

THE ABNORMAL NORMAL DISTRIBUTION.

- ① Gaussians are ubiquitous
- ② Information Entropy
- ③ Gauss' derivation
- ④ Error cancellation
- ⑤ Ubiquity explained
- ⑥ Properties - Gaussians are special
- ⑦ Note on naming
- ⑧ History?
- ⑨ Herschel derivation

The extreme abnormality of the normal distribution 1.

①

Gaussians are everywhere. At the heart of modern physics:

Path integral: $\int e^{i\hbar (KE - PE)} d\phi$

$\uparrow$
 $\frac{1}{2} m \phi^2$

Feynman diagrams
are: Taylor expansion of
the PE term...
ie can do Gaussian integral,
interaction terms are tricky...

Gaussian distribution is also very common
model - stock price, swap rate... but most common application in history
of science is as an "error law" (erf(x)!).

└ assumed frequency distribution of errors.

But, are errors Gaussian? "Willy" Feller ~1955 after dinner speech
anecdote (Jaynes pg 198)... Even 1983... (George Barnard)
└ quality control guy.

② Groundwork: Gaussian form has higher "information entropy" than any
other for a given variance. Information entropy is really "amount of
uncertainty". Want to assign a distribution that reflects our state of
knowledge as fully as possible, which is ignorance, other than, say,
average size of the errors. how "spread out" a dist'n is... *

Amount of uncertainty is related to # ways of getting a given
distribution - we should assign one that is realizable in a number
of ways so much greater than that of nearest neighbour that it is
overwhelmingly the most likely.

* want to be as "impartial" as
possible, as true to the principle

namer → Keynes 1921 of ~~ignorance~~ indifference as possible...

Amount of uncertainty in distribution $p(x)$ is:

2.

$$H_p = - \int_{-\infty}^{\infty} p(x) \log p(x) dx. \quad (\text{Shannon 1948})$$

If know nothing other than range (or finite set of possibilities in discrete case), variational calculus gives us same dist'n as intuition: $p(x) = \frac{1}{b-a}$.

$$L_0[p(x)] = - \int_a^b p(x) \log p(x) dx + \lambda \left(\int_a^b p(x) dx - S(b) \right)$$

↑
Lagrange multiplier

$$\delta L_0[p(x)] = - \int_a^b \left(\log p(x) + \frac{p(x)}{p(x)} - \lambda \right) \delta(p(x)) dx = 0$$

Add more constraints,

still have the $\log p(x)$ term which leads to a family of exponential distributions including just about everything. Poisson, Γ , χ^2 , Binomial etc etc....

So, Gaussian:

$$L_2[p(x)] = \int_{-\infty}^{\infty} \left[-p(x) \log p(x) + \alpha(p(x) - 1) + \beta(x p(x) - \mu) + \gamma(x^2 p(x) - \nu) \right] dx$$

x dx

actually $\delta(x)$

$$\int_a^b e^{\lambda-1} dx = 1$$

$$\therefore e^{\lambda-1} = \frac{1}{b-a}$$

$$\delta L_2[p(x)] = 0 \Rightarrow -\log p(x) - 1 + \alpha + \beta x + \gamma x^2 = 0$$

$$\therefore p(x) = e^{\gamma x^2 + \beta x + \alpha - 1}$$

Have "derived" Gaussian form as dist'n best representing knowledge of variance. Exhibit A. Now for Gauss' derivation: 3.

③ Estimate θ from $n+1$ observations $(x_0 \dots x_n)$ by max Likelihood,

$$p(x|\theta) = \prod_{i=0}^n f(x_i|\theta)$$

θ and x have same scale

Write as $\frac{d}{d\theta} \log p(x|\theta) = \sum_{i=0}^n \frac{d}{d\theta} \log f(x_i|\theta) = 0$.

then write $\log f(x|\theta) = g(\theta - x) = g(u)$ $u = \theta - x$

"location" parameter

so MLE condition becomes $\sum_i \frac{dg(\hat{\theta} - x_i)}{d\theta} = 0$ for MLE $\hat{\theta}$.

