

A General Framework for Empirical Bayes Estimation in Discrete Linear Exponential Family

Trambak Banerjee

trambak@ku.edu

Analytics, Information and Operations Management
University of Kansas
Lawrence, KS 66045, USA

Qiang Liu

lqiang@cs.utexas.edu

Computer Science
University of Texas at Austin
Austin, Texas 78712, USA

Gourab Mukherjee

gourab@usc.edu

Data Sciences and Operations
University of Southern California
Los Angeles, CA 90089, USA

Wenguang Sun

wenguans@marshall.usc.edu

Data Sciences and Operations
University of Southern California
Los Angeles, CA 90089, USA

Editor: Jon McAuliffe

Abstract

We develop a Nonparametric Empirical Bayes (NEB) framework for compound estimation in the discrete linear exponential family, which includes a wide class of discrete distributions frequently arising from modern big data applications. We propose to directly estimate the Bayes shrinkage factor in the generalized Robbins' formula via solving a convex program, which is carefully developed based on a RKHS representation of the Stein's discrepancy measure. The new NEB estimation framework is flexible for incorporating various structural constraints into the data driven rule, and provides a unified approach to compound estimation with both regular and scaled squared error losses. We develop theory to show that the class of NEB estimators enjoys strong asymptotic properties. Comprehensive simulation studies as well as analyses of real data examples are carried out to demonstrate the superiority of the NEB estimator over competing methods.

Keywords: Asymptotic Optimality; Empirical Bayes; Power Series Distributions; Shrinkage estimation; Stein's discrepancy

1. Introduction

Shrinkage methods, exemplified by the seminal work of James and Stein (1961), have received renewed attention in modern large-scale inference problems (Efron, 2012; Fourdrinier et al., 2018). Under this setting, the classical Normal means problem has been extensively studied (Brown, 2008; Jiang and Zhang, 2009; Brown and Greenshtein, 2009; Efron, 2011; Xie et al., 2012; Weinstein et al., 2018). However, in a variety of applications, the observed data are often discrete. For instance, in the News Popularity study discussed in Section 5, the goal is to estimate the popularity of a large number of news items based on their frequencies of being shared in social media platforms such as Facebook and LinkedIn. Another important application scenario arises from genomics research, where estimating the expected number of mutations across a large number of genomic locations can help identify key drivers or inhibitors of a given phenotype of interest.

We mention two main limitations of existing shrinkage estimation methods. First, the methodology and theory developed for continuous variables, in particular for Normal means problem, may not be directly applicable to discrete models. Second, existing methods have focused on the squared error loss. However, the scaled loss (Clevenson and Zidek, 1975), which effectively reflects the asymmetries in decision making [cf. Equation (3)], is a more desirable choice for many discrete models such as Poisson, where the scaled loss corresponds to the local Kulback-Leibler distance. The scaled loss also provides a more desirable criterion in a range of sparse settings, for example, when the goal is to estimate the rates of rare outcomes in Binomial distributions (Fourdrinier and Robert, 1995). Much research is needed for discrete estimation problems under various loss functions. This article develops a general framework for empirical Bayes estimation for the discrete linear exponential (DLE) family, also known as the family of discrete power series distributions (Noack, 1950), under both regular and scaled squared error losses.

The DLE family includes a wide class of popular members such as the Poisson, Binomial, Negative Binomial and Geometric distributions. Let Y be a non-negative integer valued random variable. Then Y is said to belong to a DLE family if its probability mass function (pmf) is of the form

$$p(y|\sqrt{\cdot}) = \frac{a_y}{g(\sqrt{\cdot})^y}, \quad y \in \{0, 1, 2, \dots\}, \quad (1)$$

where a_y and $g(\sqrt{\cdot})$ are known functions such that $a_y > 0$ is independent of $\sqrt{\cdot}$ and $g(\sqrt{\cdot})$ is a normalizing factor that is differentiable at every $\sqrt{\cdot}$. Special cases of DLE include the Poisson(λ) distribution with $a_y = (\lambda^y / y!)$, $\sqrt{\cdot} = \lambda$ and $g(\sqrt{\cdot}) = \exp(\sqrt{\cdot})$, and the Binomial(m, q) distribution with $a_y = \binom{m}{y} q^y (1-q)^{m-y}$, $\sqrt{\cdot} = q/(1-q)$ and $g(\sqrt{\cdot}) = (1 + \sqrt{\cdot})^m$. Suppose $Y_1, \dots, Y_n$ obey the following hierarchical model

$$Y_i | \sqrt{\cdot}_i \stackrel{\text{i.i.d.}}{\sim} \text{DLE}(\sqrt{\cdot}_i), \quad \sqrt{\cdot}_i \stackrel{\text{i.i.d.}}{\sim} G(\cdot), \quad (2)$$

where $G(\cdot)$ is an unspecified prior distribution on $\sqrt{\cdot}_i$. The problem of interest is to estimate $\sqrt{\cdot} = (\sqrt{\cdot}_1, \dots, \sqrt{\cdot}_n)$ based on $Y = (Y_1, \dots, Y_n)$. Empirical Bayes (EB) approaches to this compound decision problem date back to the famous Robbins' formula (Robbins, 1956) under the Poisson model. In the terminology of Efron (2014, 2019) there are two main modeling strategies for such EB estimation, namely, g-modeling and f-modeling strategies. The main

goal under g-modeling is to model the prior distribution G of $\sqrt{\cdot}$ using, for example, Non-parametric Maximum Likelihood estimation (NPMLE) techniques (Kiefer and Wolfowitz, 1956; Laird, 1978) or by modeling G as a low dimensional exponential family (Efron, 2016). With an estimate $\hat{G}$ of G , one can then derive an estimate of $\sqrt{\cdot}$ by plugging $\hat{G}$ into the Bayes rule for various loss functions (see for example Jiang and Zhang (2009); Koenker and Mizera (2014); Gu and Koenker (2017)). The f-modeling strategy, on the other hand, first starts from a particular form of the Bayes rule and then directly estimates the unknown marginal pmf $p(\cdot)$ of Y using, for instance, the observed empirical frequencies (Robbins, 1956), the smoothness-adjusted estimator of Brown et al. (2013), kernel density estimation techniques (Brown and Greenshtein, 2009) or through maximum likelihood estimation in flexible exponential family models (Efron, 2012).

This article develops a general non-parametric empirical Bayes (NEB) framework for compound estimation in discrete models. We first derive generalized Robbins' formula (GRF) for the DLE model (2), and then implement GRF via solving a convex program which is carefully developed based on a reproducing kernel Hilbert space (RKHS) representation of Stein's discrepancy measure and leads to a class of efficient NEB shrinkage estimators. Our work is related to the aforementioned f-modeling strategy however, in contrast with existing f-modeling approaches that estimate $p(y)$, the proposed NEB estimation framework directly produces estimates of Bayes shrinkage factors that are ratios of the marginal pmf $p(y)$ and appear in the GRF for the DLE model (2). We develop theories to show that the NEB estimator is $\sqrt{p}/\sqrt{n}$ consistent up to certain logarithmic factors and enjoys superior risk properties. Simulation studies are conducted to illustrate that the efficiency gain of the NEB estimator over existing approaches, such as Brown et al. (2013), Koenker and Mizera (2014); Koenker and Gu (2017), Efron (2016), is substantial in many settings.

There are several advantages of the proposed NEB estimation framework. First, in contrast with existing methods such as the smoothness-adjusted Poisson estimator in Brown et al. (2013), our methodology covers a much wider range of distributions and presents a unified approach to compound estimation in discrete models. Second, our proposed convex program directly produces stable estimates of optimal Bayes shrinkage factors and can easily incorporate various structural constraints into the decision rule. By contrast, the three-step estimator in Brown et al. (2013), which involves smoothing, Rao-Blackwellization and monotonicity adjustments, is complicated, computationally intensive and sometimes unstable (as the numerator and denominator of the ratio are computed separately). Third, the RKHS representation of Stein's discrepancy measure provides a new analytical tool for developing theories such as asymptotic optimality and convergence rates. Finally, the NEB estimation framework is robust to departures from the true model due to its utilization of a generic quadratic program that does not rely on the specific form of a particular DLE family. Our numerical results in Section 4 demonstrate that the NEB estimator has a better risk performance than competitive approaches of Efron (2011), Brown et al. (2013) and Efron (2016) under a mis-specified Poisson model.

An alternative approach to compound estimation in discrete models, as suggested and investigated by Brown et al. (2013), is to employ variance stabilizing transformations, which converts the discrete problem to a classical normal means problem. This allows estimation via Tweedie's formula for normal variables (Efron, 2011). However, there are several drawbacks of this approach compared to our NEB framework. First, Tweedie's formula is not

applicable to scaled error loss whereas our methodology is built upon the generalized Robbins' formula, which covers both regular and scaled squared error losses. Second, there can be information loss in conventional data processing steps such as standardization, transformation and continuity approximation. While investigating the impact of information loss on compound estimation is of great interest, it is desirable to develop methodologies directly based on generalized Robbins' formula that is specifically derived and tailored for discrete variables. Finally, our NEB framework provides a convenient tool for developing asymptotic theories. By contrast, convergence rates are yet to be developed for normality inducing transformations, which can be highly non-trivial.

The rest of the paper is organized as follows. In Section 2, we introduce our estimation framework while Section 3 presents a theoretical analysis of the NEB estimator. The numerical performance of our method is investigated using both simulated and real data in Sections 4 and 5 respectively. Section 6 concludes with a discussion. Additional technical details and proofs are relegated to the Appendices.

2. A General Framework for Compound Estimation in DLE Family

This section describes the proposed NEB framework for compound estimation in discrete models. We first introduce in Section 2.1 the generalized Robbins' formula for the DLE family (2), then propose in Section 2.2 a convex optimization approach for its practical implementation. Details regarding tuning parameter selection are discussed in Section 2.3.

2.1 Generalized Robbins' formula for DLE models

Denote $\hat{\theta} = (\hat{\theta}_1, \dots, \hat{\theta}_n)$ to be an estimator of θ based on Y . Consider a class of loss functions

$$\ell^{(k)}(\hat{\theta}_i, \theta_i) = \sqrt{\theta_i}^{-k} (\hat{\theta}_i - \theta_i)^2 \quad (3)$$

for $k \in \{0, 1\}$, where $\ell^{(0)}(\hat{\theta}_i, \theta_i)$ is the usual squared error loss, and $\ell^{(1)}(\hat{\theta}_i, \theta_i) = \sqrt{\theta_i}^{-1} (\hat{\theta}_i - \theta_i)^2$ corresponds to the scaled squared error loss (Clevenson and Zidek, 1975; Fourdrinier and Robert, 1995). In compound estimation, one is concerned with the average loss

$$L_n^{(k)}(\hat{\theta}, \theta) = \frac{1}{n} \sum_{i=1}^n \ell^{(k)}(\hat{\theta}_i, \theta_i).$$

The associated risk is denoted $R_n^{(k)}(\hat{\theta}, \theta) = E_Y \ell^{(k)}(\hat{\theta}, \theta)$. Let $G(\theta)$ denote the joint distribution of $(\theta_1, \dots, \theta_n)$. The Bayes estimator $\hat{\theta}^{(k)}$ that minimizes the Bayes risk $R_n^{(k)}(\hat{\theta}, \theta)$ is given by Lemma 1.

Lemma 1 (Generalized Robbins' formula). Consider the DLE Model (2). Let $p(\cdot) = p(\cdot | \theta)$ be the marginal pmf of Y . Define for $k \in \{0, 1\}$,

$$w_p^{(k)}(y_i) = \frac{p(y_i - k)}{p(y_i + 1 - k)}, \text{ for } y_i = k, k + 1, \dots.$$

Then the Bayes estimator that minimizes the risk $B_n^{(k)}(\check{\gamma})$ is given by $\hat{\gamma}_{(k)} = \{ \hat{\gamma}_{(k),i}(y_i) : 1 \leq i \leq n \}$, where

$$\hat{\gamma}_{(k),i}(y_i) = \begin{cases} \frac{a_{y_i-k}/a_{y_i+1-k}}{w_p^{(k)}(y_i)}, & \text{for } y_i = k, k+1, \dots \\ 0, & \text{for } y_i < k \end{cases}. \quad (4)$$

Remark 1. Under the squared error loss ($k = 0$) with $Y_i | \check{\gamma}_i \leftarrow \text{Poi}(\check{\gamma}_i)$ and $a_{y_i} = (y_i!)^{-1}$, Lemma 1 yields

$$\hat{\gamma}_{(0),i}(y_i) = (y_i + 1) \frac{p(y_i + 1)}{p(y_i)}, \quad (5)$$

which recovers the classical Robbins' formula (Robbins, 1956). In contrast, under the scaled loss, we have

$$\hat{\gamma}_{(1),i}(y_i) = y_i \frac{p(y_i)}{p(y_i - 1)} \text{ for } y_i > 0 \text{ and } \hat{\gamma}_{(1),i}(y_i) = 0 \text{ otherwise.} \quad (6)$$

Under scaled error loss the estimator (5) can be much outperformed by (6) (and vice versa under the regular loss). We develop parallel results for the two types of loss functions.

Next we discuss related works for implementing Robbins' formula under the empirical Bayes (EB) estimation framework. Inspecting (4) and (5), we can view a_{y_i-k}/a_{y_i+1-k} as a naive and known estimator of $\check{\gamma}_i$. The ratio functional $w_p^{(k)}(y_i)$, which is unknown in practice, represents the optimal shrinkage factor that depends on $p(\cdot)$. Hence under the f-modeling strategy a simple EB approach, as done in the classical Robbins' formula, is to estimate $w_p^{(k)}(y)$ by plugging-in empirical frequencies: $\hat{w}_n^{(0)}(y) = \hat{p}_n(y)/\hat{p}_n(y+1)$, where $\hat{p}_n(y) = n^{-1} \sum_{i=1}^n I(y_i = y)$. It is noted by Brown et al. (2013) that this plug-in estimator can be highly inefficient especially when $\check{\gamma}_i$ are small. Moreover, the numerator and denominator in $w_p^{(0)}(y)$ are estimated separately, which may lead to unstable ratios. Brown et al. (2013) showed that Robbins' formula can be dramatically improved by imposing additional smoothness and monotonicity adjustments. An alternative approach is to estimate G using NPMLE under appropriate shape constraints. However, efficient estimation of G may not directly translate into an efficient estimation of the ratio functional $w_p^{(k)}(y)$. We recast the compound estimation problem as a convex program, which directly produces consistent estimates of the ratio functionals

$$w_p^{(k)} = (w_p^{(k)}(y_1), \dots, w_p^{(k)}(y_n))^O$$

from data. The estimators are shown to enjoy superior numerical and theoretical properties. Unlike existing f-modeling works that are limited to squared loss and specific members in the DLE family, our method can handle a wide range of discrete distributions and various types of loss functions in a unified framework.

2.2 Shrinkage estimation by convex optimization

This section focuses on the scaled squared error loss ($k = 1$). Methodologies and theories for the case with the squared error loss ($k = 0$) can be derived similarly; details are provided

in Appendix A.1. We first introduce some notations and then present the NEB estimator in Definition 1.

Suppose Y is a non-negative integer-valued random variable with pmf $p(\cdot)$. Define

$$h_0^{(1)}(y) = \begin{cases} 1, & \text{if } y = 0 \\ 1 - w_p^{(1)}(y), & \text{if } y \in \{1, 2, \dots\}. \end{cases} \quad (7)$$

Let $K(y, y^0) = \exp\{-\frac{1}{2}(y - y^0)^2\}$ be the positive definite Gaussian kernel with bandwidth parameter 2κ where κ is a compact subset of $\mathbb{R}^+$ bounded away from 0. Given observations $y = (y_1, \dots, y_n)$ from model (2), let $h_0^{(1)} = (h_0^{(1)}(y_1), \dots, h_0^{(1)}(y_n))$. Define operators ${}_y K(y, y^0) = K(y + 1, y^0) - K(y, y^0)$ and

$${}_{y, y^0} K(y, y^0) = {}_{y^0} {}_y K(y, y^0) = {}_y {}_{y^0} K(y, y^0).$$

Consider the following $n \times n$ matrices, which are needed in the definition of the NEB estimator:

$$K = n^{-2} [K(y_i, y_j)]_{ij}, \quad K = n^{-2} [{}_y K(y_i, y_j)]_{ij}, \quad {}_2 K = n^{-2} [{}_y {}_{y_j} K(y_i, y_j)]_{ij}.$$

Definition 1 (NEB estimator). Consider the DLE model (2) with loss $\ell^{(1)}(\sqrt{\cdot}, \cdot)$. For any fixed 2κ , let $\hat{h}_n^{(1)}(\cdot) = (\hat{h}_1^{(1)}(\cdot), \dots, \hat{h}_n^{(1)}(\cdot))$ be the solution to the following quadratic optimization problem:

$$\min_{h \in H_n} h^T K h + 2h^T K \mathbf{1} + \mathbf{1}^T {}_2 K \mathbf{1}, \quad (8)$$

where $H_n = \{h = (h_1, \dots, h_n) : Ah \leq b, Ch = d\}$ is a convex set and A, C, b and d are known real matrices and vectors that enforce linear constraints on the components of h . Define $\hat{w}_i^{(1)}(\cdot) = 1 - \hat{h}_i^{(1)}(\cdot)$. Then the NEB estimator is given by $\hat{h}_{(1)}^{neb}(\cdot) = (\hat{h}_{(1),i}^{neb}(\cdot) : 1 \leq i \leq n)$, where

$$\hat{h}_{(1),i}^{neb}(\cdot) = \frac{a_{y_i-1}/a_{y_i}}{\hat{w}_i^{(1)}(\cdot)}, \quad \text{if } y_i \in \{1, 2, \dots\},$$

and $\hat{h}_{(1),i}^{neb}(\cdot) = 0$ if $y_i = 0$.

Remark 2. In problem (8) the linear inequality constraints $Ah \leq b$ can be used to impose structural constraints on the NEB decision rule $\hat{h}_{(1)}^{neb}(\cdot)$. These structural constraints, which may take the form of monotonicity constraints (Brown et al., 2013; Koenker and Mizera, 2014), have been shown to be effective for stabilizing the estimator and hence improving the accuracy. For instance, a monotonicity constraint on $\hat{h}_{(1),i}^{neb}(\cdot)$ will imply $\hat{h}_{(1),(1)}^{neb}(\cdot) \leq \dots \leq \hat{h}_{(1),(n)}^{neb}(\cdot)$ for $y_{(1)} \leq y_{(2)} \leq \dots \leq y_{(n)}$. In particular, when $Y_i | \sqrt{\cdot} \leftarrow \text{Poi}(\sqrt{\cdot})$ then $\hat{h}_{(1),i}^{neb}(\cdot) = y_i / \{1 - \hat{h}_i^{(1)}(\cdot)\}$ and the monotonicity constraints in this setting will imply

$$h_{(i)}^{(1)}(\cdot) + \frac{y^{(i)}}{y^{(i+1)}} h_{(i+1)}^{(1)}(\cdot) \leq \frac{y^{(i)}}{y^{(i+1)}} \leq 1, \text{ for } 1 \leq i \leq (n-1)$$

These $n-1$ linear inequality constraints may be imposed with an $(n-1) \times n$ matrix A and an $n-1$ column vector b such that for $1 \leq i \leq (n-1)$ and $1 \leq r \leq n$,

$$A(i, r) = \begin{cases} 1, & \text{when } y_r = y_{(i)} \\ y_{(i)}/y_{(i+1)}, & \text{when } y_r = y_{(i+1)} \\ 0, & \text{otherwise} \end{cases}$$

and $b_i = y_{(i)}/y_{(i+1)} - 1$.

Moreover, when $y_i = 0$ we set $\tilde{h}_{(1),i}^{\text{neb}}(\cdot) = 0$ by convention (see lemma 1). The equality constraints $Ch = d$ accommodate such boundary conditions along with instances of ties for which we require $\tilde{h}_i^{(1)}(\cdot) = \tilde{h}_j^{(1)}(\cdot)$ whenever $y_i = y_j$.

Next we provide some insights on why the optimization criterion (8) works; theories are developed in Section 3 to establish the properties of the NEB estimator rigorously. Denote $h_0^{(1)}$ and $\tilde{h}^{(1)}$ as the ratio functionals corresponding to pmfs p and $\tilde{p}$, respectively and let $M_{\hat{n}}(h) = h^T K h + 2h^T K \mathbf{1} + \mathbf{1}^T \mathbf{2} K \mathbf{1}$. Suppose Y_i are i.i.d. samples obeying $p(y)$. Theorem 1 shows that

$$\hat{M}_{\hat{n}}(\tilde{h}) = M(\tilde{h}) + O_p\left(\frac{\log^2 n}{n^{1/2}}\right),$$

where $\hat{M}_{\hat{n}}(\tilde{h})$ is the objective function in (8) and $M(\tilde{h})$, also denoted $S[\tilde{p}](p)$, is the kernelized Stein's discrepancy (KSD). Roughly speaking, the KSD measures how different one distribution p is from another distribution $\tilde{p}$, with $S[\tilde{p}](p) = 0$ if and only if $\tilde{p} = p$. A key feature of the KSD is that $S[\tilde{p}](p)$ can be equivalently represented by the discrepancy between the corresponding ratio functionals $h_0^{(1)}$ and $\tilde{h}^{(1)}$. Hence, optimizing (8) is asymptotically equivalent to finding $\tilde{h}^{(1)}$ that is as close as possible to the true underlying $h_0^{(1)}$, which corresponds to the optimal shrinkage factor in the compound estimation problem. Theorems 2 and 3 demonstrate that (8) is an effective convex program in the sense that the minimizer $h_{\hat{n}}$ is P -consistent with respect to $h_0^{(1)}$, and the resultant NEB estimator converges to the Bayes estimator.

