

# UNIVERSITY OF BIRMINGHAM

**School of Computer Science**

**Artificial Intelligence 1**

Main Summer Examinations 2023

Time allowed: 2 hours

[Answer all questions]

## Question 1 Clustering

Suppose we have five observations,  $\mathbf{x}^{(1)}, \mathbf{x}^{(2)}, \mathbf{x}^{(3)}, \mathbf{x}^{(4)}$  and  $\mathbf{x}^{(5)}$ , for which we compute the following dissimilarity matrix

$$D = \begin{matrix} & \mathbf{x}^{(1)} & \mathbf{x}^{(2)} & \mathbf{x}^{(3)} & \mathbf{x}^{(4)} & \mathbf{x}^{(5)} \\ \mathbf{x}^{(1)} & 0 & 0.3 & 0.6 & 0.7 & 0.8 \\ \mathbf{x}^{(2)} & & 0 & 0.1 & 0.2 & 0.4 \\ \mathbf{x}^{(3)} & & & 0 & 0.45 & 0.25 \\ \mathbf{x}^{(4)} & & & & 0 & 0.9 \\ \mathbf{x}^{(5)} & & & & & 0 \end{matrix},$$

- (a) Based on the given dissimilarity matrix, hierarchically cluster the observations using **complete linkage**. Sketch the dendrogram, clearly illustrating the height at which each cluster fusion occurs. [10 marks]

Shay

1) a) ~~Handwritten scribbles~~

|       | $x_1$ | $x_2$ | $x_3$ | $x_4$ | $x_5$ |
|-------|-------|-------|-------|-------|-------|
| $x_1$ | 0     | 0.3   | 0.6   | 0.7   | 0.8   |
| $x_2$ |       | 0     | 0.1   | 0.2   | 0.4   |
| $x_3$ |       |       | 0     | 0.45  | 0.25  |
| $x_4$ |       |       |       | 0     | 0.9   |
| $x_5$ |       |       |       |       | 0     |

$x_2, x_3$  @ 0.1

|           | $x_1$ | $x_2 x_3$       | $x_4$ | $x_5$ |
|-----------|-------|-----------------|-------|-------|
| $x_1$     | 0     | 0.6             | 0.7   | 0.8   |
| $x_2 x_3$ |       | 0               | 0.45  | 0.4   |
| $x_4$     |       | <del>0.45</del> | 0     | 0.9   |
| $x_5$     |       |                 |       | 0     |

$x_2 x_3, x_5$

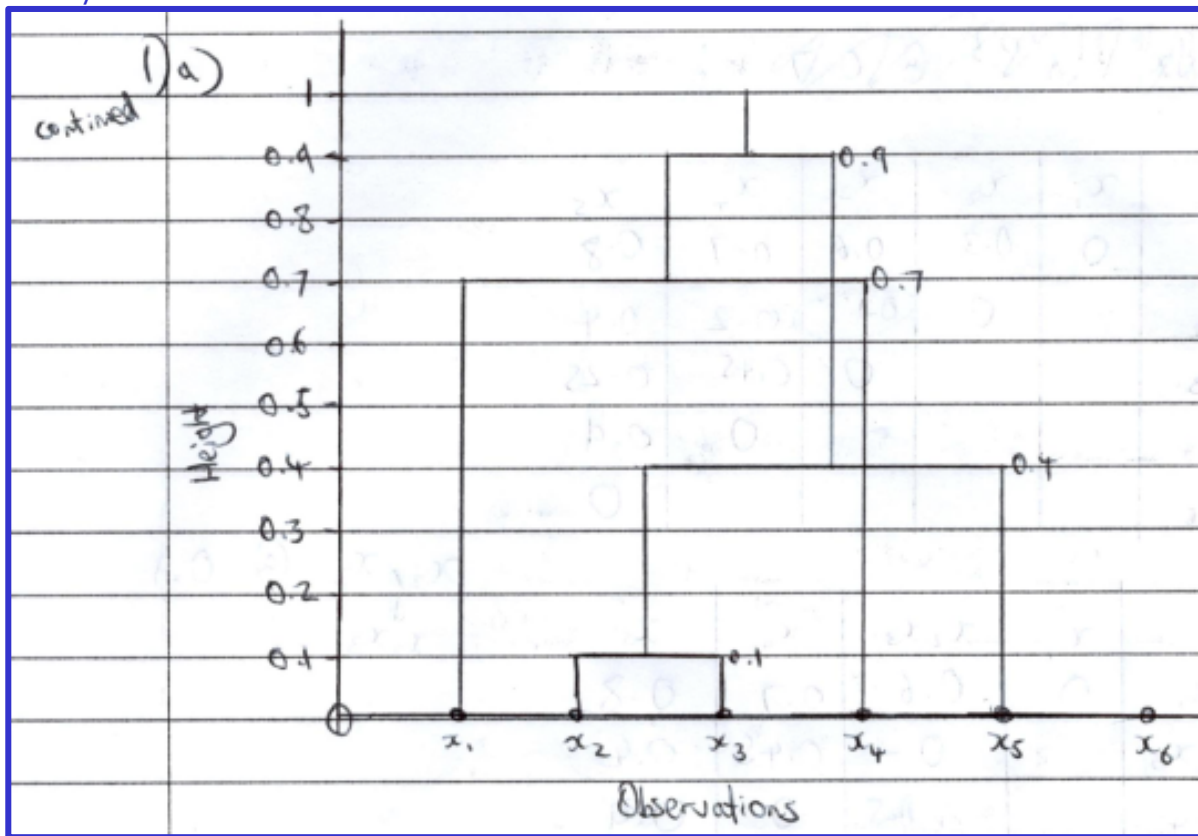
|               | $x_1$ | $x_2 x_3 x_5$ | $x_4$ |
|---------------|-------|---------------|-------|
| $x_1$         | 0     | 0.8           | 0.7   |
| $x_2 x_3 x_5$ |       | 0             | 0.4   |
| $x_4$         |       | 0.4           | 0     |

