

Statistics & Probability

76

What is error?

Two kinds: Random & Systematic

Random error is the scatter you get from repeated measurements

Systematic error is systematic: it happens to the same degree each time you make a measurement.

Multiple measurements will average out (reduce) random error, but won't affect systematic error!

If you do many measurements, the error in the avg. will be completely dominated by systematic error, but this is very difficult to estimate!

Be wary of error estimates - always look to see how it is calculated.

As an engineer, a detailed understanding of statistics and error is critical; You must master this material!

OK, now for some elementary statistics:

The most important probability distribution is the Gaussian or normal distribution:

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

$p(x) \equiv$ probability density

$=$ probability of meas. being in interval $[x, x+dx]$
 $\frac{dx}{dx}$

We can define the cumulative probability:

$$P(x) = \int_{-\infty}^x p(x') dx'$$

(78)

This is the probability that a meas. is less than some value x

The normal distribution is characterized by 2 quantities:

$\mu \equiv$ mean of distribution

$\sigma \equiv$ population standard deviation

About 69% of observations lie within 1σ of μ , 95% within 2σ , and 99% within 3σ .

In working w/ statistics, we define an expectation value

$E(x) \equiv$ what you expect to get if you do a meas. many times and average it together!

Mathematically:

$$E(x) \equiv \int_{-\infty}^{\infty} x \cdot p(x) dx$$

$\uparrow$ meas. quantity $\uparrow$ prob. density

$$= \int_{-\infty}^{\infty} \frac{x}{\sqrt{2\pi}\sigma} e^{-\frac{(x-\mu)^2}{2\sigma^2}} dx$$

$$= \int_{-\infty}^{\infty} \frac{x-\mu}{\sqrt{2\pi}\sigma^2} e^{-\frac{(x-\mu)^2}{2\sigma^2}} dx + \mu \int_{-\infty}^{\infty} p(x) dx$$

Let $z \equiv x - \mu$

$$\therefore E(x) = \underbrace{\int_{-\infty}^{\infty} \frac{z}{\sqrt{2\pi}\sigma^2} e^{-\frac{z^2}{2\sigma^2}} dz}_{\text{odd}} + \mu \underbrace{\int_{-\infty}^{\infty} p(x) dx}_{=1}$$

The integral of an odd function over an even domain is zero, so first integral vanishes!

Thus $E(x) = \mu$

What about the mean of N observations?

$$\bar{x} = \frac{1}{N} \sum_{i=1}^N x_i$$

$$\begin{aligned} \therefore E(\bar{x}) &= E\left(\frac{1}{N} \sum_{i=1}^N x_i\right) \\ &= \frac{1}{N} \sum_{i=1}^N E(x_i) = \frac{1}{N} \sum_{i=1}^N \mu = \mu \end{aligned}$$

This may have been self-evident, but we can use the same approach for the variance:

What is $E((x-\mu)^2)$?

$$\begin{aligned} E((x-\mu)^2) &= \int_{-\infty}^{\infty} (x-\mu)^2 p(x) dx \\ &= \int_{-\infty}^{\infty} \frac{(x-\mu)^2}{\sqrt{2\pi} \sigma} e^{-\frac{(x-\mu)^2}{2\sigma^2}} dx \end{aligned}$$

Let $z = (x-\mu)$

$$\therefore E((x-\mu)^2) = \int_{-\infty}^{\infty} \frac{z^2}{\sqrt{2\pi}\sigma} e^{-\frac{z^2}{2\sigma^2}} dz \quad \text{Ex}$$

Integrate by parts:

$$= - \frac{\sigma^2 z}{\sqrt{2\pi}\sigma} e^{-\frac{z^2}{2\sigma^2}} \Big|_{-\infty}^{\infty} + \sigma^2 \int_{-\infty}^{\infty} \frac{1}{\sqrt{2\pi}\sigma} e^{-\frac{z^2}{2\sigma^2}} dz$$

$\underbrace{\hspace{10em}}_{\text{vanishes at } \pm\infty} \quad \underbrace{\hspace{10em}}_{=1}$

$$\therefore E((x-\mu)^2) = \sigma^2$$

What about the variance of the mean of N measurements?

$$E((\bar{x} - \mu)^2) = \sigma_{\bar{x}}^2$$

$$= E\left(\left[\frac{1}{N} \sum_{i=1}^N x_i - \mu\right]^2\right)$$

$$= \frac{1}{N^2} E\left(\left[\sum_{i=1}^N x_i - N\mu\right]^2\right) = \frac{1}{N^2} E\left(\left(\sum_{i=1}^N (x_i - \mu)\right)^2\right)$$

$$= \frac{1}{N^2} E\left(\sum_{i=1}^N (x_i - \mu)^2 + \sum_{i=1}^N \sum_{j \neq i} (x_i - \mu)(x_j - \mu)\right)$$

(82)

$$= \frac{1}{N} E \left(\frac{1}{N} \sum_{i=1}^N (x_i - \mu)^2 \right)$$

$$+ E \left(\frac{1}{N^2} \sum_{i=1}^N \sum_{j \neq i}^N (x_i - \mu)(x_j - \mu) \right)$$

$$= \frac{1}{N^2} \sum_{i=1}^N E((x_i - \mu)^2)$$

$$+ \frac{1}{N^2} \sum_{i=1}^N \sum_{j \neq i}^N E((x_i - \mu)(x_j - \mu))$$

First expectation is just σ^2

second term is zero provided x_i & x_j are independent

$$E((x_i - \mu)(x_j - \mu)) = \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} (x_i - \mu)(x_j - \mu) p(x_i) p(x_j) dx_i dx_j$$

$$= \int_{-\infty}^{\infty} (x_i - \mu) p(x_i) dx_i \int_{-\infty}^{\infty} (x_j - \mu) p(x_j) dx_j = 0 \quad i \neq j$$

Thus:

$$\sigma_{\bar{x}}^2 \equiv \frac{\sigma^2}{N}$$

That is why multiple measurements reduce the random error!

OK, how do we estimate μ and σ ?

In general we have a sample of N observations drawn from an underlying population of possible observations!

We define the sample mean

$$\bar{x} = \frac{1}{N} \sum_{i=1}^N x_i$$

We have already shown that

$$E(\bar{x}) = \mu$$

That is, $\bar{x}$ is an unbiased estimate of μ .