2.3 Bandwidth selection

The implementation of the quadratic program in (8) requires the choice of a tuning parameter γ in the Gaussian kernel. For practical applications, γ must be determined in a data-driven fashion. For infinitely divisible random variables (Klenke, 2014) such as Poisson variables, Brown et al. (2013) proposed a modified cross validation (MCV) method for choosing the tuning parameter. However, the MCV method cannot be applied to distributions with bounded support as they are not infinitely divisible (Sato and Ken-Iti, 1999) such as the Binomial distribution. To provide a unified estimation framework for the DLE family, we develop an alternative method for choosing γ . The key idea is to derive an asymptotic risk estimate $\text{ARE}^{(1)}(\gamma)$ that serves as an approximation to the true loss $L^{(1)}(\gamma, \tilde{h}_n^{(1)}(\cdot))$. Then the tuning parameter is chosen to minimize $\text{ARE}^{(1)}(\gamma)$.

The methodology based on ARE is illustrated below under the scaled loss (see definition 2) and in Appendix A.2 we provide relevant details for choosing γ under the regular squared loss $L_n^{(0)}$.

Definition 2 (ARE of $\eta_{(1)}^{\text{neb}}(\cdot)$ in the DLE model). Suppose $Y_i \mid \sqrt{y_i} \stackrel{\text{i.i.d.}}{\sim} \text{DLE}(\sqrt{\cdot})$. Under the loss $\ell^{(1)}(\sqrt{\cdot}, \cdot)$, an asymptotic risk estimate of the true loss of $\eta_{(1)}^{\text{neb}}(\cdot)$ is

$$\text{ARE}_n^{(1)}(\cdot, y) = \frac{1}{n} \sum_{i=1}^n \ell^{(1)}(\sqrt{y_i}, \eta_{(1),i}^{\text{neb}}(\cdot))^2, \text{ where}$$

$$\eta_{(1),i}^{\text{neb}}(\cdot) = \{ \eta_{(1),j_i}^{\text{neb}}(\cdot) \}^2 (a_{y_i+1}/a_{y_i}), \quad y_i = 0, 1, \dots$$

with $j_i \in \{1, \dots, n\}$ such that $y_{j_i} = y_i + 1$.

We propose the following estimate of the tuning parameter κ based on the $\text{ARE}_n^{(1)}(\cdot, y)$:

$$\hat{\kappa} = \arg \min_{\kappa \in \mathcal{K}} \text{ARE}_n^{(1)}(\cdot, y) \quad (9)$$

In practice we recommend using $\mathcal{K} = [10, 10^2]$, which worked well in all our simulations and real data analyses. In Section 3, we present Lemma 2 which provides asymptotic justifications for selecting $\hat{\kappa}$ using equation (9).

3. Theory

This section studies the theoretical properties for the NEB estimator under the Poisson and Binomial models. We first investigate the large-sample behavior of the KSD measure (Section 3.1), then turn to the performance of the estimated risk ratios $\mathbf{w}_n$ (Section 3.2), and finally establish the consistency and risk properties of the proposed estimator $\eta_{(1)}^{\text{neb}}(\cdot)$ (Section 3.3). The accuracy of the ARE criteria, which are used in choosing tuning parameter κ , will also be investigated.

3.1 Theoretical properties of the KSD measure

To provide motivation and theoretical support for Definition 1, we introduce the Kernelized Stein's Discrepancy (KSD) (Liu et al., 2016; Chwiałkowski et al., 2016) and discuss its connection to the quadratic program (8). While the KSD has been used in various contexts including goodness of fit tests (Liu et al., 2016; Yang et al., 2018), variational inference (Liu and Wang, 2016) and Monte Carlo integration (Oates et al., 2017), our theory on its connection to the compound estimation problem and empirical Bayes methodology is novel.

Assume that (Y, Y^0) are i.i.d. copies from the marginal pmf p . Consider h_0 defined in Equation (7)¹ and let $\tilde{p}$ denote a pmf on the support of Y , for which we similarly define $\tilde{h}$. The KSD, which is formally defined as

$$S[\tilde{p}](p) = E_p \left[\tilde{h}(Y) - h_0(Y) \right]^2 - K(Y, Y^0) \left[\tilde{h}(Y^0) - h_0(Y^0) \right]^2, \quad (10)$$

provides a discrepancy measure between p and $\tilde{p}$ in the sense that (a)

$$S[\tilde{p}](p) \geq 0 \text{ and } S[\tilde{p}](p) = 0 \text{ if and only if } p = \tilde{p},$$

1. In Section 3.1 we shall drop the superscript from h_0 , which is used to indicate whether the loss is scaled or regular. The simplification has no impact since the general idea holds for both types of losses and the discussion in this section focuses on the scaled loss.

and (b) informally, $S[\tilde{p}](p)$ tends to increase when there is a bigger disparity between h_0 and $\tilde{h}$ (or equivalently, between p and $\tilde{p}$).

The direct evaluation of $S[\tilde{p}](p)$ via Equation (10) is difficult because h_0 is unknown. Note that while the pmf p can be learned well from a random sample $\{Y_1, \dots, Y_n\} \leftarrow p$, we introduce an alternative representation of KSD, developed by Liu et al. (2016), in a reproducing kernel Hilbert space (RKHS) that does not directly involve unknown h_0 . Concretely, consider a positive definite kernel function $\mathbb{E}[\tilde{h}(u), \tilde{h}(v)]$ where

$$\mathbb{E}[\tilde{h}(u), \tilde{h}(v)](u, v) = \tilde{h}(u)\tilde{h}(v)K(u, v) + \tilde{h}(u) \int v K(u, v) + \tilde{h}(v) \int u K(u, v) + \int u, v K(u, v). \quad (11)$$

For i.i.d. copies (Y, Y^0) from distribution p , it can be shown that

$$\begin{aligned} S[\tilde{p}](p) &= \mathbb{E}_{(Y, Y^0) \stackrel{\text{i.i.d.}}{\sim} p} \mathbb{E}[\tilde{h}(Y), \tilde{h}(Y^0)](Y, Y^0) \\ &= \frac{1}{n(n-1)} \mathbb{E}_p \sum_{1 \leq i < j \leq n} \mathbb{E}[\tilde{h}(Y_i), \tilde{h}(Y_j)](Y_i, Y_j) \\ &:= M(\tilde{h}), \end{aligned} \quad (12)$$

where $\{Y_1, \dots, Y_n\}$ is a random sample from p . It can be similarly shown that $M(\tilde{h}) = 0$ if and only if $\tilde{h} = h_0$. Substituting the empirical distribution $\hat{p}_n$ in place of the pmf p in (12), we obtain the following empirical evaluation scheme for $S[\tilde{p}](p)$

$$S[\tilde{p}](\hat{p}_n) = \frac{1}{n^2} \sum_{i=1}^n \sum_{j=1}^n \mathbb{E}[\tilde{h}(y_i), \tilde{h}(y_j)](y_i, y_j). \quad (13)$$

Note that (13) is exactly the objective function $\hat{M}_{\tilde{p}, n}(\tilde{h})$ of the quadratic program (8).

The empirical representation of KSD (13) provides an extremely useful tool for solving the discrete compound decision problem under the EB estimation framework. A key observation is that the kernel function $\mathbb{E}[\tilde{h}(u), \tilde{h}(v)](u, v)$ depends on $\tilde{p}$ only through $\tilde{h}$. Meanwhile, the EB implementation of the generalized Robbins' formula [cf. Equations (4) and (7)] essentially boils down to the estimation of h_0 . Hence, if $S[\tilde{p}](\hat{p}_n)$ is asymptotically equal to $S[\tilde{p}](p)$, then minimizing $S[\tilde{p}](\hat{p}_n)$ with respect to the unknowns $\tilde{h} = (\tilde{h}(y_1), \dots, \tilde{h}(y_n))$ is effectively the process of finding an $\tilde{h}$ that is as close as possible to h_0 , which yields an asymptotically optimal solution to the EB estimation problem. Therefore our formulation of the NEB estimator (1) would be justified as long as we can establish the asymptotic consistency of the sample criterion $S[\tilde{p}](\hat{p}_n)$ around the population criterion $S[\tilde{p}](p)$ uniformly over $\mathcal{H}$ (Theorem 1).

Our analysis in this and the following sections will be based on the hierarchical model of equation (2): $Y_i | \sqrt{i} \stackrel{\text{i.i.d.}}{\sim} \text{DLE}(\sqrt{i})$, $\sqrt{i} \stackrel{\text{i.i.d.}}{\sim} G(\cdot)$ where $G(\cdot)$ is an unspecified prior distribution on $\sqrt{i}$. In this setup the marginal pmf of Y is $p(y) := P(Y = y) = \int p(y|\sqrt{i})dG(\sqrt{i})$. We impose the following regularity conditions that are needed in our technical analysis.

(A1) $\mathbb{E}_p |\mathbb{E}[\tilde{h}(U), \tilde{h}(V)](U, V)|^2 < 1$ for all $U \in \mathcal{K}$ where $\mathcal{K}$ is a compact subset of $\mathbb{R}^+$ bounded away from 0.

(A2) For some $\alpha \in (0, 1)$, $E_G\{\exp(\alpha\sqrt{\cdot})\} < 1$ where the expectation is taken with respect to the prior distribution G of $\sqrt{\cdot}$.

(A3) For any function g that satisfies $0 < kgk_2^2 < 1$, there exists a constant $c > 0$ such that $\lim_{n \rightarrow \infty} \frac{1}{n} \sum_{y, y^0=0}^n g(y)K(y, y^0)g(y^0) > ckgk_2^2$ for every $\alpha \in (0, 1)$.

Remark 3. Assumption (A1) is a moment condition on the kernel function related to V-statistics; see, for example, Serfling (2009). Assumption (A2) is a moment condition on the prior distribution G . In particular, it ensures that with high probability $\max(Y_1, \dots, Y_n) \leq \log n$ as $n \rightarrow \infty$. This idea is formalized in Lemma 4 in Appendix B. It is likely that assumption (A2) can be further relaxed but we do not seek the full generality here. Assumption (A3) is a standard condition which ensures that the KSD $S[\tilde{p}](p)$ is a valid discrepancy measure (Liu et al., 2016; Chwialkowski et al., 2016).

Theorem 1. If $\tilde{p}$ is a probability mass function on the support of Y then, under Assumptions (A1) and (A2), we have

$$\sup_{\alpha \in (0, 1)} \|\hat{M}_{\alpha, n}(\tilde{h}) - M(\tilde{h})\| = O_p\left(\frac{\log^2 n}{n}\right).$$

In the context of our compound estimation framework, Theorem 1 is significant because it guarantees that the empirical version of the KSD measure given by $\hat{M}_{\alpha, n}(\tilde{h})$ is asymptotically close to its population counterpart $M(\tilde{h})$ uniformly in $\alpha \in (0, 1)$. Moreover, along with the fact that $M(h_0) = 0$, Theorem 1 establishes that $\hat{M}_{\alpha, n}(h)$ is the appropriate criteria to minimize with respect to $h \in \mathcal{H}$. In Theorem 2, we further show that the resulting estimator of the ratio functionals $w_p^{(1)}$ from equation (8) are consistent.

3.2 Theoretical properties of $\hat{w}_n$

The optimization problem in (8) is defined over a convex set $H_n = \{h = (h_1, \dots, h_n) : Ah \leq b, Ch = d\}$ which is a subset of $\mathbb{R}^n$. However, the dimension of H_n , denoted by $\dim(H_n)$, is usually much smaller than n . Consider the Binomial case where $Y_i | q_i \sim \text{Bin}(m_i, q_i)$ with $q_i \in (0, 1)$, $m_i \leq m < \infty$ and $\sqrt{\cdot} = q_i/(1 - q_i)$. Here $\dim(H_n)$ is at most m since $\max(Y_1, \dots, Y_n) \leq m$. While the boundedness of the support is not always available outside the Binomial case, in most practical applications it is reasonable to assume that the distribution of $\sqrt{\cdot}$ has some finite moments, which ensures that $\dim(H_n)$ grows slower than $\log n$; see Assumption (A2). In Lemma 4 we make this precise. Moreover, as discussed in remark 2, the linear inequality constraints $Ah \leq b$ impose structural constraints on $w_p^{(1)}$. For the ensuing discussion and following Brown et al. (2013), we let these structural constraints to take the form of monotonicity constraints on the NEB decision rule. Since the Binomial and the Poisson models have the monotone likelihood ratio property, $w_p^{(1)}$ is monotone and so $h_0 \in H_n$. The next theorem establishes the asymptotic consistency of $w_n^{(1)}$.

Theorem 2. Let $K(\cdot, \cdot)$ be the positive definite Gaussian kernel with bandwidth parameter $\alpha \in (0, 1)$. If $\lim_{n \rightarrow \infty} c_n n^{-1/2} \log^2 n = 0$ then, under Assumptions (A1) - (A3), we have for any $\alpha \in (0, 1)$,

$$\lim_{n \rightarrow \infty} P\left(\frac{1}{n} \sum_{i=1}^n \hat{w}_n^{(1)}(Y_i) - w_p^{(1)}\right)^2 \leq c_n \alpha = 0, \text{ for any } \alpha > 0,$$

where $\hat{w}_n^{(1)}(\cdot) = 1 - \hat{h}_n^{(1)}(\cdot)$.

Theorem 2 shows that under the scaled squared error loss, $\hat{w}_n^{(1)}(\cdot)$, the optimizer of quadratic problem in equation (8), is a consistent estimator of $w_p^{(1)}$, the optimal shrinkage factor in the Bayes rule (Lemma 1). In particular, the aforementioned consistency result is related to the theoretical analysis of minimum KSD estimators in Barp et al. (2019). While Barp et al. (2019) establish almost sure convergence and asymptotic normality of such minimum KSD estimators, the analysis in this section is geared towards studying the asymptotic optimality of the proposed NEB estimator in the sense of Theorem 3 below. The proof of Theorem 2 is available in Appendix B.3 which also includes relevant details for proving a companion result under the regular squared error loss.

Remark 4. The estimation framework in Definition 1 may be used for producing consistent estimators for any member in the DLE family. This allows the corresponding NEB estimator to cover a much wider class of discrete distributions than previously proposed. Compared to the existing methods of Efron (2011) and Brown et al. (2013), our proposed NEB estimation framework is robust against departures from the true data generating process. This is due to the fact that the quadratic optimization problem in (8) does not rely on the specific form of the distribution of $Y|\sqrt{\cdot}$, and that the shrinkage factors are estimated in a non-parametric fashion. The robustness of the estimator is corroborated by our numerical results in Section 4.

3.3 Properties of the NEB estimator

In this section we discuss the risk properties of the NEB estimator. Let

$$\bar{\mathcal{L}}_n^{(1)}(\sqrt{\cdot}, \mathbf{w}_{(1)}^{\text{neb}}) = \frac{1}{n} \sum_{i=1}^n L_n^{(1)}(\sqrt{\cdot}, \mathbf{w}_{(1)}^{\text{neb}})$$

We begin with Lemma 2 which shows that uniformly in $2 \leftarrow$, the gap between $\text{ARE}_n^{(1)}(\cdot)$ and $E\{\bar{\mathcal{L}}_n^{(1)}(\sqrt{\cdot}, \mathbf{w}_{(1)}^{\text{neb}})\}$ is asymptotically negligible. This justifies our proposed methodology for choosing the tuning parameter in Section 2.3. In the following lemma, we let c be a sequence satisfying $\lim_{n \rightarrow \infty} c_n n^{-1/4} \log^3 n = 0$.

Lemma 2. Under Assumptions (A1) - (A3), we have

- (1). $c_n \sup_{2 \leftarrow} \text{ARE}_n^{(1)}(\cdot, Y) - \bar{\mathcal{L}}_n^{(1)}(\sqrt{\cdot}, \mathbf{w}_{(1)}^{\text{neb}}) = o_p(1)$.
- (2). $c_n \sup_{2 \leftarrow} \text{ARE}_n^{(1)}(\cdot, Y) - E\{\bar{\mathcal{L}}_n^{(1)}(\sqrt{\cdot}, \mathbf{w}_{(1)}^{\text{neb}})\} = o_p(1)$;

In Appendices B.5 and B.6 we prove Lemma 2 for the Binomial and Poisson models under both scaled squared error and squared error losses.

To analyze the quality of the data-driven bandwidth $\hat{n}$ [cf. Equation (9)], we consider an oracle loss estimator $\mathcal{L}_{(1)}^{\text{orc}} := \mathcal{L}_{(1)}^{\text{neb}}(\mathbf{w}_{(1)}^{\text{orc}})$, where

$$\mathbf{w}_{(1)}^{\text{orc}} := \arg \min_{2 \leftarrow} L_n^{(1)}(\sqrt{\cdot}, \mathbf{w}_{(1)}^{\text{neb}}).$$

The oracle bandwidth $\hat{h}_1^{orc}$ is not available in practice since it requires the knowledge of unknown $\check{\gamma}$. However, it provides a benchmark for assessing the effectiveness of the data-driven bandwidth selection procedure in Section 2.3. The following lemma shows that the loss of $\hat{\gamma}_{(1)}^{neb}$ converges in probability to the loss of $\hat{\gamma}_{(1)}^{orc}$.

Lemma 3. Under Assumptions (A1) - (A3), if $\lim_{n \rightarrow \infty} c_n n^{-1/4} \log^2 n = 0$, then for both the Poisson and Binomial models, we have

$$\lim_{n \rightarrow \infty} P \left(L_n^{(1)}(\check{\gamma}, \hat{\gamma}_{(1)}^{neb}) - L_n^{(1)}(\check{\gamma}, \hat{\gamma}_{(1)}^{orc}) + \frac{1}{n} = o_p(1) \right) = 0 \text{ for any } \epsilon > 0.$$

Obviously, the estimator $\hat{\gamma}_{(1)}^{neb}$ is lower bounded by the risk of the optimal solution $\hat{\gamma}_{(1)}^*$ (Lemma 1). Next we study the asymptotic optimality of $\hat{\gamma}_{(1)}^{neb}$, which aims to provide decision theoretic guarantees on $\hat{\gamma}_{(1)}^{neb}$ in relation to $\hat{\gamma}_{(1)}^*$. Theorem 3 establishes the optimality theory by showing that (a) the average squared error between $\hat{\gamma}_{(1)}^{neb}$ and $\hat{\gamma}_{(1)}^*$ is asymptotically small, and (b) the NEB estimator is asymptotically as good as the corresponding Bayes estimator in terms of expected loss.

Theorem 3. Under the conditions of Theorem 2, if $\lim_{n \rightarrow \infty} c_n n^{-1/2} \log^4 n = 0$, then for both the Poisson and Binomial models, we have

$$\frac{c_n}{n} \left(\hat{\gamma}_{(1)}^{neb} - \hat{\gamma}_{(1)}^* \right)^2 = o_p(1).$$

Furthermore, under the same conditions, we have,

$$\lim_{n \rightarrow \infty} E \left(L_n^{(1)}(\check{\gamma}, \hat{\gamma}_{(1)}^{neb}) - L_n^{(1)}(\check{\gamma}, \hat{\gamma}_{(1)}^*) \right) = 0.$$

In Appendix A.2, we discuss the counterpart to Theorem 3 under the squared error loss $L_n^{(0)}$.

4. Numerical Results

In this section we first discuss, in Section 4.1, the implementation details of the convex program (8) and bandwidth selection process (9) (see also (19) in Appendix A.2). Then we investigate the numerical performance of the NEB estimator for Poisson, Binomial and Negative Binomial compound decision problems, respectively in Sections 4.2, 4.3 and 4.4. In each case, we consider both regular and scaled squared losses. Our numerical results demonstrate that the efficiency gain of the NEB estimator over competitive methods is substantial in many settings.

We have developed an R package, `npeb`, to implement the NEB estimator in definition 1 (and definition 3 in Appendix A.1). Moreover, the R code that reproduces the numerical results in this section can be downloaded from the following link: https://github.com/trambakbanerjee/DLE_paper.

4.1 Implementation Details

For a fixed η we use the R-package CVXR (Fu et al., 2017) to solve the optimization problem in Equations (8) (and (15) in Appendix A.1). As discussed in remark 2 of section 2.2, under the scaled squared error loss ($k = 1$) the linear inequality constraints, given by $Ah \leq b$, ensure that the resulting decision rule $\hat{h}^{eb}(\cdot)$ is monotonic, while the equality constraints

$Ch = d$ handle boundary cases that involve $y_i = 0$ and ties. Moreover, since $w^{(1)}(y) > 0$, the inequality constraints also ensure that $h_i < 1$ whenever $y_i > 0$. Implementation under the squared error loss ($k = 0$) follows along similar lines and the inequality constraints in this case ensure that $h_i + y_i > 0$ whenever $y_i = 0$.

A data-driven choice of the tuning parameter η is obtained by first solving problems (8) and (15) over a grid of η values, i.e. $\{\eta_1, \dots, \eta_s\}$, and then computing the corresponding asymptotic risk estimate $ARE_n^{(k)}(\eta_j)$ for $j = 1, \dots, s$. Then $\hat{\eta}_k$ is chosen according to

$$\hat{\eta}_k := \arg \min_{\eta \in \{\eta_1, \dots, \eta_s\}} ARE_n^{(k)}(\eta),$$

where $k \in \{0, 1\}$. For all simulations and real data analyses considered in this paper, we have fixed $s = 10$ and employed an equi-spaced grid over $[10, 10^2]$.

4.2 Simulations: Poisson Distribution

In this section we consider the Poisson compound decision problem and generate $Y_i | \sqrt{i} \stackrel{\text{i.i.d.}}{\sim} \text{Poi}(\sqrt{i})$ for $i = 1, \dots, n$. We vary n from 500 to 5000 in increments of 500 and simulate $\sqrt{i}$ from the following four different scenarios:

Scenario 1: $\sqrt{i} \stackrel{\text{i.i.d.}}{\sim} \text{Unif}(0.5, 15)$.

