

4. Parameterschätzung

4. Parameterschätzung

Aufgabenstellung der Parameterschätzung

Gegeben:

- **Mathematisches Modell** einer WV $p(\mathbf{m}|\theta)$ für eine beobachtbare Größe $\mathbf{m}$
- **Stichprobe** $D = \{\mathbf{m}_1, \dots, \mathbf{m}_N\}$ für festen, aber unbekannten Parametervektor θ .

Gesucht:

- „**Wahrer**“ **Wert** des Parametervektors θ .

Ansatz:

- **Schätzfunktion** $\hat{\theta}(D)$ (Schätzer), welche die Stichprobe D statistisch auswertet und in gewissem Sinne dem „wahren“ Wert θ möglichst nahe kommt.

4. Parameterschätzung

Methodiken der Parameterschätzung

Likelihoodmethodik (Klassische Statistik)

- θ wird als unbekannte, konstante, **nicht-stochastische Größe** angesehen.
- **Maximum-Likelihood-Schätzung**: Man wählt $\hat{\theta}(D)$ so, dass die gegebenen Beobachtungen $D = \{\mathbf{m}_1, \dots, \mathbf{m}_N\}$ maximal wahrscheinlich werden.

Bayes'sche Methodik (Bayes'sche Statistik)

- θ wird als **Zufallsvariable** angesehen und über eine WV beschrieben.
- Z.B.: **Maximum-A Posteriori-Schätzung**: Man wählt $\hat{\theta}(D)$ als den Wert θ mit der höchsten A Posteriori Wahrscheinlichkeit nach Beobachtung von D .

4. Parameterschätzung – Einschub: Bedeutung von Wahrscheinlichkeit

Axiome nach Kolmogorov:

A1: Nichtnegativität $\Pr(A) \geq 0$

A2: Normierung $\Pr(\text{Sicheres Ereignis}) = 1$

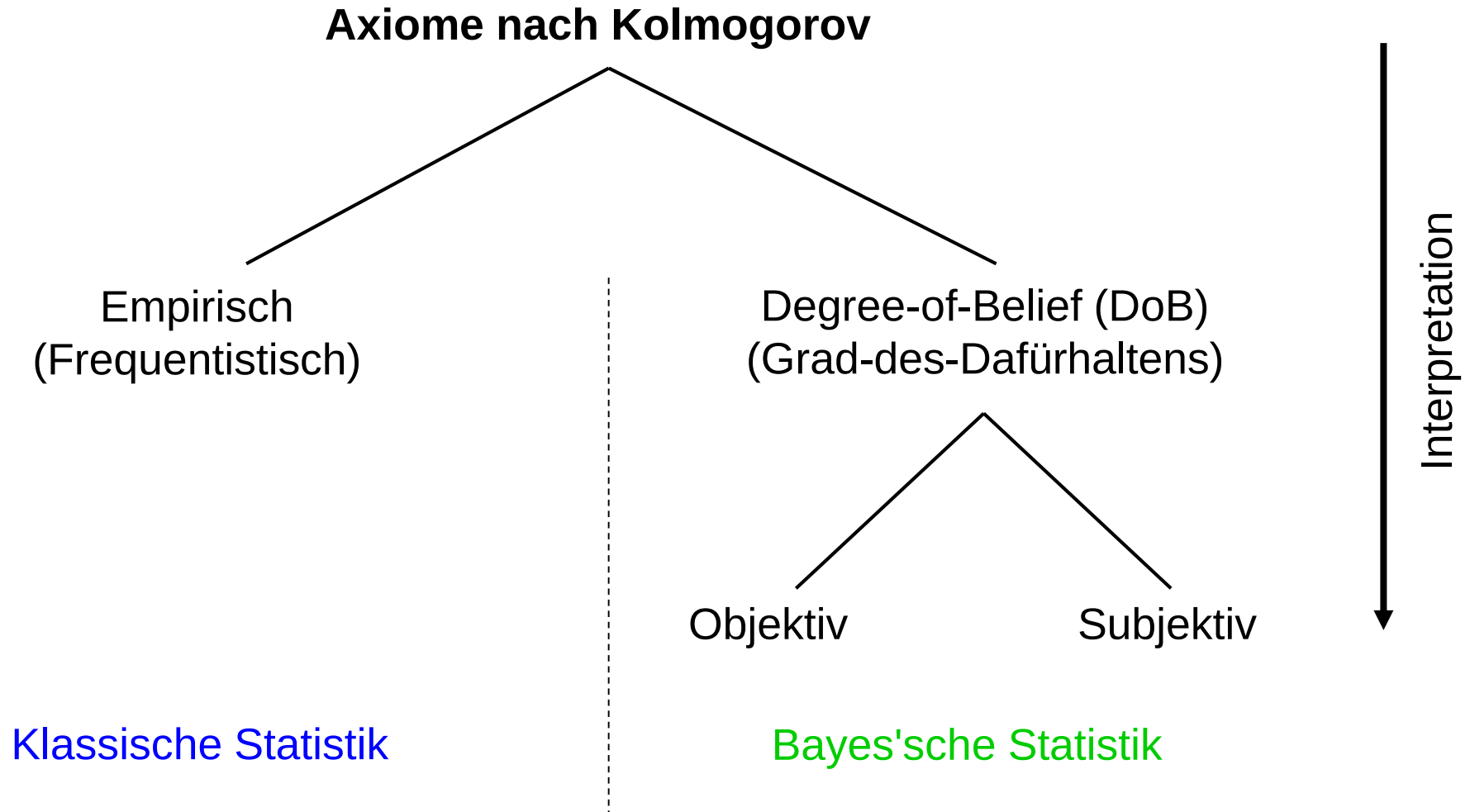
A3: Additivität $A \cap B = \emptyset \Rightarrow \Pr(A \cup B) = \Pr(A) + \Pr(B)$

Bedingte Wahrscheinlichkeit: $\Pr(A|B) := \frac{\Pr(A \cap B)}{\Pr(B)}$

Die Kolmogorovschen Axiome legen fest, **wie** mit Wahrscheinlichkeiten gerechnet wird (Syntax), machen aber keine Aussage darüber, **was** Wahrscheinlichkeiten bedeuten (Semantik).

4. Parameterschätzung – Einschub: Bedeutung von Wahrscheinlichkeit

Mit den Axiomen verträgliche Interpretationen bezüglich der Bedeutung von Wahrscheinlichkeit:



4. Parameterschätzung

Parameterschätzung im Kontext der Mustererkennung

- Schätzung der Parameter der klassenspezifischen WVen $p(\mathbf{m}|\theta_i, \omega_i)$ $i = 1, \dots, c$ aus den Daten D .

Annahmen:

- Struktur der klassenspezifischen WVen ist gegeben. Nur die Parameter θ_i sind unbekannt.
- Überwachtes Lernen: die Klassenzugehörigkeiten der einzelnen Stichproben in D sind bekannt. $\rightarrow$ Partitionierung von D in $D_1, \dots, D_c$ mit den zur jeweiligen Klasse gehörenden Stichproben.
- Die Stichproben in D_j haben keinen Einfluss auf $p(\mathbf{m}|\theta_i, \omega_i)$ $i \neq j$.

Konsequenz:

- Schätzung der Parameter der klassenspezifischen WVen zerfällt in c separate Parameterschätzungsaufgaben, die unabhängig voneinander behandelt werden können. $\rightarrow$ Klassenunterscheidende Kennzeichnung kann in diesem Kapitel weitgehend unterdrückt werden.

4. Parameterschätzung

Eigenschaften von Schätzern: Erwartungstreue

Klassische Statistik

Erwartungstreue (unbiased estimator): $E\{\hat{\theta}\} = \theta$

$$E\{\hat{\theta}\} = \int_{\Theta} \hat{\theta} p(\hat{\theta}) d\hat{\theta} = \int_{\mathbb{M}} \hat{\theta}(\mathbf{m}) p(\mathbf{m} | \theta) d\mathbf{m}$$

Bias (systematischer Fehler): $b(\theta) = E\{\hat{\theta}\} - \theta$

Bemerkung: Simple Korrektur durch Subtraktion nur möglich, falls b nicht vom Parameter θ abhängt.

Bayes'sche Statistik

Erwartungstreue (unbiased estimator): $E\{\hat{\theta}\} = E\{\theta\}$

$$E\{\hat{\theta}\} = \int_{\Theta} \hat{\theta} p(\hat{\theta}) d\hat{\theta} = \int_{\Theta} \int_{\mathbb{M}} \hat{\theta}(\mathbf{m}) p(\mathbf{m}, \theta) d\mathbf{m} d\theta$$

Bias: $b = E\{\hat{\theta}\} - E\{\theta\}$

4. Parameterschätzung

Eigenschaften von Schätzern: Varianz

Varianz (stochastischer Fehler): $\text{Var}\{\hat{\boldsymbol{\theta}}\} = \text{E}\{(\hat{\boldsymbol{\theta}} - \text{E}\{\hat{\boldsymbol{\theta}}\})^2\} = \text{E}\{\hat{\boldsymbol{\theta}}^2\} - \text{E}\{\hat{\boldsymbol{\theta}}\}^2$

Schätzverfahren mit minimaler Varianz werden als **effiziente (wirksame) Schätzverfahren** bezeichnet.

Cramer-Rao-Schranke (CRB) für erwartungstreue Schätzer (klassische Statistik) für einen skalaren Parameter:

$D = \{\mathbf{m}_1, \dots, \mathbf{m}_N\}$, Beobachtungen $\mathbf{m}_i$ seien stochastisch unabhängig.

$$\text{Var}\{\hat{\theta}(D)\} \geq \frac{1}{N \text{E}\left\{\left(\frac{\partial \ln p(\mathbf{m}|\theta)}{\partial \theta}\right)^2\right\}} = \frac{1}{N \int_{\mathbf{M}} \left(\frac{\partial \ln p(\mathbf{m}|\theta)}{\partial \theta}\right)^2 p(\mathbf{m} | \theta) d\mathbf{m}}$$

Details und Verallgemeinerungen siehe z.B.: K. Kroschel, „Statistische Informationstechnik“, Springer 2004

- CRB ist eine **untere Schranke** für die **Varianz aller erwartungstreuen Schätzer**.
- Effiziente Schätzverfahren erreichen die CRB.

