

$$\begin{aligned} A[X+Y] &= A[X] + b, E[X+Y] = E[X] + E[Y], \text{ if } X \perp Y \\ E[XY^T] &= E[X]E[Y^T], E[X|Y=y] = \int xp_{X|Y}(x|y)dx \\ E_Y[E_X[X|Y]] &= E[X], E[g(X)] = \int g(x) \cdot p(x)dx, \\ \text{Cov}[X, Y] &\doteq E[(X - E[X])(Y - E[Y])^T] = E[XY^T] - E[X]E[Y^T], \\ \text{Cov}[X, Y] &= \text{Cov}[Y, X]^T, \text{Cov}[AX+b, BY+d] = A\text{Cov}[X, Y]B^T, \\ X \perp Y &\Rightarrow \text{Cov}[X, Y] = 0, \text{Var}[X] = \text{Cov}[X, X], \text{Var}[AX+b] = \\ &= A\text{Var}[X]A^T, \text{Var}[X+Y] = \text{Var}[X] + \text{Var}[Y] - 2\text{Cov}[X, Y], \\ \text{Var}[X|Y] &\doteq E[(X - E[X|Y])(X - E[X|Y])^T | Y] \\ \text{Var}[X] &= E_Y[\text{Var}_X[X|Y]] + \text{Var}_Y[E_X[X|Y]] \end{aligned}$$

C.of variable: X with density p_X and g diff.able and invertible function: then $Y = g(X)$, $p_Y(y) = p_X(g^{-1}(y)) \cdot |\det(Dg^{-1}(y))|$
 $\mathcal{N}(x; \mu, \sigma^2) \doteq \frac{1}{\sqrt{2\pi\sigma^2}} \exp(-((x - \mu)^2)/(2\sigma^2))$

$$\begin{aligned} \mathcal{N}(x; \mu, \Sigma) &\doteq 1/(\sqrt{(2\pi)^n \det(\Sigma)}) \exp(-\frac{1}{2}(x - \mu)^T \Sigma^{-1}(x - \mu)) \\ \log \mathcal{N}(x; \mu, \Sigma) &= -(1/2)x^T \Sigma^{-1}x + \mu^T \Sigma^{-1}x + \text{const} \end{aligned}$$

Gaussians are independent $\Leftrightarrow$ their Cov = 0. For *index sets* A, B:
 $X_A \sim \mathcal{N}(\mu_A, \Sigma_{AA})$, and $X_A | X_B = x_b \sim \mathcal{N}(\mu_{A|B}, \Sigma_{A|B})$ where

$$\begin{aligned} \mu_{A|B} &\doteq \mu_A + \Sigma_{AB}\Sigma_{BB}^{-1}(x_B - \mu_B), \Sigma_{A|B} \doteq \Sigma_{AA} - \Sigma_{AB}\Sigma_{BB}^{-1}\Sigma_{BA} \\ AX + b &\sim \mathcal{N}(A\mu + b, A\Sigma A^T), X = \Sigma^{1/2}Y + \mu \text{ for } Y \sim \mathcal{N}(0, I) \\ X \text{ and } X' &\text{ independent GRVs then } X + X' \sim \mathcal{N}(\mu + \mu', \Sigma + \Sigma') \\ \text{For jointly Gaussian RVs } X_A, X_B, &\text{ independent } \varepsilon \sim \mathcal{N}(0, \Sigma_{A|B}): \end{aligned}$$

$$\begin{aligned} X_A &= AX_B + b + \varepsilon \text{ for } A \doteq \Sigma_{AB}\Sigma_{BB}^{-1}, \text{ and } b \doteq \mu_A - \Sigma_{AB}\Sigma_{BB}^{-1}\mu_B \\ (\text{Jensen's}): \text{convex function } g: \mathbb{R} &\rightarrow \mathbb{R}, (\text{e.g. } S[u]) \ g(E[X]) \leq E[g(X)] \\ S[u] &\doteq -\log(u), H[p] \doteq E_{x \sim p}[S[p(x)]] = E_{x \sim p}[-\log p(x)] \end{aligned}$$

$$\begin{aligned} H[p|q] &\doteq E_{x \sim p}[S[q(x)]] = E_{x \sim p}[-\log q(x)] = H[p] + \text{KL}(p||q) \\ \text{KL}(p||q) &\doteq H[p|q] - H[p] = E_{\theta \sim p}[S[q(\theta)] - S[p(\theta)]] = E_{\theta \sim p}[\log \frac{p(\theta)}{q(\theta)}] \\ (\text{Gibb's}): \text{KL}(p||q) &\geq 0, \text{KL}(p||q)=0 \Leftrightarrow p = q \text{ a.s.}, \text{KL}(p||q) \neq \text{KL}(q||p) \\ H[X | Y = y] &\doteq E_{x \sim p(x|y)}[-\log p(x|y)], \\ H[X | Y] &\doteq E_{y \sim p(y)}[H[X | Y = y]] = E_{(x,y) \sim p(x,y)}[-\log p(x|y)] \\ H[X, Y] &\doteq E_{(x,y) \sim p(x,y)}[-\log p(x, y)] \end{aligned}$$

$$\begin{aligned} H[X, Y] &= H[Y] + H[X | Y] = H[X] + H[Y | X] = H[Y, X] \\ \text{Bayes' rule for entropy: } H[X | Y] &= H[Y | X] + H[X] - H[Y] \\ \text{"Information never hurts" principle: } H[X | Y] &\leq H[X] \\ 0 \leq I(X; Y) &\doteq H[X] - H[X | Y] = H[X] + [Y] - H[X, Y] = I(Y; X) \\ I(X; Y|Z) &\doteq H[X|Z] - H[X|Y, Z] = I(Y; X|Z) = I(X; Y, Z) - I(X; Z) \end{aligned}$$