Scenario 2: $\sqrt{i} \stackrel{\text{i.i.d.}}{\sim} \text{Gamma}(10, 2)$.

In the next two scenarios we consider departures from the usual Poisson model and simulate our data from the Conway-Maxwell-Poisson distribution (Shmueli et al., 2005) $\text{CMP}(\sqrt{i}, \eta)$. The CMP distribution is a generalization of some well-known discrete distributions. With $\eta < 1$, CMP represents a discrete distribution that has longer tails than the Poisson distribution with parameter $\sqrt{i}$.

Scenario 3: We simulate $\sqrt{i} \stackrel{\text{i.i.d.}}{\sim} 0.5 \cdot \{10\} + 0.5 \text{Gamma}(5, 2)$ for each i and let

$$Y_i | \sqrt{i} \stackrel{\text{i.i.d.}}{\sim} 0.8 \text{Poi}(\sqrt{i}) + 0.2 \text{CMP}(\sqrt{i}, \eta),$$

where we fix $\eta = 0.8$ for the CMP distribution.

Scenario 4.1: In this scenario we conduct estimation under the scaled squared error loss.

We let $\sqrt{i} \stackrel{\text{i.i.d.}}{\sim} 0.5 \cdot \{5\} + 0.5 \cdot \{15\}$, $\eta_i | \sqrt{i} = 0.8 \mathbb{I}(\sqrt{i} = 5) + 1 \mathbb{I}(\sqrt{i} = 15)$ and simulate Y_i from the CMP distribution with parameters $\sqrt{i}$ and η_i . Thus, about half of the samples arise from a Poisson distribution with mean 15 while remaining are realizations from a $\text{CMP}(5, 0.8)$.

Scenario 4.2: We consider estimation under the squared error loss and let $\check{\cdot}$ to be an equi-spaced vector of length n in $[1, 5]$. We simulate Y_i from the CMP distribution with parameters $\check{\cdot}_i$ and η fixed at 0.8.

For each scenario, the following competing estimators of $\check{\cdot}_i$ are considered:

1. the proposed estimator, denoted NEB and the oracle NEB estimator $\check{\cdot}_{(k)}^{or} := \check{\cdot}_{(k)}^{neb}(\check{\cdot}^{orc})$, denoted NEB OR;
2. the estimator of Poisson means from Brown et al. (2013), denoted BGR;
3. Tweedie's formula for the Poisson model, denoted TF OR;
4. Tweedie's formula for the Normal means problem based on transformed data, denoted TF Gauss. The approach using transformation was suggested by Brown et al. (2013).
5. the estimator of Poisson means from Koenker and Gu (2017), denoted KM;
6. the estimator of Poisson means based on the g-modeling approach of Efron (2016), denoted Deconv.

The risk performance of the TF OR method relies heavily on the choice of a suitable bandwidth parameter $h > 0$. We use the oracle loss estimate h^{orc} , which is obtained by minimizing the true loss $L^{(0)}_{\check{\cdot}, n}$. The TF Gauss methodology is only applicable for the Normal means problem, and uses a variance stabilization transformation on Y_i to get $Z_i = \sqrt{2} \sqrt{Y_i + 0.25}$. The Z_i are then treated as approximate Normal random variables with mean μ_i and variances 1. To estimate the normal means μ we rely on g-modeling and use NPMLE. Finally, $\check{\cdot}_i$ are estimated as $0.25\hat{\mu}^2$. It is important to note that along with the NEB estimator, BGR and TF OR are based on f-modeling while the rest in the preceding list of six competitors are based on g-modeling. Moreover, BGR, TF OR and TF Gauss only focus on the regular squared error loss $L^{(0)}$. Nevertheless, in our simulation we assess the performance of these estimators for estimating $\check{\cdot}$ under both $L^{(0)}$ and $L^{(1)}$.

Table 1: Poisson compound decision problem under scaled squared error loss: Risk ratios $R_n^{(1)}(\check{\cdot}, \cdot)/R_n^{(1)}(\check{\cdot}, \check{\cdot}_{(1)}^{neb})$ at $n = 5000$ for estimating $\check{\cdot}$.

Method	Scenario			
	1	2	3	4.1
KM	0.94	1.00	1.11	1.00
Deconv	1.00	1.06	1.03	1.11
TF Gauss	1.03	1.03	1.23	1.18
TF OR	1.00	1.02	1.28	1.10
BGR	1.22	1.07	1.28	1.25
NEB	1.00	1.00	1.00	1.00
NEB OR	1.00	1.00	0.98	1.00

Table 2: Poisson compound decision problem under squared error loss: Risk ratios $R_n^{(0)}(\check{\cdot}, \cdot)/R_n^{(0)}(\check{\cdot}, \check{\cdot}_{(0)}^{neb})$ at $n = 5000$ for estimating $\check{\cdot}$.

Method	Scenario			
	1	2	3	4.2
KM	1.00	1.01	1.59	1.21
Deconv	1.02	1.08	1.43	1.21
TF Gauss	1.00	1.01	1.51	1.08
TF OR	1.07	1.03	1.66	1.12
BGR	1.01	1.02	1.55	1.15
NEB	1.00	1.00	1.00	1.00
NEB OR	1.00	1.00	0.90	1.00

The performances of these six estimators are presented in figures 1 and 2 wherein the risk $R_n^{(k)}(\check{\cdot}, \cdot)$ is estimated using 50 Monte Carlo repetitions for varying n . Tables 1 and 2

EB Estimation in Discrete Linear Exponential Family

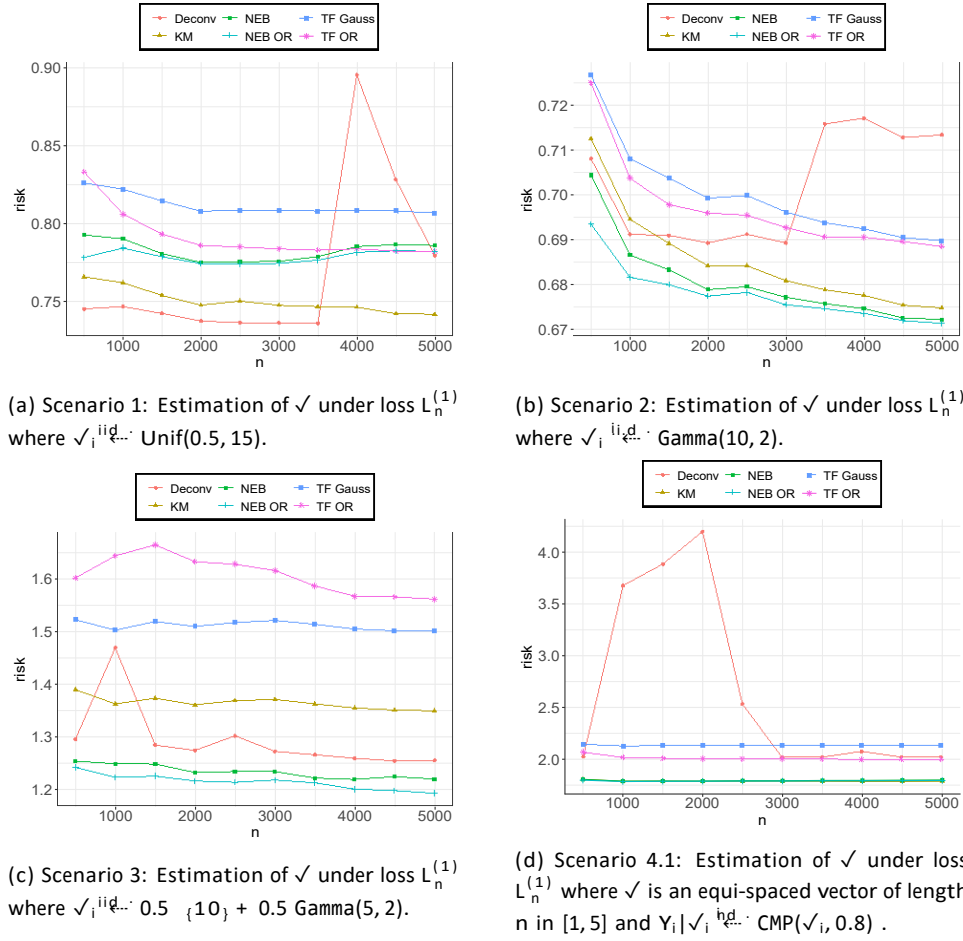


Figure 1: Poisson compound decision problem under scaled squared error loss: Risk estimates of the various estimators for scenarios 1, 2, 3 and 4.1.

report the ratios $R_n^{(k)}(\sqrt{\cdot}, \cdot) / R_n^{(k)}(\sqrt{\cdot}, \text{neb})$ of the average risks at $n = 5000$ and for $k = 1, 0$ respectively, where a risk ratio bigger than 1 indicates a smaller estimation risk for the NEB estimator. For BGR the modified cross validation approach of choosing the bandwidth parameter was extremely slow in our simulations and we therefore report its risk performance only at $n = 5000$.

Figure 1 and table 1 present the risk performances of the competing estimators under the scaled squared error loss. Under scenarios 1 and 2 all estimators, with the exception of BGR in scenario 1 (table 1), exhibit competitive risk performance. For scenarios 3 and 4.1, which represent departures from the Poisson model, the NEB estimator demonstrates a substantially better performance than TF Gauss, TF OR and BGR. We note that KM and Deconv are competitive in scenarios 4.1 and 3, respectively, which indicates that along with the NEB estimator these g-modeling based approaches are potentially robust to misspecifications of the Poisson model considered in scenarios 3 and 4.1. Figure 1 reveals that the risk profile of Deconv is affected by its poorer estimates of $\sqrt{\cdot}$ at various sample sizes

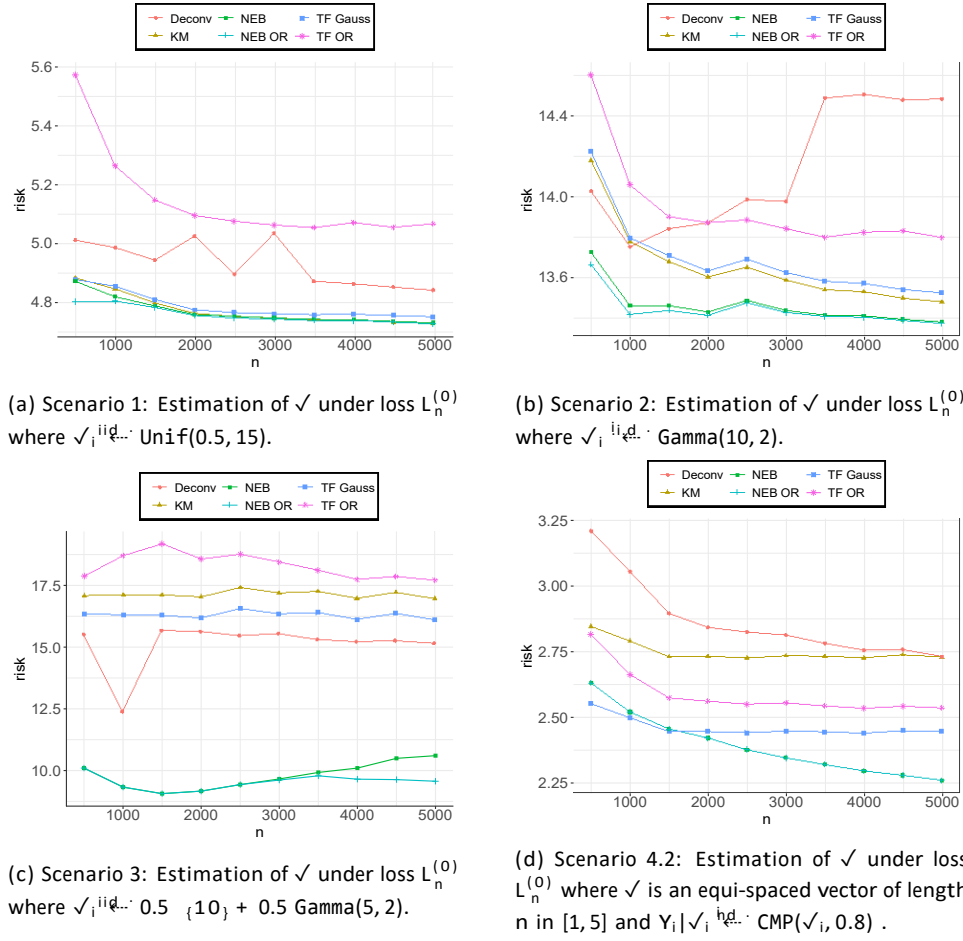


Figure 2: Poisson compound decision problem under squared error loss: Risk estimates of the various estimators for scenarios 1, 2, 3 and 4.2.

and especially at the smaller sample sizes for scenarios 3 and 4.1. This behavior continues to appear even when the number of Monte Carlo repetitions are increased.

The risk performance of the competing estimators under the squared error loss is presented in figure 2 and table 2. Under scenarios 1 and 2 all estimators continue to exhibit a competitive performance. BGR, in particular, demonstrates a substantially improved performance now that estimation is conducted under squared error loss. Scenarios 3 and 4.2 consider departures from the Poisson model and in these settings the NEB estimator has a substantially better risk performance than all other competing methods considered here. We note that in scenarios 3, 4.1 and 4.2 the NEB estimator is robust to departures from the Poisson model. Proposition 7 in Barp et al. (2019) guarantees that, in general, the influence function of minimum KSD estimators, such as the NEB estimator, is bounded under data corruption and the behavior of the proposed NEB estimator in scenarios 3, 4.1 and 4.2 is potentially an example of such robustness property of minimum KSD estimators.

4.3 Simulations: Binomial Distribution

In this section we consider the Binomial compound decision problem and generate $Y_i | q_i \stackrel{i.i.d.}{\sim} \text{Bin}(m_i, q_i)$ for $i = 1, \dots, n$. We vary n from 500 to 5000 in increments of 500 and simulate $\sqrt{q_i} = q_i / (1 - q_i)$ from the following four different scenarios:

Scenario 1: $q_i \stackrel{i.i.d.}{\sim} \text{Unif}(0.1, 0.7)$ and $m_i = 10$.

Scenario 2: $\sqrt{q_i} \stackrel{i.i.d.}{\sim} (1/3) \cdot \{0.5\} + (1/3) \cdot \{1\} + (1/3) \cdot \{2\}$ and $m_i = 10$.

Scenario 3: $q_i \stackrel{i.i.d.}{\sim} \text{Beta}(1, 6)$ and $m_i = 10$.

Scenario 4: $\sqrt{q_i} \stackrel{i.i.d.}{\sim} \text{Exp}(2)$ and $m_i = 5$.

Unlike scenarios 1 and 3, the data generating process in scenarios 2 and 4 directly sample the odds. Moreover the compound estimation problem in scenarios 3 and 4 is challenging because in these settings the distribution of $\sqrt{q_i}$ has a mean that is substantially smaller in magnitude to the mean of $\sqrt{q_i}$ in scenarios 1 and 2. For example in scenario 2 the mean of $\sqrt{q_i}$ is about 1.16 while that in scenario 4 is 0.5. We consider the following five competing estimators of $\sqrt{q_i}$:

1. the proposed estimator NEB and its oracle version NEB OR;
2. Tweedie's formula for Binomial log odds, denoted TF OR;
3. Tweedie's formula for the Normal means problem based on transformed data, denoted TF Gauss.
4. the estimator of Binomial odds from Koenker and Gu (2017), denoted KM;
5. the estimator of Binomial odds from Efron (2016), denoted Deconv.

For TF OR, analogous to the Poisson case, we continue to use the oracle loss estimate h^{orc} as a choice for the bandwidth parameter. Since the TF Gauss methodology is only applicable for the Normal means problem, it uses a variance stabilization transformation on Y_i to get $Z_i = \arcsin((Y_i + 0.25)/(m_i + 0.5))$. The Z_i are then treated as approximate Normal random variables with mean μ_i , variances $(4m_i)^{-1}$, and estimate of the means μ_i 's are obtained using NPMLE. Finally, q_i is estimated as $\{\sin(\hat{\mu}_i)\}^2$. We note that the competitors TF OR and TF Gauss to our NEB estimator do not directly estimate the odds $\sqrt{q_i}$. For instance under the squared error loss, TF Gauss estimates the success probabilities q_i while TF OR estimates $\log \sqrt{q_i}$. Nevertheless, in this simulation experiment we assess the performance of these two estimators for estimating the odds under both squared error loss and its scaled version.

The simulation results are presented in Figures 3 and 4 wherein the risks of various estimators are calculated by averaging over 50 Monte Carlo repetitions for varying n . Tables 3 and 4 report the risk ratios $R^{(k)}(\sqrt{\cdot})/R^{(k)}(\sqrt{\cdot}, \text{neb}_{(k)})$ at $n = 5000$ and for $k = 1, 0$ respectively, where a risk ratio bigger than 1 indicates a smaller estimation risk for the NEB estimator.

Under the scaled squared error loss (figure 3 and table 3) KM and Deconv demonstrate a superior risk performance for scenarios 1 and 2 while the NEB outperforms them for the

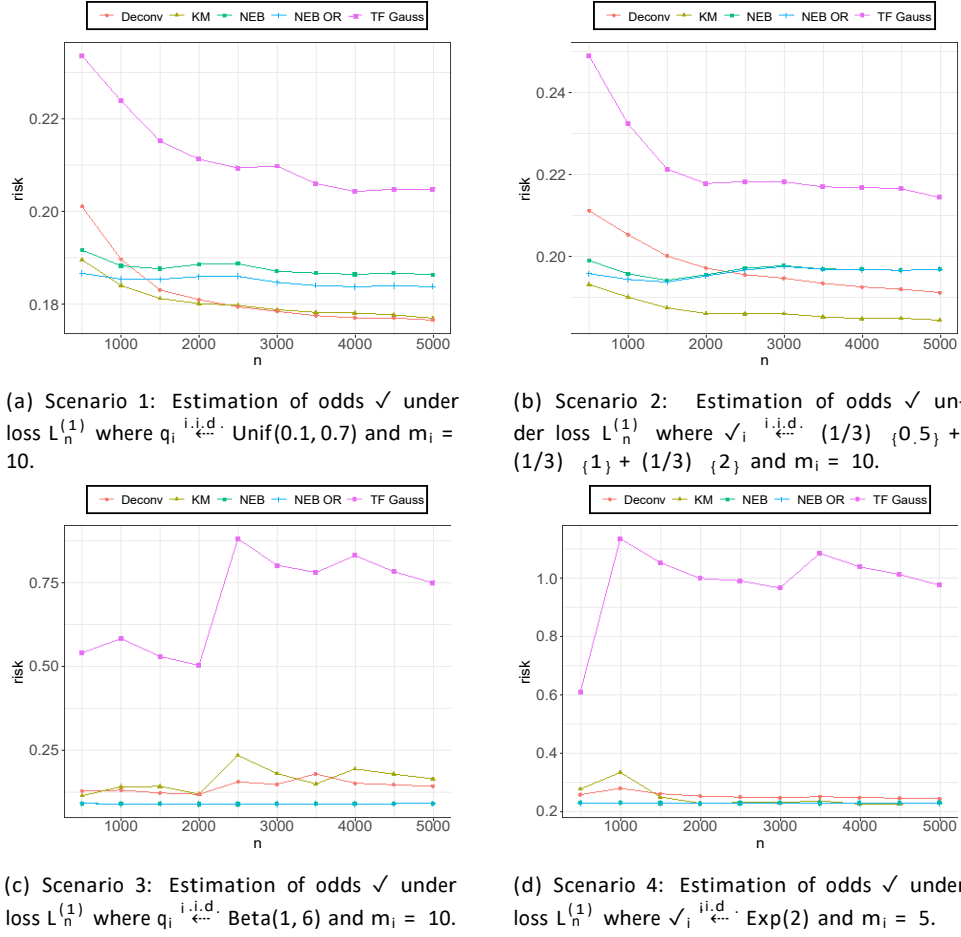
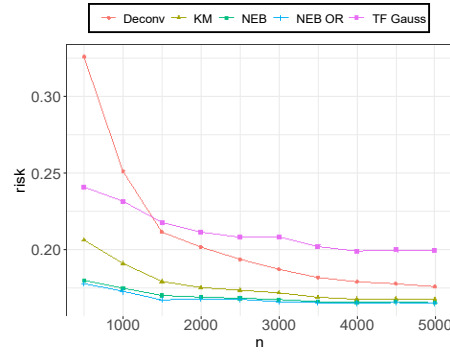


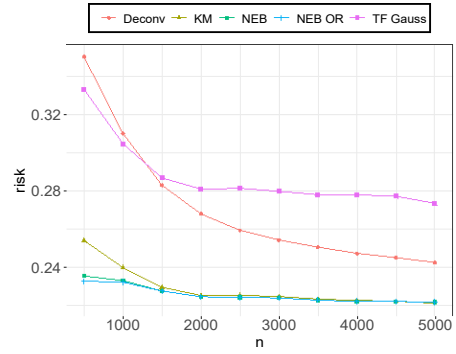
Figure 3: Binomial compound decision problem under scaled squared error loss: Risk estimates of the various estimators for Scenarios 1 to 4.

challenging settings of scenarios 3 and 4. The two Tweedie's formula based estimators, TF Gauss and TF OR, exhibit relatively poorer performance which is not surprising because these two estimators are designed to estimate q_i and $\log \sqrt{\cdot}$ under loss $L_n^{(0)}$. For the squared error loss (figure 4 and table 4) the simulation results reveal that with the exception of scenario 3, the NEB estimator and KM demonstrate competitive risk performance. Scenario 3, along with scenario 4, is a challenging setting wherein the mean of the distribution of $\sqrt{\cdot}$ is substantially smaller in magnitude to the mean of $\sqrt{\cdot}$ in scenarios 1 and 2. Across the four scenarios, TF OR exhibits the poorest performance and appears to suffer from the fragmented approach of estimating the gradient of the log density $\log p(y)$ wherein $p(y)$ and its first derivative with respect to y are estimated separately using a Gaussian kernel with common bandwidth h^{orc} . Between the two g-modeling based approaches considered in this section, Deconv exhibits a relatively poorer risk performance than KM and for scenario 3 in particular the average risk of Deconv is substantially larger.

