

LECTURE 5: PROPERTIES OF A RANDOM SAMPLE

MECO 7312.

INSTRUCTOR: DR. KHAI CHIONG

OCTOBER 7, 2025

1. Random Samples

Suppose n number of observations are randomly sampled from the density $f(x)$. The first observation is X_1 , the second observation X_2 , and so on. The outcome is a *random sample* described by the random variables $X_1, \dots, X_n$.

Random sampling is the default assumption that we maintained throughout.¹ It means that $X_1, \dots, X_n$ are independent and identically distributed (i.i.d) with the pdf $f(x)$. That is, the density of each $X_1, \dots, X_n$ is the same and equals $f(x)$. Moreover $X_1, \dots, X_n$ are mutually independent.

The joint pdf of $X_1, \dots, X_n$ is:

$$f(x_1, \dots, x_n) = f(x_1) \cdots f(x_n) = \prod_{i=1}^n f(x_i)$$

1.1. Statistic

Let $X_1, \dots, X_n$ be a random sample from a population. Let $T(X_1, \dots, X_n)$ be a function that maps the sample space of $(X_1, \dots, X_n)$ into $\mathbb{R}$. Then $Y = T(X_1, \dots, X_n)$ is called a statistic. A statistic is a random variable. The pdf of $Y = T(X_1, \dots, X_n)$ is called the *sampling distribution* of Y .

For example, the *sample mean* is the statistic defined by:

$$\bar{X} = \frac{X_1 + \cdots + X_n}{n}$$

The *sample variance* is the statistic defined by:

¹If our survey and experiment is not randomized correctly, then random sampling is violated. For instance, when a certain type of individuals opt out of the survey. The population also needs to be large enough, so that surveying the first household (and hence “removing” this household from the population to be surveyed next) has no effect on the population.

$$S^2 = \frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2$$

1.2. Unbiased estimator

A statistic $T(X_1, \dots, X_n)$ is an unbiased estimator of a population parameter θ if $\mathbb{E}[T(X_1, \dots, X_n)] = \theta$.

Let $\mathbb{E}[X_i] = \mu$, then $\bar{X}$ is an unbiased estimator (or statistic) of μ .

$$\begin{aligned} \mathbb{E}[\bar{X}] &= \mathbb{E}\left[\frac{X_1 + \dots + X_n}{n}\right] \\ &= \frac{1}{n} \mathbb{E}\left[\sum_{i=1}^n X_i\right] \\ &= \frac{1}{n} \sum_{i=1}^n \mathbb{E}[X_i] \\ &= \frac{1}{n} n\mu = \mu \end{aligned}$$

Moreover, the statistic $\bar{X}$ has a variance of $\frac{\sigma^2}{n}$, where $\text{Var}(X_i) = \sigma^2$.

$$\begin{aligned} \text{Var}[\bar{X}] &= \text{Var}\left[\frac{X_1 + \dots + X_n}{n}\right] \\ &= \frac{1}{n^2} \sum_{i=1}^n \text{Var}[X_i] \\ &= \frac{1}{n^2} n\sigma^2 = \frac{\sigma^2}{n} \end{aligned}$$

The sample variance S^2 is an unbiased estimator of σ^2 .

$$\begin{aligned}
\mathbb{E}[S^2] &= \mathbb{E} \left[\frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2 \right] \\
&= \frac{1}{n-1} \mathbb{E} \left[\sum_{i=1}^n X_i^2 - 2\bar{X} \sum_{i=1}^n X_i + \sum_{i=1}^n \bar{X}^2 \right] \\
&= \frac{1}{n-1} \mathbb{E} \left[\sum_{i=1}^n X_i^2 - 2n\bar{X}^2 + n\bar{X}^2 \right] \quad \text{using } \sum X_i = n\bar{X} \\
&= \frac{1}{n-1} (n \mathbb{E}[X_i^2] - n \mathbb{E}[\bar{X}^2])
\end{aligned}$$

Since we have $\mathbb{E}[X_i^2] = \text{Var}(X_i) + \mathbb{E}[X_i]^2 = \sigma^2 + \mu^2$, and $\mathbb{E}[\bar{X}^2] = \text{Var}(\bar{X}) + \mathbb{E}[\bar{X}]^2 = \sigma^2/n + \mu^2$,

$$\mathbb{E}[S^2] = \frac{n}{n-1}(\sigma^2 + \mu^2) - \frac{n}{n-1} \left(\frac{\sigma^2}{n} + \mu^2 \right) = \sigma^2$$