Now for the variance:

We define the sample variance:

$$S_x^2 \equiv \frac{1}{N-1} \sum_{i=1}^N (x_i - \bar{x})^2$$

$$= \frac{1}{N-1} \sum_{i=1}^N (x_i^2 - 2x_i\bar{x} + \bar{x}^2)$$

$$= \frac{1}{N-1} \left\{ \sum_{i=1}^N x_i^2 - 2\bar{x} \underbrace{\sum_{i=1}^N x_i}_{\equiv N\bar{x}} + N\bar{x}^2 \right\}$$

$$= \frac{1}{N-1} \sum_{i=1}^N (x_i^2 - \bar{x}^2)$$

Note that to calculate this we don't have to store all the x_i 's!

We just need to keep N , $\sum_{i=1}^N x_i$

and $\sum_{i=1}^N x_i^2 \Rightarrow$ that's how a calculator

with only a few memories can calc. S_x^2 .

OK, how is S_x^2 related to σ^2 ? 85

We need the following identities:

$$\sum_{i=1}^N (x_i - \mu)^2 = \sum_{i=1}^N (x_i - \bar{x})^2 + N(\bar{x} - \mu)^2$$

(this is because $\sum_{i=1}^N (x_i - \bar{x}) = 0$ by the definition of $\bar{x}$)

Also, we had:

$$E((\bar{x} - \mu)^2) \equiv \sigma_{\bar{x}}^2 = \frac{\sigma^2}{N}$$

$$\text{and } E((x_i - \mu)^2) = \sigma^2$$

Thus:

$$E(S_x^2) \equiv E\left(\frac{1}{N-1} \sum_{i=1}^N (x_i - \bar{x})^2\right)$$

$$= \frac{1}{N-1} \left\{ E\left(\sum_{i=1}^N (x_i - \mu)^2\right) - N E((\bar{x} - \mu)^2) \right\}$$

$$= \frac{1}{N-1} \left\{ N\sigma^2 - \frac{N}{N}\sigma^2 \right\} = \overset{(86)}{\sigma^2}!$$

So S_x^2 is an unbiased estimate for σ^2 !

The $N-1$ part comes about because we are calculating the mean from the same sample that we are calculating the variance.

We have reduced the number of degrees of freedom by one.

Propagation of Errors:

Usually when doing experimental calculations you must combine several measurements to get the final result.

Each individual meas. has some error associated with it. How do these combine?

(87)

Let x_i, y_i be random variables
w/ mean $\bar{x}, \bar{y}$ and st. dev.
 s_x & s_y

Let

$$z_i = c_1 x_i + c_2 y_i$$

where c_1 & c_2 are csts.

$$\therefore \bar{z} = \frac{1}{N} \sum_{i=1}^N z_i = c_1 \bar{x} + c_2 \bar{y}$$

What about the variance?

$$s_z^2 = \frac{1}{N-1} \sum_{i=1}^N (z_i - \bar{z})^2$$

$$= \frac{1}{N-1} \left\{ \sum_{i=1}^N (z_i^2 - \bar{z}^2) \right\}$$

Substituting in:

$$z_i^2 = c_1^2 x_i^2 + c_2^2 y_i^2 + 2c_1 c_2 x_i y_i$$

$$\bar{z}^2 = c_1^2 \bar{x}^2 + c_2^2 \bar{y}^2 + 2c_1 c_2 \bar{x} \bar{y}$$

(88)

$$\begin{aligned} \therefore S_z^2 &\equiv \frac{1}{N-1} \sum_{i=1}^N \left(C_1^2 (x_i^2 - \bar{x}^2) + C_2^2 (y_i^2 - \bar{y}^2) \right. \\ &\quad \left. + 2 C_1 C_2 (x_i y_i - \bar{x} \bar{y}) \right) \\ &= C_1^2 S_x^2 + C_2^2 S_y^2 + \frac{2 C_1 C_2}{N-1} \sum_{i=1}^N (x_i y_i - \bar{x} \bar{y}) \end{aligned}$$

Let's look at this last term. We define:

$$\begin{aligned} S_{xy}^2 &\equiv \frac{1}{N-1} \sum_{i=1}^N (x_i y_i - \bar{x} \bar{y}) \\ &= \frac{1}{N-1} \sum_{i=1}^N (x_i - \bar{x})(y_i - \bar{y}) \end{aligned}$$

This is the covariance of x and y . If they are independent, then:

$$E(S_{xy}^2) = 0$$

So for addition:

$$S_z^2 = C_1^2 S_x^2 + C_2^2 S_y^2 + 2 C_1 C_2 S_{xy}^2$$

If we subtract two variables we get a similar result:

$$z_i = x_i - y_i$$

$$\text{then } \bar{z} = \bar{x} - \bar{y}; \quad s_z^2 = s_x^2 + s_y^2 - 2s_{xy}$$

↑
note (-) sign!

A positive covariance (note: s_{xy} may be + or -) reduces the error in the difference between variables.

OK, how about multiplication and division?

$$\text{Let } z_i = \frac{x_i}{y_i}$$

where x_i and y_i are normally distributed.

Note that z is not normally distributed! It only approaches this if $\frac{s_x}{\bar{x}}$ and $\frac{s_y}{\bar{y}} \ll 1$!

(80)

Let's look at this more closely.

$$\text{Let } x_i = \bar{x} + x'_i$$

$\uparrow$ mean $\uparrow$ deviation

$$\text{and } y_i = \bar{y} + y'_i$$

Suppose $\frac{s_x}{\bar{x}} \ll 1$ and $\frac{s_y}{\bar{y}} \ll 1$

This means that x'_i and y'_i are small:

$$z_i = \frac{x'_i}{y'_i} = \frac{\bar{x}_i + x'_i}{\bar{y} + y'_i} = \frac{\bar{x}}{\bar{y}} \frac{(1 + \frac{x'_i}{\bar{x}})}{(1 + \frac{y'_i}{\bar{y}})}$$

$$\approx \frac{\bar{x}}{\bar{y}} \left(1 + \frac{x'_i}{\bar{x}}\right) \left(1 - \frac{y'_i}{\bar{y}} + O\left(\left(\frac{y'_i}{\bar{y}}\right)^2\right)\right)$$

$$\approx \frac{\bar{x}}{\bar{y}} \left(1 + \frac{x'_i}{\bar{x}} - \frac{y'_i}{\bar{y}} + O\left(\frac{x'_i y'_i}{\bar{x} \bar{y}}, \left(\frac{y'_i}{\bar{y}}\right)^2\right)\right)$$

We ignore the higher order terms!