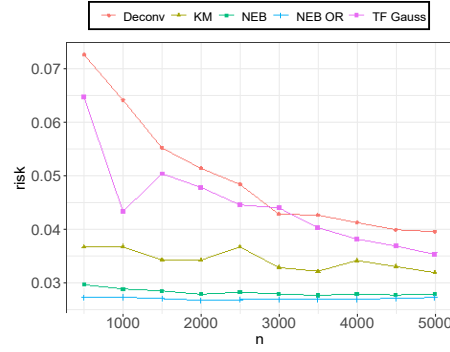
EB Estimation in Discrete Linear Exponential Family



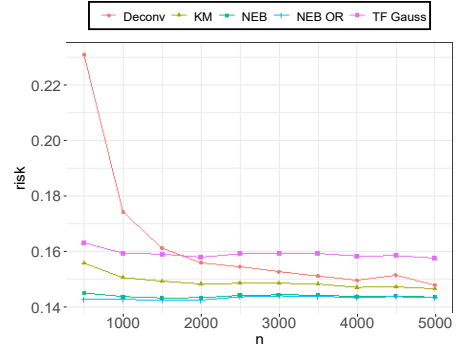
(a) Scenario 1: Estimation of odds $\checkmark$ under loss $L_n^{(0)}$ where $q_i \stackrel{i.i.d.}{\leftarrow} \text{Unif}(0.1, 0.7)$ and $m_i = 10$.



(b) Scenario 2: Estimation of odds $\checkmark$ under loss $L_n^{(0)}$ where $\checkmark_i \stackrel{i.i.d.}{\leftarrow} (1/3) \{0.5\} + (1/3) \{1\} + (1/3) \{2\}$ and $m_i = 10$.



(c) Scenario 3: Estimation of odds $\checkmark$ under loss $L_n^{(0)}$ where $q_i \stackrel{i.i.d.}{\leftarrow} \text{Beta}(1, 6)$ and $m_i = 10$.



(d) Scenario 4: Estimation of odds $\checkmark$ under loss $L_n^{(0)}$ where $\checkmark_i \stackrel{i.i.d.}{\leftarrow} \text{Exp}(2)$ and $m_i = 5$.

Figure 4: Binomial compound decision problem under squared error loss: Risk estimates of the various estimators for Scenarios 1 to 4.

Table 3: The Binomial compound decision problem under scaled squared error loss: Risk ratios $R_n^{(1)}(\checkmark, \cdot)/R_n^{(1)}(\checkmark, \frac{\text{neb}}{1})$ at $n = 5000$ for estimating $\checkmark$.

Method	Scenario			
	1	2	3	4
KM	0.95	0.94	1.85	1.00
Deconv	0.95	0.97	1.59	1.06
TF Gauss	1.01	1.09	8.52	4.30
TF OR	> 10	> 10	> 10	> 10
NEB	1.00	1.00	1.00	1.00
NEB OR	0.99	1.00	1.00	1.00

Table 4: The Binomial compound decision problem under the squared error loss: Risk ratios $R_n^{(0)}(\checkmark, \cdot)/R_n^{(0)}(\checkmark, \frac{\text{neb}}{0})$ at $n = 5000$ for estimating $\checkmark$.

Method	Scenario			
	1	2	3	4
KM	1.01	1.00	1.15	1.02
Deconv	1.06	1.09	1.42	1.03
TF Gauss	1.21	1.23	1.27	1.10
TF OR	> 10	> 10	> 10	> 10
NEB	1.00	1.00	1.00	1.00
NEB OR	1.00	1.00	0.98	1.00

4.4 Simulations: Negative Binomial Distribution

In this section we investigate the numerical performance of the NEB estimator for compound decision problems involving the Negative Binomial (NB) distribution. We generate observations $Y_i | q_i \stackrel{i.i.d.}{\sim} \text{NB}(\text{Binom}(r_i, q_i))$ for $i = 1, \dots, n$ and vary n from 500 to 5000 in increments of 500. Here the goal is to estimate $\sqrt{q_i} = 1 - q_i$ and we consider the following three different scenarios for simulating q_i for $i = 1, \dots, n$:

Scenario 1: $q_i \stackrel{i.i.d.}{\sim} 0.4 \cdot \{0.5\} + 0.6 \cdot \text{Beta}(1, 1)$ and fix $r_i = 3$.

Scenario 2: $q_i \stackrel{i.i.d.}{\sim} (1/3) \cdot \{0.5\} + (1/3) \cdot \{0.7\} + (1/3) \cdot \{0.9\}$ and fix $r_i = 5$.

Scenario 3: $q_i \stackrel{i.i.d.}{\sim} \text{Beta}(5, 2)$ and fix $r_i = 10$.

In scenarios 2 and 3 the median $\sqrt{q_i}$ is substantially smaller than 0.5 which represents a challenging estimation setting for the following competing estimators:

1. the proposed estimator, denoted NEB and the oracle version NEB OR;
2. Tweedie's formula for $\log \sqrt{q_i}$ under the NB model, denoted TF OR;
3. Tweedie's formula for the Normal means problem based on transformed data, denoted TF Gauss;
4. the naive estimator $1 - (r_i - 1)/(r_i + Y_i - 1)$ of $\sqrt{q_i}$ where $(r_i - 1)/(r_i + Y_i - 1)$ is the minimum variance unbiased estimator (MVUE) of q_i .

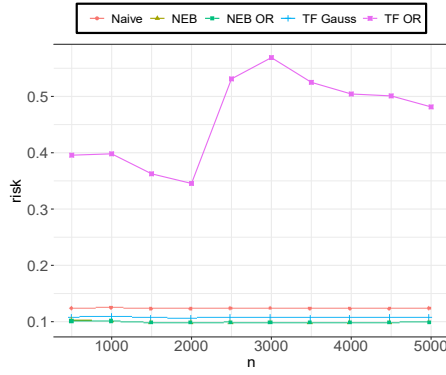
We continue to use the oracle loss estimate h^{orc} as the bandwidth choice for TF OR. For TF Gauss we use a variance stabilization transformation on Y_i to get $Z_i = \sqrt{2} \arcsin \sqrt{Y_i/r_i}$. The Z_i are then treated as approximate Normal random variables with mean μ_i and variances $1/r_i$. To estimate the normal means μ_i we rely on g-modeling and use NPMLE. Finally, $\sqrt{q_i}$ are estimated as $1 - \{1 + [\sinh(0.5\hat{\mu}_i)]^2\}^{-1/2}$. It is important to note that unlike the NEB estimator, the remaining competing estimators only focus on the regular squared error loss $L_n^{(0)}$. Nevertheless, in our simulation we assess the performance of these estimators for estimating $\sqrt{q_i}$ under both $L_n^{(0)}$ and $L_n^{(1)}$.

Figure 5 and tables 5, 6 report the performance of the competing estimators of $\sqrt{q_i}$ for the NB compound estimation problem. Under the squared error loss table 6 and right panel of figure 5 reveal that across all sample sizes performance of the NEB estimator is substantially better than the competing estimators considered in this experiment. In this setting TF OR is the next best while the naive estimator of $\sqrt{q_i}$ is outperformed by the three shrinkage estimators. Under the scaled squared error loss (table 5 and left panel of figure 5), the NEB estimator continues to be better than the competing estimators although the performance of TF Gauss is impressive given that it is based on Normality transformed data under the usual squared error loss.

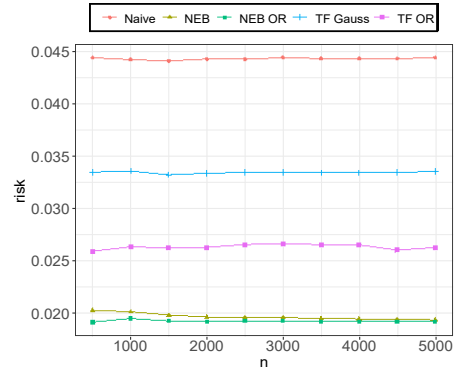
5. Real Data Analyses

This section illustrates two real data applications that use the proposed method for estimating Juvenile Delinquency rates from Poisson models and news popularity from Binomial models.

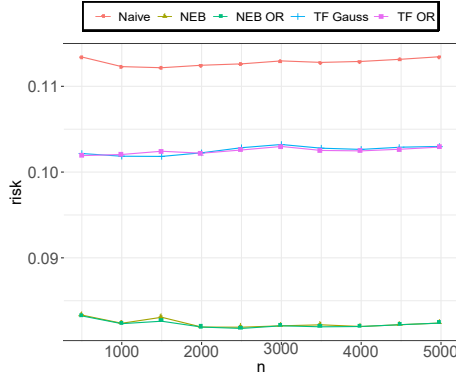
EB Estimation in Discrete Linear Exponential Family



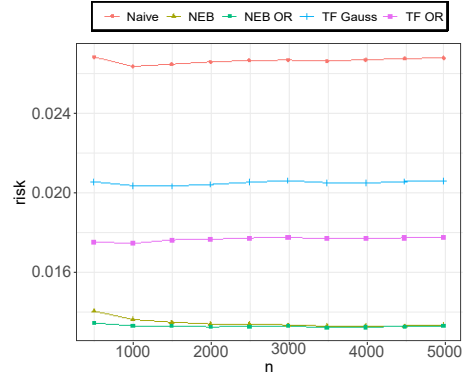
(a) Scenario 1: Estimation of $\sqrt{I} = 1$ q_i under loss $L_n^{(1)}$ where $q_i \stackrel{i.i.d.}{\sim} 0.4 \{0.5\} + 0.6 \text{Beta}(1, 1)$ and $r_i = 3$.



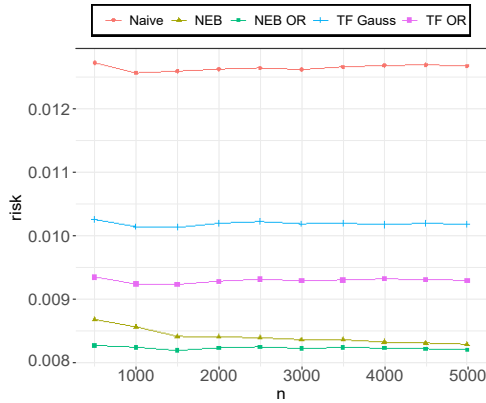
(b) Scenario 1: Estimation of $\sqrt{I} = 1$ q_i under loss $L_n^{(0)}$ where $q_i \stackrel{i.i.d.}{\sim} 0.4 \{0.5\} + 0.6 \text{Beta}(1, 1)$ and $r_i = 3$.



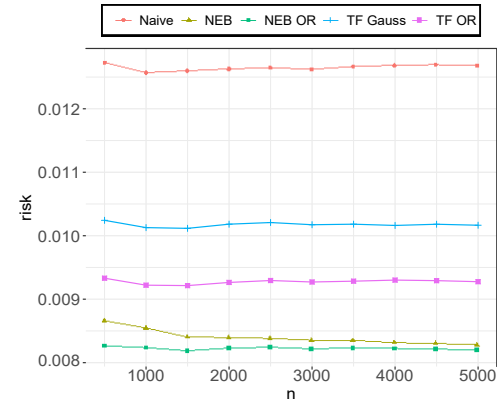
(c) Scenario 2: Estimation of $\sqrt{I} = 1$ q_i under loss $L_n^{(1)}$ where $q_i \stackrel{i.i.d.}{\sim} (1/3) \{0.5\} + (1/3) \{0.7\} + (1/3) \{0.9\}$ and $r_i = 5$.



(d) Scenario 2: Estimation of $\sqrt{I} = 1$ q_i under loss $L_n^{(0)}$ where $q_i \stackrel{i.i.d.}{\sim} (1/3) \{0.5\} + (1/3) \{0.7\} + (1/3) \{0.9\}$ and $r_i = 5$.



(e) Scenario 3: Estimation of $\sqrt{I} = 1$ q_i under loss $L_n^{(1)}$ where $q_i \stackrel{i.i.d.}{\sim} \text{Beta}(5, 2)$ and $r_i = 10$.



(f) Scenario 3: Estimation of $\sqrt{I} = 1$ q_i under loss $L_n^{(0)}$ where $q_i \stackrel{i.i.d.}{\sim} \text{Beta}(5, 2)$ and $r_i = 10$.

Figure 5: Negative Binomial compound decision problem: Risk estimates of the various estimators for Scenarios 1 to 3.

Table 5: The NB compound decision problem under scaled squared error loss: Risk ratios $R_n^{(1)}(\checkmark, \cdot)/R_n^{(1)}(\checkmark, \binom{neb}{1})$ at $n = 5000$ for estimating $\checkmark$.

Method	Scenario		
	1	2	3
Naive	1.27	1.38	1.23
TF Gauss	1.09	1.25	1.14
TF OR	4.94	1.25	1.11
NEB	1.00	1.00	1.00
NEB OR	1.00	1.00	1.00

Table 6: The NB compound decision problem under the squared error loss: Risk ratios $R_n^{(0)}(\checkmark, \cdot)/R_n^{(0)}(\checkmark, \binom{neb}{0})$ at $n = 5000$ for estimating $\checkmark$.

Method	Scenario		
	1	2	3
Naive	2.31	2.02	1.53
TF Gauss	1.74	1.55	1.23
TF OR	1.36	1.33	1.12
NEB	1.00	1.00	1.00
NEB OR	0.99	1.00	0.99

5.1 Estimation of Juvenile Delinquency rates

We consider an application for analysis of the Uniform Crime Reporting Program (UCRP) Database (US Department of Justice and Federal Bureau of Investigation, 2014) that holds county-level counts of arrests and offenses ranging from robbery to weapons violations in 2012. The database is maintained by the National Archive of Criminal Justice Data (NACJD) and is one of the most widely used database for research related to factors that affect juvenile delinquency (JD) across the United States (see for example (Aizer and Doyle Jr, 2015; Damm and Dustmann, 2014; Koski et al., 2018)). A preliminary and important goal in these analyses is to estimate the JD rates based on observed arrest data and determine the counties that are amongst the worst or least affected. However with almost 3,000 counties being evaluated the observed JD counts are susceptible to selection bias, wherein some of the data points are in the extremes merely by chance and traditional estimators may underestimate or overestimate the corresponding delinquency rates, especially in counties with fewer total number of arrests across all age groups.

For the purpose of our analyses, we use the 2012 UCRP data that spans $n = 3,178$ counties in the U.S. and consider estimating the JD rate $\checkmark_i$ for county $i = 1, \dots, n$. The observed data for county i in the year 2012 is the pair (y_{i1}, m_{i1}) which represent, respectively, the number of juvenile arrests and total arrests in county i during that year. We assume that $Y_{i1} | m_{i1}, \checkmark_i \stackrel{\text{i.i.d.}}{\sim} \text{Poi}(m_{i1}\checkmark_i)$ and use the following six competing estimators of $\checkmark = (\checkmark_1, \dots, \checkmark_n)$ from section 4.2: NEB, BGR, KM, TF OR, TF Gauss and Deconv. To assess the performance of the aforementioned estimators we consider predicting the 2014 county level JD counts $Y_2 = (Y_{12}, \dots, Y_{n2})$ and compare their prediction performance under both $L_n^{(0)}$ and $L_n^{(1)}$ losses. In particular for any estimate $\hat{\checkmark}_i$ of $\checkmark_i$, the 2014 predicted JD counts are $\hat{Y}_2 = (\hat{\checkmark}_1 m_{12}, \dots, \hat{\checkmark}_n m_{n2})$ where m_{i2} is the total number of arrests in county i during 2014. The prediction performance of $\hat{\checkmark}$ is then evaluated under loss $L_n^{(k)}(Y_2, \hat{Y}_2)$ for $k \in \{0, 1\}$. The data were cleaned prior to any analyses which ensured that all counties in the year 2012 had at least one arrest (juvenile or not). This resulted in $n = 2803$ counties where all methods are applied to. Let $Y_{2, (k)}^{\hat{\checkmark}}$ denote the n vector of predicted JD counts for 2014 using $\hat{\checkmark}_{(k)}$. Table 7 reports the loss ratios $L_n^{(k)}(Y_2, \hat{Y}_2)/L_n^{(k)}(Y_2, Y_{2, (k)}^{\hat{\checkmark}})$ where a ratio bigger than 1 indicates a smaller prediction loss for $\hat{\checkmark}_{(k)}$. We see that under the scaled squared

Table 7: Loss ratios of the competing methods for predicting Y_2 .

(n = 2,803)		Loss ratios	
Method	k = 1	k = 0	
BGR	1.09	0.97	
KM	1.04	1.01	
Deconv	3.18	1.01	
TF Gauss	1.07	1.00	
TF OR	1.19	1.08	
NEB	1.00	1.00	

error loss ($k = 1$) all five competing estimators to NEB exhibit loss ratios bigger than 1 while BGR outperforms all others under the squared error loss ($k = 0$). Under this loss, however, the NEB estimator continues to provide a better prediction accuracy than TF OR and demonstrates a competitive performance against KM, Deconv and TF Gauss.

5.2 News popularity in social media platforms

Journalists and editors often face the critical task of assessing the popularity of various news items and determining which articles are likely to become popular; hence existing content generation resources can be efficiently managed and optimally allocated to avenues with maximum potential. Due to the dynamic nature of the news articles, popularity is usually measured by how quickly the article propagates (frequency) and the number of readers that the article can reach (severity) through social media platforms like Twitter, Youtube, Facebook and LinkedIn. As such predicting these two aspects of popularity based on early trends is extremely valuable to journalists and content generators (Bandari et al., 2012). In this section, we assess the popularity of several news items based on their frequency

Table 8: Loss ratios of the competing methods for estimating $\sqrt{\cdot}$. News article genre: Economy and social media: Facebook

(n = 3,972)		Loss Ratios	
Method	k = 1	k = 0	
NEB	1.00	1.00	
KM	6.98	> 10	
Deconv	> 10	> 10	
TF Gauss	4.13	3.33	
TF OR	> 10	> 10	

Table 9: Loss ratios of the competing methods for estimating $\sqrt{\cdot}$. News article genre: Microsoft and social media: LinkedIn

(n = 3,850)		Loss Ratios	
Method	k = 1	k = 0	
NEB	1.00	1.00	
KM	> 10	> 10	
Deconv	> 10	> 10	
TF Gauss	9.26	7.34	
TF OR	> 10	> 10	

of propagation and analyze a dataset from Moniz and Torgo (2018) that holds 48 hours worth of social media feedback data on a large collection of news articles since the time of first publication. For the purposes of our analysis, we consider two popular genres of news from this data set: Economy and Microsoft, and examine how frequently these articles

were shared in Facebook and LinkedIn, respectively, over a period of 48 hours from the time of their first publication. Each news article in the data has a unique identifier and 16 consecutive time intervals, each of length 180 minutes, to detect whether the article was shared at least once in that time interval. Let $Z_{ij} = 1$ if article i was shared in time interval j and 0 otherwise, where $i = 1, \dots, n$ and $j = 1, \dots, 16$. Suppose $q_{ij} \in [0, 1]$ denotes the probability that news article i is shared in interval j . Note that in general q_{ij} depends on j since the popularity of any news article evolves with time and therefore Z_{ij} are not independently distributed for $j = 1, \dots, 16$. However for the purposes of this analysis, we let $q_{ij} = q_i$ for $j = 1, \dots, 16$ and assume that for each i , Z_{ij} are independent realizations from $\text{Ber}(q_i)$. It then follows that $Y_{i1} = \sum_{j=1}^8 Z_{ij} \stackrel{\text{i.i.d.}}{\sim} \text{Bin}(8, q_i)$. To assess the popularity of article i we estimate its odds of sharing in the remaining 8 time intervals ($j = 9, \dots, 16$) and consider the following 5 estimators from section 4.3: NEB, KM, Deconv, TF Gauss and TF OR. Tables 8 and 9 report the loss ratios $L_n^{(k)}(\check{\cdot}) / L_n^{(k)}(\check{\cdot}, \hat{\cdot}_{(k)}^{\text{neb}})$ for any estimator of $\check{\cdot}$ where $\check{\cdot}_i = Y_{i2} / (8 - Y_{i2})$ and $Y_{i2} = \sum_{j=9}^{16} Z_{ij}$. We observe that all four competitors to the NEB estimator exhibit loss ratios substantially bigger than 1 under both the losses. The relatively poorer performance of KM and Deconv in this example stems from the fact that in this application $Y_{i1} = 8$ for several news articles. For those news articles both KM and Deconv return disproportionately bigger estimates of $\check{\cdot}_i$ which explains their relatively larger estimation loss.

6. Discussion

In this paper we propose a Nonparametric Empirical Bayes framework for compound estimation in the discrete linear exponential family. The proposed estimator is consistent and presents a unified framework for compound estimation in the DLE family by estimating the Bayes shrinkage factors via a convex program that can easily incorporate various structural constraints, such as monotonicity, into the data driven decision rule. Our numerical evidence suggests that across many settings the NEB estimator has a substantially better risk performance than the competing approaches considered here.