0.1 Information theoretic quantities about Gaussians
Entropy $H[\mathcal{N}(\mu, \Sigma)] = \frac{1}{2} \log \det(2\pi e \Sigma) = \frac{1}{2} \log((2\pi e) \det(\Sigma))$

$$\begin{aligned} \text{KL } p &\doteq \mathcal{N}(\mu_p, \Sigma_p) \text{ and } q \doteq \mathcal{N}(\mu_q, \Sigma_q), \text{ then } \text{KL}(p||q) = \\ &\frac{1}{2} \left(\text{tr}(\Sigma_q^{-1}\Sigma_p) + (\mu_p - \mu_q)^T \Sigma_q^{-1}(\mu_p - \mu_q) - d + \log \frac{\det(\Sigma_q)}{\det(\Sigma_p)} \right) \end{aligned}$$

M.I. Let $X \sim \mathcal{N}(\mu, \Sigma)$ and $Y = X + \varepsilon$ where $\varepsilon \sim \mathcal{N}(0, \sigma_\varepsilon^2 I)$ then:
 $I(X; Y) = I(Y; X) = H[Y] - H[\varepsilon] = 1/2 \log(\det(\sigma_\varepsilon^{-2}\Sigma + I))$

1 Linear regression point estimates

$$y_i = w^T x_i + \varepsilon_i \text{ where } \varepsilon_i \sim \mathcal{N}(0, \sigma_n^2) \Leftrightarrow y_i | x_i, w \sim \mathcal{N}(w^T x_i, \sigma_n^2)$$

$$\hat{w}_{\text{ls}} \doteq \arg \min_{w \in \mathbb{R}^d} \sum_{i=1}^n (y_i - w^T x_i)^2 = \arg \min_{w \in \mathbb{R}^d} \|y - Xw\|_2^2$$

$$\hat{w}_{\text{MLE}} \doteq \arg \min_{\theta \in \Theta} -\sum_{i=1}^n \log p(y_i | x_i, \theta) = \hat{w}_{\text{LS}} = (X^T X)^{-1} X^T y$$

$$\hat{w}_{\text{ridge}} \doteq \arg \min_{w \in \mathbb{R}^d} \|y - Xw\|_2^2 + \lambda \|w\|_2^2 = (X^T X + \lambda I)^{-1} X^T y$$

$$\hat{w}_{\text{MAP}} \doteq \arg \max_{w \in \mathbb{R}^d} \log p(y_{1:n} | x_{1:n}, w) + \log p(w)$$

$$= \arg \min_{w \in \mathbb{R}^d} \|y - Xw\|_2^2 + \frac{\sigma_n^2}{\sigma_p^2} \|w\|_2^2 = \hat{w}_{\text{ridge}}$$

Given $\hat{w}$, $f^* \doteq \hat{w}^T x^*$ and $y^* \doteq f^* + \varepsilon$, $y^* \sim \mathcal{N}(f^*, \sigma_n^2 I)$.

2 Bayesian linear regression (probabilistic inference)

If $w \sim \mathcal{N}(\mu_p, \Sigma_p)$, $w | x_{1:n}, y_{1:n} \sim \mathcal{N}(\mu, \Sigma)$ where

$$\begin{aligned} \Sigma &\doteq (\sigma_n^{-2} X^T X + \Sigma_p^{-1})^{-1}, \mu \doteq (\Sigma_p^{-1} \mu_p + \sigma_n^{-2} X^T y) = \hat{w}_{\text{MAP}} \\ \text{by the properties Gaussians, closed form predictive posterior:} \\ w^T x^* &\sim f^* \sim \mathcal{N}(\mu^T x^*, x^{*T} \Sigma x^*), y^* \sim \mathcal{N}(\mu^T x^*, x^{*T} \Sigma x^* + \sigma_n^2 I) \end{aligned}$$

$$\begin{aligned} \textbf{3 Optimal action for given cost function } C(y, a): a^* \\ (y - a)^2: E_y[Y|x] &= \int y \cdot p(y|x) dy = \mu_{Y|X=x} \\ c_1 \max(y - a, 0) + c_2 \max(a - y, 0): &\mu_{Y|X=x} + \sigma_{Y|X} \Phi^{-1}(c_1/(c_1 + c_2)) \end{aligned}$$

$$\begin{aligned} \textbf{4 Function-space view and non-linear regression} \\ \text{Given design matrix } \Phi \in \mathbb{R}^{n \times e}, \text{ from weights prior we get:} \\ f | X \sim \mathcal{N}(\Phi E[w], \Phi \text{Var}[w] \Phi^T) &= \mathcal{N}(0, \sigma_p^2 \Phi \Phi^T) \text{ where} \\ \sigma_p^2 \Phi \Phi^T &\doteq K \in \mathbb{R}^{n \times n} \text{ is the kernel matrix } (n \times n \text{ instead of } e \times e). \end{aligned}$$

5 Gaussian Processes

$$\begin{aligned} f &\sim \mathcal{GP}(\mu, k) \text{ for any set } A \doteq \{x_1, \dots, x_m\} \subseteq \mathcal{X}, f_A \sim \mathcal{N}(\mu_A, K_{AA}) \\ f^* | x^* &\sim \mathcal{N}(\mu(x^*), k(x^*, x^*)), y^* | x^* \sim \mathcal{N}(\mu(x^*), k(x^*, x^*) + \sigma_n^2) \end{aligned}$$

$$\begin{aligned} \text{For } f &\sim \mathcal{GP}(\mu, k) \text{ can define } g(x) = f(x) - \mu(x), g \sim \mathcal{GP}(0, k). \\ \text{Can write joint prior distribution of } y_{1:n} \text{ and } f^* \text{ at a test point } x^* \\ \text{as } [y, f^*]^T | x^*, x_{1:n} &\sim \mathcal{N}(\tilde{\mu}, \tilde{K}), \text{ where} \end{aligned}$$

$$\tilde{\mu} \doteq \begin{bmatrix} \mu_A \\ \mu(x^*) \end{bmatrix} \tilde{K} \doteq \begin{bmatrix} K_{AA} + \sigma_n^2 I & k_{x^*, A} \\ k_{x^*, A}^T & k(x^*, x^*) \end{bmatrix} k_{x^*, A}^T \doteq [k(x, x_1), \dots, k(x, x_n)]$$