The bias of a statistics with respect to the population parameter θ is $\mathbb{E}[T(X_1, \dots, X_n)] - \theta$. For instance, if we were to use the sample variance formula $\hat{\sigma}^2 = \frac{1}{n} \sum_{i=1}^n (X_i - \bar{X})^2$, then:

$$\begin{aligned}
\mathbb{E}[\hat{\sigma}^2] &= \mathbb{E} \left[\frac{1}{n} \sum_{i=1}^n (X_i - \bar{X})^2 \right] \\
&= \frac{n-1}{n} \sigma^2
\end{aligned}$$

Therefore the bias in using the statistic $\hat{\sigma}^2 = \frac{1}{n} \sum_{i=1}^n (X_i - \bar{X})^2$ in estimating the population parameter σ^2 is $\frac{1}{n}\sigma^2$, which diminishes as the sample size, n increases.

2. Sampling distribution

Let $X_1, \dots, X_n$ be a random sample from a $N(\mu, \sigma^2)$ distribution. Consider the sample mean $\bar{X} = \frac{X_1 + \dots + X_n}{n}$, which is a random variable. The sampling distribution of $\bar{X}$ is Normally distributed as $\mathcal{N}(\mu, \frac{\sigma^2}{n})$. In another words, the pdf of $\bar{X}$ is $\mathcal{N}(\mu, \frac{\sigma^2}{n})$. This can be proven using moment generating functions.

Sampling distribution is important because it describes the distribution of an estimator (or statistic) due to its' sampling variation. For example, the sample mean

we calculated from a sample is a noisy estimate of the true population mean, and through the sampling distribution we get a sense of how noisy this estimate is.

On the other hand, for the sample variance, S^2 , the sampling distribution of $(n-1)S^2/\sigma^2$ is a chi squared distribution with $n-1$ degrees of freedom. Refer to the companion code for how we can use simulation to obtain the sampling distribution.

2.1. Chi squared distribution

Notation: χ_p^2 denotes a chi squared distribution with p degrees of freedom.

If X is a $N(0, 1)$ random variable, then $X^2 \sim \chi_1^2$. If $X_1, \dots, X_n$ are independent $N(0, 1)$, then $Z = X_1^2 + \dots + X_n^2$ is distributed according to χ_n^2 . The pdf of χ_p^2 is $f(x) = \frac{1}{\Gamma(p/2)2^{p/2}} x^{\frac{p}{2}-1} e^{-\frac{x}{2}}$, for $x > 0$. Recall that $\Gamma(\cdot)$ is the Gamma function. In fact, χ_p^2 is a special case of the Gamma distribution.

The mean of χ_p^2 is p , while the variance of χ_p^2 is $2p$. Because $(n-1)S^2/\sigma^2$ is a chi-squared distribution with $n-1$ degrees of freedom, we have:

$$S^2 \sim \frac{\sigma^2}{n-1} \chi_{n-1}^2$$

As such, $\mathbb{E}[S^2] = \sigma^2$ and $\text{Var}(S^2) = 2\sigma^4/(n-1)$.

To see intuitively why the sample variance has a χ_{n-1}^2 distribution, notice that,

$$\frac{(n-1)S^2}{\sigma^2} = \sum_{i=1}^n \left(\frac{X_i - \bar{X}}{\sigma} \right)^2$$

If we know that the population mean is μ , then we have $\frac{X_i - \mu}{\sigma} \sim N(0, 1)$, and as such $\sum_{i=1}^n \left(\frac{X_i - \mu}{\sigma} \right)^2 \sim \chi_n^2$. Since we do not know μ and we estimate μ by $\frac{X_1 + \dots + X_n}{n}$, we lose one degree of freedom.

2.2. Student's t distribution

Recall that the sampling distribution of $\bar{X}$ is $N(\mu, \frac{\sigma^2}{n})$. Therefore, $\frac{\bar{X} - \mu}{\sigma/\sqrt{n}} \sim N(0, 1)$.

What if we replace σ with the sample standard deviation S ? What is the sampling distribution of $\frac{\bar{X} - \mu}{S/\sqrt{n}}$, where S is the sample standard deviation?

