

The Hot Stove Effect in Venture Capital: How Adaptive Sampling Generates a Negativity Bias

Maxim Chzhan-Vin-Zin
maksimchzanvinzin@gmail.com

ABSTRACT

Venture capital (VC) decision-making operates under bounded rationality and sparse, noisy feedback. This paper examines the Hot Stove Effect, a negativity bias that arises when the intensity of future sampling depends on the outcome of early sampling. Using a stochastic simulation of a two-period investment problem, we analyse how an adaptive sampling policy, under which a favourable initial signal raises and an unfavourable one lowers the number of subsequent draws, distorts the beliefs of learners ranging from naive averages to Bayesian updaters. We derive a closed-form expression for the expected belief and show that it is strictly negative whenever the commitment following a positive signal exceeds that following a negative one. The bias is attenuated, but not eliminated, by Bayesian updating with a correctly centred prior: a rational learner remains negatively biased in the mean because the sample size is endogenous to the sign of the early signal. Independently of the updating rule, a majority of learners (about 58.5% in the baseline) end below the true value, because the underestimation share is governed by the sign of the accumulated evidence and not by the learning rule. At low volatility the resulting belief distribution is negatively skewed, with its mode near the true value and a heavy pessimistic tail; for the Bayesian learner the asymmetry reverses as volatility rises, while the underestimation share does not move at all. The findings indicate that the systematic undervaluation of neutral opportunities in VC follows structurally from payoff-contingent sampling and persists even under rational inference.

INTRODUCTION

Decision-making in firms reflects a coalition of participants with disparate demands, limited attention, and conflicting incentives (Cyert & March, 1963). This paper addresses a question within that tradition: why do reasonably successful and adaptive organisations exhibit behaviour that appears more systematically biased than that of their individual members? The question is sharpest where stakes are high and feedback is sparse, as in venture capital (VC), and it is illustrated by a natural experiment.

The London Underground Paradox

The question is illustrated by a natural experiment. In February 2014 a strike on the London Underground closed a subset of stations for approximately two days, forcing many commuters to abandon habitual

routes and try alternatives. Analysing the episode, Larcom, Rauch, and Willems (2017) found that a significant fraction of commuters did not revert to their prior routes once the strike ended: forced experimentation had revealed faster journeys that these commuters had previously ignored.

Why had rational commuters used suboptimal routes for years? The answer lies in the cost of experimentation. Trying an unknown route risks a penalty (a longer journey), whereas the habitual route yields a predictable outcome. Repeated avoidance of that risk sustains a suboptimal equilibrium. The commuters were constrained by the structure of their own experiential learning.

Venture Capital as a "Hot Stove"

In venture capital this dynamic is amplified by high uncertainty. Decision-makers operate under *bounded rationality*: they work within strict cognitive, informational, and computational limits. Rather than optimising every choice, VCs "satisfice," evaluating startups under sparse data and time pressure.

These conditions make learning from experience essential, yet experience is often a poor teacher because of the Hot Stove Effect. Named after Mark Twain's observation that "we should be careful to get out of an experience only the wisdom that is in it---and stop there; lest we be like the cat that sits down on a hot stove-lid. She will never sit down on a hot stove-lid again---and *that is well*; but also she will never sit down on a cold one any more" (Twain, 1897, ch. XI), the effect describes an asymmetry in error correction. The qualification carries the argument: the avoidance is adaptive, and that is precisely why it is costly. Errors of overestimation (judging a weak startup to be strong) tend to be corrected, because the investor commits capital and eventually observes the failure. Errors of underestimation (judging a strong startup to be weak) lead to avoidance: the investor passes, and unless the company becomes a widely visible success, no corrective feedback is received.

Research Objectives and Contribution

This paper connects theories of organisational learning with a computational model of belief formation, and studies how a payoff-contingent sampling policy biases beliefs. The mechanism was introduced by Denrell and March (2001) and formalised for a soft, propensity-based sampling rule by Denrell (2007). Our contribution is one new result and two refinements of existing ones. The refinements are these. We derive a single closed-form expression for the expected belief of a one-parameter family of learners spanning naive averaging and Bayesian updating under a hard threshold rule, where Denrell (2007) treats a soft, propensity-based rule; and we show that Bayesian shrinkage towards a correctly centred prior attenuates the mean bias monotonically in the shrinkage constant but never removes it, sharpening into a comparative static what Denrell and Le Mens (2011) establish qualitatively for rational sampling. The new result is the third, and it is narrower than a separation of the mean bias from the majority-underestimation property, which Denrell (2007) already draws. We show that the majority property is *invariant*: because the sign of the belief equals the sign of the accumulated evidence for every shrinkage constant, the share of learners below the truth is not merely above one half but numerically identical across the entire estimator family and unchanged by the payoff volatility, whereas the mean bias moves with both. Rationality repairs part of the first quantity and none of the second. Together these

results characterise the systematic undervaluation of neutral verticals in VC as a product of payoff-maximising sampling rather than of preference or of inferential failure.

LITERATURE REVIEW

This section situates the Hot Stove Effect within organisational economics, behavioural decision theory, and models of experiential learning, and identifies the gap that the present study addresses.

Bounded Rationality and Internal Organisation

The analysis of internal organisation is rooted in the concept of *bounded rationality*, introduced by Simon (1955), who argued that decision-makers operate under limits on cognition, information, and computation. In high-pressure environments such as VC, investors cannot compute the expected value of every possible deal. They satisfice, choosing options that are good enough given simplified representations of a complex environment, and they rely on heuristics that reduce cognitive load. Such heuristics are efficient, but they produce systematic errors when the environment is noisy.

Exploration versus Exploitation

Organisations may be viewed as adaptive systems that balance two processes: *exploration* (variation, experimentation, risk-taking) and *exploitation* (refinement of existing competencies). March (1991) described the tension between these forces. Exploitation yields immediate, reinforcing feedback: when a firm invests in a familiar sector and observes a moderate return, the loop is short. Exploration involves long delays and uncertain outcomes. Excessive exploitation risks obsolescence in a local optimum; excessive exploration wastes resources without capturing the returns to discovery.