$$\begin{aligned} \text{Conditioning on } y_{1:n} \text{ yields: } f | x_{1:n}, y_{1:n} &\sim \mathcal{GP}(\mu', k') \text{ where} \\ \mu'(x) &\doteq \mu(x) + k_{x, A}^T (K_{AA} + \sigma_n^2 I)^{-1} (y_A - \mu_A) \\ k'(x, x') &\doteq k(x, x') - k_{x, A}^T (K_{AA} + \sigma_n^2 I)^{-1} k_{x', A} \end{aligned}$$

$$\text{Adding noise: } y^* | x^*, x_{1:n}, y_{1:n} \sim \mathcal{N}(\mu'(x^*), k'(x^*, x^*) + \sigma_n^2)$$

6 Kernel functions

$$\begin{aligned} \bullet \text{ Symmetric: } k(x, x') &= k(x', x) \text{ for any } x, x' \in \mathcal{X} \\ \bullet \text{ Positive semi-definite: } K_{AA} &\text{ is positive semi-def for any } A \subseteq \mathcal{X} \\ \text{Fact: a matrix is positive (semi-)definite} &\Leftrightarrow \lambda_i \geq (\geq) 0 \text{ and } \det(A) = \prod_i \lambda_i \Rightarrow \text{a matrix is positive (semi-)definite only if } \det(A) > (\geq) 0. \\ k(x, x') &\doteq k_1(x, x') + k_2(x, x') \text{ and } k(x, x') \doteq k_1(x, x') k_2(x, x') \\ k(x, x') &\doteq c \cdot k_1(x, x') \text{ for any } c > 0 \text{ "output scale"} \\ k(x, x') &\doteq f(k_1(x, x')) \text{ for } f = \exp \text{ or any pos. coeff. polynomial} \\ \bullet \text{ stationary: } k(x - x') &= k(x, x') \bullet \text{ isotropic: } k(\|x - x'\|) = k(x, x') \\ \text{Isotropy} &\Rightarrow \text{stationarity, in } \mathbb{R} \Leftrightarrow \text{due to symmetry of kernel.} \end{aligned}$$

6.1 Maximum marginal likelihood in GP regression

$$\begin{aligned} \hat{\theta} &= \arg \max_{\theta} -\frac{n}{2} \log 2\pi - \frac{1}{2} \log \det(K_y^{(\theta)}) - \frac{1}{2} y^T (K_y^{(\theta)})^{-1} y \\ \frac{\partial}{\partial \theta_j} \log p(y_{1:n} | x_{1:n}, \theta) &= \frac{1}{2} \text{tr} \left((\alpha \alpha^T - K_y^{(\theta)}) \frac{\partial K_y^{(\theta)}}{\partial \theta_j} \right) [\alpha \doteq K_y^{(\theta)} y] \end{aligned}$$

In general non-convex, might have multiple local optima.

7 Laplace Approximation

$$\begin{aligned} p(\theta | y) &\approx q(\theta) = \mathcal{N}(\hat{\theta}; \hat{\Lambda}^{-1}) \text{ where } \hat{\theta} = \arg \max p(\theta | y) \text{ and } \\ \hat{\Lambda} &\doteq -H_{\psi(\hat{\theta})} = -H_{\theta} \log p(\theta | x_{1:n}, y_{1:n})|_{\theta=\hat{\theta}} \text{ is obtained by} \end{aligned}$$

second-order Taylor of $\psi(\theta) \doteq \log p(\theta | x_{1:n}, y_{1:n})$ around $\hat{\theta}$.

7.0.1 Laplace for Bayesian Logistic Regression

$$\begin{aligned} \sigma(z) &\doteq 1/(1 + \exp(-z)) \in (0, 1), \text{ where } z = w^T x, \text{ i.e.} \\ p(y | x, w) &= \begin{cases} \sigma(w^T x) & \text{if } y=1 \\ 1 - \sigma(w^T x) & \text{if } y=-1 \end{cases} = \sigma(y \cdot w^T x) \end{aligned}$$

$$\begin{aligned} \text{Bayesian logistic regression} &= \text{BLR with } y | x, w \sim \text{Bern}(\sigma(w^T x)) \\ \hat{w} &= \arg \min_w \frac{1}{2\sigma_p^2} \|w\|_2^2 + \sum_{i=1}^n \log(1 + \exp(-y_i w^T x_i)) = \end{aligned}$$

$$\begin{aligned} \text{regularized logistic regression } (\log(1 + \exp(-y_i w^T x_i)) &= \text{logistic loss}) \Rightarrow \text{can use SGD: } \theta_{t+1} \leftarrow \theta_t - \eta_t \frac{1}{m} \sum_{i=1}^m \nabla_{\theta} l(\theta_t; x_i) \\ w &\leftarrow w(1 - 2\lambda \eta_t) + \eta_t y x \sigma(-y w^T x) \text{ as } \nabla_w l_{\log}(w^T x; y) = -y x \sigma(-y w^T x) \end{aligned}$$

$$\Lambda = -H_w \log p(w | x_{1:n}, y_{1:n})|_{w=\hat{w}} = \sigma_p^{-2} I + X^T \text{diag}_{i \in [n]} \{\pi_i(1 - \pi_i)\} X$$

where $\pi_i = \sigma(\hat{w}^T x_i) = P(Y = 1 | \hat{w}, x_i)$. Making predictions:

$$p(y^* | x^*, x_{1:n}, y_{1:n}) \approx \int p(y^* | x^*, w) q_{\lambda}(w) dw = \int \sigma(y^* f^*) \cdot$$

$\mathcal{N}(f^*; \hat{w}^T x^*, x^{*T} \Lambda^{-1} x^*) d f^*$ can be approximated by num. quadrature methods s.a. Gauss-Hermite quadrature.

