

INDIAN BAYESIAN SOCIETY

NEWSLETTER

Vol. XXIII, No. 1, May, 2026

Advisory Committee
Members:

Tapas Samanta
Professor
Applied Statistics Unit
Indian Statistical Institute
Kolkata

Umesh Singh
Retired Professor
Department of Statistics
Banaras Hindu University
Varanasi

A. A. Khan
Professor
Department of Statistics
& Operations Research
Aligarh Muslim University
Aligarh

Mohan Delampady
Professor
Theoretical Statistics and
Mathematics Unit
Indian Statistical
Institute, Bangalore

A. Loganathan
Retired Professor
Department of Statistics
Manonmaniam
Sundaranar University
Tirunelveli

S. K. Upadhyay
Retired Professor
Department of Statistics
Banaras Hindu University
Varanasi

From the Desk of Professor Dipak Dey

A Self-Organized Criticality Model of Extreme Events and Cascading Disasters: A New Statistical View

Dear Fellow Statisticians,

It is striking to some extent that statisticians get involved in applied problems in a rather arbitrary fashion based on personal connections and whether the statistician's personal methodological research fits with the applied problem. Are statisticians sufficiently involved in the most important scientific problems in the world and if not, is there some mechanism that could be developed by which we as a profession can make our expertise available to the scientists tackling those problems? How do we close the gap between the sophistication of methods developed in our field and the simplicity of methods used by many practitioners?

In this issue, I will introduce a relatively new topic of research, which is extremely important for our society and is rarely touched by statisticians.

Introduction

Extreme events and catastrophic failures—including infrastructure collapse (e.g., bridges, buildings, roads), climate-related disasters, systemic power outages, and organizational breakdowns—are increasing in frequency, scale, and complexity. Conventional statistical approaches often rely on assumptions of independence, linearity, and thin-tailed variation. These assumptions may be inadequate for modern systems that are highly interconnected, adaptive, and vulnerable to cascading disruptions.

This article presents **Self-Organized Criticality (SOC)** as a scientifically grounded framework for understanding extreme risks. By integrating concepts from physics, probability, spatio-temporal statistics, environmental science, and computational

Contact Address: S. K. Upadhyay, Department of Statistics, Banaras Hindu University, Varanasi-221 005.

Phone (O): (91-542) 6702892; (R): (91-542) 2214000, 2210700; Fax: (91-542) 2368174; e-mail: ibs2003@gmail.com.

data analysis, SOC provides a unified explanation for how small perturbations may trigger failures of arbitrary magnitude.

Why Traditional Models Often Fail?

Many classical risk models assume:

- Independent component failures
- Gaussian or light-tailed variability
- Linear cause-and-effect relationships
- Rare extremes as outliers

However, real systems frequently exhibit:

- Strong dependence across components
- Heavy-tailed event distributions
- Nonlinear feedback mechanisms
- Cascading failures across networks
- Extreme events as intrinsic outcomes

These features motivate a broader modeling paradigm.

What is Self-Organized Criticality?

Self-Organized Criticality describes how complex systems naturally evolve toward a **critical state** where they remain apparently stable, yet increasingly fragile. In this regime:

- Small disturbances can trigger events of any size
- Large events are not frequent but unavoidable
- No characteristic event scale exists
- System-wide behavior follows the laws of scale-invariant laws

A common statistical signature is:

$$P(S > s) \sim s^{-\alpha}$$

where: S = event size, s = threshold, α = critical exponent governing tail behavior.

Lower values of α imply heavier tails and greater systemic fragility.

Physical Intuition: The Sandpile Paradigm

The classical sandpile model remains the simplest illustration of SOC:

1. Grains are added sequentially.
2. The pile steepens while appearing stable.
3. Internal stress accumulates.
4. One additional grain may trigger a minor slip or a major avalanche.

The key lesson is profound: systems may drift toward instability through ordinary incremental inputs.

SOC in Interconnected Systems

Many modern systems can be represented as networks with nodes and dependencies:

- Power grids
- Transportation systems
- Supply chains
- Financial markets
- Communication networks
- Climate-linked infrastructure

Typical SOC dynamics include:

- Gradual stress accumulation
- Threshold-triggered local failure
- Redistribution of load to neighbors

- Recursive cascades
- Emergent large-scale collapse

Thus, localized failures can propagate into system-wide disruption.

