

Байесовский выбор моделей: Вариационный ЕМ-алгоритм (завершение). Гауссовские Процессы для учета эволюции модели.

Александр Адуенко

18е ноября 2025

Содержание предыдущих лекций

- Формула Байеса и формула полной вероятности;
- Определение априорных вероятностей и selection bias;
- (Множественное) тестирование гипотез
- Экспоненциальное семейства. Достаточные статистики.
- Наивный байесовский классификатор. Связь целевой функции и вероятностной модели.
- Линейная регрессия: связь МНК и w_{ML} , регуляризации и w_{MAP} .
- Свойство сопряженности априорного распределения правдоподобию.
- Прогноз для одиночной модели:

$$p(\mathbf{y}_{test} | \mathbf{X}_{test}, \mathbf{X}_{train}, \mathbf{y}_{train}) = \int p(\mathbf{y}_{test} | \mathbf{w}, \mathbf{X}_{test}) p(\mathbf{w} | \mathbf{X}_{train}, \mathbf{y}_{train}) d\mathbf{w}.$$

- Связь апостериорной вероятности модели и обоснованности
- Обоснованность: понимание и связь со статистической значимостью.
- Логистическая регрессия: проблемы ML-оценки w и связь априорного распределения с отбором признаков.
- ЕМ-алгоритм. Использование ЕМ-алгоритма для отбора признаков в байесовской линейной регрессии.
- Вариационный ЕМ-алгоритм. Смесь моделей логистической регрессии.

ЕМ-алгоритм: воспоминание

Пусть $\mathbf{D} = (\mathbf{X}, \mathbf{y})$ – наблюдаемые переменные, $\mathbf{Z}$ – скрытые переменные.
 $p(\mathbf{D}, \mathbf{Z}|\Theta) = p(\mathbf{D}|\mathbf{Z}, \Theta)p(\mathbf{Z}|\Theta)$.

Вопрос 1: как решить задачу $p(\mathbf{D}|\Theta) = \int p(\mathbf{D}, \mathbf{Z}|\Theta)d\mathbf{Z} \rightarrow \max_{\Theta}$?

ЕМ-алгоритм

$$\begin{aligned} \text{Введем } F(q, \Theta) &= - \int q(\mathbf{Z}) \log q(\mathbf{Z})d\mathbf{Z} + \int q(\mathbf{Z}) \log p(\mathbf{D}, \mathbf{Z}|\Theta)d\mathbf{Z} = \\ &= - \int q(\mathbf{Z}) \log q(\mathbf{Z})d\mathbf{Z} + \int q(\mathbf{Z}) \log p(\mathbf{Z}|\mathbf{D}, \Theta)d\mathbf{Z} + \int \log p(\mathbf{D}|\Theta)q(\mathbf{Z})d\mathbf{Z} = \\ &= \log p(\mathbf{D}|\Theta) - \int q(\mathbf{Z}) \log \frac{q(\mathbf{Z})}{p(\mathbf{Z}|\mathbf{D}, \Theta)}d\mathbf{Z} = \log p(\mathbf{D}|\Theta) - D_{\text{KL}}(q||p(\mathbf{Z}|\mathbf{D}, \Theta)). \end{aligned}$$

Идея 1: $p(\mathbf{D}|\Theta) \rightarrow \max_{\Theta}$ заменим на $F(q, \Theta) \rightarrow \max_{q, \Theta}$.

Идея 2: Пошагово оптимизируем по Θ и q , то есть

1 Е-шаг: $q^s = F(q, \Theta^{s-1}) \rightarrow \max_{q \in Q}$;

2 М-шаг: $\Theta^s = F(q^s, \Theta) \rightarrow \max_{\Theta}$.

Вопрос: Зачем $q \in Q$? Как Е-шаг был выполнен при максимизации обоснованности для модели линейной регрессии?

$$F(q, \Theta) = \int q(\mathbf{Z}) \log \frac{p(\mathbf{D}, \mathbf{Z}|\Theta)}{q(\mathbf{Z})} d\mathbf{Z} = \log p(\mathbf{D}|\Theta) - D_{\text{KL}}(q||p(\mathbf{Z}|\mathbf{D}, \Theta)).$$

$$\text{Е-шаг. } \int q(\mathbf{Z}_k) \log \frac{q(\mathbf{Z}_k)}{\frac{1}{C} e^{\mathbb{E}_{q \setminus k} \log p(\mathbf{D}, \mathbf{Z}|\Theta)}} d\mathbf{Z}_k \rightarrow \min_{q(\mathbf{Z}_k)}.$$