4. Parameterschätzung

Beweisskizze CRB: $\mathbf{m} \in \mathbb{R}^d, \theta \in \mathbb{R}, N = 1$

Voraussetzungen: $\frac{\partial}{\partial \theta} \int p(\mathbf{m} | \theta) d\mathbf{m} = \int \frac{\partial p(\mathbf{m} | \theta)}{\partial \theta} d\mathbf{m}$

$\hat{\theta}(\mathbf{m})$ ist erwartungstreu: $E\{\hat{\theta}(\mathbf{m})\} = \theta$

$$\frac{\partial}{\partial \theta} \int (\hat{\theta}(\mathbf{m}) - \theta) p(\mathbf{m} | \theta) d\mathbf{m} = \int \frac{\partial}{\partial \theta} (\hat{\theta}(\mathbf{m}) - \theta) p(\mathbf{m} | \theta) d\mathbf{m} = 0$$

mit $\frac{\partial \hat{\theta}(\mathbf{m})}{\partial \theta} = 0$ folgt:

$$-\int p(\mathbf{m} | \theta) d\mathbf{m} + \int (\hat{\theta}(\mathbf{m}) - \theta) \frac{\partial p(\mathbf{m} | \theta)}{\partial \theta} d\mathbf{m} = 0$$

mit $\frac{\partial p(\mathbf{m} | \theta)}{\partial \theta} = \frac{\partial \ln p(\mathbf{m} | \theta)}{\partial \theta} p(\mathbf{m} | \theta)$ folgt:

$$\int (\hat{\theta}(\mathbf{m}) - \theta) p(\mathbf{m} | \theta) \frac{\partial \ln p(\mathbf{m} | \theta)}{\partial \theta} d\mathbf{m} = 1$$

4. Parameterschätzung

$$\int (\hat{\theta}(\mathbf{m}) - \theta) p(\mathbf{m} | \theta) \frac{\partial \ln p(\mathbf{m} | \theta)}{\partial \theta} d\mathbf{m} = 0$$

Mit der Schwarzschen Ungleichung: $\int x^2(\mathbf{m}) d\mathbf{m} \int y^2(\mathbf{m}) d\mathbf{m} \geq \left(\int x(\mathbf{m}) y(\mathbf{m}) d\mathbf{m} \right)^2$

und $\left(\int (\hat{\theta}(\mathbf{m}) - \theta) \sqrt{p(\mathbf{m} | \theta)} \sqrt{p(\mathbf{m} | \theta)} \frac{\partial \ln p(\mathbf{m} | \theta)}{\partial \theta} d\mathbf{m} \right)^2 = 1$ folgt:

$$\underbrace{\int (\hat{\theta}(\mathbf{m}) - \theta)^2 p(\mathbf{m} | \theta) d\mathbf{m}}_{\text{Var}\{\hat{\theta}(\mathbf{m})\}} \underbrace{\int \left(\frac{\partial \ln p(\mathbf{m} | \theta)}{\partial \theta} \right)^2 p(\mathbf{m} | \theta) d\mathbf{m}}_{\text{E}\left\{ \left(\frac{\partial \ln p(\mathbf{m} | \theta)}{\partial \theta} \right)^2 \right\} : \text{Fisher-Information}} \geq 1$$

→ Cramer-Rao-Schranke

4. Parameterschätzung

Plausibilität: Wie stark ändert sich $\mathbf{m}$ mit θ ?

$$\begin{aligned} \mathbf{m} \sim p(\mathbf{m} | \theta) &\xrightarrow{\frac{\partial}{\partial \theta}} \frac{\partial p(\mathbf{m} | \theta)}{\partial \theta} \xrightarrow{\cdot \frac{1}{p(\mathbf{m} | \theta)}} \frac{\partial p}{\partial \theta} \frac{1}{p} = \frac{\partial \ln p(\mathbf{m} | \theta)}{\partial \theta} \\ &\xrightarrow{(\)^2} \left(\frac{\partial \ln p(\mathbf{m} | \theta)}{\partial \theta} \right)^2 \xrightarrow{\mathbb{E}\{ \}} \mathbb{E} \left\{ \left(\frac{\partial \ln p(\mathbf{m} | \theta)}{\partial \theta} \right)^2 \right\} = J(\theta) \end{aligned}$$

Quantifizierung der Empfindlichkeit von $\mathbf{m}$ bezüglich Veränderungen von θ .

4. Parameterschätzung

Eigenschaften von Schätzern: Konsistenz

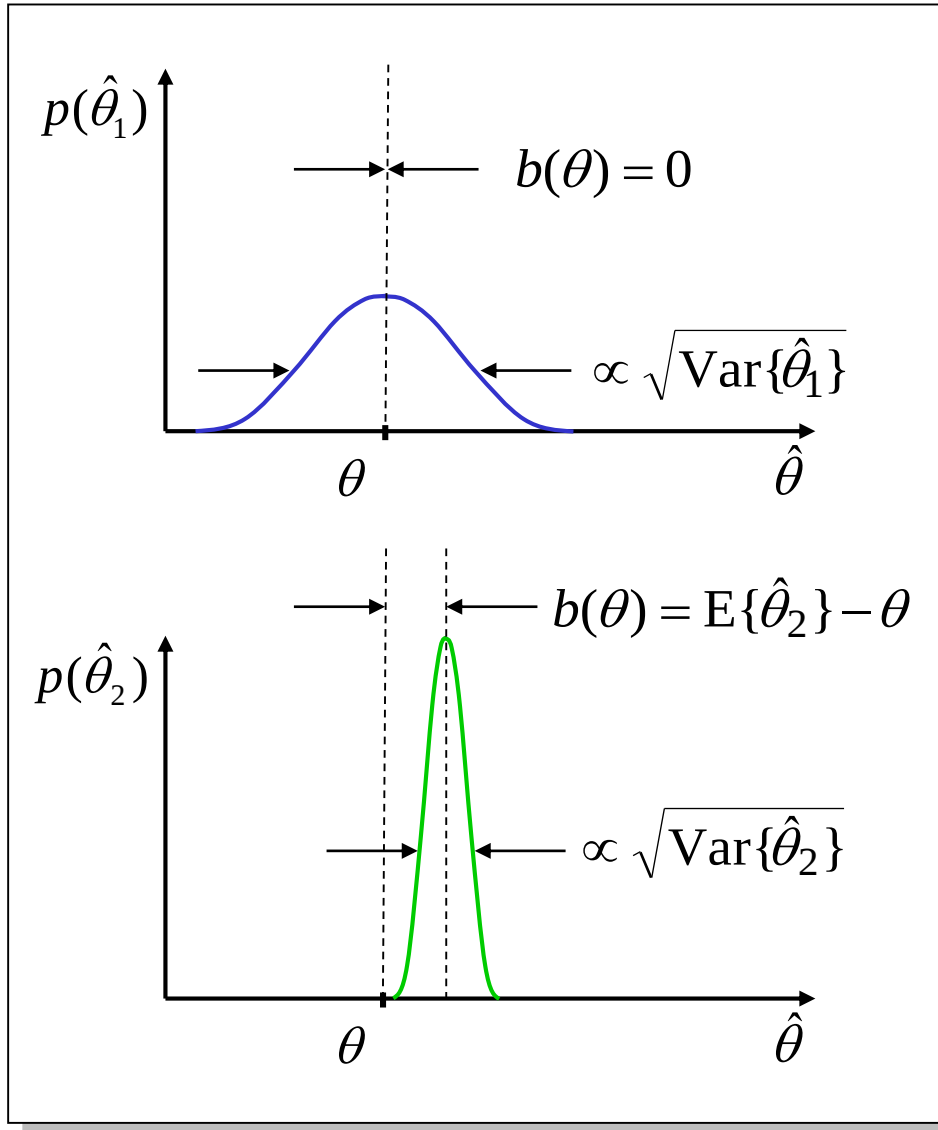
Ein Schätzverfahren heißt **konsistent**, falls mit wachsender Zahl N der zur Schätzung herangezogenen Beobachtungen die Wahrscheinlichkeit, mit der die Schätzung vom wahren Wert um ein beliebiges ε abweicht, gegen Null konvergiert:

$$\lim_{N \rightarrow \infty} P\left\{ \|\hat{\boldsymbol{\theta}}(D) - \boldsymbol{\theta}\| \geq \varepsilon \right\} = 0, \quad \forall \varepsilon > 0$$

$$D = \{\mathbf{m}_1, \dots, \mathbf{m}_N\}$$

4. Parameterschätzung

Diskussion: Erwartungstreue, Varianz, mittlerer quadratischer Fehler



θ : sei wahrer Wert

Schätzer $\hat{\theta}_1$: erwartungstreu

Schätzer $\hat{\theta}_2$: nicht erwartungstreu, aber:

$$\sqrt{\text{Var}\{\hat{\theta}_2\}} < \sqrt{\text{Var}\{\hat{\theta}_1\}}$$

$\text{Var}\{.\}$: **stochastischer Fehler**

Bias b : **systematischer Fehler**

$E\{(\hat{\theta} - \theta)^2\} = b^2(\theta) + \text{Var}\{\hat{\theta}\}$:
mittlerer quadratischer Fehler

Welcher Schätzer ist besser?

4. Parameterschätzung

Beispiel: N stochastisch unabhängig gewonnene Stichproben einer Zufallsvariablen m , für die gilt: $m \sim p(m | \mu)$ (unabhängige identisch verteilte Stichprobe: i.i.d.). Es gelte:

$$\mu = E\{m\}$$

$$\sigma^2 = \text{Var}\{m\} = E\{(m - \mu)^2\}$$

Erwartungswert μ soll durch arithmetischen Mittelwert geschätzt werden.

$$\hat{\mu} = \frac{1}{N} \sum_{k=1}^N m_k$$

$$E\{\hat{\mu}\} = E\left\{\frac{1}{N} \sum_{k=1}^N m_k\right\} = \frac{1}{N} \sum_{k=1}^N E\{m_k\} = \frac{1}{N} \sum_{k=1}^N \mu = \mu \quad \rightarrow \text{Erwartungstreue}$$

4. Parameterschätzung

Beispiel: Fortsetzung

Varianz des Schätzers:

$$\text{Var}\{\hat{\mu}\} = \text{E}\{(\hat{\mu} - \mu)^2\} = \text{E}\left\{\left(\frac{1}{N} \sum_{k=1}^N m_k - \mu\right)^2\right\} = \frac{1}{N^2} \sum_{k=1}^N \sum_{l=1}^N \text{E}\{(m_k - \mu)(m_l - \mu)\}$$

$$\text{Var}\{\hat{\mu}\} = \frac{1}{N^2} \sum_{k=1}^N \text{E}\{(m_k - \mu)^2\} = \frac{1}{N^2} \sum_{k=1}^N \sigma^2 = \frac{1}{N} \sigma^2$$

- Varianz verschwindet mit wachsendem $N \rightarrow$ Konsistenz aufgrund der Ungleichung von Tschebyscheff: $P\left\{|\hat{\theta} - \text{E}\{\hat{\theta}\}| \geq \varepsilon\right\} \leq \text{Var}\{\hat{\theta}\} / \varepsilon^2, \quad \forall \varepsilon > 0$
- Abnahme der Standardabweichung des Schätzers $\sqrt{\text{Var}\{\hat{\mu}\}}$ mit $1/\sqrt{N}$ ist typisch für die meisten praktisch relevanten Aufgabenstellungen.

4.1. Maximum Likelihood Schätzung

Gegeben:

N stochastisch unabhängig gewonnene Stichproben eines Zufallsvektors $\mathbf{m}$, für den gilt: $\mathbf{m} \sim p(\mathbf{m} | \boldsymbol{\theta})$ (unabhängige identisch verteilte Stichprobe: i.i.d.).

