

Modeling Adaptive Success: A Discrete Hill-Type Hazard Approach in Education

Leon Monsoon

Independent Researcher, Dallas, TX

leon.monsoon.research@gmail.com

November 15, 2025

Abstract

Traditional models of student performance and item mastery often assume a constant probability of success across attempts—for example, the geometric distribution. Such models cannot capture the “aha” moment where performance transitions from repeated failure to rapid success. In this work, we adapt a discrete Hill-type hazard (equivalently, a log-logistic hazard function) to the problem of first success over repeated attempts in educational settings. We use this trial-indexed success probability, $\tau(n; h, K) = n^h / (n^h + K^h)$, as the discrete hazard of a non-homogeneous geometric process governing the trial of first success. The resulting model—which we refer to as the Learner’s Tau framework—is governed by a steepness parameter h and a midpoint parameter K and is designed to capture adaptive success in systems where the probability of success increases over time, such as human learning, reinforcement learning, and task mastery.

We derive the probability mass function, analyze key mathematical properties, and introduce a set of interpretable summary quantities—including the discrete hazard function, a Difficulty Ratio K/h , an Early Success Probability (the CDF evaluated at the midpoint), the Mean Time to Mastery, and the Peak Effort Trial where effort-weighted success is maximized. Using numerical simulations and maximum likelihood estimation, we demonstrate that the model can recover underlying parameters and distinguish between qualitatively different learner profiles. Empirical validation using educational data from the Cognitive Tutor system demonstrates that the model captures learning dynamics significantly better than a memoryless geometric baseline for non-trivial tasks ($\Delta\text{AIC} = 8.0$ and 16.6 for two algebra skills), while appropriately deferring to the simpler model when ceiling effects preclude observable learning. Rather than proposing a brand-new functional form, this work adapts a classical Hill-type hazard to an explicitly discrete, learning-centric context and highlights its utility for modeling adaptive success trajectories in education.

1 Introduction

Success is rarely immediate. In real-world learning, mastery is often preceded by repeated failure, gradual refinement, and eventual breakthroughs. Yet most discrete probabilistic models treat success as a memoryless or constant-probability event, ignoring the dynamic, nonlinear progression that defines most human and system learning [6, 10, 11].

This paper studies a discrete learning model built from a Hill-type hazard function, applied to the problem of first success over repeated attempts in education. We treat the trial-indexed success probability

$$\tau(n; h, K) = \frac{n^h}{n^h + K^h}$$

as a discrete analogue of the continuous Hill equation and log-logistic hazard, and we use it to induce a non-homogeneous geometric distribution for the trial of first success. The resulting model—which we refer to as the **Learner’s Tau framework**—is governed by a steepness parameter h and a midpoint parameter K and is designed to capture adaptive success in systems where the probability of success increases over time, such as human learning, reinforcement learning, and task mastery.

At its core, the Learner’s Tau distribution is defined by two parameters:

- h : the **steepness** of the learning curve (how sharply success probability increases),
- K : the **midpoint**, the trial at which the probability of success reaches 50%.

These parameters allow the model to represent a broad range of learner types: from slow, patient learners who take many attempts before improving, to fast-adapting learners who gain traction quickly.

1.1 Motivation and Real-World Need

Existing models like the geometric, negative binomial, and beta-geometric distributions serve important roles in failure modeling, lifetime analysis, and reliability studies [1, 2, 3, 4]. However, they often assume:

- static success probabilities,
- constant hazard rates,
- or specific decay structures unrelated to learning.

In contrast, human learning and adaptive algorithms frequently exhibit:

- success becoming more likely over time,
- progress that follows an S-shaped curve,
- early failure that is not only common but expected [6, 7, 12].

The Learner’s Tau distribution fills this gap by providing a mathematically tractable model that is:

- discrete and interpretable,
- centered around adaptive success,
- flexible enough to represent multiple learning dynamics.

1.2 Contribution of This Paper

This paper formally introduces and analyzes the Learner’s Tau distribution as a discrete Hill-type hazard model for item mastery. Our contributions include:

1. A clear definition of the model, including its PMF, success function (τ), and parameter roles, obtained by adapting a classical Hill/log-logistic hazard to a discrete first-success setting.
2. A set of *derived quantities*—the discrete hazard function, a Difficulty Ratio K/h , an Early Success Probability (CDF at the midpoint), the Mean Time to Mastery, and the Peak Effort Trial—that package standard probabilistic objects in learning-centric terms.

3. Empirical validation on real educational data, demonstrating improved model fit compared to the geometric baseline for skills exhibiting progressive learning ($\Delta\text{AIC} \geq 8$), and appropriate model selection behavior for ceiling-effect tasks.
4. Estimation methods based on maximum likelihood and joint confidence regions, together with bootstrap-based uncertainty quantification.
5. A discussion of potential applications in education, adaptive testing, and reinforcement learning where interpretable learning curves are desired.

Our goal is not to claim a novel functional form, but to demonstrate that a discrete Hill-type hazard, when paired with standard estimation tools and applied to educational data, provides an interpretable and practically useful model of adaptive success.

2 Related Work and Motivation

The landscape of discrete probability distributions is rich with models designed to describe time-to-event outcomes, yet relatively few are equipped to handle learning-driven success. Most classical distributions operate under assumptions that, while useful in modeling failures, often fail to capture the adaptive dynamics inherent in success. In this section, we review existing models commonly used for first success or event timing and identify the specific limitations that motivated the development of the Learner’s Tau distribution.

2.1 The Geometric Distribution and the Memoryless Assumption

The geometric distribution is a standard model for the number of trials until first success. Defined by a constant success probability p , its probability mass function (PMF) is [1, 2]:

$$P(N = n) = (1 - p)^{n-1}p, \quad n = 1, 2, 3, \dots \quad (1)$$

This model assumes that the probability of success is constant on each trial, an assumption known as the **memoryless property**. That is, for any $n, k > 0$,

$$P(N > n + k \mid N > n) = P(N > k). \quad (2)$$

In many contexts such as coin tosses, radioactive decay, or Bernoulli trials, this assumption is appropriate. However, in learning scenarios, the assumption is often unrealistic. Learners do not attempt tasks with the same likelihood of success each time. With each failure comes information, and with each repetition, the learner refines their internal model of the task. In effect, the probability of success increases over time—something the geometric model cannot represent.

2.2 Negative Binomial and Beta-Geometric Models

The negative binomial distribution generalizes the geometric model to represent the number of failures before a fixed number of successes. While this adds flexibility, it still relies on the assumption of a **fixed** success probability per trial and primarily models cumulative success, not the timing of first success when learning is occurring [3].

The **beta-geometric distribution** introduces variability in the success probability by modeling p as a beta-distributed random variable. This allows for overdispersion and accounts for **heterogeneity across individuals**, but it still does not model learning as a process where the

probability of success increases within an individual over time. Rather, it represents variability across a population with fixed but unknown p , not a dynamically evolving success curve for a single learner. Moreover, the relationship between the beta parameters and learning progress is indirect, and modeling interventions or adaptive training is not straightforward [4].

2.3 Weibull and General Failure Models

In reliability analysis, the Weibull distribution and its discrete analogs are commonly used to model failure times with increasing or decreasing hazard rates [5, 4]. The discrete Weibull can model non-constant probabilities, and while it is more flexible than the geometric or negative binomial distributions, it is still largely focused on **decay and time-to-failure**. Its increasing hazard rate can resemble learning, but it lacks a clear learning-centric interpretation: there is no direct mapping between Weibull parameters and the learner’s rate of improvement or learning curve dynamics.

In addition, most existing discrete models related to failure or success timing:

- do not have intuitive parameters representing “learning steepness” or “midpoint of success”,
- do not naturally model S-shaped growth curves,
- and offer little to no interpretive framework for characterizing different learners.

2.4 The Need for a Learning-Centric Model

In many learning processes, failure precedes success but not in a memoryless way. The more one tries, the better one tends to become. The chance of success on attempt 15 is not the same as on attempt 1. There is a qualitative and quantitative shift in confidence, skill, and strategy. Traditional models that assume constant or decaying probabilities fail to capture this behavior.

We seek a model that can:

- represent an **adaptive** probability of success over time,
- be parameterized by interpretable constructs such as **steepness** and **midpoint**,
- offer both theoretical tractability and empirical flexibility,
- and support diagnostic tools, simulation engines, and cognitive models [6, 7, 14].

The Learner’s Tau distribution is designed to address these requirements by explicitly encoding a non-constant, trial-indexed success probability derived from a classical Hill-type hazard.

3 Definition of the Learner’s Tau Distribution

The Learner’s Tau distribution is a discrete probability model for the probability of **first success** on trial n , under the assumption that the probability of success **increases** with each attempt. It is governed by two parameters: h , capturing the **steepness** of the learner’s improvement, and K , representing the **midpoint** of the learning process. Unlike traditional memoryless or fixed-probability models, this distribution is explicitly constructed to mirror the dynamics of adaptive, iterative learning over time.

In this section, we formally define the probability mass function (PMF), cumulative distribution function (CDF), and success probability function (the τ -function) for the Learner’s Tau model. We also summarize key properties and behavior over the domain of positive integers.

3.1 The Success Probability Function (τ Function)

Let $N \in \mathbb{N}^+$ be the random variable denoting the trial on which the first success occurs. For each trial n , we define a function $\tau(n; h, K)$, referred to as the **Tau function**, which represents the **probability of success on the n th attempt**, conditional on all previous attempts having failed:

$$\tau(n; h, K) = \frac{n^h}{n^h + K^h}. \quad (3)$$

Here:

- $h > 0$ is the **steepness parameter**. It controls how rapidly the success probability grows with each additional trial. A higher value of h results in a **sharper increase**, mimicking a “breakthrough” or “aha moment” dynamic.
- $K > 0$ is the **midpoint parameter**. It marks the trial n at which the learner reaches a **50% chance of success**, i.e.,

$$\tau(K; h, K) = \frac{K^h}{K^h + K^h} = 0.5.$$

This parameter shifts the curve left or right and can be thought of as a **skill acquisition delay** or point of inflection in the learning process.

The function $\tau(n; h, K)$ is monotonically increasing in n , bounded between 0 and 1:

$$\tau(1; h, K) = \frac{1}{1 + K^h}, \quad \text{and} \quad \lim_{n \rightarrow \infty} \tau(n; h, K) = 1. \quad (4)$$

For $h > 1$, it defines an S-shaped curve reminiscent of learning curves observed in psychology, education, and cognitive neuroscience [6, 12].

