

Martingale

A Martingale is a sequence of random variables $X_1, X_2, \dots$, such that

$$E[X_{n+1} | X_1, X_2, \dots, X_n] = X_n \quad \forall n$$

The expected value of X_{n+1} is X_n given X_1 to X_n

A martingale can also exist in continuous time,

$$E[X_{t+s} | X_t, X_{[0,t]}] = X_t$$

A submartingale has the property that

$$E[X_{n+1} | X_1, \dots, X_n] \geq X_n \quad \forall n$$

A supermartingale has the property that

$$E[X_{n+1} | X_1, \dots, X_n] \leq X_n \quad \forall n$$

Time-to-Event Data

Time-to-event data is a random vector of variables $T_1, \dots, T_n$ such that each r.v. i corresponds to the time it takes for person i to have the event of interest

Often, for some people, say person j , the event of interest is not observed during the observation period (the time for which person j is "watched" to see if they have an event). When this occurs, we say person j 's outcome is censored.

Hazard Function

The Hazard Function represents the instantaneous likelihood (density) that a person has an event at time t given that they have survived up to time T

Let T_j be a random variable representing the survival time of individual j , T_j having pdf $f(t)$ and CDF $F(t)$

The Hazard Function $h(t)$ is defined as

$$h(t) = \frac{f(t)}{1 - F(t)}$$

$1 - F(t) = P(T \geq t)$, which is often called the Survival Function $S(t) = 1 - F(t)$

$$h(t) = \frac{f(t)}{S(t)}$$

By chain rule, $-\frac{d}{dt} \log(S(t)) = \frac{f(t)}{S(t)} = h(t)$

We call $H(t) = \int_0^t h(u) du$ the Cumulative Hazard Function.
It can be shown that $S(t) = \exp(-H(t))$

Kaplan-Meier Curve

The Kaplan-Meier "Curve" is a non-parametric estimator of the survival function that accommodates censored data. It is NPMLE. Whenever an event occurs among one of the iid time to events $T_1, \dots, T_n$ in the sample, say an event occurs at time a , the survival curve shrinks by a factor of the proportion of people removed by the event. When someone is censored, that person is removed without changing the KM curve.

Example:

time	$\frac{d}{n}$	$\frac{m}{n}$	$(1 - \frac{d}{n})$	$(1 - \frac{m}{n})$	$\prod (1 - \frac{d}{n})(1 - \frac{m}{n})$	KM C
6	$\frac{1}{8}$	0	$\frac{7}{8}$	1	$\frac{7}{8}$	$\frac{7}{8}$
6	$\frac{1}{7}$	0	$\frac{6}{7}$	1	$\frac{6}{8}$	$\frac{6}{8}$
6*	0	$\frac{1}{6}$	1	$\frac{5}{6}$	$\frac{5}{8}$	$\frac{6}{8}$
7	$\frac{1}{5}$	0	$\frac{4}{5}$	1	$\frac{4}{8}$	0.6
8*	0	$\frac{1}{4}$	1	$\frac{3}{4}$	$\frac{3}{8}$	0.6
9	$\frac{1}{3}$	0	$\frac{2}{3}$	1	$\frac{2}{8}$	0.4
10	$\frac{1}{2}$	0	$\frac{1}{2}$	1	$\frac{1}{8}$	0.2

Sandwich Variance

The "Sandwich" variance of a parameter estimate $\hat{\beta}$ obtained with maximum likelihood estimation is an estimate of the variance-covariance matrix of $\hat{\beta}$ $(\hat{\beta}_1, \hat{\beta}_2, \dots, \hat{\beta}_j)$ of the following form:

$$\widehat{\text{Var}}(\hat{\beta}) = \underbrace{\left(H(\hat{\beta}) \right)^{-1}}_{\text{bread, Hessian}} \underbrace{\left(\sum_{i=1}^n \frac{\partial \ell_i(\beta)}{\partial \beta} \left(\frac{\partial \ell_i(\beta)}{\partial \beta} \right)^T \right) \bigg|_{\beta=\hat{\beta}}}_{\text{Meat (p x p, each j x k is } \sum_{i=1}^n \frac{\partial \ell_i(\beta_j)}{\partial \beta_j} \frac{\partial \ell_i(\beta_k)}{\partial \beta_k})} \underbrace{\left(H(\hat{\beta}) \right)^{-1}}_{\text{bread, Hessian}}$$

