

Peatükk 9

Osamudeli hindamisest

Vaatame lineaarset mudelit,

$$\mathbf{y} = \mathbf{X}\boldsymbol{\beta} + \varepsilon$$

mille parameetervektor on jagatud kaheks osaks $\boldsymbol{\beta}^T = (\boldsymbol{\beta}_1^T | \boldsymbol{\beta}_2^T)$. Samuti saame jagada ka mudelimaatriksi kaheks parameetervektori tükile vastavaks osaks, $\mathbf{X} = (\mathbf{X}_1 | \mathbf{X}_2)$, ning soovi korral võime algse mudeli kirja panna kujul

$$\mathbf{y} = \mathbf{X}_1\boldsymbol{\beta}_1 + \mathbf{X}_2\boldsymbol{\beta}_2 + \varepsilon,$$

Soovime leida, millised näevad välja hinnangud $\boldsymbol{\beta}_1$ -le ja $\boldsymbol{\beta}_2$ -le. Antud juhul teeme arutelu läbi eeldusel, et parameetervektor on hinnatav (ehk maatriks $\mathbf{X}^T\mathbf{X}$ on pööratav).

Esmalt juhime tähelepanu mõnele maatriksalgebra tulemusele. Blokkmaatriksi pöördmaatriksit on võimalik leida järgmise valemi abil:

$$\begin{pmatrix} \mathbf{A} & \mathbf{B} \\ \mathbf{C} & \mathbf{D} \end{pmatrix}^{-1} = \begin{pmatrix} (\mathbf{A} - \mathbf{B}\mathbf{D}^{-1}\mathbf{C})^{-1} & -(\mathbf{A} - \mathbf{B}\mathbf{D}^{-1}\mathbf{C})^{-1}\mathbf{B}\mathbf{D}^{-1} \\ -(\mathbf{D} - \mathbf{C}\mathbf{A}^{-1}\mathbf{B})^{-1}\mathbf{C}\mathbf{A}^{-1} & (\mathbf{D} - \mathbf{C}\mathbf{A}^{-1}\mathbf{B})^{-1} \end{pmatrix}.$$

Paneme kirja, milline näeb välja parameetervektori $\boldsymbol{\beta}$ hinnang:

$$\begin{aligned} \hat{\boldsymbol{\beta}} &= (\mathbf{X}^T\mathbf{X})^{-1}\mathbf{X}^T\mathbf{y} \\ &= \begin{pmatrix} \mathbf{X}_1^T\mathbf{X}_1 & \mathbf{X}_1^T\mathbf{X}_2 \\ \mathbf{X}_2^T\mathbf{X}_1 & \mathbf{X}_2^T\mathbf{X}_2 \end{pmatrix}^{-1} \begin{pmatrix} \mathbf{X}_1^T \\ \mathbf{X}_2^T \end{pmatrix} \mathbf{y} \end{aligned}$$

Leiame, milline näeb välja pöördmaatriks, kasutades ülaltoodud valemit blokkmaatriksi pöördmaatriksi leidmiseks:

$$\begin{aligned}
(\mathbf{A} - \mathbf{B}\mathbf{D}^{-1}\mathbf{C})^{-1} &= (\mathbf{X}_1^T \mathbf{X}_1 - \mathbf{X}_1^T \mathbf{X}_2 (\mathbf{X}_2^T \mathbf{X}_2)^{-1} \mathbf{X}_2^T \mathbf{X}_1)^{-1} \\
&= (\mathbf{X}_1^T (\mathbf{I} - \mathbf{P}_{\mathbf{X}_2}) \mathbf{X}_1)^{-1}
\end{aligned}$$

$$-(\mathbf{A} - \mathbf{B}\mathbf{D}^{-1}\mathbf{C})^{-1} \mathbf{B}\mathbf{D}^{-1} = (\mathbf{X}_1^T (\mathbf{I} - \mathbf{P}_{\mathbf{X}_2}) \mathbf{X}_1)^{-1} \mathbf{X}_1^T \mathbf{X}_2 (\mathbf{X}_2^T \mathbf{X}_2)^{-1}$$

