

Lognormal regression – three views

1. Index form

The per-observation reading: what happens for each row i of your data.

$$\begin{aligned} y_i \mid \mu_i, \sigma_i &\sim \text{Lognormal}(\mu_i, \sigma_i^2) \iff \log(y_i) \mid \mu_i, \sigma_i \sim \text{Normal}(\mu_i, \sigma_i^2) \\ \mu_i &= \beta_0 + \beta_1 x_i \\ \log(\sigma_i) &= \gamma_0 \end{aligned}$$

Coefficient reading. On the response scale, $\exp(\hat{\beta})$ is the multiplicative effect on the geometric mean (the median) of the response: a one-unit increase multiplies the median by $\exp(\hat{\beta})$.

2. Matrix form

The same model in matrix notation. The structural contract every textbook past chapter 4 switches to.

$$\begin{aligned} \log(\mathbf{y}) \mid \boldsymbol{\mu}, \boldsymbol{\sigma} &\sim \mathcal{N}(\boldsymbol{\mu}, \text{diag}(\boldsymbol{\sigma}^2)) \\ \boldsymbol{\mu} &= \mathbf{X}\boldsymbol{\beta} \\ \log(\boldsymbol{\sigma}) &= \mathbf{X}_\sigma \boldsymbol{\gamma} \end{aligned}$$

3. Worked observation ($i = 1$)

The index-form equations evaluated at the first observation of your data, with coefficient estimates plugged in.

$$\hat{\mu}_1 = 2.02 - 0.0136 \cdot 0.409 \approx 2.01$$

Observed $\log y_1 = 1.56$ (log scale); predicted log-scale mean $\hat{\mu}_1 = 2.01$; log-scale residual $\hat{\varepsilon}_1^{(\log)} = -0.448$ (observed = mean + residual). Back-transform to the response-scale mean with $\mathbb{E}[y_1] = \exp(\hat{\mu}_1 + \hat{\sigma}_1^2/2)$.

And the σ submodel at $i = 1$:

$$\log \hat{\sigma}_1 = -0.764 \Rightarrow \hat{\sigma}_1 = \exp(-0.764) \approx 0.466$$

Stacking the same response equation for all $n = 100$ observations:

$$\underbrace{\begin{bmatrix} 1.56 \\ 2.07 \\ 1.38 \\ 2.54 \\ 1.84 \\ \vdots \\ 1.8 \\ 1.71 \end{bmatrix}}_{\log(\mathbf{y})_{100 \times 1} \text{ (observed, log scale)}} = \underbrace{\begin{bmatrix} 1 & 0.409 \\ 1 & 1.69 \\ 1 & 1.59 \\ 1 & -0.331 \\ 1 & -2.29 \\ \vdots & \vdots \\ 1 & -0.0506 \\ 1 & -0.306 \end{bmatrix}}_{\mathbf{X}_{100 \times 2}} \underbrace{\begin{bmatrix} 2.02 \\ -0.0136 \end{bmatrix}}_{\hat{\boldsymbol{\beta}}_{2 \times 1} \text{ (estimated)}} + \underbrace{\begin{bmatrix} -0.448 \\ 0.0769 \\ -0.611 \\ 0.519 \\ -0.212 \\ \vdots \\ -0.221 \\ -0.317 \end{bmatrix}}_{\hat{\boldsymbol{\varepsilon}}_{100 \times 1}^{(\log)} \text{ (log-scale residual)}}$$

And the σ submodel (no observed counterpart – σ 's job is to describe the spread of the log-scale residual $\hat{\varepsilon}^{(\log)}$, i.e. the SD of $\log y$):

$$\log \underbrace{\begin{bmatrix} 0.466 \\ 0.466 \\ 0.466 \\ 0.466 \\ 0.466 \\ \vdots \\ 0.466 \\ 0.466 \end{bmatrix}}_{\boldsymbol{\sigma}_{100 \times 1}} = \underbrace{\begin{bmatrix} 1 \\ 1 \\ 1 \\ 1 \\ 1 \\ \vdots \\ 1 \\ 1 \end{bmatrix}}_{\mathbf{X}_{\sigma, 100 \times 1}} \underbrace{\begin{bmatrix} -0.764 \end{bmatrix}}_{\boldsymbol{\gamma}_{1 \times 1}}$$