8 Variational inference

$$\begin{aligned} q^* &\doteq \arg \min_{q \in \mathcal{Q}} \text{KL}(q || p) = \arg \min_{\lambda \in \Lambda} \text{KL}(q_{\lambda} || p) \\ \bullet \text{ forward KL}(p||q) &(\text{q conservative of uncertainty and "includes" } p) \\ \bullet \text{ reverse KL}(q||p) &(\text{q underestimates uncertainty, "excluding" } q) \end{aligned}$$

$$\begin{aligned} \text{In V.I: } \text{KL}(q_{\lambda} || p(\cdot | x_{1:n}, y_{1:n})) &= \log p(y_{1:n} | x_{1:n}) - \mathcal{L}(q_{\lambda}, p; \mathcal{D}_n) \\ \mathcal{L}(q_{\lambda}, p; \mathcal{D}_n) &= E_{\theta \sim q_{\lambda}} [\log p(\theta, y_{1:n} | x_{1:n})] + H(q_{\lambda}) \end{aligned}$$

$$\begin{aligned} &= E_{\theta \sim q_{\lambda}} [\log p(y_{1:n} | x_{1:n}, \theta)] - \text{KL}(q_{\lambda} || p(\theta)) \leftarrow \text{prior over } \theta \\ \text{Evidence} &= \text{const. w.r.t. } q \Rightarrow \min \text{KL} \equiv \max \text{ELBO.} \end{aligned}$$

$$\nabla_{\lambda} L(q_{\lambda}, p; \mathcal{D}_n) = \nabla_{\lambda} E_{\theta \sim q} [\log p(y_{1:n} | x_{1:n}, \theta)] - \nabla_{\lambda} \text{KL}(q || p(\cdot))$$

KL and its gradient can typically be computed exactly. For ∇_{λ} of exp. log-likelihood: **score gradients / reparametrization trick.**

8.1 Reparametrization trick

Given $\varepsilon \sim \phi$ independent of λ + diff.able and invertible function g parameterized by λ . We let $\theta \doteq g(\varepsilon; \lambda)$. Then $\varepsilon = g^{-1}(\theta; \lambda)$, and $\theta \sim q(\cdot | \lambda) \equiv q_{\lambda}(\theta) = \phi(\varepsilon) \cdot |\det(D_{\varepsilon} g(\varepsilon; \lambda))|^{-1}$

For "nice" function f holds $E_{\theta \sim q_{\lambda}} [f(\theta)] = E_{\varepsilon \sim \phi} [f(g(\varepsilon; \lambda))]$

As ϕ doesn't depend on λ : $\nabla_{\lambda} E_{\theta \sim q_{\lambda}} [f(\theta)] = E_{\varepsilon \sim \phi} [\nabla_{\lambda} f(g(\varepsilon; \lambda))]$

8.2 Reparametrization for Gaussians

$$\begin{aligned} q_{\lambda}(\theta) &\doteq \mathcal{N}(\theta; \mu, \Sigma) \Rightarrow \theta = g(\varepsilon; \lambda) \doteq C\varepsilon + \mu \text{ where } CC^T = \Sigma \\ \varepsilon &= g^{-1}(\theta; \lambda) = C^{-1}(\theta - \mu) \text{ and } \phi(\varepsilon) = q_{\lambda}(\theta) \cdot |\det(C)| \end{aligned}$$

8.3 Black Box Stochastic Variational Inference

$$\begin{aligned} \nabla_{\lambda} L(q_{\lambda}, p; \mathcal{D}_n) &= \nabla_{\lambda} E_{\theta \sim q_{\lambda}} [\log p(y_{1:n} | x_{1:n}, \theta)] - \nabla_{\lambda} \text{KL}(q_{\lambda} || p(\cdot)) \\ &\approx \frac{n}{m} \sum_{j=1}^m \nabla_{C, \mu} \log p(y_{i_j} | x_{i_j}, C\varepsilon^{(j)} + \mu) - \nabla_{C, \mu} \text{KL}(q_{C, \mu} || p(\cdot)) \end{aligned}$$

where $\varepsilon^{(j)} \stackrel{\text{iid}}{\sim} \mathcal{N}(0, I)$ and $i_j \stackrel{\text{iid}}{\sim} \text{Unif}([n])$.

9 Bayesian Deep Learning $\theta \sim \mathcal{N}(0, \sigma_p^2 I)$

Homoscedastic likelihood: $y|x, \theta \sim \mathcal{N}(f(x; \theta), \sigma_n^2)$, then

$$\hat{\theta}_{\text{MAP}} = \arg \max_{\theta} \log p(\theta) + \sum_{i=1}^n \log p(y_i | x_i, \theta)$$

$$= \arg \min_{\theta} \frac{1}{2\sigma_p^2} \|\theta\|_2^2 + \frac{1}{2\sigma_n^2} \sum_{i=1}^n (y_i - f(x_i; \theta))^2 = \text{Ridge MSE}$$

Heteroscedastic likelihood $y|x, \theta \sim \mathcal{N}(f_{\mu}(x; \theta), f_{\sigma^2}(x; \theta))$

$$\hat{\theta}_{\text{MAP}} = \arg \min_{\theta} \frac{1}{2\sigma_p^2} \|\theta\|_2^2 + \frac{1}{2} \sum_{i=1}^n \left[\log f_{\sigma^2}(x_i; \theta) + \frac{(y_i - f_{\mu}(x_i; \theta))^2}{f_{\sigma^2}(x_i; \theta)} \right]$$

can attribute some loss to large variance (aleatoric uncertainty). To model epistemic uncertainty too, perform probabilistic inference:

$$\begin{aligned} p(y^* | x^*, x_{1:n}, y_{1:n}) &= \int p(y^* | x^*, \theta) p(\theta | x_{1:n}, y_{1:n}) d\theta \\ p(\theta | x_{1:n}, y_{1:n}) &= (p(\theta) \prod_{i=1}^n p(y_i | x_i, \theta)) / \int p(\theta) \prod_{i=1}^n p(y_i | x_i, \theta) d\theta \end{aligned}$$