$x_1, x_4$

|               | $x_1 x_4$ | $x_2 x_3 x_5$ |
|---------------|-----------|---------------|
| $x_1 x_4$     | 0         | 0.9           |
| $x_2 x_3 x_5$ |           | 0             |

$x_1 x_4, x_2 x_3 x_5$

$x_1 x_2 x_3 x_4 x_5$  @ 0.9



- (b) Assume that a clustering algorithm produces the following two clusters:  $C_1 = (x^{(1)}, x^{(2)})$  and  $C_2 = (x^{(3)}, x^{(4)}, x^{(5)})$ . For each of the observations in cluster  $C_1$ , compute the Silhouette coefficient using the information from the dissimilarity matrix. Comment on the suitability of their assignment to cluster  $C_1$ .

Note: The Silhouette coefficient of an observation  $x$  is defined as:

$$SC(x) = \frac{b(x) - a(x)}{\max\{b(x), a(x)\}},$$

where  $a(x)$  is the average dissimilarity of  $x$  with respect to other observations in its cluster, and  $b(x)$  is the minimum average dissimilarity of  $x$  with respect to all clusters to which it does not belong. [10 marks]

Phoebe

$$C_1 = (\vec{x}^{(1)}, \vec{x}^{(2)})$$

$$\begin{aligned} \underline{x^{(1)}}: \quad a(x^{(1)}) &= \frac{0.3}{1} = 0.3 \\ b(x^{(1)}) &= \frac{0.6 + 0.7 + 0.8}{3} = 0.7 \\ SC(x^{(1)}) &= \frac{0.4}{0.7} = 0.5714 \end{aligned}$$

$$a(\vec{z}) = \frac{\text{sum of all distances from } \vec{z} \text{ in cluster}}{|\mathcal{C}_1| - 1}$$

$$b(\vec{z}) = \min \frac{\text{sum of all distances in cluster } \mathcal{C}_k}{|\mathcal{C}_k|}$$

$$SC(\vec{z}) = \frac{b(\vec{z}) - a(\vec{z})}{\max\{a(\vec{z}), b(\vec{z})\}}$$

for 1 example

$$\begin{aligned} \underline{x^{(2)}}: \quad a(x^{(2)}) &= 0.3 \\ b(x^{(2)}) &= \frac{0.1 + 0.2 + 0.4}{3} = 0.2333 \end{aligned}$$

$$SC(x^{(2)}) = \frac{-0.0667}{0.3} = -0.222$$

$$b) C_1 = (x_1, x_2) \quad C_2 = (x_3, x_4, x_5)$$

$$SC(x) = \frac{b(x) - a(x)}{\max\{b(x), a(x)\}}$$

~~$$a(x_1) = \dots$$~~  
~~$$a(x_2) = \dots$$~~

and

$$a(x_1) = 0.3$$

$$a(x_2) = 0.3$$

$$b(x_1) = (0.6 + 0.7 + 0.8) / 3$$

$$= 0.7$$

$$b(x_2) = (0.1 + 0.2 + 0.4) / 3$$

$$= 0.23$$

$$SC(x_1) = \frac{0.7 - 0.3}{0.7} = 0.571$$

✓ well suited to its cluster.