Let $X_1, \dots, X_n$ be a random sample from a $N(\mu, \sigma^2)$ distribution. The statistic $\frac{\bar{X} - \mu}{S/\sqrt{n}}$ has *Student's t distribution with $n - 1$ degrees of freedom*.²

Let us sketch the derivation of the pdf of the Student's t distribution. By multiplying both the numerator and denominator by σ , we get

$$(1) \quad \frac{\bar{X} - \mu}{S/\sqrt{n}} = \frac{(\bar{X} - \mu)/(\sigma/\sqrt{n})}{\sqrt{S^2/\sigma^2}}$$

The numerator $(\bar{X} - \mu)/(\sigma/\sqrt{n})$ is distributed $N(0, 1)$. While the denominator $\sqrt{S^2/\sigma^2}$ is distributed as $\sqrt{V/(n-1)}$ where $V \sim \chi_{n-1}^2$. Moreover, the sample variance and mean are independent (see Theorem 5.3.1 in Casella-Berger). Therefore, the Student's t distribution with $n - 1$ degrees of freedom is distributed as $\frac{U}{\sqrt{V/n-1}}$, where $U \sim N(0, 1)$ and $V \sim \chi_{n-1}^2$, with U and V independent. The pdf of the Student's t distribution can then be obtained by applying the bivariate change-of-variables formula.

The t-distribution is symmetric and bell-shaped, like the normal distribution, but has heavier tails, meaning that it is more prone to producing values that fall far from its mean. As the degree of freedom $n - 1$ increases, the t-distribution converges in distribution to the standard Normal.

The pdf of the t-distribution with degree of freedom v is:

$$f_v(x) = \frac{\left(\frac{v}{v+x^2}\right)^{\frac{v+1}{2}}}{\sqrt{v} \text{Beta}\left(\frac{v}{2}, \frac{1}{2}\right)}$$

It has mean (median, mode) of zero, but its variance is $\frac{v}{v-2}$.

3. Order Statistics

Let $X_1, \dots, X_n$ be a random sample from a population with distribution $f(x)$.

Let $X_{(n)}$ be the largest value from the sample $X_1, \dots, X_n$. That is, $X_{(n)} = \max_{1 \leq i \leq n} X_i = \max(X_1, \dots, X_n)$. The random variable $X_{(n)}$ is called the largest order statistics. While $X_{(1)} = \min(X_1, \dots, X_n)$ is called the first-orders statistics.

What is the distributions of $X_{(1)}$ and $X_{(n)}$? Order statistics plays a big role in Auction Theory. In a second-price auction, the participant with the largest bid wins the auction but pays the second-largest bid. If bids are assumed to be realized

²William Gosset (under the pseudonym of Student) in the 1900s proposed substituting the unknown σ^2 with the sample variance S^2 , and derived the resulting sampling distribution.

i.i.d from a distribution, then the revenue from the auction is distributed as the second-largest order statistics.

To obtain the distributions of the order statistics, we start with the cdf.

$$\begin{aligned} P(X_{(n)} \leq x) &= P(\max(X_1, \dots, X_n) \leq x) \\ &= \text{Probability that all } n \text{ observations are smaller than } x \\ &= F(x)^n \end{aligned}$$

Hence the cdf of $X_{(n)}$ is $F(x)^n$, and the pdf is $nF(x)^{n-1}f(x)$.

The smallest-order statistic is:

$$\begin{aligned} P(X_{(1)} \leq x) &= P(\min(X_1, \dots, X_n) \leq x) \\ &= 1 - P(\min(X_1, \dots, X_n) > x) \\ &= 1 - \text{Probability that all } n \text{ observations are larger than } x \\ &= 1 - (1 - F(x))^n \end{aligned}$$

Which leads to the density $n(1 - F(x))^{n-1}f(x)$.

In general, the density of $X_{(k)}$ can be derived as follows:

$$\begin{aligned} P(X_{(n-1)} \leq x) &= \text{Probability that exactly } n-1 \text{ observations are smaller than } x \\ &\quad + \text{Probability that exactly } n \text{ observations are smaller than } x \\ &= \binom{n}{n-1} F(x)^{n-1} (1 - F(x)) + F(x)^n \end{aligned}$$

$$\begin{aligned} P(X_{(n-2)} \leq x) &= \text{Probability that exactly } n-2 \text{ observations are smaller than } x \\ &\quad + \text{Probability that exactly } n-1 \text{ observations are smaller than } x \\ &\quad + \text{Probability that exactly } n \text{ observations are smaller than } x \\ &= \binom{n}{n-2} F(x)^{n-2} (1 - F(x))^2 + \binom{n}{n-1} F(x)^{n-1} (1 - F(x)) + F(x)^n \end{aligned}$$

3.1. Order statistic for the Uniform distribution

Suppose $X_1, \dots, X_n$ is a random sample from the Uniform distribution $U[0, 1]$.