Real World Relevance

Recent disasters reveal a repeating pattern:

- Aging infrastructure weakened over time
- Environmental stressors intensified loads
- Protective systems exceeded design capacity
- Communication delays worsened outcomes
- Interdependent systems failed sequentially

Natural hazards often become catastrophic when they interact with pre-existing structural fragility.

Interpreting the Critical Exponent

The exponent α can be viewed as a quantitative stress indicator:

- **Higher** α : failures tend to remain localized
- **Intermediate** α : broad range of cascade sizes possible
- **Lower** α : extreme system-wide failures become more probable

Monitoring changes in α may provide an early warning signal of increasing vulnerability.

Applications

1. Power Systems

Electric grids operate under continuously varying demand, weather exposure, and operational constraints.

SOC models suggest that persistent declines in stability indicators may precede major blackouts.

2. Air Transportation

Hub-and-spoke networks are efficient but fragile. Minor disruptions at critical hubs can propagate nationally through delay cascades.

3. Climate and Environmental Risk

Climate change may amplify SOC dynamics by increasing both stress frequency and propagation intensity, producing compound disasters across multiple sectors.

4. Athlete Safety and Human Performance

Athletes may also be modeled as adaptive systems balancing:

- workload
- recovery
- hydration
- heat stress
- sleep
- injury burden
- psychological pressure

Collapse may occur not from one extreme exposure, but from cumulative stress crossing a physiological threshold.

Toward an Early Warning Framework

An SOC-based monitoring system would aim to:

- Track system-wide stress indicators
- Estimate tail behavior dynamically
- Detect transitions toward criticality

- Simulate never-before-seen extreme scenarios
- Support preventive intervention before collapse

This shifts risk management from reactive response to anticipatory resilience.

Implications for Statistics

SOC invites statisticians to develop new methodologies for:

- Inference under heavy tails
- Dependent extremes
- Dynamic network cascades
- Bayesian critical-state monitoring
- Real-time anomaly detection
- Uncertainty quantification for rare disasters

This is a fertile research frontier for Bayesian science.

Conclusion

Extreme events are often not random anomalies. They may be the expected consequences of complex systems operating near critical thresholds.

Self-Organized Criticality offers:

- A coherent theory of cascading disasters
- A statistical explanation for heavy-tailed extremes
- Practical tools for early warning and prevention
- A bridge between physical science and modern data analytics

Understanding fragility is the first step toward resilience. The SOC perspective encourages us to view disasters not as isolated

surprises, but as emergent outcomes of interconnected systems under accumulated stress.

Future Directions

- Robust Bayesian inference for SOC parameters
- Sequential estimation of critical exponents
- Integration with sensor and streaming data
- Infrastructure resilience analytics
- Athlete protection monitoring systems
- Climate-linked systemic risk forecasting

References

Per Bak (1996), *How Nature Works: The Science of Self-Organized Criticality*. Copernicus Books. <https://www.amazon.com/How-Nature-Works-self-organized-criticality/dp/038798738X>.

M. E. J. Newman (2005), Power laws, Pareto distributions and Zipf's law. *Contemporary Physics*, 46 (5), September–October 2005, 323 - 351.

Didier Sornette (2006), *Critical Phenomena in Natural Sciences*. Springer-Verlag. <https://link.springer.com/book/10.1007/3-540-33182-4>.

National Center for Catastrophic Sport Injury Research. <https://nccsir.unc.edu/>.

Sandpile Model. https://itp.uni-frankfurt.de/~gros/StudentProjects/WS22_23.SandpileModel/.

– Dipak Dey
University of Connecticut
Storrs, CT
Dipak.dey@uconn.edu

What is Statistical Evidence?

Dedicated to the memories of Professor D.
A. S. Fraser and Professor Irwin Guttman

Professor Michael Evans
Department of Statistical Sciences
University of Toronto
mevansthree.evans@utoronto.ca

I had the good fortune to write an article for the Indian Bayesian Society Newsletter in 2015. In that Newsletter there were also contributions by Professor J. K. Ghosh and by Professor Malay Ghosh, so I felt very honored to have my article included. Prof. Malay Ghosh's article was entitled "A Bayesian 21st Century". His article commented on some historical prejudices that had existed in the statistical community against Bayesian inference, and how this was changing to a more positive view. I definitely think that this is the case, and very much so since 1977 when I started my career as a university professor. Actually, I was trained as a frequentist statistician and that is how I addressed statistical problems early in my career and I'll confess to having that negative prejudice when I graduated.

