

Bayes rule:
$$P(X|Y) = P(Y|X) \cdot P(X) \cdot P(Y)$$

$$= \alpha P(Y|X) \cdot P(X)$$

Marginalization:

$$P(\phi) = \sum_{w: w \models \phi} P(w)$$

Summing out:

$$P(X) = \sum_Y P(X, Y)$$

$$P(X) = \sum_{e \in E} P(X|e) \cdot \underline{P(e)}$$

Normalization

$$P(X|E) = P(X, E) \cdot P(E)$$

$$= \alpha \cdot P(X, E)$$

because $\sum_{y \in Y} P(y|E) \cdot P(E) = 1$

Absolute independence: $P(X, Y) = P(X) \cdot P(Y)$

Conditional independence: $P(X, Y|Z) = P(X|Z) \cdot P(Y|Z)$
 $P(X|Y, Z) = P(X|Z)$ etc...

Chain rule for Bayesian networks:

known from the network

$$P(x_1 - x_n) = \prod_i P(x_i | \text{Parent}(x_i))$$

From samples to probability

$$P(x_1 - x_n) = \lim_N (N(x_1 - x_n) / N)$$

rejection sampling: $\rightarrow$ taking only those satisfying evidence

$$P(x|e) = P(x, e) \cdot P(e)$$

likelihood sampling:

Fix evidence variables and sample only remaining

We get:

$$P(z, e) = \prod_i (z_i | \text{Parents}(z_i))$$

We need weighting:

$$w(z, e) = \prod_j (e_j | \text{parents}(e_j))$$

$$P(x, e) = \alpha \cdot N(x, e) \cdot w(x, e)$$

Filtering:

Must define f such that $P(X_{t+1} | e_{1:t+1}) = f(e_{t+1}, P(X_t | e_{1:t}))$

$$\begin{aligned} P(X_{t+1} | e_{1:t+1}) &= P(X_{t+1} | e_{1:t}, e_{t+1}) && P(X|Y) = \alpha \cdot P(Y|X) \cdot P(X) \\ &= \alpha \cdot P(e_{t+1} | X_{t+1}, e_{1:t}) \cdot P(X_{t+1} | e_{1:t}) \\ &= \alpha \cdot P(e_{t+1} | X_{t+1}) \cdot P(X_{t+1} | e_{1:t}) \\ &= \alpha \cdot P(e_{t+1} | X_{t+1}) \cdot \sum_{X_t} P(X_{t+1} | X_t) \cdot \underbrace{P(X_t | e_{1:t})}_{f_{1:t}} \end{aligned}$$

$$f_{1:t+1} = \alpha \cdot \text{Forward}(f_{1:t}, e_{t+1})$$

Prediction:

$$P(X_{t+h+1} | e_{1:t}) = \sum_{X_{t+h}} P(X_{t+h+1} | X_{t+h}) \cdot P(X_{t+h} | e_{1:t})$$

Uverjajo se, da stacionarni distribuciji damon Markovskijev procesom

Smoothing:

$$\begin{aligned} P(X_n | e_{1:t}) &= P(X_n | e_{1:n}, e_{n+1:t}) \\ &= \alpha \cdot P(e_{n+1:t} | X_n, e_{1:n}) \cdot P(X_n | e_{1:n}) \\ &= \alpha \cdot \underbrace{P(e_{n+1:t} | X_n)}_{b_{n+1:t}} \cdot \underbrace{P(X_n | e_{1:n})}_{f_{1:n}} \end{aligned}$$

$$\begin{aligned}
P(e_{k+1:t} | X_k) &= \sum_{X_{k+1}} P(e_{k+1:t} | X_k, X_{k+1}) \cdot P(X_{k+1} | X_k) \\
&= \sum_{X_{k+1}} P(e_{k+1:t} | X_{k+1}) \cdot P(X_{k+1} | X_k) \\
&= \sum_{X_{k+1}} P(e_{k+1}, e_{k+2:t} | X_{k+1}) \cdot P(X_{k+1} | X_k) \\
&= \sum_{X_{k+1}} P(e_{k+1} | X_{k+1}) \cdot \underbrace{P(e_{k+2:t} | X_{k+1})}_{b_{k+2:t}} \cdot P(X_{k+1} | X_k)
\end{aligned}$$

Using backward alg.

$$b_{k+1:t} = \text{backward}(b_{k+2:t}, e_{k+1})$$

$$b_{t+1:t} = 1 = P(e_{t+1:t} | X_t) = P(1 | X_t)$$

Most likely explanations

Viterbi alg. (going through graph)

$$\text{argmax}_{X_{1:t}} P(X_{1:t} | e_{1:t})$$

$$\max_{X_1 \dots X_t} P(X_1, \dots, X_t, X_{t+1} | e_{1:t+1})$$

$$= \alpha \cdot P(e_{t+1} | X_{t+1}) \cdot \max_{X_t} (P(X_{t+1} | X_t) \cdot \max_{X_1 \dots X_{t-1}} P(X_1 \dots X_{t-1} | e_{1:t-1}))$$