Connection to existing functions. The functional form in Equation (3) is not new in itself. In continuous settings it appears as the classical Hill equation [15], widely used in biochemistry to model saturation and cooperativity, and as the hazard function of the log-logistic family in survival analysis and reliability modeling [4]. Our contribution is to work with a discrete, trial-indexed version of this Hill-type hazard, embed it in a non-homogeneous geometric construction for first success, and interpret the parameters (h, K) and derived quantities in explicitly learning-centric terms.

3.2 Probability Mass Function (PMF)

Using the Tau function, we define the PMF of the Learner’s Tau distribution as the probability of failing on trials $1, \dots, n-1$ and succeeding on trial n :

$$P(N = n; h, K) = \left[\prod_{i=1}^{n-1} (1 - \tau(i; h, K)) \right] \cdot \tau(n; h, K), \quad n = 1, 2, 3, \dots \quad (5)$$

where the product term is defined as 1 if $n = 1$. This expression models a process in which:

- the learner fails on trials 1 through $n-1$,
- then succeeds on trial n , with the success probability at trial n determined by $\tau(n; h, K)$.

The PMF is inherently **non-stationary**; the success probability is not constant but evolves across trials. This makes the distribution suited to modeling learning, rather than static failure or decay.

3.3 Cumulative Distribution Function (CDF)

The cumulative distribution function (CDF) of the Learner's Tau distribution is obtained by summing the PMF:

$$F(n; h, K) = P(N \leq n) = \sum_{k=1}^n P(N = k; h, K). \quad (6)$$

Equivalently, $F(n; h, K)$ is one minus the probability of failing on all trials from 1 to n :

$$F(n; h, K) = 1 - \prod_{i=1}^n (1 - \tau(i; h, K)). \quad (7)$$

This representation is a closed-form expression in terms of a finite product. For large n , both forms can be efficiently evaluated numerically.

3.4 Properness of the Distribution

For any $h > 0$ and $K > 0$, the PMF in Equation (5) defines a proper distribution on $\mathbb{N}^+$, i.e.,

$$\sum_{n=1}^{\infty} P(N = n; h, K) = 1.$$

[Sketch] By construction,

$$P(N > n) = \prod_{i=1}^n (1 - \tau(i; h, K)),$$

so $\sum_{n=1}^{\infty} P(N = n) = 1 - \lim_{n \rightarrow \infty} P(N > n)$. Thus it is enough to show that $P(N > n) \rightarrow 0$ as $n \rightarrow \infty$. For large i we have $\tau(i; h, K) \rightarrow 1$ and, more concretely,

$$\tau(i; h, K) = \frac{1}{1 + (K/i)^h} \geq 1 - c i^{-h}$$

for some constant $c > 0$ and all sufficiently large i . Hence $1 - \tau(i; h, K) \leq c i^{-h}$ and

$$\log P(N > n) = \sum_{i=1}^n \log(1 - \tau(i; h, K)) \leq \sum_{i=i_0}^n \log(c i^{-h}),$$

which diverges to $-\infty$ as $n \rightarrow \infty$. Therefore $P(N > n) \rightarrow 0$ and the PMF sums to 1.

3.5 Key Properties

Several important characteristics of the distribution are summarized below:

1. **Support:** The distribution is defined on the positive integers: $n \in \mathbb{N}^+ = \{1, 2, 3, \dots\}$.
2. **Monotonicity of $\tau(n)$:** For fixed $h, K > 0$, the function $\tau(n; h, K)$ is monotonically increasing in n , approaches 1 asymptotically, and never exceeds it.
3. **Unimodality of PMF:** Over the parameter ranges considered in this paper, the PMF is empirically unimodal: it has a single peak corresponding to the most likely trial of success (denoted n_{peak}). The location and sharpness of the peak depend on h and K .

4. **Mean and variance:** The distribution generally lacks simple closed-form expressions for the mean ($\mathbb{E}[N]$) and variance ($\text{Var}(N)$). These moments can be computed using numerical methods such as truncated summation or Monte Carlo simulation. For example,

$$\mathbb{E}[N] \approx \sum_{n=1}^{N_{\max}} n \cdot P(N = n; h, K), \quad (8)$$

$$\text{Var}(N) \approx \sum_{n=1}^{N_{\max}} (n - \mathbb{E}[N])^2 \cdot P(N = n; h, K), \quad (9)$$

where $N_{\max}$ is a sufficiently large integer such that $P(N = n)$ is negligible for $n > N_{\max}$.

3.6 Comparison to Other Distributions

Unlike the geometric distribution (which assumes a constant p), or the negative binomial (which models cumulative successes with constant p), the Learner's Tau distribution is designed specifically for **learning and adaptive success**. It belongs to a class of models where the **success probability is an explicit function of the trial number**, rather than a fixed scalar or a random variable representing population heterogeneity.

3.7 Visualization of the Distribution

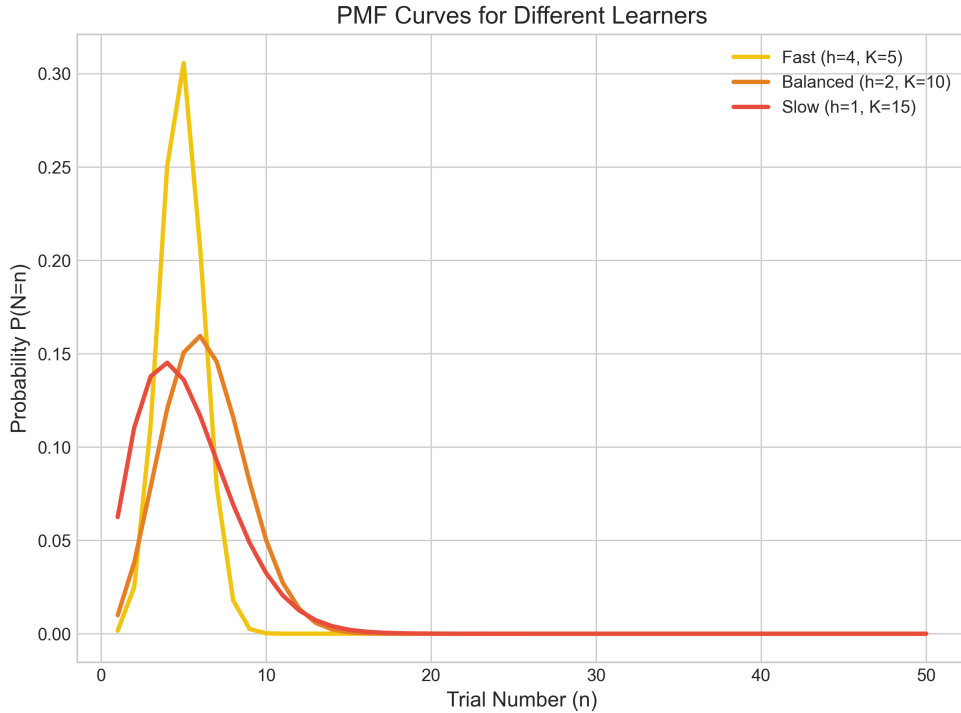


Figure 1: Probability mass functions (PMFs) for three Learner's Tau profiles: Fast ($h = 4, K = 5$), Balanced ($h = 2, K = 10$), and Slow ($h = 1, K = 15$). The shapes illustrate different learning dynamics.

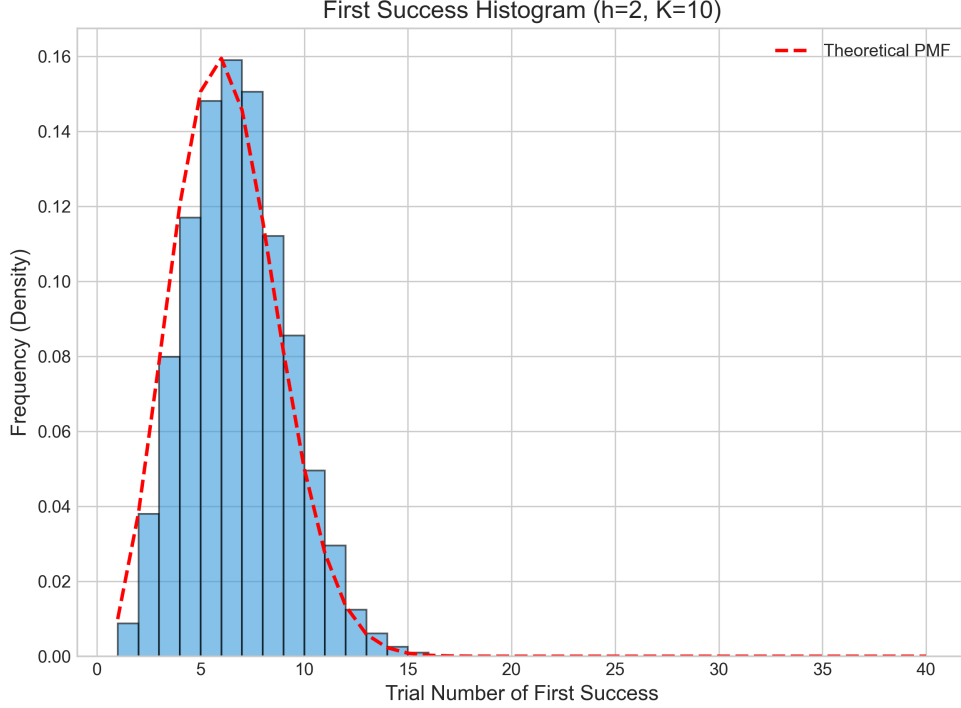


Figure 2: Histogram of first success trials from a simulation of 10,000 learners with parameters $h = 2, K = 10$. Most learners succeed near the midpoint ($K = 10$), with a tail of persistent learners taking longer.

4 Intuition Behind Parameters

The Learner’s Tau distribution was intentionally designed with parameters that are not only mathematically meaningful, but also interpretable in behavioral and educational settings. In this section, we explore the roles of h and K , their effects on the success probability curve $\tau(n)$, and how varying these parameters captures different learner archetypes and real-world behaviors.

4.1 The Steepness Parameter: h

The parameter $h > 0$ controls the **steepness** or **acceleration** of learning. Mathematically, it affects the curvature of the success probability function $\tau(n)$, transforming a gradual incline into a steep ramp as h increases.