We conclude this article with two open issues. First, in large scale compound estimation problems, one is often interested in constructing confidence intervals for the EB shrinkage estimators. Recall that for any coordinate i , $\hat{\cdot}_{(k),i}^{\text{neb}}$ is non-linear and a biased estimate of $\check{\cdot}_i$. While the CLT for minimum KSD estimators in Barp et al. (2019) will be important for deriving the asymptotic distribution of the NEB estimator, the main challenge in constructing confidence intervals lies in accounting for the bias in $\hat{\cdot}_{(k),i}^{\text{neb}}$ for estimating $\check{\cdot}_i$. In the absence of any information on the prior G , it is not immediately clear how to accurately characterize this bias. Notable recent developments include “de-biasing” the EB estimator (Ignatiadis and Wager, 2019) or assuming a Normal distribution on G but using a carefully constructed larger critical value to account for the bias due to shrinkage (Armstrong et al., 2020). Secondly, the NEB estimation framework handles both regular and scaled squared error losses and it is desirable to construct such empirical Bayes shrinkage estimators for other asymmetric losses such as the Linex loss (Varian, 1975) and the Generalized absolute loss function (Koenker and Bassett Jr, 1978). As part of our future research, we will be interested in pursuing these aforementioned directions.

Acknowledgments

We thank the Action Editor and three anonymous referees whose comments have substantially improved the quality of the paper. G. Mukherjee's work was partly supported by the NSF grant DMS-1811866. W. Sun's work was partly supported by the NSF grant DMS-1712983. Q. Liu's work is supported in part by NSF CRII 1830161 and NSF CAREER 1846421.

Supplementary Material for “EB Estimation in Discrete Linear Exponential Family”

In this supplement, we first present in Appendix A the results for the NEB estimator under the squared loss, then in Appendix B we provide the proofs and technical details of all theories in the main text and Appendix A.

A. Results Under the Squared Error Loss

A.1 The NEB estimator

In this section we discuss the estimation of $w_p^{(0)}$ that appear in lemma 1 under the usual squared error loss ($k = 0$). Let Y be a non-negative integer-valued random variable with probability mass function (pmf) p and define

$$h_0^{(0)}(y) = \frac{y + 1}{w_p^{(0)}(y)} \quad y, y \in \{0\} \cup \mathbb{N} \quad (14)$$

Suppose $K(y, y^0) = \exp\{-0.5^{-1}(y - y^0)^2\}$ be the positive definite Gaussian kernel with bandwidth parameter $2 \leftarrow \leftarrow$ where $\leftarrow$ is a compact subset of $\mathbb{R}^+$ bounded away from 0. Given observations $y = (y_1, \dots, y_n)$ from model (2), let $h_0^{(0)} = (h_0^{(0)}(y_1), \dots, h_0^{(0)}(y_n))$ and define the following $n \rightarrow n$ matrices: $n^2 K = [K(y_i, y_j)]_{ij}$, $n^2_0 K = [\frac{y_i}{y_j} K(y_i, y_j + 1)]_{ij}$ and $n^2_{-2} K = [\frac{y_i}{y_j} K(y_i, y_j)]_{ij}$ where $\frac{y_i}{y_j} K(y, y^0) = K(y + 1, y^0) - K(y, y^0)$ and $\frac{y_i}{y_j} K(y, y^0) = \frac{y^0}{y} K(y, y^0) = \frac{y}{y^0} K(y, y^0)$.

Definition 3 (NEB estimator of $\sqrt{\cdot}$). Consider the DLE Model (2) with loss $\psi^{(0)}(\sqrt{\cdot})$. For a fixed $2 \leftarrow$, let $\hat{w}^{(0)}(\cdot) = (y_i + 1)/(\hat{y}_i + h^{(0)}(\cdot))$ and $\hat{h}^{(0)}(\cdot) = \hat{h}^{(0)}(\cdot), \dots, \hat{h}^{(0)}(\cdot)$ be

the solution to the following quadratic optimization problem:

$$\min_{h \in H_n} h^T K h + 2h^T K y + y^T_{-2} K y, \quad (15)$$

where $H_n = \{h = (h_1, \dots, h_n) : Ah = b, Ch = d\}$ is a convex set and A, C, b, d are known real matrices and vectors that enforce linear constraints on the components of h . Then the NEB estimator for a fixed $\leftarrow$ is given by $\hat{w}^{neb}_{(0)}(\cdot) = \hat{w}^{neb}_{(0),i}(\cdot) : 1 \leq i \leq n$, where

$$\hat{w}^{neb}_{(0),i}(\cdot) = \frac{a_{y_i}/a_{y_i+1}}{\hat{w}_i^{(0)}(\cdot)}, \quad \text{if } y_i \in \{0, 1, 2, \dots\}$$

Remark 5. In problem (15) the linear inequality constraints $Ah \leq b$ can be used to impose structural constraints on the NEB decision rule $\eta_{(0)}^{eb}(\cdot)$. The structural constraints may take the form of monotonicity constraints so that $\eta_{(0),(1)}^{eb}(\cdot) \leq \dots \leq \eta_{(0),(n)}^{eb}(\cdot)$ for $y_{(1)} \leq y_{(2)} \leq \dots \leq y_{(n)}$. In particular, when $Y_i | \sqrt{\cdot} \leftarrow \text{Poi}(\sqrt{\cdot})$ then $\eta_{(0),i}^{eb}(\cdot) = y_i + h^{(0)}(\cdot)$ and the monotonicity constraints in this setting will imply

$$\hat{h}_{(i)}^{(0)}(\cdot) + \hat{h}_{(i+1)}^{(0)}(\cdot) \geq y_{(i)} - y_{(i+1)}, \text{ for } 1 \leq i \leq (n-1)$$

These $n-1$ linear inequality constraints may be imposed with an $(n-1) \times n$ matrix A and an $n-1$ column vector b such that for $1 \leq i \leq (n-1)$ and $1 \leq r \leq n$,

$$A(i, r) = \begin{cases} 1, & \text{when } y_r = y_{(i)} \\ 1, & \text{when } y_r = y_{(i+1)} \\ 0, & \text{otherwise} \end{cases}$$

and $b_i = y_{(i)} - y_{(i+1)}$.

The equality constraints $Ch = d$ may accommodate instances of ties for which we require $\hat{h}_i^{(0)}(\cdot) = \hat{h}_j^{(0)}(\cdot)$ whenever $y_i = y_j$.

Theorem 4. Let $K(\cdot, \cdot)$ be the positive definite Gaussian kernel with bandwidth parameter $h_n \rightarrow 0$. If $\lim_{n \rightarrow \infty} c_n n^{-1/2} \log^4 n = 0$ then, under assumptions (A1)–(A3), we have for any $h_n \rightarrow 0$,

$$\lim_{n \rightarrow \infty} P \left(\frac{1}{n} \sum_{i=1}^n \hat{w}_n^{(0)}(\cdot) - w_p^{(0)}(\cdot) \right)^2 \leq c_n \frac{1}{n} = 0, \text{ for any } h_n \rightarrow 0$$

> 0 where $w_n^{(0)}(\cdot) = [(Y_i + \hat{1})/(h^{(0)}(\cdot) + Y_i)]_i$.

We now provide some motivation behind the minimization problem in definition 3 for estimating the ratio functionals $w_p^{(0)}$. Let $\hat{M}_{p,n}(h) = h^T K h + 2h^T K y + y^T K y$ be the objective function in equation (15). Suppose $\tilde{p}$ be a probability mass function on the support of Y and define

$$S[\tilde{p}](p) = E_p \left(\tilde{h}^{(0)}(Y) - h_0^{(0)}(Y) \right) K(Y + 1, Y_0 + 1) (\tilde{h}^{(0)}(Y_0) - h_0^{(0)}(Y_0)) \quad (16)$$

where $h_0^{(0)}, \tilde{h}^{(0)}$ are as defined in equation (14) and Y, Y_0 are i.i.d copies from the marginal distribution that has mass function p . $S[\tilde{p}](p)$ in equation (16) is the Kernelized Stein's Discrepancy (KSD) measure that can be used to distinguish between two distributions with mass functions $p, \tilde{p}$ such that $S[\tilde{p}](p) = 0$ and $S[\tilde{p}](p) = 0$ if and only if $p = \tilde{p}$ (Liu et al., 2016; Chwialkowski et al., 2016; Yang et al., 2018). Moreover for i.i.d. copies (Y, Y_0) from p , it can be shown that

$$S[\tilde{p}](p) = \frac{1}{n(n-1)} E_p \sum_{1 \leq i < j \leq n} [\tilde{h}^{(0)}(Y_i) - \tilde{h}^{(0)}(Y_j)](Y_i, Y_j)$$

where $Y = (Y_1, \dots, Y_n)$ is a random sample from the marginal distribution with mass function p and under the squared error loss $\int [\tilde{h}^{(0)}(u) - \tilde{h}^{(0)}(v)](u, v)$ is the positive definite

kernel function

$$\tilde{h}^{(0)}(u)\tilde{h}^{(0)}(v)K(u, v) + \tilde{h}^{(0)}(u)v_{-v}K(u+1, v) + \tilde{h}^{(0)}(v)u_{-u}K(u, v+1) + uv_{-u,v}K(u, v). \quad (17)$$

An empirical evaluation scheme for $S[\tilde{p}](p)$ is given by $S[\tilde{p}](\hat{p}_n)$ where

$$S[\tilde{p}](\hat{p}_n) = \frac{1}{n^2} \sum_{i,j=1}^n [\tilde{h}^{(0)}(y_i), \tilde{h}^{(0)}(y_j)](y_i, y_j) \quad (18)$$

where $\hat{p}_n$ is the empirical CDF. Note that $\sum [\tilde{h}^{(0)}(u), \tilde{h}^{(0)}(v)](u, v)$ in equation (17) involves $\tilde{p}$ only through $\tilde{h}^{(0)}$ and may analogously be denoted by $\sum [\tilde{h}(u), \tilde{h}(v)](u, v)$ where we have dropped the superscript from $\tilde{h}$ that indicates that the loss in question is the regular squared error loss. This slight abuse of notation is harmless as the discussion in this section is geared towards the squared error loss only.

Under the compound estimation framework of model (2), our goal is to estimate $h_0^{(0)}$. To do that we minimize $S[\tilde{p}](\hat{p}_n)$ in equation (18) with respect to the unknowns $\tilde{h} = (\tilde{h}(y_1), \dots, \tilde{h}(y_n))$. Note that $S[\tilde{p}](\hat{p}_n)$ is exactly the objective function $\hat{M}_{\tilde{p},n}(\tilde{h})$ of the quadratic program (15) with optimisation variables $h_i \in \tilde{h}(y_i)$ for $i = 1, \dots, n$.

A.2 Bandwidth choice and asymptotic properties

We propose the following asymptotic risk estimate $ARE_n^{(0)}(\cdot)$ of the true loss of $\tilde{h}_0^{(0)}$ in the DLE model (2).

Definition 4 (ARE of $\tilde{h}_0^{(0)}$ in the DLE model). Suppose $Y_i \mid \sqrt{c_n} \xrightarrow{d} \text{DLE}(\sqrt{c_n})$. Under the loss $\ell^{(0)}(\sqrt{c_n}, \cdot)$ an asymptotic risk estimate of the true loss of $\tilde{h}_0^{(0)}$ is

$$ARE_n^{(0)}(\cdot, y) = \frac{1}{n} \sum_{i=1}^n [\tilde{h}_0^{(0)}(y_i)]^2 - 2 \sum_{i=1}^n \tilde{h}_0^{(0)}(y_i)$$

where

$$\tilde{h}_0^{(0)}(y_i) = \tilde{h}_0^{(0),j_i}(\cdot)(a_{y_i-1}/a_{y_i}), \quad y_i = 1, 2, \dots$$

with $j_i \in \{1, \dots, n\}$ such that $y_{j_i} = y_i - 1$.

An estimate of the tuning parameter based on $ARE_n^{(0)}(\cdot, y)$ is given by:

$$\hat{c}_n = \arg \min_{2^{\kappa}} ARE_n^{(0)}(\cdot, y) \quad (19)$$

where a choice of $\kappa = [10, 10^2]$ worked well in the simulations and real data analyses of sections 4 and 5. Lemma 2 continues to provide the large-sample properties of the proposed $ARE_n^{(0)}$ criteria for the Poisson and Binomial distributions provided c_n is a sequence that satisfies $\lim_{n \rightarrow \infty} c_n n^{-1/4} \log^4 n = 0$.

To analyze the quality of the estimates $\hat{c}_n$ obtained from equation (19), we consider an oracle loss estimator $\tilde{h}_0^{(0),orc} = \tilde{h}_0^{(0),orc}(\cdot)$ where

$$\tilde{h}_0^{(0),orc} = \arg \min_{2^{\kappa}} L_n^{(0)}(\sqrt{c_n}, \tilde{h}_0^{(0),orc}(\cdot))$$

and Lemma 3 establishes the asymptotic optimality of $\hat{\theta}$ obtained from equation (19). In theorem 5 below we provide decision theoretic guarantees on the NEB estimator and show that the average squared error between $\hat{\theta}_{(0)}^{neb}(\hat{\theta})$ and $\hat{\theta}_{(0)}$ is asymptotically small.

Theorem 5. Under the conditions of Theorem 4, if $\lim_{n \rightarrow \infty} c_n n^{-1/2} \log^6 n = 0$ then, for the Poisson and the Binomial model,

$$\frac{c_n}{n} \|\hat{\theta}_{(0)}^{neb}(\hat{\theta}) - \hat{\theta}_{(0)}\|_2^2 = o_p(1).$$

Furthermore, under the same conditions, we have,

$$\lim_{n \rightarrow \infty} E \left[L_n^{(0)}(\sqrt{\cdot}, \hat{\theta}_{(0)}^{neb}(\hat{\theta})) - L_n^{(0)}(\sqrt{\cdot}, \hat{\theta}_{(0)}) \right] = 0.$$

B. Technical Details and Proofs

We will begin this section with some notations and then state three lemmata that will be used in proving the statements discussed in Section 3.

Let $c_0, c_1, \dots$ denote some generic positive constants which may vary in different statements. Let $D_n = \{0, 1, 2, \dots, d \log n\}$ where $d \log n$ denotes the smallest integer greater or equal to x . Given a random sample $(Y_1, \dots, Y_n)$ from model (2) denote B_n to be the event $\{\max_{1 \leq i \leq n} Y_i \leq C \log n\}$ where C is the constant given by lemma 4 below under assumption (A2).

Lemma 4. Assumption (A2) implies that with probability tending to 1 as $n \rightarrow \infty$,

$$\max(Y_1, \dots, Y_n) \leq C \log n$$

where $C > 0$ is a constant depending on θ .

Our next lemma below is a statement on the pointwise Lipschitz stability of the optimal solution $\hat{\theta}_n^{(k)}(\theta)$ under perturbations on the parameter $\theta \in \mathbb{R}^2$. See, for example, Bonnans and Shapiro (2013) for general results on the stability and sensitivity of parametrized optimization problems.

Lemma 5. Let $\hat{\theta}_n^{(k)}(\theta)$ be the solution to problems (8) and (15), respectively, for $k \in \{0, 1\}$ and for some $\theta \in \mathbb{R}^2$. Then, under Assumption (A3), there exists a constant $L > 0$ such that for any $\theta \in \mathbb{R}^2$ the solution $\hat{\theta}_n^{(k)}(\theta)$ to problems (8) and (15) satisfies

$$\|\hat{\theta}_n^{(k)}(\theta) - \hat{\theta}_n^{(k)}(\theta_0)\|_2 \leq L \|\theta - \theta_0\|_2$$

Lemma 6. Suppose $Y \mid \sqrt{\cdot} \stackrel{hd}{\rightarrow} \text{DLE}(\sqrt{\cdot})$. Then the following hold:

$$\begin{aligned} E_{Y \mid \sqrt{\cdot}} \left[\frac{1}{n} \sum_{i=1}^n \left[\hat{\theta}_{(1)}^{(Y_i+1)} \right]^2 \frac{a_{Y_i+1}}{a_{Y_i}} \right] &= \frac{[\hat{\theta}_{(1)}^{(Y)}]^2}{o} = 0 \text{ and} \\ E_{Y \mid \sqrt{\cdot}} \left[\hat{\theta}_{(0)}^{(Y-1)} \frac{a_{Y-1}}{a_Y \sqrt{\cdot}} \right] &= \sqrt{\hat{\theta}_{(0)}^{(Y)}} = 0. \end{aligned}$$

The proofs of Lemmata 4, 5 and 6 are available in appendix B.8.

But from assumption (A1), $E_p \sup_{2 \leftarrow} |\tilde{h}(U), \tilde{h}(V)|(U, V)| < 1$ and together with assumption (A2) and proof of lemma 4, it follows that the terms in the display above are $O(n^{-\frac{1}{2}})$ for some $\frac{1}{2} > \frac{1}{2}$.

Now fix an $\epsilon > 0$ and let $c_n = n^p / \log^2 n$. Since $E_p \sup_{2 \leq n} |\hat{M}_{\cdot, n}(\tilde{h}) - \overline{M}(\tilde{h})|$ is $O(\log^2 n / n^p)$ there exists a finite constant $M > 0$ and an N_1 such that $c_n E_p \sup_{2 \leq n} |\hat{M}_{\cdot, n}(\tilde{h}) - \overline{M}(\tilde{h})| \leq M$ for all $n \geq N_1$. Moreover since $\sup_{2 \leq n} |M(\tilde{h}) - \overline{M}(\tilde{h})| \rightarrow 0$ as $n \rightarrow \infty$, there exists an N_2 such that $\sup_{2 \leq n} |M(\tilde{h}) - \overline{M}(\tilde{h})| \leq M/c_n$ for all $n \geq N_2$. Thus with $t = 4M/\epsilon$, we have $P(c_n \sup_{2 \leq n} |\hat{M}_{\cdot, n}(\tilde{h}) - \overline{M}(\tilde{h})| > t) < \epsilon$ for all $n \geq \max(N_1, N_2)$ which suffices to prove the desired result.

B.3 Proofs of Theorems 2 and 4

We will first prove Theorem 2. Note that from equation (7),

$$\hat{\mathbf{w}}_n^{(1)}(\cdot) - \mathbf{w}_p^{(1)} \Big|_2^2 = \hat{\mathbf{h}}_n^{(1)}(\cdot) - \mathbf{h}_0^{(1)} \Big|_2^2.$$

Now from assumption (A3) and for any $\epsilon > 0$, there exists a $\delta > 0$ such that for any 2

$$\begin{aligned} \leftarrow, \frac{h}{h} c_n h^{(1)}()_0 \quad \hat{ }_n \quad h^{(1)}_0 \quad \bullet \quad \frac{i}{?} P_{c_n} M \\ (h^{(1)}) \quad M \quad (h^{(1)})^o \quad i \end{aligned}$$

But the right hand side is upper bounded by the sum of $P_{c_n} \hat{M}_{,n}(\hat{h}_n^{(1)}) / 3$, $\hat{M}_{,n}(h_0^{(1)}) / 3$ and $P_{c_n} \hat{M}_{,n}(h_0^{(1)}) / 3$. From theorem 1, the first and third terms go to zero as $n \rightarrow \infty$ while the second term is zero since $\hat{M}_{,n}(\hat{h}_n^{(1)}) \rightarrow \hat{M}_{,n}(h_0^{(1)})$ as $h_0^{(1)} \rightarrow H_n$. This proves the statement of theorem 2.

To prove theorem 4 first note that from equation (14),

$$\mathbf{w}_n^{(0)}(\cdot) \mathbf{w}_p^{(0)} = \frac{X_n^n \mathbf{Y}_i + 1}{\prod_{j=1}^n [\mathbf{Y}_i + \hat{h}_{n,i}^{(0)}(\cdot)] [\mathbf{Y}_i + h_{0,i}^{(0)}]} \mathbf{h}_{n,i}^{(0)}(\cdot) \mathbf{h}_{0,i}^{(0)}.$$

From assumption (A2) and Lemma 4, there exists a constant $c_0 > 0$ such that for large n , $\max_{1 \leq i \leq n} (Y_i + 1) \leq c_0 \log n$ with high probability. Moreover for $i = 1, \dots, n$, since $\hat{w}_{n,i}^{(0)}(\cdot) > 0$ for every $\cdot \in \mathbb{R}$ and $w_{p,i}^{(0)} > 0$, equation (14) implies $\hat{h}_{n,i}^{(0)}(\cdot) + Y_i > 0$ and $h_{0,i}^{(0)} + Y_i > 0$. Thus, conditional on the event $\{\max_{1 \leq i \leq n} (Y_i + 1) \leq c_0 \log n\}$ we have for some constant $c_1 > 0$,

$$\hat{w}_n^{(0)}(\cdot) - w_p^{(0)} \leq c_1 \log^2 n \cdot h_n^{(0)}(\cdot) - h_0^{(0)} \leq$$

and for any $\epsilon > 0$,

$$P^h \frac{c_n}{n \log^2 n} \hat{w}_n^{(0)}(\cdot) \leq w_p^{(0)} \leq P^i \frac{c_n}{n} \hat{h}_n^{(0)}(\cdot) \leq h_0^{(0)} \leq P^i.$$

The proof of the statement of theorem 4 then follows from the proof of theorem 2 above and lemma 4.

B.4 Proofs of Theorems 3 and 5

We will prove Theorem 3 while the proof of Theorem 5 will follow using similar arguments and Theorem 4.

Note that,

$$\frac{\hat{\eta}_{(1)}^{neb}(\hat{\eta})}{\hat{\eta}_{(1)}^2} = \frac{X^n}{2} \sum_{i=1}^n \frac{a_{Y_i} - 1/a_{Y_i}}{\hat{w}_{n,i}(\hat{\eta})w_{p,i}} \hat{w}_{n,i}^{(1)}(\hat{\eta}) w_{p,i}^{(1)} \hat{\eta}_{(1)}^2.$$

Now, $\hat{w}_{n,i}^{(1)}(\hat{\eta}) > 0$ for every $i \in \mathbb{N}$ and $w_{p,i}^{(1)} > 0$. This fact along with assumption (A2) and lemma 4 imply that there exists a constant $c_0 > 0$ such that $\|\hat{\eta}_{(1)}^{neb}(\hat{\eta}) - \hat{\eta}_{(1)}^{(1)}\|_2^2 \leq c_0 \log^2 n k w_n^{(1)}(\hat{\eta}) w_p^{(1)} k_2^2$. The first result thus follows from the above inequality and Theorem 2.