Levinthal and March (1993) argued that the myopia of learning tilts this balance toward exploitation. Organisations favour the short run over the long run (temporal myopia), the near neighbourhood over the distant (spatial myopia), and success over failure: under failure myopia, "organizational learning oversamples successes and undersamples failures" (Levinthal & March, 1993, p. 110), so that the risks of failure are underestimated. The third form is the mirror image of the mechanism studied here. It operates on the alternatives an organisation continues to pursue, where accumulated success breeds confidence; the Hot Stove Effect operates on those it has stopped pursuing, where an early failure is never revised. In the present context this implies that VCs gravitate toward familiar bets with large effective sample sizes and avoid the high-variance regions of the opportunity space.

Risk Aversion versus Information Sampling

A distinction central to this study is that between *risk aversion*, a property of preferences, and the Hot Stove Effect, a property of the information set. Standard accounts, including Prospect Theory (Kahneman & Tversky, 1979), attribute conservatism to the utility function: losses loom larger than equivalent gains. The Hot Stove Effect is distinct. As shown by Denrell (2007), a risk-neutral agent with a linear utility function can behave as if risk-averse, not because it fears loss, but because it stops gathering information

about alternatives that yield early negative feedback. The bias is located in the information set rather than in preferences. This paper builds on that distinction: VCs need not be averse to risk to undervalue it; they can be misled by their own sampling policy.

A related structural bias operates on the chosen alternatives rather than the rejected ones. Harrison and March (1984) showed that unbiased estimation errors, combined with the act of selecting the apparently best option, produce a systematic tendency toward postdecision disappointment. The Hot Stove Effect is the counterpart for rejected options: the errors that suppress future sampling are precisely those that are never corrected.

The Hot Stove Effect: Asymmetry in Error Correction

The Hot Stove Effect is a negativity bias arising from the adaptive character of learning. Errors of overestimation are likely to be corrected because the learner continues to sample the alternative. Errors of underestimation lead to avoidance or reduced sampling, so the learner is not exposed to the information required to correct the mistake.

The asymmetry is especially damaging in VC because feedback is noisy. A strong startup can fail through bad timing (a false negative), and a weak one can succeed through luck (a false positive). Because investors withdraw funding from failing companies, they routinely observe false positives resolve into failures, but rarely observe a false negative correct itself. This one-sided data collection reinforces scepticism while penalising optimism, a pattern consistent with the asymmetric-feedback account of misplaced cynicism (Fetchenhauer & Dunning, 2010).

Knowledge Gap

Three results bound the present contribution, and it is worth stating precisely where each one stops. Denrell and March (2001) establish the mechanism itself: the reproduction of successful actions inherent in adaptive processes biases a learner against alternatives that initially appear worse than they are, and the bias is attenuated whenever that reproduction is slowed, made imprecise, or recalled less reliably. Denrell (2007) supplies the analytical theory for a soft, propensity-based rule, and carries it further than is sometimes recognised. He shows that a risk-neutral learner may come to prefer a sure thing to a symmetric alternative of equal expected value "even if the decision maker follows an optimal policy of learning" (p. 177), and he establishes the majority result directly: for "more than 50% of all decision makers" the posterior expectation of the uncertain alternative is negative, and this holds "for all $t > 1$ " (p. 181), with 64.7 per cent underestimating in a ten-period Bernoulli bandit solved by dynamic programming. Denrell and Le Mens (2011) complete the argument by proving that interested sampling generates systematic judgemental error even under fully rational processing, dispensing with the naivety assumption on which earlier sampling accounts rested.

The separation of the majority property from the mean bias is therefore not new, and we do not claim it. What is new is its *invariance*. Denrell's results establish that the underestimation share exceeds one half under particular rules, and the share he reports varies with the rule. Proposition 1 establishes something stronger for the threshold model: because the sign of the belief equals the sign of the accumulated

evidence for every $a \geq 0$, the share is not merely above one half but numerically identical across the entire family of estimators, from naive averaging to any degree of Bayesian shrinkage, and invariant to the payoff volatility σ . The mean bias, by contrast, moves with both. The contribution is thus a dissociation: one statistic of the belief distribution responds to rationality and to volatility, the other is fixed by the sign of the accumulated evidence and responds to neither.

METHODS

We use a computational simulation, supported by exact analytical results, to characterise the belief bias induced by payoff-contingent sampling.

Rationale for a Simulation

The study uses synthetic data generated by stochastic simulation rather than a field dataset, for three reasons. First, *isolation of mechanism*: in field data it is not possible to separate luck from judgement or market timing, whereas a simulation fixes the true mean u explicitly. Second, *observation of counterfactuals*: field data do not reveal the returns of investments that were not made, whereas the simulation records the foregone payoffs exactly. Third, *scale*: establishing a statistical regularity to high precision requires many independent trials, well beyond the historical volume of any single firm.

Model and Parameters

The design uses a two-period learning setup ($t = 1, 2$) that represents sequential investment rounds (for example, a seed round followed by a Series A).

Interpretation of a draw. A draw is a diligence signal on one company within the vertical: a reviewed pitch, a meeting with the management team, a reference call, or a completed technical assessment. The scout sample of $k = 2$ corresponds to the two-stage initial screen documented by Gompers, Gornall, Kaplan, and Strebulaev (2020), in which an opportunity is first evaluated by the originating partner and, if it survives, generates at least one meeting with the management team before the firm decides whether to take it further. The commitment m is the number of further signals the firm gathers after that decision. The magnitudes are disciplined by the same survey rather than chosen freely: for every closed deal the median firm considers 101 opportunities and carries only about ten as far as a partners' meeting, so an opportunity that survives the initial screen attracts roughly an order of magnitude more subsequent attention than one that does not, which is the ratio $h/l = 10$ used below. That table also supplies an empirical counterpart to σ : information-technology firms consider 151 opportunities per investment against 78 in healthcare, consistent with wider screening in the higher-variance sector. Under this reading the two periods are a screening decision and a diligence decision within a single fund cycle rather than a seed round followed by a Series A, which keeps the timescale of the mechanism and that of the model consistent, since diligence information arrives in weeks whereas exit outcomes resolve over years.

