

# Cheat Sheet for Probability Theory

Ingmar Land

## 1 Scalar-valued Random Variables

Consider two real-valued random variables (RV)  $X$  and  $Y$  with the individual probability distributions  $p_X(x)$  and  $p_Y(y)$ , and the joint distribution  $p_{X,Y}(x, y)$ . The probability distributions are probability mass functions (pmf) if the random variables take discrete values, and they are probability density functions (pdf) if the random variables are continuous. Some authors use  $f()$  instead of  $p()$ , especially for continuous RVs.

In the following, the RVs are assumed to be continuous. (For discrete RVs, the integrals have simply to be replaced by sums.)

- Marginal distributions:

$$p_X(x) = \int p_{X,Y}(x, y) \, dy \qquad p_Y(y) = \int p_{X,Y}(x, y) \, dx$$

- Conditional distributions:

$$p_{X|Y}(x|y) = \frac{p_{X,Y}(x, y)}{p_Y(y)} \qquad p_{Y|X}(y|x) = \frac{p_{X,Y}(x, y)}{p_X(x)}$$

for  $p_X(x) \neq 0$  and  $p_Y(y) \neq 0$

- Bayes' rule:

$$p_{X|Y}(x|y) = \frac{p_{X,Y}(x, y)}{\int p_{X,Y}(x', y) \, dx'} \qquad p_{Y|X}(y|x) = \frac{p_{X,Y}(x, y)}{\int p_{X,Y}(x, y') \, dy'}$$

- Expected values (expectations):

$$\begin{aligned} E[g_1(X)] &:= \int g_1(x) p_X(x) \, dx \\ E[g_2(Y)] &:= \int g_2(y) p_Y(y) \, dy \\ E[g_3(X, Y)] &:= \int g_3(x, y) p_{X,Y}(x, y) \, dx \, dy \end{aligned}$$

for any functions  $g_1(\cdot)$ ,  $g_2(\cdot)$ ,  $g_3(\cdot, \cdot)$

- Some special expected values:

- Means (mean values):

$$\mu_X := E[X] = \int x p_X(x) dx \quad \mu_Y := E[Y] = \int y p_Y(y) dy$$

- Variances:

$$\begin{aligned} \sigma_X^2 &\equiv \Sigma_{XX} := E[(X - \mu_X)^2] = \int (x - \mu_X)^2 p_X(x) dx \\ \sigma_Y^2 &\equiv \Sigma_{YY} := E[(Y - \mu_Y)^2] = \int (y - \mu_Y)^2 p_Y(y) dy \end{aligned}$$

Remark: The variance measures the “width” of a distribution. A small variance means that most of the probability mass is concentrated around the mean value.

- Covariance:

$$\begin{aligned} \sigma_{XY} &\equiv \Sigma_{XY} := E[(X - \mu_X)(Y - \mu_Y)] \\ &= \int (x - \mu_X)(y - \mu_Y) p_{X,Y}(x, y) dx dy \end{aligned}$$

Remark: The covariance measures how “related” two RVs are. Two independent RVs have covariance zero.

- Correlation coefficient:

$$\rho_{XY} := \frac{\sigma_{XY}}{\sigma_X \sigma_Y}$$

- Relations:

$$\begin{aligned} E[X^2] &= \Sigma_{XX} + \mu_X^2 \\ E[Y^2] &= \Sigma_{YY} + \mu_Y^2 \\ E[X \cdot Y] &= \Sigma_{XY} + \mu_X \cdot \mu_Y \end{aligned}$$

- Proof of last relation:

$$\begin{aligned} E[XY] &= E[((X - \mu_X) + \mu_X)((Y - \mu_Y) + \mu_Y)] \\ &= E[(X - \mu_X)(Y - \mu_Y)] - E[(X - \mu_X)\mu_Y] - E[\mu_X(Y - \mu_Y)] + \\ &\quad + E[\mu_X\mu_Y] \\ &= \Sigma_{XY} - (E[X] - \mu_X)\mu_Y - (E[Y] - \mu_Y)\mu_X + \mu_X\mu_Y \\ &= \Sigma_{XY} + \mu_X\mu_Y \end{aligned}$$

This method of proof is typical.