Dann gilt:

$$p(D | \boldsymbol{\theta}) = \prod_{k=1}^N p(\mathbf{m}_k | \boldsymbol{\theta}) \quad D = \{\mathbf{m}_1, \dots, \mathbf{m}_N\}$$

$p(D|\boldsymbol{\theta})$ als Funktion von $\boldsymbol{\theta}$ wird als **Likelihoodfunktion** bezeichnet.

Maximum Likelihood Ansatz: Man schätze $\boldsymbol{\theta}$ derart, dass die **Wahrscheinlichkeit** (Wahrscheinlichkeitsdichte) $p(D|\boldsymbol{\theta})$ **der vorliegenden Beobachtungen D maximal** wird.

$$\hat{\boldsymbol{\theta}}_{\text{ML}} = \arg \max_{\boldsymbol{\theta}} \{ p(D | \boldsymbol{\theta}) \}$$

4.1. Maximum Likelihood Schätzung

Wegen der Monotonie der Logarithmusfunktion, kann man äquivalent die logarithmierte Likelihoodfunktion maximieren; das bringt oft analytische Vorteile.

$$l(\boldsymbol{\theta}) := \ln(p(D | \boldsymbol{\theta})) \quad \text{Loglikelihoodfunktion}$$

$$l(\boldsymbol{\theta}) = \sum_{k=1}^N \ln(p(\mathbf{m}_k | \boldsymbol{\theta}))$$

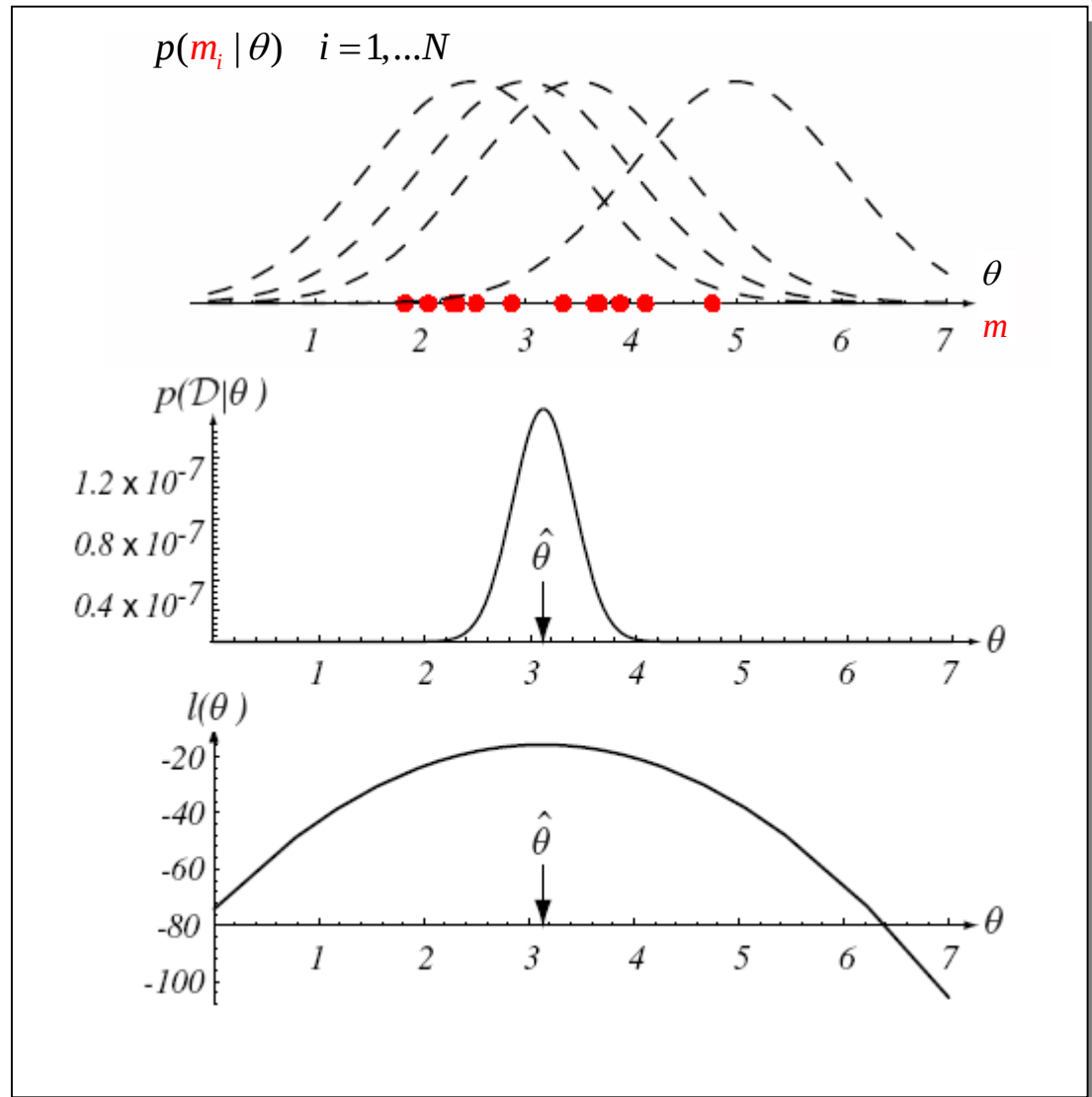
$$\hat{\boldsymbol{\theta}}_{\text{ML}} = \arg \max_{\boldsymbol{\theta}} \{l(\boldsymbol{\theta})\}$$

Notwendig für Maximum:

$$\nabla_{\boldsymbol{\theta}} l(\boldsymbol{\theta}) = \sum_{k=1}^N \nabla_{\boldsymbol{\theta}} \ln(p(\mathbf{m}_k | \boldsymbol{\theta})) \stackrel{!}{=} \mathbf{0} \quad \text{mit} \quad \nabla_{\boldsymbol{\theta}} \equiv \left[\frac{\partial}{\partial \theta_1}, \dots, \frac{\partial}{\partial \theta_c} \right]^T$$

4.1. Maximum Likelihood Schätzung

Beispiel: rote Punkte sind Stichproben einer normalverteilten Zufallsvariablen m mit bekannter Standardabweichung σ und unbekanntem, zu schätzendem Erwartungswert θ .



Quelle: R. O. Duda, P. E. Hart, D. G. Stork: Pattern Classification

4.1. Maximum Likelihood Schätzung

Beispiel: $\mathbf{m} \sim \mathcal{N}(\mathbf{m}; \boldsymbol{\mu}, \boldsymbol{\Sigma})$, $\boldsymbol{\mu}$ unbekannt

$$\ln p(\mathbf{m}_k | \boldsymbol{\mu}) = -\frac{1}{2}(\mathbf{m}_k - \boldsymbol{\mu})^T \boldsymbol{\Sigma}^{-1}(\mathbf{m}_k - \boldsymbol{\mu}) - \frac{d}{2} \ln 2\pi - \frac{1}{2} \ln |\boldsymbol{\Sigma}|$$

$$\nabla_{\boldsymbol{\mu}} \ln p(\mathbf{m}_k | \boldsymbol{\mu}) = \boldsymbol{\Sigma}^{-1}(\mathbf{m}_k - \boldsymbol{\mu})$$

$$\sum_{k=1}^N \nabla_{\boldsymbol{\mu}} \ln p(\mathbf{m}_k | \boldsymbol{\mu}) = \sum_{k=1}^N \boldsymbol{\Sigma}^{-1}(\mathbf{m}_k - \hat{\boldsymbol{\mu}}) = 0$$

$$\sum_{k=1}^N (\mathbf{m}_k - \hat{\boldsymbol{\mu}}) = 0$$

$$\Rightarrow \hat{\boldsymbol{\mu}}_{\text{ML}} = \frac{1}{N} \sum_{k=1}^N \mathbf{m}_k$$

4.1. Maximum Likelihood Schätzung

Beispiel: $m \sim N(m; \mu, \sigma^2)$, μ, σ^2 unbekannt

$$\boldsymbol{\theta} = (\mu, \sigma^2)^T = (\theta_1, \theta_2)^T$$

$$\ln p(m_k | \boldsymbol{\theta}) = -\frac{1}{2\theta_2}(m_k - \theta_1)^2 - \frac{1}{2} \ln 2\pi\theta_2$$

$$\sum_{k=1}^N \nabla_{\boldsymbol{\theta}} \ln p(m_k | \boldsymbol{\theta}) = \sum_{k=1}^N \left[\frac{1}{\theta_2}(m_k - \theta_1), -\frac{1}{2\theta_2} + \frac{1}{2\theta_2^2}(m_k - \theta_1)^2 \right]^T \stackrel{!}{=} \mathbf{0} \Leftrightarrow$$

$$\left. \begin{aligned} \sum_{k=1}^N \frac{1}{\hat{\theta}_2}(m_k - \hat{\theta}_1) &= 0 \\ \sum_{k=1}^N \left(-\frac{1}{2\hat{\theta}_2} + \frac{1}{2\hat{\theta}_2^2}(m_k - \hat{\theta}_1)^2 \right) &= 0 \end{aligned} \right\} \Leftrightarrow$$

ML-Schätzer:

$$\hat{\theta}_1 = \hat{\mu}_{\text{ML}} = \frac{1}{N} \sum_{k=1}^N m_k$$

$$\hat{\theta}_2 = \hat{\sigma}_{\text{ML}}^2 = \frac{1}{N} \sum_{k=1}^N (m_k - \hat{\mu})^2$$

4.1. Maximum Likelihood Schätzung

Beispiel: $\mathbf{m} \sim \mathcal{N}(\mathbf{m}; \boldsymbol{\mu}, \boldsymbol{\Sigma})$, $\boldsymbol{\mu}, \boldsymbol{\Sigma}$ unbekannt

$$\boldsymbol{\theta} = (\boldsymbol{\mu}, \boldsymbol{\Sigma})^T$$

$$\ln p(\mathbf{m}_k | \boldsymbol{\theta}) = -\frac{1}{2}(\mathbf{m}_k - \boldsymbol{\mu})^T \boldsymbol{\Sigma}^{-1}(\mathbf{m}_k - \boldsymbol{\mu}) - \frac{d}{2} \ln 2\pi - \frac{1}{2} \ln |\boldsymbol{\Sigma}|$$

$$\sum_{k=1}^N \nabla_{\boldsymbol{\theta}} \ln p(\mathbf{m}_k | \boldsymbol{\theta}) =$$

$$\sum_{k=1}^N \left[\boldsymbol{\Sigma}^{-1}(\mathbf{m}_k - \boldsymbol{\mu}), \quad -\frac{1}{2} \frac{\partial}{\partial \boldsymbol{\Sigma}} (\mathbf{m}_k - \boldsymbol{\mu})^T \boldsymbol{\Sigma}^{-1}(\mathbf{m}_k - \boldsymbol{\mu}) - \frac{1}{2} \frac{\partial}{\partial \boldsymbol{\Sigma}} \ln |\boldsymbol{\Sigma}| \right]^T \stackrel{!}{=} \mathbf{0}$$

$$\Leftrightarrow \begin{cases} \sum_{k=1}^N \hat{\boldsymbol{\Sigma}}^{-1}(\mathbf{m}_k - \hat{\boldsymbol{\mu}}) = \mathbf{0} \\ \sum_{k=1}^N \left[-\frac{1}{2} \frac{\partial}{\partial \hat{\boldsymbol{\Sigma}}} (\mathbf{m}_k - \hat{\boldsymbol{\mu}})^T \hat{\boldsymbol{\Sigma}}^{-1}(\mathbf{m}_k - \hat{\boldsymbol{\mu}}) - \frac{1}{2} \frac{\partial}{\partial \hat{\boldsymbol{\Sigma}}} \ln |\hat{\boldsymbol{\Sigma}}| \right] = \mathbf{0} \end{cases}$$