Payoff distribution. The payoff of a startup vertical is a random variable $x \sim N(u, \sigma^2)$. We set the true mean $u = 0$, representing a neutral vertical whose expected return equals its opportunity cost. This is the hardest case for a learner to distinguish and isolates the bias. We vary the standard deviation $\sigma \in \{1, 5\}$; a higher σ represents a deep-tech or moonshot sector with high variance.

Adaptive sampling policy. In period $t = 1$ the agent takes a small scout sample of size $k = 2$ and computes the initial mean $\bar{x}_1$. It then chooses a commitment level m by comparing $\bar{x}_1$ to a threshold $c = 0$:

$$m = \begin{cases} h = 10 & \text{if } \bar{x}_1 > c \quad (\text{high commitment}) \\ l = 1 & \text{if } \bar{x}_1 \leq c \quad (\text{low commitment}). \end{cases}$$

In period $t = 2$ the agent takes m additional draws from the same distribution.

Throughout, $\mathbb{1}\{A\}$ denotes the indicator of the event A , equal to one when A occurs and zero otherwise.

Learners. Let $S = \sum_{i=1}^{k+m} x_i$ be the accumulated evidence. We consider a one-parameter family of belief estimators indexed by a shrinkage constant $a \geq 0$:

$$b = \frac{S}{k + m + a}.$$

Two members of this family are studied. The *averaging learner* ($a = 0$) reports the simple sample mean. The *Bayesian learner* adopts a conjugate prior $N(\mu_0, \tau_0^2)$ with $\mu_0 = 0$ and known variance σ^2 ; its posterior mean equals $S / (k + m + \sigma^2 / \tau_0^2)$, so it corresponds to $a = \sigma^2 / \tau_0^2$. We set $\tau_0^2 = 1$, representing an analyst's prior on the vertical's mean return whose scale does not move with sector volatility. The prior mean is centred at the true value, which is the most favourable specification for the Bayesian learner.

The Bayesian learner is also assumed to know the payoff variance σ^2 exactly, so that only the mean is uncertain. This is the standard conjugate specification and is again the case most favourable to that learner, since uncertainty about σ^2 would make the amount of shrinkage itself random.

Simulation Algorithm

For each agent i in a population of $N = 10 \times h(0.08em)000 \times h(0.08em)000$:

+ *Initialization.* Draw k values from $N(0, \sigma^2)$ and compute $\bar{x}_{1,i}$. + *Policy.* If $\bar{x}_{1,i} > c$ set $m_i = h$; otherwise set $m_i = l$. + *Second period.* Draw m_i further values from the same distribution. + *Belief update.* Compute the averaging belief $S_i / (k + m_i)$ and the Bayesian posterior mean $S_i / (k + m_i + \sigma^2 / \tau_0^2)$. + *Aggregation.* Store the beliefs and compute the population statistics.

The simulation is executed over 10^7 agents with a fixed random seed (20260201). The code that reproduces every reported statistic and Figure 1 accompanies the paper.

Analytical Results

The central object is the expected belief. Theorem 1 gives it in closed form for the entire family of estimators.

Theorem 1. Let $x_1, \dots, x_{k+m}$ be i.i.d. $N(u, \sigma^2)$ with $u = 0$, let the sample size follow the adaptive rule $m = h \mathbb{1}\{\bar{x}_1 > c\} + l \mathbb{1}\{\bar{x}_1 \leq c\}$, and let the belief be $b = S / (k + m + a)$ with shrinkage $a \geq 0$ and $S = \sum_{i=1}^{k+m} x_i$. Then

$$E[b] = -\frac{\sigma\sqrt{k}(h-l)}{\sqrt{2\pi}(k+h+a)(k+l+a)} e^{-\frac{c^2 k}{2\sigma^2}}.$$

In particular $E[b] < 0$ whenever $h > l$.

Proof. Condition on $\bar{x}_1$. The m second-period draws are independent of $\bar{x}_1$ with mean $u = 0$, so $E[S | \bar{x}_1] = k \bar{x}_1$, and the denominator $k + m + a$ is determined by $\bar{x}_1$. Hence $E[b | \bar{x}_1] = k \bar{x}_1 / (k + m(\bar{x}_1) + a)$, and taking expectations over $\bar{x}_1 \sim N(0, \sigma^2/k)$,

$$E[b] = \frac{k}{k+h+a} E[\bar{x}_1 \mathbb{1}\{\bar{x}_1 > c\}] + \frac{k}{k+l+a} E[\bar{x}_1 \mathbb{1}\{\bar{x}_1 \leq c\}].$$

For $X \sim N(0, s^2)$ with $s = \sigma / \sqrt{k}$ one has $E[X \mathbb{1}\{X > c\}] = s \varphi(c/s)$ and $E[X \mathbb{1}\{X \leq c\}] = -s \varphi(c/s)$, where φ is the standard normal density. Substituting and using $ks = \sigma\sqrt{k}$,

$$E[b] = \sigma\sqrt{k}\varphi(c/s) \left[\frac{1}{k+h+a} - \frac{1}{k+l+a} \right] = -\sigma\sqrt{k}\varphi(c/s) \frac{h-l}{(k+h+a)(k+l+a)}.$$

Finally $\varphi(c/s) = (2\pi)^{-1/2} e^{-(c^2/(2s^2))} = (2\pi)^{-1/2} e^{-(c^2 k/(2\sigma^2))}$, which gives the stated result.

Corollary 1 (Averaging learner). For $a = 0$ and $c = 0$, $E[b] = -\sigma\sqrt{k}(h-l) / (\sqrt{2\pi}(k+h)(k+l))$. With $k = 2$, $h = 10$, $l = 1$ this equals -0.141σ .