$$SC(x_2) = \frac{0.23 - 0.3}{0.3} = -0.233$$

x would be more optimally clustered in a different cluster.



## Question 2 Supervised Learning

- (a) The following pseudo-code represents one iteration through the training set for gradient descent applied to univariate linear regression.

```

1 cost = 0;
2 w0 = 0;
3 w1 = 0;
4 for j=1 to size(trainingSet) do
5     f = w0 + w1*x(j);
6     cost = cost + (y(j) - f)2;
7     w0 = w0 - α * (f - y(j));
8     w1 = w1 - α * (f - y(j)) * x(j);

```

Assume that the value of the learning rate,  $\alpha$  is 1.

Give the numerical values of 'cost', 'w0', and 'w1' at the end of the execution of this pseudo-code for the following training set:  $\{(-3, -1), (1, 1), (2, 5)\}$ . **Show all your working.** [10 marks]

Shay

|                                   |          |        |        |
|-----------------------------------|----------|--------|--------|
| 2) a)                             | cost = 0 | w0 = 0 | w1 = 0 |
| (-3, -1)                          |          |        |        |
| $f = 0 + 0(-3) = 0$               |          |        |        |
| $cost = 0 + ((-1) - 0)^2 = 1$     |          |        |        |
| $w0 = 0 - 1(0 - (-1)) = -1$       |          |        |        |
| $w1 = 0 - 1(0 - (-1))(-3) = 3$    |          |        |        |
| cost = 1      w0 = -1      w1 = 3 |          |        |        |
| (1, 1)                            |          |        |        |
| $f = -1 + 3(1) = 2$               |          |        |        |
| $cost = 1 + (1 - 2)^2 = 2$        |          |        |        |
| $w0 = -1 - 1(2 - 1) = -2$         |          |        |        |
| $w1 = 3 - 1(2 - 1)(1) = 2$        |          |        |        |
| cost = 2      w0 = -2      w1 = 2 |          |        |        |
| (2, 5)                            |          |        |        |
| $f = -2 + 2(2) = 2$               |          |        |        |
| $cost = 2 + (5 - 2)^2 = 11$       |          |        |        |
| $w0 = -2 - 1(2 - 5) = 1$          |          |        |        |
| $w1 = 2 - 1(2 - 5)(2) = 8$        |          |        |        |
| cost = 11      w0 = 1      w1 = 8 |          |        |        |

```

1 cost = 0;
2 w0 = 0;
3 w1 = 0;
4 for j=1 to size(trainingSet) do
5     f = w0 + w1*x(j);
6     cost = cost + (y(j) - f)2;
7     w0 = w0 - α * (f - y(j));
8     w1 = w1 - α * (f - y(j)) * x(j);

```

- (b) Suppose that you want to use a k-Nearest Neighbours classification, but the distance metric is not explicitly specified to you. Instead, you are given a "black box" where you input a set of points  $\mathbf{x}^{(1)}, \mathbf{x}^{(2)}, \dots, \mathbf{x}^{(N)}$  and a new point  $\mathbf{x}$ , and the black box outputs the nearest neighbour of  $\mathbf{x}$  and its corresponding class label.