Compute variational posterior $p(\theta | x_{1:n}, y_{1:n}) \approx q_{\lambda} \doteq q(\cdot | \lambda)$ and $p(y^* | x^*, x_{1:n}, y_{1:n}) \approx \int p(y^* | x^*, \theta) q(\theta | \lambda) d\theta = E_{\theta \sim q_{\lambda}(\cdot)} [p(y^* | x^*, \theta)]$

$$\approx \frac{1}{m} \sum_{i=1}^m \mathcal{N}(y^*; f_{\mu}(x^*; \theta^{(i)}), f_{\sigma^2}(x^*; \theta^{(i)})) \text{ for } \theta^{(i)} \stackrel{\text{iid}}{\sim} q_{\lambda}$$

Can estimate: $E[y^* | x^*, x_{1:n}, y_{1:n}] \approx \frac{1}{m} \sum_{i=1}^m f_{\mu}(x^*; \theta^{(i)}) = \bar{\mu}(x^*)$

$$\text{Var}[y^* | x^*, x_{1:n}, y_{1:n}] = E_{\theta} [\text{Var}_{y^*}[y^* | x^*, \theta]] + \text{Var}_{\theta} [E_{y^*}[y^* | x^*, \theta]]$$

$$\approx \frac{1}{m} \sum_{i=1}^m f_{\sigma^2}(x^*, \theta^{(i)}) + \frac{1}{m} \sum_{i=1}^m (f_{\mu}(x^*, \theta^{(i)}) - \bar{\mu}(x^*))^2$$

9.0.1 Monte Carlo Dropout / Dropout regularization

Dropout = variational inference with variational posterior $q(\theta | \lambda) \doteq \prod_{j=1}^d q_j(\theta_j | \lambda_j)$ where $q_j(\theta_j | \lambda_j) \doteq p\delta_0(\theta_j) + (1 - p)\delta_{\lambda_j}(\theta_j)$.

$$p(y^* | x^*, x_{1:n}, y_{1:n}) \approx E_{\theta \sim q_{\lambda}} [p(y^* | x^*, \theta)] \approx \frac{1}{m} \sum_{i=1}^m p(y^* | x^*, \theta^{(i)})$$

where $\theta^{(i)}$ is a NN with weights λ , each set to 0 with probability p .

0.2 Stochastic weight averaging (SWA)
 $[\theta^{(1)}, \dots, \theta^{(T)}]$ T iter of SGD, approx posterior $q(\theta) \sim \mathcal{N}(\mu_{1:d}, I\sigma_{1:d}^2)$
 $\mu_i = 1/T \sum_{j=1}^T \theta_i^{(j)}$, and $\sigma_i^2 = 1/T \sum_{j=1}^T (\theta_i^{(j)} - \mu_i)^2$

9.0.3 Probabilistic ensembles
Sample m D_j uniformly from D , for each obtain MAP estimate $\theta^{(j)}$, approximate posterior: $p(y^*|x^*, x_{1:n}, y_{1:n}) \approx \sum_{i=1}^m p(y^*|x^*, \theta^{(i)})$

9.1 Calibration
Split $D_{\text{val}} = B_1 \cup \dots \cup B_m$ s.t. $\hat{p}_i \in [(j-1)/m, j/m] \forall i \in B_j$.
 $\text{Acc}_j = 1/|B_j| \sum_{i \in B_j} \mathbb{1}\{\hat{y}_i = y_i\}$, $\text{Conf}_j = 1/|B_j| \sum_{i \in B_j} \hat{p}_i$ OR $\rightarrow$
 $\text{freq}_j = 1/|B_j| \sum_{i \in B_j} \mathbb{1}\{y_i = 1\}$, $\text{conf}_j = 1/|B_j| \sum_{i \in B_j} \hat{p}(y_i = 1|x_i)$
Calibrated $\Leftrightarrow p(\hat{y} = y|\hat{p} = p) = p$, i.e. $\text{Acc}_j \approx \text{Conf}_j$ or $\text{freq}_j \approx \text{conf}_j$.
ECE: **weighted** (on $|B_j|$) avg of $|\text{Acc}(\text{freq}) - \text{C}(\text{conf})|$, MCE is max.

10 Active learning: find $S \subseteq D$ maximizing info gain
 $I(S) := I(f_S; y_S) = H[f_S] - H[f_S | y_S] = \frac{1}{2} \log \det(I + \sigma_n^{-2} K_S)$
Theorem: $I(S)$ is monotone submodular.
marginal gain: $\Delta_F(x | A) \doteq F(A \cup \{x\}) - F(A)$. Given $A \subseteq B \subseteq \mathcal{X}$
submodular: $\Delta_F(x | A) \geq \Delta_F(x | B)$ | monotone: $F(A) \leq F(B)$
Theorem: greedily maximizing monotone submodular F yields constant-factor approx.: $F(S_n) \geq (1 - \frac{1}{e}) \max_{S \subseteq \mathcal{X}, |S|=n} F(S)$.

10.1 Uncertainty sampling for $S_t = \{x_1, \dots, x_t\}$
 $x_{t+1} \doteq \arg \max \Delta_I(x | S_t) = \arg \max \frac{1}{2} \log(1 + \sigma_t^2(x)/\sigma_n^2)$
where $\sigma_t^2(x)$ is predicted variance of the GP conditioned on y_{S_t} .
Homoscedastic noise $\rightarrow \arg \max \sigma_t^2(x)$ “uncertainty sampling”.
Heteroscedastic noise $\rightarrow \arg \max \sigma_t^2(x)/\sigma_n^2(x)$.