Where $H(\hat{\beta}) = \frac{-\partial^2 \ell(\beta)}{\partial \beta \partial \beta^T} \bigg|_{\beta=\hat{\beta}}$, also p x p,

each j x k is $\frac{-\partial^2 \ell(\beta)}{\partial \beta_j \partial \beta_k} \bigg|_{\beta=\hat{\beta}}$

Discrete Hazard Function

If the time to event is distributed discretely, with mass at $t = a_j$ is $F_j = P(T = a_j)$ $j=1, \dots$, the discrete hazard func. is

$$P(T = a_j | T \geq a_j) = \frac{F_j}{F_j + F_{j+1} + F_{j+2} + \dots}$$

$$\text{So } F_i = h_i \prod_{j=1}^{i-1} (1 - h_j) \text{ and } S(t) = \prod_{j: a_j \leq t} (1 - h_j)$$

(much like Kaplan-Meier curve)

The cumulative hazard is $H(a_j) = \sum_{i=1}^j h_i$

Stieglis notation gives general F and H :

$$F(t) = \int_0^t dF(u) \quad H(t) = \int_0^t \frac{dF(u)}{1 - F(u-)}$$

$F(u-)$ is $\lim_{\epsilon \rightarrow 0^+} F(u - \epsilon)$, left hand limit of F at u

Exponential & Weibull

If $\lambda(t)$ is constant, $\lambda(t) = \lambda$, then the event times must have exponential distribution with parameter λ . Some info about exponential:

$$f(t) = \lambda e^{-\lambda t} \quad \text{"memoryless"} \leftrightarrow h(t) = \lambda$$

$$S(t) = e^{-\lambda t}$$

$$H(t) = -\log(S(t)) = \lambda t$$

A more general distribution that is exponential when $K=1$ is the Weibull distribution, K, λ

$$f(t) = K\lambda(\lambda t)^{K-1} e^{-(\lambda t)^K}$$

$$h(t) = K\lambda(\lambda t)^{K-1}$$

$$S(t) = e^{-(\lambda t)^K}$$

$$H(t) = (\lambda t)^K \quad (\lambda > 0, K > 0)$$

The plot of $\log(-\log(S(t)))$ vs $\log(t)$ is linear with slope K and intercept $K \log(\lambda)$. $(\lambda t)^K$ is unit exponential. The hazard function is monotone increasing if $K > 1$, constant if $K = 1$, and monotone decreasing if $K < 1$.

Extreme Value Distribution

If Y is Weibull, $Y = \log(T)$ has an "extreme value distribution"

Write $y = \log(t)$, $w = -\log(\lambda)$, and $\sigma = \frac{1}{K}$, then

$$(\lambda t)^K = (e^{y-w})^K = \exp\left(\frac{y-w}{\sigma}\right) \quad \text{thus}$$

$Y = \frac{\log T - w}{\sigma}$ has the standard extreme value distribution, with survival function

$$P(Y > x) = e^{-e^x} \quad \text{and density } e^{(x-e^x)}$$

So $\log T$ is in a location-scale family

Integrated Hazard Transform

Let T be a survival time with arbitrary continuous survivor function $S_T(t)$ and cum hazard function $H(t)$. $H(T)$ is the integrated hazard transform of T . The survivor function of $U = H(T)$ is $S_U(u) = P(U > u) = P(H(T) > u) = P(T > H^{-1}(u))$, assuming H has an inverse. This is the probability that T is larger than T 's own hazard function at u . Now, since T has survivor function S , we have $= S_T(H^{-1}(u))$. Now, recall that $S(t) = \exp(-H(t))$, so $S_U(u) = S_T(H^{-1}(u)) = \exp(-H(H^{-1}(u))) = e^{-u}$

Thus $H_T(T)$ always has exponential distribution with $\lambda=1$. This transform is similar to the probability integral transform that says that $F_X(X)$ must be uniform $(0,1)$ (provided X is continuous)

Censoring

Censoring occurs when for some units the outcome of interest is not observed, but the unit is known not to have had an event in the time interval $(0, c)$ where c is called the censoring time.