This is all very fancy, depends on unknown dist'n f (ie g)...

What if $\hat{\theta} = \frac{1}{(n+1)} \sum_{i=0}^n x_i$? Turns out requiring arithmetic mean is

MLE uniquely determines $f(x|\theta)$ as Gaussian.

Haven't constrained to a particular set of values x - can try it out with an extreme, if unlikely, realization, to probe things:

$$x_0 = (n+1)u, \quad x_1, \dots, x_n = 0$$

then $\hat{\theta} = u$ and $\hat{\theta} - x_0 = -nu$ and MLE condition is:

$$\hat{\theta} - x_i = u \quad i > 0. \quad \frac{d}{d\theta} g(-nu) + n \frac{d}{d\theta} g(u) = 0$$

$n=1$ case: $\frac{dg(-u)}{d\theta} = - \frac{dg(u)}{d\theta}$, ie $\frac{dg}{d\theta}$ is antisymmetric,

so have $\frac{dg(nu)}{d\theta} = n \frac{dg(u)}{d\theta}$ so $\frac{dg}{d\theta}$ is linear.

and $\frac{d}{d\theta} = \frac{d}{du}$, so $\frac{dg(u)}{du} = au + b \rightarrow g(u) = \frac{1}{2}au^2 + b$

redefine!

4.

$$\log f(x|\theta) = g(u) = g(\theta - x) = \frac{1}{2} a u^2 + b \quad \text{scale param}$$

$$\text{so } \underline{f(x|\theta) = e^{\frac{1}{2} a (x-\theta)^2 + b}}$$

solve for $b(a)$ via normalization, which will diverge unless $a < 0$.

But assumed a special sample. Necessary form for MLE $\leftrightarrow$ arithmetic mean.
 Suppose instead it is "the general solution"? No, assume the distribution in MLE condition and solve for $\hat{\theta}$:

$$f(x|\theta) = \frac{1}{\sqrt{2\pi\sigma^2}} e^{-\frac{(x-\theta)^2}{2\sigma^2}}$$

$$\sum_i \left. \frac{dg(\theta-x_i)}{d\theta} \right|_{\hat{\theta}} = \sum_i \frac{x_i - \hat{\theta}}{2\sigma^2} = 0 = \frac{1}{2\sigma^2} \left(\sum_i x_i - \hat{\theta}(n+1) \right)$$

$$\therefore \underline{\hat{\theta} = \frac{1}{n+1} \sum_{i=0}^n x_i}$$

remember Daniel Bernoulli's arrows...

ie Gaussian form is sufficient for MLE $\leftrightarrow$ arithmetic mean, as well.

④ Nothing 100% fundamental about arithmetic mean as an estimator though. There are lots of ^{possible} estimators, varying by distribution. Can consider:

$$\beta(\underline{y}) = \sum_{i=1}^n w_i(\underline{y}) y_i$$

Easier now if $n+1 = n'$
 $x_0 \dots x_n = y_1 \dots y_{n'}$
 $n' \rightarrow n$

and assess quality of estimate given $f(y|\theta)$ Gaussian

I.I.D. $\rightarrow p(\underline{y}|\theta) = f(y_1)f(y_2)\dots f(y_n)$

How assess quality? $\searrow$

$$\Psi = \langle (\beta(\underline{y}) - \theta)^2 \rangle \quad \text{where } \langle h(\underline{y}) \rangle = \int_{-\infty}^{\infty} h(\underline{y}) p(\underline{y}|\theta) d\underline{y}$$

Define i^{th} error: $\epsilon_i = y_i - \theta \rightarrow \left\langle \left(\sum_{i=1}^n w_i(\underline{y}) \epsilon_i \right)^2 \right\rangle$

$$\therefore \Psi = \sum_{i,j=1}^n w_i w_j \underbrace{\langle e_i e_j \rangle}_{\sigma^2 \delta_{ij}}$$

$$= \sigma^2 \sum_i w_i^2$$

$\sigma^2 \delta_{ij}$ IID Gaussian,
drop y -dependence of weights for now.

