



Research article

Distributionally robust automation under overdispersed count risk: Ambiguity sets, tail protection, and oracle-relative policy gap

Jinho Cha^{1,*}, Long Pham², Minh Ngo Thi³, Ho Thi Hien⁴ and Hai Thi Bich Vu⁵

¹ Department of Computer Science, Gwinnett Technical College, Georgia, USA

² Department of Decision Sciences and Economics, College of Business, Texas A&M University–Corpus Christi, Texas, USA

³ Department of Finance and Accounting, East Asia University of Technology, Bac Ninh, Vietnam

⁴ Department of Economics and Business Administration, Nghe An University, Nghe An Province, Vietnam

⁵ School of Economics and International Business, Foreign Trade University, Hanoi, Vietnam

* **Correspondence:** Email: jcha@gwinnettech.edu.

Abstract: Automation under clustered count risk can be fragile when a fitted model understates dispersion or upper-tail exposure. We developed a distributionally robust sequential automation framework in which a protective intervention is selected based on fitted or ambiguity-adjusted exposure and the previously implemented level. The control may represent hedging, monitoring, preventive maintenance, or automated incident response. Convex implementation and adjustment costs produce history-dependent decisions without a Bellman recursion or specified transition kernel. A local ambiguity set perturbs the fitted negative binomial (NB) mean and dispersion. We established existence and uniqueness, monotonicity, an activation condition, and adjustment-induced policy persistence. Balanced experiments include favorable, neutral, adverse, and reverse misspecification, revealing protection benefits and over-intervention costs of robustness. Using 437 monthly Standard & Poor’s 500 (S&P 500) jump counts, the out-of-sample illustration shows a variance-to-mean ratio of 4.811 and that the NB model fits substantially better than the Poisson benchmark. Robust policies reduce tail loss by 1.90%–17.11%; total-cost estimates also improve, although bootstrap intervals include zero. The financial application illustrates the framework, while the theory extends to other clustered count-risk settings.

Keywords: distributionally robust optimization; negative binomial (NB) model; overdispersed count risk; robust automation; ambiguity calibration

Mathematics Subject Classification: 90C15, 90C17, 90C39, 60E05, 62P20

1. Introduction

Automation increasingly governs environments in which risk arrives as discrete, clustered events rather than smooth fluctuations. System failures, cyber incidents, financial jumps, service disruptions, safety violations, and demand shocks all share this feature. In these settings, the decision-maker must choose a protective automation level before the realized count is fully known. Too little automation leaves the system exposed to rare but severe event clusters; too much creates unnecessary implementation, monitoring, and adjustment costs.

The central operational risk is therefore not overdispersion alone, but *decision fragility caused by acting on an imperfectly fitted overdispersed model*. Count variance may substantially exceed the mean, and extreme realizations may cluster across time, as documented in the literature on volatility persistence, jump concentration, and time-varying tail risk [1–3]. A comprehensive review of ARCH-based evidence in finance is provided by Bollerslev et al. [4]. Recent evidence on systemic fragility and early-warning signals further links extreme-risk clustering to persistent instability and cross-sectional amplification [5–7]. A statistically plausible count model may therefore still be operationally dangerous if it understates policy-relevant tail exposure.

The Negative Binomial (NB) model is a natural benchmark, as its Gamma–Poisson mixture structure captures excess variation while remaining parsimonious and interpretable [8–10]. Flexible long-tailed count models, including the Sichel family, show that tail shape may vary beyond what a single NB specification can represent [11–13]. These alternatives motivate tail-shape misspecification rather than serve as final policy benchmarks. The decision problem is thus sharper than model selection: *how should a protective automation rule be calibrated when the fitted count law may either understate or overstate tail-relevant exposure?*

This question connects count-data modeling, adjustment frictions, automated execution, and robust decision-making. Dynamic investment and switching models show how adjustment costs create delayed responses and history dependence [14, 15], while persistence and long-memory mechanisms strengthen the importance of gradual policy adjustment [16, 17]. In parallel, algorithmic execution, market microstructure, blockchain systems, and digital platforms demonstrate how observed information is converted into event-contingent action [18–20]. These streams motivate automation as a sequential protective control rather than a one-time adoption choice.

Robust optimization and distributionally robust optimization (DRO) provide the natural tools for this problem [21–23]. Work on moment and dispersion ambiguity is particularly relevant because the automation decision depends on the tail-sensitive risk-damage functional induced by the fitted count law [24, 25]. Related dynamic approaches include robust dynamic programming and distributionally robust Markov decision processes [26–28]. The proposed framework robustifies the current-period count-risk law while linking consecutive actions through adjustment cost; it does not require a Bellman recursion or ambiguity over a transition kernel.

The closest predecessor is Cha et al. [29], which establishes that known overdispersed NB risk can generate automation thresholds and hysteresis. The central question is different: What happens when the fitted NB law itself is uncertain? The fitted law is treated as a nominal model, and a local ambiguity set perturbs its mean and dispersion. Automation intensity is then chosen to balance worst-case residual loss, implementation cost, and adjustment friction. Threshold shifts and policy persistence arise as consequences of ambiguity rather than as the primary object of the model.

This sequential formulation is intentionally lighter than a fully specified dynamic program. The previous action carries operational memory through the adjustment term, while the current fitted or ambiguity-adjusted exposure drives the new decision. It therefore captures gradual adjustment, delayed reversal, and policy persistence without requiring the analyst to estimate a complete state-transition model. This feature is useful in applications where count-risk parameters can be updated reliably but long-horizon transition dynamics are difficult to identify.

The framework has three ingredients. First, a fitted NB law summarizes clustered count risk. Second, a continuous automation action reduces residual exposure but is costly to implement and adjust. Third, the fitted law is not treated as truth: The decision maker protects against plausible local misspecification. Depending on the application, the action can represent hedging intensity, exposure reduction, liquidity or margin protection, monitoring escalation, preventive maintenance, or automated incident response. This interpretation makes the model relevant to both financial and industrial systems in which protective actions must be updated repeatedly as risk estimates change.

Three contributions follow. First, the framework formulates a distributionally robust sequential automation problem centered on a fitted NB count law. Second, the analysis establishes existence and uniqueness, monotonicity in ambiguity-adjusted exposure, an explicit activation condition, and adjustment-induced policy persistence. Third, the computational design evaluates the downstream consequences of model error through total cost, tail loss, automation intensity, switching, under- and over-intervention, and oracle-relative policy gaps. These metrics expose both sides of robustness: protection when the nominal model understates risk and avoidable intervention when it overstates risk.

The computational design is deliberately balanced. It includes favorable, neutral, adverse, and reverse misspecification environments rather than only stress scenarios aligned with the ambiguity set. The comparisons cover Poisson plug-in, NB plug-in, robust NB, and oracle policies across ambiguity radii, dispersion and long-tail perturbations, adjustment costs, and regime patterns. The results show that robustness is not monotonically desirable: Moderate ambiguity can reduce under-intervention and tail loss, whereas excessive ambiguity can create costly over-intervention. Accordingly, the ambiguity radius should be calibrated to the direction and severity of model error rather than increased mechanically.

The balanced scenarios also answer an important design concern. Because a one-sided ambiguity set can appear favorable when the data-generating process is perturbed in the same direction, the experiment includes environments in which the nominal model is accurate or already conservative. These cases reveal the price of robustness directly through higher implementation cost, over-automation, and policy gaps. The comparison therefore distinguishes genuine protection from an advantage created mechanically by the stress design.

An out-of-sample financial illustration uses 437 monthly S&P 500 jump counts. The sample variance-to-mean ratio is 4.811, and the NB specification decisively improves on the Poisson benchmark. Across the tested ambiguity radii, robust policies reduce tail loss by 1.90%–17.11%. Total-cost point estimates also improve, although the moving-block-bootstrap intervals include zero. The empirical analysis is therefore interpreted as a disciplined financial application rather than universal external validation. Its broader role is to demonstrate how the theoretical framework can be transferred to reliability, insurance, cyber-risk, service-disruption, safety-event, and other clustered count-risk settings.

Recent studies in *Journal of Industrial and Management Optimization* have examined irreversible investment, green technology adoption, and unreliable production and inventory systems [30–32]. Related work has considered regime-switching financial games and digital transformation decisions [33, 34]. The

contribution to this stream lies in linking overdispersed count modeling, ambiguity-aware optimization, and sequential protective automation in a unified decision-oriented framework.

The remainder is organized as follows: Section 2 reviews the related literature. Sections 3 and 4 introduce the baseline automation model and count-risk specifications. Section 5 formulates the robust problem, and Section 6 develops the structural results. Sections 7 and 8 present the computational and empirical analyses. Section 9 concludes with implications and limitations.

2. Literature review

2.1. Clustered events and overdispersed count risk

The starting point for this study is that many operational hazards arrive as event counts, while their intensity and dispersion vary over time. Financial econometrics documents the broader mechanism clearly: Volatility is persistent, jumps cluster in episodes, and tail compensation changes across market states [3, 35, 36]. Recent work on systemic fragility and early-warning signals further links concentrated extreme events to persistent instability and cross-sectional amplification [5–7]. These findings motivate a count-risk representation in which rare events are neither independent nor adequately summarized by a constant Poisson intensity.

Count-data research has long treated excess variation as a structural feature rather than a residual anomaly. When the variance of observed counts exceeds the mean, the equidispersion restriction of the Poisson model is often rejected, and mixed-Poisson models provide a natural alternative [8–10]. A broader overview of overdispersion is provided by Dean and Lundy [37]. The NB model is especially useful because its Gamma–Poisson representation separates mean exposure from dispersion in a parsimonious and interpretable way. Related actuarial and applied-count studies use such models for claims, accidents, health events, and heterogeneous operational exposure [13, 38, 39].

Recent applied work has also developed alternative one-parameter, compound, mixed-Poisson, and related flexible distributional models for heterogeneous risk data in radiation monitoring, public health, and agriculture [40–42]. Additional applications cover cytogenetics, biology, and clinical remission analysis [43–45]. Statistical-process-control applications provide another extension [46]. These studies broaden the available statistical toolkit for heterogeneous count, monitoring, and reliability-related data. Our focus is the downstream sequential decision consequences of ambiguity around a fitted NB law rather than by proposing a new count distribution.

Overdispersion, however, does not uniquely determine upper-tail behavior. Flexible mixed-Poisson families, including the Sichel distribution, can produce longer and differently shaped tails than a standard NB law [11–13]. Crash-count evidence further shows that dispersion accuracy matters for long-tail behavior [39]. This distinction matters for decisions because two models with similar central fit may imply different probabilities of costly high-count events. Accordingly, the present study uses the NB model as the main policy benchmark while treating long-tailed alternatives as evidence that the fitted tail structure may be misspecified.

The relevant gap is therefore not the absence of statistical models for overdispersion, but the limited understanding of how estimation error in an overdispersed count law propagates into a protective operational decision. Statistical adequacy and decision adequacy need not coincide: A model may fit average counts well yet produce a poor action when the loss function places special weight on the upper tail.

2.2. *Event-contingent automation and digital execution*

Future work explains why count risk naturally enters automated decision systems. In algorithmic trading and market microstructure, execution rules convert observed order flow, inventory, and price-impact signals into actions [18, 19, 47]. Related studies examine high-frequency trading, toxic arbitrage, and liquidity effects [48–50]. The speed–sophistication trade-off provides an additional perspective [51]. In blockchain and decentralized systems, verifiable states trigger programmable transfers, settlement, and contract execution [20, 52, 53]. Further research addresses decentralized finance, token valuation, and digital economics [54–56]. Smart-contract-based financial markets provide a related example [57]. Across these settings, action is commonly event-contingent rather than a smooth response to infinitesimal shocks.

This execution logic creates a close match between discrete-risk models and automation. A burst of failures, cyber alerts, claims, or financial jumps changes the operational state in a way that can trigger stronger monitoring, hedging, liquidity protection, maintenance, or incident response. A purely diffusion-based representation may capture continuous variability while obscuring the number and clustering of actionable events. Count models instead describe the arrival structure that automated systems must process.

The literature has examined many execution mechanisms, but usually takes the underlying risk representation as given. Attention shifts to the interface between estimation and action: An automated rule may be internally consistent with its fitted count model and still be poorly calibrated if that model understates or overstates tail-relevant exposure. This interface is where distributional uncertainty becomes operationally consequential.

2.3. *Sequential adjustment and policy persistence*

Protective automation is also rarely chosen once and held fixed forever. Decision-makers revise controls as new event data arrive, but changing an implemented action is itself costly. The investment and switching literature shows that adjustment frictions generate delayed responses, inaction, and history dependence [14, 15]. Persistence in financial and operational states strengthens these effects because recent conditions remain informative about near-term risk [16, 17, 58]. Regime-switching evidence provides a complementary account of persistent state dependence [59, 60].

Cha et al. [29] applied this logic to known NB risk and showed that overdispersion can alter automation thresholds and hysteresis. That result provides the structural benchmark for the present study. Our question begins one step later: If the count law used to construct the action is estimated rather than known, how should the sequential rule respond to uncertainty around that estimate?

The formulation is deliberately parsimonious. The current action depends on the current fitted or ambiguity-adjusted risk exposure and on the previous action through a convex adjustment cost. This captures gradual adaptation, delayed reversal, and policy persistence without requiring an infinite-horizon Bellman equation or a fully estimated transition kernel. The distinction is important: The model is sequential, but its ambiguity concerns the current count-risk law rather than the transition dynamics of a Markov decision process.

2.4. Robust optimization and decision-oriented model evaluation

Robust control and robust optimization provide the conceptual foundation for protecting decisions against model error [21,61,62]. Distributionally robust optimization formalizes this principle by optimizing over a set of plausible distributions around a nominal or estimated law [22, 23, 63]. A comprehensive review is given by Rahimian and Mehrotra [64]. Moment-, dispersion-, and tail-based ambiguity sets are particularly relevant when limited data identify some distributional features more reliably than others [24, 25, 65]. Stochastic-dominance constraints offer another data-driven formulation [66].

Dynamic extensions include robust dynamic programming and distributionally robust Markov decision processes [26–28]. Average- and Blackwell-optimal extensions have also been studied [67]. Recent work has developed moment-based, optimal-transport, Sinkhorn, and multistage data-driven ambiguity models [68–70]. Multistage performance is examined separately by Zhang and Sun [71]. These approaches establish that ambiguity-set design should reflect the information actually supported by the data.

Our framework combines this robust perspective with decision-oriented model evaluation. Predict-then-optimize research emphasizes that a forecast or fitted distribution should be judged by the downstream decision it induces, not only by likelihood or predictive error [72]. Policies are therefore compared through realized cost, tail loss, adjustment behavior, under- and over-intervention, and oracle-relative gaps. This metric set is essential because robustness has two possible consequences: It can correct harmful under-reaction when the nominal law is optimistic, but it can also create expensive over-reaction when the nominal law is already accurate or conservative.

The literature thus leaves a focused gap. Existing count models describe overdispersion and long tails; execution studies explain event-triggered action; adjustment models explain persistence; and robust optimization protects against distributional error. What remains missing is a unified framework that traces local uncertainty in a fitted overdispersed count law through a sequential automation decision and evaluates both its protection benefit and its over-intervention cost. The next section develops that framework.

3. Model setup

In this section, we introduce the baseline automation environment and define the notation used throughout the theoretical, computational, and empirical analyses. The formulation follows the NB automation logic of Cha et al. [29], but assigns a different role to the fitted count law. In Cha et al. [29], the NB risk law provides the structural basis for threshold and hysteresis behavior. Here, the fitted count law is not treated as the true decision environment, but as a nominal model around which an ambiguity set is constructed for robust policy evaluation.

The modeling approach follows a decision-oriented optimization perspective: The fitted count distribution is evaluated not only by statistical fit, but by the automation decision and policy cost it induces. This perspective is related to predict-then-optimize models, where estimation quality is assessed through downstream optimization performance [72]. Table 1 summarizes the main notation used throughout the model, structural analysis, computational experiments, and empirical illustration.

The notation table is intentionally comprehensive so that the distinction between primitive random quantities, fitted objects, policy controls, ambiguity-set parameters, and evaluation metrics remains explicit throughout the manuscript. Before presenting the structural results, the modeling assumptions are separated from computational specifications. Compactness of the action set ensures existence of a

period decision; convexity of residual exposure, implementation cost, and adjustment friction yields a well-behaved one-dimensional optimization problem; strict convexity gives uniqueness; and differentiability is used only when first-order conditions, activation thresholds, or comparative-static formulas are invoked. These assumptions are also operationally interpretable: Automation reduces residual exposure, higher automation is increasingly costly, and rapid changes in implemented automation create adjustment friction.

Table 1. Notations used in the distributionally robust automation framework.

Symbol	Type	Description
Panel A: Time, count risk, and distributional primitives		
t	Index	Decision period in the automation model.
m	Index	Calendar month in the empirical jump-count illustration.
X_t	Random variable	Nonnegative integer-valued risk count realized in period t .
X_m	Random variable	Monthly jump count in the empirical illustration.
$\mathbb{Z}_+$	Set	Set of nonnegative integers.
P_{θ_t}	Distribution	Count-risk distribution in period t with parameter vector θ_t .
P_t^*	Distribution	True data-generating count distribution in period t .
$\widehat{P}_t$	Distribution	Fitted count distribution used by the decision-maker in period t .
$\widehat{\theta}_t$	Parameter vector	Estimated parameter vector of the fitted count model.
$\mathcal{F}$	Set	Family of admissible count distributions.
μ	Parameter	Mean parameter of the Poisson or NB model.
r	Parameter	NB dispersion parameter. Smaller r implies stronger overdispersion.
(λ, χ, ψ)	Parameter vector	Parameters of flexible long-tailed count alternatives, such as the Sichel/Poisson–GIG model.
Λ	Latent variable	Latent Poisson intensity in mixed-Poisson representations.
VMR	Statistic	Variance-to-mean ratio used to diagnose overdispersion.
Panel B: Automation decision, loss, and cost functions		
a_t	Control	Automation intensity chosen in period t .
$\mathcal{A} = [0, 1]$	Set	Feasible set for automation intensity.
$q(x)$	Function	Primitive damage from risk count x .
τ	Parameter	Tail threshold in the damage function.
κ	Parameter	Tail-penalty coefficient in $q(x) = x + \kappa(x - \tau)_+$.
$(x - \tau)_+$	Function	Positive-part tail exceedance, defined as $\max\{x - \tau, 0\}$.
$R(a)$	Function	Residual-exposure function after automation intensity a .
ξ	Parameter	Automation effectiveness parameter in examples such as $R(a) = \exp(-\xi a)$.
$\ell(x, a)$	Function	Post-automation loss from count x under automation intensity a .
$L(a; P)$	Function	Expected residual risk loss under count law P .
$C(a)$	Function	Implementation cost of automation intensity a .
c_1, c_2	Parameters	Linear and quadratic implementation-cost coefficients.
$H(a_t, a_{t-1})$	Function	Adjustment friction from changing automation from a_{t-1} to a_t .
η	Parameter	Adjustment-cost coefficient.