If for unit i , T_i is the event time and C_i is the censoring time, let X_i be

$$X_i \equiv \min(T_i, C_i)$$

$$\delta = \begin{cases} 1 & \text{if } T_i \leq C_i \\ 0 & \text{or} \end{cases}$$

C_i is viewed as constant, not random, and thus independent of T_i , "non-informative" censoring

Type I censoring: All C_i are equal

Type II censoring: Censor all remaining subj after d events

Restricted Mean Survival & Residual Life

The restricted mean survival, which gives the mean survival of all individuals up to time ℓ , is defined as

$$RMS(\ell) = \int_0^{\ell} S(t) dt$$

The conditional distribution of the remaining life U of an individual known to have survived to t has density

$$f_{U|t}(u|t) = \frac{f(u+t)}{S(t)}, \quad u \geq 0$$

The expected value of U as a function of t is called the mean residual life function, and is given by

$$r(t) = E[T - t | T \geq t] = \frac{\int_{u=t}^{\infty} (u - t) f(u) du}{S(t)}$$

Note that $r(0) = E[T]$, and expectation of total life $E[T | T \geq t]$ of T given the person survived to t is just $r(t) + t$

Pyke's formula: $Var(T) = E[r(T)]$

Pareto & Log-Logistic

Pareto:

$$S(t) = \left(\frac{K}{K+Pt} \right)^K \quad h(t) = \frac{Kp}{K+Pt}$$

if $K > 1$:

$$E[T] = \frac{K}{p(K-1)}$$

Provides a model for heterogeneous Pop. where each lifetime is exp. dist. w/ diff λ

Log-logistic:

The standard logistic density over $-\infty < y < \infty$ is

$$f(y) = \frac{e^y}{(1+e^y)^2} \quad \text{and survival function} \quad \frac{1}{1+e^y}$$

As a location scale family: $f(y|\sigma, u) = \frac{\frac{1}{\sigma} e^{(y-u)/\sigma}}{(1+e^{(y-u)/\sigma})^2}$

When $P = e^{-u}$ and $K = \frac{1}{\sigma}$, exp(logistic) has

$$S(t) = \frac{1}{1+(pt)^K} \quad f(t) = \frac{Kp(pt)^{K-1}}{(1+(pt)^K)^2} \quad h(t) = \frac{Kp(pt)^{K-1}}{1+(pt)^K}$$

Standard log-logistic is: $f(t) = \frac{1}{(1+t)^2}$, $K=1$, $p=1$

Scale & Proportional Hazard Families

A Scale family has a distribution that, when multiplied by a constant, the random variable is still in that distributional family. A scale family with scale parameter B satisfies

$S(t|B) = S_0(e^B t)$ where S_0 is the survival function of the "standard" version where the scale is 1

Likewise, the hazard satisfies $h(t, B) = e^B h_0(e^B t)$

A Proportional hazards family with "standard" hazard h_0 has parameter B such that all other hazards in family have

$$h(t) = e^B h_0(t) \quad \text{implying} \quad S(t) = S_0(t)^{e^B}$$

Scale and Prop-hazards families based on S_0 are identical iff S_0 is the survival function of a Weibull

Nelson Aalen Plot

Simple Graphical Procedure for estimating the integrated hazard function $H(t)$ from a sample of iid variables subject to non-informative censoring

Analogous to Kaplan-Meier estimator of $S(t)$

Data are in form (x_i, δ_i) , $x_i = \min(t_i, c_i)$,
 $\delta_i = 1_{[t_i \leq c_i]}$

Put data in order $X_{(1)} < \dots < X_{(m)}$ are observed,
 $d_{(j)} = \#\{i: t_i = t_{(j)}, \delta_i = 1\}$ is the number of observed failures at $X_{(j)}$, let $r(t) = \#\{i: t_i \geq t\}$ be the size of the risk set at time t and $r_{(j)} = r(t_{(j)})$

The Nelson-Aalen estimator $\hat{H}(t)$ of $H(t)$ is

$$\hat{H}(t) = \sum_{j: t_{(j)} \leq t} \frac{d_{(j)}}{r_{(j)}}$$

Step function with jumps at the ordered failure times $t_{(j)}$. It can be shown that $\hat{H}(t)$ is consistent for $H(t)$.