Prof. Don Fraser was my supervisor and, at that time, his primary research direction was involved with what he called structural inference. Structural inference was a formalization of some aspects of Fisher's fiducial inference. Structural inference is not far from Bayesian inference as it can be considered as such, provided one accepts the use of improper priors. In the case of a structural model, the prior is the right Haar prior on the group associated with this model. Structural inference struck me then, and still does, as a very beautiful theory and it convinced me, as a former student of mathematics, that statistical inference was a subject worth studying further. Professor Fraser has since passed on, but he remained active in statistics well into his 90's. I'll always be grateful to him for convincing me that the

foundations of statistics was an important subject. For those interested in structural inference, I recommend his book [Fraser \(1979\)](#). This book also seems to be one of the first to begin to mathematically formalize a theory of statistical inference that did not sit within decision theory, as no loss function was incorporated.

As I worked at my research in the field of statistics, I became convinced that some of the basic issues arising in the foundations could not be resolved within the context of what is commonly called frequentist statistics. I will have more to say about frequentism, and its role in inference, later in the article. As such, I began to investigate other approaches and, due largely to the influence of another mentor, Professor Irwin Guttman, I slowly began to see that Bayesian inference offered a natural solution to many problems. Professor Guttman was also a student of Professor Fraser and he was definitely a Bayesian. His guidance was pivotal for me beginning to think much more seriously about Bayesian inference. I will admit that I found this a somewhat confusing endeavor at times. There seemed to be a plethora of different Bayesian approaches and perhaps this is reflected in I. J. Good's claim that there are 46,656 different types of Bayesians, see [Good \(1971\)](#). Sadly, Professor Guttman recently passed away, and so our joint work has ended, but I will always be indebted to him.

Today I am definitely a Bayesian, but maybe of an even additional variety to the ones Good counted. This is because for me the central concept for the subject of statistics is *statistical evidence*. Certainly the words evidence or statistical evidence are commonly used as part of statistical analyses. So, it is very natural to ask, what is statistical evidence? By this I mean, can a precise definition of this concept be provided? Certainly, I am not the first to consider this question and, in fact, I think the answer has been

clear for many years and well-discussed by others. I will state the answer via a principle of inference that I call the *principle of evidence*.

Let us suppose we have a probability model $(\Omega, \mathcal{A}, P)$ and we are interested in whether or not an unobserved $\omega \in \Omega$ actually satisfies $\omega \in A$ where $P(A) > 0$, and we are told unambiguously that $\omega \in C$ where also $P(C) > 0$. Then by the basic principle of conditional probability, our beliefs in the truth of $\omega \in A$ are now changed from $P(A)$ to $P(A|C)$. This leads to the following.

Principle of Evidence: If $P(A|C) > P(A)$, then there is evidence in favor of A being true, if $P(A|C) < P(A)$, then there is evidence that A is false, and if $P(A|C) = P(A)$, then there is no evidence either way.

Clearly $P(A|C) = P(A)$ iff the events A and C are independent so indeed, in that case, observing that C is true can tell us nothing about the truth of A . I do not know of a simpler characterization of evidence than this principle. It seems obvious that, if what we have learned to be true has caused our beliefs to increase (decrease), then what has been observed to be true has provided evidence in favor of (against) the truth of the event in question. After all, that is what evidence does, it changes our beliefs. This basic characterization has long been studied in the philosophy of science under the rubric, confirmation theory.