10.2 Active learning for classification (BALD)
Uncertainty sampling: $x_{t+1} \doteq \arg \max_{x \in \mathcal{X}} H[y_x | x_{1:t}, y_{1:t}]$
Mutual information distinguishes between aleatoric and epistemic:
 $x_{t+1} \doteq \arg \max_{x \in \mathcal{X}} I(y_x; \theta | x, x_{1:t}, y_{1:t})$
 $= \arg \max_{x \in \mathcal{X}} H(y_x | x, x_{1:t}, y_{1:t}) - \mathbb{E}_{\theta|x_{1:t}, y_{1:t}} H(y_x | \theta, x)$

- looks for points where different models predict different labels
- penalizes points where models are not confident

= points x where the models *disagree*. Requires approx. posterior.

11 Bayesian optimization
Find $\arg \max_{x \in D} f^*(x)$ choosing $x_1, \dots, x_T \in D$ and observing $y_t = f^*(x_t) + \varepsilon_t$. **Cumulative regret for time horizon T :**
 $R_T \doteq \sum_{t=1}^T (\max_x f^*(x) - f^*(x_t)) = T \max_x f^*(x) - \sum_{t=1}^T f^*(x_t)$
GOAL: $\lim_{T \rightarrow \infty} R_T/T = 0 \Rightarrow \lim_{T \rightarrow \infty} \max_t f(x_t) = f^*(x^*)$

11.1 Upper confidence sampling (optimism)
Maximize upper confidence bound $x_t = \arg \max_{\mu_{t-1}} (x) + \beta_t \sigma_{t-1}(x)$.
Non-convex, can use Lipschitz-optimization (low-D) or gradient ascent (high-D). β_t scales confidence bounds to enclose f^* :
Theorem if $f^* \sim \mathcal{GP}$, and we choose β_t correctly, UCB yields
 $R_T/T = O^*(\sqrt{\gamma_T/T})$ where $\gamma_T = \max_{S \subseteq \mathcal{X}, |S| \leq T} I(f_S; y_S)$. This means that if γ_T is sublinear in T , we achieve sublinear regret:
Linear ker: $\gamma_T = O(d \log T)$, Gaussian kernel: $\gamma_T = O((\log T)^{d+1})$
Matérn kernel ($\nu > \frac{1}{2}$): $\gamma_T = O(T^{\frac{d}{2\nu+d}} (\log T)^{\frac{2\nu}{2\nu+d}})$, Independent kernel: γ_T **linear** in T . “Smoother” f = more diminishing returns.

Theorem (frequentist confidence intervals) assuming $f^* \in \mathcal{H}_k(\mathcal{X})$.
 $R_T/T = O^*(\sqrt{(\beta_T \gamma_T)/T})$ where $\beta_T = O(\|f^*\|_k + \gamma_T \log^3 T)$
 β_t depends on γ_T and on “complexity” of f^* (its norm in $\mathcal{H}_k(\mathcal{X})$).

11.2 Thompson sampling
 $x_{t+1} \in \arg \max_{x \in D} f(x)$, where $\tilde{f} \sim p(f | x_{1:t}, y_{1:t})$ sampled.

12 Diffusion Generative Models
Markov chain is a *stochastic process* with: • a prior $p(X_1)$

• transition probabilities $p(X_{t+1}|X_t)$ with **Markov property**
 $X_{t+1} \perp X_{1:t-1} | X_t \Leftrightarrow p(X_{1:T}) = p(X_1) \prod_{t=1}^T p(X_{t+1} | X_t)$
Under conditions like *ergodicity*, over time the probability of being in state x becomes constant and equals the stationary distribution
 $p_\infty: X_T \rightarrow X_\infty$, i.e. $\lim_{t \rightarrow \infty} p(X_t = x) = p_\infty(X' = x)$.

12.1 Forward (noising) process
 x_0 = data, we want x_T to be Gaussian noise. Forward process is:
 $q(x_{1:T}|x_0) = \prod_{t=1}^T q(x_t|x_{t-1})$, $q(x_t|x_{t-1}) = \mathcal{N}(x_t; \sqrt{\alpha_t}x_{t-1}, \beta_t I)$,
i.e. $x_t \leftarrow \sqrt{\alpha_t}x_{t-1} + \varepsilon_t$ where $\alpha_t = 1 - \beta_t$, $\varepsilon_t \sim \mathcal{N}(0, \beta_t I)$.
t-step marginal: $q(x_t|x_0) = \mathcal{N}(x_t; \sqrt{\bar{\alpha}_t}x_0, (1 - \bar{\alpha}_t)I)$
where $\bar{\alpha}_t = \prod_{s=1}^t \alpha_s$. Schedule $\beta_t \in (0, 1)$ to be monotonically increasing, implying $\bar{\alpha}_t \rightarrow 0$ and $X_T \rightarrow \mathcal{N}(0, I)$ for $T \rightarrow \infty$.