4.1. Maximum Likelihood Schätzung

Beispiel, Fortsetzung

$$\sum_{k=1}^N \left[-\frac{1}{2} \frac{\partial}{\partial \hat{\Sigma}} (\mathbf{m}_k - \hat{\mu})^T \hat{\Sigma}^{-1} (\mathbf{m}_k - \hat{\mu}) - \frac{1}{2} \frac{\partial}{\partial \hat{\Sigma}} \ln |\hat{\Sigma}| \right] = \mathbf{0}$$

$$-\frac{\partial}{\partial \hat{\Sigma}} \sum_{k=1}^N [(\mathbf{m}_k - \hat{\mu})^T \hat{\Sigma}^{-1} (\mathbf{m}_k - \hat{\mu})] - N \frac{\partial}{\partial \hat{\Sigma}} \ln |\hat{\Sigma}| = \mathbf{0}$$

$$\bar{\mathbf{m}} = \frac{1}{N} \sum_{k=1}^N \mathbf{m}_k$$

$$\frac{\partial}{\partial \hat{\Sigma}} \sum_{k=1}^N [(\mathbf{m}_k - \bar{\mathbf{m}} + \bar{\mathbf{m}} - \hat{\mu})^T \hat{\Sigma}^{-1} (\mathbf{m}_k - \bar{\mathbf{m}} + \bar{\mathbf{m}} - \hat{\mu})] + N \frac{\partial}{\partial \hat{\Sigma}} \ln |\hat{\Sigma}| = \mathbf{0}$$

$$\frac{\partial}{\partial \hat{\Sigma}} \sum_{k=1}^N [(\mathbf{m}_k - \bar{\mathbf{m}})^T \hat{\Sigma}^{-1} (\mathbf{m}_k - \bar{\mathbf{m}})] + N \frac{\partial}{\partial \hat{\Sigma}} (\bar{\mathbf{m}} - \hat{\mu})^T \hat{\Sigma}^{-1} (\bar{\mathbf{m}} - \hat{\mu}) + N \frac{\partial}{\partial \hat{\Sigma}} \ln |\hat{\Sigma}| = \mathbf{0}$$

Skalar

$$\frac{\partial}{\partial \hat{\Sigma}} \text{tr} \left[\hat{\Sigma}^{-1} \sum_{k=1}^N (\mathbf{m}_k - \bar{\mathbf{m}})(\mathbf{m}_k - \bar{\mathbf{m}})^T \right] + N \frac{\partial}{\partial \hat{\Sigma}} (\bar{\mathbf{m}} - \hat{\mu})^T \hat{\Sigma}^{-1} (\bar{\mathbf{m}} - \hat{\mu}) + N \frac{\partial}{\partial \hat{\Sigma}} \ln |\hat{\Sigma}| = \mathbf{0}$$

$$N \frac{\partial}{\partial \hat{\Sigma}} \text{tr} (\hat{\Sigma}^{-1} \tilde{\mathbf{S}}) + N \frac{\partial}{\partial \hat{\Sigma}} (\bar{\mathbf{m}} - \hat{\mu})^T \hat{\Sigma}^{-1} (\bar{\mathbf{m}} - \hat{\mu}) + N \frac{\partial}{\partial \hat{\Sigma}} \ln |\hat{\Sigma}| = \mathbf{0}$$

$$\text{mit } \tilde{\mathbf{S}} := \frac{1}{N} \mathbf{S} = \frac{1}{N} \sum_{k=1}^N (\mathbf{m}_k - \bar{\mathbf{m}})^T (\mathbf{m}_k - \bar{\mathbf{m}})$$

4.1. Maximum Likelihood Schätzung

Beispiel, Fortsetzung

$$\frac{\partial}{\partial \hat{\Sigma}} \text{tr}(\hat{\Sigma}^{-1} \tilde{\mathbf{S}}) + \frac{\partial}{\partial \hat{\Sigma}} (\bar{\mathbf{m}} - \hat{\boldsymbol{\mu}})^T \hat{\Sigma}^{-1} (\bar{\mathbf{m}} - \hat{\boldsymbol{\mu}}) - \frac{\partial}{\partial \hat{\Sigma}} \ln |\hat{\Sigma}^{-1}| = \mathbf{0}$$

$$\left. \begin{aligned} \frac{\partial}{\partial \hat{\Sigma}} \text{tr}(\mathbf{V} \tilde{\mathbf{S}}) + \frac{\partial}{\partial \hat{\Sigma}} (\bar{\mathbf{m}} - \hat{\boldsymbol{\mu}})^T \mathbf{V} (\bar{\mathbf{m}} - \hat{\boldsymbol{\mu}}) - \frac{\partial}{\partial \hat{\Sigma}} \ln |\mathbf{V}| &= \mathbf{0} \\ \frac{\partial \ln p(\mathbf{m}_k | \boldsymbol{\theta})}{\partial \hat{\Sigma}} &= \frac{\partial \ln p(\mathbf{m}_k | \boldsymbol{\theta})}{\partial \mathbf{V}} \frac{\partial \mathbf{V}}{\partial \hat{\Sigma}} \\ \frac{\partial \mathbf{V}}{\partial \hat{\Sigma}} &= -\hat{\Sigma}^{-2} \end{aligned} \right\} \Rightarrow$$

$$\left. \begin{aligned} \frac{\partial}{\partial \mathbf{V}} \text{tr}(\mathbf{V} \tilde{\mathbf{S}}) + \frac{\partial}{\partial \mathbf{V}} (\bar{\mathbf{m}} - \hat{\boldsymbol{\mu}})^T \mathbf{V} (\bar{\mathbf{m}} - \hat{\boldsymbol{\mu}}) - \frac{\partial}{\partial \mathbf{V}} \ln |\mathbf{V}| &= \mathbf{0} \\ \frac{\partial}{\partial \mathbf{X}} \ln |\mathbf{X}| &= 2\mathbf{X}^{-1} - \text{diag}(\mathbf{X}^{-1}) \\ \frac{\partial}{\partial \mathbf{X}} \text{tr}(\mathbf{A}\mathbf{X}) &= \frac{\partial}{\partial \mathbf{X}} \text{tr}(\mathbf{X}\mathbf{A}) = \mathbf{A}^T + \mathbf{A} - \text{diag}(\mathbf{A}) \end{aligned} \right\} \Rightarrow$$

$$2\tilde{\mathbf{S}} - \text{diag}(\tilde{\mathbf{S}}) + \frac{\partial}{\partial \mathbf{V}} (\bar{\mathbf{m}} - \hat{\boldsymbol{\mu}})^T \mathbf{V} (\bar{\mathbf{m}} - \hat{\boldsymbol{\mu}}) - 2\hat{\Sigma} + \text{diag}(\hat{\Sigma}) = \mathbf{0}$$

4.1. Maximum Likelihood Schätzung

Beispiel, Fortsetzung

$$\begin{cases} \hat{\boldsymbol{\mu}} = \frac{1}{N} \sum_{k=1}^N \mathbf{m}_k \\ 2\tilde{\mathbf{S}} - \text{diag}(\tilde{\mathbf{S}}) + \frac{\partial}{\partial \mathbf{V}} (\bar{\mathbf{m}} - \hat{\boldsymbol{\mu}})^T \mathbf{V} (\bar{\mathbf{m}} - \hat{\boldsymbol{\mu}}) - 2\hat{\boldsymbol{\Sigma}} + \text{diag}(\hat{\boldsymbol{\Sigma}}) = \mathbf{0} \end{cases}$$

$$\Leftrightarrow \begin{cases} \hat{\boldsymbol{\mu}} = \frac{1}{N} \sum_{k=1}^N \mathbf{m}_k \\ 2\tilde{\mathbf{S}} - 2\hat{\boldsymbol{\Sigma}} + \text{diag}(\hat{\boldsymbol{\Sigma}} - \tilde{\mathbf{S}}) = \mathbf{0} \end{cases}$$

$\Leftrightarrow$

ML-Schätzer:

$$\hat{\boldsymbol{\mu}}_{\text{ML}} = \frac{1}{N} \sum_{k=1}^N \mathbf{m}_k$$

$$\hat{\boldsymbol{\Sigma}}_{\text{ML}} = \frac{1}{N} \sum_{k=1}^N (\mathbf{m}_k - \hat{\boldsymbol{\mu}})(\mathbf{m}_k - \hat{\boldsymbol{\mu}})^T$$

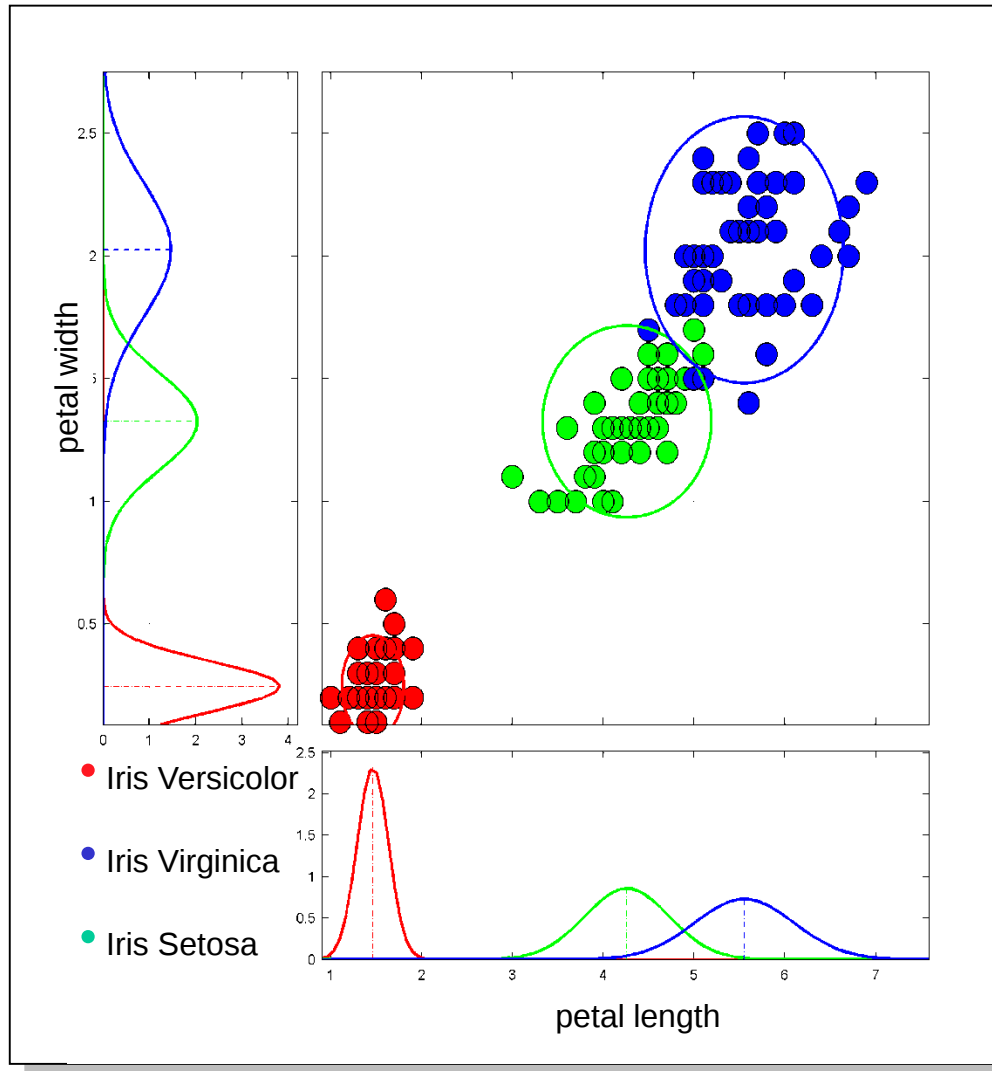
Differentiation nach Vektoren und Matrizen, siehe z.B.:

T. K. Moon; W. C. Stirling: Mathematical Methods and Algorithms for Signal Processing, Prentice Hall, 2000.

C. Voigt; J. Adamy: Formelsammlung der Matrizenrechnung, Oldenbourg Verlag, 2007.

4.1. Maximum Likelihood Schätzung

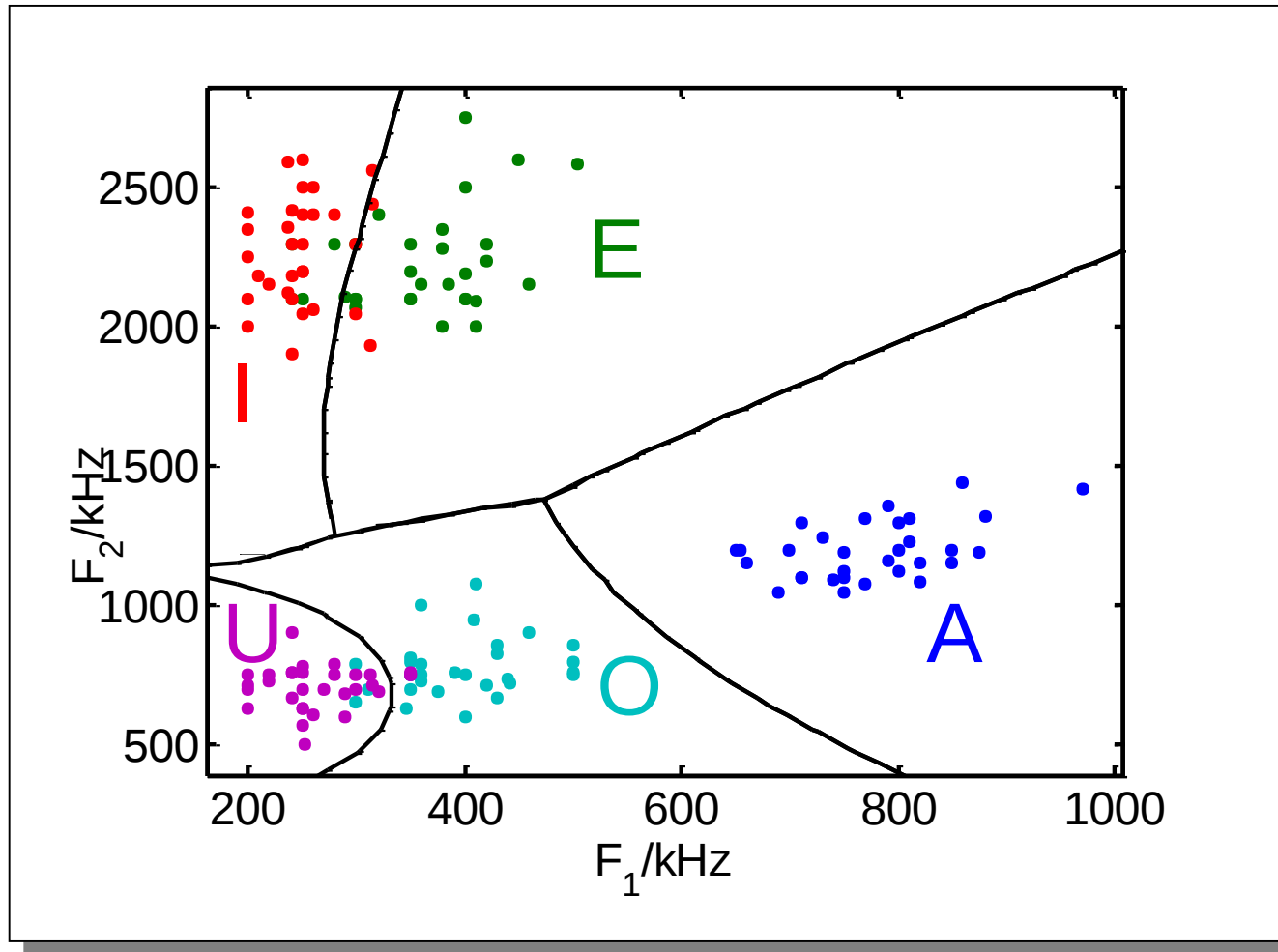
Beispiel: Iris Datensatz nach R. A. Fisher, Modell: Gauß'sche WDFen mit unkorrelierten Merkmalen, Maximum Likelihood Schätzung der Parameter



Ellipsen kennzeichnen die 2σ - Grenzen der Wahrscheinlichkeitsdichten.

4.1. Maximum Likelihood Schätzung

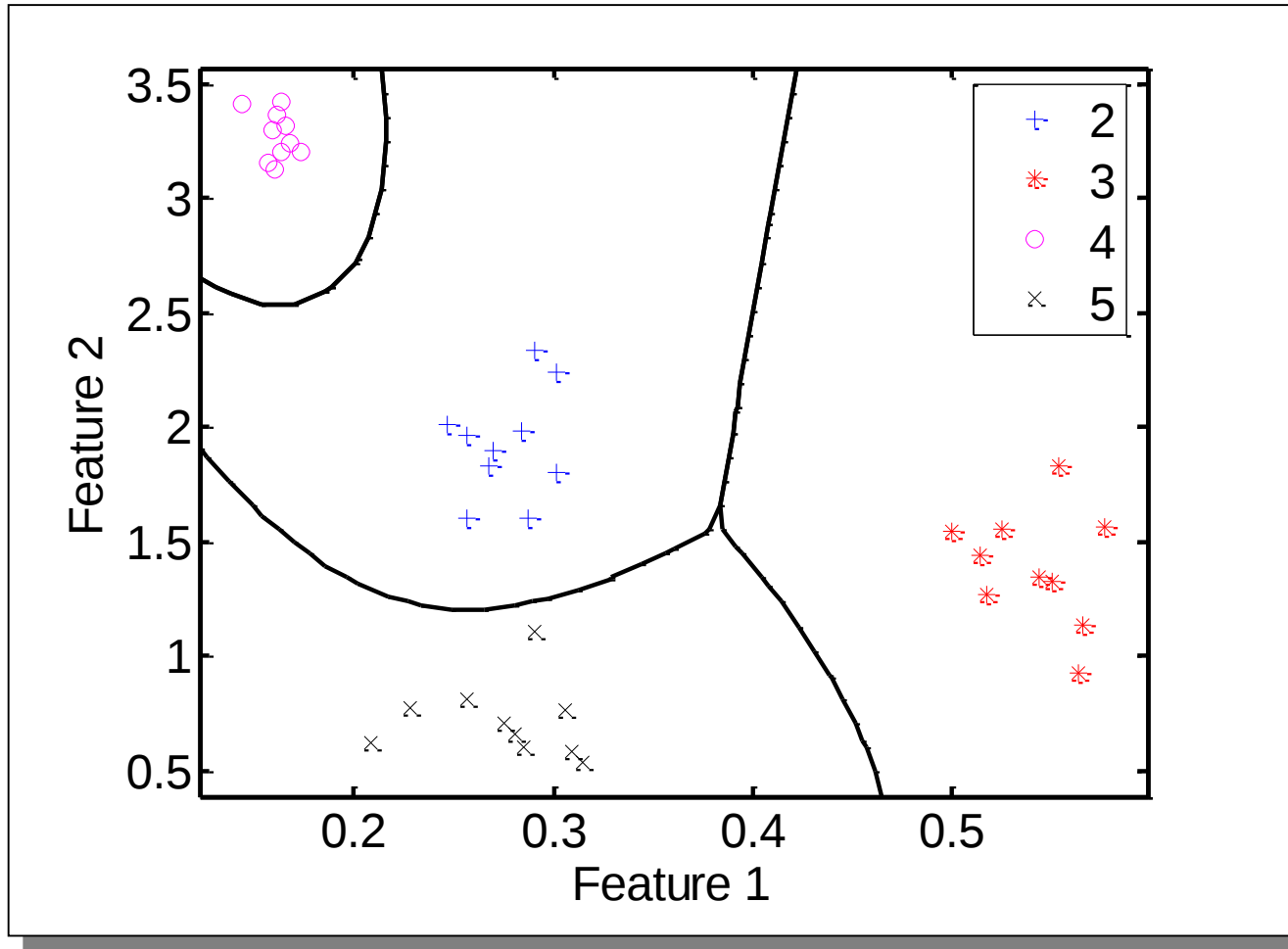
Beispiel: Vokalerkennung



Bayes'sche Klassifikation, Modell: Normalverteilung für klassenbedingte Merkmals-WDFen, Parameterschätzung: Maximum Likelihood

4.1. Maximum Likelihood Schätzung

Beispiel: Grifferkennung Hand



Bayes'sche Klassifikation, Modell: Normalverteilung für klassenbedingte Merkmals-WDFen, Parameterschätzung: Maximum Likelihood

4.2. Bayes'sche Bestimmung der AP-Klassenwahrscheinlichkeiten

Fundamentalgröße der Bayes'schen Klassifikation:

A Posteriori Wahrscheinlichkeitsverteilungen der Klassen

$$P(\omega_i | \mathbf{m}) = \frac{p(\mathbf{m} | \omega_i)P(\omega_i)}{p(\mathbf{m})}$$

Problem: Wenn unbekannt → Müssen aus den Daten D geschätzt werden.

Ansatz:

$$P(\omega_i | \mathbf{m}, D) = \frac{p(\mathbf{m} | \omega_i, D)P(\omega_i | D)}{\sum_{j=1}^c p(\mathbf{m} | \omega_j, D)P(\omega_j | D)}$$

Die $p(\mathbf{m} | \omega_i, D)$ approximieren die $p(\mathbf{m} | \omega_i)$.