Behavior:

- When $h \ll 1$: The increase in success probability is very gradual, creating a shallow learning curve. Success remains unlikely for many trials before marginal gains accumulate.
- When $h \approx 1$: The learning curve behaves somewhat like a logistic function, with slow early growth that accelerates near the midpoint K .
- When $h \gg 1$: The learning curve becomes more step-like: success probability remains low until a sharp rise around the midpoint K . This mimics **breakthrough learning** or an “aha” moment.

Real-world examples:

- A low h might describe a learner engaging with a complex skill (e.g., mastering jazz piano) that rewards consistent, incremental improvement.
- A high h might describe a learner who experiences a sharp change in performance after a single conceptual insight (e.g., finally understanding recursion in programming).

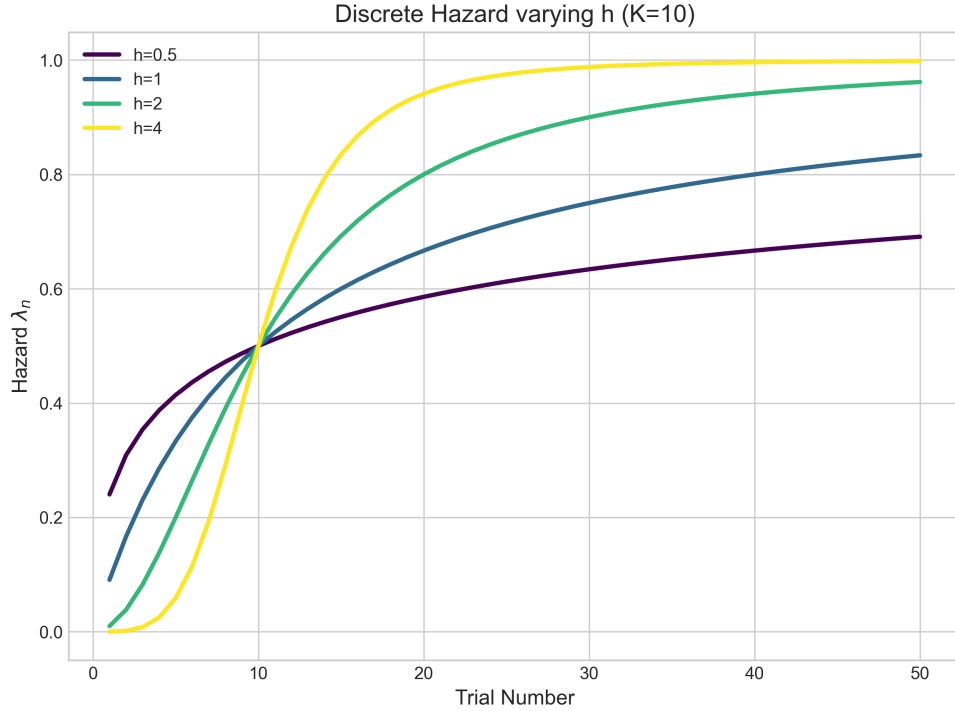


Figure 3: Tau function curves ($\tau(n)$) for varying steepness parameter h (with fixed midpoint $K = 10$). Higher h leads to a sharper increase in success probability around the midpoint.

4.2 The Midpoint Parameter: K

The parameter $K > 0$ governs the **horizontal shift** of the τ -function. Specifically, it determines the trial $n = K$ at which the learner's probability of success crosses 50%. It acts as a temporal anchor indicating when meaningful progress is expected.

Behavior:

- A low K (e.g., 2–5): Learners are expected to succeed relatively early in the process.
- A medium K (e.g., 10–15): Learners show delayed but measurable improvement.
- A high K (e.g., 20+): Learners are expected to struggle for a longer period before improvement emerges.

Mathematically, increasing K delays the rise of $\tau(n)$, pushing the probability curve to the right.

Real-world examples:

- A student preparing for a difficult standardized exam may show little success in early attempts but gradually internalize strategies, with performance improving closer to $K = 15$.
- In contrast, someone playing a new video game may grasp the mechanics quickly (low K), succeeding after only a few trials.

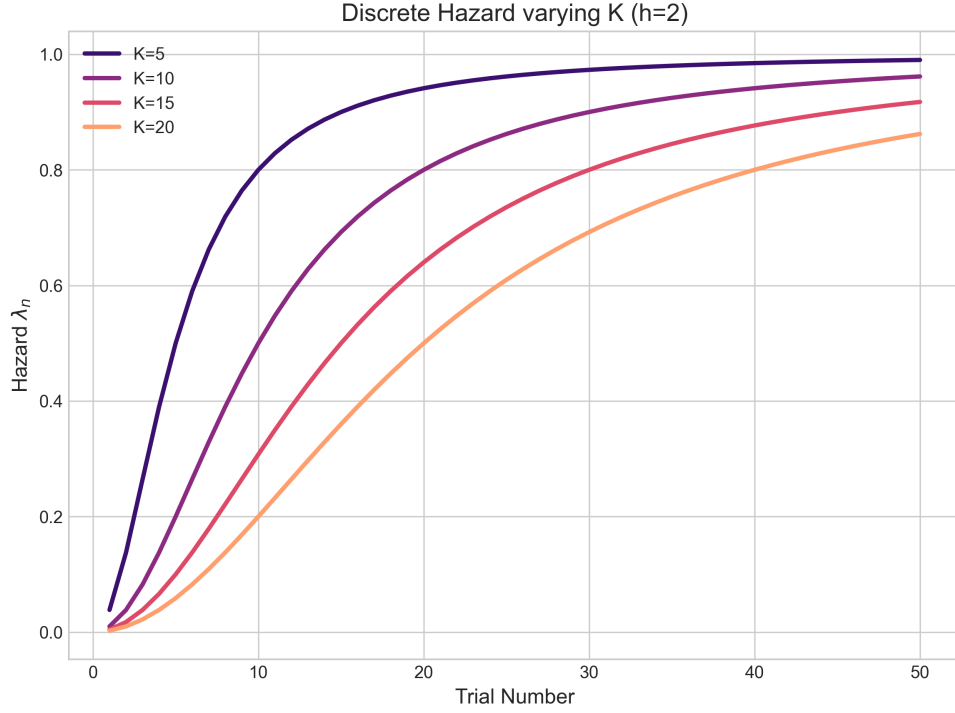


Figure 4: Tau function curves ($\tau(n)$) for varying midpoint parameter K (with fixed steepness $h = 2$). Higher K shifts the learning curve to the right, delaying the onset of significant success probability.

4.3 Combined Interpretation of h and K

The interaction between h and K defines the overall shape, timing, and nature of the learning curve. By adjusting them jointly, we can simulate a spectrum of learner behaviors:

Table 1: Learner archetypes based on combinations of steepness (h) and midpoint (K) parameters.

Learner Type	h	K	Description
Fast Learner	High	Low	Rapid improvement after few trials
Slow, Patient Learner	Low	High	Gradual, delayed mastery
Breakthrough Learner	High	Medium	Sharp success after extended struggle
Consistent Improver	Medium	Medium	Steady gains across time

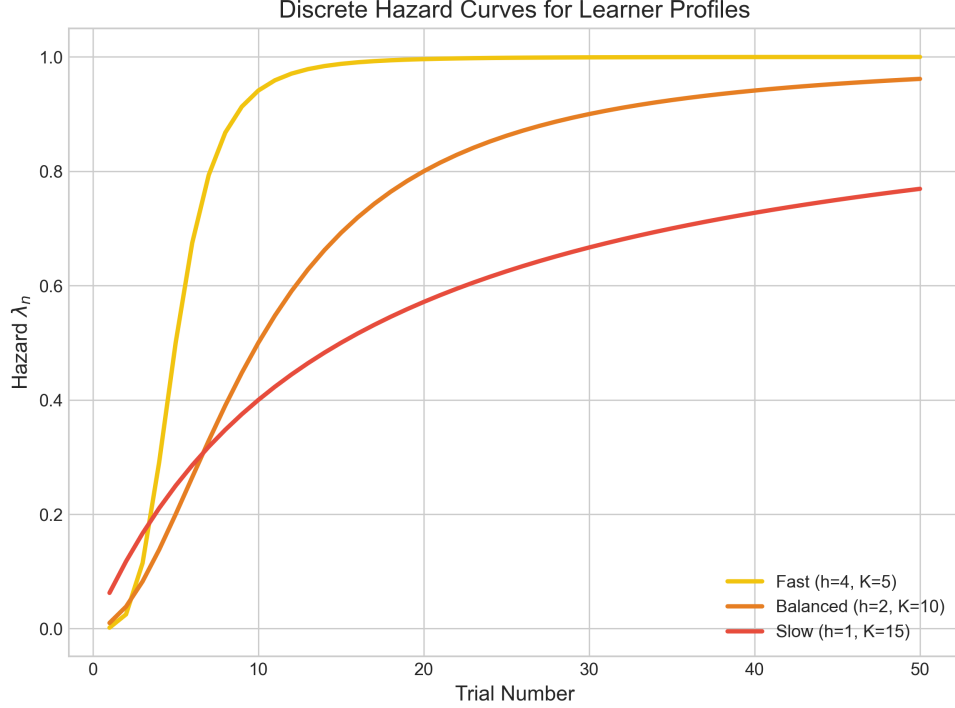


Figure 5: Tau function curves (equivalently, discrete hazard curves; see Section 5.1) for three learner profiles: Fast ($h = 4, K = 5$), Balanced ($h = 2, K = 10$), and Slow ($h = 1, K = 15$). This illustrates the diverse learning trajectories the model can capture.

4.4 $\tau(n)$ as a Qualitative Map

Although the Learner’s Tau distribution is defined mathematically, the shape of $\tau(n)$ can also aid qualitative interpretation. For example, a steep change near K corresponds to a sharp transition from mostly failing to mostly succeeding; a very gradual curve reflects slow accumulation of skill. In educational or cognitive diagnostic settings, such shapes can help distinguish qualitatively different learning trajectories, even when learners reach similar end performance.

5 Derived Quantities and Learning-Centric Summaries

The Learner’s Tau distribution is a model for adaptive success over time. To facilitate interpretation and comparison across learners or tasks, we introduce a suite of **five derived quantities** built from standard probabilistic objects. These provide learning-centric summaries of the hazard, timing, difficulty, and effort associated with the first-success process.

Overview of the five quantities:

Table 2: Overview of five derived quantities for the Learner’s Tau distribution. Each is a standard probabilistic object, repackaged for learning-curve interpretation.