Полный алгоритм

Пошагово оптимизируем по Θ и $q(\mathbf{Z}_k)$, $k = 1, \dots, K$, то есть

1 Е-шаг: $\log q(\mathbf{Z}_k^s) \propto \mathbb{E}_{q \setminus k} \log p(\mathbf{D}, \mathbf{Z}|\Theta^{s-1})$;

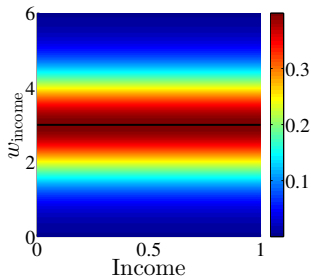
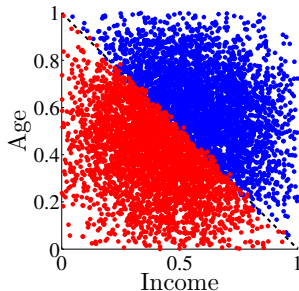
2 М-шаг: $\mathbb{E}_{q^s} \log p(\mathbf{D}, \mathbf{Z}|\Theta) \rightarrow \max_{\Theta}$.

Вопрос 1: Зачем нужна факторизация? Чем полученные итеративные формулы лучше формул полного EM-алгоритма?

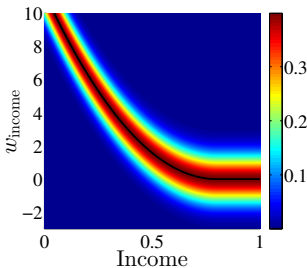
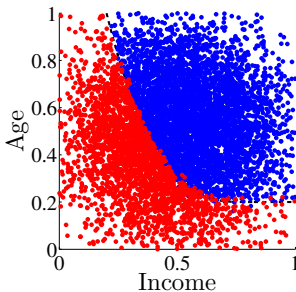
Вопрос 2: Как понять, что в конкретной задаче формулы Е и М-шагов выписаны верно?

Нарушение свойства $p(\mathbf{w}|\mathbf{x}_i) = p(\mathbf{w})$

Предполагаемый результат



Реальные данные



Вопрос: как можно учесть указанную нелинейность в модели?

Смесь моделей логистической регрессии

Вероятностная модель генерации данных

- Веса моделей в смеси π получены из априорного распределения $p(\pi|\mu)$;
- Векторы параметров моделей $\mathbf{w}_k$ получены из нормального распределения $p(\mathbf{w}_k|\mathbf{A}_k) = \mathcal{N}(\mathbf{w}_k|\mathbf{0}, \mathbf{A}_k^{-1})$, $k = 1, \dots, K$;
- Для каждого объекта $\mathbf{x}_i$ выбрана модель f_{k_i} , которой он описывается, причем $p(k_i = k) = \pi_k$;
- Для каждого объекта $\mathbf{x}_i$ класс y_i определен в соответствии с моделью f_{k_i} : $y_i \sim \text{Be}(\sigma(\mathbf{w}_{k_i}^\top \mathbf{x}_i))$.

Совместное правдоподобие модели

$$p(\mathbf{y}, \mathbf{w}_1, \dots, \mathbf{w}_K, \pi | \mathbf{X}, \mathbf{A}_1, \dots, \mathbf{A}_K, \mu) = p(\pi | \mu) \prod_{k=1}^K \mathcal{N}(\mathbf{w}_k | \mathbf{0}, \mathbf{A}_k^{-1}) \prod_{i=1}^m \left(\sum_{l=1}^K \pi_l \sigma(y_i \mathbf{w}_l^\top \mathbf{x}_i) \right).$$

Введем матрицу скрытых переменных $\mathbf{Z} = \|z_{ik}\|$, где $z_{ik} = 1 \iff k_i = k$.

$$p(\mathbf{y}, \mathbf{w}_1, \dots, \mathbf{w}_K, \pi, \mathbf{Z} | \mathbf{X}, \mathbf{A}_1, \dots, \mathbf{A}_K, \mu) = p(\pi | \mu) \prod_{k=1}^K \mathcal{N}(\mathbf{w}_k | \mathbf{0}, \mathbf{A}_k^{-1}) \prod_{i=1}^m \prod_{l=1}^K \left(\pi_l \sigma(y_i \mathbf{w}_l^\top \mathbf{x}_i) \right)^{z_{il}}.$$

Пусть $p(\boldsymbol{\pi}|\boldsymbol{\mu}) = \text{Dir}(\boldsymbol{\mu}) = \frac{\Gamma(\sum_k \mu_k)}{\prod_l \Gamma(\mu_l)} \prod_k \pi_k^{\mu_k-1}$.