(9)

$$\bar{z} = \frac{1}{N} \sum_{i=1}^N \frac{x_i}{y_i} \approx \frac{\bar{x}}{\bar{y}} \frac{1}{N} \sum_{i=1}^N \left(1 + \frac{x_i - \bar{x}}{\bar{x}} + \frac{y_i - \bar{y}}{\bar{y}} \right)$$

Sum to zero

$$\therefore \bar{z} \approx \frac{\bar{x}}{\bar{y}}$$

Now for S_z^2 :

$$S_z^2 \approx \left(\frac{\bar{x}^2}{\bar{y}^2} \right) \left(\frac{S_x^2}{\bar{x}^2} + \frac{S_y^2}{\bar{y}^2} - 2 \frac{S_{xy}}{(\bar{x} \bar{y})} \right)$$

$\uparrow$
 $\approx \bar{z}^2$

Thus:

$$\frac{S_z^2}{\bar{z}^2} \approx \frac{S_x^2}{\bar{x}^2} + \frac{S_y^2}{\bar{y}^2} - 2 \frac{S_{xy}}{\bar{x} \bar{y}}$$

so for division you add the fractional
or relative variance!

(92)

For multiplication you get the same:

$$z_i \equiv x_i y_i$$

$$\bar{z} \approx \bar{x} \bar{y} \quad \text{provided} \quad \frac{s_x}{\bar{x}}, \frac{s_y}{\bar{y}} \ll 1$$

and

$$\frac{s_z^2}{\bar{z}^2} \approx \frac{s_x^2}{\bar{x}^2} + \frac{s_y^2}{\bar{y}^2} + 2 \frac{s_{xy}}{\bar{x} \bar{y}}$$

↑
note (+) sign

What about a more general functional relationship?

Suppose $z_i = f(x_i, y_i)$ where f is some nasty function

$$\text{Let } x_i = \bar{x} + x'_i, \quad y_i = \bar{y} + y'_i$$

We shall expand f in a 2-D

(93)

Taylor series (Note: the T.S. can be generalized to n -dimensions)

So:

$$f(x_i, y_i) = f(\bar{x}, \bar{y}) + x'_i \left. \frac{\partial f}{\partial x} \right|_{\bar{x}, \bar{y}} + y'_i \left. \frac{\partial f}{\partial y} \right|_{\bar{x}, \bar{y}} + O(x_i'^2, y_i'^2, x_i' y_i')$$

(Ignore 2nd order terms!)

$$\therefore \bar{z} \equiv \frac{1}{N} \sum_{i=1}^N f(x_i, y_i)$$

$$\approx f(\bar{x}, \bar{y}) + \left. \frac{\partial f}{\partial x} \right|_{\bar{x}, \bar{y}} \frac{1}{N} \sum_{i=1}^N x'_i + \left. \frac{\partial f}{\partial y} \right|_{\bar{x}, \bar{y}} \frac{1}{N} \sum_{i=1}^N y'_i$$

Sum to zero!

$$= f(\bar{x}, \bar{y})$$

The error is $O(x_i'^2, y_i'^2, x_i' y_i' \cdot \nabla \nabla f)$

↑
dyadic

(94)

Similarly, we get for the variance:

$$s_z^2 \approx \left(\frac{\partial f}{\partial x} \bigg|_{\bar{x}, \bar{y}} \right)^2 s_x^2 + \left(\frac{\partial f}{\partial y} \bigg|_{\bar{x}, \bar{y}} \right)^2 s_y^2 + 2 \left(\frac{\partial f}{\partial x} \bigg|_{\bar{x}, \bar{y}} \right) \left(\frac{\partial f}{\partial y} \bigg|_{\bar{x}, \bar{y}} \right) s_{xy}$$

again providing that the higher order terms can be neglected.

This formula reduces to the earlier ones for addition, subtr., mult. & div.

Only for addⁿ & subtr. is it exact because $\nabla \nabla f \equiv 0$ for this case and the higher order terms vanish.

Let's generalize this result to an n -dimensional (n -variable) system

Let $\underline{x} = x_1, x_2, \dots, x_n$

$$z = f(\underline{x})$$

we need to find the variance of each variable x_1, x_2 , etc. and the covariance of all pairs $S_{x_1 x_2}^2$, etc.

We can define the matrix of variance & covariance

$$V_{ij} = \frac{N}{N-1} (\overline{x_i x_j} - \bar{x}_i \bar{x}_j)$$

↑
number of elements contributing
to each x_i, x_j

Now if $\frac{V_{ii}}{\bar{x}_i^2} \ll 1$ for all i

(this corresponds to a small error in each variable, e.g. $\frac{S_x^2}{\bar{x}^2} \ll 1$, $\frac{S_y^2}{\bar{y}^2} \ll 1$, etc.)

And if we define:

$$\nabla f = \left(\frac{\partial f}{\partial x_1}, \frac{\partial f}{\partial x_2}, \frac{\partial f}{\partial x_3}, \dots, \frac{\partial f}{\partial x_n} \right)$$

Then we get the result:

$$S_z^2 = (\nabla f) \underset{\sim}{V} (\nabla f)^T$$

Which gives us an easy way of calculating the variance in $\underline{z} = f(\underline{x})$!

Note: in this derivation the subscripts i and j denote different variables entirely, not different measurements contributing to the same average.

Let's look at the example $f(\underline{x}) = x_1 x_2$

Let $x_1 \equiv x$, $x_2 \equiv y$

$$\therefore V_{ij} = \frac{N}{N-1} \left(\overline{X_i X_j} - \bar{X}_i \bar{X}_j \right)$$

$$\begin{aligned} \text{so } V_{xx} &= \frac{N}{N-1} \left(\overline{XX} - \bar{X} \bar{X} \right) \\ &= \frac{1}{N-1} \sum^N (X^2 - \bar{X}^2) = S_x^2 \end{aligned}$$

(9.12)

$$V_{xy} = \frac{N}{N-1} (\overline{xy} - \bar{x} \bar{y}) = \frac{1}{N-1} \sum^N (xy - \bar{x} \bar{y})$$

$$\equiv S_{xy}^2$$

$$V \approx \begin{pmatrix} S_x^2 & S_{xy}^2 \\ S_{xy}^2 & S_y^2 \end{pmatrix}$$

$$\nabla f = \left(\frac{\partial f}{\partial x}, \frac{\partial f}{\partial y} \right) \equiv (y, x)$$

$$\nabla f \cdot V \cdot (\nabla f)^T = (y \ x) \cdot \begin{pmatrix} y \\ x \end{pmatrix}$$

$$= (y S_x^2 + x S_{xy}^2, y S_{xy}^2 + x S_y^2) \begin{pmatrix} y \\ x \end{pmatrix}$$

$$= y^2 S_x^2 + xy S_{xy}^2 + xy S_{xy}^2 + x^2 S_y^2$$

$$= y^2 S_x^2 + x^2 S_y^2 + 2xy S_{xy}^2$$

So

$$\frac{S_z^2}{(xy)^2} = \frac{S_x^2}{x^2} + \frac{S_y^2}{y^2} + 2 \frac{S_{xy}^2}{xy} \quad !$$