Name	Definition (standard object)	Learning-curve interpretation
Discrete hazard λ_n	$\tau(n)$ (discrete hazard)	Trial-wise success probability given no prior success
Difficulty Ratio K/h	Simple ratio of parameters	Heuristic index of delayed vs. sharp learning
Early Success Probability	$F(\lfloor K \rfloor)$ (CDF at midpoint)	Fraction of first successes occurring by trial K
Mean Time to Mastery	$\mathbb{E}[N]$ (mean time to success)	Expected number of trials required to succeed
Peak Effort Trial	$\arg \max_n \{nP(N = n)\}$ and its value $nP(N = n)$	Trial where effort-weighted success is maximized

Below we define each quantity and discuss its interpretation.

5.1 Discrete Hazard Function

Formula:

$$\lambda_n = P(N = n \mid N \geq n) = \tau(n; h, K) = \frac{n^h}{n^h + K^h}. \quad (10)$$

Interpretation: λ_n is the discrete hazard function of the process: the probability of success on trial n given that no success has occurred earlier. In our parametrization it coincides with the τ -function. We emphasize its role as a hazard in order to connect directly to standard survival and reliability terminology.

Use: The sequence $\{\lambda_n\}$ can be visualized as a learning curve and compared across learners or tasks to identify **early** versus **late** improvement.

5.2 Difficulty Ratio

Formula:

$$\text{Difficulty Ratio}(h, K) = \frac{K}{h}. \quad (11)$$

Interpretation: The ratio K/h is a deliberately simple, dimensionless summary of the two parameters. In this parameterization most of the “delay” in learning is driven by K , while h controls how sharp the transition is once learning accelerates. Larger values of K/h correspond to later and/or more gradual onset of high success probability.

Use: The Difficulty Ratio can be used descriptively to compare tasks or learners: higher ratios suggest more challenging or slowly mastered items.

5.3 Early Success Probability

Formula:

$$\text{Early Success Probability}(h, K) = F(\lfloor K \rfloor; h, K) = \sum_{n=1}^{\lfloor K \rfloor} P(N = n; h, K), \quad (12)$$

where $P(N = n; h, K)$ is the PMF from Equation (5) and F is the CDF.

Interpretation: This quantity is the **cumulative probability** that a learner succeeds **on or before** the theoretical midpoint K . It measures how much of the first-success mass lies in the “early” portion of the process, where “early” is defined relative to the midpoint:

- High Early Success Probability: a large fraction of learners succeed by trial K .
- Low Early Success Probability: a substantial fraction of learners must persist beyond trial K to succeed.

Use: This measure quantifies how frontloaded or backloaded success is relative to K . It can inform instructional design decisions about whether an item tends to be mastered quickly or requires extended practice before most learners succeed.

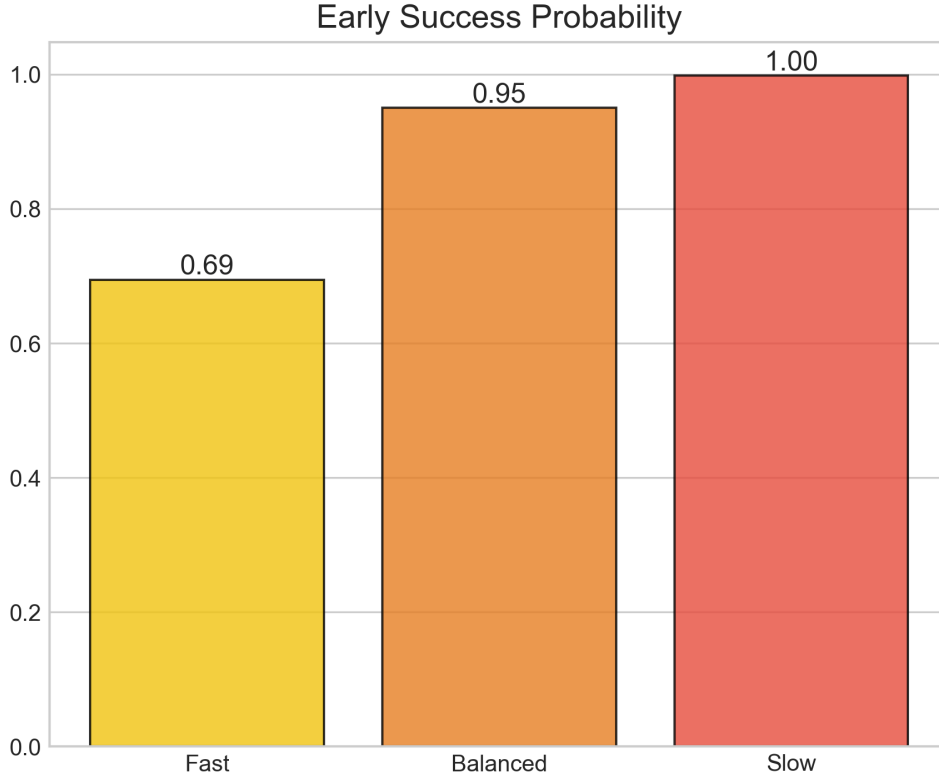


Figure 6: Early Success Probability (CDF at the midpoint K) for Fast ($h = 4, K = 5$), Balanced ($h = 2, K = 10$), and Slow ($h = 1, K = 15$) learners. Higher values indicate a greater fraction of first successes occurring on or before trial K .

5.4 Mean Time to Mastery

Formula:

$$\text{Mean Time to Mastery}(h, K) = \mathbb{E}[N] \approx \sum_{n=1}^{N_{\max}} n \cdot P(N = n; h, K). \quad (13)$$

Interpretation: The Mean Time to Mastery (MTM) is simply the mean of the first-success distribution, often called the mean time to success in reliability or survival settings. It captures the average number of attempts required for mastery.

- Low MTM: learners succeed quickly on average.
- High MTM: success typically requires many repeated trials.

Use: MTM maps directly to concepts like training length, resource allocation, or potential for fatigue. It measures the overall investment required for success.



Figure 7: Mean Time to Mastery (expected number of trials to first success) for Fast ($h = 4, K = 5$), Balanced ($h = 2, K = 10$), and Slow ($h = 1, K = 15$) learners. Higher values indicate a longer average learning journey.

5.5 Peak Effort Trial

Formula: Let $I(n) = n \cdot P(N = n; h, K)$ denote an effort-weighted success function. Define

$$n^*(h, K) = \arg \max_{n \geq 1} \{I(n)\}, \quad I^*(h, K) = I(n^*(h, K)) = \max_{n \geq 1} \{n \cdot P(N = n; h, K)\}. \quad (14)$$

Interpretation: The **Peak Effort Trial** $n^*(h, K)$ identifies the trial at which the effort-weighted success probability $I(n)$ is largest. This is not necessarily the mode n_{peak} of the PMF, but rather the trial where success is both sufficiently likely and represents a substantial accumulation of attempts. The value $I^*(h, K)$ summarizes the magnitude of this peak.

Use: The Peak Effort Trial provides a natural focus point on the learning trajectory: before n^* many learners are still accumulating experience, while after n^* the remaining learners tend to be increasingly rare. Interventions scheduled just before n^* may therefore have the highest leverage.

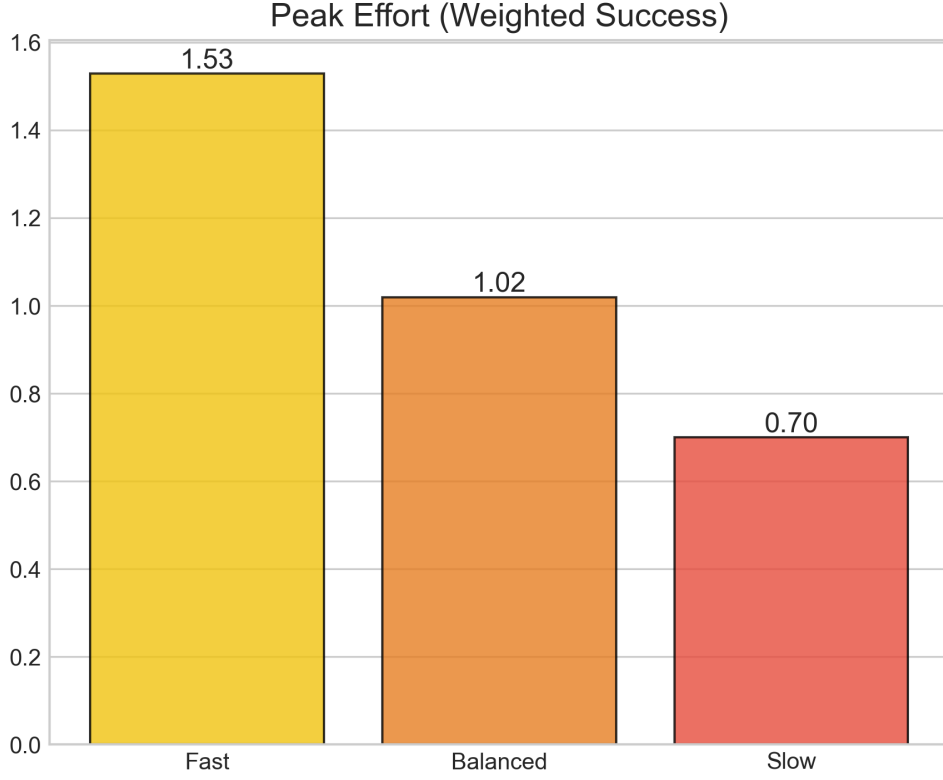


Figure 8: Effort-weighted peak $I^* = \max_n \{nP(N = n)\}$ for Fast ($h = 4, K = 5$), Balanced ($h = 2, K = 10$), and Slow ($h = 1, K = 15$) learners. The associated Peak Effort Trial n^* indicates where success is both likely and effortful.

5.6 Combined Summary Table

Table 3 summarizes these derived quantities for three representative learner profiles. MTM and the effort-weighted peak are approximate, based on numerical calculation.

Table 3: Summary of derived quantities for three learner profiles.

Learner	Hazard λ_{10}	Diff. K/h	Ratio	Early $F(\lfloor K \rfloor)$	Succ.	MTM $\approx \mathbb{E}[N]$	Effort I^*	Peak
Fast ($h = 4, K = 5$)	0.94	1.25		0.69		~ 4.9	~ 1.53	
Balanced ($h = 2, K = 10$)	0.50	5.00		0.95		~ 6.3	~ 1.02	
Slow ($h = 1, K = 15$)	0.40	15.00		1.00		~ 5.2	~ 0.70	

Figure 9 provides a visual comparison by plotting normalized versions of these summaries.