To prove the second part of the theorem, note that $|L_n^{(1)}(\check{\eta}, \hat{\eta}_{(1)}) - L_n^{(1)}(\check{\eta}, \hat{\eta}_{(1)}^{neb}(\hat{\eta}))|$ equals

$$\frac{r}{L_n^{(1)}(\check{\eta}, \hat{\eta}_{(1)})} - \frac{r}{L_n^{(1)}(\check{\eta}, \hat{\eta}_{(1)}^{neb}(\hat{\eta}))} = \frac{r}{L_n^{(1)}(\check{\eta}, \hat{\eta}_{(1)}) + L_n^{(1)}(\check{\eta}, \hat{\eta}_{(1)}^{neb}(\hat{\eta}))}$$

and Triangle inequality implies

$$\begin{aligned} \frac{r}{L_n^{(1)}(\check{\eta}, \hat{\eta}_{(1)})} - \frac{r}{L_n^{(1)}(\check{\eta}, \hat{\eta}_{(1)}^{neb}(\hat{\eta}))} &\leq \frac{r}{L_n^{(1)}(\check{\eta}, \hat{\eta}_{(1)})} \leq \frac{r}{\frac{1}{n} \sum_{i=1}^n \frac{a_{Y_i} - 1/a_{Y_i}}{\hat{w}_{n,i}(\hat{\eta})w_{p,i}} \hat{w}_{n,i}^{(1)}(\hat{\eta}) w_{p,i}^{(1)} \hat{\eta}_{(1)}^2} \\ &\leq \frac{c_0}{n} \frac{\hat{\eta}_{(1)}^{neb}(\hat{\eta})}{\hat{\eta}_{(1)}^2} = O_p \left(\frac{\log^2 n}{n^{1/4}} \right) \end{aligned} \quad (22)$$

from the first part of theorem 3. Thus, it follows from equation (22) that

$$\frac{r}{L_n^{(1)}(\check{\eta}, \hat{\eta}_{(1)}^{neb}(\hat{\eta}))} \leq \frac{r}{L_n^{(1)}(\check{\eta}, \hat{\eta}_{(1)})} + O_p \left(\frac{\log^2 n}{n^{1/4}} \right)$$

and

$$L_n^{(1)}(\check{\eta}, \hat{\eta}_{(1)}) - L_n^{(1)}(\check{\eta}, \hat{\eta}_{(1)}^{neb}(\hat{\eta})) \leq 4 \frac{r}{L_n^{(1)}(\check{\eta}, \hat{\eta}_{(1)})} - \frac{r}{L_n^{(1)}(\check{\eta}, \hat{\eta}_{(1)})} = \frac{r}{L_n^{(1)}(\check{\eta}, \hat{\eta}_{(1)})} \quad (23)$$

Now from assumption (A2) and the proof of Lemma 4, $L_n^{(1)}(\check{\eta}, \hat{\eta}_{(1)}) \geq c_1 \log^2 n$ for some constant $c_1 > 0$. Therefore, together with equations (22) and (23) we have

$$L_n^{(1)}(\check{\eta}, \hat{\eta}_{(1)}) - L_n^{(1)}(\check{\eta}, \hat{\eta}_{(1)}^{neb}(\hat{\eta})) = O_p \left(\frac{\log^3 n}{n^{1/4}} \right).$$

Define $Z_n(\hat{\eta}) = L_n^{(1)}(\check{\eta}, \hat{\eta}_{(1)}) - L_n^{(1)}(\check{\eta}, \hat{\eta}_{(1)}^{neb}(\hat{\eta}))$. We have already shown that $Z_n(\hat{\eta}) \rightarrow 0$ in probability as $n \rightarrow \infty$. Moreover, under the Poisson model with $\hat{\eta}_{(1),i}^{neb}(\hat{\eta}) \geq Y_i/d_1$

and

$\hat{\eta}_{(1),i} \geq Y_i/w_p^{(1)}(Y_i)$ where $w_p^{(1)}(Y_i) > 0$, we have for some positive constants c_0, c_1

$$|Z_n(\hat{\eta})| \leq \frac{1}{n} \sum_{i=1}^n c_0 Y_i + c_1 Y_i^2 := U_n. \quad (24)$$

Now, under the Poisson model and from assumption (A2), $\sup_n E(U_n^{1+}) < 1$ for some $\epsilon > 0$. Thus $\{U_n\}$ is uniformly integrable. Therefore, from equation (24), $\{Z_n(\hat{\gamma})\}$ is uniformly integrable and along with the fact that $Z_n(\hat{\gamma}) \rightarrow 0$ in probability as $n \rightarrow \infty$, we have $E|Z_n(\hat{\gamma})| \rightarrow 0$ as $n \rightarrow \infty$. This proves the desired result under the Poisson model. For the Binomial model, with $m < 1$, the result continues to hold since $\hat{\gamma}_{(1),i}^{(1)} \rightarrow m/d_1$ and $\hat{\gamma}_{(1),i}^{(1)} \rightarrow m/w_p^{(1)}(Y_i)$ where $w_p^{(1)}(Y_i) > 0$. Thus, $|Z_n(\hat{\gamma})| < 1$ and so $\{Z_n(\hat{\gamma})\}$ is still uniformly integrable.

B.5 Proof of Lemma 2 - Binomial model

We will first prove the two statements of lemma 2 under the scaled squared error loss and conditional on the event B_n which is the event that $\{\max_{1 \leq i \leq n} Y_i \leq C \log n\}$. Under assumption (A2), lemma 4 guarantees that B_n holds with high probability. Throughout the proof, we shall denote $d_1 := \inf_{2 \leq k} \inf_{1 \leq i \leq n} (1 - h_{n,i}^{(k)}(\gamma)) > 0$, $d_2 := \inf_{2 \leq k} \inf_{1 \leq i \leq n} \hat{w}_{n,i}^{(0)}(\gamma) > 0$ and assume $m < 1$. Moreover, we will use the fact that under the Binomial model, $|h_{n,i}^{(k)}(\gamma)| \leq 1$ uniformly in $2 \leq k$. This is a consequence of $d_1 > 0$, $d_2 > 0$ and $m < 1$. The proof for the squared error loss will follow from similar arguments and we will highlight only the important steps.

Proof of statement 1 (Binomial model) for the scaled squared error loss ($k = 1$)

First note that under the Binomial model, $y_i \leq m$ and $a_{y_i+1}/a_{y_i} = (m - y_i)/(y_i + 1) \leq m$. Now,

$$\begin{aligned} \sup_{2 \leq k} \text{ARE}_n^{(1)}(\gamma, Y) & \rightarrow_n^{(1)}(\gamma, \hat{\gamma}_{(1)}^{(1)}) = \sup_{2 \leq k} \frac{1}{n} \sum_{i=1}^n \left[\hat{\gamma}_{(1),j_i}^{(1)}(\gamma) \right]^2 \frac{a_{y_i+1}}{a_{y_i}} \frac{[\hat{\gamma}_{(1),i}^{(1)}(\gamma)]^2}{i} \\ & \leq \sup_{2 \leq k} \frac{m}{n} \sum_{i=1}^n \left[\hat{\gamma}_{(1),j_i}^{(1)}(\gamma) \right]^2 \frac{[\hat{\gamma}_{(1),j_i}^{(1)}(\gamma)]^2}{i} + \sup_{2 \leq k} \frac{c_0}{n} \sum_{i=1}^n \left[\hat{\gamma}_{(1),i}^{(1)}(\gamma) \right]^2 \frac{[\hat{\gamma}_{(1),i}^{(1)}(\gamma)]^2}{i} \\ & + \frac{1}{n} \sum_{i=1}^n \left[\hat{\gamma}_{(1),j_i}^{(1)}(\gamma) \right]^2 \frac{a_{y_i+1}}{a_{y_i}} \frac{[\hat{\gamma}_{(1),i}^{(1)}(\gamma)]^2}{i} =: T_1 + T_2 + T_3. \end{aligned} \quad (25)$$

Here we have used the fact that $a_{y_i+1}/a_{y_i} \leq m$ and since $\sqrt{i} > 0$, $1/\sqrt{i} < c_0$ for some positive constant c_0 . Consider the term T_3 in equation (25) above and define

$$V_i = \left[\hat{\gamma}_{(1),j_i}^{(1)}(\gamma) \right]^2 \frac{a_{y_i+1}}{a_{y_i}} \frac{[\hat{\gamma}_{(1),i}^{(1)}(\gamma)]^2}{i}.$$

Note that from lemma 6, $EV_i = 0$. Moreover, $\sqrt{V_i}$ are independent and $E|V_i|^2 < 1$ since $|V_i| \leq c_1 m^3$ for some constant $c_1 > 0$. The bound on $|V_i|$ follows from the fact that for any i , $\hat{\gamma}_{(1),i}^{(1)}(\gamma) \leq \hat{\gamma}_{(1),i}^{(1)}(\gamma) \leq m/w_p^{(1)}(\gamma)$ where $w_p^{(1)}(\gamma) > 0$. So, T_3 is $O_p(n^{-1/2})$.

We now consider the second term T_2 in equation (25) and define $Z_n(\cdot) = \frac{1}{n} \sum_{i=1}^n \{ \hat{\theta}_{(1),i}^{neb}(\cdot) \}$. Note that under the Binomial model $\hat{\theta}_{(1),i}^{neb}(\cdot) \in \mathbb{M}/d_1$ and so for any $\theta \in \mathbb{M}$

$$\frac{c_0}{n} \sum_{i=1}^n [\hat{\theta}_{(1),i}^2 - [\hat{\theta}_{(1),i}^{neb}(\cdot)]^2] \leq c_2 Z_n(\cdot) \leq \frac{c_2}{n} \sum_{i=1}^n \hat{\theta}_{(1),i}^{neb}(\cdot) \leq \frac{c_2}{n} \sum_{i=1}^n \hat{\theta}_{(1),i}^2 \quad (26)$$

for some constant $c_2 > 0$. The last term in the inequality above is $O_p(\log^2 n/n^{1/4})$ from the first part of theorem 3. Next for a perturbation θ^0 of θ such that $(\theta, \theta^0) \in \mathbb{M} := [l, u]$, we will bound the increments $|Z_n(\theta) - Z_n(\theta^0)|$. To that effect, note that

$$n Z_n(\theta) - Z_n(\theta^0) \leq \sum_{i=1}^n \hat{\theta}_{(1),i}^{neb}(\theta) - \sum_{i=1}^n \hat{\theta}_{(1),i}^{neb}(\theta^0) \leq \frac{m}{d_1^2} h_n^{(1)}(\theta) - \hat{h}_n^{(1)}(\theta^0).$$

Now from lemma 5 we know that

$$k \hat{h}_n^{(1)}(\theta) - \hat{h}_n^{(1)}(\theta^0) k_1 \leq n^{1/2} c_1 \sup_{h \in N(\hat{h}_n^{(1)}(\theta^0))} k r_{\hat{h}_n^{(1)}(\theta)}^2, \hat{M}_{\theta,n}(h) + o(1) k_2.$$

However, under the Binomial model with $m < 1$, $|\hat{h}_{n,i}^{(1)}(\cdot)| < 1$ uniformly in $\theta \in \mathbb{M}$. Thus, the supremum in the display above is finite. So,

$$Z_n(\theta) - Z_n(\theta^0) \leq \frac{c_3}{n} \theta^0.$$

Thus $Z_n(\cdot)$ is Lipschitz continuous in $\theta \in \mathbb{M}$ and, along with the compactness of $\mathbb{M}$, it implies that there exists a $(\theta_1, \theta_2) \in \mathbb{M}$ such that $\sup_{\theta \in \mathbb{M}} Z_n(\theta) = Z_n(\theta_1)$ and $\inf_{\theta \in \mathbb{M}} Z_n(\theta) = Z_n(\theta_2)$. Therefore, taking the supremum with respect to θ in equation (26) and using the first part of theorem 3 suffice to prove that T_2 is $O_p(\log n/n^{1/4})$. Finally, the first term T_1 in equation (26) is $O_p(\log n/n^{1/4})$ which follows using similar arguments for the term T_2 . Therefore, we have the desired result that $\sup_{\theta \in \mathbb{M}} |\text{ARE}^{(1)}(\theta, Y)| \xrightarrow{P} |\text{ARE}^{(1)}(\theta, Y)|$ is $O_p(\log^2 n/n^{1/4})$.

Proof of statement 2 (Binomial model) for the scaled squared error loss ($k = 1$)

From Triangle inequality, $\sup_{\theta \in \mathbb{M}} |\text{ARE}_n^{(1)}(\theta, Y)| \leq E_n^{(1)}(\sqrt{\cdot}, \hat{\theta}_{(1)}^{neb}(\cdot)) \leq T_1 + T_2 + T_3 + T_4$ where $T_1 = \sup_{\theta \in \mathbb{M}} |\text{ARE}_n^{(1)}(\theta, Y)| \leq E_n^{(1)}(\sqrt{\cdot}, \hat{\theta}_{(1)}^{neb}(\cdot))$, $T_2 = \sup_{\theta \in \mathbb{M}} |E_n^{(1)}(\sqrt{\cdot}, \hat{\theta}_{(1)}^{neb}(\cdot)) - E_n^{(1)}(\sqrt{\cdot}, \hat{\theta}_{(1)}^{neb}(\cdot))|$, $T_3 = |E_n^{(1)}(\sqrt{\cdot}, \hat{\theta}_{(1)}^{neb}(\cdot)) - E_n^{(1)}(\sqrt{\cdot}, \hat{\theta}_{(1)}^{neb}(\cdot))|$ and $T_4 = \sup_{\theta \in \mathbb{M}} |E_n^{(1)}(\sqrt{\cdot}, \hat{\theta}_{(1)}^{neb}(\cdot)) - E_n^{(1)}(\sqrt{\cdot}, \hat{\theta}_{(1)}^{neb}(\cdot))|$.

From statement 1 of lemma 2, T_1 is $O_p(\log^2 n/n^{1/4})$. Moreover, under the Binomial model $|E_n^{(1)}(\sqrt{\cdot}, \hat{\theta}_{(1)}^{neb}(\cdot))| \leq c_0 m^2$ and so T_3 is $O_p(n^{-1/2})$.

We will now consider the term T_2 . Define $Z_n(\cdot) = E_n^{(1)}(\sqrt{\cdot}, \hat{\theta}_{(1)}^{neb}(\cdot)) - E_n^{(1)}(\sqrt{\cdot}, \hat{\theta}_{(1)}^{neb}(\cdot))$ and note that for the Binomial model the proof of the second part of theorem 3 implies that $|Z_n(\cdot)|$ is $O_p(\log^2 n/n^{1/4})$ for any $\theta \in \mathbb{M}$. Moreover, for $(\theta, \theta^0) \in \mathbb{M}$

$$Z_n(\theta) - Z_n(\theta^0) \leq \frac{c_0}{n} \sum_{i=1}^n \hat{\theta}_{(1),i}^{neb}(\theta) - \sum_{i=1}^n \hat{\theta}_{(1),i}^{neb}(\theta^0) \leq \frac{c_1}{n} h_n^{(1)}(\theta) - \hat{h}_n^{(1)}(\theta^0).$$

Now using Lemma 5 and the fact that under the Binomial model with $m < 1$, $|\hat{h}_{n,i}^{(1)}(\cdot)| < 1$ uniformly in $\kappa \in \mathcal{K}$, we conclude

$$Z_n(\kappa) - Z_n(\kappa_0) \leq \frac{C_3}{n} |\kappa - \kappa_0|.$$

Thus $Z_n(\cdot)$ is Lipschitz continuous in $\kappa \in \mathcal{K}$ and, along with the compactness of $\mathcal{K}$, it implies that there exists a $(\kappa_1, \kappa_2) \in \mathcal{K}$ such that $\sup_{\kappa \in \mathcal{K}} Z_n(\kappa) = Z_n(\kappa_1)$ and $\inf_{\kappa \in \mathcal{K}} Z_n(\kappa) = Z_n(\kappa_2)$. Therefore, T_2 is $O_p(\log^2 n/n^{1/4})$.

Our desired result will follow if we can now show that $T_4 \rightarrow 0$ as $n \rightarrow \infty$. To do that we consider the proof of the second part of theorem 3 which shows that for any $\kappa \in \mathcal{K}$, $Z_n(\kappa) \rightarrow 0$ in probability as $n \rightarrow \infty$. Moreover, under the Binomial model with $\hat{p}_{(1),i}^{neb}(\cdot) \leq m/d_1$ and $\hat{p}_{(1),i}^{neb}(\cdot) \leq m/w_p^{(1)}(y_i)$ where $w_p^{(1)}(y_i) > 0$, we have $|Z_n(\kappa)| < 1$. Thus $\{Z_n(\kappa)\}$ is uniformly integrable and along with the fact that $Z_n(\kappa) \rightarrow 0$ in probability as $n \rightarrow \infty$, we have $E|Z_n(\kappa)| \rightarrow 0$ as $n \rightarrow \infty$. Therefore, for any $\kappa \in \mathcal{K}$ we have shown that $E \hat{\gamma}_n^{(1)}(\sqrt{\cdot}, \hat{p}_{(1)}^{neb}(\cdot)) - E \hat{\gamma}_n^{(1)}(\sqrt{\cdot}, \hat{p}_{(1)}^{neb}(\cdot)) \rightarrow 0$ as $n \rightarrow \infty$. To prove the result uniformly in κ we note that $Z_n(\cdot)$ is Lipschitz continuous in $\kappa \in \mathcal{K}$ and $\mathcal{K}$ is compact. So $T_4 \leq E \sup_{\kappa \in \mathcal{K}} |Z_n(\kappa)| = E|Z_n(\kappa^*)|$ where $\kappa^* \in \mathcal{K}$ is such that $\sup_{\kappa \in \mathcal{K}} |Z_n(\kappa)| = |Z_n(\kappa^*)|$. Thus $T_4 \rightarrow 0$ as $n \rightarrow \infty$ which completes our proof.

Proof of statement 1 (Binomial model) for the squared error loss ($k = 0$)

Under the Binomial model, $y_i \leq m$ and $a_{y_i-1}/a_{y_i} = y_i/(m - y_i + 1) \leq m$. Now,

$$\begin{aligned} \sup_{\kappa \in \mathcal{K}} \text{ARE}_n^{(0)}(\sqrt{\cdot}, \gamma) &= \hat{\gamma}_n^{(0)}(\sqrt{\cdot}, \hat{p}_{(0)}^{neb}(\cdot)) = \sup_{\kappa \in \mathcal{K}} \frac{2}{n} \sum_{i=1}^n X_i^n \hat{\gamma}_i^{neb}(\hat{p}_{(0),i}^{neb}(\cdot)) \sqrt{\hat{p}_{(0),i}^{neb}(\cdot)} \frac{a_{y_i-1}}{a_{y_i}} \\ &\leq \frac{2}{n} \sup_{\kappa \in \mathcal{K}} \sum_{i=1}^n X_i^n \hat{\gamma}_i^{neb}(\hat{p}_{(0),i}^{neb}(\cdot)) \hat{p}_{(0),i}^{neb}(\cdot) + \frac{2}{n} \sum_{i=1}^n X_i^n \hat{\gamma}_i^{neb}(\hat{p}_{(0),i}^{neb}(\cdot)) \sqrt{\hat{p}_{(0),i}^{neb}(\cdot)} \frac{a_{y_i-1}}{a_{y_i}} \\ &\quad + \frac{2m}{n} \sup_{\kappa \in \mathcal{K}} \sum_{i=1}^n X_i^n \hat{\gamma}_i^{neb}(\hat{p}_{(0),i}^{neb}(\cdot)) \hat{p}_{(0),i}^{neb}(\cdot) := T_1 + T_2 + T_3. \end{aligned} \quad (27)$$

Consider the term T_2 in equation (27) above and define

$$V_i = \hat{\gamma}_i^{neb}(\hat{p}_{(0),i}^{neb}(\cdot)) \sqrt{\hat{p}_{(0),i}^{neb}(\cdot)} \frac{a_{y_i-1}}{a_{y_i}}.$$

Note that from lemma 6 and conditional on $\sqrt{\cdot}$, $E_{Y_i|\sqrt{\cdot}} V_i = 0$. Moreover, V_i are independent and $|V_i| \leq c_0 \sqrt{y_i}$ for some constant $c_0 > 0$. The bound on $|V_i|$ follows from the fact that for any i , $\hat{p}_{(0),i}^{neb}(\cdot) = \hat{p}_{(0),i}^{neb}(\gamma_i) \leq m/w_p^{(0)}(y_i)$ where $w_p^{(0)}(y_i) > 0$. Thus, applying Hoeffding's inequality to $n^{-1} \sum_{i=1}^n V_i$ we get that T_2 is $O_p(k\sqrt{k_2}/n)$. Now, from assumption (A2) the distribution of $\sqrt{\cdot}$ has finite second moments which implies T_2 is $O_p(n^{-1/2})$.