- Conditional expectations:

$$\mathbb{E}[g(X)|Y = y] := \int g(x)p_{X|Y}(x|y) \, dx$$

Law of total expectation:

$$\begin{aligned} \mathbb{E}[\mathbb{E}[g(X)|Y]] &= \int \mathbb{E}[g(X)|Y = y]p_Y(y) \, dy \\ &= \int \int g(x)p_{X|Y}(x|y) \, dx p_Y(y) \, dy \\ &= \int g(x) \int p_{X|Y}(x|y)p_Y(y) \, dy \, dx \\ &= \int g(x)p_X(x) \, dx \\ &= \mathbb{E}[g(X)] \end{aligned}$$

- Special conditional expectations:

– Conditional mean:

$$\mu_{X|Y=y} := \mathbb{E}[X|Y = y] = \int xp_{X|Y}(x|y) \, dx$$

– Conditional variance:

$$\Sigma_{X|Y=y} := \mathbb{E}[(X - \mu_X)^2|Y = y] = \int (x - \mu_X)^2 p_{X|Y}(x|y) \, dx$$

– Relation:

$$\mathbb{E}[X^2|Y = y] = \Sigma_{X|Y=y} + \mu_{X|Y=y}^2$$

- Sum of two random variables: Let  $Z := X + Y$ ; then

$$p_Z(z) = p_X(z) * p_Y(z)$$

where  $*$  denotes convolution. The proof uses the characteristic functions.

- The RVs are called independent if

$$p_{X,Y}(x, y) = p_X(x) \cdot p_Y(y).$$

This condition is equivalent to the condition that

$$\mathbb{E}[g_1(X) \cdot g_2(Y)] = \mathbb{E}[g_1(X)] \cdot \mathbb{E}[g_2(Y)]$$

for all (!) functions  $g_1(\cdot)$  and  $g_2(\cdot)$ .

- The RVs are called uncorrelated if

$$\sigma_{XY} \equiv \Sigma_{XY} = \mathbb{E}[(X - \mu_X)(Y - \mu_Y)] = 0.$$

Remark: If RVs are independent, they are also uncorrelated. The reverse holds only for Gaussian RVs (see below).

- Two RVs  $X$  and  $Y$  are called orthogonal if  $\mathbb{E}[XY] = 0$ .

Remark: The RVs with finite energy,  $\mathbb{E}[X^2] < \infty$ , form a vector space with scalar product  $\langle X, Y \rangle = \mathbb{E}[XY]$  and norm  $\|X\| = \sqrt{\mathbb{E}[X^2]}$ . (This is used in MMSE estimation.)

These relations for scalar-valued RVs are generalized to vector-valued RVs in the following.

## 2 Vector-valued Random Variables

Consider two real-valued vector-valued random variables (RV)

$$\underline{X} = \begin{bmatrix} X_1 \\ X_2 \end{bmatrix}, \quad \underline{Y} = \begin{bmatrix} Y_1 \\ Y_2 \end{bmatrix}$$

with the individual probability distributions  $p_{\underline{X}}(\underline{x})$  and  $p_{\underline{Y}}(\underline{y})$ , and the joint distribution  $p_{\underline{X}, \underline{Y}}(\underline{x}, \underline{y})$ . (The following considerations can be generalized to longer vectors, of course.) The probability distributions are probability mass functions (pmf) if the random variables take discrete values, and they are probability density functions (pdf) if the random variables are continuous. Some authors use  $f()$  instead of  $p()$ , especially for continuous RVs.

In the following, the RVs are assumed to be continuous. (For discrete RVs, the integrals have simply to be replaced by sums.)

Remark: The following matrix notations may seem to be cumbersome at the first glance, but they turn out to be quite handy and convenient (once you got used to).

- Marginal distributions, conditional distributions, Bayes' rule, expected values work as in the scalar case.
- Some special expected values:
  - Mean vector (vector of mean values):

$$\underline{\mu}_{\underline{X}} := \mathbb{E}[\underline{X}] = \mathbb{E} \begin{bmatrix} X_1 \\ X_2 \end{bmatrix} = \begin{bmatrix} \mathbb{E}[X_1] \\ \mathbb{E}[X_2] \end{bmatrix} = \begin{bmatrix} \mu_{X_1} \\ \mu_{X_2} \end{bmatrix}$$

- Covariance matrix (auto-covariance matrix):

$$\begin{aligned}
 \Sigma_{\underline{X}\underline{X}} &:= \mathbb{E}[(\underline{X} - \underline{\mu}_X)(\underline{X} - \underline{\mu}_X)^\top] \\
 &= \mathbb{E} \left[ \begin{bmatrix} X_1 - \mu_1 \\ X_2 - \mu_2 \end{bmatrix} \begin{bmatrix} X_1 - \mu_1 & X_2 - \mu_2 \end{bmatrix} \right] \\
 &= \begin{bmatrix} \mathbb{E}[(X_1 - \mu_1)(X_1 - \mu_1)] & \mathbb{E}[(X_1 - \mu_1)(X_2 - \mu_2)] \\ \mathbb{E}[(X_2 - \mu_2)(X_1 - \mu_1)] & \mathbb{E}[(X_2 - \mu_2)(X_2 - \mu_2)] \end{bmatrix} \\
 &= \begin{bmatrix} \Sigma_{X_1 X_1} & \Sigma_{X_1 X_2} \\ \Sigma_{X_2 X_1} & \Sigma_{X_2 X_2} \end{bmatrix}
 \end{aligned}$$