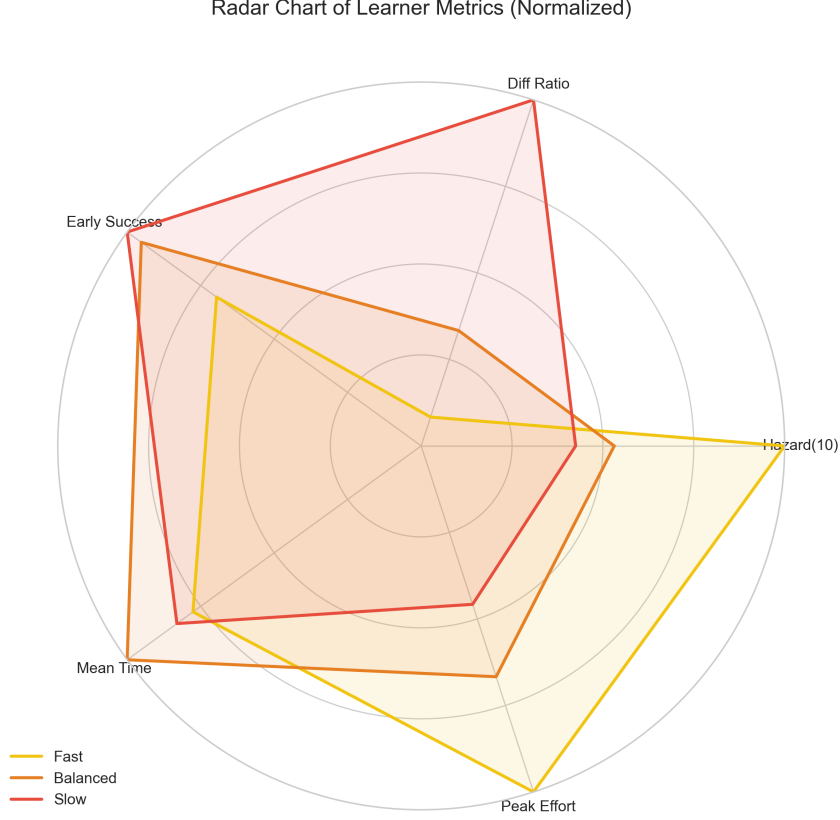


Figure 9: Radar chart comparing key derived quantities (normalized for visualization) across the Fast, Balanced, and Slow learner profiles.

6 Estimation and Inference

For a distribution to have scientific and practical value, it must be **estimable from data**. The Learner’s Tau distribution is designed to be fitted to observed first-success trials via standard likelihood-based and resampling methods [8, 9].

In this section, we develop methods to estimate the parameters h and K from data, including:

- maximum likelihood estimation (MLE),
- nonparametric bootstrapping for uncertainty quantification,
- and confidence ellipses to visualize joint estimation precision.

6.1 Maximum Likelihood Estimation (MLE)

Let $D = \{n_1, n_2, \dots, n_m\}$ be a set of observed first-success trial numbers from m independent learners, assumed to be drawn from a Learner’s Tau distribution with parameters (h, K) . The likelihood function is the product of the PMFs for each observation:

$$\mathcal{L}(h, K \mid D) = \prod_{j=1}^m P(N = n_j; h, K). \quad (15)$$

Substituting the PMF from Equation (5), we obtain

$$\mathcal{L}(h, K \mid D) = \prod_{j=1}^m \left[\left(\prod_{i=1}^{n_j-1} (1 - \tau(i; h, K)) \right) \tau(n_j; h, K) \right]. \quad (16)$$

For numerical stability and optimization, we work with the log-likelihood:

$$\log \mathcal{L}(h, K \mid D) = \sum_{j=1}^m \log P(N = n_j; h, K), \quad (17)$$

which expands to

$$\log \mathcal{L}(h, K \mid D) = \sum_{j=1}^m \left[\sum_{i=1}^{n_j-1} \log (1 - \tau(i; h, K)) + \log (\tau(n_j; h, K)) \right], \quad (18)$$

where $\tau(i; h, K) = i^h / (i^h + K^h)$.

Because the parameter space is low dimensional, the log-likelihood surface can be examined directly. In our pedagogical examples we use coarse grids to visualize this surface, but for estimation in practice we recommend standard numerical optimization routines. In applied work, a natural choice is to minimize the negative log-likelihood using quasi-Newton methods such as BFGS or constrained optimizers available in `scipy.optimize`. This yields maximum likelihood estimates $(\hat{h}, \hat{K})$ along with convergence diagnostics.

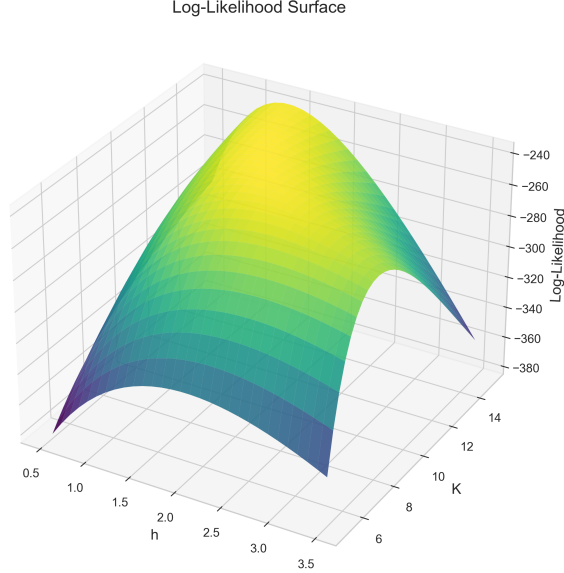


Figure 10: Example log-likelihood surface for Learner’s Tau parameters (h, K) based on simulated data. The peak of the surface corresponds to the maximum likelihood estimates $(\hat{h}, \hat{K})$.

6.2 Bootstrap Estimation and Uncertainty Quantification

MLE provides point estimates $(\hat{h}, \hat{K})$, but it does not directly quantify estimation uncertainty. To address this, we apply **nonparametric bootstrapping** [9]:

1. Given observed data $D = \{n_1, \dots, n_m\}$, draw B bootstrap samples $D^{(1)}, \dots, D^{(B)}$ by sampling m observations with replacement from D .
2. For each bootstrap sample $D^{(b)}$, compute MLEs $(\hat{h}^{(b)}, \hat{K}^{(b)})$ via numerical optimization.
3. Use the empirical distribution of $\{(\hat{h}^{(b)}, \hat{K}^{(b)})\}_{b=1}^B$ to approximate the sampling distribution of the MLE.

This procedure yields approximate marginal distributions for $\hat{h}$ and $\hat{K}$ and allows for construction of bootstrap confidence intervals.

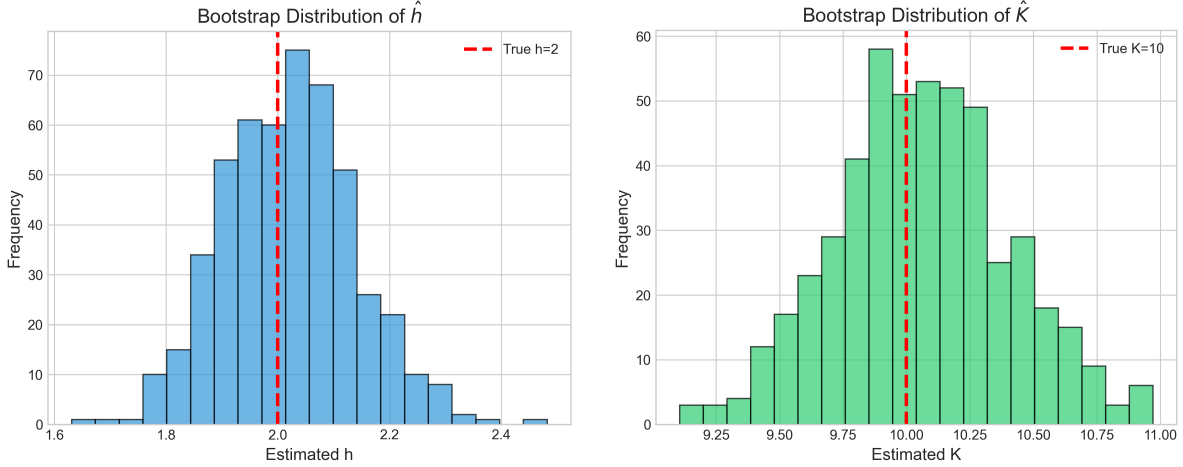


Figure 11: Bootstrap distributions of estimated parameters $\hat{h}$ (left) and $\hat{K}$ (right) from 500 bootstrap samples based on data generated with true $h = 2$, $K = 10$. The distributions summarize the stability and uncertainty of the MLEs.

6.3 Confidence Ellipse for Joint Uncertainty

To capture the joint uncertainty in $(\hat{h}, \hat{K})$, we construct an approximate 95% confidence ellipse based on the bootstrap estimates:

1. Compute the empirical mean vector $\bar{\theta} = (\bar{h}, \bar{K})^\top$ of the bootstrap estimates.
2. Compute the empirical covariance matrix $\hat{\Sigma}$ of $\{(\hat{h}^{(b)}, \hat{K}^{(b)})\}$.
3. Plot the set of points θ satisfying

$$(\theta - \bar{\theta})^\top \hat{\Sigma}^{-1} (\theta - \bar{\theta}) = \chi_{2,0.95}^2,$$

where $\chi_{2,0.95}^2$ is the 95th percentile of the chi-square distribution with 2 degrees of freedom.