Variational problem:

$$L(\underline{w}) = \sigma^2 \sum_{i=1}^n \left(w_i^2 + \lambda \left(w_i - \frac{1}{n} \right) \right) \rightarrow 2w_i + \lambda = 0$$

$$w_i = -\frac{\lambda}{2} \text{ (constant)}$$

$$\sum_i w_i = -\frac{\lambda}{2} n = 1 \text{ so } w_i = \frac{1}{n}$$

$$\therefore \beta(\underline{y}) = \frac{1}{n} \sum_i y_i$$

⑤ Take stock:

- (i) Gaussian is maximum entropy distribution under mean and variance constraints. Most spread out a distribution can be under the constraints. Assigns maximum possible probability to each value, indifferently, while conserving variance. Most honestly represents our state of knowledge. Ignores distribution of errors beyond first 2 moments...
 - (ii) Requiring the arithmetic mean to maximize likelihood uniquely determines a Gaussian form. more opportunity for error cancellation
 - (iii) Estimate of parameter via arithmetic mean is more accurate (minimizes expected square error Ψ) for a Gaussian than for any other distribution.
- So a Gaussian is the distribution we assign to errors if we want the most accurate estimates from arithmetic means, given the general magnitude of the noise.

⑥ "Gaussian form really is odd:

A: Product of 2 Gaussians is Gaussian

B: Fourier Transform of G is G .

C: Convolution of 2 G 's is G

D: G is limiting form of various others: binomial, Poisson, beta...

E: Any hump raised to higher and higher powers $\rightarrow G$.

F: Any hump under repeated convolution $\rightarrow G$ (central limit theorem)

CLT is usually given as justification for Gaussian noise assumption, but it's really a special case of something deeper, as as E, D.