Continued on next page

Symbol	Type	Description
Panel C: Ambiguity sets and robust optimization quantities		
$d(P, Q)$	Distance	Distributional or parameter-space discrepancy between count laws P and Q .
Δ_t	Scalar	Count-model misspecification level, $d(P_t^*, \widehat{P}_t)$.
ρ_t	Parameter	Ambiguity radius in period t .
$\mathcal{P}_t(\widehat{P}_t, \rho_t)$	Set	Ambiguity set of plausible count distributions around $\widehat{P}_t$.
$\mathcal{P}^{NB}(\widehat{\mu}, \widehat{r}, \rho)$	Set	Within-family NB ambiguity set.
$\mathcal{P}^{LT}(\widehat{P}, \rho)$	Set	Conceptual long-tail perturbation set used to motivate tail-shape misspecification.
$M_q(P)$	Functional	Risk-damage functional, $\mathbb{E}_P[q(X)]$.
$\overline{M}_q(\widehat{P}_t, \rho_t)$	Functional	Worst-case risk-damage functional over the ambiguity set.
$\Phi(a; \widehat{P}, \rho)$	Function	One-period robust objective before exploiting separability.
$\Psi(a; m, a_-)$	Function	Sequential period objective $R(a)m + C(a) + H(a, a_-)$.
K_q	Constant	Lipschitz constant linking distributional error to risk-damage error.
Panel D: Sequential policies and policy evaluation metrics		
m_t^π	Functional	Risk-damage exposure supplied to the period problem under policy π .
π^{PI}	Policy	Generic sequential plug-in policy based on a fitted count distribution.
π^P	Policy	Poisson plug-in sequential policy.
π^{NB}	Policy	NB plug-in sequential policy.
$\pi^{R,NB}$	Policy	Robust NB sequential policy.
π^*	Policy	Oracle-informed sequential policy using the true current-period count law.
τ^{PI}	Threshold	Plug-in activation threshold.
$\tau^R(\rho)$	Threshold	Robust activation threshold at ambiguity radius ρ .
$G_t(a_t, a_{t-1}; P_t^*)$	Function	Expected period cost under the true current-period law, including adjustment friction.
$\text{Cost}_T(\pi)$	Metric	Total realized cost of policy π over horizon T .
$\text{Gap}_T(\pi)$	Metric	Oracle-relative realized cost difference over horizon T .
$\text{TailLoss}_T(\pi)$	Metric	Tail-loss metric under policy π .
$\bar{a}(\pi)$	Metric	Average automation intensity under policy π .
$\text{Switch}_T(\pi)$	Metric	Frequency of meaningful automation changes.
$\text{UnderAuto}_T(\pi)$	Metric	Frequency below the oracle action by more than ε_a .
$\text{OverAuto}_T(\pi)$	Metric	Frequency above the oracle action by more than ε_a .
$\text{ImplCost}_T(\pi)$	Metric	Cumulative implementation cost under policy π .
$\text{AdjCost}_T(\pi)$	Metric	Cumulative adjustment cost under policy π .
ε_a	Parameter	Tolerance for meaningful automation differences.

Continued on next page

Symbol	Type	Description
Panel E: Empirical jump-count construction		
P_t	Data series	Daily adjusted closing level of the S&P 500 index.
r_t	Data series	Daily log return, $r_t = \log P_t - \log P_{t-1}$.
$ r_t $	Data series	Absolute daily log return.
q_α	Threshold	Empirical α -quantile of $ r_t $.
α	Parameter	Jump threshold quantile; baseline value is 0.97.
J_t	Indicator	Jump-day indicator, $\mathbf{1}\{ r_t > q_\alpha\}$.
w	Index	Rolling estimation window.
$\widehat{P}_w^P$	Distribution	Poisson distribution fitted in rolling window w .
$\widehat{P}_w^{NB}$	Distribution	NB distribution fitted in rolling window w .
$\widehat{M}_q^{(w)}$	Functional	Fitted tail-damage functional in rolling window w .
$\log \widehat{L}$	Statistic	Maximized log-likelihood of a fitted count model.
AIC	Statistic	Akaike information criterion.
BIC	Statistic	Bayesian information criterion.

Note: The automation intensity a_t is the control variable, while the fitted count law $\widehat{P}_t$ is treated as a nominal model for policy construction. Plug-in policies treat $\widehat{P}_t$ as correct, whereas the robust NB policy optimizes over an ambiguity set $\mathcal{P}_t(\widehat{P}_t, \rho_t)$. The oracle policy uses the true count law P_t^* only for benchmarking in simulation. Flexible long-tailed count alternatives motivate tail-shape misspecification, but Sichel plug-in and robust Sichel policies are not retained as main benchmarks in the final computational comparison. In the empirical illustration, the true count law is unobserved, so feasible policies are compared using realized cost, tail loss, average automation intensity, switching frequency, and cost decomposition. Oracle-relative quantities are used only in simulation, where the true current-period data-generating law is known to the experimenter.

3.1. Decision environment

Time is indexed by $t = 0, 1, 2, \dots$. At each time t , the decision-maker faces a nonnegative integer-valued risk count

$$X_t \in \mathbb{Z}_+, \quad (3.1)$$

representing the number of risk events realized during the period. Examples include operational incidents, system failures, cyber alerts, financial jump events, service disruptions, safety violations, or clustered demand shocks. The count distribution is denoted by

$$P_{\theta_t} \in \mathcal{F}, \quad (3.2)$$

where θ_t is a vector of distributional parameters, and $\mathcal{F}$ is a family of admissible count distributions.

The Poisson model is used as an equidispersed benchmark, and the NB model is used as the main overdispersed count law for policy construction. Flexible long-tailed count alternatives are discussed in Section 4 only as motivation for tail-shape misspecification. The final computational comparison does not retain Sichel plug-in or robust Sichel policies as main benchmarks.

The decision-maker chooses an automation intensity

$$a_t \in \mathcal{A} = [0, 1]. \quad (3.3)$$

The value $a_t = 0$ corresponds to no automation or minimal automated intervention, whereas $a_t = 1$ corresponds to full automation intensity. Intermediate values represent partial automation, such as increased monitoring frequency, stronger algorithmic screening, additional automated verification, or more aggressive event-triggered response.

3.2. Risk loss and automation effect

Let $q : \mathbb{Z}_+ \rightarrow \mathbb{R}_+$ denote the primitive damage from a realized risk count. The function $q(x)$ is assumed to be nondecreasing in x . To allow tail events to receive greater weight than ordinary count variation, the following specification is used:

$$q(x) = x + \kappa(x - \tau)_+, \quad \kappa \geq 0, \quad \tau \in \mathbb{Z}_+, \quad (3.4)$$

where $(x - \tau)_+ = \max\{x - \tau, 0\}$. The parameter κ measures the additional penalty assigned to high-count realizations beyond the tail threshold τ . Automation reduces residual exposure to the realized count. The post-automation loss is written as

$$\ell(x, a) = R(a)q(x), \quad (3.5)$$

where $R : [0, 1] \rightarrow \mathbb{R}_+$ is the residual-exposure function.

Assumption 3.1 (Residual exposure). The function $R(a)$ is continuous, nonincreasing, and convex on $[0, 1]$, with $R(0) = 1$ and $R(1) \geq 0$.

Two representative specifications are $R(a) = 1 - \xi a$ for $0 < \xi \leq 1$, and $R(a) = \exp(-\xi a)$ for $\xi > 0$. The first specification gives a linear residual-exposure rule, whereas the second captures diminishing residual exposure while keeping loss strictly positive. The structural analysis requires only the properties stated in Assumption 3.1. For every count law P considered below, assume $\mathbb{E}_P[q(X)] < \infty$. Given such a distribution P , the expected residual risk loss under automation intensity a is

$$L(a; P) = \mathbb{E}_P[\ell(X, a)] = \mathbb{E}_P[R(a)q(X)]. \quad (3.6)$$

Under the separable loss in (3.5), the count law affects the automation problem through the risk-damage functional

$$M_q(P) = \mathbb{E}_P[q(X)]. \quad (3.7)$$

3.3. Implementation and adjustment costs

Automation is costly. Let $C : [0, 1] \rightarrow \mathbb{R}_+$ denote the implementation cost of automation.

Assumption 3.2 (Implementation cost). The function $C(a)$ is continuous, nondecreasing, and convex on $[0, 1]$, with $C(0) = 0$.

A standard specification is

$$C(a) = c_1 a + c_2 a^2, \quad c_1, c_2 \geq 0. \quad (3.8)$$

Convexity captures the idea that higher levels of automation require disproportionately greater implementation effort, monitoring resources, computational capacity, or organizational adjustment. In the sequential problem, changing the automation level across periods may also be costly. Let $H(a_t, a_{t-1})$ denote the adjustment friction from moving from a_{t-1} to a_t .

Assumption 3.3 (Adjustment friction). The function $H(a, a')$ is continuous and nonnegative on $[0, 1]^2$, satisfies $H(a, a) = 0$, and is increasing in $|a - a'|$.

Two useful specifications are

$$H(a_t, a_{t-1}) = \eta |a_t - a_{t-1}|, \quad \eta \geq 0, \quad (3.9)$$

and

$$H(a_t, a_{t-1}) = \eta (a_t - a_{t-1})^2, \quad \eta \geq 0. \quad (3.10)$$

The first specification captures proportional switching friction, whereas the second captures smooth convex adjustment cost.

3.4. Static plug-in automation problem

Suppose first that the decision-maker has fitted a count distribution

$$\widehat{P}_t = P_{\widehat{\theta}_t} \in \mathcal{F} \quad (3.11)$$

and treats it as the true risk law for period t . The corresponding single-period plug-in automation problem is

$$a_t^{PI} \in \arg \min_{a \in [0,1]} \left\{ \mathbb{E}_{\widehat{P}_t}[\ell(X, a)] + C(a) \right\}. \quad (3.12)$$

Equivalently, using (3.6),

$$a_t^{PI} \in \arg \min_{a \in [0,1]} \left\{ L(a; \widehat{P}_t) + C(a) \right\}. \quad (3.13)$$

The plug-in policy is natural when the fitted count model is trusted. However, it can be fragile when the fitted law understates dispersion or tail probabilities. A fitted model may describe average counts reasonably well while still misrepresenting the tail region that determines severe automation losses. Whenever first-order conditions are invoked, R and C are assumed to be continuously differentiable on $(0, 1)$. Under this regularity, the first-order condition for an interior plug-in solution is

$$\frac{\partial}{\partial a} L(a; \widehat{P}_t) + C'(a) = 0. \quad (3.14)$$

When $\ell(x, a) = R(a)q(x)$, (3.14) becomes

$$R'(a) \mathbb{E}_{\widehat{P}_t}[q(X)] + C'(a) = 0. \quad (3.15)$$

Thus, the fitted distribution affects automation through the expected risk-damage functional $M_q(\widehat{P}_t)$. If the fitted model underestimates high-count tail probabilities, then $\mathbb{E}_{\widehat{P}_t}[q(X)]$ may be too small, leading the plug-in policy to choose insufficient automation.

3.5. Sequential plug-in automation with adjustment friction

The plug-in policy is implemented sequentially. At the beginning of period t , the decision-maker observes the fitted count law $\widehat{P}_t$ and the previously implemented automation level a_{t-1}^{PI} . The current automation decision is

$$a_t^{PI} \in \arg \min_{a \in [0,1]} \left\{ L(a; \widehat{P}_t) + C(a) + H(a, a_{t-1}^{PI}) \right\}. \quad (3.16)$$

Under the separable loss $\ell(x, a) = R(a)q(x)$, this becomes

$$a_t^{PI} \in \arg \min_{a \in [0,1]} \{R(a)M_q(\widehat{P}_t) + C(a) + H(a, a_{t-1}^{PI})\}. \quad (3.17)$$

The formulation is sequential because the fitted count law and automation decision are updated over time. It is path-dependent because the previous automation level enters the current objective through adjustment friction. It is not an infinite-horizon dynamic-programming problem and does not include a Bellman continuation value or require a specified transition kernel.

3.6. Oracle-informed sequential benchmark and realized policy loss

Let P_t^* denote the true current-period count law. The oracle-informed benchmark solves the same sequential adjustment problem as the feasible policies but uses the true current-period risk exposure:

$$a_t^* \in \arg \min_{a \in [0,1]} \{L(a; P_t^*) + C(a) + H(a, a_{t-1}^*)\}. \quad (3.18)$$

The oracle does not observe future realizations and does not use a true transition law. Its informational advantage is limited to knowledge of P_t^* when the current decision is made.

For a realized count x_t , define the realized period cost

$$g_t(x_t, a_t, a_{t-1}) = R(a_t)q(x_t) + C(a_t) + H(a_t, a_{t-1}). \quad (3.19)$$

Its conditional expectation under the true current-period law is

$$G_t(a_t, a_{t-1}; P_t^*) = L(a_t; P_t^*) + C(a_t) + H(a_t, a_{t-1}). \quad (3.20)$$

For a sequential policy π , define the oracle-relative realized cost difference

$$\text{Gap}_T(\pi) = \sum_{t=1}^T [g_t(X_t, a_t^\pi, a_{t-1}^\pi) - g_t(X_t, a_t^*, a_{t-1}^*)]. \quad (3.21)$$

This is a simulation-based policy-evaluation metric. It is not online-learning regret and, because compared policies generate different previous-action paths, it need not be nonnegative in every finite replication. The manuscript therefore reports Monte Carlo means and uncertainty measures rather than interpreting each realized value as a deterministic regret guarantee.

3.7. Count-model misspecification

The key source of fragility is count-model misspecification. Let

$$\Delta_t = d(P_t^*, \widehat{P}_t) \quad (3.22)$$

measure the discrepancy between the true count law and the fitted count law under a discrepancy $d(\cdot, \cdot)$. Wasserstein distance is used for the policy-loss bound in Section 6, whereas the implemented NB ambiguity set is defined directly in parameter space. A small value of Δ_t indicates that the fitted model is close to the true distribution, whereas a large value indicates dispersion or tail-shape misspecification. The quantities Δ_t and ρ_t serve different roles: Δ_t is an ex post discrepancy from the unknown true

law used in policy-loss analysis, whereas ρ_t is an ex ante, user-selected robustness parameter used to construct feasible decisions. No equality between the two quantities is assumed.

The distinction between dispersion misspecification and tail-shape misspecification is important. Dispersion misspecification occurs when the fitted model misstates the variance-to-mean relationship within a given count family. Tail-shape misspecification occurs when the fitted family itself is too restrictive to capture the probability of high-count events. The NB model is useful for analyzing the first type of error. Flexible long-tailed count alternatives motivate the second type of error, but they are not retained as main plug-in policy benchmarks in the final computational comparison.

The robust framework developed in Section 5 addresses this issue by replacing the single fitted law $\widehat{P}_t$ with an ambiguity set

$$\mathcal{P}_t(\widehat{P}_t, \rho_t), \quad (3.23)$$

where ρ_t controls the degree of distributional uncertainty. This ambiguity-set formulation follows the distributionally robust optimization tradition, in which decisions are protected against misspecified nominal distributions, sampling uncertainty, and out-of-sample distributional perturbations [22, 23, 63]. A broader review of these frameworks is provided by Rahimian and Mehrotra [64]. Because the present model focuses on dispersion and tail-risk uncertainty in a fitted count distribution, it is also closely related to robust optimization models with moment and dispersion ambiguity [24, 25]. The main structural question is how this ambiguity changes the risk-damage functional, robust automation incentives, persistence under adjustment friction, and oracle-relative policy gap.

4. Overdispersed count specifications

In this section, we introduce the count-risk specifications used to construct and interpret the automation policies. The Poisson model is used as an equidispersed benchmark, and the NB model is used as the main overdispersed count law for policy construction. Flexible long-tailed count alternatives, such as the Sichel model, are discussed only to motivate the possibility of tail-shape misspecification. They are not treated as final plug-in policy benchmarks. The distinction matters because the main objective is not to compare many flexible count models, but to evaluate whether a parsimonious fitted NB automation rule can be protected against local dispersion and upper-tail misspecification.

4.1. Poisson benchmark

The Poisson distribution is the canonical benchmark for count data. Let $X \sim \text{Poisson}(\mu)$ with $\mu > 0$. Its probability mass function is $\mathbb{P}(X = x) = e^{-\mu}\mu^x/x!$ for $x = 0, 1, 2, \dots$, and it satisfies $\mathbb{E}[X] = \mu$ and $\text{Var}(X) = \mu$. Thus, the conditional variance equals the conditional mean. This equidispersion restriction makes the Poisson model analytically convenient, but it can be too restrictive when risk events are clustered or when high-count realizations occur more often than the Poisson law predicts.

The Poisson model is not treated as a realistic description of clustered risk. Instead, it serves as a benchmark for evaluating the policy cost of ignoring overdispersion. A Poisson plug-in automation policy is expected to underestimate tail exposure when the true count process is overdispersed. This benchmark therefore clarifies the operational value of moving from an equidispersed plug-in rule to a NB plug-in rule and, ultimately, to a robust NB policy.

4.2. NB specification

The NB model provides the main overdispersed count specification used for policy construction. Under the mean–dispersion parameterization, $X \sim \text{NB}(\mu, r)$, where $\mu > 0$ is the mean and $r > 0$ is the dispersion parameter. Its probability mass function is

$$\mathbb{P}(X = x) = \frac{\Gamma(x+r)}{\Gamma(r)x!} \left(\frac{r}{r+\mu}\right)^r \left(\frac{\mu}{r+\mu}\right)^x, \quad x = 0, 1, 2, \dots \quad (4.1)$$

Under this parameterization, $\mathbb{E}[X] = \mu$ and $\text{Var}(X) = \mu + \mu^2/r$, so the variance-to-mean ratio is $\text{Var}(X)/\mathbb{E}[X] = 1 + \mu/r$. A smaller value of r therefore implies stronger overdispersion and greater tail exposure relative to the Poisson benchmark.

The NB model also admits a Gamma–Poisson mixture representation: $X | \Lambda \sim \text{Poisson}(\Lambda)$ and $\Lambda \sim \text{Gamma}(r, r/\mu)$, where the Gamma distribution is parameterized by shape and rate. This mixture interpretation is useful for automation modeling because it links observed count overdispersion to latent intensity heterogeneity. In operational terms, the realized risk count is Poisson conditional on a latent risk intensity, while that intensity varies across periods or regimes.