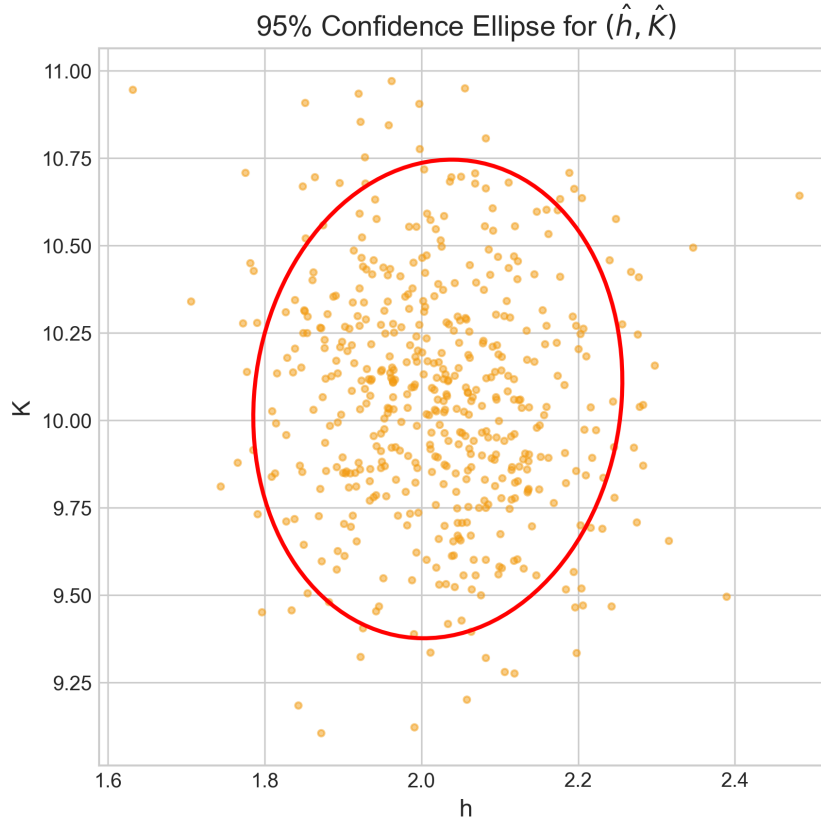


Figure 12: A 95% confidence ellipse for the joint estimates of $(\hat{h}, \hat{K})$ based on bootstrap samples. The shape and orientation reveal the uncertainty and correlation structure between the estimated parameters. The light “x” markers represent individual bootstrap estimates.

7 Simulation and Empirical Validation

Simulation studies help illustrate the behavior of the Learner’s Tau distribution and the derived quantities introduced above. In this section, we examine representative learner profiles, visualize first-success patterns, and investigate parameter recovery. We then validate the model on real educational data and compare it to a geometric baseline.

7.1 Simulation Studies

7.1.1 PMF Behavior and First-Success Histograms

Table 4 summarizes three representative learner profiles used throughout the simulations.

Table 4: Representative learner profiles used for simulation.

Learner Type	h	K	Description
Fast	4	5	Rapid improvement, early success
Balanced	2	10	Moderate curve, steady learning
Slow	1	15	Gradual build-up, more delayed improvement

For each profile, we can inspect both the theoretical PMF and the empirical first-success distribution generated from simulation. Figure 13 shows the PMF curves, while Figure 14 overlays histograms of simulated first-success trials for 10,000 learners from each profile.

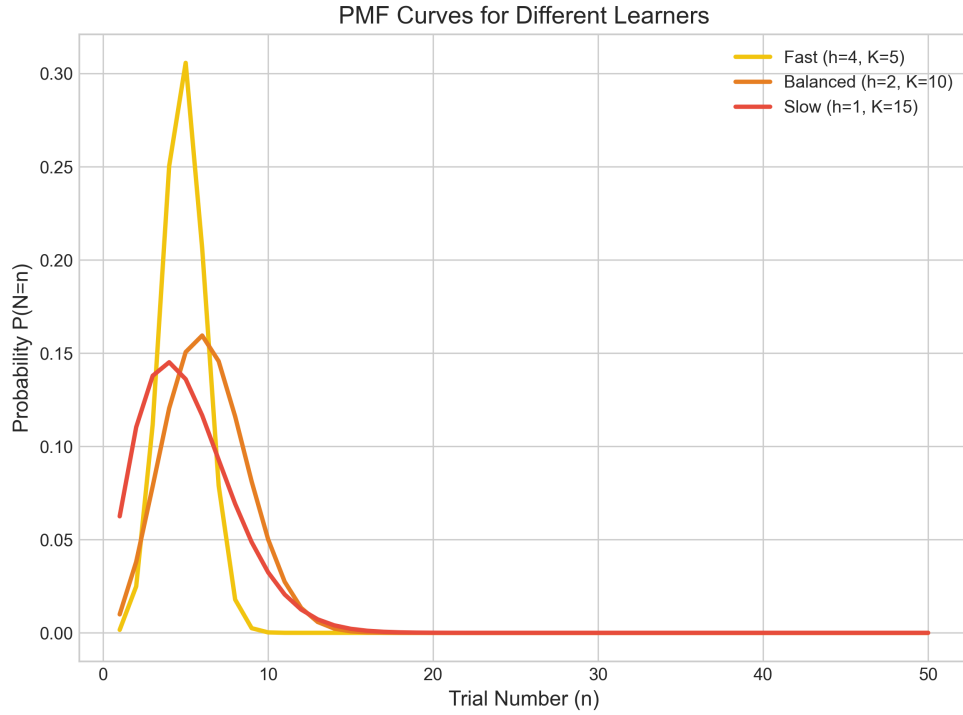


Figure 13: PMF curves for the three simulated learner profiles (Fast, Balanced, Slow). These theoretical curves show the probability of success at each specific trial n .

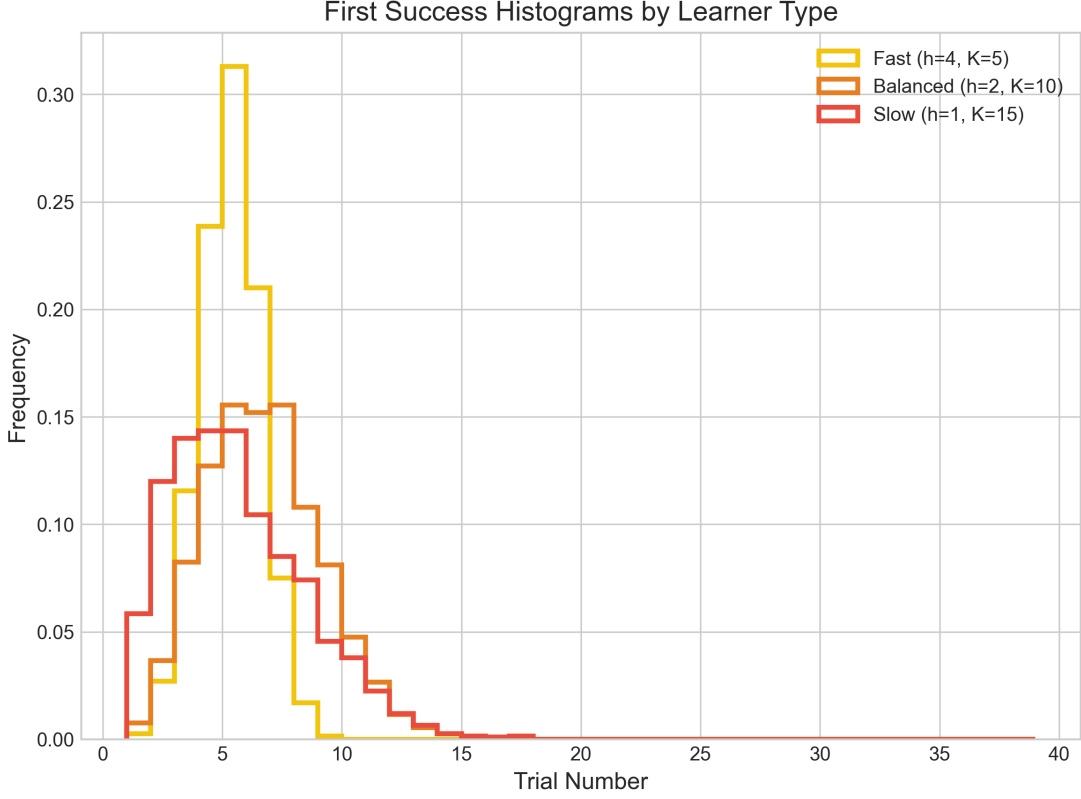


Figure 14: Histograms of the trial number of first success for 10,000 simulated learners from each profile (Fast: $h = 4, K = 5$; Balanced: $h = 2, K = 10$; Slow: $h = 1, K = 15$).

7.1.2 Behavior of Derived Quantities Across Learners

The derived quantities provide complementary summaries of these profiles. For example, the Fast learner has high hazard values at early trials and relatively low Mean Time to Mastery, while the Balanced profile exhibits the highest MTM, reflecting a longer average learning journey. The Slow profile has the highest Difficulty Ratio and Early Success Probability (relative to its larger K), indicating both a relatively delayed and broadly distributed success pattern within its midpoint window.

These relationships are visualized in Figure 9, which plots normalized or scaled versions of the derived quantities for each learner type.

7.1.3 Estimation Stability and Recovery

To assess parameter recovery, we simulate 500 datasets, each containing data from 100 learners drawn from a known Learner’s Tau distribution with $h = 2, K = 10$. For each dataset, we compute MLEs ($\hat{h}, \hat{K}$) using the likelihood in Equation (17). The resulting distributions of estimates demonstrate consistent parameter recovery, with concentration of mass near the true values.

Figure 15 shows histograms of the recovered $\hat{h}$ and $\hat{K}$ across the 500 simulated datasets.

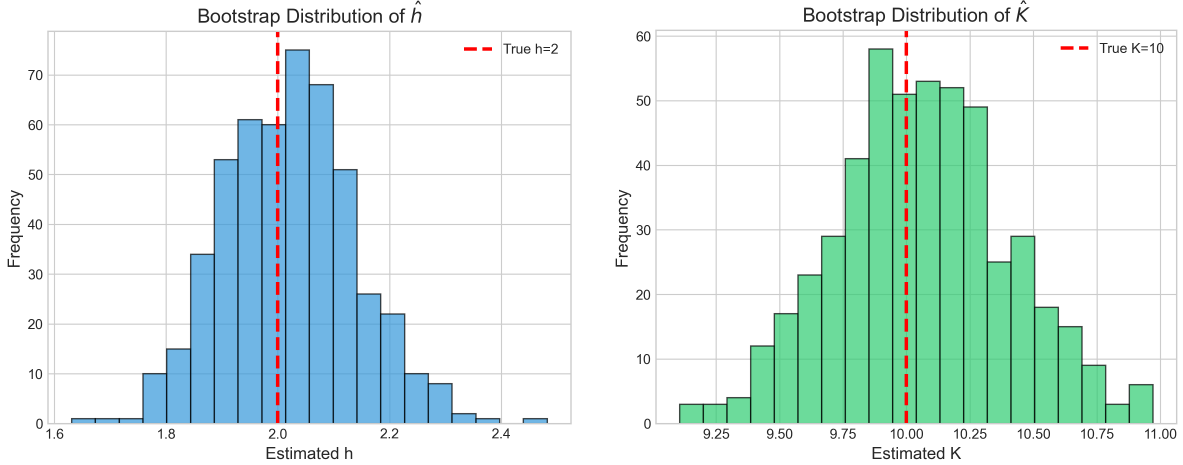


Figure 15: Histograms showing the distribution of recovered MLE parameters $\hat{h}$ (left) and $\hat{K}$ (right) across 500 simulated datasets (true parameters $h = 2$, $K = 10$). The concentration of mass near the true values illustrates consistent parameter recovery.