- Covariance matrix (cross-covariance matrix):

$$\begin{aligned}
 \Sigma_{\underline{X}\underline{Y}} &:= \mathbb{E}[(\underline{X} - \underline{\mu}_X)(\underline{Y} - \underline{\mu}_Y)^\top] \\
 &= \begin{bmatrix} \Sigma_{X_1 Y_1} & \Sigma_{X_1 Y_2} \\ \Sigma_{X_2 Y_1} & \Sigma_{X_2 Y_2} \end{bmatrix}
 \end{aligned}$$

Remark: This matrix contains the covariance of each element of the first vector with each element of the second vector.

- Relations:

$$\begin{aligned}
 \mathbb{E}[\underline{X}\underline{X}^\top] &= \Sigma_{\underline{X}\underline{X}} + \underline{\mu}_X \underline{\mu}_X^\top \\
 \mathbb{E}[\underline{X}\underline{Y}^\top] &= \Sigma_{\underline{X}\underline{Y}} + \underline{\mu}_X \underline{\mu}_Y^\top
 \end{aligned}$$

Remark: This result is not too surprising when you know the result for the scalar case.

### 3 Gaussian Random Variables

- A Gaussian RV  $X$  with mean  $\mu_X$  and variance  $\sigma_X^2$  is a continuous random variable with a Gaussian pdf, i.e., with

$$p_X(x) = \frac{1}{\sqrt{2\pi\sigma_X^2}} \cdot \exp\left(-\frac{(x - \mu_X)^2}{2\sigma_X^2}\right)$$

The often used symbolic notation

$$X \sim \mathcal{N}(\mu_X, \sigma_X^2)$$

may be read as:  $X$  is (distributed) Gaussian with mean  $\mu_X$  and variance  $\sigma_X^2$ .

- A Gaussian distribution with mean zero and variance one is called a normal distribution:

$$p(x) = \frac{1}{\sqrt{2\pi}} e^{-\frac{x^2}{2}}.$$

Some authors use the term “normal distribution” equivalently with “Gaussian distribution”. So, use the term “normal distribution” with caution.

The integral of a normal pdf cannot be solved in closed form and is therefore often expressed using the Q-function

$$Q(z) := \frac{1}{\sqrt{2\pi}} \int_z^\infty e^{-\frac{x^2}{2}} dx$$

Notice the limits of the integral. The integral of any Gaussian pdf can also be expressed using the Q-function, of course.

- A vector-valued RV  $\underline{X}$  with mean  $\underline{\mu}_X$  and covariance matrix  $\Sigma_{\underline{X}\underline{X}}$ ,

$$\underline{X} = \begin{bmatrix} X_1 \\ X_2 \end{bmatrix}, \quad \underline{\mu}_X = \begin{bmatrix} \mu_{X_1} \\ \mu_{X_2} \end{bmatrix}, \quad \Sigma_{\underline{X}\underline{X}} = \begin{bmatrix} \Sigma_{X_1 X_1} & \Sigma_{X_1 X_2} \\ \Sigma_{X_2 X_1} & \Sigma_{X_2 X_2} \end{bmatrix},$$

is called Gaussian if its components have a jointly Gaussian pdf

$$p_{\underline{X}}(\underline{x}) = \frac{1}{\sqrt{2\pi}^2 |\Sigma_{\underline{X}\underline{X}}|^{1/2}} \cdot \exp\left(-\frac{1}{2}(\underline{X} - \underline{\mu}_X)^T \Sigma_{\underline{X}\underline{X}} (\underline{X} - \underline{\mu}_X)\right).$$

(There is nothing to understand, this is just a definition.) The marginal distributions and the conditional distributions of a Gaussian vector are again Gaussian distributions. (But this can be proved.)

The corresponding symbolic notation is

$$\underline{X} \sim \mathcal{N}(\underline{\mu}_X, \Sigma_{\underline{X}\underline{X}})$$

This can be generalized to longer vectors, of course.

- Gaussian RVs are completely described by mean and covariance. Therefore, if they are uncorrelated, they are also independent. This holds only for Gaussian RVs. (Cf. above.)