This combination of flexibility and interpretability makes the specification well suited to the present framework. It captures excess variance while preserving a direct interpretation of dispersion. In the computational experiments, the fitted law is used to construct both the main plug-in policy and the robust policy centered on that law. Accordingly, it is the operational count specification being robustified, rather than merely one candidate among a broad set of competing final plug-in policies.

4.3. Long-tailed count alternatives

Flexible long-tailed count models motivate the possibility that a fitted NB law may understate upper-tail exposure. The Sichel distribution provides one such example. It was originally developed as a flexible mixed-Poisson model for long-tailed count phenomena and remains useful when tail behavior is not adequately captured by simpler Poisson or NB specifications [11, 12]. It can be represented as a Poisson mixture in which the latent intensity follows a generalized inverse Gaussian (GIG) distribution:

$$X | \Lambda \sim \text{Poisson}(\Lambda), \quad \Lambda \sim \text{GIG}(\lambda, \chi, \psi), \quad (4.2)$$

where $\lambda \in \mathbb{R}$, $\chi > 0$, and $\psi > 0$.

The density of Λ is

$$f_{\text{GIG}}(\lambda_0) = \frac{(\psi/\chi)^{\lambda/2}}{2K_\lambda(\sqrt{\chi\psi})} \lambda_0^{\lambda-1} \exp\left\{-\frac{1}{2}(\chi\lambda_0^{-1} + \psi\lambda_0)\right\}, \quad \lambda_0 > 0, \quad (4.3)$$

where $K_\lambda(\cdot)$ denotes the modified Bessel function of the second kind. Integrating out the latent intensity yields the Sichel probability mass function

$$\mathbb{P}(X = x) = \frac{1}{x!} \left(\frac{\psi}{\chi}\right)^{\lambda/2} \frac{K_{\lambda+x}(\sqrt{\chi(\psi+2)})}{K_\lambda(\sqrt{\chi\psi})} \left(\frac{\chi}{\psi+2}\right)^{(\lambda+x)/2}, \quad x = 0, 1, 2, \dots \quad (4.4)$$

The parameters of the Sichel model jointly govern dispersion and tail shape. Its generalized inverse Gaussian mixing structure is more flexible than the Gamma mixing structure underlying the NB model,

allowing the Sichel distribution to represent a wider range of heavy-tail behavior. Here, this flexibility is used only to motivate the possibility of tail-shape misspecification. The final policy comparison therefore does not include Sichel plug-in or robust Sichel policies. Instead, long-tailed misspecification is examined through stress environments that assess whether robust NB automation improves on the NB plug-in benchmark. Accordingly, the Sichel model serves a diagnostic and motivational role rather than a final policy-construction benchmark.

4.4. Tail-risk functionals

The automation problem depends on the count distribution through expected risk-damage functionals. For any count law P , define

$$M_q(P) = \mathbb{E}_P[q(X)], \quad (4.5)$$

where $q(x)$ is the primitive damage function introduced in Section 3.2. When

$$q(x) = x + \kappa(x - \tau)_+, \quad (4.6)$$

the risk-damage functional becomes

$$M_q(P) = \mathbb{E}_P[X] + \kappa \mathbb{E}_P[(X - \tau)_+]. \quad (4.7)$$

The first term captures average event volume, while the second term captures tail exposure above the threshold τ . Two count distributions with similar means may therefore imply different automation decisions if their tail expectations differ.

This distinction is central to the robust automation problem. A fitted NB model may capture overdispersion through the dispersion parameter r , but the resulting value of $M_q(\widehat{P})$ may still be too small if the fitted law understates the exceedance term

$$\mathbb{E}_P[(X - \tau)_+]. \quad (4.8)$$

The robust formulation in Section 5 addresses this issue by replacing the fitted value $M_q(\widehat{P})$ with a worst-case risk-damage functional over an ambiguity set around the fitted NB law.

4.5. Tail ordering and policy relevance

To compare count laws in a policy-relevant way, the following tail-ordering concept is used. For two count distributions P_1 and P_2 , P_2 is said to have greater q -tail exposure than P_1 if

$$\mathbb{E}_{P_2}[q(X)] \geq \mathbb{E}_{P_1}[q(X)]. \quad (4.9)$$

When $q(x) = x + \kappa(x - \tau)_+$, this ordering compares both mean exposure and tail exposure. If the means are matched, the ordering is driven by the expected exceedance term in (4.8).

This ordering is useful because the plug-in single-period automation problem depends on the fitted distribution through $M_q(P)$. Under the separable loss structure

$$\ell(x, a) = R(a)q(x), \quad (4.10)$$

the single-period objective can be written as

$$R(a)M_q(P) + C(a). \quad (4.11)$$

Thus, for a fixed cost function and residual-exposure function, a distribution with greater q -tail exposure generally creates stronger incentives for automation. The exact effect depends on the curvature of R and C , but the economic mechanism is direct: Higher tail exposure increases the marginal benefit of reducing residual risk.

This observation motivates the long-tail stress design in the computational experiments. The purpose is not to select a more flexible plug-in count model. Rather, the purpose is to evaluate whether a robust NB policy can protect against tail-relevant misspecification when the fitted NB law understates high-count exposure.

4.6. Parameter estimation and model comparison

For a sample of counts $x_1, \dots, x_n$, the fitted distribution is denoted by $\widehat{P} = P_{\widehat{\theta}}$. The framework does not require one particular estimator. Maximum likelihood is used for full-sample empirical model-fit comparisons, while the rolling simulation implementation uses a transparent moment estimator for the NB parameters:

$$\widehat{\mu} = \bar{x}, \quad \widehat{r} = \begin{cases} \frac{\bar{x}^2}{s^2 - \bar{x}}, & s^2 > \bar{x}, \\ +\infty, & s^2 \leq \bar{x}, \end{cases} \quad (4.12)$$

where $\bar{x}$ and s^2 are the sample mean and sample variance. In numerical work, a large finite value is used when the fitted variance does not exceed the fitted mean, so that the NB law approaches its Poisson limit.

For likelihood-based comparisons, the fitted parameter may instead be obtained from

$$\widehat{\theta} \in \arg \max_{\theta \in \Theta} \sum_{i=1}^n \log p_{\theta}(x_i), \quad (4.13)$$

and statistical fit is summarized by the log-likelihood, Akaike information criterion (AIC), and Bayesian information criterion (BIC) [73, 74]:

$$\text{AIC} = -2 \log \widehat{L} + 2k, \quad \text{BIC} = -2 \log \widehat{L} + k \log n. \quad (4.14)$$

Statistical fit is supporting evidence rather than the final decision criterion. The computational analysis evaluates fitted and robust exposures through downstream outcomes: realized total cost, tail loss, implementation cost, adjustment cost, switching frequency, under-automation, over-automation, and the oracle-relative realized cost difference in simulation. This is consistent with decision-oriented predict-then-optimize evaluation [72].

4.7. Role of the count specifications in the robust model

The count specifications play distinct roles. The Poisson law is an equidispersed benchmark. The NB law is the fitted model used for both plug-in and robust policy construction. Flexible long-tailed count models motivate cross-shape stress environments in which the fitted NB law understates upper-tail exposure; they are not additional final policy classes. Section 5 defines the local parametric NB ambiguity set and the adverse-direction exposure used by the robust sequential policy.

5. Distributionally robust automation

In this section, we formulate the distributionally robust automation problem. The plug-in policies introduced in Section 3 treat a fitted count law as if it were correct. In contrast, the robust policy recognizes that the fitted law may understate dispersion, tail exposure, or the probability of high-count risk events. Although the formulation can be written for a general count family, the implemented robust policy uses a local ambiguity set around the fitted NB mean and dispersion parameters. Thus, the purpose is not to compare competing flexible plug-in count policies, but to robustify a parsimonious and interpretable NB automation rule.

5.1. Ambiguity around a fitted NB law

Let

$$\widehat{P}_t^{NB} = P_{\widehat{\mu}_t, \widehat{r}_t}^{NB} \quad (5.1)$$

denote the fitted NB count law at time t , where $\widehat{\mu}_t$ and $\widehat{r}_t$ are estimated from historical or rolling count data. The true count law P_t^* is not observed directly and may differ from $\widehat{P}_t^{NB}$ because of finite-sample error, regime changes, dispersion misspecification, or tail-risk underestimation.

This uncertainty is represented through a local ambiguity set around the fitted NB law:

$$\mathcal{P}_t^{NB}(\widehat{P}_t^{NB}, \rho_t) = \{P \in \mathcal{F}^{NB} : d(P, \widehat{P}_t^{NB}) \leq \rho_t\}, \quad (5.2)$$

where $d(\cdot, \cdot)$ is a distributional or parameter-space discrepancy, $0 \leq \rho_t < 1$ is the ambiguity radius, and $\mathcal{F}^{NB}$ is the admissible NB neighborhood used for robust policy construction.

A convenient parameter-space version of (5.2) is

$$\mathcal{P}_t^{NB}(\widehat{\mu}_t, \widehat{r}_t, \rho_t) = \left\{ P_{\mu, r}^{NB} : \mu \in [\widehat{\mu}_t(1 - \rho_t), \widehat{\mu}_t(1 + \rho_t)], r \in \left[\frac{\widehat{r}_t}{1 + \rho_t}, \widehat{r}_t(1 + \rho_t) \right] \right\}. \quad (5.3)$$

The rectangle is two-sided in both parameters. Its upper-mean/lower- r corner represents the adverse direction relevant to the increasing convex damage functional used here: The upper mean increases event intensity, while the lower dispersion parameter increases overdispersion. The distinction between the two-sided parameter set and the adverse corner used for policy construction is maintained explicitly below.

The ambiguity radius ρ_t measures the degree of protection against distributional error. When $\rho_t = 0$, the ambiguity set collapses to the fitted NB law and the robust problem reduces to the NB plug-in problem. When $\rho_t > 0$, the parameter rectangle admits both upward and downward local perturbations, while the implemented adverse-direction exposure protects against the upper-mean/lower-dispersion corner.

The parameters μ and r are perturbed because they carry the main operational meaning in the fitted NB law: μ controls event intensity, whereas r controls overdispersion and tail exposure. This local parametric construction is intentionally auditable in rolling implementation, especially when the estimation window is not large enough to support a fully nonparametric ambiguity set reliably. Other ambiguity formulations are possible, including moment sets, Wasserstein balls, ϕ -divergence neighborhoods, zero-inflated or hurdle extensions, and ambiguity sets over transition dynamics. Those

alternatives require different calibration targets and are therefore treated as extensions rather than as additional benchmarks in the present decision-focused comparison.

This modeling choice follows the distributionally robust optimization literature, in which ambiguity sets are used to protect decisions against sampling error, misspecified nominal distributions, and plausible out-of-sample distributional perturbations [22,23,63]. A broader review of these frameworks is provided by Rahimian and Mehrotra [64]. Recent operations research work further develops ambiguity-set formulations based on moment information, dispersion information, stochastic dominance, optimal transport, and Sinkhorn regularization [24, 25, 66]. Related contributions address dimensionality reduction, optimal transport, and Sinkhorn regularization [68–70]. The present formulation differs from generic DRO applications because the ambiguity set is placed on a fitted count-risk law that directly determines the risk-damage functional, robust automation incentives, and oracle-relative policy loss.

5.2. One-period robust benchmark without adjustment friction

For a fitted NB law $\widehat{P}_t^{NB}$, the one-period robust benchmark without adjustment friction is

$$a_t^{R,NB} \in \arg \min_{a \in [0,1]} \sup_{P \in \mathcal{P}_t^{NB}(\widehat{P}_t^{NB}, \rho_t)} \{ \mathbb{E}_P[\ell(X, a)] + C(a) \}. \quad (5.4)$$

Using the expected residual loss $L(a; P)$ defined in (3.6), this can be written as

$$a_t^{R,NB} \in \arg \min_{a \in [0,1]} \left\{ \sup_{P \in \mathcal{P}_t^{NB}(\widehat{P}_t^{NB}, \rho_t)} L(a; P) + C(a) \right\}. \quad (5.5)$$

The robust policy differs from the plug-in policy in (3.13) because it minimizes worst-case expected residual loss rather than expected residual loss under a single fitted model. The difference is policy-relevant when high-count events are costly. If the fitted NB law understates tail exposure, then the plug-in policy may select insufficient automation, whereas the robust policy accounts for nearby laws with larger tail-relevant risk.

Under the separable loss structure $\ell(x, a) = R(a)q(x)$ in (3.5), the worst-case residual loss becomes

$$\sup_{P \in \mathcal{P}_t^{NB}(\widehat{P}_t^{NB}, \rho_t)} R(a) \mathbb{E}_P[q(X)]. \quad (5.6)$$

Since $R(a) \geq 0$, (5.6) is equivalent to

$$R(a) \sup_{P \in \mathcal{P}_t^{NB}(\widehat{P}_t^{NB}, \rho_t)} \mathbb{E}_P[q(X)]. \quad (5.7)$$

Thus, the robust problem depends on the ambiguity set through the worst-case risk-damage functional

$$\overline{M}_q^{NB}(\widehat{P}_t^{NB}, \rho_t) = \sup_{P \in \mathcal{P}_t^{NB}(\widehat{P}_t^{NB}, \rho_t)} \mathbb{E}_P[q(X)]. \quad (5.8)$$

For the numerical policy, define the adverse-corner exposure

$$\widetilde{M}_{q,t}^{NB}(\rho_t) = M_q\left(P_{\widehat{\mu}_t(1+\rho_t), \widehat{\tau}_t/(1+\rho_t)}^{NB}\right). \quad (5.9)$$

When $M_q(P_{\mu,r}^{NB})$ is nondecreasing in μ and nonincreasing in r on the parameter rectangle, the implemented exposure equals the exact supremum in (5.8). Otherwise, (5.9) is interpreted as a transparent directional stress rule rather than as an unrestricted distributional supremum. The computational experiments report this implementation explicitly.

The corresponding one-period robust benchmark can therefore be written as

$$a_t^{R,NB} \in \arg \min_{a \in [0,1]} \left\{ R(a) \widetilde{M}_{q,t}^{NB}(\rho_t) + C(a) \right\}. \quad (5.10)$$

This expression clarifies the role of distributional robustness. The ambiguity set replaces the fitted risk-damage estimate

$$M_q(\widehat{P}_t^{NB}) = \mathbb{E}_{\widehat{P}_t^{NB}}[q(X)] \quad (5.11)$$

with the implemented adverse-corner counterpart $\widetilde{M}_{q,t}^{NB}(\rho_t)$. The larger this ambiguity-adjusted functional is, the stronger the marginal incentive to increase automation. The policy effect of robustness is therefore mediated through the risk-damage functional rather than through an arbitrary increase in automation intensity.

5.3. Sequential robust automation with adjustment friction

The robust policy is recalculated sequentially as the fitted NB parameters evolve. It solves

$$a_t^{R,NB} \in \arg \min_{a \in [0,1]} \left\{ \sup_{P \in \mathcal{P}_t^{NB}(\widehat{P}_t^{NB}, \rho_t)} L(a; P) + C(a) + H(a, a_{t-1}^{R,NB}) \right\}. \quad (5.12)$$

Under the separable loss, this is

$$a_t^{R,NB} \in \arg \min_{a \in [0,1]} \left\{ R(a) \widetilde{M}_{q,t}^{NB}(\rho_t) + C(a) + H(a, a_{t-1}^{R,NB}) \right\}. \quad (5.13)$$

The previous automation level makes the policy path history-dependent. A larger ambiguity-adjusted exposure may increase the current automation level, whereas the adjustment cost discourages abrupt changes and can delay reversal after automation has increased. This effect is termed adjustment persistence, not as an optimal hysteresis band, because the formulation does not solve an infinite-horizon switching problem.

5.4. Practical interpretation and selection of the ambiguity radius

The radius has a direct percentage interpretation. For example, $\rho = 0.10$ allows a mean perturbation of $\pm 10\%$, while the adverse corner raises the fitted mean by 10% and reduces r by the factor $1/1.10$. The radius should not be interpreted as a universal preference for more conservative automation. It is a tuning parameter governing the trade-off between tail protection and excess implementation cost.

In an applied setting, a candidate radius can be selected on a temporally separate validation sample. One decision-oriented rule is

$$\widehat{\rho} \in \arg \min_{\rho \in \mathcal{R}} \left\{ \widehat{\text{Cost}}_{\text{val}}(\rho) + \omega \widehat{\text{TailLoss}}_{\text{val}}(\rho) \right\}, \quad \omega \geq 0, \quad (5.14)$$

where $\mathcal{R}$ is a prespecified radius grid. An alternative is to minimize validation cost subject to a tail-loss ceiling. The experiments use a broad common radius grid across adverse, correctly specified,

low-dispersion, and reverse-misspecification environments to show that the preferred radius is scenario-dependent rather than mechanically increasing.

For implementation, the radius should be selected together with an over-automation diagnostic. A larger radius is appropriate when tail-loss, liquidity, safety, or service-continuity constraints dominate average cost. A smaller radius is preferable when implementation budgets are tight, the fitted count process is stable, or rolling diagnostics indicate weak dispersion. In all cases, the selected value should be justified by validation performance, scenario analysis, and the marginal cost of additional intervention.

5.5. Plug-in, robust, and oracle-informed sequential policies

All policies solve the same sequential adjustment problem

$$a_t^\pi \in \arg \min_{a \in [0,1]} \{R(a)m_t^\pi + C(a) + H(a, a_{t-1}^\pi)\}, \quad (5.15)$$

and differ only in the risk-damage exposure supplied to that problem:

$$m_t^P = M_q(\widehat{P}_t^P), \quad (5.16)$$

$$m_t^{NB} = M_q(\widehat{P}_t^{NB}), \quad (5.17)$$

$$m_t^{R,NB} = \widetilde{M}_{q,t}^{NB}(\rho_t), \quad (5.18)$$

$$m_t^\star = M_q(P_t^\star). \quad (5.19)$$

Thus, the comparison separates the cost of ignoring overdispersion, the cost of relying on a single fitted overdispersed law, and the benefit or cost of robustifying the fitted NB rule.

5.6. Oracle-relative realized cost difference

Using the realized period cost in (3.19), the primary performance metric is the oracle-relative realized cost difference in (3.21). It measures the downstream cost of using a fitted or robust exposure rather than the true current-period exposure. No compared policy uses a continuation value or an estimated transition model. This is a decision-oriented model-evaluation criterion rather than online-learning regret [72].