The largest-order statistic has the pdf $f(x) = nx^{n-1}$ with the support on $x \in (0, 1)$. Further, we can calculate the expectation with respect to this density, which is $\mathbb{E}[X_{(n)}] = \frac{n}{n+1}$. Therefore, when $n = 20$, the maximum out of n uniform random values is $20/21 \approx 0.952$, *on average*. In fact, $f(x) = nx^{n-1}$ is the Beta distribution with parameters $\alpha = n$ and $\beta = 1$. In general, the k -th order statistic of a Uniform distribution is $Beta(k, n + 1 - k)$.

As an exercise, derive the expectation of the largest-order statistic of the Standard Normal distribution in terms of n , the sample size.

4. Convergence and consistency

4.1. Convergence in probability

A sequence of random variables $(Y_n)_{n \geq 1}$ *converges in probability* to a constant θ if, for every $\varepsilon > 0$,

$$\lim_{n \rightarrow \infty} \Pr(|Y_n - \theta| \geq \varepsilon) = 0,$$

equivalently, $\Pr(|Y_n - \theta| < \varepsilon) \rightarrow 1$. In words, the probability that Y_n is arbitrarily close to θ converges to one as n increases. We write $Y_n \xrightarrow{p} \theta$.

A statistic $Y_n = T(X_1, \dots, X_n)$ is a *consistent* estimator of θ if $Y_n \xrightarrow{p} \theta$ as $n \rightarrow \infty$. Equivalently, $Y_n = \theta + o_p(1)$.³

The Weak Law of Large Numbers (WLLN) state that under some assumptions, the sample mean is a consistent estimator of the population mean. Let $X_1, \dots, X_n$ be i.i.d. with mean μ and variance $0 < \sigma^2 < \infty$. The sample mean

$$\bar{X}_n = \frac{1}{n} \sum_{i=1}^n X_i$$

satisfies, for any $\varepsilon > 0$,

$$\Pr(|\bar{X}_n - \mu| \geq \varepsilon) \leq \frac{\text{Var}(\bar{X}_n)}{\varepsilon^2} = \frac{\sigma^2}{n \varepsilon^2} \xrightarrow[n \rightarrow \infty]{} 0.$$

Due to Chebyshev's. Thus $\bar{X}_n \xrightarrow{p} \mu$. As the sample size becomes larger, the probability that the sample mean deviates from μ by ε decreases to zero. The sequence defined by $p_n = P(|\bar{X}_n - \mu| \geq \varepsilon)$ converges to 0 as the sample size increases. Moreover, $\mathbb{E}[\bar{X}_n] = \mu$, so $\bar{X}_n$ is both unbiased and consistent for μ .

³We say that $Y_n = o_p(1)$ iff $Y_n \xrightarrow{p} 0$; $Y_n = o_p(a_n)$ means $Y_n/a_n \xrightarrow{p} 0$.

4.2. Sufficient conditions for consistency

Proposition 1 (MSE criterion). *If $\text{MSE}(Y_n) \equiv \mathbb{E}[(Y_n - \theta)^2] \rightarrow 0$, then $Y_n \xrightarrow{p} \theta$.*

Proof. By Chebyshev applied to $Y_n - \theta$,

$$\Pr(|Y_n - \theta| \geq \varepsilon) \leq \frac{\mathbb{E}[(Y_n - \theta)^2]}{\varepsilon^2} = \frac{\text{MSE}(Y_n)}{\varepsilon^2} \rightarrow 0. \quad \square$$

Using the bias–variance decomposition,

$$\text{MSE}(Y_n) = (\mathbb{E}[Y_n] - \theta)^2 + \text{Var}(Y_n),$$

we immediately obtain the common corollaries:

Corollary 1 (Asymptotic unbiasedness + vanishing variance): If $\mathbb{E}[Y_n] \rightarrow \theta$ (*bias* $\rightarrow 0$) and $\text{Var}(Y_n) \rightarrow 0$, then Y_n is consistent.