Corollary 2 (Bayesian learner). For $a = \sigma^2 / \tau_0^2$ and $c = 0$, $E[b] = -\sigma\sqrt{k}(h-l) / (\sqrt{2\pi}(k+h+\sigma^2/\tau_0^2)(k+l+\sigma^2/\tau_0^2))$. With $\tau_0^2 = 1$ this equals -0.098 at $\sigma = 1$ and -0.025 at $\sigma = 5$.

Remark 1 (Monotone attenuation). For $h > l$, $|E[b]|$ is strictly positive and strictly decreasing in a , since the denominator $(k+h+a)(k+l+a)$ increases in a while the numerator is constant. Rational shrinkage toward a correctly centred prior therefore reduces the mean bias but never eliminates it.

The second object of interest is the share of learners who end below the true value. Proposition 1 shows that, unlike the mean, this share does not depend on the learning rule.

Proposition 1. For $c = u = 0$, $\text{sign}(b) = \text{sign}(S)$ for every $a \geq 0$. Consequently $P(b < u) = P(S < 0)$ is identical across learning rules, and because $S = \sigma \cdot W$ where W depends only on standardised draws and the sign of $\bar{x}_1$, it is invariant to σ .

The complete-avoidance limit. The classic hot stove is total: the burnt cat takes no further draws at all, whereas $l = l$ keeps a thin corrective channel open. Setting $l = 0$ in Theorem 1 gives $E[b] = -\sigma \sqrt{(k) h / (\sqrt{(2 \pi)(k + h + a)(k + a)})}$, which at the baseline is -0.235 for the averaging learner and -0.145 for the Bayesian learner at $\sigma = 1$, against -0.141 and -0.098 when a single corrective draw survives: increases of 67 and 48 per cent. The underestimation share rises from 0.585 to 0.683 , so the invariance of Proposition 1 is a property of the interior case and not of the limit. The introduction motivates the model with a market in which a passed deal generates no feedback whatever; that market is the $l = 0$ limit, and the baseline $l = l$ is the more conservative specification. The estimates reported below therefore understate the bias in markets where rejections are never revisited.

RESULTS

The statistics below are derived from 10^7 simulated agents and match the analytical results reported under *Analytical Results* above.

Expected Bias

When the second-period sample size increases with the initial signal, the expected belief is negative for both learners. In the baseline ($u = 0$, $\sigma = 1$) the averaging learner settles at $E[b] = -0.141$ and the Bayesian learner at $E[b] = -0.098$, both confirming Corollaries 1–2 to within Monte Carlo error.

The mechanism is a weighting asymmetry. A positive initial signal triggers ten further draws; if these regress toward the mean of zero, the initial positive error is diluted by the large denominator $k + h + a$. A negative initial signal triggers only one further draw, so the initial negative error is preserved by the small denominator $k + l + a$. The asymmetry in the denominator therefore corrects optimism more strongly than pessimism.

A Two-Agent Illustration

Consider two agents exploring a vertical with true value 0 .

Agent A (a false positive). The first two draws are favourable $(+2, +2)$, giving $\bar{x}_1 = +2$. The agent commits high capital ($m = 10$). Over the ten additional draws the results average to approximately zero, so the final averaging belief is $(2 + 2 + 0 \times 10) / 12 \approx 0.33$. The initial optimism is corrected downward.

Agent B (a false negative). The first two draws are unfavourable $(-2, -2)$, giving $\bar{x}_1 = -2$. The agent commits low capital ($m = 1$) and takes one further draw, which is approximately zero. The final averaging belief is $(-2 - 2 + 0) / 3 \approx -1.33$. The initial pessimism is barely corrected.

The asymmetry. Agent A moves from +2 to +0.33 (strong correction); Agent B moves from -2 to -1.33 (weak correction). Agent B remains pessimistic because it denied itself the sample size needed to dilute a poor start. When many such agents are aggregated, the under-corrected pessimists dominate the lower tail and pull the population mean below zero.

Sensitivity to Volatility

The dependence of the bias on volatility differs sharply between the two learners, and it is not the simple monotone amplification suggested by an informal reading of the mechanism.

For the averaging learner, Corollary 1 gives a bias that is *exactly linear* in σ : it scales from -0.141 at $\sigma = 1$ to -0.705 at $\sigma = 5$, a factor of exactly five. Within the model the averaging bias therefore grows linearly in σ without bound; under heavy-tailed payoffs the effect persists, though the linear scaling need not (see *Heavy Tails and the Limits of the Gaussian Model*).

For the Bayesian learner the dependence is non-monotone. As σ grows, the shrinkage constant $a = \sigma^2 / \tau_{u0}^2$ grows quadratically, so the correctly centred prior increasingly dominates the noisy data and pulls beliefs back toward the truth. The magnitude first rises to a maximum of about 0.103 near $\sigma \approx 1.3$ and then decays (0.098 at $\sigma = 1$, 0.025 at $\sigma = 5$; Figure 1b). A rational learner with a fixed-scale prior is therefore better protected as volatility rises.

Sensitivity to the prior specification. The attenuation just described depends on holding the prior variance fixed at $\tau_{u0}^2 = 1$ while the payoff variance grows, so that $a = \sigma^2 / \tau_{u0}^2$ increases quadratically and a prior of fixed scale comes to dominate increasingly noisy data. The opposite specification is equally natural: if the analyst's uncertainty about a vertical scales with that vertical's volatility, so that $\tau_{u0}^2 = \sigma^2$, then $a = 1$ for every σ and Theorem 1 gives $E[b] = -\sigma \sqrt{k(h-1)} / (\sqrt{2\pi}(k+h+1)(k+l+1)) = -0.098 \sigma$, exactly linear in σ , as for the averaging learner. The non-monotonicity of Figure 1b is therefore a property of the fixed-scale prior and not of Bayesian updating as such. Which specification is appropriate is an empirical question about investment committees, and the proportional prior has the better claim: an analyst covering a moonshot vertical is plausibly more uncertain, not equally certain, than one covering an established vertical. We retain the fixed-scale case as the baseline because it is the specification most favourable to the Bayesian learner, and therefore the most conservative setting in which to claim that shrinkage does not remove the bias.