(97A)

Finally, we examine vector functions of vector variables:

Let $\underline{z} = f(\underline{x})$ be a vector function of $\underline{x}$. We define:

$$\underline{\mu}_z \equiv E(\underline{z}), \quad \underline{\mu}_x \equiv E(\underline{x})$$

and the matrix of covariance for each:

$$\underline{\Sigma}_x^2 \equiv E[(\underline{x} - \underline{\mu}_x)(\underline{x} - \underline{\mu}_x)^T]$$

$$\underline{\Sigma}_z^2 \equiv E[(\underline{z} - \underline{\mu}_z)(\underline{z} - \underline{\mu}_z)^T]$$

we have taken $\underline{z}$ and $\underline{x}$ to be column vectors.

Suppose we know $\underline{x}$ and $\underline{\Sigma}_x^2$. We wish to calculate $\underline{\Sigma}_z^2$.

We do a multivariable Taylor series expansion of $\underline{z}$ about $\underline{\mu}_x$:

(97B)

$$\begin{aligned} \underline{z} = \underline{f}(\underline{x}) &= \underline{f}(\underline{\mu}_x) + \left. \nabla \underline{f} \right|_{\underline{\mu}_x} \cdot (\underline{x} - \underline{\mu}_x) \\ &\quad + \frac{1}{2} \nabla \nabla \underline{f} : (\underline{x} - \underline{\mu}_x)(\underline{x} - \underline{\mu}_x)^T + \dots \end{aligned}$$

If the deviation between $\underline{x}$ and $\underline{\mu}_x$ is small (e.g. that $\Sigma_{\underline{x}}^2$ is small),

or $\underline{f}(\underline{x})$ is linear in $\underline{x}$, then we can ignore the quadratic or higher order terms in the Taylor series!

Thus:

$$\underline{z} \approx \underline{f}(\underline{\mu}_x) + \left. \nabla \underline{f} \right|_{\underline{\mu}_x} \cdot (\underline{x} - \underline{\mu}_x)$$

and thus:

$$\underline{\mu}_z \equiv E(\underline{z}) \approx \underline{f}(\underline{\mu}_x) + \left. \nabla \underline{f} \right|_{\underline{\mu}_x} \cdot E(\underline{x} - \underline{\mu}_x)$$

but $E(\underline{x} - \underline{\mu}_x) \equiv 0$ by definition!

Thus $\underline{\mu}_z \approx \underline{f}(\underline{\mu}_x) + \text{quadratic terms!}$

Putting this together, we get:

$$\tilde{z} - \mu_{\tilde{z}} \approx \left. \nabla f \right|_{\mu_x} (\tilde{x} - \mu_x)$$

So:

$$\sum_{\tilde{z}}^2 = E\left((\tilde{z} - \mu_{\tilde{z}})(\tilde{z} - \mu_{\tilde{z}})^T\right)$$

$$\approx E\left[(\nabla f)(\tilde{x} - \mu_x)(\tilde{x} - \mu_x)^T(\nabla f)^T\right]$$

$$= (\nabla f) E((\tilde{x} - \mu_x)(\tilde{x} - \mu_x)^T)(\nabla f)^T$$

$$= (\nabla f) \sum_{\tilde{x}}^2 (\nabla f)^T$$

provided that $\frac{1}{2} \nabla \nabla f : \sum_{\tilde{x}}^2 \ll f$

so that we can neglect the quadratic term in the Taylor series expansion.

We will use this formulation again in calculating linear regression error!

Linear Regression

(50)

In HW problem, wrote routine to fit a parabola to 3 points. In this case there was only one answer.

Often in data analysis you want to fit a curve to many data points

How do we do this? Use (usually) linear regression

This does not mean that we fit data with a line, rather that we use an algorithm which reduces the problem to a set of linear equations!

Let's look at a simple problem: measure mass transfer coefficients

Suppose we have a membrane between two reservoirs



membrane area = A

(5d)

We start with some concentration in reservoir ① of C_0 and we assume that reservoir ② is maintained at a constant C_{eq} . If the volume of reservoir ① is V , then:

$$V \frac{dC}{dt} = -Ah(C - C_{eq})$$

where $C \Big|_{t=0} = C_0$

We can solve this equation:

$$C = C_{eq} + (C_0 - C_{eq}) e^{-\frac{(htA)}{V}}$$

We wish to determine h by measuring C as a function of time.

In order for us to use linear regression, the model must be linear in the modelling parameter,

In this case, the modelling parameter is h . We can rewrite the eqn:

$$(C - C_{eq}) = (C_0 - C_{eq}) e^{-\frac{hA}{V}t}$$

$$\ln(C - C_{eq}) = \ln(C_0 - C_{eq}) - \frac{hA}{V}t$$

So if we plot $\ln(C - C_{eq})$ vs $\frac{At}{V}$

we get a line w/ intercept of $\ln(C_0 - C_{eq})$ and a slope of h !

How do we get the best fit line?

We look at the deviation of the points from the line, and try to minimize this in some way.

Let's take:

$$y \equiv \ln(C - C_{eq}), \quad a \equiv \frac{Ah}{V}$$

$$b \equiv \ln(C_0 - C_{eq})$$

We want to fit a series of points (t_i, y_i) by the model $y = at + b$

(53)

We shall use the method of least squares. We form the sum:

$$\text{Sum} = \sum_{i=1}^N (y_i - (at_i + b))^2$$

which is the sum of the squared distance between data pt. and model in the y direction.

We pick a & b so that this is a minimum

we let:

$$\left. \begin{array}{l} \frac{\partial \text{sum}}{\partial a} = 0 \\ \frac{\partial \text{sum}}{\partial b} = 0 \end{array} \right\} 2 \text{ eqns for } a, b!$$

$$\frac{\partial \text{sum}}{\partial a} = \sum_{i=1}^N -2(y_i - (at_i + b))t_i$$

$$= \sum_{i=1}^N -2y_it_i + 2a \sum_{i=1}^N t_i^2 + 2b \sum_{i=1}^N t_i = 0$$

and:

$$\frac{\partial \text{sum}}{\partial b} = \sum_{i=1}^N -2(y_i - (at_i + b))$$

$$= \sum_{i=1}^N -2y_i + 2a \sum_{i=1}^N t_i + 2Nb = 0$$

This is a pair of linear equations for a & b !