$$\begin{aligned}
(\mathbf{D} - \mathbf{C}\mathbf{A}^{-1}\mathbf{B})^{-1} &= (\mathbf{X}_2^T \mathbf{X}_2 - \mathbf{X}_2^T \mathbf{X}_1 (\mathbf{X}_1^T \mathbf{X}_1)^{-1} \mathbf{X}_1^T \mathbf{X}_2)^{-1} \\
&= (\mathbf{X}_2^T (\mathbf{I} - \mathbf{P}_{\mathbf{X}_1}) \mathbf{X}_2)^{-1}
\end{aligned}$$

$$-(\mathbf{D} - \mathbf{C}\mathbf{A}^{-1}\mathbf{B})^{-1} \mathbf{C}\mathbf{A}^{-1} = (\mathbf{X}_2^T (\mathbf{I} - \mathbf{P}_{\mathbf{X}_1}) \mathbf{X}_2)^{-1} \mathbf{X}_2^T \mathbf{X}_1 (\mathbf{X}_1^T \mathbf{X}_1)^{-1}$$

ja lõplikuks hinnanguks saame

$$\begin{pmatrix} \hat{\beta}_1 \\ \hat{\beta}_2 \end{pmatrix} = \begin{pmatrix} (\mathbf{X}_1^T (\mathbf{I} - \mathbf{P}_{\mathbf{X}_2}) \mathbf{X}_1)^{-1} \mathbf{X}_1^T (\mathbf{I} - \mathbf{P}_{\mathbf{X}_2}) \mathbf{y} \\ (\mathbf{X}_2^T (\mathbf{I} - \mathbf{P}_{\mathbf{X}_1}) \mathbf{X}_2)^{-1} \mathbf{X}_2^T (\mathbf{I} - \mathbf{P}_{\mathbf{X}_1}) \mathbf{y} \end{pmatrix} \quad (9.1)$$

Interpretatsioon: parameetervektori esimese poole hinnangu saamiseks peaksime eemaldama muutujate $\mathbf{X}_2$ mõju muutjatest $\mathbf{X}_1$ ja seejärel hindama saadud jääkide mõju sõltuvale tunnusele.

Alljärgnevalt toome ära ühe arutluskäigu, mille tulemuste ekslik tõlgendamine viib sageli tõsiste eksimusteni reaalse andmete analüüsimisel. Uurime nimelt, kuna saame mudelimaatriksit $\mathbf{X}_1$ kasutades head (nihketa) hinnangud parameetritele β_1 (ja kuna peame tingimata kasutama ka teist poolt mudelimaatriksist ($\mathbf{X}_2$ -te). Täpsemalt: olgu $\mathbf{y} = \mathbf{X}_1\beta_1 + \mathbf{X}_2\beta_2 + \varepsilon$. Kui hindame parameetervektori β_1 kasutades lihtsamat mudelit ($\mathbf{y} = \mathbf{X}_1\beta_1 + \varepsilon$), siis millal saame nihketa hinnangu meid huvitavataele parameetritele (vaatamata sellele, et kasutasime vale mudelit)? Vaatame seda küsimust veidi lähemalt:

$$\begin{aligned}
\mathbb{E}((\mathbf{X}_1^T \mathbf{X}_1)^{-1} \mathbf{X}_1 \mathbf{y}) &= \mathbb{E}((\mathbf{X}_1^T \mathbf{X}_1)^{-1} \mathbf{X}_1 (\mathbf{X}_1 \beta_1 + \mathbf{X}_2 \beta_2 + \varepsilon)) \\
&= \beta_1 + (\mathbf{X}_1^T \mathbf{X}_1)^{-1} \mathbf{X}_1 \mathbf{X}_2 \beta_2.
\end{aligned}$$

Saadud avaldis võrdub β_1 -ga siis, kui $\mathbf{X}_2 \perp \mathbf{X}_1$ või kui $\beta_2 = 0$ (kui mudelist väljajäänud tunnused on kas ortogonaalsed — või sõltumatud — mudelisse sattunud tunnustest või kui mudelist väljajäänud tunnused tegelikult uuritava tunnuse keskväärtust ei mõjuta).