We now consider the second term T_1 in equation (27) and define $Z_n(\kappa) = n^{-1} \sum_{i=1}^n \hat{\gamma}_i^{neb}(\hat{p}_{(0),i}^{neb}(\cdot)) \sqrt{\hat{p}_{(0),i}^{neb}(\cdot)}$. For any $\kappa \in \mathcal{K}$, we have

$$\frac{2}{n} \sum_{i=1}^n X_i^n \hat{\gamma}_i^{neb}(\hat{p}_{(0),i}^{neb}(\cdot)) \hat{p}_{(0),i}^{neb}(\cdot) \leq C_3 Z_n(\kappa) \leq \frac{C_3}{n} \sum_{i=1}^n \hat{\gamma}_i^{neb}(\hat{p}_{(0),i}^{neb}(\cdot)) \hat{p}_{(0),i}^{neb}(\cdot) \quad (28)$$

for some constant $c_3 > 0$. From assumption (A2) and the proof of theorem 5, the last term in the inequality above is $O_p(\log n/n^{1/4})$. Next for a perturbation δ of θ such that $(\theta, \delta) \in \mathcal{K} = [\delta_L, \delta_U]$, we will bound the increments $|Z_n(\theta) - Z_n(\theta_0)|$. To that effect, note that

$$Z_n(\theta) - Z_n(\theta_0) \leq \sqrt{2} \frac{\text{neb}(\theta)}{(0)} + \frac{1}{2} \sqrt{2} c_4 \sqrt{2} \frac{h_n^{(0)}(\theta) - h_n^{(0)}(\theta_0)}{2}.$$

Now from lemma 5 we know that

$$k \hat{h}_n^{(0)}(\theta) - \hat{h}_n^{(0)}(\theta_0) k_2 \leq c_1 \left| \sup_{h \in N(\hat{h}_n^{(0)}(\theta_0))} k r_{\hat{h}_n^{(0)}(\theta)}^2, \hat{M}_{0,n}(h) + o(1) k_2, \right.$$

and the supremum in the display above is finite since $|\hat{h}_{n,i}^{(0)}(\theta)| < 1$ uniformly in $\theta \in \mathcal{K}$. So, from assumption (A2) and Lemma 4

$$Z_n(\theta) - Z_n(\theta_0) \leq c_5 \frac{\log n}{\sqrt{n}} \left| \theta - \theta_0 \right|,$$

for n sufficiently large. Thus $Z_n(\cdot)$ is Lipschitz continuous in $\theta \in \mathcal{K}$ and, along with the compactness of $\mathcal{K}$, it implies that there exists a $(\theta_1, \theta_2) \in \mathcal{K}$ such that $\sup_{\theta \in \mathcal{K}} Z_n(\theta) = Z_n(\theta_1)$ and $\inf_{\theta \in \mathcal{K}} Z_n(\theta) = Z_n(\theta_2)$. Therefore, taking the supremum with respect θ in equation (28) and using the first part of theorem 5 suffice to prove that T_1 is $O_p(\log n/n^{1/4})$. Finally, the third term T_3 in equation (28) is $O_p(\log n/n^{1/4})$ which follows using similar arguments for the term T_1 . Therefore, we have the desired result that $\sup_{\theta \in \mathcal{K}} |\text{ARE}_n^{(0)}(\theta, Y)| \xrightarrow{P}^{(0)}(\sqrt{2} \frac{\text{neb}(\theta)}{(0)})$ is $O_p(\log^3 n/n^{1/4})$.

Proof of statement 2 (Binomial model) for the squared error loss ($k = 0$)

From Triangle inequality, $\sup_{\theta \in \mathcal{K}} |\text{ARE}_n^{(0)}(\theta, Y) - E \xrightarrow{P}^{(0)}(\sqrt{2} \frac{\text{neb}(\theta)}{(0)})| \leq T_1 + T_2 + T_3 + T_4$ where $T_1 = \sup_{\theta \in \mathcal{K}} |\text{ARE}_n^{(0)}(\theta, Y) - \xrightarrow{P}^{(0)}(\sqrt{2} \frac{\text{neb}(\theta)}{(0)})|$, $T_2 = \sup_{\theta \in \mathcal{K}} |\xrightarrow{P}^{(0)}(\sqrt{2} \frac{\text{neb}(\theta)}{(0)}) - \xrightarrow{P}^{(0)}(\sqrt{2} \frac{\text{neb}(\theta_0)}{(0)})|$, $T_3 = |\xrightarrow{P}^{(0)}(\sqrt{2} \frac{\text{neb}(\theta_0)}{(0)}) - E \xrightarrow{P}^{(0)}(\sqrt{2} \frac{\text{neb}(\theta_0)}{(0)})|$ and $T_4 = \sup_{\theta \in \mathcal{K}} |E \xrightarrow{P}^{(0)}(\sqrt{2} \frac{\text{neb}(\theta)}{(0)}) - E \xrightarrow{P}^{(0)}(\sqrt{2} \frac{\text{neb}(\theta_0)}{(0)})|$.

Under squared error loss, statement 1 of lemma 2 implies T_1 is $O_p(\log^3 n/n^{1/4})$. Moreover, under the Binomial model, $|\xrightarrow{P}^{(0)}(\sqrt{2} \frac{\text{neb}(\theta_0)}{(0)})| \leq c_0 n^{-1} k \sqrt{k_1}$. Thus, applying Hoeffding's inequality to T_3 and using assumption (A2) we get that T_3 is $O_p(n^{-1/2})$.

We will now consider the term T_2 . Define $Z_n(\theta) = \xrightarrow{P}^{(0)}(\sqrt{2} \frac{\text{neb}(\theta)}{(0)}) - \xrightarrow{P}^{(0)}(\sqrt{2} \frac{\text{neb}(\theta_0)}{(0)})$ and note that for the Binomial model, the proof of the second part of theorem 5 implies that $|Z_n(\theta)|$ is $O_p(\log^3 n/n^{1/4})$ for any $\theta \in \mathcal{K}$. Moreover, for $(\theta, \theta_0) \in \mathcal{K}$

$$Z_n(\theta) - Z_n(\theta_0) \leq \frac{c_0}{n} \sqrt{k_2} \frac{\text{neb}(\theta)}{(0)} - \frac{\text{neb}(\theta_0)}{(0)} \leq \frac{c_1}{n} \sqrt{k_2} \frac{\hat{h}_n^{(0)}(\theta) - \hat{h}_n^{(0)}(\theta_0)}{2}.$$

Now under the Binomial model and along with Lemmata 4 and 5, we conclude

$$Z_n(\theta) - Z_n(\theta_0) \leq c_2 \frac{\log n}{\sqrt{n}} \left| \theta - \theta_0 \right|$$

for n sufficiently large. Thus $Z_n(\cdot)$ is Lipschitz continuous in $\mathcal{K}$ and, along with the compactness of $\mathcal{K}$, it implies that there exists a $(\eta_1, \eta_2) \in \mathcal{K}$ such that $\sup_{\mathcal{K}} Z_n(\cdot) = Z_n(\eta_1)$ and $\inf_{\mathcal{K}} Z_n(\cdot) = Z_n(\eta_2)$. Therefore, T_2 is $O_p(\log^3 n/n^{1/4})$.

Our desired result will follow if we can now show that $T_4 \rightarrow 0$ as $n \rightarrow \infty$. We already know from the proof of the second part of theorem 5 that for any $\eta \in \mathcal{K}$, $Z_n(\eta) \rightarrow 0$ in probability as $n \rightarrow \infty$. Moreover, under the Binomial model with $\pi_{(0),i}^{neb}(\cdot) \leq c_3 m$ and $\hat{\pi}_{(0),i} \geq m/w_p^{(0)}(y_i)$ where $w_p^{(0)}(y_i) > 0$, we have

$$|Z_n(\eta)| \leq \frac{c_4}{n} \sum_{i=1}^n |\sqrt{y_i}| := U_n.$$

Now, from assumption (A2), $\sup_n E(U_n^{1+}) < \infty$ for some $\gamma > 0$. Thus $\{U_n\}$ is uniformly integrable. Therefore, $\{Z_n(\eta)\}$ is uniformly integrable and along with the fact that $Z_n(\eta) \rightarrow 0$ in probability as $n \rightarrow \infty$, we have $E|Z_n(\eta)| \rightarrow 0$ as $n \rightarrow \infty$. Therefore, for any $\eta \in \mathcal{K}$ we have shown that $|E \rightarrow_n^{(0)}(\sqrt{\cdot}, \pi_{(0)}^{(0)}) - E \rightarrow_n^{(0)}(\sqrt{\cdot}, \pi_{(0)}^{(0)}(\eta))| \rightarrow 0$ as $n \rightarrow \infty$. To prove the result uniformly in $\mathcal{K}$ we note that $Z_n(\cdot)$ is Lipschitz continuous in $\mathcal{K}$ and $\mathcal{K}$ is compact. So $T_4 \leq E \sup_{\mathcal{K}} |Z_n(\eta)| = E|Z_n(\eta^*)|$ where $\eta^* \in \mathcal{K}$ is such that $\sup_{\mathcal{K}} |Z_n(\eta)| = |Z_n(\eta^*)|$. Thus $T_4 \rightarrow 0$ as $n \rightarrow \infty$ which completes our proof.

B.6 Proof of Lemma 2 - Poisson model

Here we will prove the two statements of lemma 2 under the Poisson model. As in the Binomial case, the statements will be proved first under the scaled squared error loss and conditional on the event B_n which is the event that $\{\max_{1 \leq i \leq n} Y_i \leq C \log n\}$. Under assumption (A2), lemma 4 guarantees that B_n holds with high probability. Throughout the proof, we will denote $d_1 := \inf_{\mathcal{K}} \inf_{1 \leq i \leq n} (1 - \hat{h}_{n,i}^{(1)}(\cdot)) > 0$ and $d_2 := \inf_{\mathcal{K}} \inf_{1 \leq i \leq n} \hat{w}_{n,i}^{(0)}(\cdot) > 0$. Moreover, in the proof we will use $|\hat{h}_{n,i}^{(k)}(\cdot)| < c_0 \log n$ uniformly in $\mathcal{K}$ which is a consequence of lemma 4. The proof for the squared error loss will follow from similar arguments and we will highlight only the important steps.

Proof of statement 1 (Poisson model) for the scaled squared error loss ($k = 1$)

First note that under the Poisson model $a_{Y_i+1}/a_{Y_i} = 1/(Y_i + 1) \leq 1$. Now,

$$\begin{aligned} \sup_{\mathcal{K}} \text{ARE}_n^{(1)}(\cdot, Y) &\rightarrow_n^{(1)}(\sqrt{\cdot}, \pi_{(1)}^{neb}(\cdot)) = \sup_{\mathcal{K}} \frac{1}{n} \sum_{i=1}^n \frac{[\pi_{(1),j_i}^{neb}(\cdot)]^2}{a_{Y_i}} \frac{[\hat{\pi}_{(1),i}^{neb}(\cdot)]^2}{a_{Y_i+1}} \\ &\leq \sup_{\mathcal{K}} \frac{1}{n} \sum_{i=1}^n \frac{[\pi_{(1),j_i}^{neb}(\cdot)]^2}{a_{Y_i}} \frac{[\hat{\pi}_{(1),i}^{neb}(\cdot)]^2}{a_{Y_i}} + \sup_{\mathcal{K}} \frac{c_1}{n} \sum_{i=1}^n \frac{[\hat{\pi}_{(1),i}^{neb}(\cdot)]^2}{a_{Y_i}} \frac{[\pi_{(1),i}^{neb}(\cdot)]^2}{a_{Y_i}} \\ &\quad + \frac{1}{n} \sum_{i=1}^n \frac{[\hat{\pi}_{(1),i}^{neb}(\cdot)]^2}{a_{Y_i}} \frac{[\pi_{(1),i}^{neb}(\cdot)]^2}{a_{Y_i}} \doteq T_1 + T_2 + T_3. \end{aligned} \quad (29)$$

Here we have used the fact that $a_{y_i+1}/a_{y_i} \geq 1$ and since $\sqrt{y_i} > 0$, $1/\sqrt{y_i} < c_1$ for some positive constant c_1 . Consider the term T_3 in equation (29) above and define

$$V_i = \left[\hat{\theta}_{(1),j_i} \right]^2 \frac{a_{y_i+1}}{a_{y_i}} - \frac{\left[\hat{\theta}_{(1),i} \right]^2}{y_i}.$$

Note that from lemma 6, $EV_i = 0$. Moreover, V_i are independent and $|V_i| \leq c_2 \log^2 n$ for some constant $c_2 > 0$. The bound on $|V_i|$ follows from the fact that for any i and conditional on B_n , $\hat{\theta}_{(1),i} = \hat{\theta}_{(1),i}(y_i) \leq c_0 \log n / w_p^{(1)}(y_i)$ where $w_p^{(1)}(y_i) > 0$. Thus, applying Hoeffding's inequality to $n^{-1} \sum_{i=1}^n V_i$ we get that T_3 is $O_p(\log^2 n / \sqrt{n})$.

We now consider the second term T_2 in equation (29) and define $Z_n(\cdot) = n^{-1} \sum_{i=1}^n \{ \hat{\theta}_{(1),i}^{neb}(\cdot) \}$. Note that under the Poisson model $\hat{\theta}_{(1),i}^{neb}(\cdot) \leq Y_i/d_1$ and so for any $\gamma \in \mathcal{K}$

$$\frac{c_1}{n} \sum_{i=1}^n \left[\hat{\theta}_{(1),i} \right]^2 - \left[\hat{\theta}_{(1),i}^{neb}(\gamma) \right]^2 \leq c_3 \log n Z_n(\gamma) \leq \frac{c_3 \log n}{n} \hat{\theta}_{(1),1}^{neb}(\gamma) \quad (30)$$

for some constant $c_3 > 0$. In equation (30), we have used the fact that under assumption (A2) and lemma 4, $Y_i \leq c_0 \log n$ with high probability. Moreover, the last term in the inequality in (30) is $O_p(\log n / n^{1/4})$ since $k_{(1)}^{neb}(\gamma) \leq n^{1/2} k_{(1)}^{neb}(\gamma)$ and $n^{-1/2} k_{(1)}^{neb}(\gamma) \leq O_p(\log^2 n / n^{1/4})$ from the first part of theorem 3. Next for a perturbation of γ such that $(\gamma, \gamma') \in \mathcal{K}$, we will bound the increments $|Z_n(\gamma) - Z_n(\gamma')|$. To that effect, note that conditional on B_n

$$n Z_n(\gamma) - Z_n(\gamma') \leq \sum_{i=1}^n \left[\hat{\theta}_{(1),i}^{neb}(\gamma) - \hat{\theta}_{(1),i}^{neb}(\gamma') \right] \leq \frac{\log n}{d_1^2} \left[\hat{h}_n^{(1)}(\gamma) - \hat{h}_n^{(1)}(\gamma') \right].$$

Now from lemma 5 we know that

$$k \hat{h}_n^{(1)}(\gamma) - \hat{h}_n^{(1)}(\gamma') \leq n^{1/2} c_1 \left| \sup_{h \in \mathcal{H}(\hat{h}^{(1)}(\gamma))} k r_{h,n}^2 \hat{M}_{\gamma,n}(h) + o(1) k_2 \right|.$$

However, under the Poisson model, assumption (A2) and lemma 4, $|\hat{h}_{n,i}^{(1)}(\gamma)| < c_4 \log n$ uniformly in $\gamma \in \mathcal{K}$. Thus for n sufficiently large, the supremum in the display above is much smaller than $\log n$. So,

$$Z_n(\gamma) - Z_n(\gamma') \leq c_5 \frac{\log^2 n}{\sqrt{n}} |\gamma - \gamma'|.$$

Thus $Z_n(\cdot)$ is Lipschitz continuous in $\gamma \in \mathcal{K}$ and, along with the compactness of $\mathcal{K}$, it implies that there exists a $(\gamma_1, \gamma_2) \in \mathcal{K}$ such that $\sup_{\gamma \in \mathcal{K}} Z_n(\gamma) = Z_n(\gamma_1)$ and $\inf_{\gamma \in \mathcal{K}} Z_n(\gamma) = Z_n(\gamma_2)$. Therefore, taking the supremum with respect to γ in equation (30) and using the first part of theorem 3 suffice to prove that T_2 is $O_p(\log n / n^{1/4})$. Finally, the first term T_1 in equation (30) is $O_p(\log n / n^{1/4})$ which follows using similar arguments for the term T_2 . Therefore, we have the desired result that $\sup_{\gamma \in \mathcal{K}} |\text{ARE}^{(1)}(\gamma, \gamma')| \xrightarrow{p} 0$ is $O_p(\log^2 n / n^{1/4})$.

Proof of statement 2 (Poisson model) for the scaled squared error loss ($k = 1$)

From Triangle inequality, $\sup_{2 \leftarrow} |\text{ARE}_n^{(1)}(\cdot, Y) - E \rightarrow_n^{(1)}(\sqrt{\cdot}, \hat{\cdot}_{(1)}^{neb})| \leq T_1 + T_2 + T_3 + T_4$ where $T_1 = \sup_{2 \leftarrow} |\text{ARE}_n^{(1)}(\cdot, Y) - \rightarrow_n^{(1)}(\sqrt{\cdot}, \hat{\cdot}_{(1)}^{neb})|$, $T_2 = \sup_{2 \leftarrow} |\rightarrow_n^{(1)}(\sqrt{\cdot}, \hat{\cdot}_{(1)}^{neb}) - E \rightarrow_n^{(1)}(\sqrt{\cdot}, \hat{\cdot}_{(1)}^{neb})|$, $T_3 = \sup_{2 \leftarrow} |E \rightarrow_n^{(1)}(\sqrt{\cdot}, \hat{\cdot}_{(1)}^{neb}) - E \rightarrow_n^{(1)}(\sqrt{\cdot}, \hat{\cdot}_{(1)}^{neb})|$ and $T_4 = \sup_{2 \leftarrow} |E \rightarrow_n^{(1)}(\sqrt{\cdot}, \hat{\cdot}_{(1)}^{neb}) - E \rightarrow_n^{(1)}(\sqrt{\cdot}, \hat{\cdot}_{(1)}^{neb})|$.

From statement 1 of lemma 2, T_1 is $O_p(\log^3 n/n^{1/4})$ for the Poisson model. Moreover, under the Poisson model and conditional on B_n , $|\rightarrow_n^{(1)}(\sqrt{\cdot}, \hat{\cdot}_{(1)}^{neb})| \leq c_0 \log^2 n$. Thus, applying Hoeffding's inequality to T_3 we get that T_3 is $O_p(\log^2 n/n^{1/2})$.

We will now consider the term T_2 . Define $Z_n(\cdot) = \rightarrow_n^{(1)}(\sqrt{\cdot}, \hat{\cdot}_{(1)}^{neb}) - E \rightarrow_n^{(1)}(\sqrt{\cdot}, \hat{\cdot}_{(1)}^{neb})$ and note that from the proof of the second part of theorem 3 $|Z_n(\cdot)|$ is $O_p(\log^3 n/n^{1/4})$ for any $2 \leftarrow$. Moreover, for $(\cdot, \cdot) \in 2 \leftarrow$ and conditional on B_n

$$Z_n(\cdot) - Z_n(\cdot_0) \leq \frac{c_0 \log n}{n} \frac{neb(\cdot)}{(1)} - \frac{neb(\cdot_0)}{(1)} + \frac{c_1 \log^2 n}{n} \frac{h_n^{(1)}(\cdot)}{h_n^{(1)}(\cdot_0)}.$$

Thus, from Lemma 5, the Poisson model and assumption (A2) we have

$$Z_n(\cdot) - Z_n(\cdot_0) \leq c_3 \frac{\log^3 n}{n} + o_p(1).$$

Therefore, $Z_n(\cdot)$ is Lipschitz continuous in $2 \leftarrow$ and, along with the compactness of $2 \leftarrow$, it implies that there exists a $(\cdot_1, \cdot_2) \in 2 \leftarrow$ such that $\sup_{2 \leftarrow} Z_n(\cdot) = Z_n(\cdot_1)$ and $\inf_{2 \leftarrow} Z_n(\cdot) = Z_n(\cdot_2)$. Thus, T_2 is $O_p(\log^3 n/n^{1/4})$.

Our desired result will follow if we can now show that $T_4 \rightarrow 0$ as $n \rightarrow \infty$. To do that we consider the proof of the second part of theorem 3 which shows that for any $2 \leftarrow$, $Z_n(\cdot) \rightarrow 0$ in probability as $n \rightarrow \infty$. Moreover, under the Poisson model with $\frac{neb_{(1),i}}{(1)} \leq Y_i/d_1$ and $\hat{\cdot}_{(1),i} \leq Y_i/w_p^{(1)}(Y_i)$ where $w_p^{(1)}(Y_i) > 0$, we have for some positive constants c_0, c_1

$$|Z_n(\cdot)| \leq \frac{1}{n} \sum_{i=1}^n c_0 Y_i + c_1 Y_i^2 := U_n. \quad (31)$$

Now, the Poisson model and assumption (A2) imply that $\sup_n E(U_n^{1+}) < \infty$ for some $\gamma > 0$. Thus $\{U_n\}$ is uniformly integrable. Therefore, from equation (31), $\{Z_n(\cdot)\}$ is uniformly integrable and along with the fact that $Z_n(\cdot) \rightarrow 0$ in probability as $n \rightarrow \infty$, we have $E|Z_n(\cdot)| \rightarrow 0$ as $n \rightarrow \infty$. So, for any $2 \leftarrow$ we have shown that $|E \rightarrow_n^{(1)}(\sqrt{\cdot}, \hat{\cdot}_{(1)}^{neb}) - E \rightarrow_n^{(1)}(\sqrt{\cdot}, \hat{\cdot}_{(1)}^{neb})| \rightarrow 0$ as $n \rightarrow \infty$. To prove the result uniformly in $2 \leftarrow$ we note that $Z_n(\cdot)$ is Lipschitz continuous in $2 \leftarrow$ and $2 \leftarrow$ is compact. Therefore, $T_4 \leq E \sup_{2 \leftarrow} |Z_n(\cdot)| = E|Z_n(\cdot^*)|$ where $\cdot^* \in 2 \leftarrow$ is such that $\sup_{2 \leftarrow} |Z_n(\cdot)| = |Z_n(\cdot^*)|$. Thus $T_4 \rightarrow 0$ as $n \rightarrow \infty$ which completes our proof.