Persistence of Bias under Bayesian Updating

Bayesian updating does not remove the negativity bias. Corollary 2 and Remark 1 show that the posterior mean remains strictly negative whenever $h > l$: a rational learner that conditions on a fixed neutral vertical underestimates it in expectation, because the sample size is endogenous to the sign of the early signal.

The share of learners who underestimate is governed by a different quantity. By Proposition 1, $P(b < u) = P(S < 0)$ for both learners, so the underestimation share is identical across updating rules and invariant to σ . In the simulation it equals 0.585 in every cell of Table 1: a majority of learners, whether naive or Bayesian, end below the true value.

If the true value is itself drawn from the prior ($u \sim N(0, \tau u_0^2)$), the tower property gives $E[b - u] = 0$ for any data-dependent stopping rule, and the median is unbiased as well, $P(b < u) = 0.500$ (verified by simulation). The exact zero-mean result for Bayesian learners holds only under this prior-predictive experiment. It does not describe a fixed neutral vertical, for which both the mean and the majority remain below the truth.

Summary of Findings

Table 1: Summary statistics ($N = 10^7$; parameters $k = 2, h = 10, l = 1, c = 0, u = 0$).

Learner	(k, h, l)	σ	$E[b]$	Median	Skew	$P(b < 0)$
Averaging	(2, 10, 1)	1	-0.141	-0.083	-0.65	0.585
Averaging	(2, 10, 1)	5	-0.705	-0.417	-0.65	0.585
Bayesian	(2, 10, 1)	1	-0.098	-0.070	-0.42	0.585
Bayesian	(2, 10, 1)	5	-0.025	-0.075	+0.59	0.585

The $E[b]$ column reports the closed form of Corollaries 1–2 (-0.141, -0.705, -0.098, -0.025); medians and $P(b < 0)$ are simulation estimates ($N = 10^7$), which reproduce every $E[b]$ entry to within 2×10^{-3} . The Monte Carlo standard error of $P(b < 0)$ is about 1.6×10^{-4} ; the four values coincide up to this error, as required by Proposition 1.

That tolerance is a bound rather than an estimate of sampling error: the Monte Carlo standard error of the sample mean, $sd(b) / \sqrt{N}$, ranges from 1.1×10^{-4} to 6.9×10^{-4} across the four cells, and the largest observed deviation from the closed form is 7.6×10^{-4} . Skewness values are simulation estimates; the averaging figure does not vary with σ because the standardised belief distribution is independent of it.

The Distribution of Beliefs

Figure 1a shows the belief distributions at $\sigma = 1$. Both are negatively skewed (sample skewness -0.65 for the averaging learner and -0.42 for the Bayesian learner), with the mode near the true value and a heavy tail on the pessimistic side. The shape reflects the mechanism directly: learners with a favourable initial signal take ten further draws and their beliefs concentrate near the truth, forming the central peak; learners with an unfavourable initial signal take a single further draw and their beliefs scatter widely into pessimism, forming the left tail. At this volatility the mean is dragged below the median by this tail, and 58.5% of the mass lies below the true value. Figure 1b contrasts the exactly linear bias of the averaging learner with the attenuated, non-monotone bias of the Bayesian learner.

The direction of this asymmetry is not a robust feature of the model. For the averaging learner the standardised belief distribution does not depend on σ , so the sample skewness is -0.65 at every volatility. For the Bayesian learner it is not: the skewness is -0.42 at $\sigma = 1$, approximately zero near $\sigma \approx 2$, and +0.59 at $\sigma = 5$. The reversal follows from the fixed-scale prior that also produces the non-monotone mean bias reported under *Sensitivity to Volatility*. Because the shrinkage constant enters both denominators additively, it compresses the ratio $(k+h+a) / (k+l+a)$ from 3.25 at $\sigma = 1$ to 1.32 at $\sigma = 5$, progressively

removing the denominator asymmetry that stretches the pessimistic tail; the numerator asymmetry survives, since the high-commitment group accumulates $k + h = 12$ draws against $k + l = 3$. At high volatility the optimistic group is therefore the more dispersed of the two, and the heavy tail changes sides. The mean is then pulled above rather than below the median, as the $\sigma = 5$ row of Table 1 records (-0.025 against -0.075). The mode moves with them: computed from the exact mixture density it lies at $+0.019$ for the averaging learner and -0.015 for the Bayesian learner at $\sigma = 1$, but at -0.136 for the Bayesian learner at $\sigma = 5$. The ordering of the three statistics therefore inverts exactly where the skewness changes sign, from mean < median < mode in the three negatively skewed cells to mode < median < mean in the fourth. The mode is in no cell exactly at the true value; at $\sigma = 1$ it is displaced by less than five per cent of a standard deviation, which is why Figure 1a shows a peak that is visually indistinguishable from the truth. The underestimation share is untouched by the reversal and remains 0.585 , in accordance with Proposition 1: the shape of the belief distribution depends on the prior specification, while the majority-underestimation property does not.