For the principle of evidence to apply it is necessary to have a full probability model. In the proper Bayesian context this applies and for the moment suppose that all probabilities are discrete. In this case, we have the inference base $(\{f(\cdot|\theta) : \theta \in \Theta\}, \pi, x)$ comprised of the model $\{f(\cdot|\theta) : \theta \in \Theta\}$, a collection of probability densities on a sample space $\mathcal{X}$ with respect to some support measure and indexed by a model parameter $\theta \in$

Θ , a prior π on Θ , and finally the observed data x . This leads to a full probability model for (θ, x) with density $\pi(\theta)f(x|\theta)$ and then the principle of conditional probability leads to the posterior $\pi(\theta|x) = \pi(\theta)f(x|\theta)/m(x)$, where $m(x) = \sum_{\theta \in \Theta} \pi(\theta)f(x|\theta)$ is the prior predictive density of the data. Generally, there is a parameter of interest $\psi = \Psi(\theta) \in \Psi(\Theta)$ and our goal is to make inference (hypothesis assessment, estimation) about this quantity. The inference base for θ is then replaced by the inference base for ψ , namely,

$$(\{m(\cdot|\psi) : \psi \in \Psi(\Theta)\}, \pi_\Psi, x)$$

where $m(x|\psi) = \sum_{\theta \in \Psi^{-1}\{\psi\}} \pi(\theta|\psi)f(x|\theta)$, $\pi(\cdot|\psi)$ is the conditional prior on θ given that $\Psi(\theta) = \psi$, namely, it is the prior on the nuisance parameters, and $\pi_\Psi(\psi) = \sum_{\theta} \pi(\theta)$ ($\theta \in \Psi^{-1}\{\psi\}$) is the marginal prior on ψ . All inferences about ψ are then based on $(\{m(\cdot|\psi) : \psi \in \Psi(\Theta)\}, \pi_\Psi, x)$. It is immediate from the principle of evidence that there is evidence in favor of (against, none either way) ψ when $\pi_\Psi(\psi|x) > (<, =) \pi_\Psi(\psi)$.

It isn't sufficient to only determine whether there is evidence in favor or against some value of ψ , it is also necessary to say something about the strength of the evidence, namely, is it strong, moderate or weak evidence. This seems best addressed via a posterior probability that measures how strongly we believe what the evidence says. So, for example, suppose we have evidence in favor of ψ_0 being true, so $\pi_\Psi(\psi_0|x) > \pi_\Psi(\psi_0)$ and, in the discrete context, $\pi_\Psi(\psi_0|x)$ is the posterior probability of ψ_0 . If this posterior probability is large, then we have strong evidence in favor of ψ_0 , etc. Actually quoting $\pi_\Psi(\psi_0|x)$ doesn't seem adequate in many contexts because $\pi_\Psi(\psi_0|x)$ may be small simply because the cardinality of $\Psi(\Theta)$ is large. Furthermore, since evidence is measured by change in belief, it seems more appropriate to look at how much beliefs have changed when considering strength. To address this a numerical measure of evidence

seems essential and one possibility is the *relative belief ratio*

$$RB_{\Psi}(\psi_0 | x) = \frac{\pi_{\Psi}(\psi_0 | x)}{\pi_{\Psi}(\psi_0)}. \quad (1)$$

Note that this is a valid measure of evidence because there is a cut-off, determined by the principle of evidence, so that $RB_{\Psi}(\psi_0 | x) > (<= 1)$ means evidence in favor of (against, none either way) ψ_0 . Also, if $RB_{\Psi}(\psi_1 | x) > RB_{\Psi}(\psi_2 | x) > 1$, then there is more evidence in favor of ψ_1 than ψ_2 , because the data has lead to a greater proportional increase in belief for ψ_1 than for ψ_2 and when $RB_{\Psi}(\psi_1 | x) < RB_{\Psi}(\psi_2 | x) < 1$, then there is more evidence against ψ_1 than against ψ_2 , because the data has lead to a greater proportional decrease in belief for ψ_1 than for ψ_2 . A natural measure of the strength of the evidence provided by $RB_{\Psi}(\psi_0 | x)$, is then given by the posterior probability

$$\begin{aligned} Str_{\Psi}(\psi_0 | x) = \\ \Pi_{\Psi}(RB_{\Psi}(\psi | x) \leq RB_{\Psi}(\psi_0 | x) | x). \end{aligned} \quad (2)$$

Further, if there is evidence in favor of ψ_0 and $Str_{\Psi}(\psi_0 | x) \approx 1$, then there is small belief that the true value of ψ has a larger relative belief ratio than ψ_0 , and so there is strong evidence in favor of ψ_0 . If there is evidence against ψ_0 and $Str_{\Psi}(\psi_0 | x) \approx 0$, then there is large belief that the true value has a larger relative belief ratio and so we have strong evidence against ψ_0 , etc.