We have the solution:

$$a\bar{t}^2 + b\bar{t} = \bar{y}\bar{t} \quad (\text{let } \bar{x} \equiv \frac{1}{N} \sum_{i=1}^N x_i)$$

$$a\bar{t} + b = \bar{y}$$

$$\therefore a = \frac{\bar{y}\bar{t} - \bar{y}\bar{t}}{\bar{t}^2 - \bar{t}^2}$$

$$\text{and } b = \bar{y} - a\bar{t}$$

which can be used to get the mass transfer coefficient!

This is not the only way to approach the problem. Let's look at it as a system of equations!

We are trying to satisfy:

$$y_1 = a t_1 + b$$

$$y_2 = a t_2 + b$$

$$\vdots$$

$$y_N = a t_N + b$$

Let's write this in matrix form:

$$\begin{pmatrix} y_1 \\ y_2 \\ \vdots \\ y_N \end{pmatrix} \approx \begin{pmatrix} t_1 & 1 \\ t_2 & 1 \\ \vdots & \vdots \\ t_N & 1 \end{pmatrix} \begin{pmatrix} a \\ b \end{pmatrix}$$

we shall use X to be the vector containing the modelling parameters:

$$X_1 \equiv a \quad X_2 \equiv b$$

If we let the matrix of modelling functions be $\underline{A}$ (e.g. $A_{11} = t_1$, etc.)

then we want to satisfy:

$$\underline{y} \approx \underline{A} \underline{x}$$

Actually, we want to minimize the residual:

$$\underline{r} = \underline{y} - \underline{A} \underline{x}$$

Or, for least squares, adjust $\underline{x}$ s.t.:

$$\|\underline{r}\|_2 \equiv \|\underline{y} - \underline{A} \underline{x}\|_2 \text{ is as small as possible.}$$

↑
Euclidean norm

We want to pick $\underline{x}$ s.t.

$$\underline{r}^T \underline{r} \equiv (\underline{y} - \underline{A} \underline{x})^T (\underline{y} - \underline{A} \underline{x})$$

is a minimum.

This is exactly equivalent to what we did before!

Let's multiply this out:

$$r^2 \equiv \underbrace{y^T}_{\sim} \underbrace{y}_{\sim} - \underbrace{y^T}_{\sim} \underbrace{A}_{\sim} \underbrace{x}_{\sim} - \underbrace{x^T}_{\sim} \underbrace{A^T}_{\sim} \underbrace{y}_{\sim} + \underbrace{x^T}_{\sim} \underbrace{A^T}_{\sim} \underbrace{A}_{\sim} \underbrace{x}_{\sim}$$

We take the gradient of this w.r.t. $\underline{x}$ and set it equal to zero!

$$\nabla_{\underline{x}} r^2 = 0 = -2 \underbrace{A^T}_{\sim} \underbrace{y}_{\sim} + 2 \underbrace{A^T}_{\sim} \underbrace{A}_{\sim} \underbrace{x}_{\sim} = 0$$

Thus we have the problem:

$$\underbrace{A^T}_{\sim} \underbrace{A}_{\sim} \underbrace{x}_{\sim} = \underbrace{A^T}_{\sim} \underbrace{y}_{\sim} \quad !$$

This gives us a set of n eq'ns for the n unknowns in $\underline{x}$!

These are known as the normal equations.

Linear Regression Error

(98)

Let's apply all this back to linear regression. Suppose we have some data and we are trying to fit it, and calc. the uncertainty in the modelling parameters.

We measure g by tossing a lead weight in the air and meas. its pos'n as a $f^n(t)$

We get:

t_i (sec)	h_i (m)
0	0.3
1	14.8
2	20.7
3	15.9
4	2.0

We fit this to the equation:

$$h = h_0 + v_0 t + \frac{1}{2} g t^2$$

We want to both calculate g , and determine the uncertainty in g .

First let's get g . We have the problem:

$$h_i = (1) h_0 + (t) v_0 + (-\frac{1}{2}t^2) g$$

$\uparrow \quad \quad \quad \uparrow \quad \quad \quad \uparrow$
 $x_1 \quad \quad \quad x_2 \quad \quad \quad x_3$

So:

$$\begin{pmatrix} 1 & 0 & 0 \\ 1 & 1 & -0.5 \\ 2 & 2 & -2 \\ 3 & 3 & -\frac{9}{2} \\ 4 & 4 & -8 \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} = \begin{pmatrix} 0.3 \\ 14.8 \\ 20.7 \\ 15.9 \\ 2.0 \end{pmatrix}$$

$\nwarrow \hat{h}$

Using the normal eq's we get:

$$\underset{\sim}{A}^T \underset{\sim}{A} \underset{\sim}{x} = \underset{\sim}{A}^T \underset{\sim}{h}$$

$$\underset{\sim}{x} = [\underset{\sim}{A}^T \underset{\sim}{A}]^{-1} \underset{\sim}{A}^T \underset{\sim}{h}$$

(100)

Or numerically:

$$\tilde{x} = \begin{pmatrix} .886 & .257 & -.086 & -.143 & .086 \\ -.771 & .186 & .571 & .386 & -.371 \\ -.286 & .143 & .286 & .143 & -.286 \end{pmatrix} \begin{pmatrix} h_1 \\ h_2 \\ h_3 \\ h_4 \\ h_5 \end{pmatrix}$$

This yields:

$$x_1 = 0.197, \quad x_2 = 19.7 \quad x_3 = 9.64$$

Now we want to examine the error in the position measurements h_i .

To do this we look at the deviation:

$$S_h^2 \approx \frac{1}{N-3} \sum_{i=1}^N (h_i - \underset{\substack{\uparrow \\ \text{using fitted} \\ \text{parameters}}}{h(t_i)}})^2$$

$\uparrow$
because 3
adjustable
parameters
det. from data

The quantity $h_i - h(t_i)$ is just the residual; (10)

$$\tilde{v} = \tilde{h} - \tilde{A} \tilde{x} !$$

$$\therefore S_h^2 = \frac{1}{N-3} \|\tilde{v}\|_2^2 = 0.110$$

→ you should really have more pts for accuracy!

OK, how do we use this to get the error in $q(x_3)$?