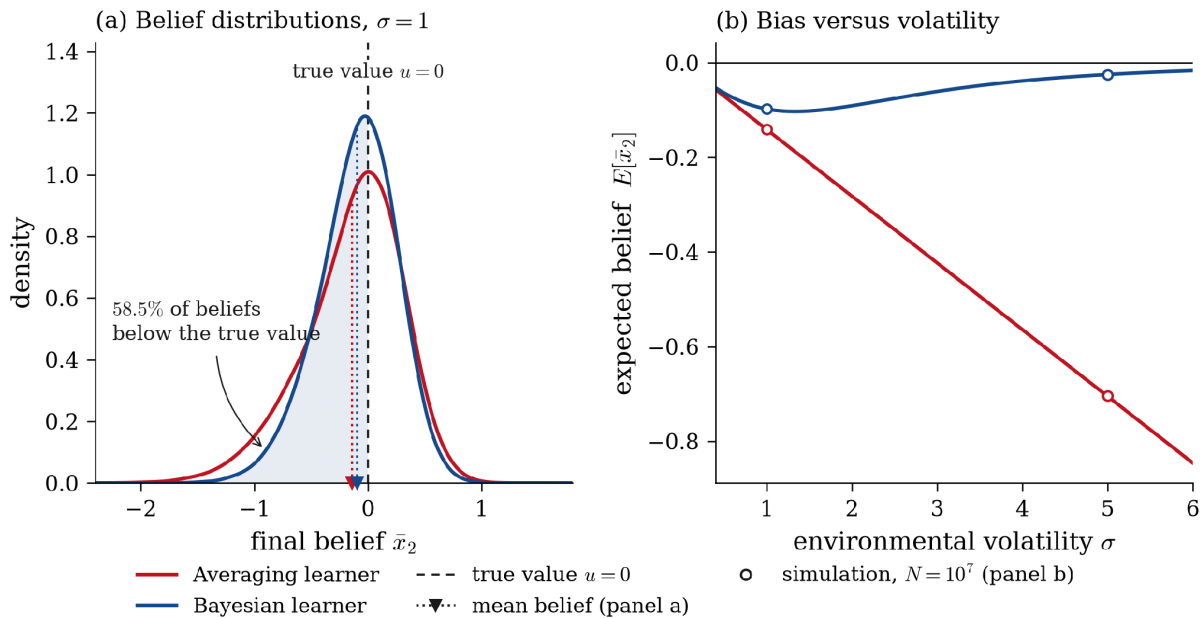
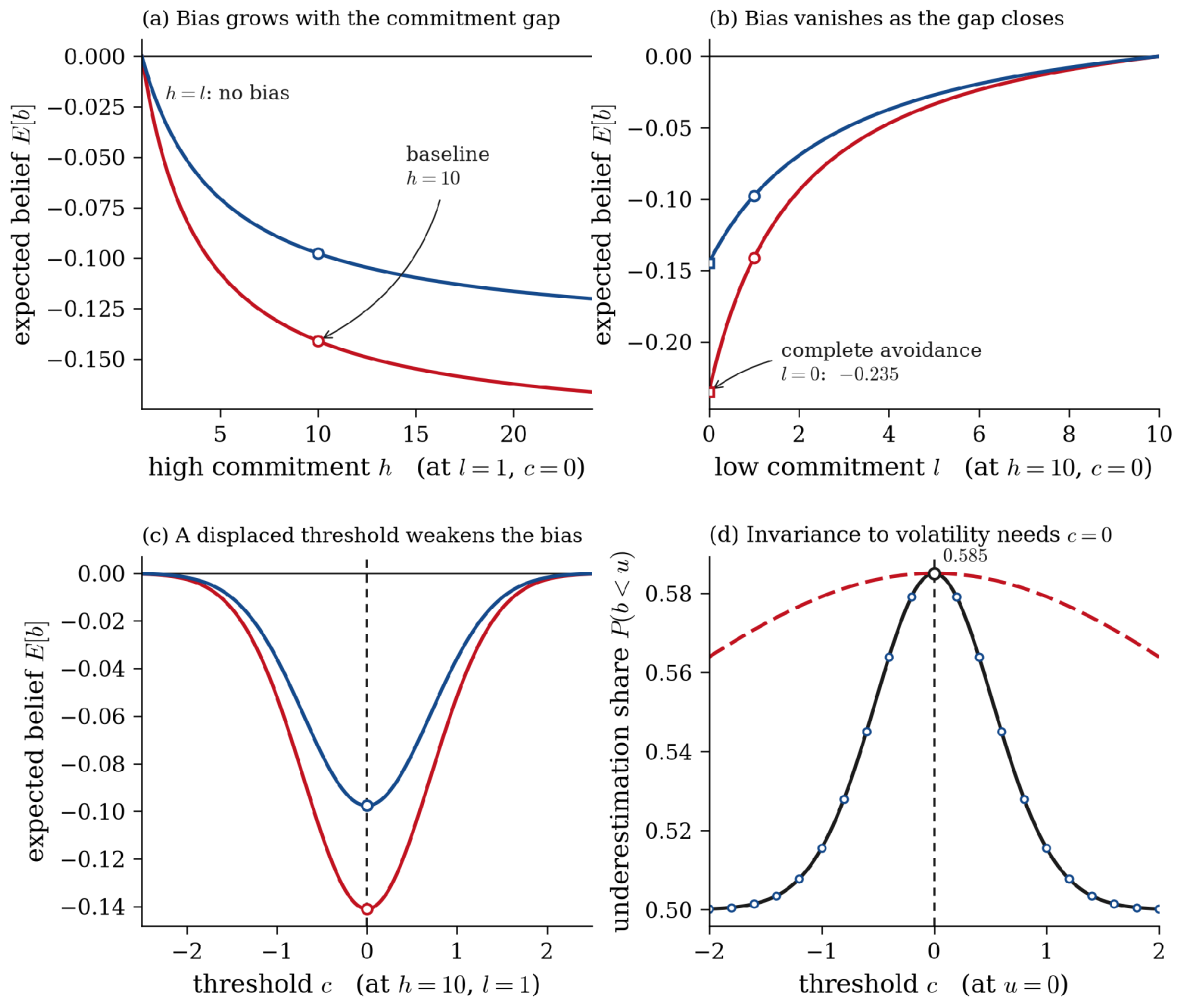


Figure 1: (a) Belief densities at $\sigma = 1$ for the averaging and Bayesian learners. Both peak near the true value $u = 0$ but carry a heavy pessimistic tail; 58.5% of beliefs fall below the truth for either learner. (b) Expected belief as a function of volatility σ . The averaging bias (Corollary 1) is exactly linear; the Bayesian bias (Corollary 2) is attenuated and non-monotone. Open circles are simulation estimates ($N = 10^7$). Panel (b) holds the prior variance fixed at $\tau u_0^2 = 1$; under a prior whose variance scales with σ^2 the Bayesian bias is exactly linear in σ (*Sensitivity to Volatility*).