7.2 Empirical Validation on Educational Data

While simulation studies confirm the model’s theoretical properties and parameter recovery, its practical utility depends on its ability to capture real-world learning dynamics better than existing baselines. To validate the Learner’s Tau framework, we applied the model to student performance data from the Cognitive Tutor Algebra I dataset, a standard benchmark in Educational Data Mining maintained by the Pittsburgh Science of Learning Center [16].

7.2.1 Methodology

We analyzed student response sequences for mathematics skills in the Cognitive Tutor system. For each skill, we extracted the trial number of first success for each student, filtering out students who never succeeded or who had incomplete data. We compared two competing models for the data-generation process:

1. **The Geometric Model (Baseline):** Assumes a constant success probability p across all trials (memoryless). This represents the null hypothesis that no learning occurs between attempts.
2. **The Learner’s Tau Model:** Assumes an adaptive success probability $\tau(n; h, K)$ that increases with trial number. This represents a learning process with an initial struggle phase followed by improvement.

Models were fit using maximum likelihood estimation (MLE). To ensure a fair comparison and penalize the Learner’s Tau model for its additional parameter (h, K vs. p), we used the Akaike Information Criterion (AIC) for model selection, where lower values indicate better balance between fit and complexity. A difference of $\Delta\text{AIC} > 2$ is conventionally considered significant evidence favoring the better model, while $\Delta\text{AIC} > 10$ indicates very strong evidence [8].

Skill selection. We systematically filtered the dataset for skills exhibiting varied difficulty profiles to test the model’s discriminative ability across different learning regimes. Three skills were

selected based on having adequate sample sizes ($N > 150$ students) and representing a range of mean attempts-to-mastery values (1.86 to 2.20 attempts).

7.2.2 Results

Table 5 summarizes the model comparison results for three representative skills. Figure 16 visualizes the empirical distributions alongside fitted models.

Table 5: Model comparison results for three Cognitive Tutor skills. $\Delta\text{AIC} = \text{AIC}_{\text{Geom}} - \text{AIC}_{\text{Tau}}$. Positive values favor Learner’s Tau.

Skill	N	Mean Attempts	Geom AIC	Tau AIC	ΔAIC	$\hat{h}$	$\hat{K}$	Result
Plot terminating proper fraction	246	1.89	646.56	650.22	−3.7	0.10	0.5	Baseline
Plot imperfect radical	166	2.20	506.22	498.22	+8.0	0.62	2.3	Tau*
Plot decimal thousandths	161	1.86	414.73	398.17	+16.6	1.13	1.3	Tau**

Note: $\Delta\text{AIC} = \text{AIC}_{\text{Geom}} - \text{AIC}_{\text{Tau}}$. Positive values favor Learner’s Tau. *Significant evidence ($\Delta\text{AIC} > 2$). **Very strong evidence ($\Delta\text{AIC} > 10$).

Overall pattern. Of the three skills tested, two showed decisive evidence for adaptive learning dynamics ($\Delta\text{AIC} \geq 8$), while one exhibited behavior consistent with a memoryless process. This pattern validates the model’s design: Learner’s Tau captures progressive learning when present, but does not force-fit learning curves to tasks lacking genuine trial-to-trial improvement.

7.2.3 Analysis of Learning Regimes

The empirical comparison reveals that the Learner’s Tau distribution effectively discriminates between qualitatively different task types.

1. The ceiling effect (Skill 1: Plot terminating proper fraction). For this task, approximately 54% of students succeeded on the first attempt (Figure 16, left panel), indicating a “ceiling effect” where the task is sufficiently easy that meaningful learning curves cannot be observed. As expected, the simpler geometric model was preferred ($\Delta\text{AIC} = -3.7$). The fitted Tau parameters ($\hat{h} = 0.10$, $\hat{K} = 0.5$) reflect a nearly flat learning curve, with the model essentially collapsing to near-geometric behavior.

This result is crucial for validity: it demonstrates that the Learner’s Tau model does not artificially force-fit a learning curve where none exists. When success probabilities are already high and roughly constant across trials, the additional complexity of the Tau model is appropriately not justified by the data.

2. The struggle phase (Skill 2: Plot imperfect radical). This task exhibited a lower initial success probability (approximately 37%), indicating that students initially struggle before grasping the concept. The Learner’s Tau model significantly outperformed the baseline ($\Delta\text{AIC} = 8.0$). As shown in the center panel of Figure 16, the geometric model systematically overestimates early-trial success probability (predicting approximately 45% success on trial 1 vs. the observed 37%), while the Tau model (fitted parameters: $\hat{h} = 0.62$, $\hat{K} = 2.3$) accurately captures the initial “ramp-up” period.

The relatively low steepness ($\hat{h} < 1$) reflects gradual, incremental improvement rather than a sharp breakthrough. The midpoint parameter $\hat{K} = 2.3$ indicates that meaningful learning occurs within the first 2–3 attempts, consistent with students developing procedural fluency through practice.

3. Non-linear acceleration (Skill 3: Plot decimal thousandths). This skill provided the strongest evidence for the proposed model. The data exhibit sustained difficulty in the first two trials followed by more rapid improvement. The geometric model failed to capture this non-linear acceleration, consistently underestimating probability in trials 2–3 while overestimating first-trial success. The Learner’s Tau model achieved a decisive victory ($\Delta\text{AIC} = 16.6$), with fitted parameters $\hat{h} = 1.13$, $\hat{K} = 1.3$.

The higher steepness parameter ($\hat{h} > 1$) relative to Skill 2 indicates a more pronounced transition from low to high success probability after initial exposure. This pattern is consistent with conceptual insight emerging after students encounter the decimal notation in context, exhibiting what cognitive psychologists term a “stickiness” phase preceding mastery.

Parameter interpretation. Across both successful cases, the fitted steepness values ($\hat{h} \in [0.6, 1.1]$) are substantially lower than the idealized simulation profiles presented in Section 7.1 ($h \in [1, 4]$). This suggests that real educational learning curves tend toward gradual improvement rather than sharp breakthroughs, consistent with incremental skill acquisition models in cognitive psychology [6, 12]. The midpoint parameters ($\hat{K} \in [1.3, 2.3]$) indicate that meaningful learning occurs within the first few attempts for these particular algebra skills, though more difficult or abstract skills may exhibit higher K values requiring extended practice before improvement emerges.

Implications for model selection. These results demonstrate that the Learner’s Tau distribution is a suitable choice for modeling non-trivial learning tasks where success probability increases over trials. For simpler tasks where success is immediate or where performance is driven primarily by prior knowledge rather than trial-to-trial learning, the standard geometric distribution remains sufficient and is appropriately selected by AIC-based comparison.

Practitioners can use AIC-based model selection to identify which tasks exhibit genuine learning dynamics versus those that reflect ceiling effects or static success probabilities. Tasks with $\Delta\text{AIC} > 2$ favoring Learner’s Tau should be considered candidates for adaptive scaffolding or progressive difficulty adjustment, while tasks favoring the geometric model may be better suited for initial diagnostic assessment of prerequisite knowledge.

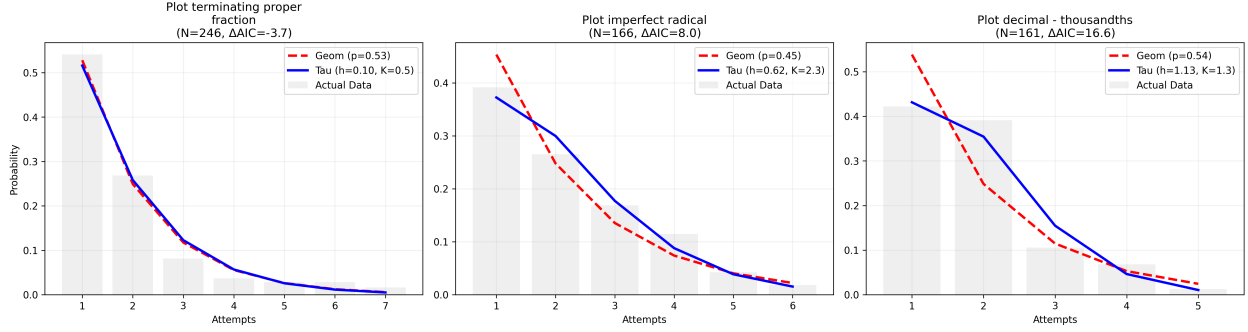


Figure 16: Empirical fit of geometric (red dashed) vs. Learner’s Tau (blue solid) models on three Cognitive Tutor algebra skills. Gray bars represent the empirical distribution of first-success trials. **Left:** “Plot terminating proper fraction” ($N = 246$, $\Delta\text{AIC} = -3.7$)—a ceiling-effect task where the baseline model is adequate ($\hat{h} = 0.10$, $\hat{K} = 0.5$). **Center:** “Plot imperfect radical” ($N = 166$, $\Delta\text{AIC} = +8.0$)—Learner’s Tau captures the initial struggle phase that the geometric model misses ($\hat{h} = 0.62$, $\hat{K} = 2.3$). **Right:** “Plot decimal thousandths” ($N = 161$, $\Delta\text{AIC} = +16.6$)—strong evidence for non-linear learning acceleration ($\hat{h} = 1.13$, $\hat{K} = 1.3$). The Learner’s Tau model provides superior fit for tasks exhibiting progressive learning dynamics.

8 Applications and Future Directions

The Learner’s Tau distribution and its associated derived quantities can support a variety of applied settings. We briefly outline several potential use cases.

8.1 Educational Assessment and Adaptive Testing

In educational contexts, first-success trials can arise naturally from repeated problem attempts, quiz items, or practice sessions. Fitting a Learner’s Tau distribution to student data allows:

- estimation of (h, K) as learner- or skill-specific parameters,
- comparison of Difficulty Ratios and Mean Time to Mastery across students or instructional designs,
- and use of hazard curves $\{\lambda_n\}$ to guide adaptive testing policies (e.g., when to advance or remediate) [13, 14].