For continuous parameters the problem of measuring statistical evidence is approached via limits. So, if $B_{\epsilon}(\psi_0)$ is a sequence of neighborhoods that converges (nicely) to ψ_0 as $\epsilon \rightarrow 0$, then we define

$$\begin{aligned} RB_{\Psi}(\psi_0 | x) &= \lim_{\epsilon \rightarrow 0} RB_{\Psi}(B_{\epsilon}(\psi_0) | x) \\ &= \frac{\pi_{\Psi}(\psi_0 | x)}{\pi_{\Psi}(\psi_0)} \end{aligned}$$

where the final equality holds whenever π_{Ψ} is continuous and positive at ψ_0 . So, the basic form (1) is maintained generally. The

strength of the evidence is then measured as in (2).

The above discussion focuses on the assessment of the hypothesis $H_0 : \Psi(\theta) = \psi_0$. There is still more to this story as it is relatively easy to incorporate the *difference that matters* δ into the assessment that arises with continuous parameters. By this is meant that the hypothesis now assessed is $H_0 : \Psi(\theta) \in (\psi_0 - \delta/2, \psi_0 + \delta/2)$, at least when ψ is real-valued, as values differing by less than $\delta/2$ are equivalent in terms of the application. The value of δ is application dependent, although several values can be considered as part of assessing sensitivity to the choice. For assessing the strength of the evidence, a grid of values $\dots, \psi_0 - \delta, \psi_0, \psi_0 + \delta, \dots$, is formed, with their prior and posterior probabilities determined from the corresponding intervals that surround them. This can be generalized to what is called a *regular discretization*, see [Evans and Jang \(2024\)](#), for more general spaces. Note that this discretization is relatively easy in Bayesian contexts and it avoids the “large sample paradox” associated with p-values for point nulls. For, in the continuous case, all such point nulls are false, so if you take a large enough sample size, a p-value will give evidence against even though the difference detected is of no practical importance. In general, it seems cumbersome to incorporate such a δ with p-values, but with relative belief it is easy. This discretization even simplifies computations as implementation only requires samples from the prior and posterior and there is no need to estimate continuous densities.

For estimation problems it seems clear that a possible estimate is $\psi(x) = \arg \sup RB_{\Psi}(\psi | x)$, as this is the value with the most evidence in its favor. Notice that this is the MLE of ψ based on the model $\{m(\cdot | \psi) : \psi \in \Psi(\Theta)\}$. To assess the accuracy of the estimate there is the plausible region $Pl_{\Psi}(x) = \{\psi : RB_{\Psi}(\psi | x) > 1\}$, namely, the set of val-

ues having evidence in their favor. The size of $Pl_{\Psi}(x)$, together with its posterior probability, provide an assessment of how accurate we believe the estimate is. Notice that $Pl_{\Psi}(x)$ is a likelihood region for the model $\{m(\cdot|\psi) : \psi \in \Psi(\Theta)\}$. In fact, relative belief inferences always respect the likelihood ordering. An advantage that relative belief inferences have over likelihood inferences, is that they are based on a clear definition of statistical evidence.

There is yet another consideration that needs to be taken into account. The Bartlett paradox, often seen as part of the Jeffreys-Lindley paradox, can cause a real problem if we are not careful. This problem arises when diffuse priors are used when assessing a hypothesis $H_0 : \Psi(\theta) = \psi_0$, namely, as the prior becomes more diffuse, in some problems $RB_{\Psi}(\psi_0|x) \rightarrow \infty$. So, for a diffuse enough prior, evidence in favor of H_0 will be obtained even when it is meaningfully false. This problem is not avoided by discretization. There is a natural way to deal with problems like this and, moreover, it is an essential aspect of a statistical argument. We are referring here to experimental design which, at its most basic, answers the question, how much data do we need to collect to ensure that the inferences drawn are reliable? In the context of relative belief, this is answered by the computation of what are called the biases, see, [Evans and Guo \(2021\)](#). Basically there are two biases, *bias against* and *bias in favor*, for the hypothesis assessment and the estimation problems. For example, for the hypothesis problem $H_0 : \Psi(\theta) = \psi_0$, then

$$\begin{aligned} \text{bias against}_{\Psi}(\psi_0) &= M(RB_{\Psi}(\psi_0|x) \leq 1 | \psi_0), \\ \text{bias in favor}_{\Psi}(\psi_0) &= \sup_{\psi: d(\psi, \psi_0) \geq \delta} M(RB_{\Psi}(\psi_0|x) \geq 1 | \psi). \end{aligned}$$