5.7. Distributional misspecification and policy loss

The performance gap between plug-in and oracle policies is driven by distributional misspecification. Let

$$\Delta_t = d(P_t^\star, \widehat{P}_t^{NB}) \quad (5.20)$$

measure the discrepancy between the true count law and the fitted NB law. Suppose the risk-damage functional is Lipschitz with respect to d , meaning that there exists $K_q > 0$ such that

$$|\mathbb{E}_P[q(X)] - \mathbb{E}_Q[q(X)]| \leq K_q d(P, Q) \quad (5.21)$$

for all relevant P, Q . Then

$$|M_q(P_t^\star) - M_q(\widehat{P}_t^{NB})| \leq K_q \Delta_t. \quad (5.22)$$

Inequality (5.22) shows how distributional error translates into policy-relevant risk-damage error. A fitted model may have a small average prediction error but still produce a large policy error if it understates the tail component of $q(X)$. This is especially important when $q(x) = x + \kappa(x - \tau)_+$, because the exceedance term magnifies errors in the upper tail of the count distribution.

The robust policy reduces this sensitivity by replacing the fitted functional $M_q(\widehat{P}_t^{NB})$ with the worst-case functional $\overline{M}_q^{NB}(\widehat{P}_t^{NB}, \rho_t)$. If the true law belongs to the ambiguity set,

$$P_t^* \in \mathcal{P}_t^{NB}(\widehat{P}_t^{NB}, \rho_t), \quad (5.23)$$

then

$$M_q(P_t^*) \leq \overline{M}_q^{NB}(\widehat{P}_t^{NB}, \rho_t). \quad (5.24)$$

Thus, robust automation protects against underestimation of the risk-damage functional whenever the ambiguity set contains the true law.

5.8. Within-family robustness and long-tail stress

The final computational design uses within-family NB robustness. In this case, the ambiguity set is centered on a fitted NB model and allows nearby mean and dispersion values within the same family, as in (5.3). This design reflects the main operational goal of the study: to evaluate whether a parsimonious fitted NB automation policy can be protected against local misspecification.

The experiments do not use Sichel plug-in policies as final benchmarks. Instead, long-tail misspecification is represented through stress environments that evaluate whether robust NB policies remain stable when the true count process has heavier tails than the fitted NB law suggests. This distinction is important. The purpose is not to select the most flexible count distribution, but to measure the policy value of robustifying an interpretable NB automation rule.

This design separates two sources of policy loss. The first is overdispersion neglect, measured by comparing Poisson plug-in and NB plug-in policies. The second is local misspecification of a fitted overdispersed law, measured by comparing NB plug-in and robust NB policies. Long-tail stress environments then test whether this robustification remains valuable when tail exposure exceeds what the fitted NB law suggests.

5.9. Transition to structural analysis

The robust formulation has three structural implications. First, increasing the ambiguity radius increases the worst-case risk-damage functional $\overline{M}_q^{NB}(\widehat{P}_t^{NB}, \rho_t)$. Second, this increase can shift automation activation incentives toward earlier intervention. Third, when adjustment friction is present, the same ambiguity can make automation more persistent after activation by increasing the cost of underreacting to tail-relevant risk.

In Section 6, these implications are formalized for the robust NB automation model.

6. Structural results

In this section, we study the sequential period problem implemented in the computational and empirical analyses. For a current risk-damage exposure $m \geq 0$ and previous automation level $a_- \in [0, 1]$,

define

$$J(a; m, a_-) = R(a)m + C(a) + H(a, a_-), \quad a \in [0, 1]. \quad (6.1)$$

The policy map is denoted by

$$a^*(m, a_-) \in \arg \min_{a \in [0, 1]} J(a; m, a_-). \quad (6.2)$$

6.1. Existence and uniqueness of the period decision

Proposition 6.1 (Existence and uniqueness). *Suppose R is continuous and convex, C is continuous and strictly convex, and $H(\cdot, a_-)$ is continuous and convex for every $a_- \in [0, 1]$. Then, (6.2) has a unique solution for every $m \geq 0$ and $a_- \in [0, 1]$.*

Proof. The objective is continuous on the compact action set $[0, 1]$, so a minimizer exists. Strict convexity of C , together with convexity of mR and $H(\cdot, a_-)$, makes the objective strictly convex, so the minimizer is unique. $\square$

For the numerical specification

$$R(a) = e^{-\xi a}, \quad C(a) = c_1 a + c_2 a^2, \quad H(a, a_-) = \eta(a - a_-)^2, \quad (6.3)$$

strict convexity holds whenever $c_2 + \eta > 0$.

6.2. Adverse-direction NB exposure

For $0 \leq \rho < 1$, let

$$\Theta^{NB}(\widehat{\mu}, \widehat{r}, \rho) = \left\{ (\mu, r) : \widehat{\mu}(1 - \rho) \leq \mu \leq \widehat{\mu}(1 + \rho), \quad \frac{\widehat{r}}{1 + \rho} \leq r \leq \widehat{r}(1 + \rho) \right\}. \quad (6.4)$$

Define the adverse-corner exposure

$$\widetilde{M}_q^{NB}(\widehat{\mu}, \widehat{r}, \rho) = M_q(P_{\widehat{\mu}(1+\rho), \widehat{r}/(1+\rho)}^{NB}). \quad (6.5)$$

Assumption 6.2 (Adverse-corner monotonicity). Over the parameter region used for policy construction, $M_q(P_{\mu, r}^{NB})$ is nondecreasing in μ and nonincreasing in r .

Proposition 6.3 (Corner characterization). *Under Assumption 6.2,*

$$\sup_{(\mu, r) \in \Theta^{NB}(\widehat{\mu}, \widehat{r}, \rho)} M_q(P_{\mu, r}^{NB}) = \widetilde{M}_q^{NB}(\widehat{\mu}, \widehat{r}, \rho). \quad (6.6)$$

Moreover, the right-hand side is nondecreasing in ρ .

Proof. The upper mean endpoint and lower r endpoint maximize the functional under the stated coordinatewise monotonicity. As ρ increases, the upper mean endpoint increases and the lower r endpoint decreases, so the corner exposure cannot decrease. $\square$

Assumption 6.2 is transparent and testable numerically on the parameter rectangle. If it is not imposed, the same formula is retained as a directional stress implementation and is not described as an exact supremum.

6.3. Monotone automation response

Proposition 6.4 (Monotonicity in risk exposure). *Suppose R is nonincreasing and the minimizer in Proposition 6.1 is unique. Then, $a^*(m, a_-)$ is nondecreasing in m for every fixed a_- .*

Proof. For $a_2 > a_1$ and $m_2 > m_1$,

$$[J(a_2; m_2, a_-) - J(a_1; m_2, a_-)] - [J(a_2; m_1, a_-) - J(a_1; m_1, a_-)] = (m_2 - m_1)[R(a_2) - R(a_1)] \leq 0.$$

Thus, J has decreasing differences in (a, m) . Economically, a larger exposure m increases the penalty from leaving residual exposure unreduced. Because $R(a)$ is nonincreasing, higher automation becomes relatively more attractive as m increases. The monotone comparative statics result for minimization, together with uniqueness, implies that the minimizer is nondecreasing in m . $\square$

Corollary 6.5 (Monotonicity in the ambiguity radius). *Under Assumption 6.2, the robust period action*

$$a^{R,NB}(\rho, a_-) = a^*(\widetilde{M}_q^{NB}(\widehat{\mu}, \widehat{r}, \rho), a_-)$$

is nondecreasing in ρ for fixed fitted parameters and fixed a_- .

6.4. Adjustment-induced path dependence

Proposition 6.6 (Dependence on previous automation). *Let $H(a, a_-) = \eta(a - a_-)^2$ with $\eta > 0$, and suppose the unique solution is interior and twice differentiable. Then,*

$$\frac{\partial a^*(m, a_-)}{\partial a_-} = \frac{2\eta}{R''(a^*)m + C''(a^*) + 2\eta} > 0. \quad (6.7)$$

Proof. The first-order condition is $R'(a)m + C'(a) + 2\eta(a - a_-) = 0$. Implicit differentiation with respect to a_- gives (6.7). The denominator is the second derivative of the strictly convex period objective and is therefore positive. $\square$

The result establishes policy persistence rather than an optimal entry–exit hysteresis band: A higher previously implemented automation level raises the current period action, other things equal.

6.5. Explicit activation condition

Under (6.3), the right derivative of the objective at zero is

$$J'_+(0; m, a_-) = -\xi m + c_1 - 2\eta a_-. \quad (6.8)$$

Because the objective is strictly convex,

$$a^*(m, a_-) > 0 \iff m > \frac{c_1 - 2\eta a_-}{\xi}. \quad (6.9)$$

If the right-hand side is negative, positive automation is selected for every feasible $m \geq 0$, including $m = 0$. Hence, higher exposure promotes earlier activation, while a larger previous action lowers the exposure required to maintain positive automation.

6.6. Directional robustness: Benefit and cost

Let $m^* = M_q(P^*)$, $\widehat{m} = M_q(\widehat{P}^{NB})$, and $\widetilde{m}(\rho) = \widetilde{M}_q^{NB}(\widehat{\mu}, \widehat{r}, \rho)$. The following ordering result connects the theoretical policy map to the balanced experimental design.

Proposition 6.7 (Directional correction and amplification). *Fix a_- . Under Proposition 6.4:*

1. If $\widehat{m} \leq \widetilde{m}(\rho) \leq m^*$, then

$$a^*(\widehat{m}, a_-) \leq a^*(\widetilde{m}(\rho), a_-) \leq a^*(m^*, a_-).$$

Thus directional robustness partially corrects under-automation.

2. If $m^* \leq \widehat{m} \leq \widetilde{m}(\rho)$, then,

$$a^*(m^*, a_-) \leq a^*(\widehat{m}, a_-) \leq a^*(\widetilde{m}(\rho), a_-).$$

Thus, directional robustness amplifies over-automation.

Proof. Both claims follow directly from the monotonicity of the policy map in the exposure argument. The first ordering corresponds to a nominal model that understates the policy-relevant exposure; the robust exposure then moves the action toward the oracle-informed action. The second ordering corresponds to a nominal model that is already conservative; the same adverse-direction adjustment moves the action farther above the oracle-informed action and can increase over-automation. $\square$

The proposition does not claim that a robust policy always has lower realized cost. Under correct specification, $m^* = \widehat{m}$, every positive radius is conservative. Under reverse misspecification, $m^* < \widehat{m}$, the adverse-direction rule can increase the implementation cost and the frequency of over-automation. Accordingly, the preferred radius is generally scenario-dependent and may equal zero.

6.7. One-period plug-in policy-loss bound

The damage function $q(x) = x + \kappa(x - \tau)_+$ is $(1 + \kappa)$ -Lipschitz. Therefore, for count laws with finite first moments, Kantorovich–Rubinstein duality gives

$$|M_q(P) - M_q(Q)| \leq (1 + \kappa)W_1(P, Q). \quad (6.10)$$

Since Assumption 3.1 implies $0 \leq R(a) \leq 1$, conditionally on a common previous action a_- , the standard plug-in decision argument yields

$$J(a^{PI}; m^*, a_-) - J(a^*; m^*, a_-) \leq 2(1 + \kappa)W_1(P^*, \widehat{P}). \quad (6.11)$$

Proof. Let $\widehat{J}(a) = R(a)M_q(\widehat{P}) + C(a) + H(a, a_-)$ and $J^*(a) = R(a)M_q(P^*) + C(a) + H(a, a_-)$. Uniformly in a ,

$$|J^*(a) - \widehat{J}(a)| \leq |M_q(P^*) - M_q(\widehat{P})| \leq (1 + \kappa)W_1(P^*, \widehat{P}).$$

Apply this bound at the plug-in and oracle minimizers and use plug-in optimality. $\square$

This is a one-period decision guarantee conditional on the same previous action. It is not an infinite-horizon or pathwise regret guarantee, because different sequential policies generally generate different previous-action paths.

6.8. Implications for the experiments

The structural analysis yields testable implications. First, increasing the adverse-direction radius weakly increases the robust action for fixed fitted parameters and previous action. Second, quadratic adjustment friction produces history-dependent persistence. Third, positive robustness may improve policy alignment when the nominal exposure is too small, but it may create unnecessary automation when the nominal exposure is correct or too large. The computational design therefore includes adverse, correctly specified, low-dispersion, and reverse-misspecification environments, and reports both under-automation and over-automation, together with total cost, tail loss, implementation cost, and adjustment cost.

7. Computational experiments

In this section, we evaluate the proposed robust NB automation framework through controlled computational experiments. The purpose is to assess fitted count models through the sequential decisions and realized costs they induce, rather than through statistical fit alone. The revised design is balanced across favorable, neutral, and unfavorable directions of misspecification. It therefore evaluates both the tail-protection benefit of robustness and the possibility that an unnecessarily conservative ambiguity radius can cause over-automation and excess implementation cost.

The experiments address four questions. First, what is the policy cost of approximating overdispersed count risk by an equidispersed Poisson model? Second, how much does a fitted NB policy improve on this benchmark? Third, when does local robustification of the fitted NB law reduce downstream policy loss? Fourth, when does the same robustification become unnecessarily conservative because the nominal model is already accurate or overstates the true count-risk exposure?

The computational experiments are empirically calibrated. The monthly S&P 500 jump-count series constructed in Section 8 provides the baseline scale, dispersion, and regime frequency. Controlled simulation is then used so that the current true count law is known in each period and an oracle-informed sequential benchmark can be evaluated. Flexible long-tailed count behavior is represented through a mixture stress environment rather than through a separate Sichel plug-in policy. The main policy comparison therefore focuses on Poisson plug-in, NB plug-in, robust NB, and oracle-informed sequential policies.

7.1. Common sequential implementation and performance measures

The computational study implements the plug-in, robust, and oracle-informed policies defined in Section 5.5. All policies therefore share the same sequential adjustment structure and differ only in the risk-damage exposure supplied to the current-period optimization problem. Their paths remain history-dependent through the previously implemented action and the adjustment-cost term, but none uses a Bellman continuation value or an estimated state-transition kernel.

At each evaluation date, the Poisson and NB plug-in policies are constructed from estimates obtained using the preceding rolling window. The robust NB policy uses the same rolling NB estimate but evaluates the risk-damage functional at the upper-mean and lower- r adverse corner defined in Section 5. In particular, its period- t exposure is

$$m_t^{R,NB}(\rho) = M_q\left(P_{\widehat{\mu}_t(1+\rho), \widehat{r}_t/(1+\rho)}^{NB}\right). \quad (7.1)$$

The oracle-informed benchmark uses the true current-period count law and regime but neither observes future count realizations nor optimizes a continuation value.

For the baseline experiments, the functional specifications introduced earlier are evaluated at $\tau = 2$, $\kappa = 2$, $\xi = 1.5$, $c_1 = 0.05$, $c_2 = 0.40$, and $\eta = 0.10$. Each replication uses a rolling estimation window of $W = 120$ periods and an evaluation horizon of $T = 360$ periods. The reported results are based on $N_{MC} = 200$ Monte Carlo (MC) replications. A fixed master seed generates a replication-level seed table, and common random numbers are used whenever policies, ambiguity radii, or cost settings are compared within the same data-generating environment. This matched design reduces simulation noise and isolates policy differences from path-specific variation.

For a realized count path, the period cost of policy π is

$$g_t^\pi = R(a_t^\pi)q(X_t) + C(a_t^\pi) + H(a_t^\pi, a_{t-1}^\pi), \quad (7.2)$$

and its cumulative realized cost is

$$\text{Cost}_T(\pi) = \sum_{t=1}^T g_t^\pi. \quad (7.3)$$

The corresponding oracle-relative cost difference is

$$\text{Gap}_T(\pi) = \text{Cost}_T(\pi) - \text{Cost}_T(\pi^*). \quad (7.4)$$

Because competing policies generally generate different adjustment paths, $\text{Gap}_T(\pi)$ is interpreted as a realized benchmark difference rather than as an infinite-horizon dynamic-regret guarantee.

To distinguish protection against extreme counts from overall economic performance, tail loss is measured separately as

$$\text{TailLoss}_T(\pi) = \sum_{t=1}^T R(a_t^\pi)(X_t - \tau)_+. \quad (7.5)$$

Automation intensity and policy persistence are summarized by

$$\bar{a}_T(\pi) = \frac{1}{T} \sum_{t=1}^T a_t^\pi, \quad (7.6)$$

$$\text{Switch}_T(\pi) = \sum_{t=2}^T \mathbf{1}\{|a_t^\pi - a_{t-1}^\pi| > \varepsilon_a\}, \quad (7.7)$$

and

$$\text{MeanAbsChange}_T(\pi) = \frac{1}{T-1} \sum_{t=2}^T |a_t^\pi - a_{t-1}^\pi|. \quad (7.8)$$

The switching and action-change measures provide computational counterparts to the adjustment-induced persistence result established in Section 6.

Under- and over-automation are assessed relative to the oracle-informed current-period action using the economically meaningful tolerance $\varepsilon_a = 0.025$:

$$\text{UnderAuto}_T(\pi) = \frac{1}{T} \sum_{t=1}^T \mathbf{1}\{a_t^\pi < a_t^* - \varepsilon_a\}, \quad (7.9)$$

$$\text{OverAuto}_T(\pi) = \frac{1}{T} \sum_{t=1}^T \mathbf{1}\{a_t^\pi > a_t^* + \varepsilon_a\}. \quad (7.10)$$

The second measure is particularly important in the neutral and reverse misspecification environments, where robustness may create excessive intervention rather than additional protection.

Finally, total cost is decomposed into implementation and adjustment components:

$$\text{ImplCost}_T(\pi) = \sum_{t=1}^T C(a_t^\pi), \quad (7.11)$$

$$\text{AdjCost}_T(\pi) = \sum_{t=1}^T H(a_t^\pi, a_{t-1}^\pi). \quad (7.12)$$

For matched comparisons with the nominal NB policy, the analysis also reports the robust-minus-NB total-cost difference and the NB-minus-robust tail-loss difference. Together, these measures reveal whether robustness reduces extreme-event exposure sufficiently to compensate for its additional implementation and adjustment costs.

The numerical implementation is computationally parsimonious. For each policy-period pair, the optimization is a deterministic one-dimensional bounded minimization over $a \in [0, 1]$. Under the exponential residual-exposure and quadratic cost specification, the implementation evaluates the boundary actions $a = 0$ and $a = 1$ and solves the scalar first-order condition for any interior candidate. The previous action initializes the sequential path, and fixed absolute tolerances are used for scalar root solving and objective comparison. For T periods and K ambiguity radii, policy evaluation scales linearly as $O(TK)$, apart from rolling count-model estimation and evaluation of M_q . Fixed master seeds, replication-level seeds, common random numbers, and moving-block bootstrap resampling are used to support reproducibility. Processed data, simulation outputs, and analysis code are available from the corresponding author upon reasonable request. The common computational settings and policy inputs used throughout the numerical analysis are summarized in Table 2.