$$p(\mathbf{y}, \mathbf{w}_1, \dots, \mathbf{w}_K, \boldsymbol{\pi}, \mathbf{Z}|\mathbf{X}, \mathbf{A}_1, \dots, \mathbf{A}_K, \boldsymbol{\mu}) \propto \prod_{k=1}^K \pi_k^{\mu_k-1} \prod_{k=1}^K \sqrt{\det \mathbf{A}_k} \exp(-\frac{1}{2} \mathbf{w}_k^\top \mathbf{A}_k \mathbf{w}_k) \prod_{i=1}^m \prod_{l=1}^K \left(\pi_l \sigma(y_i \mathbf{w}_l^\top \mathbf{x}_i) \right)^{z_{il}}.$$

$$(\boldsymbol{\pi}^*, \mathbf{w}_1^*, \dots, \mathbf{w}_K^*) = \arg \max_{\boldsymbol{\pi}, \mathbf{w}_1, \dots, \mathbf{w}_K} p(\mathbf{y}, \mathbf{w}_1, \dots, \mathbf{w}_K, \boldsymbol{\pi}|\mathbf{X}, \mathbf{A}_1, \dots, \mathbf{A}_K, \boldsymbol{\mu}).$$

E-шаг. $\log q(\mathbf{Z}) \propto \sum_{i=1}^m \sum_{l=1}^K z_{il} (\log \pi_l + \log \sigma(y_i \mathbf{w}_l^\top \mathbf{x}_i))$, откуда

$$\gamma_{ik} = p(z_{ik} = 1) \propto \pi_k \sigma(y_i \mathbf{w}_k^\top \mathbf{x}_i).$$

M-шаг. $E_q \log p(\mathbf{y}, \mathbf{w}_1, \dots, \mathbf{w}_K, \boldsymbol{\pi}, \mathbf{Z}|\mathbf{X}, \mathbf{A}_1, \dots, \mathbf{A}_K, \boldsymbol{\mu}) \rightarrow \max_{\boldsymbol{\pi}, \mathbf{w}_1, \dots, \mathbf{w}_K}$.

$$\mathbf{w}_k^* = \arg \max_{\mathbf{w}_k} \left[-\frac{1}{2} \mathbf{w}_k^\top \mathbf{A}_k \mathbf{w}_k + \sum_{i=1}^m \gamma_{ik} \log \sigma(y_i \mathbf{w}_k^\top \mathbf{x}_i) \right].$$

$$\boldsymbol{\pi}^* = \arg \max_{\boldsymbol{\pi}} \sum_{k=1}^K \log \pi_k \left(\underbrace{\sum_{i=1}^m \gamma_{ik}}_{\gamma_k} + \mu_k - 1 \right) \implies \pi_k \propto \max(0, \gamma_k + \mu_k - 1).$$

Получение апостериорного распределения

Вопрос: как получить $p(\mathbf{w}_1, \dots, \mathbf{w}_k, \boldsymbol{\pi} | \mathbf{X}, \mathbf{y}, \mathbf{A}_1, \dots, \mathbf{A}_K, \boldsymbol{\mu})$?

$$p(\mathbf{y}, \mathbf{w}_1, \dots, \mathbf{w}_K, \boldsymbol{\pi} | \mathbf{X}, \mathbf{A}_1, \dots, \mathbf{A}_K, \boldsymbol{\mu}) \propto \prod_{k=1}^K \pi_k^{\mu_k - 1} \prod_{k=1}^K \sqrt{\det \mathbf{A}_k} \exp(-\frac{1}{2} \mathbf{w}_k^\top \mathbf{A}_k \mathbf{w}_k) \prod_{i=1}^m \left(\sum_{l=1}^K \pi_l \sigma(y_i \mathbf{w}_l^\top \mathbf{x}_i) \right).$$

Идея: найдем $q(\mathbf{Z}, \mathbf{w}_1, \dots, \mathbf{w}_K, \boldsymbol{\pi}) = q(\mathbf{Z})q(\mathbf{w}_1, \dots, \mathbf{w}_K)q(\boldsymbol{\pi})$, наиболее близкое к $p(\mathbf{w}_1, \dots, \mathbf{w}_k, \boldsymbol{\pi}, \mathbf{Z} | \mathbf{X}, \mathbf{y})$.