We had:

$$\tilde{x} = \left\{ \left[\tilde{A}^T \tilde{A} \right]^{-1} \tilde{A}^T \right\} \tilde{h}$$

↑

$$\equiv \tilde{K}$$

$$\text{Thus } x_i \equiv \sum_{j=1}^N K_{ij} h_j$$

i 'th row of $\tilde{K}$!

(122)

We can thus use the formula for a linear combination of random variables!

$$S_{X_i}^2 = \sum_{j=1}^N K_{ij}^2 S_{h_j}^2 = 0.0315$$

↑
these are all the same
in this problem!

Thus $g = 9.64 \pm 0.18$

to the 69% confidence level (1 σ)

Often we want to develop correlations for their predictive value; for example what is the expected ht. of the lead wt. at some time t ?

$$h_m = X_1 + X_2 t - \frac{1}{2} t^2 X_3 !$$

But what is the error? We have a linear combination of X_1, X_2, X_3 each of which have some error.

(103)

You might expect that:

$$S_{h_m}^2 = S_{x_1}^2 + (t)^2 S_{x_2}^2 + \left(-\frac{1}{2}t^2\right)^2 S_{x_3}^2$$

But this is incorrect!!

The problem is that x_1, x_2, x_3 are not independent! Their covariance is non-zero!

Thus we should have:

$$\begin{aligned} S_{h_m}^2 &= S_{x_1}^2 + (t) S_{x_2}^2 + \left(-\frac{1}{2}t^2\right)^2 S_{x_3}^2 \\ &\quad + 2(1)(t) S_{x_1 x_2}^2 + 2(1)\left(-\frac{1}{2}t^2\right) S_{x_1 x_3}^2 \\ &\quad + 2(t)\left(-\frac{1}{2}t^2\right) S_{x_2 x_3}^2 \end{aligned}$$

How do we calculate the covariance?

Recall:

$$x_i = \sum_{j=1}^N K_{ij} h_j$$

Where $\underline{\underline{K}} \equiv [\underline{\underline{A}}^T \underline{\underline{A}}]^{-1} \underline{\underline{A}}^T$

Thus we let:

$$x_1 = \sum_{j=1}^N \alpha_j h_j$$

where α is 1st row
of $\tilde{K}$

and

$$x_2 = \sum_{j=1}^N \beta_j h_j$$

where β is the 2nd
row of $\tilde{K}$

Now we need to determine the
covariance

$$\sigma_{x_1 x_2}^2 = E((x_1 - \mu_{x_1})(x_2 - \mu_{x_2}))$$

↑ ↑
value for x_1 & x_2 expected
for an infinite sample

$$\mu_{x_1} = \sum_{j=1}^N \alpha_j \mu_{h_j}$$

↑
from infinite sample

$$\mu_{x_2} = \sum_{j=1}^N \beta_j \mu_{h_j}$$

So:

$$\sigma_{x, x_2}^2 \equiv E \left(\left(\sum_{j=1}^N \alpha_j (h_j - \mu_{h_j}) \right) \left(\sum_{j=1}^N \beta_j (h_j - \mu_{h_j}) \right) \right)$$

Now if all of the h_j are independent then

$$E((h_j - \mu_{h_j})(h_i - \mu_{h_i})) = \underline{\underline{0}}$$

for $i \neq j$

$$\therefore \sigma_{x, x_2}^2 = E \left(\sum_{j=1}^N \alpha_j \beta_j (h_j - \mu_{h_j})^2 \right)$$

$$= \sum_{j=1}^N \alpha_j \beta_j \sigma_h^2$$

↑
provided all the h_j have the same uncertainty

$$\therefore \sigma_{x, x_2}^2 = \underline{\underline{\alpha}} \underline{\underline{\beta}}^T \sigma_h^2$$

So for this problem if we let

$$\underline{\underline{\gamma}} \equiv \text{3rd row of } \underline{\underline{K}}$$

(106)

$$x_1 = \underset{\sim}{\alpha} \underset{\sim}{h}, \quad x_2 = \underset{\sim}{\beta} \underset{\sim}{h}, \quad x_3 = \underset{\sim}{\gamma} \underset{\sim}{h}$$

$$S_h^2 = \frac{1}{N-3} \|\underset{\sim}{r}\|_2^2, \quad \underset{\sim}{r} = \underset{\sim}{h} - \underset{\sim}{A} \underset{\sim}{x}$$

$$S_{x_1}^2 = (\underset{\sim}{\alpha} \underset{\sim}{\alpha}^T) S_h^2$$

$$S_{x_2}^2 = (\underset{\sim}{\beta} \underset{\sim}{\beta}^T) S_h^2 \quad \text{etc.}$$

and:

$$h_m = (1)x_1 + (t)x_2 + \left(-\frac{1}{2}t^2\right)x_3$$

$$\begin{aligned} \sigma_{h_m}^2 &= (1)^2 \sigma_{x_1}^2 + (t)^2 \sigma_{x_2}^2 + \left(-\frac{1}{2}t^2\right)^2 \sigma_{x_3}^2 \\ &\quad + 2(1)(t) \underset{\sim}{\alpha} \underset{\sim}{\beta}^T \sigma_h^2 + 2(1)\left(-\frac{1}{2}t^2\right) \underset{\sim}{\alpha} \underset{\sim}{\gamma}^T \sigma_h^2 \\ &\quad + 2(t)\left(-\frac{1}{2}t^2\right) \underset{\sim}{\beta} \underset{\sim}{\gamma}^T \sigma_h^2 \end{aligned}$$

$$= \left\| \left[(1) \underset{\sim}{\alpha} + (t) \underset{\sim}{\beta} + \left(-\frac{1}{2}t^2\right) \underset{\sim}{\gamma} \right] \right\|_2^2 \sigma_h^2$$

If we represent the vector of modelling functions as $\underline{a}(t)$, we can collapse this:

$$\underline{a}(t) \equiv (1, t, -\frac{1}{2}t^2)$$

Then

$$\sigma_h^2 = \|\underline{a} \underline{K} \underline{a}\|_2^2 \sigma_h^2$$

Let's generalize the regression analysis we developed above.