4.2. Bayes'sche Bestimmung der AP-Klassenwahrscheinlichkeiten

Vereinfachende Annahmen:

- A Priori Wahrscheinlichkeiten seien gegeben oder ausreichend genau bestimmt worden: $P(\omega_i) = P(\omega_i | D)$
- Die Stichproben in D_k haben **keinen Einfluss** auf $p(\mathbf{m} | \omega_i, D)$ $i \neq k$.

$$P(\omega_i | \mathbf{m}, D) = \frac{p(\mathbf{m} | \omega_i, D_i)P(\omega_i)}{\sum_{j=1}^c p(\mathbf{m} | \omega_j, D_j)P(\omega_j)}$$

- WDFen $p(\mathbf{m} | \omega_i)$ haben eine bekannte parametrisierte Form $p(\mathbf{m} | \boldsymbol{\theta}_i, \omega_i)$ mit $p(\boldsymbol{\theta}_i | \omega_i)$ bekannt.

$$p(\mathbf{m} | \omega_i, D_i) = \int_{\Theta_i} p(\mathbf{m}, \boldsymbol{\theta}_i | \omega_i, D_i) d\boldsymbol{\theta}_i = \int_{\Theta_i} p(\mathbf{m} | \boldsymbol{\theta}_i, \omega_i, \cancel{D_i}) p(\boldsymbol{\theta}_i | \omega_i, D_i) d\boldsymbol{\theta}_i$$

→

$$p(\mathbf{m} | \omega_i, D_i) = \int_{\Theta_i} p(\mathbf{m} | \boldsymbol{\theta}_i, \omega_i) p(\boldsymbol{\theta}_i | \omega_i, D_i) d\boldsymbol{\theta}_i$$

4.2. Bayes'sche Bestimmung der AP-Klassenwahrscheinlichkeiten

Die Gesamtaufgabe zerfällt in c Teilaufgaben; eine pro Klasse. Für jede Klasse ω_i gilt es, die klassenspezifische WDF

$$p(\mathbf{m} | \omega_i, D_i) = \int_{\Theta_i} p(\mathbf{m} | \boldsymbol{\theta}_i, \omega_i) p(\boldsymbol{\theta}_i | \omega_i, D_i) d\boldsymbol{\theta}_i \quad (*)$$

anhand der Daten D_i zu bestimmen.

Zur einfacheren Schreibweise wird im Folgenden die Kennzeichnung der Klassen unterdrückt. Die Überlegungen gelten aber immer nur für eine Klasse ω_i .

→ Vereinfachte Schreibweise für (*):

$$p(\mathbf{m} | D) = \int_{\Theta} p(\mathbf{m} | \boldsymbol{\theta}) p(\boldsymbol{\theta} | D) d\boldsymbol{\theta}$$

Da $p(\mathbf{m} | \boldsymbol{\theta})$ als bekannt vorausgesetzt ist, gilt es im Folgenden, $p(\boldsymbol{\theta} | D)$ zu bestimmen.

4.2. Bayes'sche Bestimmung der AP-Klassenwahrscheinlichkeiten

Beispiel: $m \sim N(m; \mu, \sigma^2)$, μ unbekannt

Annahme: $\mu \sim N(\mu; \mu_0, \sigma_0^2)$ (μ_0, σ_0^2 sind sogenannte Hyperparameter)

$$p(\mu | D) = \frac{p(D | \mu) p(\mu)}{\int p(D | \mu) p(\mu) d\mu} = \alpha(D) \prod_{k=1}^N p(m_k | \mu) p(\mu)$$

$$p(\mu | D) = \alpha(D) \underbrace{\prod_{k=1}^N \frac{1}{\sqrt{2\pi}\sigma} \exp\left[-\frac{1}{2}\left(\frac{m_k - \mu}{\sigma}\right)^2\right]}_{p(m_k | \mu)} \underbrace{\frac{1}{\sqrt{2\pi}\sigma_0} \exp\left[-\frac{1}{2}\left(\frac{\mu - \mu_0}{\sigma_0}\right)^2\right]}_{p(\mu)}$$

$$= \alpha'(D) \exp\left[-\frac{1}{2}\left(\sum_{k=1}^N \left(\frac{\mu - m_k}{\sigma}\right)^2 + \left(\frac{\mu - \mu_0}{\sigma_0}\right)^2\right)\right]$$

$$= \alpha''(D) \exp\left[-\frac{1}{2}\left[\left(\frac{N}{\sigma^2} + \frac{1}{\sigma_0^2}\right)\mu^2 - 2\left(\frac{1}{\sigma^2} \sum_{k=1}^N m_k + \frac{\mu_0}{\sigma_0^2}\right)\mu\right]\right]$$

($\rightarrow$... kann nur eine Gauß'sche WDF sein! Quadratische Ergänzung $\rightarrow$...)

4.2. Bayes'sche Bestimmung der AP-Klassenwahrscheinlichkeiten

Beispiel, Fortsetzung

Diese WDF lässt sich schreiben als: $p(\mu | D) = N(\mu; \mu_N, \sigma_N^2)$

$$p(\mu | D) = \frac{1}{\sqrt{2\pi}\sigma_N} \exp\left[-\frac{1}{2}\left(\frac{\mu - \mu_N}{\sigma_N}\right)^2\right]$$

$$\mu_N := \left(\frac{N\sigma_0^2}{N\sigma_0^2 + \sigma^2}\right)\hat{\mu}_N + \frac{\sigma^2}{N\sigma_0^2 + \sigma^2}\mu_0 \quad \sigma_N^2 := \frac{\sigma_0^2\sigma^2}{N\sigma_0^2 + \sigma^2}$$

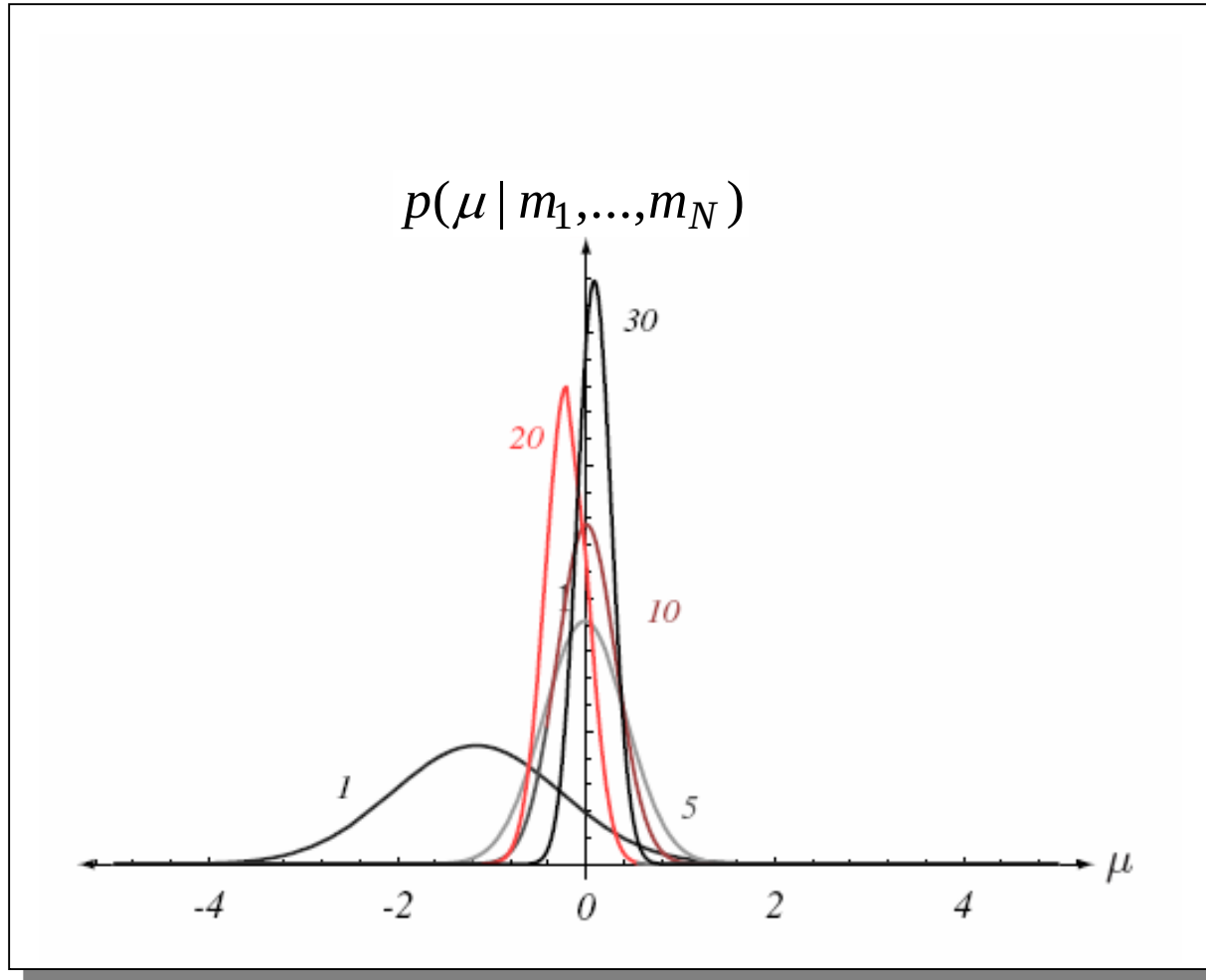
$$\hat{\mu}_N := \frac{1}{N} \sum_{k=1}^N m_k$$

Diskussion:

- A Priori + Empirische Information $\rightarrow$ A Posteriori WDF $p(\mu | D)$
- μ_N die beste Schätzung (MAP) für μ nach N beobachteten Werten.
- σ_N^2 beschreibt die Schätzunsicherheit (stochastischer Schätzfehler)
- μ_N liegt zwischen $\hat{\mu}_N$ und μ_0

4.2. Bayes'sche Bestimmung der AP-Klassenwahrscheinlichkeiten

Beispiel, Fortsetzung



Quelle: R. O. Duda, P. E. Hart, D. G. Stork: Pattern Classification

A Posteriori Verteilungsdichte für wachsende Anzahl von Beobachtungen.