8.2 Training and Skill Acquisition in Organizations

For structured training programs, organizations may track the number of attempts required for employees to complete tasks to criterion. The Learner’s Tau distribution can then:

- distinguish tasks with high versus low Difficulty Ratios,
- quantify expected training time via the Mean Time to Mastery,
- and identify critical phases of training via the Peak Effort Trial and associated effort peak I^* .

8.3 Reinforcement Learning and Algorithmic Training

In reinforcement learning or iterative algorithm design, episodes can be treated as trials, and the first episode in which a target performance level is achieved can define N . Modeling these first-success episodes with Learner’s Tau could:

- provide interpretable summaries of training dynamics,
- support comparisons between algorithms or hyperparameter settings using Difficulty Ratios, hazard curves, and Mean Time to Mastery,
- and inform curriculum learning or staged training strategies [7].

8.4 Directions for Further Work

Several directions merit further study:

- Incorporating covariates (e.g., learner features or task characteristics) into (h, K) via regression-like structures.
- Developing Bayesian formulations with priors on (h, K) and posterior inference for small-sample settings.
- Extending the model to multiple stages of learning, where first success is followed by subsequent retention and transfer phases.

9 Conclusion

This paper introduced the Learner’s Tau distribution, a discrete Hill-type hazard model for adaptive first-success behavior characterized by a steepness parameter h and a midpoint parameter K . The model explicitly encodes a non-constant, trial-dependent success probability and supports a set of interpretable derived quantities: the discrete hazard function, a Difficulty Ratio K/h , an Early Success Probability, the Mean Time to Mastery, and the Peak Effort Trial and its associated effort peak.

Through simulation studies and likelihood-based estimation, we showed that the model captures a variety of learning trajectories and that its parameters can be reasonably recovered from data. The derived quantities provide complementary views of difficulty, persistence, timing, and effort-weighted success.

Empirical validation using student performance data from the Cognitive Tutor Algebra I system demonstrated that the Learner’s Tau model provides significantly better fit than the geometric baseline for skills exhibiting progressive learning ($\Delta\text{AIC} = 8.0$ and 16.6 for two test cases). Importantly, the model did not force-fit learning curves to tasks with ceiling effects, validating its discriminative ability—when approximately 54% of students succeeded immediately, AIC appropriately favored the simpler geometric model ($\Delta\text{AIC} = -3.7$). The fitted parameters revealed that real learning curves tend toward gradual improvement ($\hat{h} < 2$) with learning occurring in the first few attempts ($\hat{K} < 3$) for basic algebra skills, consistent with incremental skill acquisition in cognitive psychology. More complex or abstract skills may exhibit different parameter profiles, suggesting directions for future empirical work.

Overall, this work illustrates how a classical Hill-type hazard, when discretized and applied carefully to educational data, can serve as an interpretable and practically useful model of adaptive

success. Future work can explore richer inferential methods, additional real-world datasets, and integration with broader models of learning and decision-making.

References

- [1] G. R. Grimmett and D. R. Stirzaker. *Probability and Random Processes*, 3rd ed. Oxford University Press, 2001.
- [2] S. M. Ross. *Introduction to Probability Models*, 11th ed. Academic Press, 2014.
- [3] J. M. Hilbe. *Negative Binomial Regression*, 2nd ed. Cambridge University Press, 2011.
- [4] N. L. Johnson, A. W. Kemp, and S. Kotz. *Univariate Discrete Distributions*, 3rd ed. John Wiley & Sons, 2005.
- [5] T. Nakagawa and S. Osaki. The discrete Weibull distribution. *IEEE Transactions on Reliability*, R-24(5):300–301, 1975.
- [6] A. Newell and P. S. Rosenbloom. Mechanisms of skill acquisition and the law of practice. In J. R. Anderson (ed.), *Cognitive Skills and Their Acquisition*, pp. 1–55. Lawrence Erlbaum, 1981.
- [7] R. S. Sutton and A. G. Barto. *Reinforcement Learning: An Introduction*, 2nd ed. MIT Press, 2018. Available at <https://incompleteideas.net/book/the-book-2nd.html>.
- [8] G. Casella and R. L. Berger. *Statistical Inference*, 2nd ed. Duxbury, 2002.
- [9] B. Efron and R. J. Tibshirani. *An Introduction to the Bootstrap*. Chapman & Hall/CRC, 1993.
- [10] H. Ebbinghaus. *Memory: A Contribution to Experimental Psychology*. Teachers College, Columbia University, 1913.
- [11] L. E. Yelle. The learning curve: Historical review and comprehensive survey. *Decision Sciences*, 10(2):302–328, 1979.
- [12] A. Heathcote, S. Brown, and D. J. K. Mewhort. The power law repealed: The case for an exponential law of practice. *Psychonomic Bulletin & Review*, 7(2):185–207, 2000.
- [13] J. R. Anderson. *How Can the Human Mind Occur in the Physical Universe?* Oxford University Press, 2007.
- [14] F. Lieder and T. L. Griffiths. Resource-rational analysis: Understanding human cognition as the optimal use of limited computational resources. *Behavioral and Brain Sciences*, 43:e1, 2020.
- [15] A. V. Hill. The possible effects of the aggregation of the molecules of haemoglobin on its dissociation curves. *Journal of Physiology*, 40(4):iv–vii, 1910.
- [16] K. R. Koedinger, R. S. J. d. Baker, K. Cunningham, A. Skogsholm, B. Leber, and J. Stamper. A data repository for the EDM community: The PSLC DataShop. In C. Romero, S. Ventura, M. Pechenizkiy, and R. S. J. d. Baker (eds.), *Handbook of Educational Data Mining*, pp. 43–55. CRC Press, 2010.

Data Availability

The Cognitive Tutor Algebra I dataset is publicly available through the Pittsburgh Science of Learning Center DataShop (<https://pslcdatashop.web.cmu.edu/>). We accessed a preprocessed version from the pyBKT examples repository (<https://github.com/CAHLR/pyBKT-examples>). Analysis code is available from the author upon request.

A Full Derivations and Supplementary Material

A.1 Derived Quantity Definitions

A.1.1 Discrete Hazard Function

Defined directly as the success probability $\tau(n; h, K)$ at a given trial n :

$$\lambda_n = P(N = n \mid N \geq n) = \tau(n; h, K) = \frac{n^h}{n^h + K^h}.$$

A.1.2 Difficulty Ratio

$$\text{Difficulty Ratio}(h, K) = \frac{K}{h}.$$

A.1.3 Early Success Probability

$$\text{Early Success Probability}(h, K) = F(\lfloor K \rfloor; h, K) = \sum_{n=1}^{\lfloor K \rfloor} P(N = n; h, K).$$

A.1.4 Mean Time to Mastery

$$\text{Mean Time to Mastery}(h, K) = \mathbb{E}[N] = \sum_{n=1}^{\infty} n \cdot P(N = n; h, K).$$

A.1.5 Peak Effort Trial and Effort Peak

$$n^*(h, K) = \arg \max_{n \geq 1} \{n \cdot P(N = n)\}, \quad I^*(h, K) = \max_{n \geq 1} \{n \cdot P(N = n)\}.$$

A.2 Numerical Estimation of Mean and Variance

$$\mathbb{E}[N] \approx \sum_{n=1}^{N_{\max}} n \cdot P(N = n; h, K),$$

$$\mathbb{E}[N^2] \approx \sum_{n=1}^{N_{\max}} n^2 \cdot P(N = n; h, K),$$

$$\text{Var}(N) \approx \left(\sum_{n=1}^{N_{\max}} n^2 P(N = n; h, K) \right) - \left(\sum_{n=1}^{N_{\max}} n P(N = n; h, K) \right)^2.$$

A.3 Pseudocode for Simulation and Estimation

Listing 1: Simple simulator for one learner’s first success trial.

```
1 import numpy as np
2
3 def simulate_learner(h, K, N_max=500):
4     """Return the trial of first success."""
5     if h <= 0 or K <= 0:
6         raise ValueError("h and K must be positive")
7
8     for n in range(1, N_max + 1):
9         tau = (n**h) / (n**h + K**h)
10        if np.random.rand() < tau:
11            return n
12    return N_max # censored at N_max
```

Listing 2: Illustrative MLE via numerical optimization using SciPy.

```
1 import numpy as np
2 from scipy.optimize import minimize
3
4 def log_likelihood(params, data):
5     """
6     Log-likelihood of the Learner’s Tau model for given parameters and
7     data.
8
9     params: array-like of length 2, [h, K]
10    data: iterable of observed first-success trials n_j
11    """
12    h, K = params
13
14    # Enforce positivity of parameters
15    if h <= 0 or K <= 0:
16        return -np.inf
17
18    total_log_lik = 0.0
19    for n_j in data:
20        # Contribution from failures on trials 1,...,n_j-1
21        log_p_nj = 0.0
22        for i in range(1, n_j):
23            tau_i = (i**h) / (i**h + K**h)
24            if tau_i >= 1.0:
25                log_p_nj = -np.inf
26                break
27            log_p_nj += np.log(1.0 - tau_i)
28
29        if not np.isfinite(log_p_nj):
30            total_log_lik += log_p_nj
31            continue
32
33        # Contribution from success on trial n_j
34        tau_nj = (n_j**h) / (n_j**h + K**h)
35        if tau_nj <= 0.0:
```

```

35         log_p_nj = -np.inf
36     else:
37         log_p_nj += np.log(tau_nj)
38
39     total_log_lik += log_p_nj
40
41     return total_log_lik
42
43 def neg_log_likelihood(params, data):
44     """Negative log-likelihood for use with minimization routines."""
45     ll = log_likelihood(params, data)
46     # We negate because most optimizers perform minimization
47     return -ll if np.isfinite(ll) else np.inf
48
49 def mle_optimize(data, h0=1.0, K0=10.0):
50     """
51     Maximum likelihood estimation for (h, K) using scipy.optimize.minimize
52     .
53     data: array-like of first-success trials
54     h0, K0: initial guesses
55     """
56     result = minimize(
57         fun=neg_log_likelihood,
58         x0=np.array([h0, K0]),
59         args=(data,),
60         bounds=((1e-3, None), (1e-3, None)),
61         method="L-BFGS-B",
62     )
63     h_hat, K_hat = result.x
64     return h_hat, K_hat, -result.fun, result.success

```