Table 2. Common computational settings and policy inputs.

Component	Baseline specification
Evaluation horizon	$T = 360$ periods
Rolling estimation window	$W = 120$ periods
Monte Carlo replications	$N_{MC} = 200$ for final reported results
Action set	$a_t \in [0, 1]$
Damage function	$q(x) = x + 2(x - 2)_+$
Residual exposure	$R(a) = \exp(-1.5a)$
Implementation cost	$C(a) = 0.05a + 0.40a^2$
Adjustment cost	$H(a, a_-) = 0.10(a - a_-)^2$
Meaningful action tolerance	$\varepsilon_a = 0.025$
Rolling estimation	Poisson sample mean and NB method of moments
Period optimization	One-dimensional bounded minimization over $a_t \in [0, 1]$ of the common sequential adjustment objective
Randomization	Fixed replication-level seeds and common random numbers for matched comparisons

Notes: Pilot runs may use fewer replications, but all manuscript tables and figures are based on the final replication count. The policy classes differ only in the exposure m_t^π supplied to the common sequential decision problem.

7.2. *Balanced data-generating environments and empirical calibration*

The simulation design contains neutral, lower-risk, reverse, and adverse misspecification environments. This balance is important because the robust NB ambiguity rule moves the fitted model toward a higher mean and stronger overdispersion. Evaluating that rule only under adverse stress would not reveal the possible cost of unnecessary conservatism.

The empirical calibration anchor is the monthly S&P 500 jump-count series in Section 8. Let μ_0 and r_0 denote its full-sample NB method-of-moments mean and dispersion estimates. A persistent two-state regime process $Z_t \in \{0, 1\}$ distinguishes calm and stress periods. The baseline regime-specific parameters are

$$\mu_t^0 = \mu_0 b_\mu(Z_t), \quad r_t^0 = r_0 b_r(Z_t), \quad (7.13)$$

where $b_\mu(0) = 0.75$, $b_\mu(1) = 2.10$, $b_r(0) = 1.35$, and $b_r(1) = 0.55$. The transition probabilities are calibrated to represent persistent calm and stress states, with baseline values $\Pr(Z_t = 0 \mid Z_{t-1} = 0) = 0.94$ and $\Pr(Z_t = 1 \mid Z_{t-1} = 1) = 0.82$. The initial stress probability is set equal to the empirical frequency of months in which the jump count exceeds its empirical 80th percentile. The study considers the following five data-generating environments.

Correctly specified NB environment (CorrectNB). The true count law is

$$X_t \mid Z_t \sim \text{NB}(\mu_t^0, r_t^0). \quad (7.14)$$

The fitted family is therefore correctly specified, although finite rolling samples still produce estimation error. This environment evaluates the cost of robustification when no systematic adverse misspecification is imposed.

Near-Poisson environment (NearPoisson). The true mean retains the empirically calibrated regime pattern, but the dispersion parameter is set to a large value r_{NP} :

$$X_t | Z_t \sim \text{NB}(\mu_t^0, r_{NP}), \quad r_{NP} = 100. \quad (7.15)$$

Because $\text{Var}(X_t) = \mu_t^0 + (\mu_t^0)^2/r_{NP}$, this environment is close to Poisson at the empirical count scale. It evaluates whether robust NB creates unnecessary automation when overdispersion is weak.

Reverse NB misspecification (ReverseNB). This environment makes the true count law less risky than the nominal policy input. For severity parameter $\delta \in (0, 1)$ and reverse-dispersion multiplier $\gamma_R > 0$, the true parameters are

$$\mu_t^* = \mu_t^0(1 - \delta), \quad r_t^* = r_t^0(1 + \gamma_R\delta), \quad (7.16)$$

and

$$X_t | Z_t \sim \text{NB}(\mu_t^*, r_t^*). \quad (7.17)$$

To maintain a controlled reverse misspecification after rolling estimation, the policy-facing nominal NB parameters are obtained from the rolling estimate $(\widehat{\mu}_t, \widehat{r}_t)$ and shifted in the opposite direction:

$$\widehat{\mu}_t^{\text{nom}} = \widehat{\mu}_t(1 + \delta), \quad \widehat{r}_t^{\text{nom}} = \frac{\widehat{r}_t}{1 + \gamma_R\delta}. \quad (7.18)$$

The baseline specification uses $\gamma_R = 3.5$. This environment directly tests whether robustification compounds an already conservative nominal model, thereby increasing over-automation and implementation cost.

NB stress environment (NBStress). The true law remains within the NB family but moves toward higher mean and stronger overdispersion:

$$\mu_t^* = \mu_t^0(1 + 0.75\delta), \quad r_t^* = \frac{r_t^0}{1 + 3.5\delta}, \quad (7.19)$$

with $X_t | Z_t \sim \text{NB}(\mu_t^*, r_t^*)$. This is the adverse within-family environment retained from the original computational design.

Long-tail mixture environment (LongTailMix). The true law preserves the empirically calibrated NB core while adding an upper-tail component:

$$P_t^* = (1 - \omega_t) \text{NB}(\mu_{c,t}, r_{c,t}) + \omega_t \text{NB}(\mu_{h,t}, r_{h,t}), \quad (7.20)$$

where $(\mu_{c,t}, r_{c,t})$ follow the NBStress core, $\omega_t = \min\{0.35, 0.05 + \delta\}$, $\mu_{h,t} = \mu_{c,t}(1 + 6\delta + Z_t)$, and $r_{h,t} = r_{c,t}/(1 + 8\delta + 2Z_t)$. This environment evaluates adverse cross-shape misspecification relative to the fitted NB law. It is therefore treated as a stress data-generating process rather than as a separate Sichel policy benchmark.

For the severity-indexed environments, the baseline grid is $\delta \in \{0.10, 0.20, 0.30\}$. CorrectNB and NearPoisson are evaluated without an additional severity shift. The principal balanced comparison

includes CorrectNB, NearPoisson, ReverseNB at $\delta = 0.30$, NBStress at $\delta = 0.30$, and LongTailMix at $\delta = 0.30$. The full severity grid is retained for supplementary analysis.

In the principal comparison, robust NB policies with $\rho = 0.10$ and $\rho = 0.20$ are reported alongside the Poisson plug-in, NB plug-in, and oracle-informed policies. A separate ambiguity-radius experiment evaluates the broader grid $\rho \in \{0, 0.025, 0.05, 0.10, 0.15, 0.20, 0.30, 0.40\}$ under all five environments. Table 3 summarizes the direction, purpose, and expected robustness effect of each data-generating environment.

Table 3. Balanced data-generating environments for the computational experiments.

Environment	Direction relative to nominal NB	Purpose	Expected robustness effect
CorrectNB	Correct family with no systematic directional shift	Assess the cost of unnecessary protection under correct specification	A zero or small radius may minimize total cost
NearPoisson	Substantially weaker overdispersion	Evaluate robustness when the true law is close to equidispersion	Larger radii may increase over-automation with limited tail benefit
ReverseNB	Nominal mean and tail exposure exceed the true law	Test reverse misspecification and the cost of unnecessary robustness	Robustification may amplify over-automation and implementation cost
NBStress	Higher mean and stronger overdispersion	Evaluate adverse within-family misspecification	A positive radius may reduce under-automation and tail loss
LongTailMix	Heavier upper tail outside the fitted NB shape	Evaluate adverse cross-shape misspecification	A positive radius may provide the strongest tail protection

Notes: The expected effects are hypotheses implied by the directional analysis in Section 6; they are not imposed as conclusions. All policies are evaluated on matched count paths within each environment.

Before the final Monte Carlo runs, the data-generating-process implementation is checked through three diagnostics. First, the simulated mean and variance are compared with the analytical NB or mixture moments. Second, the ordering of the true risk-damage functionals is verified for the reverse, neutral, and adverse environments. Third, at $\rho = 0$, the robust NB policy is required to coincide numerically with the NB plug-in policy. These checks ensure that subsequent figures and tables reflect the intended misspecification directions rather than coding artifacts.

7.3. Main balanced policy comparison

The sequential policies are next compared across the five balanced environments introduced in Section 7.2. The comparison is constructed to evaluate both sides of distributional robustness. CorrectNB retains the correctly specified NB family while preserving finite-window parameter uncertainty. NearPoisson represents weak dispersion, and ReverseNB represents a conservative nominal model that overstates the current risk environment. NBStress and LongTailMix represent, respectively, within-family dispersion stress and upper-tail misspecification. The last two environments therefore favor protective correction only if the fitted model materially understates policy-relevant exposure; the first three prevent the experiment from mechanically favoring robustness.

Each environment includes the Poisson plug-in policy, the NB plug-in policy, robust NB policies with $\rho \in \{0.10, 0.20\}$, and the oracle-informed sequential benchmark. All policies are evaluated on identical simulated paths within each replication. The experiment uses $N_{MC} = 200$ replications, an evaluation horizon of $T = 360$, and a rolling estimation window of $W = 120$. Monte Carlo standard errors are computed across replication-level outcomes. Policy differences are computed from matched replication-level contrasts under common random numbers.

Table 4 reports the complete policy comparison. The Poisson benchmark performs poorly when the count process is strongly overdispersed or long-tailed. Under NBStress, its mean total cost is 354.30, compared with 324.05 for the NB plug-in policy. Under LongTailMix, the corresponding values are 675.91 and 581.78. These differences show that an equidispersed event-count model can materially understate the protective response required by clustered financial-risk events.

The robust policies improve performance when the fitted NB law understates risk. Under NBStress, total cost falls from 324.05 under NB plug-in to 320.70 for $\rho = 0.10$ and 318.55 for $\rho = 0.20$. Tail loss falls from 60.16 to 56.78 and 54.11, while under-automation declines from 36.92% to 30.21% and 24.03%. The benefit is larger under LongTailMix: Total cost falls from 581.78 to 571.36 and 563.21, and tail loss falls from 141.02 to 136.08 and 132.18. These results indicate that robustification is especially valuable when model error affects the upper tail rather than only the conditional mean.

The neutral and reverse environments reveal the cost of excessive prudence. Under NearPoisson, robustification lowers tail loss but raises total cost from 180.86 under NB plug-in to 181.37 and 182.71. Under ReverseNB, the increase is more pronounced: Total cost rises from 188.03 to 191.92 and 196.09, while the over-automation rate increases from 77.29% to 80.10% and 82.98%. Thus, a policy may reduce tail exposure and still be economically inferior because the additional implementation burden exceeds the protection benefit.

CorrectNB provides an intermediate case. The distributional family is correctly specified, but the decision-maker still estimates its parameters from a finite rolling window. Robustification therefore protects against estimation-induced tail understatement and produces modest total-cost reductions. This finding should not be interpreted as evidence that positive ambiguity is always optimal; rather, it distinguishes correct family specification from perfect knowledge of the current parameters.

Table 5 isolates the matched economic effect of robustness relative to NB plug-in. Under NBStress, $\rho = 0.20$ reduces total cost by 5.50 and tail loss by 6.04, despite raising implementation cost by 15.78. Under LongTailMix, the corresponding reductions are 18.57 and 8.84, with an additional implementation cost of 9.48. The residual loss reduction therefore more than compensates for the added intervention cost in the adverse environments.

The opposite pattern appears under NearPoisson and ReverseNB. For NearPoisson, $\rho = 0.20$ lowers tail loss by 1.05 but raises total cost by 1.85. For ReverseNB, it lowers tail loss by 2.40 but raises total cost by 8.06 and over-automation by 5.69 percentage points. These matched contrasts provide direct evidence for the directional result in Section 6.6: Robustness is beneficial when the nominal exposure is understated and costly when the nominal model is already sufficiently conservative.

Table 4. Main balanced comparison of sequential protective policies across correctly specified, low-dispersion, reverse-misspecified, NB stress, and long-tail financial count-risk environments.

Environment	Policy	Total cost	Tail loss	Oracle gap	Avg. automation	Switches	Under-auto. (%)	Over-auto. (%)
Correct NB	Poisson plug-in	276.51 (3.97)	55.40 (1.14)	40.73 (1.72)	0.545 (0.005)	16.39 (0.28)	64.01 (1.26)	26.13 (1.18)
Correct NB	NB plug-in	263.87 (3.41)	42.65 (0.84)	28.09 (1.29)	0.727 (0.006)	17.23 (0.22)	33.37 (0.93)	58.18 (1.06)
Correct NB	Robust NB, $\rho = 0.10$	262.79 (3.29)	39.86 (0.80)	27.01 (1.20)	0.771 (0.006)	17.18 (0.24)	26.85 (0.76)	63.59 (0.88)
Correct NB	Robust NB, $\rho = 0.20$	262.61 (3.18)	37.56 (0.77)	26.82 (1.11)	0.809 (0.006)	16.75 (0.27)	22.41 (0.66)	68.10 (0.72)
Correct NB	Oracle-informed benchmark	235.78 (2.61)	35.78 (0.79)	0.00 (0.00)	0.668 (0.002)	48.41 (0.64)	0.00 (0.00)	0.00 (0.00)
Near-Poisson	Poisson plug-in	180.84 (1.13)	9.64 (0.19)	10.82 (0.26)	0.554 (0.003)	1.14 (0.03)	25.90 (0.46)	69.70 (0.64)
Near-Poisson	NB plug-in	180.86 (1.13)	9.39 (0.18)	10.84 (0.26)	0.571 (0.003)	3.12 (0.12)	25.53 (0.46)	71.17 (0.58)
Near-Poisson	Robust NB, $\rho = 0.10$	181.37 (1.12)	8.84 (0.17)	11.35 (0.25)	0.612 (0.003)	3.99 (0.15)	24.72 (0.44)	73.81 (0.48)
Near-Poisson	Robust NB, $\rho = 0.20$	182.71 (1.12)	8.34 (0.16)	12.69 (0.26)	0.651 (0.003)	5.06 (0.15)	24.16 (0.43)	74.89 (0.44)
Near-Poisson	Oracle-informed benchmark	170.02 (1.01)	7.21 (0.15)	0.00 (0.00)	0.534 (0.002)	62.05 (0.74)	0.00 (0.00)	0.00 (0.00)
Reverse NB	Poisson plug-in	185.32 (2.46)	28.78 (0.69)	23.17 (1.02)	0.430 (0.003)	10.23 (0.24)	40.48 (1.03)	43.73 (1.25)
Reverse NB	NB plug-in	188.03 (2.21)	19.10 (0.41)	25.89 (0.84)	0.710 (0.006)	16.77 (0.21)	19.78 (0.46)	77.29 (0.47)
Reverse NB	Robust NB, $\rho = 0.10$	191.92 (2.19)	17.79 (0.39)	29.78 (0.82)	0.757 (0.006)	17.04 (0.20)	16.62 (0.45)	80.10 (0.46)
Reverse NB	Robust NB, $\rho = 0.20$	196.09 (2.15)	16.70 (0.37)	33.95 (0.79)	0.798 (0.006)	16.94 (0.21)	13.56 (0.45)	82.98 (0.46)
Reverse NB	Oracle-informed benchmark	162.14 (1.64)	18.03 (0.43)	0.00 (0.00)	0.500 (0.002)	61.11 (0.76)	0.00 (0.00)	0.00 (0.00)
NB stress	Poisson plug-in	354.30 (6.39)	81.37 (1.85)	57.24 (2.65)	0.606 (0.007)	21.91 (0.29)	75.24 (1.28)	17.24 (1.04)
NB stress	NB plug-in	324.05 (5.02)	60.16 (1.36)	26.99 (1.47)	0.805 (0.007)	18.46 (0.34)	36.92 (1.42)	50.37 (1.27)
NB stress	Robust NB, $\rho = 0.10$	320.70 (4.81)	56.78 (1.30)	23.64 (1.29)	0.841 (0.006)	17.76 (0.39)	30.21 (1.28)	55.94 (1.19)
NB stress	Robust NB, $\rho = 0.20$	318.55 (4.64)	54.11 (1.27)	21.49 (1.15)	0.870 (0.006)	16.79 (0.41)	24.03 (1.17)	60.48 (1.05)
NB stress	Oracle-informed benchmark	297.06 (4.21)	53.67 (1.35)	0.00 (0.00)	0.777 (0.001)	48.26 (0.65)	0.00 (0.00)	0.00 (0.00)
Long-tail mixture	Poisson plug-in	675.91 (33.15)	182.18 (10.84)	144.62 (14.84)	0.737 (0.008)	20.14 (0.37)	71.42 (1.54)	19.15 (1.14)
Long-tail mixture	NB plug-in	581.78 (26.68)	141.02 (8.79)	50.49 (8.30)	0.894 (0.005)	13.44 (0.40)	40.92 (1.63)	40.32 (1.25)
Long-tail mixture	Robust NB, $\rho = 0.10$	571.36 (25.87)	136.08 (8.54)	40.07 (7.13)	0.915 (0.005)	12.18 (0.42)	34.65 (1.55)	44.89 (1.27)
Long-tail mixture	Robust NB, $\rho = 0.20$	563.21 (25.21)	132.18 (8.33)	31.92 (6.14)	0.932 (0.004)	11.02 (0.41)	28.78 (1.44)	49.13 (1.21)
Long-tail mixture	Oracle-informed benchmark	531.30 (22.48)	119.13 (7.48)	0.00 (0.00)	0.970 (0.000)	33.27 (0.46)	0.00 (0.00)	0.00 (0.00)

Notes: Entries are Monte Carlo means with standard errors in parentheses. All policies solve the same sequential one-period adjustment problem and differ only in the exposure functional used for decision-making. The oracle-informed benchmark observes the current true count law but not future count realizations. In financial applications, automation intensity can represent hedging, exposure reduction, leverage control, liquidity buffering, margin tightening, or defensive trading intervention.

Table 5. Matched economic effects of robust NB automation relative to the NB plug-in policy.