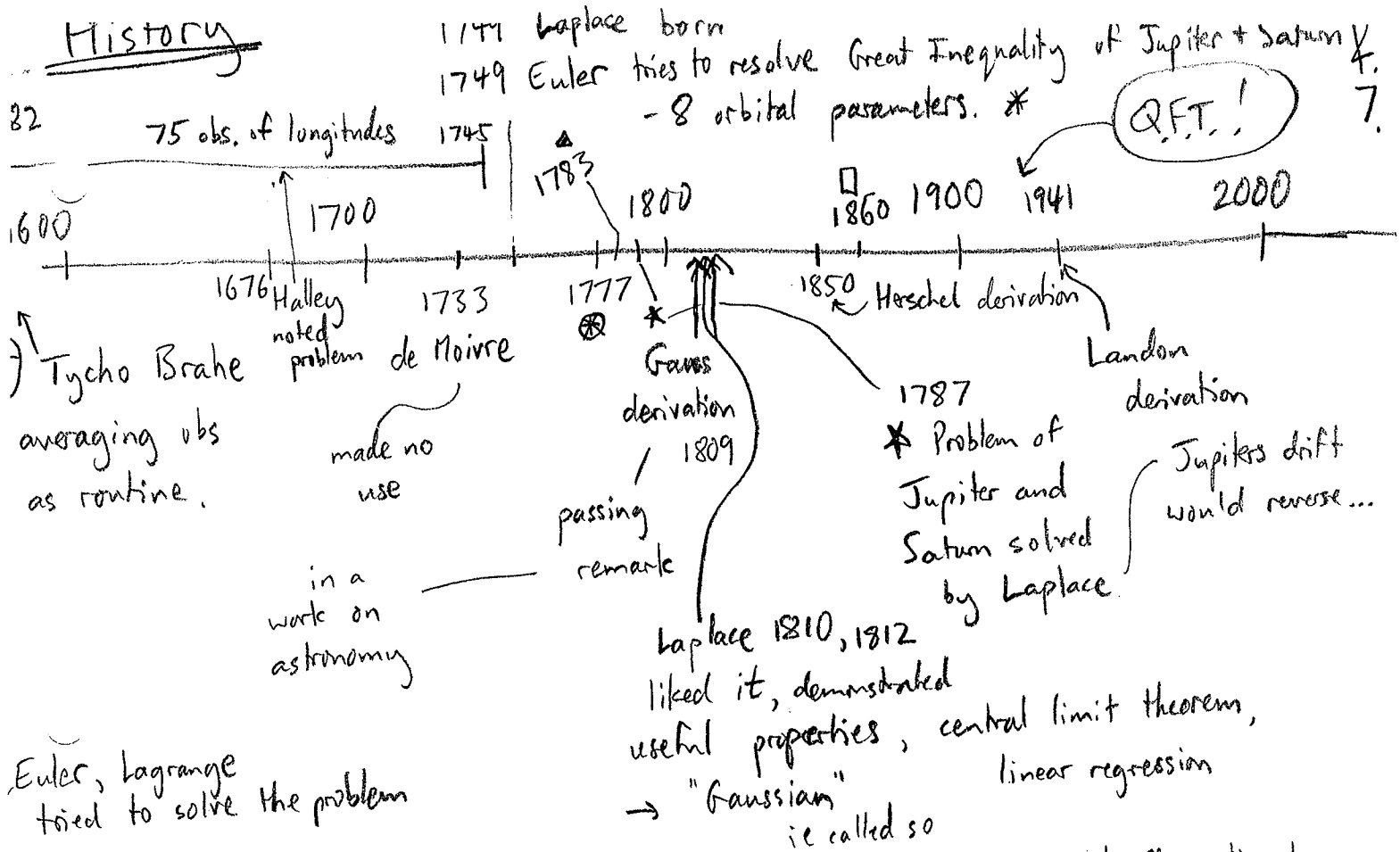
Gaussian is max-ent dist'n, so any transformation which discards information (detail, certainty, preference of some possibilities over others) while preserving variance, leads closer to Gaussian form.

⑦ Can't say this about any distribution except Gaussian. It really is weird! Not very "normal". Origin of term is unclear but probably geometrical in origin, in context of regression - Gauss described something as perpendicular to surfaces of const. probability or similar.

Jaynes suggests "central".

⑧ History?!

History



* combining the ~~errors~~ observations → "multiply the errors".

Which 8 equations to pick? Got 2, known anyway, French Ac. of Sciences prize 1749

yet: (+)

* Daniel Bernoulli says: obviously probability of distance of arrow from vertical line target is not uniform, so nonsense to weight uniformly when estimating target location from arrow positions.

▲ Laplace shows binomial → Gaussian derived main properties suggested tabulation

but failed to see Gaussian was the natural error law - still used $p(x) \propto e^{-ax^2}$

but got the combination of observations idea so *

argument (Maxwellian)
□ Maxwell 3rd version of same thing → distribution of velocities of gas molecules

P1: No preferred direction: $p(x, y) dx dy = f(x) dx \cdot f(y) dy$

P2: No preferred angle: $g(r, \theta) = g(r)$ in $p(x, y) dx dy = g(r, \theta) r dr d\theta$

$$\text{so } f(x) f(y) = g(\sqrt{x^2 + y^2})$$

Must still be true if $y=0 \rightarrow g(x) = f(x) f(0)$

$$f(x) f(y) = f(\sqrt{x^2 + y^2}) \cdot f(0)$$

$$\text{or } \frac{f(x)}{f(0)} \frac{f(y)}{f(0)} = \frac{f(\sqrt{x^2 + y^2})}{f(0)}$$

$$\frac{f(x)}{f(0)} = e^{ax^2} \text{ does the trick.}$$

↑
1850

Maxwell 1860 - 3-d version of above arg

→ Maxwellian distribution of gas molecules.