Can you construct a k-Nearest Neighbour classification algorithm based on this black box alone? The suggested algorithm should work for any given value of  $k \in \{1, 2, \dots, N\}$ . If so, describe your algorithm. If not, explain why not? **[10 marks]**

Phoebe

```

kNN ( trainingSet,  $\vec{x}$ , k ):
    labels = []
    if k > 0:
        for (int i=0; i < k; i++):
            labels.append(blackbox.label ( trainingSet,  $\vec{x}$  ))
            tempSet = trainingSet.remove ( blackbox.closestneighbour (trainingSet,  $\vec{x}$  ))
        for x in range (0, len(labels)):
             $\vec{x}$ .label = findMostPopularLabel (labels)

```

Shay

2) b) The k-Nearest Neighbour algorithm looks is ~~prede~~ initialised with a k value. This k value determines the number of neighbors to consider when classifying a data point. It will classify a data point based on what the majority of the its k neighbors are classified as. With the black box, we won't be able to construct

With a set of points, let's ~~image~~ we can pass in and a new point  $x$ , we can pass them all into the black box to get the closest point to  $x$ . Say our set of initial points is  $\{(1), (2), (5), (7)\}$  and our new point  $x$  is  $(8)$ , we would receive  $(7)$  from the black box output if the black box uses simple ~~Euclidean~~ Manhattan distance. Then, pass in the set of points excluding the  $(7)$  we know of, so  $\{(1), (2), (5)\}$  and the new point  $(8)$ , to get the second nearest neighbour. Repeat k times to get the k closest neighbors, and then run the rest of the algorithm as normal, ~~then~~ classifying the new point  $(8)$  based on a majority vote of its k nearest neighbours.

### Question 3 Optimisation

(a) Give one strength and one weakness of Hill Climbing. **Justify** your answer. [10 marks]

| +ve                                                                                                                                                                                                                                                                                              | -ve                                                                                                                                                                                                                            |
|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <ul style="list-style-type: none"><li>- no heuristics needed</li><li>- greedy algorithm</li><li>- less computing power</li><li>- easier to implement in code</li><li>- lower time complexity<br/>since greedy compared to smth like simulated annealing</li><li>- low space complexity</li></ul> | <ul style="list-style-type: none"><li>- gets stuck in local min/max or plateaus</li><li>- not complete (e.g. asymptote)</li><li>- can get stuck in infinite/loopy paths</li><li>- only looks at immediate neighbours</li></ul> |

Shay

|                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                        |
|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <p>3a) One strength of Hill Climbing is that it isn't very computationally heavy. We just need to generate direct neighbours and then do a single comparison for each neighbour to see if it improves the solution. Furthermore, this simplicity will lead us to an optimal solution in the event that there are no suboptimal local minima/maxima.</p> <p>One weakness of Hill Climbing is that most of the time it will not reach an optimal solution. Because it only checks direct neighbouring states, the algorithm can easily get stuck in a local minimum or maximum, and returning a suboptimal solution.</p> |
|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|

### Optimisation Problem

where  $\mathbf{x}$  is a vector of size  $N$ , and  $\forall i \in \{1, \dots, N\}$ ,  $x_i = 0$  if the builder does not accept job  $i$  and  $x_i = 1$  otherwise; and  $f(\mathbf{x})$  is a function that calculates the income.

**[10 marks]**

b) Does the **representation** reflect the intended problem?

yes. design variable is variable, where the builder either accepts or doesn't accept the job.

Does the **initialisation** reflect the intended problem?

yes. it will seem "random" how builders decide to accept their jobs, and we must check each  $N$  builders

Do the **neighbourhood operators** reflect the intended problem?

yes.

- if the neighbour's  $q(z) > 0$  (broken constraint) we make it so they've "not accepted" the job (i.e.  $x_i^j = 0$ )

**function** SIMULATED-ANNEALING(*problem, schedule*) **returns** a solution state

*current*  $\leftarrow$  *problem*.INITIAL

**for**  $t = 1$  **to**  $\infty$  **do**

*T*  $\leftarrow$  *schedule*( $t$ )