So, the bias against ψ_0 is the prior probability of not getting evidence in favor of ψ_0

when it is true, and the bias in favor of ψ_0 , is the largest prior probability of getting evidence in favor of ψ_0 when it is meaningfully false. In the context of the Bartlett paradox it is the case that the bias in favor converges to 1 as the prior becomes more and more diffuse. The cure for this is simple, however, as both biases converge to 0 with sample size. The reliability of the inferences demands that these probabilities be small as they refer to errors in the inferences. Would you believe the results of a statistical analysis that found evidence against ψ_0 when bias against $_{\Psi}(\psi_0) = 0.9$? Our guess is that most would say that this is pretty much a foregone conclusion and advocate that the experimenter collect more data. Similarly, if bias in favor $_{\Psi}(\psi_0)$ was very large, it is unlikely that finding evidence in favor of ψ_0 would be viewed as a reliable result. As such, it seems you are free to choose a very diffuse prior, but don't do this arbitrarily as you will need to pay the price of having to collect more data to ensure reliability. There are similar bias calculations for the estimation problem, but now these refer to the coverage probabilities of the plausible region $Pl_{\Psi}(x)$.

It is notable that these bias calculations are frequentist probabilities with respect to the model $\{m(\cdot|\psi) : \psi \in \Psi(\Theta)\}$. This makes for a nice unification of Bayesian and frequentist ideas: Bayes is for inference and frequentism is for design and error control to ensure reliability. For relative belief inference, both ways of thinking are essential and there is no conflict between them.

The relative belief ratio is not the only valid measure of evidence via the principle of evidence. For example, the Bayes factor is also a valid measure of evidence with cut-off 1. A natural question then is, why not base the theory of inference on the Bayes factor? Going back to the discussion leading up to the statement of the principle of evidence, we see that the Bayes factor in favor of A

having observed C is given by

$$BF(A|C) = \frac{P(A|C)}{P(A^c|C)} \bigg/ \frac{P(A)}{P(A^c)},$$

which is the ratio of the posterior odds in favor of A to the prior odds in favor of A . In fact, this shows that

$$BF(A|C) = \frac{RB(A|C)}{RB(A^c|C)},$$

and the Bayes factor can be expressed very simply in terms of the relative belief ratio. Since $RB(A|C) > 1$ iff $RB(A^c|C) < 1$, taking the ratio is not really accomplishing anything meaningful in terms of expressing evidence. It is worth noting too that Bayes factors are primarily used for hypothesis assessment, but the relative belief ratio also leads to a theory of estimation. While the Bayes factor may have uses in the context of discrete parameters, in the continuous context, a natural definition is given by the limit

$$\begin{aligned} BF_{\Psi}(\psi_0|x) &= \lim_{\epsilon \rightarrow 0} BF_{\Psi}(B_{\epsilon}(\psi_0)|x) \\ &= RB_{\Psi}(\psi_0|x), \end{aligned}$$

which establishes the equality of the two measures. Even when the Bayes factor is defined via a spike-and-slab prior, a very natural restriction for the prior on the null leads to equality of the two measures. These results, and other reasons to prefer the relative belief ratio to the Bayes factor, are documented in [Al-Labadai et al. \(2025\)](#). In any case, many of the results of relative belief theory do not depend on the particular valid measure of evidence used. For example, the definition of the plausible region only depends on the principle of evidence.

The relative belief ratio has arisen as a measure of evidence many times in the literature and that is because it is very natural and has many good, and even optimal, properties, see [Evans \(2015\)](#). For example, [Keynes \(1920\)](#) refers to it as the coefficient of influence and there have been other names

given to the ratio as well. In statistics it is sometimes referred to as the Savage-Dickey ratio. The terminology relative belief ratio seems more appropriate, however, as that is what it is, the ratio of beliefs a posteriori to beliefs a priori. Note too that the relative belief ratio is the density of the posterior with respect to the prior. What is unique about relative belief theory, is that it is an attempt to construct an entire theory of inference based upon the characterization of statistical evidence through the principle of evidence. A proper Bayesian inference base is necessary for the principle of evidence to apply and this is perhaps the primary reason the author considers himself a Bayesian. So far, it seems that attempts to define or characterize evidence without specifying a prior do not succeed, but success in that endeavor can never be ruled out. For a review of the various attempts to characterize statistical evidence see [Evans \(2025\)](#).