$$\log q(\mathbf{Z}) \propto \mathbb{E}_{q_{\mathbf{Z}}} \log p(\mathbf{y}, \mathbf{w}_1, \dots, \mathbf{w}_K, \boldsymbol{\pi}, \mathbf{Z} | \mathbf{X}, \mathbf{A}_1, \dots, \mathbf{A}_K, \boldsymbol{\mu}) \propto$$

$$\sum_{i=1}^m \sum_{k=1}^K z_{ik} \left(\mathbb{E} \log \pi_k + \mathbb{E} \log \sigma(y_i \mathbf{w}_k^\top \mathbf{x}_i) \right)$$

$$\implies p(z_{ik} = 1) \propto \exp \left(\mathbb{E} \log \pi_k + \mathbb{E} \log \sigma(y_i \mathbf{w}_k^\top \mathbf{x}_i) \right).$$

$$\log q(\boldsymbol{\pi}) \propto \mathbb{E}_{q_{\mathbf{Z}}} \log p(\mathbf{y}, \mathbf{w}_1, \dots, \mathbf{w}_K, \boldsymbol{\pi}, \mathbf{Z} | \mathbf{X}, \mathbf{A}_1, \dots, \mathbf{A}_K, \boldsymbol{\mu}) \propto$$

$$\sum_{k=1}^K \log \pi_k \left(\mu_k - 1 + \sum_{i=1}^m \mathbb{E} z_{ik} \right) \implies \boldsymbol{\pi} \sim \text{Dir}(\boldsymbol{\mu} + \boldsymbol{\gamma}).$$

Получение апостериорного распределения (продолжение)

$$\log q(\mathbf{w}_1, \dots, \mathbf{w}_K) \propto \mathbb{E}_{q_{\setminus \mathbf{w}}} \log p(\mathbf{y}, \mathbf{w}_1, \dots, \mathbf{w}_K, \boldsymbol{\pi}, \mathbf{Z} | \mathbf{X}, \mathbf{A}_1, \dots, \mathbf{A}_K, \boldsymbol{\mu}) \propto \sum_{k=1}^K \left(-\frac{1}{2} \mathbf{w}_k^\top \mathbf{A}_k \mathbf{w}_k + \sum_{i=1}^m \mathbb{E} z_{ik} \log \sigma(y_i \mathbf{w}_k^\top \mathbf{x}_i) \right) = \sum_{k=1}^K f_k(\mathbf{w}_k).$$

Вопрос 1: Какую структуру имеет $q(\mathbf{w}_1, \dots, \mathbf{w}_K)$?

Вопрос 2: Какой вид имеет распределение $q(\mathbf{w}_k)$?

Варианты аппроксимации $q(\mathbf{w}_k)$:

- Аппроксимация Лапласа: $q(\mathbf{w}_k) \approx \mathcal{N}(\mathbf{w}_k^*, \boldsymbol{\Sigma}_k^{-1})$;
- Ищем $q(\mathbf{w}_k) = \mathcal{N}(\mathbf{m}_k, \boldsymbol{\Sigma}_k^{-1})$ такое, что $D_{KL}(q \| C_k f_k(\mathbf{w}_k)) \rightarrow \min_q$
 - Численно (если число признаков n невелико);
 - С помощью VLB для сигмоиды и соответствующей верхней оценки для $D_{KL}(q \| C_k f_k(\mathbf{w}_k))$.

Вопрос 3: Как определить $\mathbf{A}_1, \dots, \mathbf{A}_K, \boldsymbol{\mu}$?

Определение гиперпараметров смеси моделей

Максимизация обоснованности

$$p(\mathbf{y}|\mathbf{X}, \mathbf{A}_1, \dots, \mathbf{A}_K, \boldsymbol{\mu}) \rightarrow \max_{\mathbf{A}_1, \dots, \mathbf{A}_K, \{\boldsymbol{\mu}\}}$$

Идея: воспользуемся вариационным EM-алгоритмом.

Е-шаг: найдем $q(\mathbf{Z}, \mathbf{w}_1, \dots, \mathbf{w}_K, \boldsymbol{\pi}) = q(\mathbf{Z})q(\mathbf{w}_1, \dots, \mathbf{w}_K)q(\boldsymbol{\pi})$, наиболее близкое к $p(\mathbf{w}_1, \dots, \mathbf{w}_k, \boldsymbol{\pi}, \mathbf{Z}|\mathbf{X}, \mathbf{y})$.

$$\text{М-шаг: } E_q \log p(\mathbf{y}, \mathbf{w}_1, \dots, \mathbf{w}_K, \boldsymbol{\pi}, \mathbf{Z}|\mathbf{X}, \mathbf{A}_1, \dots, \mathbf{A}_K, \boldsymbol{\mu}) \rightarrow \max_{\mathbf{A}_1, \dots, \mathbf{A}_K}.$$

$$\sum_{k=1}^K \left[\sum_{j=1}^n \log \alpha_{kj}^j - \alpha_{kj}^j E w_{kj}^2 \right] \rightarrow \max_{\alpha_{11}, \dots, \alpha_{Kn}}, \text{ где } \mathbf{A}_k = \text{diag}(\alpha_{k1}, \dots, \alpha_{kn}).$$

$$\alpha_{kj} = \frac{1}{E w_{kj}^2}. \text{ Hint: } \alpha_{kj} = \frac{1 - \alpha_{kj}^{\text{old}} D w_{kj}}{E^2 w_{kj}}.$$

Вопрос: Как получить $E w_{kj}$ и $D w_{kj}$?