Corollary 2 (Unbiased + vanishing variance): If $\mathbb{E}[Y_n] = \theta$ for all n and $\text{Var}(Y_n) \rightarrow 0$, then Y_n is consistent.

Let $X_1, \dots, X_n$ be i.i.d. with mean μ and variance $0 < \sigma^2 < \infty$. The Bessel-corrected sample variance $S^2 = \frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X}_n)^2$, satisfies $\mathbb{E}[S^2] = \sigma^2$. If the fourth central moment $\mu_4 = \mathbb{E}[(X - \mu)^4]$ is finite, then

$$\text{Var}(S^2) = \frac{1}{n} \left(\mu_4 - \frac{n-3}{n-1} \sigma^4 \right) \xrightarrow{n \rightarrow \infty} 0,$$

so S^2 is unbiased and consistent for σ^2 . Recall that if $X_1, \dots, X_n$ be i.i.d. Gaussian, then $\text{Var}(S^2) = 2\sigma^4/(n-1)$.⁴

We have seen that the sample mean is both a consistent and an unbiased estimator of the population mean. In general, unbiasedness and consistency are two different concepts concerning the property of an estimator. An estimator can be biased but consistent, and vice versa. For example,

- *Biased but consistent.* The biased sample variance with denominator n ,

$$\tilde{S}^2 = \frac{1}{n} \sum_{i=1}^n (X_i - \bar{X}_n)^2,$$

satisfies $\mathbb{E}[\tilde{S}^2] = \frac{n-1}{n} \sigma^2 \neq \sigma^2$, but $\tilde{S}^2 \xrightarrow{p} \sigma^2$ because both its bias and variance vanish.

⁴The fourth central moment of a normal distribution with mean μ and variance σ^2 is $3\sigma^4$.

- *Unbiased but inconsistent.* The fixed “first-ten” mean,

$$\hat{\mu}_{10} = \frac{1}{10} \sum_{i=1}^{10} X_i,$$

is unbiased for μ but $\text{Var}(\hat{\mu}_{10}) = \sigma^2/10 \not\rightarrow 0$, hence $\hat{\mu}_{10}$ is *inconsistent*.

4.3. Almost sure convergence

A sequence of random variables $Y_1, Y_2, \dots$, converges almost surely to θ if

$$P\left(\lim_{n \rightarrow \infty} Y_n = \theta\right) = 1$$

In terms of notation, we say that $Y_n \xrightarrow{a.s.} \theta$ as $n \rightarrow \infty$.

There is a notable difference between convergence almost surely and convergence in probability.

$P(\lim_{n \rightarrow \infty} Y_n = \theta) = 1$ is equivalent to: *there exists a finite number N such that for all $n > N$, the event $(|Y_n - \theta| < \epsilon)$ occurs with probability 1 (surely), for any arbitrarily small $\epsilon > 0$.*⁵

On the other hand, convergence in probability states that $\lim_{n \rightarrow \infty} P(|Y_n - \theta| < \epsilon) = 1$, which is equivalent to: *there exists a finite number N such that for all $n > N$, $P(|Y_n - \theta| < \epsilon) > 1 - \delta$, for any arbitrarily small $\epsilon > 0$ and $\delta > 0$.*

In almost-sure convergence, the event $|Y_n - \theta| < \epsilon$ is sure (with probability 1), while for convergence in probability, the event $|Y_n - \theta| < \epsilon$ occurs with probability close to one but never exactly one. Strong Law of Large Numbers: the sample mean converges almost surely to the population mean as $n \rightarrow \infty$.

Convergence almost surely is stronger than convergence in probability. The former implies the latter. But convergence in probability does not imply almost sure convergence. For example, consider the random variable $Y_n = 1$ with probability $\frac{1}{n}$, and $Y_n = 0$ with probability $1 - \frac{1}{n}$. Then Y_n converges to 0 in probability, but Y_n does not converge to 0 almost surely.

⁵Note that θ can be a random variable. So that when Y_n converges almost surely to Y , it strictly means that $P(\omega \in \Omega : \lim_{n \rightarrow \infty} Y_n(\omega) = \theta(\omega)) = 1$. Recall that a random variable is a function from some sample space Ω to the real numbers.

4.4. Convergence in distribution

A sequence of random variables $Y_1, Y_2, \dots$, converges in distribution to a random variable Y if

$$\lim_{n \rightarrow \infty} F_{Y_n}(y) = F_Y(y), \quad \text{for all } y$$

That is, the sequence of cdfs corresponding to the sequence of random variables $X_1, X_2, \dots$, converges to the cdf of X .