We have n observations b_i which we wish to fit with m parameters and modelling functions:

$$b_i \approx \sum_{j=1}^m x_j \phi_j(t_i)$$

or:

$$\underline{b} \approx \underline{A} \underline{x}$$

where $A_{ij} = \phi_j(t_i)$

We define the residual:

$$\underline{r} = \underline{b} - \underline{A} \underline{x}$$

We define the best fit modelling parameters $\underline{x}$ by:

$$\min_{\underline{x}} \|\underline{r}\|_2^2 = \min_{\underline{x}} (\underline{r}^T \underline{r})$$

$\underline{x}$ has the solution:

$$\underline{A}^T \underline{A} \underline{x} = \underline{A}^T \underline{b}$$

or

$$\underline{x} = [\underline{A}^T \underline{A}]^{-1} \underline{A}^T \underline{b} = \underline{K} \underline{b}$$

Note that $\underline{A}$ is an $n \times m$ matrix, $\underline{b}$ is a $n \times 1$ array, $\underline{x}$ is a $m \times 1$ array and $\underline{K}$ is a $m \times n$ matrix

If we recall that:

$$\sigma_z^2 = E \{ (z - \mu_z)^2 \}$$

where $\mu_z = E(z)$

Then the matrix of covariance of the regression parameters is:

$$\sum_{\tilde{x}}^2 \equiv E\{(\tilde{x} - \mu_{\tilde{x}})(\tilde{x} - \mu_{\tilde{x}})^T\}$$

Substituting $\tilde{x} = K \tilde{b}$ and $\mu_{\tilde{x}} = K \mu_{\tilde{b}}$ we get:

$$\begin{aligned}\sum_{\tilde{x}}^2 &\equiv K E\{(\tilde{b} - \mu_{\tilde{b}})(\tilde{b} - \mu_{\tilde{b}})^T\} K^T \\ &= K \sum_{\tilde{b}}^2 K^T\end{aligned}$$

Where $\sum_{\tilde{b}}^2$ is the matrix of covariance of the $\tilde{b}$ observations $\tilde{b}$.

If the observations are independent then $\sum_{\tilde{b}}^2$ will be a diagonal matrix.

If all of the observations have the same variance, then:

$$\sum_{\tilde{b}}^2 = I \sigma_b^2$$

and we get

$$\sum_{\tilde{x}}^2 = \sigma_b^2 \tilde{K} \tilde{K}^T$$

We may estimate σ_b^2 from:

$$\begin{aligned} \sigma_b^2 &\approx s_b^2 = \frac{1}{n-m} \| \tilde{A} \tilde{x} - \tilde{b} \|_2^2 \\ &= \frac{1}{n-m} \tilde{r}^T \tilde{r} \end{aligned}$$

We may calculate the error in model predictions similarly. Suppose we wish to determine the model value at K times $\tilde{b}^{(m)}$. Thus:

$$\tilde{b}^{(m)} = \tilde{A}^{(m)} \tilde{x}$$

where $A_{ij}^{(m)} = \phi_j(t_i^{(m)})$

Note that $\tilde{A}^{(m)}$ is a $K \times m$ matrix.

What is the error in $\tilde{b}^{(m)}$? We have the matrix of covariance:

$$\sum_{\sim}^2 b^{(m)} \equiv E \left\{ (b_{\sim}^{(m)} - \mu_{\sim} b^{(m)}) (b_{\sim}^{(m)} - \mu_{\sim} b^{(m)})^T \right\} \quad (111)$$

$$= \underline{\underline{A}}^{(m)} E \left\{ (x_{\sim} - \mu_{\sim} x) (x_{\sim} - \mu_{\sim} x)^T \right\} \underline{\underline{A}}^{(m)T}$$

$$= \underline{\underline{A}}^{(m)} \sum_{\sim}^2 x \underline{\underline{A}}^{(m)T}$$

where $\sum_{\sim}^2$ is the matrix of covariance

determined before! In general, only the diagonal elements of $\sum_{\sim}^2$ are of interest:

$$\sigma_{b^{(m)}}^2 = \text{Diag} \left\{ \sum_{\sim}^2 b^{(m)} \right\}$$

which yields the error in the model predictions.

As a final note on linear regression, let us examine the case where the variance of σ_b^2 is not constant!

In this case we wish to weight the points in our regression analysis differently.

(112)

This is weighted linear regression!

Suppose we have the diagonal matrix $\sum_{\tilde{b}}$. The weighted problem is just:

$$\min_{\tilde{x}} (\tilde{A}\tilde{x} - \tilde{b})^T \left(\sum_{\tilde{b}}^{-1} \right) (\tilde{A}\tilde{x} - \tilde{b})$$

$$= \min_{\tilde{x}} \tilde{r}^T \left(\sum_{\tilde{b}}^{-1} \right) \tilde{r}$$

where $\left(\sum_{\tilde{b}}^{-1} \right)$ is a diagonal matrix whose elements are just $\frac{1}{\sigma_{b_i}^2}$.

Thus we weight each point by the inverse of the variance of that point. This makes sense. Suppose the reason why σ_{b_i} varied was that m_i points were averaged together to get b_i .

Then we have $\sigma_{b_i}^2 = \frac{\sigma_b^2}{m_i}$ and

thus the above regression yields:

$$\frac{1}{2} \min_{\tilde{x}} \tilde{r}^T \begin{pmatrix} m_1 & & & \\ & m_2 & & \\ & & \ddots & \\ 0 & & & m_n \end{pmatrix} \tilde{r},$$

or the i^{th} point is counted m_i times. The result for $\tilde{x}$ would be the same (on average) as if you had used all of the points individually!

OK, what is the formula for $\tilde{x}$?

Just as before we take the gradient:

$$\nabla_{\tilde{x}} \left\{ \tilde{r}^T \left(\sum_{i=1}^n m_i \right) \tilde{r} \right\} = 0$$

which yields:

$$A^T \left(\sum_{i=1}^n m_i \right)^{-1} A \tilde{x} = A^T \left(\sum_{i=1}^n m_i \right)^{-1} b$$

$$\text{or } \tilde{x} = K_w b = \left\{ A^T \left(\sum_{i=1}^n m_i \right)^{-1} A \right\}^{-1} A^T \left(\sum_{i=1}^n m_i \right)^{-1} b$$

with error:

$$\sum_{\tilde{x}}^2 = K_w \sum_{b}^2 K_w^T$$

Thus far we have considered error propagation / estimation techniques which rely on the regression formulas specific to linear regression. There are other numerical ways to estimate the error in the parameters, however, which apply to both linear and non-linear regression!

One approach is undersampling.

Suppose we have a series of observations $b_i(t_i)$. We wish to fit this with a model:

$$b_i \approx f(\underline{x}, t_i)$$

where f may be linear or non-linear in $\underline{x}$. If it's linear, we just use linear regression. If it's non-linear, we still do regression, only now it's non-linear regression.