Замечание: можно сразу сделать использовать VLB для сигмоидной функции, чтобы получить нижнюю границу обоснованности

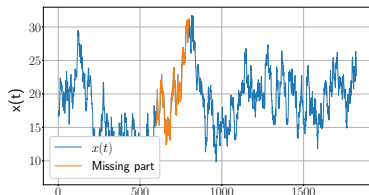
$$p(\mathbf{y}|\mathbf{X}, \mathbf{A}_1, \dots, \mathbf{A}_K, \boldsymbol{\mu}) \geq \tilde{p}(\mathbf{y}|\mathbf{X}, \mathbf{A}_1, \dots, \mathbf{A}_K, \boldsymbol{\mu})$$

и тогда на Е-шаге автоматически $q(\mathbf{w}_k)$ будет нормальным.

Учет эволюции модели во времени

Пусть у объектов есть еще метка времени, то есть наблюдаем $(\mathbf{x}_i, \mathbf{y}_i, t_i)$. Ранее имели модель $p(\mathbf{y}, \mathbf{w}|\mathbf{X}, \mathbf{A}) = p(\mathbf{y}|\mathbf{X}, \mathbf{w})p(\mathbf{w}|\mathbf{A})$, то есть зависимостью от t пренебрегали.

Вопрос 1: Как учесть наличие дополнительной информации?



Рассмотрим случайный процесс $x(t)$, $t \in T$.

$m_x(t) = \mathbb{E}x(t)$, $K_x(t, s) = \mathbb{E}x(t)x(s)$, $R_x(t, s) = \mathbb{E}\dot{x}(t)\dot{x}(s)$ – функция мат. ожидания, ковариационная и корреляционная функция.

Определение. С.п. называется **слабо стационарным**, если $m_x(t) \equiv m$, $R_x(t, s) = R_x(\tau = |t - s|)$.

Пример. Пусть $x(t)$ – температура в центре Кито.

Вопрос 2: Как восстановить пропущенные данные?

Гауссовские процессы

$x(t)$ – температура в центре Кито.

Идея: $GP(m_x(t), R_x(\tau))$, где $m_x(t) \equiv m$, $R_x(\tau) = \sigma^2 \exp(-\lambda|\tau|)$.

Рассмотрим $t_1, \dots, t_q$, тогда для GP имеем

$p(\mathbf{x}) = p(x(t_1), \dots, x(t_q)) = \mathcal{N}(\mathbf{m}, \mathbf{\Sigma})$, где

$\mathbf{m} = [m_x(t_1), \dots, m_x(t_q)]^\top$, $\mathbf{\Sigma} = \|\Sigma_{ij}\| = \|R_x(t_i - t_j)\|$.

Упражнение. $\mathbf{x} = [\mathbf{x}_1^\top, \mathbf{x}_2^\top]^\top \sim \mathcal{N}\left(\mathbf{x} | [\boldsymbol{\mu}_1^\top, \boldsymbol{\mu}_2^\top]^\top, \begin{pmatrix} \Sigma_{11} & \Sigma_{12} \\ \Sigma_{12}^\top & \Sigma_{22} \end{pmatrix}^{-1}\right)$.

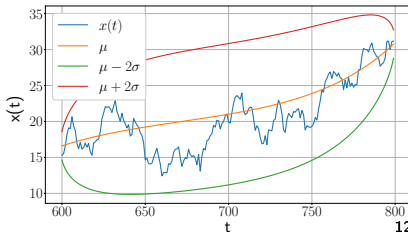
$\mathbf{x}_2 | \mathbf{x}_1 \sim \mathcal{N}(\mathbf{x}_2 | \boldsymbol{\mu}_2 - \Sigma_{22}^{-1} \Sigma_{12}(\mathbf{x}_1 - \boldsymbol{\mu}_1), \Sigma_{22}^{-1})$.

Вопрос 1: Что делать, если неизвестно m , где $\boldsymbol{\mu}_1 = m\mathbf{e}_1$, $\boldsymbol{\mu}_2 = m\mathbf{e}_2$?

Вопрос 2: Что делать, если неизвестны σ^2 и λ ?

Возможные модификации:

- Непостоянное $m_x(t)$;
- Введение разрывности $R_x(\tau) = \sigma^2(\exp(-\lambda|\tau|) + \kappa * [\tau = 0])$;
- Другая форма $R_x(\tau)$;
- $R_x(\tau) \rightarrow R_x(t_1, t_2)$.



Обозначим $r = \|x_1 - x_2\|$.