Sageli üritatakse saadud tulemust aga valesti ära kasutada. Nimelt üritatakse defineerida uut uuritavat tunnust, kust on eemaldatud nn teise blokki jäävate tunnuste mõju, $\mathbf{y}^* = \mathbf{y} - \mathbf{X}_2\beta_2$. Selliselt defineeritud tunnus ei sõltu enam mudelimaatriksisse $\mathbf{X}_2$ jäänud tunnuste mõjust ja seega võiksime uurida tunnuste $\mathbf{y}^*$ ja $\mathbf{X}_1$ vahelist seost ilma maatriksit $\mathbf{X}_2$ kasutamata. Reaalses elus pole aga β_2 teada. Praktikas theaksegi nüüd β_2 hindamisel viga — hinnatakse see vaid vektorit $\mathbf{y}$ ja maatriksit $\mathbf{X}_2$ kasutades ja saadakse seega nihkega hinnang β_2 -le, mis viib kogu edasise analüüsi omadega metsa.

Saadud tulemusest 9.1 võib olla kasu ka mõistmaks, miks ühe või teise parameetri hindamistäpsus on madal või aru saamaks mida tuleks teha saamaks võimalikult täpset parameetri hinnangut.

Vaatame juhtu, kui β_1 on vektor pikkusega 1 (tegemist on üheainsa parameetriga).

$$\begin{aligned} D(\hat{\beta}_1) &= (\mathbf{X}_1^T (\mathbf{I} - \mathbf{P}_{\mathbf{X}_2}) \mathbf{X}_1)^{-1} \sigma^2 \\ &= \frac{\sigma^2}{\mathbf{X}_1^T (\mathbf{I} - \mathbf{P}_{\mathbf{X}_2}) \mathbf{X}_1} \end{aligned}$$

Paneme tähele, et $SSE_{\mathbf{X}_1 \sim \mathbf{X}_2} := \mathbf{X}_1^T (\mathbf{I} - \mathbf{P}_{\mathbf{X}_2}) \mathbf{X}_1$ on jääkide ruutude summa mudelis, kus prognoositakse tunnuse $\mathbf{X}_1$ väärtust kasutades mudeli maatriksit $\mathbf{X}_2$. Samuti teame, et

$$\begin{aligned} R_{\mathbf{X}_1 \sim \mathbf{X}_2}^2 &= 1 - \frac{SSE_{\mathbf{X}_1 \sim \mathbf{X}_2}}{SSE_{\mathbf{X}_1 \sim 1}} \\ SSE_{\mathbf{X}_1 \sim \mathbf{X}_2} &= SSE_{\mathbf{X}_1 \sim 1} (1 - R_{\mathbf{X}_1 \sim \mathbf{X}_2}^2). \end{aligned}$$

Seega võime parameetri β_1 hinnangu dispersiooni kirja panna järgmisel kujul:

$$\begin{aligned} D(\hat{\beta}_1) &= \frac{\sigma^2}{SSE_{\mathbf{X}_1 \sim \mathbf{X}_2}} \\ &= \frac{\sigma^2}{SSE_{\mathbf{X}_1 \sim 1}} \cdot \frac{1}{1 - R_{\mathbf{X}_1 \sim \mathbf{X}_2}^2} \\ &= \frac{1}{n-1} \cdot \frac{\sigma^2}{\hat{\sigma}_{x_1}^2} \cdot \frac{1}{1 - R_{\mathbf{X}_1 \sim \mathbf{X}_2}^2}. \end{aligned}$$

Suurust $\frac{1}{1-R_{\mathbf{x}_1 \sim \mathbf{x}_2}^2}$ tuntakse hinnangu dispersiooni puhitusteguri nime all (*Variance Inflation Factor, VIF*) ja ta iseloomustab seda, kuivõrd kaotame hinnangu täpsuses seetõttu, et kasutame mitteortogonaalseid (omavahel korreleeritud) sõltuvaid tunnuseid mudeli hindamisel.