References

- Al-Labadi, L., Alzaatreh, A. and Evans, M. (2025), How to measure statistical evidence and its strength: Bayes factors or relative belief ratios?. *Canadian Journal of Statistics*, 53: e70015. <https://doi.org/10.1002/cjs.70015>
- Evans, M. (2015), *Measuring Statistical Evidence Using Relative Belief*. Monographs on Statistics and Applied Probability 144, CRC Press, Taylor & Francis Group.
- Evans, M. (2025), The Concept of Statistical Evidence, Historical Roots and Current Developments. *Encyclopedia* 2024, 4, 1201-1216. <https://doi.org/10.3390/encyclopedia4030078>
- Evans, M.; Guo, Y. (2021), Measuring and Controlling Bias for Some Bayesian Inferences and the Relation to Frequentist Criteria. *Entropy* 2021, 23, 190. <https://doi.org/10.3390/e23020190>

Evans, M.; Jang, G-H. (2024), Relative Belief Inferences from Decision Theory. *Entropy* 2024, 26, 786. <https://doi.org/10.3390/e26090786>

Fraser, D. A. S. (1979), *Inference and Linear Models*. McGraw-Hill.

Good, I. J. (1971), 46656 Varieties of Bayesians. *The American Statistician*, 25(1), 62–63.

Keynes, J. K. (1920) *A Treatise On Probability*. Wildside Press LLC.

Bayesian Theses from Banaras Hindu University in 2025-26

Analysis of Human Fertility Patterns Using Bayesian Paradigm

Shambhavi Singh

Ms. Shambhavi Singh submitted her Ph.D. thesis entitled “Analysis of Human Fertility Patterns Using Bayesian Paradigm” under the supervision of Prof. S. K. Upadhyay, Department of Statistics, Institute of Science, BHU. The thesis presents a rigorous Bayesian investigation of human fertility patterns by integrating demographic theory, probabilistic modelling, and modern computational techniques.

Fertility is a central determinant of population structure and demographic change. Although long studied in demographic research, the Bayesian framework remains relatively underutilized in this area. This thesis addresses this gap by showing how prior information, expert knowledge, and historical evidence can be systematically incorporated into fertility modelling through Bayesian inference, which is particularly relevant in demographic studies where prior insights on fertility behaviour, biological processes,

and population trends are often available.

The thesis develops flexible Bayesian models for age-specific fertility rates and extends Bayesian modelling to fecundity. Since conventional models may be inadequate for contemporary and non-standard fertility patterns, mixture-based models using lifetime distributions are proposed to capture complex structures, including bimodal fertility schedules. Fertility dynamics are further addressed by distinguishing pre-modal and post-modal variability, allowing adaptation to spatio-temporal shifts in the fertility behaviour.

It also introduces copula-based joint prior structures for demographic parameters, relaxing the usual assumption of prior independence and coherently capturing dependence among parameters in a computationally practical manner without complex conditional specifications. For fecundity analysis, a Bayesian generalized non-linear model with an asymmetric link function is used to study conception probabilities while accounting for covariates such as a woman's age, mucus characteristics, and intercourse patterns, thereby capturing the inherent asymmetry in conception processes and improving model realism.

Given the complexity of posterior distributions, Markov chain Monte Carlo methods, including Metropolis and hybrid Gibbs–Metropolis algorithms, are employed for computation. Through simulated and real-world data analyses, the proposed Bayesian models are evaluated in terms of fit, predictive performance, and comparative advantage over standard approaches. Overall, the thesis contributes to applied Bayesian demography by offering flexible, interpretable, and computationally efficient models for understanding and predicting human fertility patterns.