- $K(x_1, x_2) = \sigma^2 \exp(-\tau r^2)$ (RBF);
- $K(x_1, x_2) = \sigma^2 \exp(-\tau r)$ (Laplace);
- $K(x_1, x_2) = \sigma^2 \left(1 + \sqrt{3}r/l\right) \exp\left(-\sqrt{3}r/l\right)$ (Mattern 3/2);
- $K(x_1, x_2) = \sigma^2 \left(1 + \sqrt{5}r/l + \frac{5}{3}r^2/l^2\right) \exp\left(-\sqrt{5}r/l\right)$ (Mattern 5/2);
- $K(x_1, x_2) = \sigma^2 \exp\left(-2\frac{\sin^2(\pi r)}{l^2}\right)$ (Periodic);
- $K(x_1, x_2) = \sum_i \sigma_i^2 x_1^i x_2^i$ (Linear).

Вопрос 1: Как выбрать ядро? Какие функции задаёт каждое из вышеперечисленных?

Вопрос 2: Как получить ядро, отличное от вышеперечисленных?

Линейная регрессия с эволюцией во времени

Байесовская линейная регрессия

$(\mathbf{X}, \mathbf{y}) = \cup_{i=1}^m (\mathbf{x}_i, y_i)$ – выборка.

$$y_i = \mathbf{w}^\top \mathbf{x}_i + \varepsilon_i, \varepsilon_i \sim \mathcal{N}(\varepsilon_i | 0, \beta^{-1}), \mathbf{w} \sim \mathcal{N}(\mathbf{w} | \mathbf{0}, \mathbf{A}^{-1}).$$

$$p(\mathbf{y}, \mathbf{w} | \mathbf{X}, \mathbf{A}, \beta) = p(\mathbf{w} | \mathbf{A}) p(\mathbf{y} | \mathbf{X}, \mathbf{w}, \beta).$$

Байесовская линейная регрессия с эволюцией

$(\mathbf{X}, \mathbf{y}, \mathbf{t}) = \cup_{i=1}^m (\mathbf{x}_i, y_i, t_i)$ – выборка.

Для простоты считаем $t_1 < t_2 < \dots < t_m$.

$$y_i = \mathbf{w}_i^\top \mathbf{x}_i + \varepsilon_i, \varepsilon_i \sim \mathcal{N}(\varepsilon_i | 0, \beta^{-1}).$$

Введем матрицу $\mathbf{W} = [\mathbf{w}_1, \dots, \mathbf{w}_m]^\top = [\mathbf{v}_1, \dots, \mathbf{v}_n] \in \mathbb{R}^{m \times n}$.

Вопрос 1: как обучить модель, если у каждого объекта свой индивидуальный вектор параметров $\mathbf{w}_i$?

Идея: Априори предположим, что $\mathbf{v}_j$ получен как реализация GP $v_j(t)$.

$$\mathbf{v}_j = [v_j(t_1), \dots, v_j(t_m)]^\top, K_{v_j}(t_l, t_k) = \alpha_j^{-1} \exp(-\lambda |t_l - t_k|).$$

Тогда $p(\mathbf{w}_1, \dots, \mathbf{w}_m) = \prod_{j=1}^n \mathcal{N}(\mathbf{v}_j | \mathbf{0}, \alpha_j^{-1} \mathbf{K}^{-1})$, где

$$\mathbf{K} = \|\exp(-\lambda |t_l - t_k|)\|^{-1}.$$

Вопрос 2: что произойдет при $\lambda \rightarrow 0$?

Получение апостериорного распределения

Пусть дополнительно дана точка для прогноза $(\mathbf{x}_{m+1}, t_{m+1})$.

Найти: $p(\mathbf{w}_1, \dots, \mathbf{w}_m, \mathbf{w}_{m+1} | \mathbf{y}, \mathbf{X}, \mathbf{t}, \beta, \mathbf{A}, \lambda)$.

$$\log p(\mathbf{y}, \mathbf{w}_1, \dots, \mathbf{w}_m, \mathbf{w}_{m+1} | \mathbf{X}, \mathbf{t}, \beta, \mathbf{A}, \lambda) \propto$$
$$\frac{m}{2} \log \beta - \frac{\beta}{2} \sum_{i=1}^m (y_i - \mathbf{w}_i^\top \mathbf{x}_i)^2 + \sum_{j=1}^n \left[\frac{1}{2} \log \alpha_j + \frac{1}{2} \log \det \mathbf{K} - \frac{\alpha_j}{2} \mathbf{v}_j^\top \mathbf{K} \mathbf{v}_j \right].$$

$$\log p(\mathbf{w}_1, \dots, \mathbf{w}_m, \mathbf{w}_{m+1} | \mathbf{y}, \mathbf{X}, \mathbf{t}, \beta, \mathbf{A}, \lambda) \propto$$
$$-\frac{1}{2} \left[\sum_{j=1}^n \alpha_j \mathbf{v}_j^\top \mathbf{K} \mathbf{v}_j + \beta \sum_{i=1}^m \mathbf{w}_i^\top \mathbf{x}_i \mathbf{x}_i^\top \mathbf{w}_i - 2\beta \sum_{i=1}^m y_i \mathbf{w}_i^\top \mathbf{x}_i \right].$$