Environment	ρ	Δ total cost	Residual-loss reduction	Tail-loss reduction	Δ implementation cost	Δ adjustment cost	Δ under-auto. (pp)	Δ over-auto. (pp)
Correct NB	0.10	-1.08 (0.15)	11.16 (0.14)	2.79 (0.05)	10.08 (0.12)	0.006 (0.000)	-6.52 (0.40)	5.41 (0.42)
Correct NB	0.20	-1.26 (0.30)	20.38 (0.26)	5.09 (0.09)	19.10 (0.24)	0.011 (0.000)	-10.96 (0.57)	9.91 (0.69)
Near-Poisson	0.10	0.51 (0.03)	7.17 (0.06)	0.55 (0.01)	7.68 (0.06)	0.004 (0.000)	-0.82 (0.18)	2.65 (0.33)
Near-Poisson	0.20	1.85 (0.06)	13.67 (0.12)	1.05 (0.02)	15.51 (0.13)	0.009 (0.000)	-1.38 (0.24)	3.72 (0.46)
Reverse NB	0.10	3.90 (0.08)	6.82 (0.07)	1.31 (0.03)	10.70 (0.09)	0.006 (0.000)	-3.16 (0.20)	2.81 (0.19)
Reverse NB	0.20	8.06 (0.17)	12.52 (0.13)	2.40 (0.05)	20.57 (0.18)	0.012 (0.000)	-6.22 (0.34)	5.69 (0.31)
NB stress	0.10	-3.35 (0.25)	11.87 (0.22)	3.37 (0.07)	8.52 (0.17)	0.004 (0.000)	-6.72 (0.40)	5.57 (0.40)
NB stress	0.20	-5.50 (0.47)	21.29 (0.41)	6.04 (0.13)	15.78 (0.34)	0.008 (0.000)	-12.89 (0.60)	10.11 (0.64)
Long-tail mixture	0.10	-10.42 (1.27)	15.67 (1.26)	4.93 (0.42)	5.25 (0.19)	0.001 (0.000)	-6.27 (0.47)	4.57 (0.36)
Long-tail mixture	0.20	-18.57 (2.36)	28.05 (2.35)	8.84 (0.78)	9.48 (0.36)	0.002 (0.001)	-12.14 (0.79)	8.81 (0.55)

Notes: Entries are matched Monte Carlo means with standard errors in parentheses, computed from identical simulated paths under common random numbers. Δ denotes robust NB minus NB plug-in, except residual-loss and tail-loss reductions, which are NB plug-in minus robust NB. Thus, negative Δ total cost and positive loss reductions favor robustness. Percentage-point changes are reported for under- and over-automation. In financial applications, implementation cost can represent the expense of additional hedging, liquidity reserves, tighter margins, lower leverage, or more defensive trading controls; over-automation captures excessive intervention relative to the oracle-informed response.

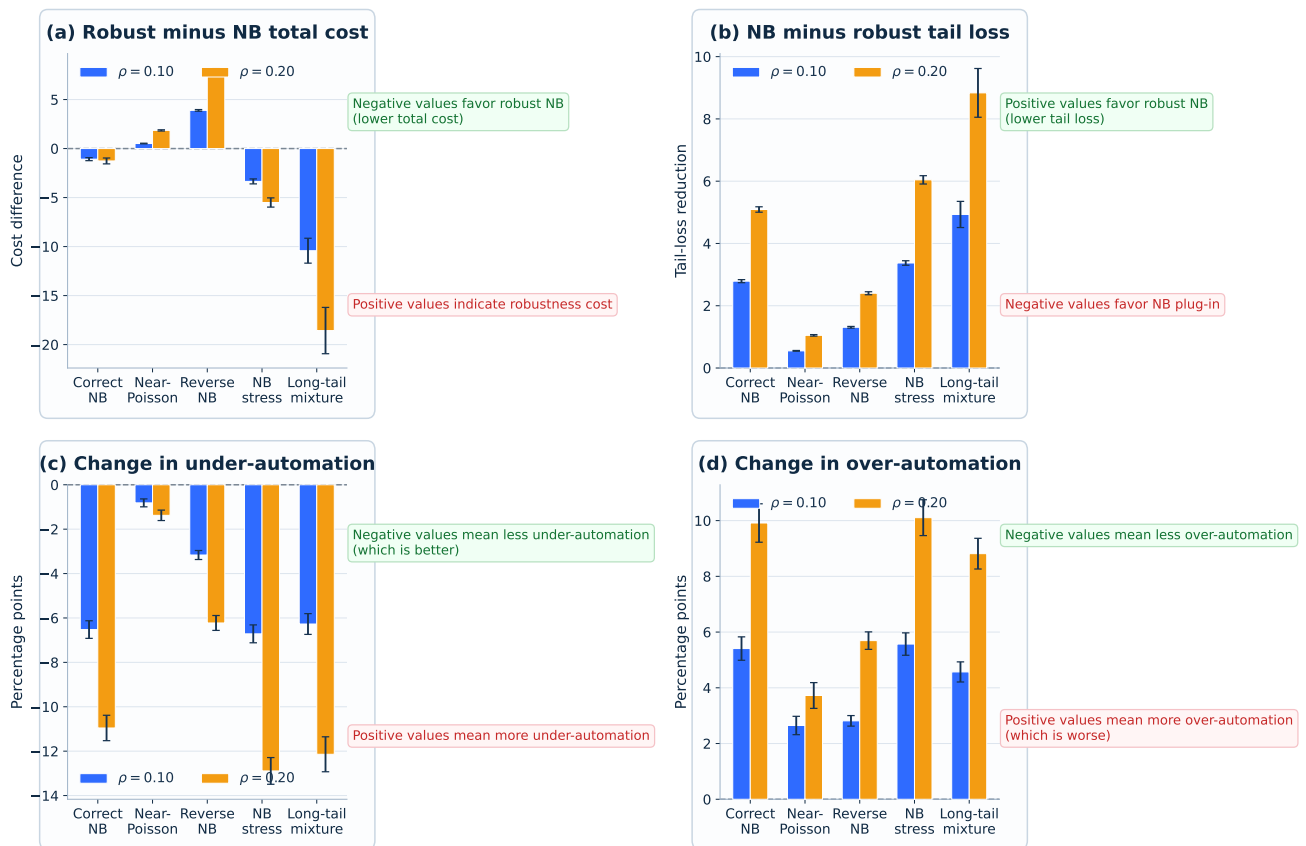


Figure 1. Matched robustness–performance trade-off for robust NB automation relative to the NB plug-in policy. Panels report (a) robust-minus-NB total cost, (b) NB-minus-robust tail loss, (c) the change in under-automation frequency, and (d) the change in over-automation frequency. Error bars are Monte Carlo standard errors of replication-level matched differences computed under common random numbers. Negative values in panel (a), positive values in panel (b), and negative values in panel (c) favor robustness. Positive values in panel (d) indicate additional intervention relative to the oracle-informed response. The caption therefore distinguishes tail protection from the economic cost of over-intervention.

Figure 1 summarizes the central tradeoff. The signs of the matched cost effects reverse across environments: Robust NB raises total cost under NearPoisson and ReverseNB, but lowers it under NBStress and LongTailMix. Tail loss decreases in every environment because robustification raises the exposure input and hence the selected automation intensity, consistently with Proposition 6.4. However, the economic value of that protection depends on whether the true environment justifies the additional intervention. Robustification also lowers under-automation while increasing over-automation, making explicit that protection and excess intervention are two sides of the same policy response.

From a financial-risk-management perspective, the adverse environments represent circumstances in which a rolling count model understates the frequency, clustering, or upper-tail severity of extreme market events. The robust action can then be interpreted as a timely increase in hedging, liquidity protection, margin discipline, or exposure reduction. By contrast, NearPoisson and ReverseNB represent calm or conservatively calibrated conditions in which the same response may create unnecessary hedging

costs, capital usage, liquidity holdings, or lost trading opportunities. The results therefore support calibrated rather than maximal robustness.

7.4. Computational findings and implementation implications

The balanced experiment establishes four findings. First, the Poisson policy is fragile when discrete financial-risk events are clustered or long-tailed, whereas the NB plug-in policy provides a substantially stronger nominal benchmark. Second, robust NB automation has positive economic value when the nominal model understates the risk-damage functional, especially under upper-tail misspecification. Third, robustification is not uniformly beneficial: Under weak dispersion or reverse misspecification, lower tail loss can be accompanied by higher total cost and more over-automation. Fourth, the ambiguity radius acts as an economically meaningful prudence parameter rather than a purely statistical tuning constant.

These findings broaden the use of the framework beyond a single market index. The count process may represent extreme-return days, volatility or liquidity breaches, margin events, clustered defaults, insurance claims, cyber incidents, or operational losses. Correspondingly, the action may represent automated hedging, leverage reduction, liquidity buffering, margin adjustment, reserve activation, or escalation of monitoring and intervention. The structural results remain unchanged across these applications because only the policy-relevant exposure input and the economic interpretation of the action must be adapted.

The empirical illustration in Section 8 complements these controlled experiments by applying the same sequential policy architecture to observed S&P 500 jump counts. In that setting, the true count law and the oracle action are unavailable, so the evaluation focuses on realized costs, tail protection, policy paths, and the empirical sensitivity of the protective response to the ambiguity radius.

Together, the five environments and the ambiguity-radius grid cover the main directions needed for sensitivity evaluation: Correct specification, near-equidispersion, reverse misspecification, within-family overdispersion stress, and cross-shape upper-tail stress. The reported metrics also separate statistical protection from economic cost by recording total cost, residual loss, tail loss, implementation cost, adjustment cost, switching, under-automation, over-automation, and oracle-relative gaps. Thus, the simulation study is not designed to show dominance of the robust policy; it is designed to identify when robustness is economically useful and when it becomes inefficient.

8. Financial application: Jump-count-driven protective control

In this section, we illustrate how the proposed robust NB framework converts a discrete financial-risk signal into sequential protective action. The application uses monthly counts of extreme S&P 500 return days. Its role is not to claim external validation for every operational setting, but to provide a transparent out-of-sample policy illustration in an environment where event counts are observable, clustered, and economically consequential.

In this application, the control $a_m \in [0, 1]$ represents the intensity of an automated protective response. Depending on the institution, this can mean stronger hedging, exposure reduction, leverage control, liquidity buffering, margin tightening, defensive trading restrictions, or escalation of automated risk monitoring. The empirical analysis retains the sequential one-period policy architecture used in the theory and computational experiments, with $a_m^\pi \in \arg \min_{a \in [0, 1]} \{R(a)m_m^\pi + C(a) + H(a, a_{m-1}^\pi)\}$. No policy

uses a Bellman continuation value or an estimated transition kernel. This makes the implementation transparent, auditable, and directly aligned with the structural results in Section 6.

The analysis proceeds in four stages: construction of monthly extreme-return counts, comparison of Poisson and NB count models, examination of rolling dispersion and policy-relevant exposure, and out-of-sample evaluation of plug-in and robust policies. The final step reports not only point estimates but also moving-block-bootstrap uncertainty, thereby separating statistically stable tail protection from more modest and less precisely estimated total-cost gains.

8.1. Financial jump-count construction

Let P_t denote the daily adjusted closing level of the S&P 500 index on trading day t , and define the daily log return as $r_t = \log P_t - \log P_{t-1}$. A jump day is identified by $J_t = \mathbf{1}\{|r_t| > q_\alpha\}$, where q_α is the empirical α -quantile of $|r_t|$. The baseline uses $\alpha = 0.97$, and the monthly jump count is $X_m = \sum_{t \in m} J_t$.

The sample covers January 1990 through May 2026 and contains 437 monthly observations. Daily adjusted closing prices are first converted to log returns, and the absolute-return threshold is computed once from the full daily sample only for constructing the event definition. All rolling model estimates and policy decisions use only information available before the evaluated month, so the out-of-sample policy comparison does not use future counts in the decision rule. The estimated threshold is $q_{0.97} = 0.02698$. The monthly count series has mean 0.629, variance 3.027, sample variance-to-mean ratio 4.811, and maximum count 18. These values provide direct evidence of overdispersion and show why an equidispersed Poisson benchmark is potentially inadequate for protective-control design.

The transformation has a clear financial interpretation. The count X_m summarizes the frequency of extreme market days in month m , while a_m controls the strength of the next protective response. A larger count can therefore justify stronger hedging, lower gross exposure, more conservative leverage limits, larger liquidity buffers, tighter margins, or stronger defensive trading restrictions.

8.2. Rolling estimation and policy-relevant instability

Using a 120-month rolling window, Poisson and NB models are fitted, denoted by $\widehat{P}_w^P$ and $\widehat{P}_w^{NB}$. For each window, the analysis computes the sample variance-to-mean ratio (VMR) and the fitted policy-relevant exposure $\widehat{M}_q^{(w)} = \mathbb{E}_{\widehat{P}_w}[q(X)]$, where $q(x) = x + \kappa(x - \tau)_+$, with $\tau = 2$ and $\kappa = 2$. The rolling variance-to-mean ratio is $\text{VMR}_w = \widehat{\text{Var}}_w(X)/\widehat{\mathbb{E}}_w[X]$.

Figure 2 reports the resulting diagnostic series. Figure 2(a) displays the monthly jump-count path in lollipop form and highlights the largest observed clusters, showing that extreme-return days arrive episodically rather than at a stable rate. Figure 2(b) plots the rolling VMR and shades the excess-dispersion region above the Poisson benchmark of one; the series remains persistently overdispersed and changes materially over time. Figure 2(c) plots the fitted Poisson and NB risk-damage exposures and shades the gap between them. The NB exposure $\widehat{M}_q^{NB}$ is consistently larger, particularly during high-clustering periods. Thus, the empirical issue is not only goodness of fit, but also the magnitude and temporal stability of the risk-damage input supplied to the sequential optimization problem.

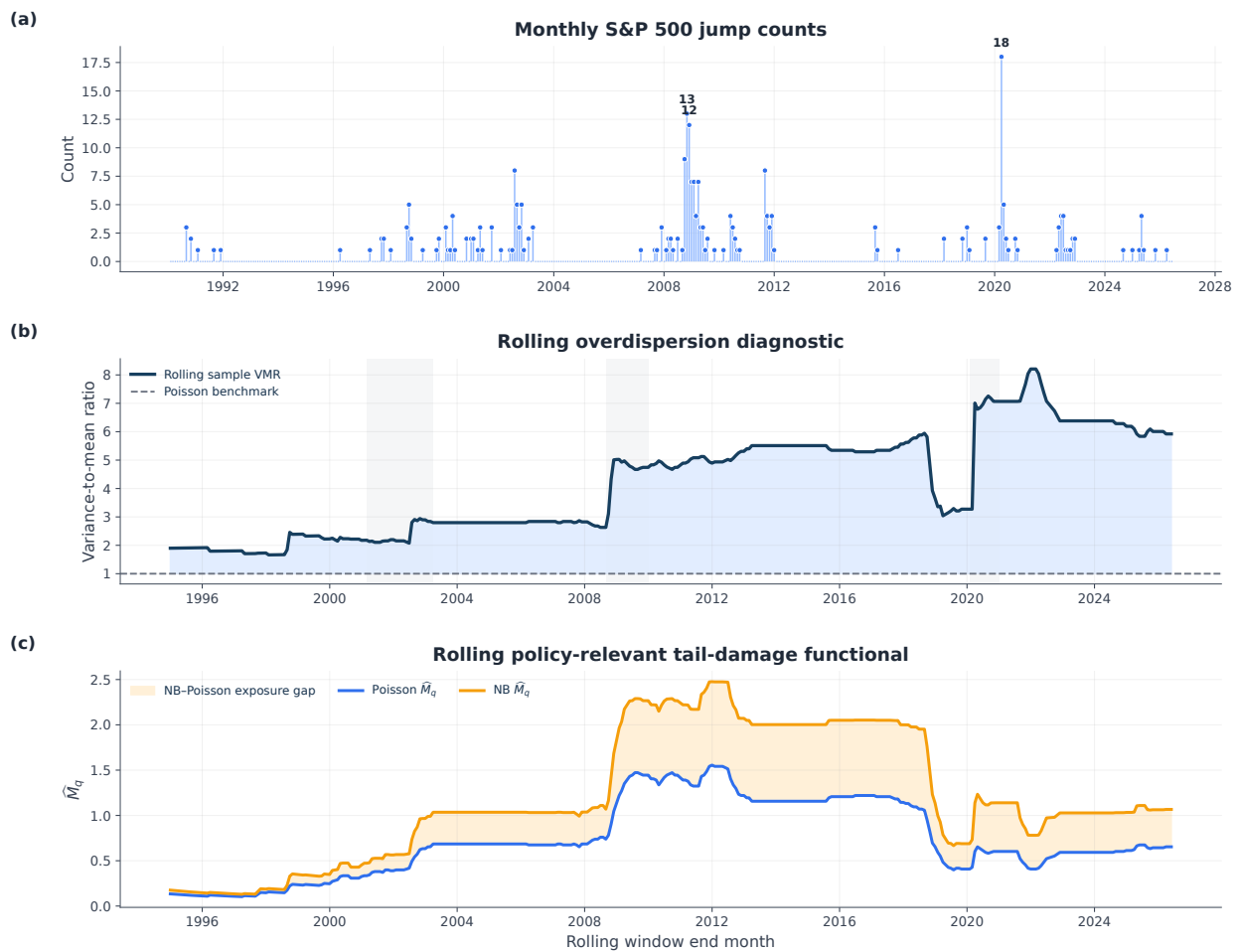


Figure 2. Rolling financial count-risk diagnostics for Standard & Poor's 500 (S&P 500) jump counts. (a) Monthly jump counts constructed from the 97th percentile of absolute daily log returns; the largest clusters are labeled. (b) Rolling 120-month variance-to-mean ratio, with the shaded region showing excess dispersion above the Poisson benchmark of one. (c) Rolling policy-relevant tail-damage exposures under fitted Poisson and NB models; the shaded band reports the NB-minus-Poisson exposure gap.

8.3. Poisson and NB model comparison

Table 6 compares the full-sample Poisson and NB fits using the maximized log-likelihood and the AIC and BIC defined in Equation (4.14).

Table 6. Empirical count-model comparison for monthly S&P 500 jump counts.

Model	Log-likelihood	AIC	BIC	Mean	VMR	$\widehat{M}_q$
Poisson	-628.6896	1259.3791	1263.4590	0.6293	1.0000	0.6905
NB	-416.4397	836.8793	845.0392	0.6293	4.6193	1.1622

Notes: The Poisson model is the equidispersed benchmark. The NB model is the nominal overdispersed law used for policy construction. The sample variance-to-mean ratio is 4.811; the fitted NB variance-to-mean ratio is 4.619.

The NB model improves the log-likelihood by more than 212 log units and reduces AIC and BIC by approximately 422.5 and 418.4, respectively. Its fitted variance-to-mean ratio is also close to the sample variance-to-mean ratio. Importantly, its fitted risk-damage exposure is 1.1622, compared with only 0.6905 under the Poisson benchmark. The empirical policy comparison is therefore deliberately demanding: It asks whether local robustification can improve decisions even after the nominal model has already been upgraded from Poisson to a well-fitting NB specification.