Bayes Analysis of Certain Medico-Social Datasets Using Selected Regression Models

Pranjal Kumar Pandey

Mr. Pranjal K. Pandey also submitted his Ph.D. thesis on the above title under the supervision of Prof. S. K. Upadhyay, Department of Statistics, Institute of Science, BHU. The thesis presents a Bayesian investigation of logistic regression models and their extensions, with emphasis on medico-social data analysis. The primary objective is to develop reliable, flexible, and interpretable Bayesian inferential frameworks for binary and multi-category response variables. In this direction, Bayesian logistic regression is considered as the foundational model and extended to polytomous logistic regression for analysing outcomes with more than two categories.

A significant feature of the thesis is the incorporation of interaction effects among co-variables within Bayesian regression models. By including interaction terms, the study captures complex relationships among explanatory variables often present in medico-social research, thereby improving both explanatory power and practical interpretability, particularly when multiple risk factors jointly influence health-related outcomes.

The Bayesian perspective is especially important for low-event and small-sample datasets frequently encountered in medical, epidemiological, and social science research, where classical likelihood-based methods may produce unstable estimates, large standard errors, or convergence difficulties. To address these limitations, the thesis employs subjectively elicited prior distributions based on expert opinion, providing additional information to stabilize estimation and improve posterior inference when observed data are limited. The study also examines the influence of prior specification on posterior

estimates, emphasizing the importance of careful prior elicitation in applied Bayesian analysis.

Another important contribution lies in the computational implementation of the proposed models. Existing Markov chain Monte Carlo techniques are suitably modified to improve applicability, efficiency, and convergence behaviour for logistic and polytomous logistic regression models, particularly in sparse-data settings. These refinements strengthen the practical usability of the Bayesian procedures developed in the thesis.

The proposed methodologies are applied to diverse medico-social datasets obtained from the Institute of Medical Sciences, Banaras Hindu University, demonstrating the relevance of Bayesian modelling in healthcare-oriented data analysis. Overall, the thesis contributes to applied Bayesian statistics by integrating prior elicitation, advanced regression modelling, computational refinement, and medico-social applications within a coherent inferential framework.

Statistical Modelling and Optimal Test Plans for Highly Reliable Products

Ankush Sharma

Mr. Ankush Sharma has submitted his Ph.D. thesis on the above title under the supervision of Prof. Sanjeev Kumar, Department of Statistics, Institute of Science, BHU. The thesis presents a rigorous study of degradation-based reliability modelling for highly reliable products operating under complex stress environments, with emphasis on Bayesian inferences, uncertainty quantification, sequential learning, and real-time reliability prediction.

The research develops stochastic process-based degradation models incorporating multiple stress factors and their interaction ef-

fects, providing flexible probabilistic frameworks for analysing degradation behaviour when failure-time data are limited. The thesis further proposes optimal accelerated degradation test plans based on D-optimality and V-optimality criteria for improving parameter estimation and reliability prediction while reducing testing cost and duration.

A major contribution is the development of an online variational Bayesian framework for sequential updating of degradation information and prediction of Remaining Useful Life, enabling computationally efficient real-time monitoring and predictive maintenance in high-dimensional settings. The work also incorporates latent-variable modelling to address delayed degradation initiation and heterogeneous degradation progression, and develops a stochastic EM-based estimation

framework under multiple stress factors.

Another significant contribution is the theoretical investigation of hierarchical latent variable models using double asymptotic theory and Laplace approximation results for posterior concentration, leading to information scaling results and robust budget allocation strategies.

Overall, the thesis integrates stochastic degradation modelling, Bayesian inference, variational approximation, latent-variable modelling, asymptotic theory, and optimal design into a unified framework for reliability analysis, with applications in aerospace, energy systems, electronics, and advanced manufacturing.

Membership Subscription Form

(For both Individual and Institutional Membership)

Name of the Individual or Institution:
Name of the Authorized Person (in case of Institutional membership):
Designation:
Department:
Address for correspondence:

Designation:

Address for correspondence:

Phone (Off.)

e-mail:

I want the information available for others

Type of Membership Sought (Please tick whichever is applicable):

Annual (Rs. 50/-)

Life (Rs. 500/-)

Institutional or Organizational (Rs. 500/-)

Place:

Signature

Please complete the above form and enclose a Cheque/Bank draft in the name of Indian Bayesian Society payable at Varanasi. Please add the appropriate Bank clearance charges for the out station cheque.

Return this form along with your membership dues to:

K 67/81 A-2

Bharat Milap Colony

Varanasi-221 001