**if**  $T = 0$  **then return** *current*

*next*  $\leftarrow$  a randomly selected successor of *current*

$\Delta E \leftarrow \text{VALUE}(\text{current}) - \text{VALUE}(\text{next})$

**if**  $\Delta E > 0$  **then** *current*  $\leftarrow$  *next*

**else** *current*  $\leftarrow$  *next* only with probability  $e^{\Delta E/T}$

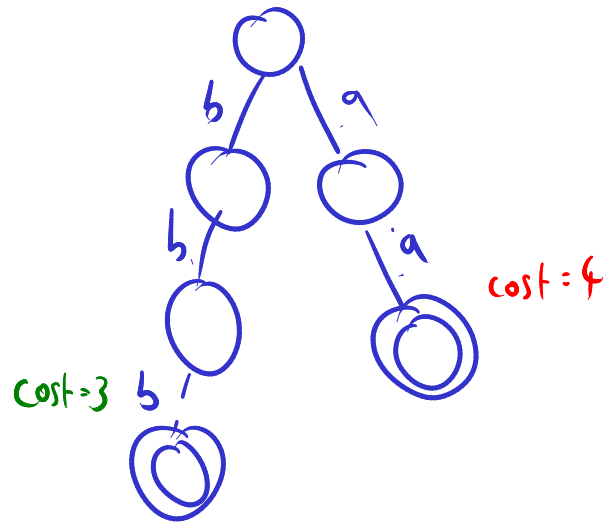
### Question 4 Search

- (a) Consider a problem with two possible actions:  $a$  and  $b$ . The cost of  $a$  is 2 and the cost of  $b$  is 1. The shallowest goal node can be at any level of the tree and the problem solution may involve a sequence of these actions.

Is breadth-first search optimal to solve this problem, given this cost function? **Justify** your answer.

**[10 marks]**

No - it can find a short path (in terms of number of actions taken) but will not necessarily find one with the lowest cost.



## Shay

4) a) A breadth-first search ~~explo~~ explores the nodes in a graph on a level by level basis. Since the cost of expanding and adding an "a" action is double that of expanding a "b" action, a breadth-first search may not be as efficient or optimal as ~~expanding the~~ iteratively expanding from the node with the smallest path cost that we haven't expanded from yet. e.g. instead of expanding a, b, aa, ab, ba, bb

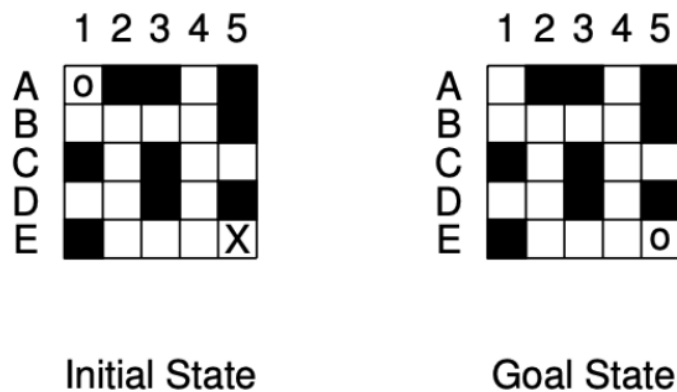
cost: 2 1 4 3 3 2 use a path  
we expand b, a, bb, ab, ba, aa finding algo  
cost: 1 2 2 3 3 4 like A\*

since then we may find the shallowest goal node ~~goal~~ ~~reaching~~ with lower cost.

(b) An agent must traverse a maze in order to find and collect a treasure. This task can be formulated as a search problem as in the following:

- Initial state: the agent is in A1.
- Goal test: the agent is in E5 (automatically collects the treasure).
- Actions:  $move(n, n')$  moves the agent from state  $n$  to a state  $n'$  that is to the left of, or the right of, or above or below state  $n$ .
- Transition model: see figure below. States depicted in black, e.g., A2, are considered obstacles and cannot be traversed by the agent.
- Path cost: each action has cost 1.

Generate the A\* tree until the goal node is found.



The heuristic to be used is the following:

$$h(n) = d_M(n, n_{goal}) = \sum_{i=1}^2 |p_i - q_i|,$$

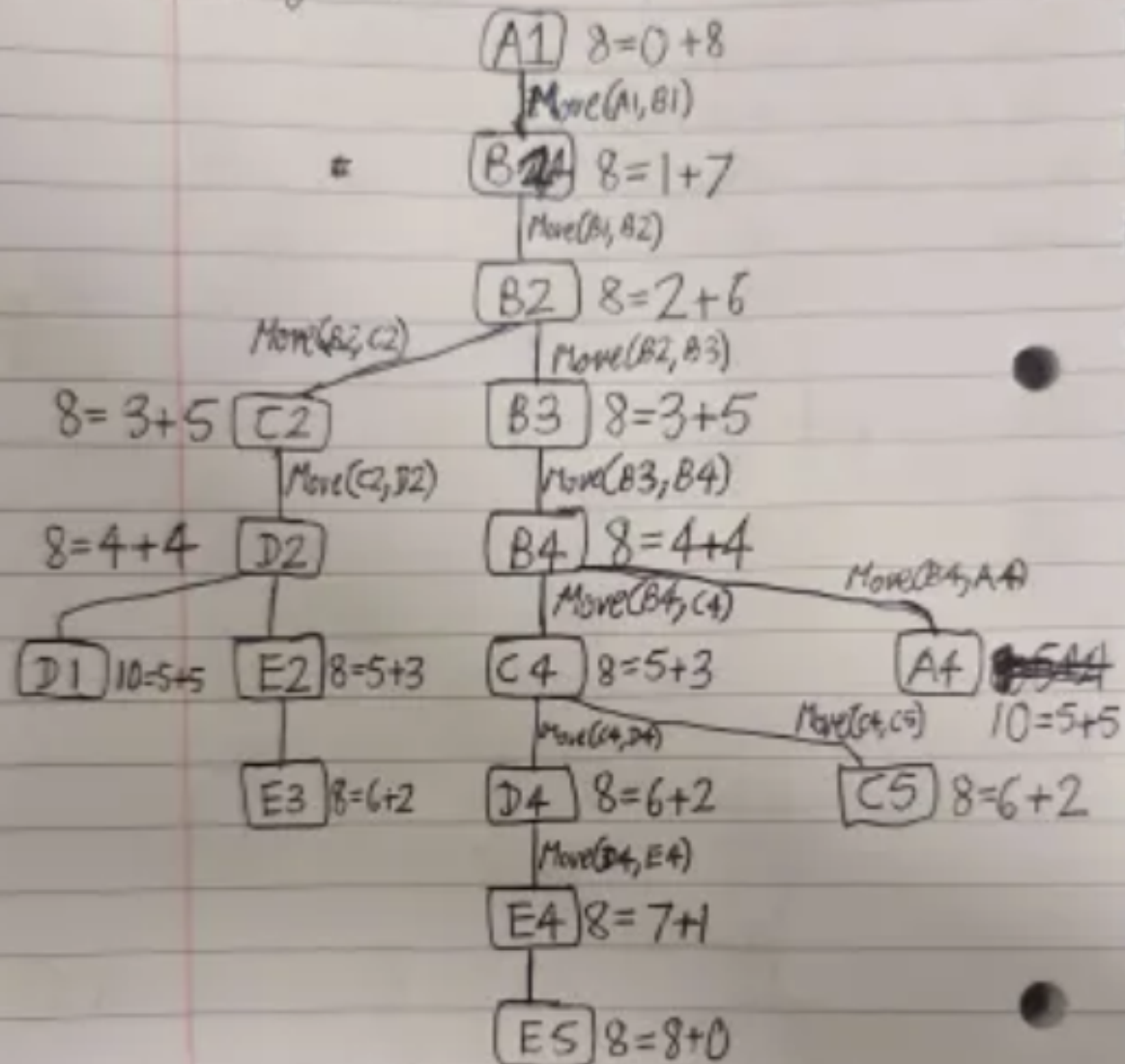
where  $d_M(n, n_{goal})$  represents the Manhattan distance between node  $n$  and the goal node  $n_{goal}$ . When comparing the first coordinate of each state, assume the following values for each of the possible letters in the grid:  $A = 1$ ,  $B = 2$ ,  $C = 3$ ,  $D = 4$  and  $E = 5$ . Order of expansion: if  $f(n)$  produces the same lowest value for two nodes, expand the nodes in alphanumerical order (e.g., A2 expanded before B1). Additionally, if  $g(n)$  also produces the same lowest value for two nodes, keep the the node that has been added last into the frontier. Recall that the evaluation function  $f(n) = g(n) + h(n)$  is the sum of the cost to reach node  $n$  and the heuristic calculated at node  $n$ .

Write down the following:

- Search tree produced by A\*, indicating the  $f(n)$ ,  $g(n)$  and  $h(n)$  values of each node  $n$ ; which nodes are in the frontier when the algorithm terminates; and which nodes have been pruned as a result of the above instructions.
- The solution retrieved by A\* and its cost.
- The sequence of nodes visited by A\* in the order they are visited. Note: you can identify a node through its state, e.g., A1 or B1.

**[10 marks]**

Visited  $[A1, B1, B2, B3, B4, C2, C4, \cancel{B1}, C5, D2, D4, E2, E3, E4, E5]$   
 $f(n) = g(n) + h(n)$



Frontier:  $\{D1, A4\}$