4.2. Bayes'sche Bestimmung der AP-Klassenwahrscheinlichkeiten

Beispiel, Fortsetzung

$$p(m | D) = \int_{\Theta} p(m | \mu) p(\mu | D) d\mu$$

$$= \int \frac{1}{\sqrt{2\pi}\sigma} \exp\left[-\frac{1}{2}\left(\frac{m-\mu}{\sigma}\right)^2\right] \frac{1}{\sqrt{2\pi}\sigma_N} \exp\left[-\frac{1}{2}\left(\frac{\mu-\mu_N}{\sigma_N}\right)^2\right] d\mu$$

$$= \frac{1}{2\pi\sigma\sigma_N} \exp\left[-\frac{1}{2} \frac{(m-\mu_N)^2}{\sigma^2 + \sigma_N^2}\right] f(\sigma, \sigma_N)$$

$$\text{mit } f(\sigma, \sigma_N) := \int \exp\left[-\frac{1}{2} \frac{\sigma^2 + \sigma_N^2}{\sigma^2 \sigma_N^2} \left(\mu - \frac{\sigma_N^2 m + \sigma^2 \mu_N}{\sigma_N^2 + \sigma^2}\right)^2\right] d\mu$$

$$p(m | D) = N(m; \mu_N, \sigma^2 + \sigma_N^2)$$

4.2. Bayes'sche Bestimmung der AP-Klassenwahrscheinlichkeiten

Beispiel: $p(\mathbf{m} | \boldsymbol{\mu}) = N(\mathbf{m}; \boldsymbol{\mu}, \boldsymbol{\Sigma})$, $\boldsymbol{\mu}$ unbekannt

Annahme: $p(\boldsymbol{\mu}) = N(\boldsymbol{\mu}; \boldsymbol{\mu}_0, \boldsymbol{\Sigma}_0)$

$$\begin{aligned} p(\boldsymbol{\mu} | D) &= \alpha \prod_{k=1}^N p(\mathbf{m}_k | \boldsymbol{\mu}) p(\boldsymbol{\mu}) \\ &= \alpha' \exp \left[-\frac{1}{2} \left(\boldsymbol{\mu}^T (N\boldsymbol{\Sigma}^{-1} + \boldsymbol{\Sigma}_0^{-1}) \boldsymbol{\mu} - 2\boldsymbol{\mu}^T \left(\boldsymbol{\Sigma}^{-1} \sum_{k=1}^N \mathbf{m}_k + \boldsymbol{\Sigma}_0^{-1} \boldsymbol{\mu}_0 \right) \right) \right] \\ &= \alpha'' \exp \left[-\frac{1}{2} (\boldsymbol{\mu} - \boldsymbol{\mu}_N)^T \boldsymbol{\Sigma}_N^{-1} (\boldsymbol{\mu} - \boldsymbol{\mu}_N) \right] \end{aligned}$$

$$p(\boldsymbol{\mu} | D) = N(\boldsymbol{\mu}; \boldsymbol{\mu}_N, \boldsymbol{\Sigma}_N)$$

$$\boldsymbol{\mu}_N = \boldsymbol{\Sigma}_0 \left(\boldsymbol{\Sigma}_0 + \frac{1}{N} \boldsymbol{\Sigma} \right)^{-1} \hat{\boldsymbol{\mu}}_N + \frac{1}{N} \boldsymbol{\Sigma} \left(\boldsymbol{\Sigma}_0 + \frac{1}{N} \boldsymbol{\Sigma} \right)^{-1} \boldsymbol{\mu}_0$$

$$\boldsymbol{\Sigma}_N = \boldsymbol{\Sigma}_0 \left(\boldsymbol{\Sigma}_0 + \frac{1}{N} \boldsymbol{\Sigma} \right)^{-1} \frac{1}{N} \boldsymbol{\Sigma}$$

$$\hat{\boldsymbol{\mu}}_N = \frac{1}{N} \sum_{k=1}^N \mathbf{m}_k$$

4.2. Bayes'sche Bestimmung der AP-Klassenwahrscheinlichkeiten

Beispiel, Fortsetzung

$$p(\mathbf{m} \mid D) = \int p(\mathbf{m} \mid \boldsymbol{\mu}) p(\boldsymbol{\mu} \mid D) d\boldsymbol{\mu}$$



$$p(\boldsymbol{\mu} \mid D) = N(\boldsymbol{\mu}_N, \boldsymbol{\Sigma}_N)$$

$$p(\mathbf{m} \mid \boldsymbol{\mu}) = N(\boldsymbol{\mu}, \boldsymbol{\Sigma})$$

$$p(\mathbf{m} \mid D) = N(\mathbf{m}; \boldsymbol{\mu}_N, \boldsymbol{\Sigma} + \boldsymbol{\Sigma}_N)$$

4.2. Bayes'sche Bestimmung der AP-Klassenwahrscheinlichkeiten

Überblick zum Verfahren; für jede Klasse $i = 1, \dots, c$ separat anzuwenden.

- $p(\mathbf{m} | \boldsymbol{\theta}_i, \omega_i)$ strukturell als bekannt angenommen; Wert $\boldsymbol{\theta}_i$ unbekannt.
- $p(\boldsymbol{\theta}_i | \omega_i)$ bekannt; enthält das A Priori Wissen über $\boldsymbol{\theta}_i$.
- Das weitere Wissen über $\boldsymbol{\theta}_i$ ist in D_i enthalten; D_i besteht aus N_i unabhängigen Stichproben $\mathbf{m}_1, \dots, \mathbf{m}_{N_i}$ bezüglich der WDF $p(\mathbf{m} | \omega_i)$.

$$p(D_i | \boldsymbol{\theta}_i, \omega_i) = \prod_{k=1}^N p(\mathbf{m}_k | \boldsymbol{\theta}_i, \omega_i)$$

$$p(\boldsymbol{\theta}_i | D_i, \omega_i) = \frac{p(D_i | \boldsymbol{\theta}_i, \omega_i) p(\boldsymbol{\theta}_i | \omega_i)}{\int p(D_i | \boldsymbol{\theta}_i, \omega_i) p(\boldsymbol{\theta}_i | \omega_i) d\boldsymbol{\theta}_i}$$

$$p(\mathbf{m} | D_i, \omega_i) = \int p(\mathbf{m} | \boldsymbol{\theta}_i, \omega_i) p(\boldsymbol{\theta}_i | D_i, \omega_i) d\boldsymbol{\theta}_i$$

- Wenn $p(\boldsymbol{\theta}_i | D_i, \omega_i)$ sich stark um einen Wert $\boldsymbol{\theta}_i = \hat{\boldsymbol{\theta}}_i$ konzentriert, dort eine ausgeprägte Spitze ausbildet, dann kann $p(\mathbf{m} | D_i, \omega_i)$ durch $p(\mathbf{m} | \hat{\boldsymbol{\theta}}_i, \omega_i)$ approximiert werden.

4.2. Bayes'sche Bestimmung der AP-Klassenwahrscheinlichkeiten

Vergleich: Bayes'sches Verfahren und Maximum Likelihood Schätzung.

Bayes'sches Verfahren:

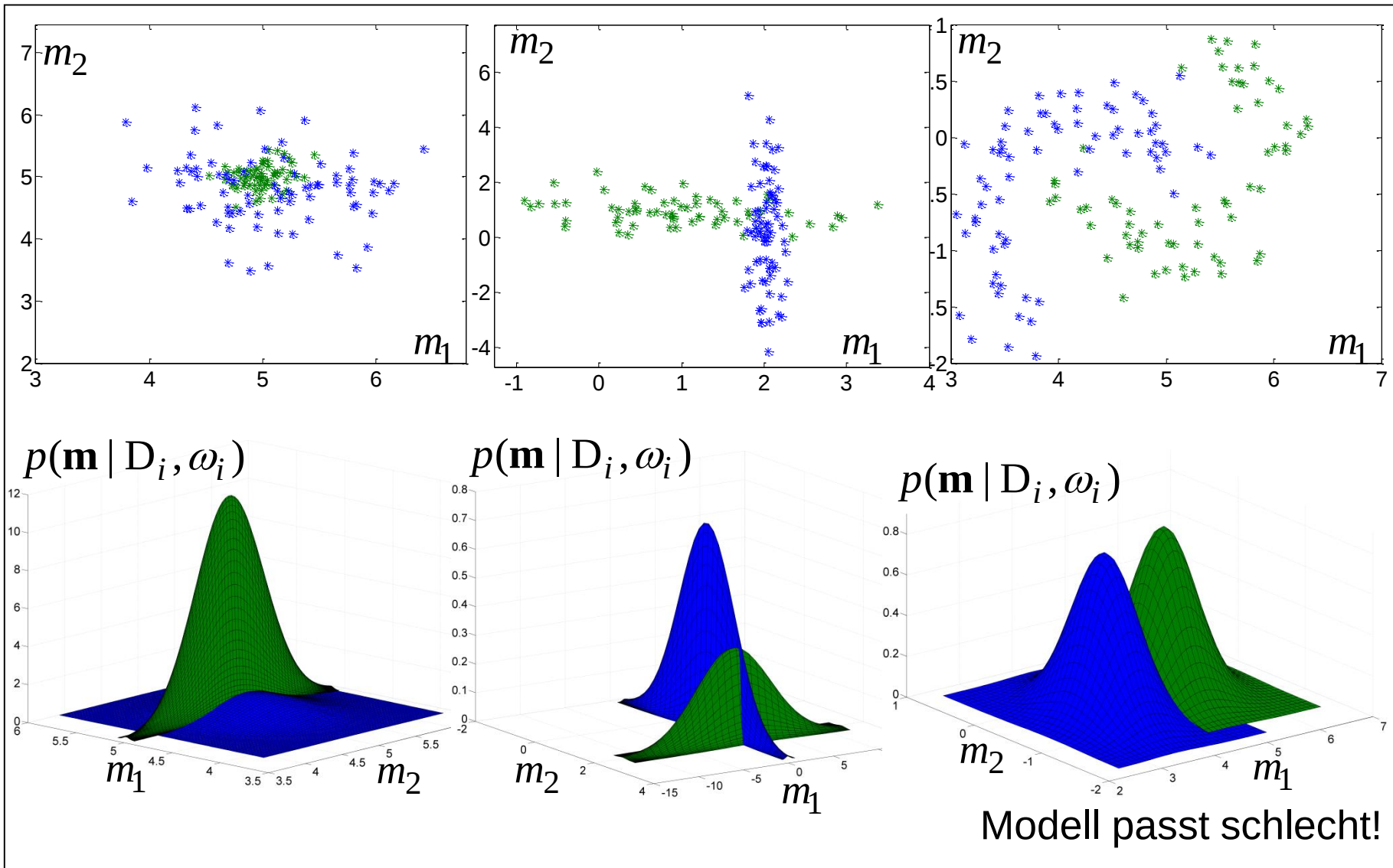
- + A Priori Wissen über θ kann eingebracht werden.
- Numerische Berechnung mehrdimensionaler Integrale erforderlich.

Maximum Likelihood Verfahren:

- + Keine mehrdimensionalen Integrale, sondern *nur* Extremwertsuche.
- A Priori Wissen über θ kann nicht eingebracht werden.