Suppose $X_1, \dots, X_n$ is a random sample from $U[0, 1]$. Consider the largest order statistic $X_{(n)} = \max\{X_1, \dots, X_n\}$.

What does $X_{(n)}$ converge to? It seems like $X_{(n)}$ would converge in probability to 1. Let's verify it.

$$\begin{aligned} P(|X_{(n)} - 1| \leq \epsilon) &= P(1 - \epsilon \leq X_{(n)} \leq 1 + \epsilon) \\ &= P(1 - \epsilon \leq X_{(n)}) \\ &= 1 - (1 - \epsilon)^n \\ &\rightarrow 1 \text{ as } n \rightarrow \infty \end{aligned}$$

However consider the random variable $n(1 - X_{(n)})$.

$$\begin{aligned} P(n(1 - X_{(n)}) \leq t) &= P(1 - t/n \leq X_{(n)}) \\ &= 1 - (1 - t/n)^n \\ &\rightarrow 1 - e^{-t} \text{ as } n \rightarrow \infty \end{aligned}$$

Therefore, $n(1 - X_{(n)})$ converges in distribution to an exponential(1) random variable, since an exponential(1) random variable has the cdf $F(x) = 1 - e^{-x}$.

Convergence in probability implies convergence in distribution. In general, convergence in distribution does not imply convergence in probability, except when the converged object is a constant.

4.5. Mean (L^p) convergence

Mean convergence or convergence in L^p : For $p \geq 1$, $Y_n \xrightarrow{L^p} \theta$ if $\mathbb{E} |Y_n - \theta|^p \rightarrow 0$.

If $Y_n \xrightarrow{L^p} \theta$ for some $p > 0$, then $Y_n \xrightarrow{p} \theta$. To see this, by Markov, for any $\epsilon > 0$, $\Pr(|Y_n - \theta| \geq \epsilon) \leq \mathbb{E} |Y_n - \theta|^p / \epsilon^p \rightarrow 0$.

4.6. Complete convergence

$Y_n \rightarrow \theta$ *completely* if for every $\varepsilon > 0$,

$$\sum_{n=1}^{\infty} \Pr(|Y_n - \theta| > \varepsilon) < \infty.$$

(Borel–Cantelli Lemma) if $Y_n \rightarrow \theta$ completely, then $Y_n \xrightarrow{a.s.} \theta$.

Theorem 1 (Hsu–Robbins). *If $X_1, X_2, \dots$ are i.i.d. with mean μ and $0 < \text{Var}(X_1) < \infty$, then for every $\varepsilon > 0$,*

$$\sum_{n=1}^{\infty} \Pr(|\bar{X}_n - \mu| > \varepsilon) < \infty,$$

i.e., $\bar{X}_n \rightarrow \mu$ completely (hence a.s.).

Erdős (1949) proved the converse, yielding the Strong Law of Large Numbers.

Theorem 2 (Hsu–Robbins–Erdős). *For i.i.d. X_i with mean μ , the following are equivalent:*

$$\text{Var}(X_1) < \infty \iff \sum_{n=1}^{\infty} \Pr(|\bar{X}_n - \mu| > \varepsilon) < \infty \text{ for all } \varepsilon > 0.$$

4.7. Relations at a glance

For any (possibly random) Y ,

$$Y_n \xrightarrow{a.s.} Y \implies Y_n \xrightarrow{p} Y \implies Y_n \xrightarrow{d} Y.$$

In addition, $Y_n \xrightarrow{L^p} Y \implies Y_n \xrightarrow{p} Y$.

If $Y \equiv c$ is a constant (degenerate r.v.), then

$$Y_n \xrightarrow{d} c \iff Y_n \xrightarrow{p} c.$$

Let $Y \sim \text{Bernoulli}(1/2)$, and let $(Y_n)_{n \geq 1}$ be i.i.d. $\text{Bernoulli}(1/2)$ independent of Y . Then the cdf of Y_n equals the cdf of Y for every n , hence $Y_n \xrightarrow{d} Y$. However, for any $0 < \varepsilon \leq 1$,

$$\Pr(|Y_n - Y| > \varepsilon) = \Pr(|Y_n - Y| = 1) = \Pr(Y_n \neq Y) = \frac{1}{2},$$

so $Y_n \not\xrightarrow{p} Y$.