12.2 Backward (denoising) process
A parameterized MC from $p_\lambda(X_T) = \mathcal{N}(0, I)$ to data. Notation:
• prior $p_\lambda(x_{1:T})$ • likelihood $p_\lambda(x_0|x_{1:T})$ • posterior $p_\lambda(x_{1:T}|x_0)$
by Markov property: $p_\lambda(x_{0:T}) = p_\lambda(x_T) \prod_{t=1}^T p_\lambda(x_{t-1}|x_t)$
transition prob: $p_\lambda(x_{t-1}|x_t) = \mathcal{N}(x_{t-1}; \mu_\lambda(x_t, t), \Sigma_\lambda(x_t, t))$
Want to maximize data evidence $\log p_\lambda(x_0) = \log \int p_\lambda(x_{0:T}) dx_{1:T}$
ELBO: can be shown that for distributions $p_\lambda(x, z)$ and $q_\theta(x, z)$:
 $\log p_\lambda(x) = \mathcal{L}(\theta, \lambda; x) + \text{KL}(q_\theta(\cdot|x) \| p_\lambda(\cdot|x))$
 $\mathcal{L}(\theta, \lambda; x) = \mathbb{E}_{z \sim q_\theta(\cdot|x)} [\log p_\lambda(x, z) - \log q_\theta(z|x)]$
 $x \equiv x_0, z \equiv x_{1:T}$, p_λ is the *backward* m. $q_\theta \equiv q$ is the *forward*.
Then: $\mathcal{L}(\lambda; x_0) = \mathbb{E}_{x_{1:T} \sim q(\cdot|x_0)} [\log p_\lambda(x_{0:T}) - \log q(x_{1:T}|x_0)] =$
 $\text{const} - \sum_{t=2}^T \mathbb{E}_{x_{1:T} \sim q(\cdot|x_0)} \left[\log \frac{q(x_{t-1}|x_t, x_0)}{p_\lambda(x_{t-1}|x_t)} + \log p_\lambda(x_0|x_1) \right] =$
w.r.t. λ $c - \sum_{t=2}^T \mathbb{E}_{x_{1:T} \sim q(\cdot|x_0)} [\text{KL}(q(\cdot|x_t, x_0) \| p_\lambda(\cdot|x_t))] + \mathbb{E}_{x_1 \sim q(\cdot|x_0)} [\log p_\lambda(x_0|x_1)]$
Maxim. the ELBO $\equiv$ minim. terms L_t , i.e. forcing the backward process $p_\lambda(x_{t-1}|x_t)$ to match reverse forward process $q(\cdot|x_t, x_0)$.
Let $L_t(x_0) = \mathbb{E}_{x_t \sim q(\cdot|x_0)} [\text{KL}(q(\cdot|x_t, x_0) \| p_\lambda(\cdot|x_t))]$:
 $q(x_{t-1}|x_t, x_0) = \mathcal{N}(x_{t-1}; \tilde{\mu}_t(x_t, x_0), \tilde{\beta}_t(x_0, x_0)I)$ for some μ_t, β_t
 $p_\lambda(x_{t-1}|x_t) = \mathcal{N}(x_{t-1}; \mu_\lambda(x_t, t), \Sigma_\lambda(x_t, t))$ is also gaussian:
 $\text{KL}(\mathcal{N}(\mu_q, \sigma_q^2 I) \| \mathcal{N}(\mu_p, \sigma_p^2 I)) \propto 1/(2\sigma_p^2) \|\mu_p - \mu_q\|^2 + \text{const}$
Set constant variance schedule $\Sigma_\lambda(x_t, t) = \sigma_t^2 I$ and get:
 $L_t(x_0) = \mathbb{E}_{x_t \sim q(\cdot|x_0)} \left[\frac{1}{2\sigma_t^2} \|\mu_\lambda(x_t, t) - \tilde{\mu}_t(x_t, x_0)\|^2 + \text{const} \right]$, where
 $\mu_\lambda(x_t, t)$ is learned approximator of reverse forward process mean
 $\tilde{\mu}_t(x_t, x_0) = x_0(\sqrt{\alpha_t}\beta_t)/(1 - \bar{\alpha}_t) + x_t(\sqrt{\alpha_t}(1 - \bar{\alpha}_{t-1}))/ (1 - \bar{\alpha}_t)$
Non-stationary, learn noise instead, as $x_t = \sqrt{\alpha_t}x_0 + \sqrt{1 - \alpha_t}\varepsilon$
where $\varepsilon \sim \mathcal{N}(0, I)$. Then, can learn to predict $\varepsilon_\lambda(x_t, t)$ instead:
 $L_t(x_0) = \frac{\beta_t^2}{2\sigma_t^2 \alpha_t(1 - \bar{\alpha}_t)} \mathbb{E}_{\varepsilon \sim \mathcal{N}(0, I)} \|\varepsilon - \varepsilon_\lambda(\sqrt{\alpha_t}x_0 + \sqrt{1 - \alpha_t}t)\|^2$.
Can be optimized via SGD by computing the gradient w.r.t. λ .
Sampling: $x_{t-1} = 1/\sqrt{\alpha_t} (x_t - (1 - \alpha_t)/(\sqrt{1 - \bar{\alpha}_t}) \varepsilon_\lambda(x_t, t)) + \sigma_t z$

12.3 Uncertainty quantification
Instead of point estimate $\lambda = \text{MLE}(D)$ set: • prior $p(\lambda)$
• approx posterior $q(\lambda|D) \approx p(\lambda|D)$ • predictive posterior
 $p(x_0|D) \approx \int p(x_0|\lambda) q(\lambda|D) d\lambda$. Then filter hallucinations by rejecting points with high uncertainty $H(p(\cdot|D))$.

13 Markov Decision Processes
 r induces sequence $R_t \doteq r(X_t, A_t)$, we usually maximize infinite horizon discounted payoff from time t : $G_t \doteq \sum_{m=0}^{\infty} \gamma^m R_{t+m}$
 $\pi(a | x) \doteq \mathbb{P}(A_t = a | X_t = x)$ induces a MC with transitions:
 $p^\pi(x'|x) \doteq \mathbb{P}(X_{t+1}^\pi = x' | X_t^\pi = x) = \sum_{a \in A} \pi(a|x) p(x'|x, a)$

$J_t(\pi) = \mathbb{E}_\pi[G_t] = \mathbb{E}[r(X_t, \pi(X_t)) + \gamma r(X_{t+1}, \pi(X_{t+1})) + \dots]$
 $V_t^\pi(x) = J(\pi | X_t = x) = \mathbb{E}_\pi[G_t | X_t = x]$
 $Q_t^\pi(x, a) = J(\pi | X_0 = x, A_0 = a) = r(x, a) + \gamma \sum_{x'} p(x'|x, a) \cdot V_{t+1}^\pi(x')$
In stationary MDP, Bellman Expectation Equation holds:
 $V^\pi(x) = r(x, \pi(x)) + \gamma \mathbb{E}_{x' \sim p(\cdot|x, \pi(x))} [V^\pi(x')]$, $Q^\pi(x) = r(x, a) + \gamma \mathbb{E}_{x' \sim p(\cdot|x, a)} [V^\pi(x')]$

Thanks to Bellman: $V^\pi = r^\pi + \gamma T^\pi V^\pi \Rightarrow V^\pi = (I - \gamma T^\pi)^{-1} r^\pi$
Can also use contraction ($\exists k < 1$ s.t. $\|f(x) - f(y)\| \leq k\|x - y\|$)
 $B^\pi: \mathbb{R}^n \rightarrow \mathbb{R}^n$, $B^\pi v \doteq r^\pi + \gamma T^\pi v$ B^π , by Banach fixed point theorem, any contraction has a unique fixed point, which can be recovered via fixed point iteration exponentially fast.