Comparative Statics and the Scope of the Invariances

Comparative statics. Theorem 1 holds for arbitrary k , h , l , c and a , so the baseline parameters can be varied at no analytical cost. Figure 2 does this over the three policy parameters, and three features prove robust. The bias is strictly negative wherever $h > l$ and vanishes exactly at $h = l$, so the baseline $(k, h, l) = (2, 10, 1)$ is not a knife-edge but an interior point of a surface that is negative throughout. Its magnitude increases in the commitment gap $h - l$ and then saturates, so the qualitative result does not turn on the precise value of h . And it is largest at $c = 0$, decaying as $e^{-(c^2 k / (2 \sigma^2))}$ in either direction: a threshold displaced from the true mean weakens the effect rather than strengthening it, because it makes the commitment decision less sensitive to the sign of the early signal.



Panels (a)-(c)

— Averaging learner — Bayesian learner ○ baseline $(k, h, l) = (2, 10, 1)$ □ complete avoidance, $l = 0$

Panel (d)

— $\sigma = 1$ - - $\sigma = 5$ ○ Bayesian learner (coincides with $\sigma = 1$)

Figure 2: Comparative statics of the expected belief at $\sigma = 1$. (a) $E[b]$ against the high commitment h at $l = 1$, $c = 0$; the bias vanishes at $h = 1$. (b) against the low commitment l at $h = 10$; the square marks the complete-avoidance limit $l = 0$. (c) against the threshold c at $(h, l) = (10, 1)$; the bias is largest at $c = 0$ and decays in either direction. (d) the underestimation share $P(b < u)$ against c at $u = 0$: the two learning rules coincide at every threshold, whereas the two volatilities agree only at $c = 0$. Open circles mark the baseline $(k, h, l) = (2, 10, 1)$. All values are closed-form or exact quadrature; no simulation is involved.

Two invariances, two conditions. Proposition 1 states two invariances, and they do not require the same conditions. The identity across learning rules follows from $\text{sign}(b) = \text{sign}(S)$, which no threshold can disturb, and it therefore holds for any c : at $c = 0.3$ the share falls to 0.572 but remains identical for the averaging learner and for a Bayesian learner at any degree of shrinkage. The invariance to volatility is the weaker of the two and requires $c = u = 0$: at $c = 0.3$ the share is 0.572 at $\sigma = 1$ against 0.585 at $\sigma = 5$. Panel (d) of Figure 2 displays both facts. A reader should not generalise the volatility invariance beyond a threshold placed at the true mean; the identity across learning rules, on which the contribution of this paper rests, survives.

DISCUSSION

The results give a computational and analytical basis for the asymmetry in error correction that characterises VC decision-making. The Hot Stove Effect is a structural byproduct of adaptive learning and payoff-maximising behaviour.

Mean Bias versus Majority Underestimation

Two distinct facts should be kept separate. The first is the mean bias: the expected belief is negative for both learners (Theorem 1), and it is only attenuated, not removed, by Bayesian updating (Remark 1). The second is the majority-underestimation property: a majority of learners end below the truth (Proposition 1), and this share is the same for naive and Bayesian learners and does not vary with volatility.

These facts have separate causes. The mean bias arises from the dilution asymmetry in the denominator $k + m + a$. The majority property arises from the sign of the accumulated evidence, which the updating rule cannot change. Bayesian shrinkage can make the mean bias small, as little as -0.025 at $\sigma = 5$, while the median remains below the truth and a majority of firms (58.5%) still underestimate the same neutral vertical. Rationality repairs part of the mean bias but leaves the majority below the truth.

Heavy Tails and the Limits of the Gaussian Model

Empirical VC returns are heavy-tailed and better described by a power law than by a Normal distribution. Two points follow. First, the negativity bias is not an artefact of Gaussian payoffs: replacing the Normal with a standardised heavy-tailed Student- t distribution leaves the averaging-learner bias negative (about -0.12 under t_3 and -0.13 under t_5 at unit variance) and slightly raises the underestimation share (about 0.59 under t_3 versus 0.585 under the Normal). Second, the welfare cost of the bias under a genuine power law,

where the upside is carried by rare extreme outcomes, is plausibly larger than under the Normal, because ceasing to sample a high-variance sector after early failures can forgo a rare outcome that would have justified all prior losses. We treat this welfare claim as a conjecture: it is not established by the present model and is left to future work, since a two-sided power law with a neutral mean requires a distinct analytical treatment.

Practitioner accounts describe the outcome distribution in these terms directly: of roughly twenty early-stage investments, most are written off entirely, three or four return the original capital with a modest premium, and one produces a tenfold or hundredfold return that covers the remainder (Strebulaev & Dang, 2024, p. 6). A distribution of that shape is precisely the case in which ceasing to sample a vertical after early failures is most costly, since the forgone outcome that would have justified the losses is the one in the extreme tail.

Consistency with Other Studies

The findings are consistent with experimental and field evidence. First, on *forced experimentation*, the model provides a mechanism for the finding of Larcom, Rauch, and Willems (2017) that a strike-induced episode of forced experimentation revealed superior routes that commuters had avoided. Second, the *BeanFest experiment*: the results mirror Fazio, Eiser, and Shook (2004), whose participants formed more negative than positive attitudes toward novel stimuli because negatively valenced objects were sampled less and their true value was never learned. Third, on *asymmetric feedback and cynicism*, the negatively skewed belief distribution obtained at low volatility aligns with Fetchenhauer and Dunning (2010), who find that individuals underestimate trustworthiness because avoidance suppresses the corrective feedback.

Implications for Limited Partners

The results bear on limited partners (LPs) who allocate capital to VC funds. Standard due diligence rewards funds that display discipline and low loss rates. Survey evidence indicates that venture capitalists rate deal selection as the most important of the three sources of value they control, ahead of sourcing and post-investment support (Gompers, Gornall, Kaplan, & Strebulaev, 2020), which places a premium on a visible record of avoided losses. The model suggests this can be counterproductive: a fund with a very low loss rate may be exhibiting the Hot Stove Effect, avoiding high-variance sectors after past losses. LPs might therefore incentivise *anti-portfolio analysis*, requiring general partners to track and report the performance of companies they passed on, which makes the foregone sampling visible; and they might *fund forced experimentation*, allocating a small share of the fund to contrarian bets that deliberately violate the firm's current heuristics.