Введем $\mathbf{u} = [\mathbf{w}_1, \dots, \mathbf{w}_{m+1}]^\top \in \mathbb{R}^{(m+1)n}$.

$$\mathbf{K}_1 = \begin{pmatrix} \beta \mathbf{x}_1 \mathbf{x}_1^\top & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} \\ \mathbf{0} & \beta \mathbf{x}_2 \mathbf{x}_2^\top & \dots & \mathbf{0} & \mathbf{0} \\ \vdots & & & & \\ \mathbf{0} & \mathbf{0} & \dots & \beta \mathbf{x}_m \mathbf{x}_m^\top & \mathbf{0} \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \mathbf{0} & \mathbf{0} \end{pmatrix}, \quad \mathbf{m} = \beta \begin{pmatrix} y_1 \mathbf{x}_1 \\ y_2 \mathbf{x}_2 \\ \vdots \\ y_m \mathbf{x}_m \\ \mathbf{0} \end{pmatrix}.$$

Получение апостериорного распределения

Введем $\mathbf{u} = [\mathbf{w}_1, \dots, \mathbf{w}_{m+1}]^\top \in \mathbb{R}^{(m+1)n}$.

$$\log p(\mathbf{w}_1, \dots, \mathbf{w}_m, \mathbf{w}_{m+1} | \mathbf{y}, \mathbf{X}, \mathbf{t}, \beta, \mathbf{A}, \lambda) \propto -\frac{1}{2} \left[\mathbf{u}^\top \Sigma^{-1} \mathbf{u} - 2\mathbf{u}^\top \mathbf{m} \right] =$$

$$-\frac{1}{2} \left[\sum_{j=1}^n \alpha_j \mathbf{v}_j^\top \mathbf{K} \mathbf{v}_j + \beta \sum_{i=1}^m \mathbf{w}_i^\top \mathbf{x}_i \mathbf{x}_i^\top \mathbf{w}_i - 2\beta \sum_{i=1}^m y_i \mathbf{w}_i^\top \mathbf{x}_i \right].$$

$$\mathbf{K}_1 = \begin{pmatrix} \beta \mathbf{x}_1 \mathbf{x}_1^\top & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} \\ \mathbf{0} & \beta \mathbf{x}_2 \mathbf{x}_2^\top & \dots & \mathbf{0} & \mathbf{0} \\ \vdots & & & & \\ \mathbf{0} & \mathbf{0} & \dots & \beta \mathbf{x}_m \mathbf{x}_m^\top & \mathbf{0} \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \mathbf{0} & \mathbf{0} \end{pmatrix}, \mathbf{m} = \beta \begin{pmatrix} y_1 \mathbf{x}_1 \\ y_2 \mathbf{x}_2 \\ \vdots \\ y_m \mathbf{x}_m \\ \mathbf{0} \end{pmatrix},$$

$$\mathbf{K}_2 = \begin{pmatrix} \mathbf{A}K_{11} & \mathbf{A}K_{12} & \dots & \mathbf{A}K_{1,m+1} \\ \mathbf{A}K_{21} & \mathbf{A}K_{22} & \dots & \mathbf{A}K_{2,m+1} \\ \vdots & & & \\ \mathbf{A}K_{m+1,1} & \mathbf{A}K_{m+1,2} & \dots & \mathbf{A}K_{m+1,m+1} \end{pmatrix}.$$

Отсюда $\mathbf{u} \sim \mathcal{N}(\mathbf{u} | (\mathbf{K}_1 + \mathbf{K}_2)^{-1} \mathbf{m}, (\mathbf{K}_1 + \mathbf{K}_2)^{-1})$.

Отбор признаков и подбор ковариационной функции

Вопрос: как определить $\mathbf{A}$, β , λ ?

Рассмотрим задачу $p(\mathbf{y}|\mathbf{X}, \mathbf{t}, \beta, \mathbf{A}, \lambda) \rightarrow \max_{\beta, \mathbf{A}, \lambda}$.

Рассмотрим $\mathbf{Z} = (\mathbf{w}_1, \dots, \mathbf{w}_{m+1})$ и воспользуемся EM-алгоритмом.

Е-шаг. $q(\mathbf{Z}) = p(\mathbf{w}_1, \dots, \mathbf{w}_m, \mathbf{w}_{m+1}|\mathbf{y}, \mathbf{X}, \mathbf{t}, \beta, \mathbf{A}, \lambda)$.

