k})\right) + \left(\sum_{k=1:L} \sum_{j=1:L^3} n_{4jk} \times \log(\theta_{4jk})\right)
 \end{aligned}$$

2. Show that with binary nodes ($L = 2$) the maximum likelihood parameter values $\hat{\theta}_{ijk}$ are the conditional frequencies observed in the data:

$$\hat{\theta}_{ijk} = \frac{n_{ijk}}{\sum_{k'} n_{ijk'}}$$

Setting $L=2$:

$$\begin{aligned}
 &\left(\sum_{i=1:3} \sum_{k=1:2} n_{i0k} \times \log(\theta_{i0k})\right) + \left(\sum_{k=1:2} \sum_{j=1:8} n_{4jk} \times \log(\theta_{4jk})\right) \\
 &= \left(\sum_{i=1:3} n_{i01} \times \log(\theta_{i01}) + n_{i02} \times \log(1 - \theta_{i01})\right) \\
 &+ \left(\sum_{j=1:8} n_{4j1} \times \log(\theta_{4j1}) + n_{4j2} \times \log(1 - \theta_{4j1})\right)
 \end{aligned}$$

Derivate and equal to zero:

$$\begin{aligned}
 \frac{\partial \log(P(X|\theta))}{\partial \theta_{ij1}} &= 0 \\
 \implies \frac{\partial(n_{ij1} \times \log(\theta_{ij1}) + n_{ij2} \times \log(1 - \theta_{ij1}))}{\partial \theta_{ij1}} &= 0 \\
 \implies \frac{n_{ij1}}{\theta_{ij1}} - \frac{n_{ij2}}{1 - \theta_{ij1}} &= 0 \\
 \implies \theta_{ij1} \times (n_{ij1} + n_{ij2}) &= n_{ij1} \\
 \implies \theta_{ij1} &= \frac{n_{ij1}}{\sum_{k'} n_{ijk'}} \\
 \implies \theta_{ijk} &= \frac{n_{ijk}}{\sum_{k'} n_{ijk'}}
 \end{aligned}$$

The maximum likelihood result holds generally for a Bayesian network with discrete variables. I have given you a specific structure with binary variables only for the sake of concreteness.

Theory: Maximum Likelihood Parameter Estimation for a Gaussian Distribution (12 marks)

Consider a Gaussian or Normal Distribution with parameters mean μ and variance σ^2 and probability density function $f(x; \mu, \sigma^2)$. Suppose we have a training set $\mathbf{x} = (x^1, \dots, x^N)$ of observed values for random variable X . Let $\hat{\mu}$ and $\hat{\sigma}^2$ be the maximum likelihood estimates of the distribution parameters estimated from the training set.

1. Write down the log-likelihood function $\mathcal{L}(\mathbf{x}; \mu, \sigma^2)$.

Considering:

$$f(x; \mu, \sigma^2) = \frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{(x-\mu)^2}{2\sigma^2}}$$

$$P(D|\mu, \sigma^2) = \prod_{i=1:N} \frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{(x^i-\mu)^2}{2\sigma^2}}$$

$$\begin{aligned} \log(P(D|\mu, \sigma^2)) &= \sum_{i=1:N} \log\left(\frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{(x^i-\mu)^2}{2\sigma^2}}\right) \\ &= N \times \log\left(\frac{1}{\sigma\sqrt{2\pi}}\right) - \frac{1}{2\sigma^2} \sum_{i=1:N} (x^i - \mu)^2 \end{aligned}$$

2. Show that the maximum likelihood estimate for the distribution mean is the sample mean: $\hat{\mu} = \frac{1}{N} \sum_{i=1}^N x^i \equiv \bar{x}$.

$$\begin{aligned} \frac{\partial(P(D|\mu, \sigma^2))}{\partial\mu} &= 0 \\ \implies \frac{\partial(N \times \log(\frac{1}{\sigma\sqrt{2\pi}}) - \frac{1}{2\sigma^2} \sum_{i=1:N} (x^i - \mu)^2)}{\partial\mu} &= 0 \\ \implies \sum_{i=1:N} (x^i - \mu) &= 0 \\ \implies \sum_{i=1:N} x^i &= N\mu \\ \implies \mu &= \frac{\sum_{i=1:N} x^i}{N} \equiv \bar{x} \end{aligned}$$

3. Show that the maximum likelihood estimate for the distribution variance is the sample variance: $\hat{\sigma}^2 = \frac{1}{N} \sum_{i=1}^N (x^i - \bar{x})^2$ assuming that $\mu = \bar{x}$.

$$\begin{aligned} \frac{\partial(P(D|\mu, \sigma))}{\partial\sigma} &= 0 \\ \implies \frac{\partial(N \times \log(\frac{1}{\sigma\sqrt{2\pi}}) - \frac{1}{2\sigma^2} \sum_{i=1:N} (x^i - \mu)^2)}{\partial\sigma} &= 0 \\ \implies -\frac{N}{\sigma} + \frac{1}{\sigma^3} \sum_{i=1:N} (x^i - \mu)^2 &= 0 \\ \implies \frac{N}{\sigma} &= \frac{1}{\sigma^3} \sum_{i=1:N} (x^i - \mu)^2 \\ \implies \sigma^2 &= \frac{1}{N} \sum_{i=1:N} (x^i - \mu)^2 \\ \implies \sigma^2 &= \frac{1}{N} \sum_{i=1:N} (x^i - \bar{x})^2 \end{aligned}$$