4.2. Bayes'sche Bestimmung der AP-Klassenwahrscheinlichkeiten

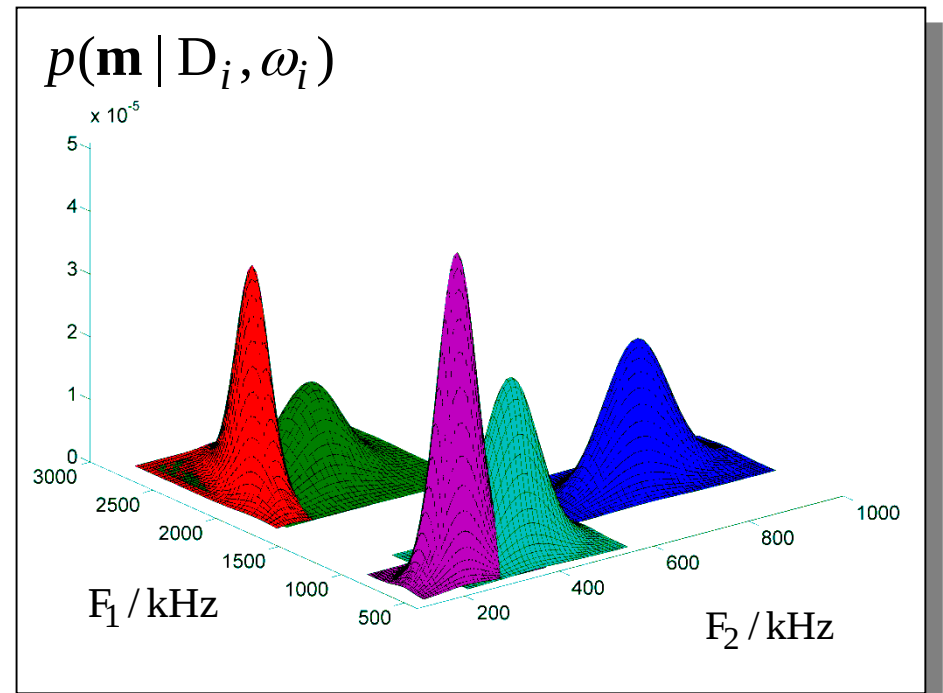
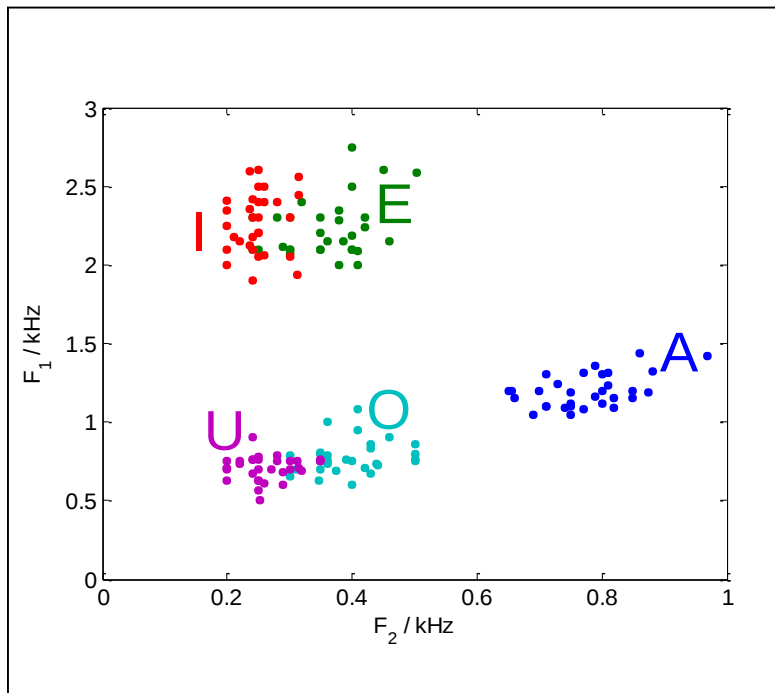
Beispiel: $d = 2$, $c = 2$, Daten und daran angepasste Normalverteilungen



4.2. Bayes'sche Bestimmung der AP-Klassenwahrscheinlichkeiten

Beispiel: Vokalerkennung, $d = 2$, $c = 5$,

Daten und daran angepasste Normalverteilungen



4.3. Bayes'sche Parameterschätzung

Wenn $p(\boldsymbol{\theta} | D)$ sich stark um einen Wert $\boldsymbol{\theta} = \hat{\boldsymbol{\theta}}$ konzentriert, dort eine ausgeprägte Spitze ausbildet, dann kann $p(\mathbf{m} | D)$ durch $p(\mathbf{m} | \hat{\boldsymbol{\theta}})$ approximiert werden.

Mit dieser Approximation wird der Rechenaufwand des Verfahrens aus Abschnitt 4.2. erheblich reduziert, da insbesondere das Integral

$$p(\mathbf{m} | D) = \int p(\mathbf{m} | \boldsymbol{\theta}) p(\boldsymbol{\theta} | D) d\boldsymbol{\theta}$$

nicht ausgewertet werden muss.

$$p(\mathbf{m} | D) = \int p(\mathbf{m} | \boldsymbol{\theta}) p(\boldsymbol{\theta} | D) d\boldsymbol{\theta} \approx \int p(\mathbf{m} | \boldsymbol{\theta}) \delta(\boldsymbol{\theta} - \hat{\boldsymbol{\theta}}) d\boldsymbol{\theta} = p(\mathbf{m} | \hat{\boldsymbol{\theta}})$$

Im Folgenden werden die wichtigsten Techniken zur **Bayes'schen Parameterschätzung** gezeigt.

4.3. Bayes'sche Parameterschätzung

θ wird als **Zufallsvariable** angesehen und über eine WV beschrieben.

Ansatz 1: Minimaler quadratischer Schätzfehler

$$E\{\|\hat{\theta}(D) - \theta\|^2\} = \int_{\mathbb{M}^N} \int_{\Theta} l(\hat{\theta}, \theta) p(\theta, D) d\theta dD \xrightarrow{!} \text{minimal} \quad dD := \prod_{i=1}^N d\mathbf{m}_i$$

Kostenfunktion:

$$l(\hat{\theta}, \theta) := (\hat{\theta} - \theta)^T (\hat{\theta} - \theta)$$

$$E\{\|\hat{\theta}(D) - \theta\|^2\} = \int_{\mathbb{M}^N} \int_{\Theta} (\hat{\theta} - \theta)^T (\hat{\theta} - \theta) p(\theta | D) d\theta p(D) dD \xrightarrow{!} \text{minimal}$$

$$\Leftrightarrow I(D) := \int_{\Theta} (\hat{\theta} - \theta)^T (\hat{\theta} - \theta) p(\theta | D) d\theta \text{ ist minimal.}$$

$$I(D) = \int_{\Theta} \sum_{j=1}^K (\hat{\theta}_j(D) - \theta_j)^2 p(\theta | D) d\theta$$

4.3. Bayes'sche Parameterschätzung

Notwendige Bedingung für Minimum:

$$\begin{aligned}\frac{\partial I(\mathbf{D})}{\partial \hat{\theta}_i(\mathbf{D})} &\stackrel{!}{=} 0 && \Leftrightarrow && \int_{\Theta} 2(\hat{\theta}_i(\mathbf{D}) - \theta_i) p(\boldsymbol{\theta} | \mathbf{D}) d\boldsymbol{\theta} = 0 \\ &&& \Leftrightarrow && \hat{\theta}_i(\mathbf{D}) = \int_{\Theta} \theta_i p(\boldsymbol{\theta} | \mathbf{D}) d\boldsymbol{\theta}\end{aligned}$$

Hinreichende Bedingung für Minimum:

$$\frac{\partial^2 I(\mathbf{D})}{\partial (\hat{\theta}_i(\mathbf{D}))^2} = 2 > 0$$

Schätzer mit minimalem mittleren quadratischen Fehler:

$$\hat{\boldsymbol{\theta}}(\mathbf{D}) = \int_{\Theta} \boldsymbol{\theta} p(\boldsymbol{\theta} | \mathbf{D}) d\boldsymbol{\theta} = \mathbb{E}\{\boldsymbol{\theta} | \mathbf{D}\}$$

$$\text{mit } p(\boldsymbol{\theta} | \mathbf{D}) = \frac{p(\mathbf{D} | \boldsymbol{\theta}) p(\boldsymbol{\theta})}{p(\mathbf{D})}$$

A Posteriori Erwartungswert der Zufallsvariablen $\boldsymbol{\theta}$

4.3. Bayes'sche Parameterschätzung

Ansatz 2: Konstante Gewichtung großer Fehler

Kostenfunktion:

$$l(\hat{\boldsymbol{\theta}}, \boldsymbol{\theta}) := \begin{cases} 0 & \text{für } \|\hat{\boldsymbol{\theta}} - \boldsymbol{\theta}\| < \Delta, \Delta > 0 \\ 1 & \text{sonst} \end{cases}$$

$$E\{l(\hat{\boldsymbol{\theta}}, \boldsymbol{\theta})\} = \int_{\mathbb{M}^N} \int_{\Theta} l(\hat{\boldsymbol{\theta}}, \boldsymbol{\theta}) p(\boldsymbol{\theta}, D) d\boldsymbol{\theta} dD \xrightarrow{!} \text{minimal}$$

$$\Leftrightarrow I(D) = \int_{\Theta} l(\hat{\boldsymbol{\theta}}, \boldsymbol{\theta}) p(\boldsymbol{\theta} | D) d\boldsymbol{\theta} \quad \text{ist minimal.}$$

$$\Leftrightarrow I(D) = 1 - \int_{\{\boldsymbol{\theta} | \|\hat{\boldsymbol{\theta}} - \boldsymbol{\theta}\| < \Delta\}} p(\boldsymbol{\theta} | D) d\boldsymbol{\theta} \quad \text{ist minimal.}$$

Für $\Delta \rightarrow 0$ folgt:

$$\hat{\boldsymbol{\theta}}(D) = \arg \max_{\boldsymbol{\theta}} \{p(\boldsymbol{\theta} | D)\}$$

MAP-Schätzer

4.3. Bayes'sche Parameterschätzung

Ansatz 2: Konstante Gewichtung großer Fehler

Wegen der Monotonie der Logarithmusfunktion, kann man äquivalent die logarithmierte A Posteriori WDF maximieren; das bringt oft analytische Vorteile.

$$\hat{\boldsymbol{\theta}}(\mathbf{D}) = \arg \max_{\boldsymbol{\theta}} \{\ln p(\boldsymbol{\theta} | \mathbf{D})\}$$

MAP-Schätzer

Falls $\ln p(\boldsymbol{\theta} | \mathbf{m})$ differenzierbar ist, lautet die notwendige Bedingung für Maxima:

$$\nabla_{\boldsymbol{\theta}} \ln(p(\boldsymbol{\theta} | \mathbf{D})) = \sum_{k=1}^N \nabla_{\boldsymbol{\theta}} \ln p(\mathbf{m}_k | \boldsymbol{\theta}) + \nabla_{\boldsymbol{\theta}} \ln p(\boldsymbol{\theta}) \stackrel{!}{=} \mathbf{0}$$

In vielen praktisch relevanten Fällen liefern MAP-Schätzung und der A Posteriori Erwartungswert das gleiche Ergebnis.

Details siehe z.B. in: K. Kroschel, „Statistische Informationstechnik“, Springer 2004

4.3. Bayes'sche Klassifikation – ergänzende Bemerkungen

Fehler bei der Bayes'schen Klassifikation

- **Bayes'scher Fehler**: ergibt sich durch die Überlappung der Verteilungsdichten; spiegelt die Diskriminanz der Merkmale wider.
- **Modellfehler**: ergibt sich durch ein nicht passendes Modell. Auswahl eines passenden Modells mit Hilfe eines Anpassungstests für WVen, z.B. Chi-Quadrat-Test; Details in statistischer Grundlagenliteratur.
- **Schätzfehler**: ergibt sich aufgrund endlicher Datensätze.