Bellman equations: $U(s) = R(s) + \gamma \cdot \max_a \sum_{s'} P(s'|a,s) \cdot U(s')$

best: $\pi^*(s) = \operatorname{argmax}_a \sum_{s'} P(s'|a,s) \cdot U(s')$

MDP:

Transition model: $P(s'|a,s)$

Reward: $R(s)$

Utility function: $U([s_0, s_1, \dots]) = R(s_0) + \gamma R(s_1) + \gamma^2 R(s_2) \dots$

Bellman: value-iteration

First set some default $U(s)$

$$U_{i+1}(s) = R(s) + \gamma \cdot \max_a \sum_{s'} P(s'|a,s) \cdot U_i(s')$$

Repeat until some threshold in max diff in one iteration

$$\pi^*(s) = \operatorname{argmax}_{\pi} U^{\pi}(s)$$

Bellman: policy-iteration

1) Policy evaluation:

$$U(s) = R(s) + \gamma \cdot \sum_{s'} P(s'|s, \pi(s)) \cdot U(s')$$

2) Policy-improvement

$$\pi_{i+1}(s) = \operatorname{argmax}_a \sum_{s'} P(s'|a,s) \cdot U^{\pi_i}(s')$$

→ pokud existuje lepší akce než mi dává policy π_i

$$H(v) = - \sum_u p(v_u) \cdot \log_2(p(v_u))$$

$$B(v) = -q \cdot \log_2 q - (1-q) \cdot \log_2 (1-q)$$

Information gain mi $\frac{p}{n}$, jak moc se mi modelí množina

$$R_{\text{remainder}}(A) = \sum_k \left(B(p_k / (n_k + p_k)) \cdot ((n_k + p_k) / (n + h)) \right)$$

$$\text{Gain}(A) = B(p / (n + p)) - R_{\text{remainder}}(A)$$

Solving optimal parameters w_0, w_1 analytically:

$$\frac{\partial \sum_j (y_j - h(x_j))^2}{\partial w_0} = 0$$

$$\frac{\partial \sum_j (y_j - h(x_j))^2}{\partial w_1} = 0$$

$$w_1 = (N \sum_j x_j y_j - \sum_j x_j \sum_j y_j) / (N \sum_j x_j^2 - (\sum_j x_j)^2)$$

$$w_0 = (\sum_j y_j - w_1 \sum_j x_j) / N$$

With gradient descent:

$$w_i = w_i - \alpha \cdot \frac{\partial \text{loss}(w_i)}{\partial w_i}$$

$$w_0 = w_0 + \alpha \cdot (\sum_j y_j - h(x_j))$$

$$w_1 = w_1 + \alpha \cdot (\sum_j y_j - h(x_j)) \cdot x_j$$

—> jde snáze z hroze
pro chybí bod
dokud nekonverguje

Finding linear separator

$$w_i = w_i + \alpha \cdot (y - h(x)) \cdot x_i$$

Gradient descent method

$$w_{ij} = w_{ij} - \alpha \cdot \frac{\partial \text{Loss}(h_w)}{\partial w_{ij}}$$

Learning in feed-forward multilayer network
and how to measure errors in hidden layers?

Backpropagate!

- initialize random values to inner weights
- calculate the network output based on examples
- calculate and modify error for each node:

$$\Delta_j = g'(in_j) \cdot (y_j - a_j)$$

↳ g je primková funkce

- propagate Δ values back

$$\Delta_i = g'(in_i) \cdot \sum_j w_{ij} \cdot \Delta_j$$

- update the weights

Nearest neighbour model:

Minkowski distance measure

$$L^p(x_j, x_q) = \left(\sum_i |x_{j,i} - x_{q,i}|^p \right)^{1/p}$$

$$P(h_i | d) = \alpha \cdot P(d | h_i) \cdot P(h_i)$$

$$P(X | d) = \sum_i P(X | d, h_i) \cdot P(h_i | d) = \sum_i P(X | d) \cdot P(h_i | d)$$

$$P(d | h_\theta) = \sum_j P(d_j | h_\theta) = \theta^c \cdot (1-\theta)^l$$

Minkowski distance

$$L^p(x_j, x_q) = \left(\sum_i |x_{ji} - x_{qi}|^p \right)^{1/p}$$

Error-weights backpropagation

- Start with some random weights
- Calculate output on the last level
- calculate the error at the output level: $\Delta_j = g'(in_j) \cdot (y_j - a_j)$
- propagate error to the previous layer: $\Delta_i = g'(in_i) \cdot \sum_j w_{ij} \cdot \Delta_j$
- update the weights

Finding the linear separator

$$w_i = w_i - \alpha \cdot (y - h(x)) \cdot x_i$$

$$Q(s, a) = Q(s, a) + \alpha \cdot (R(s) + \gamma \max_{a'} Q(s, a') - Q(s, a))$$