8.4. Out-of-sample sequential policy evaluation

At each month m , the count model is estimated using only the preceding 120 months. The policies are π^P , π^{NB} , and $\pi^{R,NB}(\rho)$, with $\rho \in \{0.025, 0.05, 0.10, 0.15, 0.20, 0.30\}$. The robust policy uses the adverse-direction parameter corner $\mu_m^{wc} = \widehat{\mu}_m(1 + \rho)$ and $r_m^{wc} = \widehat{r}_m/(1 + \rho)$, which raises the mean while reducing the NB dispersion parameter. For the increasing tail-damage functional used here, this parameter shift produces the policy-relevant adverse exposure identified in the structural analysis.

For each policy, realized total cost is $\text{Cost}(\pi) = \sum_m [R(a_m^\pi)q(X_m) + C(a_m^\pi) + H(a_m^\pi, a_{m-1}^\pi)]$, and realized tail loss is $\text{TailLoss}(\pi) = \sum_m R(a_m^\pi)(X_m - \tau)_+$. The analysis also reports average control intensity and the number of meaningful action changes satisfying $|a_m^\pi - a_{m-1}^\pi| > \varepsilon_a$. Standard errors in Table 7 are obtained from a 12-month moving-block bootstrap with 1,000 replications.

Table 7 yields three main findings. First, the Poisson policy is both less protective and more costly than the NB benchmark: Total cost is 4.55% higher, tail loss is 49.62 rather than 39.64, and average control is only 0.564 rather than 0.731. Second, increasing ρ raises average control monotonically from 0.744 to 0.844 and reduces residual and tail losses. Third, the gain is much stronger for tail protection than for total cost. At $\rho = 0.30$, tail loss falls from 39.64 to 32.86, whereas total cost falls from 252.26 to 247.25. The number of meaningful action changes also falls from 28 to 24, so the stronger protection is not produced by excessive switching.

To distinguish economic magnitude from sampling precision, Table 8 reports matched percentage reductions relative to the NB plug-in policy together with percentile 95% moving-block bootstrap intervals.

Table 7. Out-of-sample empirical policy performance.

Policy	Total cost	Residual loss	Implementation cost	Adjustment cost	Tail loss	Avg. control	Switches	Cost vs. NB (%)
Poisson plug-in	263.73 (69.10)	211.86 (69.43)	51.86 (4.34)	0.009 (0.002)	49.62 (20.09)	0.564 (0.027)	18.00 (6.13)	4.55 (5.06)
NB plug-in	252.26 (54.71)	169.35 (55.35)	82.90 (6.33)	0.016 (0.004)	39.64 (16.01)	0.731 (0.031)	28.00 (8.89)	0.00 (0.00)
Robust NB, $\rho = 0.025$	251.76 (53.63)	166.17 (54.27)	85.58 (6.46)	0.017 (0.004)	38.89 (15.70)	0.744 (0.031)	27.00 (8.88)	-0.20 (0.41)
Robust NB, $\rho = 0.05$	251.29 (52.59)	163.11 (53.24)	88.16 (6.56)	0.017 (0.004)	38.17 (15.40)	0.756 (0.032)	27.00 (8.88)	-0.38 (0.81)
Robust NB, $\rho = 0.10$	250.60 (50.64)	157.50 (51.30)	93.09 (6.70)	0.018 (0.004)	36.84 (14.83)	0.780 (0.031)	27.00 (8.95)	-0.66 (1.58)
Robust NB, $\rho = 0.15$	250.07 (48.83)	152.61 (49.53)	97.44 (6.70)	0.019 (0.004)	35.67 (14.32)	0.800 (0.031)	26.00 (8.63)	-0.87 (2.30)
Robust NB, $\rho = 0.20$	249.27 (47.26)	148.25 (47.95)	101.00 (6.54)	0.021 (0.005)	34.64 (13.85)	0.818 (0.030)	27.00 (8.59)	-1.18 (2.89)
Robust NB, $\rho = 0.30$	247.25 (44.68)	140.76 (45.28)	106.46 (5.89)	0.023 (0.005)	32.86 (13.07)	0.844 (0.026)	24.00 (7.50)	-1.99 (3.78)

Notes: Entries are point estimates with moving-block-bootstrap standard errors in parentheses. Cost vs. NB is the percentage change in realized total cost relative to the NB plug-in policy. Higher control can represent stronger hedging, exposure reduction, leverage limits, liquidity buffers, margin requirements, or defensive trading restrictions.

Table 8. Empirical ambiguity-radius trade-off with 95% moving-block-bootstrap intervals.

ρ	Total-cost reduction vs. NB (%)	Tail-loss reduction vs. NB (%)	Average control	Switches
0.025	0.20 [−0.77, 0.79]	1.90 [1.79, 2.03]	0.744	27
0.05	0.38 [−1.54, 1.54]	3.71 [3.53, 3.96]	0.756	27
0.10	0.66 [−3.14, 2.90]	7.07 [6.66, 7.52]	0.780	27
0.15	0.87 [−4.66, 4.16]	10.02 [9.06, 10.72]	0.800	26
0.20	1.18 [−5.68, 5.34]	12.63 [11.01, 13.59]	0.818	27
0.30	1.99 [−6.95, 7.37]	17.11 [14.22, 18.51]	0.844	24

Notes: Reductions are measured relative to the NB plug-in policy; positive values favor robust NB. Brackets report percentile 95% intervals from 1,000 moving-block-bootstrap replications with block length 12.

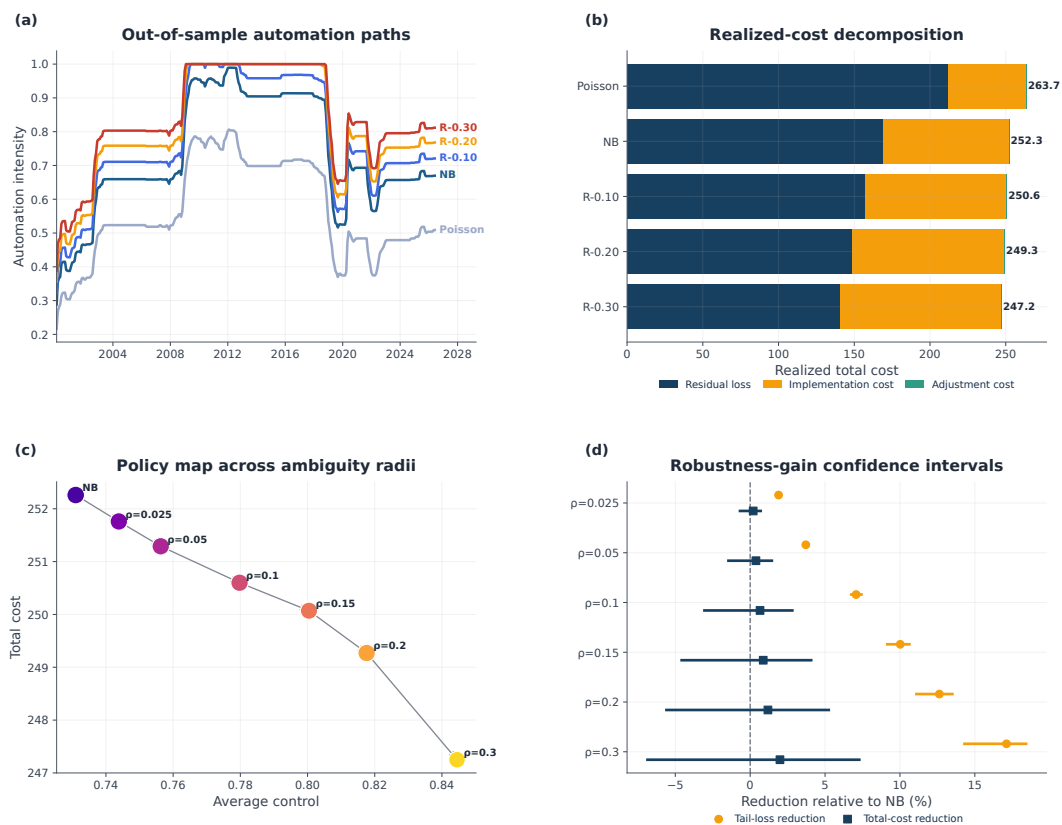


Figure 3. Out-of-sample financial policy evaluation. (a) Sequential protective-control paths under Poisson plug-in, NB plug-in, and robust NB policies, with direct labels at the path endpoints. (b) Realized-cost decomposition into residual loss, implementation cost, and adjustment cost; labels report total realized cost. (c) Policy map showing how increasing ambiguity radii move the policy from the NB plug-in benchmark in the average-control–total-cost plane; marker size reflects switching frequency. (d) Point estimates and percentile 95% moving-block-bootstrap intervals for tail-loss and total-cost reductions relative to the NB plug-in policy. The figure separates policy trajectories, cost decomposition, robustness–cost trade-offs, and bootstrap uncertainty.

The tail-loss intervals are strictly positive at every radius, rising from 1.90% [1.79, 2.03] at $\rho = 0.025$ to 17.11% [14.22, 18.51] at $\rho = 0.30$. Tail protection is therefore both economically meaningful and statistically stable in this empirical path. By contrast, the total-cost point estimates improve monotonically but their intervals include zero. The proper interpretation is consequently not that robust NB establishes statistically conclusive total-cost dominance, but that it delivers a clear tail-protection benefit while producing modest positive total-cost point estimates over the evaluated radius grid.

Figure 3(a) confirms the theoretical ordering: Stronger ambiguity adjustment produces stronger protective action, especially in periods of elevated estimated exposure. Figure 3(b) shows the economic mechanism: Residual loss falls as implementation expenditure rises, while adjustment cost remains very small. Figure 3(c) recasts the radius comparison as a policy map. Moving from the NB plug-in point toward larger ρ raises average control and lowers the total-cost point estimate; the marker sizes show that this movement is not accompanied by a systematic increase in switching frequency. Figure 3(d) separates point estimates from uncertainty. Every tail-loss interval lies above zero, whereas every total-cost interval crosses zero. The visual evidence therefore supports strong and statistically stable tail protection, but only a modest and statistically inconclusive total-cost advantage on this single empirical path.

Accordingly, the empirical evidence should be interpreted as strong evidence of improved tail protection, but not as statistically conclusive evidence of lower average total cost on this single financial series.

The numerical implementation also verifies the structural consistency of the empirical policies. The cost-accounting error is zero to numerical precision; there are no violations of monotonicity in the adverse exposure, and, when the previous action is held fixed, the first-order-condition solver produces no action-monotonicity violations. Average control is nondecreasing in ρ , and both tail loss and total-cost point estimates are nonincreasing over the tested radius grid. These checks confirm that the empirical code implements the same sequential model analyzed theoretically.

8.5. Dual industrial implications and transferability

The application supports two distinct channels of industrial impact. The first is direct financial implementation. Asset managers can interpret a_m as a hedge ratio, gross-exposure reduction, volatility-control intensity, or cash buffer. Banks and clearing institutions can interpret it as liquidity reserve, collateral tightening, margin escalation, leverage restriction, or counterparty exposure reduction. Trading systems can interpret it as a defensive execution mode, inventory limit, or monitoring escalation. In each case, the ambiguity radius makes the degree of prudence explicit and auditable rather than embedding it implicitly in an opaque risk score.

The second channel is cross-industry transfer. The same structure applies when X_t counts insurance claims, credit defaults, cyberattacks, operational-loss events, equipment failures, service outages, supply disruptions, safety violations, or adverse healthcare events. The substantive action changes—for example, reserves, inspections, preventive maintenance, staffing, isolation, or escalation—but the optimization logic is unchanged:

uncertain clustered event counts $\longrightarrow$ tail-relevant exposure $\longrightarrow$ calibrated protective intervention.

The managerial lesson is therefore two-sided. When the nominal law understates clustering or upper-tail exposure, robustification can materially reduce insufficient intervention and tail loss. When the nominal law is accurate or already conservative, the balanced experiments in Section 7 show that the

same mechanism can create over-automation and unnecessary cost. The framework does not recommend maximizing robustness mechanically; it provides a disciplined way to select protection intensity by comparing the marginal value of tail-loss reduction with the marginal cost of intervention.

9. Discussion and conclusions

The results support a decision-oriented interpretation of overdispersed count modeling. A fitted count law is not merely a forecasting device: It determines the policy-relevant exposure supplied to the automation problem. A Poisson model can understate clustered risk because it imposes equidispersion, whereas a fitted NB law captures excess variance but may still be locally misspecified in its mean, dispersion, or upper tail. The proposed robust policy treats that fitted NB law as a transparent nominal model rather than as known truth. This positioning is consistent with the broader distributionally robust optimization literature, in which decisions are protected against ambiguity around an estimated distribution [22, 23, 63]. A broader synthesis is provided by Rahimian and Mehrotra [64]. It also complements the NB automation model of Cha et al. [29]: That work established the structural role of overdispersion in threshold and hysteresis behavior, whereas the present study asks how a fitted policy should be protected when the nominal count law is uncertain.

9.1. What the balanced experiments establish

The computational evidence is deliberately two-sided. In the NB-stress and long-tail environments, the nominal law understates policy-relevant exposure; robustification then raises protection, reduces under-automation, lowers tail loss, and moves policy performance toward the oracle. In the near-Poisson and reverse-NB environments, however, the nominal law is accurate or already conservative. In those settings, a larger ambiguity radius can increase over-automation and total cost. Robustness is therefore not uniformly beneficial, and the results do not support choosing the largest feasible radius mechanically. They instead confirm the directional logic of the theory: Robustification corrects errors when the nominal model is biased toward under-protection, but can amplify errors when the nominal model is already pessimistic.

This distinction responds directly to a central concern in robust optimization: A method should be evaluated not only under perturbations aligned with its ambiguity set, but also under neutral and reverse perturbations. The balanced design shows that the robust policy is not favored by construction. Its value depends on the relationship between model error and the marginal cost of protection. Related work on moment and dispersion ambiguity similarly shows that uncertainty in distributional summaries changes the optimal protection level rather than guaranteeing uniform improvement [24, 25].

The ambiguity radius should consequently be interpreted as a managerial protection parameter. A practical calibration rule is to select the smallest radius that achieves an acceptable tail loss or under-automation target while keeping implementation expenditure and over-automation within operational limits. Scenario analysis, rolling backtests, and confidence regions for fitted count parameters can all inform this choice. When the application places a hard constraint on tail exposure, liquidity, capital, safety, or service continuity, a larger radius may be justified even when its average-cost gain is uncertain. When intervention cost dominates and the fitted model is stable, a smaller radius is more defensible. Extensions based on statistically calibrated moment, optimal-transport, or Sinkhorn ambiguity sets could link the radius more directly to sample size and estimation uncertainty [68–70].

9.2. Empirical interpretation and practical impact

The S&P 500 analysis is an out-of-sample financial illustration, not a claim of universal external validity for all operational systems. Monthly jump counts are appropriate here because they summarize clustered episodes in which protective actions may be escalated. In this setting, the automation variable can represent a composite intensity of rule-based hedging, exposure reduction, leverage control, liquidity buffering, margin tightening, defensive trading, or automated risk-monitoring escalation. The empirical analysis therefore illustrates how a discrete financial risk signal can be converted into a calibrated protective response.

The empirical findings should also be read with a precise distinction between tail protection and average cost. The sample exhibits strong overdispersion, and the NB model fits substantially better than the Poisson benchmark. Robustification then produces monotone tail-loss reductions across the evaluated radius grid, with moving-block-bootstrap intervals remaining above zero. By contrast, total-cost point estimates improve more modestly, and their intervals include zero. The strongest empirical conclusion is therefore that robust NB automation provides stable tail protection; the net average-cost advantage is favorable in point estimates but statistically less conclusive. This distinction prevents the empirical illustration from being overstated and makes the protection–cost trade-off explicit.

The industrial relevance is broader than the financial example, although the action must be interpreted application by application. The same structure can map insurance claims to reserve or reinsurance actions, defaults to credit or collateral controls, cyber incidents to monitoring and isolation, equipment failures to preventive maintenance, service disruptions to staffing or capacity buffers, and adverse healthcare events to escalation protocols. The common architecture is

clustered event counts $\longrightarrow$ tail-relevant exposure $\longrightarrow$ calibrated protective intervention.

This transferability is conceptual rather than empirical: Each industry requires its own count construction, loss function, implementation cost, and radius calibration.

9.3. Computational tractability, limitations, and future work

The implemented policy remains computationally transparent. Conditional on a fitted or worst-case exposure, each period requires the solution of a one-dimensional strictly convex problem on $[0, 1]$. With the exponential residual-exposure and quadratic cost specification, the optimizer is obtained from boundary checks and a scalar root solve. For T periods and K candidate radii, policy evaluation is therefore linear in TK apart from count-model fitting and evaluation of the tail-damage functional. This simplicity makes the method auditable and suitable for rolling implementation, while avoiding the state-space growth associated with robust dynamic programming or ambiguous transition kernels.

Several limitations remain. First, the ambiguity set is local and parametric; it does not represent every form of dependence, regime uncertainty, sampling error, structural break, or tail-shape misspecification. Severe nonlocal misspecification or abrupt changes in the count-generating process may therefore require adaptive ambiguity sets, online learning, Bayesian robust estimation, or regime-switching extensions. Second, the theory relies on compactness, continuity, monotonicity, and convexity. Nonconvex implementation choices, multiple interacting controls, capacity constraints, or high-dimensional automation actions may require different methods. Third, the empirical illustration contains one financial count series and does not observe an oracle policy. Additional applications to claims, failures, cyber events, hospital demand, environmental risk, or service disruptions are needed to

assess external validity. Fourth, direct numerical comparisons with CVaR, entropic-risk, Wasserstein-DRO, robust Markov decision process, reinforcement-learning, or Bayesian robust policies are outside the controlled comparison adopted here, because those approaches change the objective, ambiguity object, learning structure, or transition model. Such comparisons are worthwhile future work when richer datasets permit fair calibration across methods. Finally, flexible count models such as zero-inflated, hurdle, generalized Poisson, Conway–Maxwell–Poisson, and long-tailed mixed-Poisson families may be valuable when their additional parameters can be estimated stably. Long-tail stress environments are instead used to test the robustness of a parsimonious NB policy.

In conclusion, the study develops a sequential ambiguity-aware policy layer for overdispersed count risk. The theory establishes a unique period decision, an explicit activation condition, monotone response to ambiguity-adjusted exposure, and persistence induced by adjustment friction. The experiments show both sides of robustness: material protection against underestimated clustered risk and measurable over-intervention cost when the nominal model is accurate or conservative. The empirical illustration reinforces the same message: Robustness yields strong and stable tail protection, while total-cost gains are more modest. The appropriate policy is therefore not “maximum robustness,” but calibrated robustness—a radius selected to balance tail protection, intervention expenditure, and the credibility of the fitted count law.