М-шаг. $E_q \log p(\mathbf{y}, \mathbf{w}_1, \dots, \mathbf{w}_m, \mathbf{w}_{m+1}|\mathbf{X}, \mathbf{t}, \beta, \mathbf{A}, \lambda) \rightarrow \max_{\beta, \mathbf{A}, \lambda}$.

$$\frac{m}{2} \log \beta - \frac{\beta}{2} \sum_{i=1}^m E(y_i - \mathbf{w}_i^\top \mathbf{x}_i)^2 + \frac{1}{2} \sum_{j=1}^n \log \alpha_j + \frac{n}{2} \log \det \mathbf{K} -$$

$$\frac{1}{2} \sum_{j=1}^n \alpha_j E \mathbf{v}_j^\top \mathbf{K} \mathbf{v}_j \rightarrow \max_{\beta, \mathbf{A}, \lambda}.$$

$$\beta^{-1} = \frac{1}{m} \sum_{i=1}^m E(y_i - \mathbf{w}_i^\top \mathbf{x}_i)^2; \alpha_j = \frac{1}{E \mathbf{v}_j^\top \mathbf{K} \mathbf{v}_j} = \frac{1}{\text{tr}(\mathbf{K} E \mathbf{v}_j \mathbf{v}_j^\top)}.$$

Hint: $\alpha_j^{\text{new}} = \frac{1 - \alpha_j^{\text{old}} \text{tr}(\mathbf{K} E \hat{\mathbf{v}}_j \hat{\mathbf{v}}_j^\top)}{\text{tr}(\mathbf{K} (E \mathbf{v}_j)(E \mathbf{v}_j^\top))}.$

Отбор признаков и подбор ковариационной функции

M-шаг. $E_q \log p(\mathbf{y}, \mathbf{w}_1, \dots, \mathbf{w}_m, \mathbf{w}_{m+1} | \mathbf{X}, \mathbf{t}, \beta, \mathbf{A}, \lambda) \rightarrow \max_{\beta, \mathbf{A}, \lambda}.$

$$\frac{m}{2} \log \beta - \frac{\beta}{2} \sum_{i=1}^m E(y_i - \mathbf{w}_i^\top \mathbf{x}_i)^2 + \frac{1}{2} \sum_{j=1}^n \log \alpha_j + \frac{n}{2} \log \det \mathbf{K} -$$

$$\frac{1}{2} \sum_{j=1}^n \alpha_j E \mathbf{v}_j^\top \mathbf{K} \mathbf{v}_j \rightarrow \max_{\beta, \mathbf{A}, \lambda}.$$

$$\beta^{-1} = \frac{1}{m} \sum_{i=1}^m E(y_i - \mathbf{w}_i^\top \mathbf{x}_i)^2; \alpha_j = \frac{1}{E \mathbf{v}_j^\top \mathbf{K} \mathbf{v}_j} = \frac{1}{\text{tr}(\mathbf{K} E \mathbf{v}_j \mathbf{v}_j^\top)}.$$

Hint: $\alpha_j^{\text{new}} = \frac{1 - \alpha_j^{\text{old}} \text{tr}(\mathbf{K} E \hat{\mathbf{v}}_j \hat{\mathbf{v}}_j^\top)}{\text{tr}(\mathbf{K} (E \mathbf{v}_j) (E \mathbf{v}_j)^\top)}.$

$$\mathbf{B} = \sum_{j=1}^n \alpha_j E \mathbf{v}_j \mathbf{v}_j^\top, \text{ тогда } f(\lambda) = \frac{n}{2} \log \det \mathbf{K} - \frac{1}{2} \text{tr}(\mathbf{K} \mathbf{B}) \rightarrow \max_{\lambda}.$$

$$2 \frac{df}{d\lambda} = n \text{tr} \left(\frac{d\mathbf{K}}{d\lambda} \mathbf{K}^{-1} \right) - \text{tr} \left(\frac{d\mathbf{K}}{d\lambda} \mathbf{B} \right) = 0.$$

Вопрос: как получить оптимальное λ ?

- 1 Bishop, Christopher M. "Pattern recognition and machine learning". Springer, New York (2006). Pp. 78-88, 303-320.
- 2 MacKay, David JC. Bayesian methods for adaptive models. Diss. California Institute of Technology, 1992.
- 3 MacKay, David JC. "The evidence framework applied to classification networks." Neural computation 4.5 (1992): 720-736.
- 4 Gelman, Andrew, et al. Bayesian data analysis, 3rd edition. Chapman and Hall/CRC, 2013.
- 5 Дрейпер, Норман Р. Прикладной регрессионный анализ. Рипол Классик, 2007.