13.1 Policy optimization in known MDP
Bellman's theorem: π^* optimal $\Leftrightarrow$ greedy wrt. its induced value f.:
 $\pi^* = \arg \max_{a \in A} [r(x, a) + \gamma \sum_{x'} p(x'|x, a) \cdot V^{\pi^*}(x')] = \arg \max_{a \in A} Q^{\pi^*}(x, a)$
 $V^{\pi^*}(x) = \max_a [r(x, a) + \gamma \sum_{x'} p(x'|x, a) \cdot V^{\pi^*}(x')] = \max_{a \in A} Q^{\pi^*}(x, a)$

Policy iteration: Start with random π . At each iteration compute V^π and set $\pi \leftarrow \pi_{V^\pi}$. Guaranteed to monotonically improve:
 $V^{\pi_{t+1}}(x) \geq V^{\pi_t}(x) \forall x$, and either $\exists x \in X | V^{\pi_{t+1}}(x) > V^{\pi_t}(x)$ or $V^{\pi_t} \equiv V^*$. Finite MDPs: exact optimum π^* in $O^*(\frac{n^2 m}{1-\gamma})$ iterations.

Value iteration: Start with $V_0(x) = \max_a r(x, a)$, at each iteration
 $V_t(x) = \max_a Q_t(x, a) = \max_a r(x, a) + \gamma \sum_{x'} P(x'|x, a) V_{t-1}(x')$
Stop when $\|V_t - V_{t-1}\|_\infty \leq \varepsilon$, then choose greedy policy w.r.t. V_t .
Converge polynomially only to ε -optimal policy, asymptotically to π^* as Bellman update $B^* V_t = V_{t+1}$ is contraction and V^{π^*} is fixed p.

- policy iteration: requires policy *evaluation* at every iteration
- value iteration: for every state x , loop over actions a and reachable states $x' \rightarrow O(n \cdot m \cdot s)$ where s is sparsity coefficient

14 Reinforcement learning (basic algos)
 R_{max} initialize $\hat{r}(x, a) = R_{\text{max}}$, $\hat{p}(x^*|x, a) = 1$ where x^* is a “fairytale state” s.t. $\forall A \hat{p}(x^*|x, a) = 1$ and $\hat{r}(x^*, a) = R_{\text{max}}$. Follow optimal policy for $\hat{r}$ and $\hat{p}$, after “enough” transitions, recompute π .
Lemma: every T timesteps, w.h.p. R_{max} either obtains near-optimal reward or visits at least one unknown state-action pair.
Theorem: with probability $\geq 1 - \delta$, R_{max} reaches an ε -optimal policy in step $\text{poly}(|X|, |A|, T, 1/\varepsilon, \log(1/\delta), R_{\text{max}})$ steps.

TD-Learning $\hat{V}^\pi(x) \leftarrow (1 - \alpha_t) \hat{V}^\pi(x) + \alpha_t (r + \gamma \hat{V}^\pi(x'))$
 α_t sf RM and all sa pairs chosen infinitely often $\Rightarrow$ converges to V^π .
Off-p val. est. $\hat{Q}^\pi(x, a) \leftarrow (1 - \alpha_t) \hat{Q}^\pi(x, a) + \alpha_t (r + \gamma \hat{Q}^\pi(x', \pi(x')))$
 α_t sf RM and all sa pairs chosen infinitely often $\Rightarrow$ converges to Q^π .
Q-learning $\hat{Q}^*(x, a) \leftarrow (1 - \alpha_t) \hat{Q}^*(x, a) + \alpha_t (r + \gamma \max_{a'} \hat{Q}^*(x', a'))$
 α_t sf RM and all sa pairs chosen infinitely often $\Rightarrow$ converges to Q^* .
Optimistic Q-learning initialize $\hat{Q}^*(x, a) = \frac{R_{\text{max}}}{1-\gamma} \prod_{t=1}^T (1 - \alpha_t)^{-1}$
With $\text{Pr} \geq 1 - \delta$, obtains an ε -optimal policy after a number of steps $\text{poly}(|X|, |A|, 1/\varepsilon, \log(1/\delta), R_{\text{max}})$. T_{init} has similar upper bound.

Score grad. $\nabla_\theta J^T(\theta) = \nabla_\theta \mathbb{E}_{\tau \sim \Pi_\theta} [g_{0:T}] = \mathbb{E}_{\tau \sim \Pi_\theta} [g_{0:T} \nabla_\theta \log \Pi_\theta(\tau)]$
W/ s.dependent baselines: $= \mathbb{E}_\tau [(g_{0:T} - b(\tau_{0:t-1})) \nabla_\theta \log \Pi_\theta(\tau)]$
ON/OFF policy Model based: • V/P iteration • **R_{max}** • Model
Predictive Control (RH/MPC) • Random Shooting Methods • PETS (Probabilistic Ensembles Trajectory Sampling) • **LAMBDA** • **H-UCRL** • PBVI/PBPI (Point Based V/P Iteration) **Model free:** • TD-Learning • (Optimistic) Q-Learning • DQN (Deep Q-Network)
• Double DQN • REINFORCE • Actor-Critic (A2C, A3C) • GAE/GAEC • TRPO (Trust Region Policy Optimization) • PPO (Proximal Policy Optimization) • DDPG (Deep Deterministic Policy Gradients) • TD3 (Twin Delayed DDPG) • SAC (Soft Actor-Critic)
• MPO (Maximum a posteriori Policy Optimisation) • SVG (Stochastic Value Gradients) • RLHF (RL from Human Feedback)