Author contributions

Jinho Cha, Long Pham, Minh Ngo Thi, Ho Thi Hien, and Hai Thi Bich Vu contributed equally to the conceptualization, methodology, formal analysis, software, validation, investigation, data curation, visualization, writing of the original draft, and review and editing of the manuscript. All authors contributed to the mathematical modeling, computational implementation, numerical experiments, empirical analysis, and interpretation of the results. All authors reviewed and approved the final version of the manuscript.

Use of generative-AI tools declaration

The authors declare that generative-AI tools were used only to assist with language editing, formatting refinement, and drafting support during manuscript preparation. All modeling design, computational analysis, mathematical derivations, interpretation of results, and final scientific judgments were performed and verified by the authors. The authors take full responsibility for the content of this article.

Data availability statement

The empirical data used in this study are publicly available from Yahoo Finance. Daily adjusted closing prices for the S&P 500 index were obtained from this public source and processed for the empirical analysis. The computational experiments use synthetic data generated within the proposed model. The processed datasets, simulation outputs, and analysis code, including scripts for the theoretical computations, simulations, empirical policy evaluation, and figure generation, are available from the corresponding author upon reasonable request.

Conflict of interest

The authors declare no conflict of interest.

References

1. R. F. Engle, Autoregressive conditional heteroscedasticity with estimates of the variance of united kingdom inflation, *Econometrica*, **50** (1982), 987–1007. <https://doi.org/10.2307/1912773>
2. T. Bollerslev, Generalized autoregressive conditional heteroskedasticity, *J. Econom.*, **31** (1986), 307–327. [https://doi.org/10.1016/0304-4076\(86\)90063-1](https://doi.org/10.1016/0304-4076(86)90063-1)
3. T. Bollerslev, V. Todorov, Tails, fears, and risk premia, *J. Finance*, **66** (2011), 2165–2211. <https://doi.org/10.1111/j.1540-6261.2011.01695.x>
4. T. Bollerslev, R. Y. Chou, K. F. Kroner, ARCH modeling in finance: A review of the theory and empirical evidence, *J. Econom.*, **52** (1992), 5–59. [https://doi.org/10.1016/0304-4076\(92\)90064-X](https://doi.org/10.1016/0304-4076(92)90064-X)
5. G. G. Castellano, Don't call it a failure: Systemic risk governance for complex financial systems, *Law Soc. Inq.*, **49** (2024), 2245–2286. <https://doi.org/10.1017/lsi.2024.8>
6. M. Caporin, L. Garcia-Jorcano, J. A. Jimenez-Martin, Early warnings of systemic risk using one-minute high-frequency data, *Expert Syst. Appl.*, **252** (2024), 124134. <https://doi.org/10.1016/j.eswa.2024.124134>
7. D. R. Bergmann, M. A. Oliveira, Extreme risk clustering in long-memory financial series, *Chaos Solitons Fractals*, **202** (2026), 117513. <https://doi.org/10.1016/j.chaos.2025.117513>
8. J. A. Hausman, B. H. Hall, Z. Griliches, Econometric models for count data with an application to the patents–R&D relationship, *Econometrica*, **52** (1984), 909–938. <https://doi.org/10.2307/1911191>
9. A. C. Cameron, P. K. Trivedi, Regression-based tests for overdispersion in the poisson model, *J. Econom.*, **46** (1990), 347–364. [https://doi.org/10.1016/0304-4076\(90\)90014-K](https://doi.org/10.1016/0304-4076(90)90014-K)
10. R. Berk, J. M. MacDonald, Overdispersion and poisson regression, *J. Quant. Criminol.*, **24** (2008), 269–284. <https://doi.org/10.1007/s10940-008-9048-4>
11. H. S. Sichel, On a distribution representing sentence-length in written prose, *J. R. Stat. Soc. Ser. A Stat. Soc.*, **137** (1974), 25–34. <https://doi.org/10.2307/2345142>
12. H. S. Sichel, Repeat-buying and the generalized inverse gaussian–poisson distribution, *J. R. Stat. Soc. Ser. C Appl. Stat.*, **31** (1982), 193–204. <https://doi.org/10.2307/2347993>
13. R. A. Rigby, D. M. Stasinopoulos, C. Akantziliotou, A framework for modelling overdispersed count data, including the poisson-shifted generalized inverse gaussian distribution, *Comput. Stat. Data Anal.*, **53** (2008), 381–393. <https://doi.org/10.1016/j.csda.2008.07.043>
14. A. K. Dixit, R. S. Pindyck, *Investment under Uncertainty*, Princeton University Press, 1994. <https://doi.org/10.2307/j.ctt7sncv>
15. R. J. Caballero, On the sign of the investment-uncertainty relationship, *Am. Econ. Rev.*, **81** (1991), 279–288. <https://www.jstor.org/stable/2006800>
16. A. W. Lo, Long-term memory in stock market prices, *Econometrica*, **59** (1991), 1279–1313. <https://doi.org/10.2307/2938368>

17. M. K. Brunnermeier, Y. Sannikov, A macroeconomic model with a financial sector, *Am. Econ. Rev.*, **104** (2014), 379–421. <https://doi.org/10.1257/aer.104.2.379>
18. R. Almgren, N. Chriss, Optimal execution of portfolio transactions, *J. Risk*, **3** (2001), 5–39. <https://doi.org/10.21314/JOR.2001.041>
19. Á. Cartea, S. Jaimungal, J. Penalva, *Algorithmic and high-frequency trading*, Cambridge University Press, 2015. <https://books.google.com/books?id=5dMmCgAAQBAJ>
20. L. W. Cong, Z. He, Blockchain disruption and smart contracts, *Rev. Financ. Stud.*, **32** (2019), 1754–1797. <https://doi.org/10.1093/rfs/hhz007>
21. L. P. Hansen, T. J. Sargent, Robust control and model uncertainty, *Am. Econ. Rev.*, **91** (2001), 60–66. <https://doi.org/10.1257/aer.91.2.60>
22. E. Delage, Y. Ye, Distributionally robust optimization under moment uncertainty with application to data-driven problems, *Oper. Res.*, **58** (2010), 595–612. <https://doi.org/10.1287/opre.1090.0741>
23. P. Mohajerin Esfahani, D. Kuhn, Data-driven distributionally robust optimization using the wasserstein metric: Performance guarantees and tractable reformulations, *Math. Program.*, **171** (2018), 115–166. <https://doi.org/10.1007/s10107-017-1172-1>
24. K. Postek, A. Ben Tal, D. Den Hertog, B. Melenberg, Robust optimization with ambiguous stochastic constraints under mean and dispersion information, *Oper. Res.*, **66** (2018), 814–833. <https://doi.org/10.1287/opre.2017.1688>
25. L. Chen, C. Fu, F. Si, M. Sim, P. Xiong, Robust optimization with moment-dispersion ambiguity, *Oper. Res.*, **73** (2024), 3118–3138. <https://doi.org/10.1287/opre.2023.0579>
26. G. N. Iyengar, Robust dynamic programming, *Math. Oper. Res.*, **30** (2005), 257–280. <https://doi.org/10.1287/moor.1040.0129>
27. H. Xu, S. Mannor, Distributionally robust markov decision processes, *Math. Oper. Res.*, **37** (2012), 288–300. <https://doi.org/10.1287/moor.1120.0540>
28. V. Goyal, J. Grand-Clément, Robust markov decision processes: Beyond rectangularity, *Math. Oper. Res.*, **48** (2023), 203–226. <https://doi.org/10.1287/moor.2022.1259>
29. J. Cha, S. M. Han, L. Pham, J. Mollick, J. Yu, Optimal automation under overdispersed discrete risk: Thresholds and hysteresis in a Negative Binomial model, *J. Ind. Manag. Optim.*, **22** (2026), 2503–2554. <https://doi.org/10.3934/jimo.2026092>
30. S. Jeong, J. Jeon, H. K. Koo, Optimal investment under irreversible consumption and locally risk-seeking preferences, *J. Ind. Manag. Optim.*, **22** (2026), 1063–1086. <https://doi.org/10.3934/jimo.2026039>
31. F. Peng, R. Shi, Z. Zhang, S. Zhang, The irreversible green technology investment strategy of manufacturers under the Cap-and-Trade mechanism, *J. Ind. Manag. Optim.*, **22** (2026), 100–147. <https://doi.org/10.3934/jimo.2026005>
32. T. Jin, Y. Yan, H. Shen, Production and inventory rationing in an unreliable Make-to-Stock System under preventive maintenance policy, *J. Ind. Manag. Optim.*, **22** (2026), 880–910. <https://doi.org/10.3934/jimo.2026032>

33. P. Wang, G. Wang, Y. Yang, W. Mu, Equilibrium strategy in a non-performing loan securitization game with regime switching, *J. Ind. Manag. Optim.*, **22** (2026), 691–715. <https://doi.org/10.3934/jimo.2026025>
34. I. Y. Chen, C. Y. Lin, Y. T. Peng, Evaluating the optimal digital transformation strategy using hierarchical sensitivity decision model: A case study of tourism industry, *J. Ind. Manag. Optim.*, **22** (2026), 422–449. <https://doi.org/10.3934/jimo.2026016>
35. O. E. Barndorff-Nielsen, N. Shephard, Power and Bipower Variation with Stochastic Volatility and Jumps, *J. Financ. Econom.*, **2** (2004), 1–37. <https://doi.org/10.1093/jjfinec/nbh001>
36. Y. Ait-Sahalia, J. Jacod, Testing for Jumps in a discretely observed process, *Ann. Stat.*, **37** (2009), 184–222. <https://doi.org/10.1214/07-AOS568>
37. C. B. Dean, E. R. Lundy, Overdispersion, in *Wiley statsRef: statistics reference online*, John Wiley & Sons, 2016. <https://doi.org/10.1002/9781118445112.stat06788.pub2>
38. M. Denuit, J. Dhaene, M. Goovaerts, R. Kaas, *Actuarial theory for dependent risks: Measures, orders and models*, John Wiley & Sons, Chichester, 2005. <https://doi.org/10.1002/0470016450>
39. Y. Zou, L. Wu, D. Lord, Modeling over-dispersed crash data with a long tail: Examining the accuracy of the dispersion parameter in negative binomial models, *Anal. Methods Accid. Res.*, **5** (2015), 1–16. <https://doi.org/10.1016/j.amar.2014.12.002>
40. T. N. Sindhu, G. S. S. Abdalla, A. Shafiq, T. Abushal, A new statistical framework for overdispersed count data: Applications in public health, radiation dosimetry and finance, *J. Radiat. Res. Appl. Sci.*, **18** (2025), 101987. <https://doi.org/10.1016/j.jrras.2025.101987>
41. T. N. Sindhu, E. M. Almetwally, A. Shafiq, Z. Hussain, T. A. Abushal, A. Aldukeel, A compound poisson–sujit framework for over-dispersed radiation monitoring data, *J. Radiat. Res. Appl. Sci.*, **19** (2026), 102333. <https://doi.org/10.1016/j.jrras.2026.102333>
42. T. N. Sindhu, A. Shafiq, T. A. Abushal, A. A. Alkhathami, The poisson–garima distribution: Additional key features and its significance on statistical process control in agriculture, *CSIAM Trans. Appl. Math.*, **7** (2026), 352–382. <https://doi.org/10.4208/csiam-am.SO-2025-0010>
43. T. N. Sindhu, A. Shafiq, H. M. Hamouda, T. A. Abushal, Y. M. Ayid, A one-parameter poisson–emrem model for count data: Theory and applications to cytogenetic and COVID-19 studies, *Sci. Rep.*, 2026. <https://doi.org/10.1038/s41598-026-57818-2>
44. T. N. Sindhu, Y. El Khatib, T. A. Abushal, A. S. Alharthi, Modeling biological count data with the poisson–mgamma distribution: applications to radiation and yeast cell analysis, *Results Eng.*, **29** (2026), 109193. <https://doi.org/10.1016/j.rineng.2026.109193>
45. T. N. Sindhu, A. Shafiq, Y. El Khatib, T. A. Abushal, A. S. Alharthi, A length-biased sujit distribution framework: Properties, simulation-based inference, and application to clinical remission data, *Sci. Rep.*, **16** (2026), 14857. <https://doi.org/10.1038/s41598-026-42402-5>
46. T. N. Sindhu, A. Shafiq, A. Atangana, T. A. Abushal, A. A. Alkhathami, Control charts for overdispersed count data: Exploring the poisson chris–jerry distribution in agriculture and medicine, *Qual. Reliab. Eng. Int.*, **41** (2025), 1913–1935. <https://doi.org/10.1002/qre.3745>
47. A. S. Kyle, A. A. Obizhaeva, Y. Wang, Smooth trading with overconfidence and market power, *Rev. Econ. Stud.*, **85** (2018), 611–662. <https://doi.org/10.1093/restud/rdx017>

48. B. Biais, T. Foucault, S. Moinas, Equilibrium high frequency trading, available at SSRN, 2011. <https://ssrn.com/abstract=1834344>
49. T. Foucault, R. Kozhan, W. W. Tham, Toxic arbitrage, *Rev. Financ. Stud.*, **30** (2017), 1053–1094. <https://doi.org/10.1093/rfs/hhw103>
50. D. Easley, M. M. L. D. Prado, M. O'Hara, Flow toxicity and liquidity in a high-frequency world, *Rev. Financ. Stud.*, **25** (2012), 1457–1493. <https://doi.org/10.1093/rfs/hhs053>
51. D. Ladley, The high frequency trade-off between speed and sophistication, *J. Econ. Dyn. Control*, **116** (2020), 103912. <https://doi.org/10.1016/j.jedc.2020.103912>
52. B. Biais, C. Bisiere, M. Bouvard, C. Casamatta, The blockchain folk theorem, *Rev. Financ. Stud.*, **32** (2019), 1662–1715. <https://doi.org/10.1093/rfs/hhy095>
53. J. Chiu, T. V. Koepl, Blockchain-based settlement for asset trading, *Rev. Financ. Stud.*, **32** (2019), 1716–1753. <https://doi.org/10.1093/rfs/hhy122>
54. P. P. Momtaz, Decentralized finance markets for startups: Search frictions, intermediation, and the efficiency of the ICO market, *Small Bus. Econ.*, **63** (2024), 1415–1447. <https://doi.org/10.1007/s11187-024-00886-3>
55. K. Pantelidis, Active tokens and crypto-asset valuation, *Financ. Innov.*, **11** (2025), 95. <https://doi.org/10.1186/s40854-025-00752-5>
56. A. Goldfarb, C. Tucker, Digital economics, *J. Econ. Lit.*, **57** (2019), 3–43. <https://doi.org/10.1257/jel.20171452>
57. F. Schär, Decentralized finance: On blockchain- and smart contract-based financial markets, *Fed. Reserve Bank St. Louis Rev.*, **103** (2021), 153–174. <https://doi.org/10.20955/r.103.153-74>
58. R. Bansal, A. Yaron, Risks for the long run: A potential resolution of asset pricing puzzles, *J. Finance*, **59** (2004), 1481–1509. <https://doi.org/10.1111/j.1540-6261.2004.00670.x>
59. J. D. Hamilton, A new approach to the economic analysis of nonstationary time series and the business cycle, *Econometrica*, **57** (1989), 357–384. <https://doi.org/10.2307/1912559>
60. A. Ang, G. Bekaert, International asset allocation with regime shifts, *Rev. Financ. Stud.*, **15** (2002), 1137–1187. <https://doi.org/10.1093/rfs/15.4.1137>
61. E. W. Anderson, L. P. Hansen, T. J. Sargent, A quartet of semigroups for model specification, robustness, prices of risk, and model detection, *J. Eur. Econ. Assoc.*, **1** (2003), 68–123. <https://doi.org/10.1162/154247603322256774>
62. F. Barillas, L. P. Hansen, T. J. Sargent, Doubts or variability? *J. Econ. Theory*, **144** (2009), 2388–2418. <https://doi.org/10.1016/j.jet.2008.11.014>
63. W. Wiesemann, D. Kuhn, M. Sim, Distributionally robust convex optimization, *Oper. Res.*, **62** (2014), 1358–1376. <https://doi.org/10.1287/opre.2014.1314>
64. H. Rahimian, S. Mehrotra, Frameworks and results in distributionally robust optimization, *Open J. Math. Optim.*, **3** (2022), 1–85. <https://doi.org/10.5802/ojmo.15>
65. Z. Chen, M. Sim, H. Xu, Distributionally robust optimization with infinitely constrained ambiguity sets, *Oper. Res.*, **67** (2019), 1328–1344. <https://doi.org/10.1287/opre.2018.1799>

66. C. Peng, E. Delage, Data-driven optimization with distributionally robust second order stochastic dominance constraints, *Oper. Res.*, **72** (2024), 1298–1316. <https://doi.org/10.1287/opre.2022.2387>
67. J. Grand-Clément, M. Petrik, N. Vieille, Beyond discounted returns: Robust markov decision processes with average and blackwell optimality, *Oper. Res.*, 2026, 1–25. <https://doi.org/10.1287/opre.2023.0694>
68. S. Jiang, J. Cheng, K. Pan, Z. J. M. Shen, Optimized dimensionality reduction for moment-based distributionally robust optimization, *Oper. Res.*, 2025, 1–25. <https://doi.org/10.1287/opre.2023.0645>
69. S. Shafiee, L. Aolaritei, F. Dörfler, D. Kuhn, Nash equilibria, regularization, and computation in optimal transport-based distributionally robust optimization, *Oper. Res.*, 2025, 1–25. <https://doi.org/10.1287/opre.2023.0138>
70. J. Wang, R. Gao, Y. Xie, Sinkhorn distributionally robust optimization, *Oper. Res.*, 2025, 1–25. <https://doi.org/10.1287/opre.2023.0294>
71. S. X. Zhang, X. A. Sun, On distributionally robust multistage convex optimization: Data-driven models and performance, *INFORMS J. Optim.*, 2025, 1–25. <https://doi.org/10.1287/ijoo.2024.0049>
72. A. N. Elmachtoub, P. Grigas, Smart “Predict, then optimize”, *Manag. Sci.*, **68** (2022), 9–26. <https://doi.org/10.1287/mnsc.2020.3922>
73. H. Akaike, A new look at the statistical model identification, *IEEE Trans. Autom. Control*, **19** (1974), 716–723. <https://doi.org/10.1109/TAC.1974.1100705>
74. G. Schwarz, Estimating the dimension of a model, *Ann. Stat.*, **6** (1978), 461–464. <https://doi.org/10.1214/aos/1176344136>



AIMS Press

©2026 the Author(s), licensee AIMS Press. This is an open access article distributed under the terms of the Creative Commons Attribution License (<https://creativecommons.org/licenses/by/4.0>)