We want to estimate the error

and covariance in $\underline{x}$, e.g.
we wish:

$$\sum_{\underline{x}}^2 \equiv E \left\{ (\underline{x} - \underline{\mu}_x)(\underline{x} - \underline{\mu}_x)^T \right\}$$

If the # of pts N is fairly large,
we can do this with undersampling

Suppose we take every m^{th} point.
This gives us a set of N/m points
from which we can calculate $\underline{x}$. If
we do this for the m indep. sets
of N/m points, we get m estimates
for $\underline{x}$!

We can then get a sample covariance
matrix:

$$\sum_{\underline{x}}^2 \approx \frac{1}{m-1} \left\{ \overline{x_i x_j} - \bar{x}_i \bar{x}_j \right\}$$

Note that this did not require
knowing the error in $\underline{x}$! It also

works equally well for linear and non-linear regression.

Its drawbacks are that it requires a pretty big N , and it requires multiple solutions for $\hat{x}$. That can take a while for non-linear problems!

A similar approach is the bootstrap, which works pretty well with moderate sized data sets.

Suppose we have, say, 20 observations $b_i(t_i)$. These observations are a subset of all possible observations. We can produce an approximation to this infinite set by periodically replicating our original data set!

The idea is to then take samples of 20 observations from this replicated set and calculate $\hat{x}$. We then compute $\sum_{\hat{x}}^2$ as before!

(117)

This technique is appealing, but only works if N is fairly large.

This is because there is a non-zero probability of getting the same point $b_i(t_i)$ N times out of our replicated set, which leads to an ill-posed regression problem.

Fortunately, this probability goes as

$$\sim \left(\frac{1}{N}\right)^{N-1}$$

which gets pretty small fairly fast!

The regression covariance is just:

$$\sum_{i,j}^2 x_i x_j = \frac{p}{p-1} \left\{ \overline{x_i x_j} - \bar{x}_i \bar{x}_j \right\}$$

where we are solving the regression problem p times.

As a final note on statistics, we have assumed throughout that the deviation between the observations and the model is random! This is, in general, **Not True!**

Thus, you tend to underestimate your error by ignoring the non-zero covariance in your data.

It is important to always plot your residuals to see if there is a systematic deviation.

A systematic deviation means that something is missing from your model and/or there is a systematic error in your data.

Never forget that for large N parameter error is dominated by systematic error!!

Non-Linear Regression

So far we've focussed on linear regression: problems where the model is linear in the unknown modelling parameters. This is convenient, but not really necessary! Suppose we have the non-linear model:

$$b = f(\underline{x}, t)$$

where $\underline{x}$ is an array of modelling functions. We may still define the residual between observed and fitted values:

$$r_i = b_i - f(\underline{x}, t_i)$$

$\uparrow$ $\uparrow$ $\uparrow$
 residual meas model

Thus we can form the sum of squares:

$$F(\underline{x}) = \sum_{i=1}^N (r_i)^2$$

The best fitted values for $\underline{x}$ are obtained via finding the minimum of $F(\underline{x})$!

Be Wary: Non-linear optimization problems may have local minima which can trap the solver! Also, make sure you have set up the parameters so that $F(\underline{x})$ is well conditioned. The dependence on each parameter should be of similar magnitude. Let's look at an example:

Arrhenius Kinetics

$$\text{rate} \sim K_0 e^{-E/RT}$$

In senior lab you will measure the rate of oxidation of methane as a function of temperature, and use it to try to get the activation energy E for the catalyst. If we could measure this rate, holding everything but T constant, we could use linear regression. Unfortunately, the rate also depends on concentration, and that depends (experimentally) on T as well!

The exp't is written up in Chem Eng. Education, 36 (1) p. 34-40. Suppose we feed a reactor (well-mixed) a

concentration of methane C_r , at flow rate Q_r . We measure some outlet concentration C_r . The subscript "r" is because the reactor is hot: due to expansion, the actual concentration is $\frac{C}{C_r} = \frac{T_r}{T}$ from the ideal gas law - where T_r is some ref. temp ($^{\circ}\text{K}$).

The system is further complicated because the reaction is fractional order, e.g.

$$\text{rate} \sim C^n$$

where $n \neq 1$ (not first order)

we need to figure out n too!

From a mass balance, we get the rate of rxn per gram of catalyst:

$$\frac{q_r}{m} (C_{r0} - C_r) = \left(C_r \frac{T_r}{T} \right)^n K_0 e^{-\frac{E}{RT}}$$

where m is the mass of catalyst.

We can also look at the conversion ratio given by:

$$X = 1 - \frac{C_r}{C_{r0}}$$

which yields (after rearrangement):

$$\frac{X}{(1-X)^n} = \frac{m}{q_r} C_{or}^{n-1} \left(\frac{T_r}{T} \right)^n K_0 e^{-\frac{E}{RT}}$$

If you were to fix T and plot

X vs. C_{or} , you find it decreases

w/ C_{or} , thus $n < 1$ for this system.

Ok, suppose we vary C_{w0} and T , and measure C_w . How do we get the unknowns n , K_0 , and E ?

The classic approach is to do it in two steps:

1) Fix T and plot

$$\ln \left\{ \frac{q_w}{m} (C_{w0} - C_w) \right\} \text{ vs. } \ln \{ C_w \}$$

You should get a straight line with slope n !

2) With this in hand, you can then do an Arrhenius Plot of:

$$\ln \left\{ \frac{q_w}{m} \frac{(C_{w0} - C_w)}{C_w^n} \left(\frac{T_r}{T} \right)^{-n} \right\} = \ln K_0 - \frac{E}{R} \frac{1}{T}$$

Thus, the slope (vs. $\frac{1}{T}$) is just $\frac{E}{R}$ and the intercept is $\ln k_0$! 125

This works fine if the data is good. Unfortunately, this is not usually the case. The problem is that C_p is a much stronger function of T than C_{p_0} , and thus you get large errors in the first step (n), leading to large errors in the second step (E & k_0).

The fitting parameters are also strongly biased by errors at low conversions.

We can avoid this by using non-linear regression! We simply define the deviation from the model Δ :

$$\Delta = C_v - C_{v_0} + \frac{m}{\sum_r} C_v^n K_0 e^{-\frac{E}{RT}} \left(\frac{T_w}{T} \right)^n \quad (12C)$$

and minimize the objective function:

$$F(K_0, E, n) \equiv \sum \Delta_i^2$$

over all the experiments! The sensitivity of the fitting parameters to error can then be easily calculated using the non-linear error propagation/gradient method, accounting for the error in both C_v and T . Note that you should work with n , $\ln K_0$, and $\frac{E}{RT_w}$ as fitting parameters so that the optimization problem is well-conditioned!