This recommendation has a precedent in the model and a precedent in practice. Denrell (2007) shows that information about foregone payoffs increases risk taking, precisely because it restores the corrective feedback that avoidance suppresses. Bessemer Venture Partners maintains a public anti-portfolio of companies it declined, including Apple, Google, FedEx, PayPal, Airbnb, and Zoom, on the reasoning that a firm which cannot confess its mistakes cannot correct them (Strebulaev & Dang, 2024, pp. 8–9). The complete-avoidance limit derived in the Methods indicates where such disclosure matters most: as the

corrective channel closes, the underestimation share rises from 58.5 to 68.3 per cent, so the case for the recommendation rests on the limit more strongly than on the baseline.

Limitations and Future Research

Several limitations qualify the results. The model assumes learners act in isolation, whereas VCs observe the outcomes of rivals. Whether social exposure corrects the bias or deepens it is not an open question: Le Mens, Denrell, Kovács, and Karaman (2019) show that the sign of the effect turns on whom the additional exposure reaches, correcting the bias when it disproportionately reaches learners whose current estimate is low and deepening it when it concentrates on those whose estimate is already high (their Theorem 1, p. 361). The venture analogue is testable: syndicate structures that route deal flow towards partners who have already passed on a vertical should attenuate the bias, whereas co-investment that concentrates attention on deals already attracting interest should amplify it. The horizon is two periods and the policy is binary ($m \in \{l, h\}$); a continuous action space would model not only whether to invest but how much. The results are stated for a prior centred at the true value, which is the most favourable case for the Bayesian learner; a misspecified prior would shift the bias further, and its sensitivity should be characterised. Finally, the welfare consequences under heavy-tailed payoffs remain a conjecture.

CONCLUSION

Adaptive sampling policies, under which future investment and information gathering increase with initial impressions, produce a negativity bias known as the Hot Stove Effect. Errors of overestimation are corrected through subsequent trials, while errors of underestimation reduce future sampling and so persist. We formalise this in a closed-form expression for the expected belief of a one-parameter family of learners and show that the bias is strictly negative whenever a positive signal induces more sampling than a negative one.

Two conclusions follow. First, Bayesian updating with a correctly centred prior attenuates the mean bias monotonically in the shrinkage constant but does not remove it, because the sample size is endogenous to the sign of the early signal. Second, and independently, a majority of learners underestimate a neutral vertical regardless of the updating rule, because the underestimation share is fixed by the sign of the accumulated evidence. What may appear as risk aversion or a psychological defect in VC is, in part, a structural consequence of payoff-contingent sampling.

The practical implication is that encouraging risk-taking is insufficient, because the bias is structural rather than dispositional. Organisations can counter it by institutionalising forced experimentation or by using algorithmic screening in early deal flow, which maintains a sampling rate that does not fall in response to early negative feedback. Decoupling the decision to sample from the immediate outcome of the previous trial is the operative correction.

REFERENCES

Aug 2026
Vol 10, No 4.

Oxford Journal of Student Scholarship
www.oxfordjss.org

- Cyert, R. M., & March, J. G. (1963). *A Behavioral Theory of the Firm*. Englewood Cliffs, NJ: Prentice-Hall.
- Denrell, J. (2005). Why most people disapprove of me: Experience sampling in impression formation. *Psychological Review*, 112(4), 951–978.
- Denrell, J. (2007). Adaptive learning and risk taking. *Psychological Review*, 114(1), 177–187.
- Denrell, J., & Le Mens, G. (2011). Rational learning and information sampling: On the "naivety" assumption in sampling explanations of judgment biases. *Psychological Review*, 118(2), 379–392.
- Denrell, J., & March, J. G. (2001). Adaptation as information restriction: The hot stove effect. *Organization Science*, 12(5), 523–538.
- Fazio, R. H., Eiser, J. R., & Shook, N. J. (2004). Attitude formation through exploration: Valence asymmetries. *Journal of Personality and Social Psychology*, 87(3), 293–311.
- Fetchenhauer, D., & Dunning, D. (2010). Why so cynical? Asymmetric feedback underlies misguided skepticism regarding the trustworthiness of others. *Psychological Science*, 21(2), 189–193.
- Gompers, P. A., Gornall, W., Kaplan, S. N., & Strebulaev, I. A. (2020). How do venture capitalists make decisions? *Journal of Financial Economics*, 135(1), 169–190.
- Harrison, J. R., & March, J. G. (1984). Decision making and postdecision surprises. *Administrative Science Quarterly*, 29(1), 26–42.
- Kahneman, D., & Tversky, A. (1979). Prospect theory: An analysis of decision under risk. *Econometrica*, 47(2), 263–291.
- Larcom, S., Rauch, F., & Willems, T. (2017). The benefits of forced experimentation: Striking evidence from the London Underground network. *The Quarterly Journal of Economics*, 132(4), 2019–2055.
- Le Mens, G., Denrell, J., Kovács, B., & Karaman, H. (2019). Information sampling, judgment, and the environment: Application to the effect of popularity on evaluations. *Topics in Cognitive Science*, 11(2), 358–373.
- Levinthal, D. A., & March, J. G. (1993). The myopia of learning. *Strategic Management Journal*, 14(S2), 95–112.
- March, J. G. (1991). Exploration and exploitation in organizational learning. *Organization Science*, 2(1), 71–87.
- Simon, H. A. (1955). A behavioral model of rational choice. *The Quarterly Journal of Economics*, 69(1), 99–118.
- Strebulaev, I. A., & Dang, A. (2024). *The Venture Mindset: How to Make Smarter Bets and Achieve Extraordinary Growth*. New York: Portfolio/Penguin.

Twain, M. (1897). *Following the Equator: A Journey Around the World*. Hartford, CT: American Publishing Company.