Proof of statement 1 (Poisson model) for the squared error loss ($k = 0$)

First note that under the Poisson model $a_{Y_i-1}/a_{Y_i} = Y_i$. Now,

$$\sup_{2 \leftarrow} \text{ARE}_n^{(0)}(\cdot, Y) - \rightarrow_n^{(0)}(\sqrt{\cdot}, \hat{\cdot}_{(0)}^{neb}) = \sup_{2 \leftarrow} \frac{1}{n} \sum_{i=1}^n \frac{neb_{(0),i}}{(0)} \sqrt{\cdot} - \frac{neb_{(0),i}}{(0)} \frac{a_{Y_i-1}}{a_{Y_i}} \rightarrow 0$$

$$\begin{aligned} & \frac{2}{n} \sup_{2 \leftarrow} \sum_{i=1}^n \sqrt{\hat{\gamma}_i^{neb}(\hat{\gamma}_i^{(0)})} \hat{\gamma}_i^{(0)} + \frac{2}{n} \sum_{i=1}^n \sqrt{\hat{\gamma}_i^{(0)}} \hat{\gamma}_i^{neb}(\hat{\gamma}_i^{(0)}) \frac{a_{Y_i} - 1}{a_{Y_i}} \\ & + \frac{c_1 \log n}{n} \sup_{2 \leftarrow} \sum_{i=1}^n \sqrt{\hat{\gamma}_i^{neb}(\hat{\gamma}_i^{(0)})} \hat{\gamma}_i^{(0)} := T_1 + T_2 + T_3. \end{aligned} \quad (32)$$

Here we have used the fact that we have used the fact that conditional on event B_n , $a_{Y_i} - 1/a_{Y_i} \leq Y_i \leq c_0 \log n$ for some positive constant c_0 . Consider the term T_2 in equation (32) above and define

$$V_i = \sqrt{\hat{\gamma}_i^{neb}(\hat{\gamma}_i^{(0)})} \hat{\gamma}_i^{(0)} \frac{a_{Y_i} - 1}{a_{Y_i}}$$

Note that from lemma 6 and conditional on $\sqrt{\hat{\gamma}_i}$, $E_{Y_i|\sqrt{\hat{\gamma}_i}} V_i = 0$. Moreover, V_i are independent and $|V_i| \leq c_2 \sqrt{\hat{\gamma}_i} \log^2 n$ for some constant $c_2 > 0$. The bound on $|V_i|$ follows from the fact that for any i and conditional on B_n , $\hat{\gamma}_i^{neb}(\hat{\gamma}_i^{(0)}) \leq c_3 \log n / w^{(0)}(\gamma_i)$ where $w^{(0)}(\gamma_i) > 0$. Thus, applying Hoeffding's inequality to $\sum_{i=1}^n V_i$ we get that T_2 is $O_p(k \sqrt{k_2} \log^2 n / n)$. Now, from assumption (A2) the distribution of $\sqrt{\hat{\gamma}_i}$ has finite second moments which implies T_2 is $O_p(\log^2 n / \sqrt{n})$.

We now consider the T_1 in equation (32) and define

$$Z_n(\gamma) = \frac{1}{n} \sum_{i=1}^n \sqrt{\hat{\gamma}_i^{neb}(\gamma)} \hat{\gamma}_i^{(0)}.$$

Note that for any $2 \leftarrow$,

$$\frac{2}{n} \sum_{i=1}^n \sqrt{\hat{\gamma}_i^{neb}(\gamma)} \hat{\gamma}_i^{(0)} \leq c_0 Z_n(\gamma) \leq \frac{c_0}{n} \sqrt{k_2} \log^2 n \quad (33)$$

for some constant $c_0 > 0$. From assumption (A2) and the proof of theorem 5, the last term in the inequality in (33) is $O_p(\log^2 n / n^{1/4})$. Next for a perturbation γ^0 of γ such that $(\gamma, \gamma^0) \in 2 \leftarrow$, we will bound the increments $|Z_n(\gamma) - Z_n(\gamma^0)|$. To that effect, note that conditional on event B_n

$$n |Z_n(\gamma) - Z_n(\gamma^0)| \leq \sum_{i=1}^n \sqrt{\hat{\gamma}_i^{neb}(\gamma)} \hat{\gamma}_i^{(0)} - \sum_{i=1}^n \sqrt{\hat{\gamma}_i^{neb}(\gamma^0)} \hat{\gamma}_i^{(0)} \leq c_1 \log n \sum_{i=1}^n \sqrt{\hat{\gamma}_i^{(0)}} |\hat{\gamma}_i^{(0)} - \hat{\gamma}_i^{(0)}(\gamma^0)|.$$

Now from lemma 5 we know that

$$k \hat{\gamma}_n^{(0)}(\gamma) - \hat{\gamma}_n^{(0)}(\gamma^0) k_2 \leq c_1 \sup_{h \in N(\hat{\gamma}_n^{(0)}(\gamma))} k r_{h,n}^2, \hat{M}_{\gamma,n}(h) + o(1) k_2,$$

and for n sufficiently large, the supremum in the display above is much smaller than $\log n$. So, with assumption (A2)

$$|Z_n(\gamma) - Z_n(\gamma^0)| \leq c_2 \frac{\log^2 n}{\sqrt{n}}.$$

Thus $Z_n(\gamma)$ is Lipschitz continuous in $2 \leftarrow$ and, along with the compactness of $2 \leftarrow$, it implies that there exists a $(\gamma_1, \gamma_2) \in 2 \leftarrow$ such that $\sup_{2 \leftarrow} Z_n(\gamma) = Z_n(\gamma_1)$ and $\inf_{2 \leftarrow} Z_n(\gamma) =$

$Z_n(\cdot)$. Therefore, taking the supremum with respect to γ in equation (33) and using the first part of theorem 5 suffice to prove that T_1 is $O_p(\log n/n^{1/4})$. Finally, the third term T_3 in equation (33) is $O_p(\log n/n^{1/4})$ which follows using similar arguments for the term T_1 . Therefore, we have the desired result that $\sup_{\gamma \in \mathcal{K}} |ARE_n^{(0)}(\gamma, Y_n) - E_n^{(0)}(\gamma, \eta_n^{(0)})| \leq O_p(\log^3 n/n^{1/4})$.

Proof of statement 2 (Poisson model) for the squared error loss ($k = 0$)

From Triangle inequality, $\sup_{\gamma \in \mathcal{K}} |ARE_n^{(0)}(\gamma, Y) - E_n^{(0)}(\gamma, \eta_n^{(0)})| \leq T_1 + T_2 + T_3 + T_4$ where $T_1 = \sup_{\gamma \in \mathcal{K}} |ARE_n^{(0)}(\gamma, Y) - E_n^{(0)}(\gamma, \eta_n^{(0)})|$, $T_2 = \sup_{\gamma \in \mathcal{K}} |E_n^{(0)}(\gamma, \eta_n^{(0)}) - E_n^{(0)}(\gamma, \eta_n^{(0)})|$, $T_3 = \sup_{\gamma \in \mathcal{K}} |E_n^{(0)}(\gamma, \eta_n^{(0)}) - E_n^{(0)}(\gamma, \eta_n^{(0)})|$ and $T_4 = \sup_{\gamma \in \mathcal{K}} |E_n^{(0)}(\gamma, \eta_n^{(0)}) - E_n^{(0)}(\gamma, \eta_n^{(0)})|$.

Under squared error loss, statement 1 of lemma 2 implies T_1 is $O_p(\log^3 n/n^{1/4})$. Moreover, under the Poisson model and assumption (A2), $|E_n^{(0)}(\gamma, \eta_n^{(0)})| \leq c_0 n^{-1} \log^2 n k \sqrt{k_1}$. Thus, applying Hoeffding's inequality to T_3 and using assumption (A2) we get that T_3 is $O_p(\log^2 n/\sqrt{n})$.

We will now consider the term T_2 . Define $Z_n(\gamma) = E_n^{(0)}(\gamma, \eta_n^{(0)}) - E_n^{(0)}(\gamma, \eta_n^{(0)})$ and note that from the second part of theorem 5 $|Z_n(\gamma)|$ is $O_p(\log^4 n/n^{1/4})$ for any $\gamma \in \mathcal{K}$. Moreover, for $(\gamma, \eta_n^{(0)}) \in \mathcal{K}$

$$Z_n(\gamma) - Z_n(\eta_n^{(0)}) \leq \frac{1}{n} \sum_{i=1}^n \eta_n^{(0)}(\gamma) - \eta_n^{(0)}(\eta_n^{(0)}) \leq \frac{n}{2} \sqrt{k_2} + c_0 Y \sqrt{k_2}.$$

Now conditional on the event B_n , $k Y \sqrt{k_2} \leq c_1 \sqrt{n \log n}$ and assumption (A2) implies that with high probability $k \sqrt{k_2}/\sqrt{n} \leq c_2 \sqrt{\log n}$. Thus, for n sufficiently large

$$Z_n(\gamma) - Z_n(\eta_n^{(0)}) \leq \frac{c_3 \log n}{\sqrt{n}} \eta_n^{(0)}(\gamma) - \eta_n^{(0)}(\eta_n^{(0)}) \leq \frac{c_4 \log^2 n}{\sqrt{n}} \eta_n^{(0)}(\gamma) - \hat{\eta}_n^{(0)}(\eta_n^{(0)}),$$

and together with Lemma 5, we have

$$Z_n(\gamma) - Z_n(\eta_n^{(0)}) \leq c_5 \frac{\log^3 n}{\sqrt{n}} |\eta_n^{(0)}(\gamma) - \eta_n^{(0)}(\eta_n^{(0)})|.$$

Thus $Z_n(\gamma)$ is Lipschitz continuous in $\gamma \in \mathcal{K}$ and, along with the compactness of $\mathcal{K}$, it implies that there exists a $(\gamma_1, \gamma_2) \in \mathcal{K}$ such that $\sup_{\gamma \in \mathcal{K}} Z_n(\gamma) = Z_n(\gamma_1)$ and $\inf_{\gamma \in \mathcal{K}} Z_n(\gamma) = Z_n(\gamma_2)$. Therefore, T_2 is $O_p(\log^4 n/n^{1/4})$.

Our desired result will follow if we can now show that $T_4 \rightarrow 0$ as $n \rightarrow \infty$. We already know from the proof of the second part of theorem 5 that for any $\gamma \in \mathcal{K}$, $Z_n(\gamma) \rightarrow 0$ in probability as $n \rightarrow \infty$. Moreover, under the Poisson model with $\eta_n^{(0)}(\gamma) \leq (Y_i + 1)/d_2$ and $\eta_n^{(0)}(\gamma) \leq (Y_i + 1)/w_p^{(0)}(\gamma_i)$ where $w_p^{(0)}(\gamma_i) > 0$, we have

$$|Z_n(\gamma)| \leq \frac{1}{n} \sum_{i=1}^n c_0 \sqrt{Y_i} + c_1 Y_i^2 \leq U_n.$$

Now, the Poisson model and assumption (A2) imply that $\sup_n E(U_n^{1+}) < \infty$ for some $\epsilon > 0$. Thus $\{U_n\}$ is uniformly integrable. Therefore, $\{Z_n(\gamma)\}$ is uniformly integrable and

along with the fact that $Z_n(\cdot) \rightarrow 0$ in probability as $n \rightarrow \infty$, we have $E|Z_n(\cdot)| \rightarrow 0$ as $n \rightarrow \infty$. Therefore, for any $\epsilon > 0$ we have shown that $|E \int \psi_n^{(0)}(\sqrt{n}(\theta - \theta_0)) dP_n^{(0)} - E \int \psi_n^{(0)}(\sqrt{n}(\theta - \theta_0)) dP| \rightarrow 0$ as $n \rightarrow \infty$. To prove the result uniformly in θ we note that $Z_n(\cdot)$ is Lipschitz continuous in θ and Θ is compact. So $T_4 \leq E \sup_{\theta \in \Theta} |Z_n(\theta)| = E|Z_n(\theta^*)|$ where $\theta^* \in \Theta$ is such that $\sup_{\theta \in \Theta} |Z_n(\theta)| = |Z_n(\theta^*)|$. Thus $T_4 \rightarrow 0$ as $n \rightarrow \infty$ which completes our proof.

B.7 Proof of Lemma 3

The statement of this lemma follows from part (1) of Lemma 2. First note that by definition $ARE_n^{(k)}(\hat{\theta}, Y) \leq ARE_n^{(k)}(\theta^{rc}, Y)$. So for any $\epsilon > 0$ and $k \in \{0, 1\}$, the probability $P(L_n^{(k)}(\sqrt{n}(\hat{\theta} - \theta^{rc})) \geq L_n^{(k)}(\sqrt{n}(\theta^{rc} - \theta^*)) + c_n^{-1/\epsilon})$ is bounded above by

$$P(L_n^{(k)}(\sqrt{n}(\hat{\theta} - \theta^{rc})) \geq ARE_n^{(k)}(\hat{\theta}, Y) - L_n^{(k)}(\sqrt{n}(\theta^{rc} - \theta^*)) + ARE_n^{(k)}(\theta^{rc}, Y) + c_n^{-1/\epsilon}).$$

The above display converges to 0 by part (1) of Lemma 2.

B.8 Proofs of Lemmata 4, 5 and 6

Proof of Lemma 4

First note that from assumption (A2) and for some $\epsilon > 0$, $\sqrt{n}(\hat{\theta} - \theta) = O_p(n^{-(1+\epsilon)})$ with high probability. We will now prove the statement of lemma 4 for the case when $Y_i | \sqrt{n}(\hat{\theta} - \theta) \stackrel{i.i.d.}{\sim} \text{Poi}(\sqrt{n}(\hat{\theta} - \theta))$. For distributions with bounded support, like the Binomial model, the lemma follows trivially.

Under the Poisson model, we have $P(Y_i \geq \sqrt{n}(\hat{\theta} - \theta) + t) \leq \exp\{-0.5t^2/(\sqrt{n}(\hat{\theta} - \theta) + t)\}$ for any $t > 0$. The above inequality follows from an application of Bennett inequality to the Poisson MGF (see Pollard (2015)). Now consider $P(\max_{i=1, \dots, n} Y_i \geq \sqrt{n}(\hat{\theta} - \theta) + t)$ and note that since Y_i are all independent, this probability is given by $\prod_{i=1}^n [1 - \exp\{-0.5t^2/(\sqrt{n}(\hat{\theta} - \theta) + t)\}]$. Take $t = s \log n$ where $s^2/(s + \sqrt{n}(\hat{\theta} - \theta)) > 4$. Then with $\sqrt{n}(\hat{\theta} - \theta) = O_p(n^{-(1+\epsilon)})$ log n , the above probability is bounded above by $a_n = \{1 - n^{-(1+\epsilon)}\}^n$ for some $\epsilon > 0$. As $n \rightarrow \infty$, $a_n \rightarrow 1$ which proves the statement of the lemma.

Proof of Lemma 5

We begin with some remarks on the optimization problems (8) and (15). Note that the feasible set H_n in equation (8) (and (15)) is compact and independent of θ . Moreover, the optimization problem in definitions 1 and 3 is convex. Consequently, (i) for all $\theta \in \Theta$, the optimization takes place in a compact set, and (ii) the optimal solution set corresponding to any $\theta \in \Theta$ is a singleton, $\{h_n^{(k)}(\theta)\}$. Now fix an $\epsilon > 0$. Then for any $\theta \in N_{\epsilon}(\theta_0) \setminus \Theta$ there exists a $\delta > 0$ such that the optimal solution $h_n = h_n^{(k)}(\theta) \in N_{\delta}(\theta_0)$ and $M_{\hat{n}}\{h_n\} = M_{\hat{n}}\{h_n^{(k)}(\theta_0)\} \geq 0$. Moreover, we can re-write $M_{\theta, \hat{n}}\{h_n\} = M_{\theta, \hat{n}}\{h_n^{(k)}(\theta_0)\}$ as

$$\hat{M}_{\theta, n}\{h_n\} = \hat{M}_{\theta, n}\{h_n\} = \hat{M}_{\theta, n}\{\hat{h}_n^{(k)}(\theta_0)\} + \hat{M}_{\theta, n}\{\hat{h}_n^{(k)}(\theta_0)\} + \hat{M}_{\theta, n}\{h_n\} - \hat{M}_{\theta, n}\{\hat{h}_n^{(k)}(\theta_0)\}$$

The last term in the display above is negative and thus we can upper bound $\hat{M}_{o,n}\{h_n\} - \hat{M}_{o,n}\{\hat{h}_n^{(k)}(0)\}$ by

$$\hat{M}_{o,n}\{h_n\} - \hat{M}_{o,n}\{\hat{h}_n^{(k)}(0)\} = \hat{M}_{o,n}\{\hat{h}_n^{(k)}(0)\} + \hat{M}_{o,n}\{\hat{h}_n^{(k)}(0)\} - \hat{M}_{o,n}\{\hat{h}_n^{(k)}(0)\}$$

Now apply the mean value theorem with respect to h_n to the function $\hat{M}_{o,n}\{h_n\} - \hat{M}_{o,n}\{\hat{h}_n^{(k)}(0)\}$ in the display above and notice that $\hat{M}_{o,n}\{h_n\} - \hat{M}_{o,n}\{\hat{h}_n^{(k)}(0)\}$ is bounded above by

$$r_{h_n} \hat{M}_{o,n}(\bar{h}_n) - \hat{M}_{o,n}(\bar{h}_n) = o_T(h_n - \hat{h}_n^{(k)}(0))$$

where $\bar{h}_n = \hat{h}_n^{(k)}(0) + \eta\{h_n - \hat{h}_n^{(k)}(0)\}$ for some $\eta \in (0, 1)$ and $r_{h_n} \hat{M}_{o,n}(h)$ is the partial derivative of $\hat{M}_{o,n}(h)$ with respect to h . Using $r_{h_n}[\hat{M}_{o,n}(h_n) - \hat{M}_{o,n}(h_n)] = r_{h_n}^2 \hat{M}_{o,n}(h_n)(0) + o(|0|)$ we get

$$\hat{M}_{o,n}\{h_n\} - \hat{M}_{o,n}\{\hat{h}_n^{(k)}(0)\} \leq \sup_{h \in N(\hat{h}_n^{(k)}(0))} r_{h_n}^2 \hat{M}_{o,n}(h) + o(1) = o(h_n - \hat{h}_n^{(k)}(0))$$

Moreover assumption (A3) implies that

$$\hat{M}_{o,n}\{h_n\} - \hat{M}_{o,n}\{\hat{h}_n^{(k)}(0)\} \leq c h_n - \hat{h}_n^{(k)}(0)^2$$

The desired result thus follows from the above two displays with

$$L = \sup_{h \in N(\hat{h}_n^{(k)}(0))} k r_{h_n}^2 \hat{M}_{o,n}(h) + o(1) k_2 / c.$$

Proof of Lemma 6

To prove the first statement of lemma 6, note that from equation (4)

$$\hat{f}_{(1)}(y) = \frac{a_{y-1}/a_y}{w_p^{(1)}(y)}, \text{ for } y \geq 1.$$

Now let $V(y) = a_{y-1}/(a_y[w_p^{(1)}(y)]^2)$. Then,

$$E_{Y|\sqrt{h}} \hat{f}_{(1)}^2(Y) = E_{Y|\sqrt{h}} \frac{a_{y-1}^2}{a_y^2} V(Y) = \sum_{y=1}^X \frac{a_{y-1}}{\sqrt{a_y}} V(y) \frac{a_y \sqrt{h}}{g(\sqrt{h})} = \sum_{y=0}^X V(y + \frac{a_y}{g(\sqrt{h})}) = \sqrt{E_{Y|\sqrt{h}} V(Y + 1)},$$

where $V(y + 1) = a_y/(a_{y+1}[w_p^{(1)}(y + 1)]^2) = [\hat{f}_{(1)}(y + 1)]^2(a_{y+1}/a_y)$. This proves the first statement of lemma 6.

To prove the second statement, note that from equation (4)

$$\hat{f}_{(0)}(y) = \frac{a_y/a_{y+1}}{w_p^{(0)}(y)}, \text{ for } y \geq 0.$$

Let $V(y + 1) = a_y/[a_{y+1}w_p^{(0)}(y)] = \hat{f}_{(0)}(y)$. Then,

$$\sqrt{E_{Y|\sqrt{h}} V(Y + 1)} = \sum_{y=0}^X \frac{a_y}{a_{y+1}} V(y + 1) \frac{a_{y+1} \sqrt{h}}{g(\sqrt{h})} = \sum_{y=0}^X \frac{a_y}{a_{y+1}} V(y) \frac{a_{y+1}}{g(\sqrt{h})} = E_{Y|\sqrt{h}} \frac{a_y}{a_{y+1}} V(Y),$$

where $a_{-1} = 0$ and $(a_{y-1}/a_y)V(y) = (a_{y-1}/a_y) \hat{f}_{(0)}(y - 1)$. This proves the second statement of lemma 6.

References

- Anna Aizer and Joseph J Doyle Jr. Juvenile incarceration, human capital, and future crime: Evidence from randomly assigned judges. *The Quarterly Journal of Economics*, 130(2): 759–803, 2015.
- Timothy B Armstrong, Michal Kolesár, and Mikkel Plagborg-Møller. Robust empirical bayes confidence intervals. *arXiv preprint arXiv:2004.03448*, 2020.
- Roja Bandari, Sitaram Asur, and Bernardo A Huberman. The pulse of news in social media: Forecasting popularity. *ICWSM*, 12:26–33, 2012.
- Alessandro Barp, Francois-Xavier Briol, Andrew Duncan, Mark Girolami, and Lester Mackey. Minimum stein discrepancy estimators. In *Advances in Neural Information Processing Systems*, pages 12964–12976, 2019.
- J Frédéric Bonnans and Alexander Shapiro. *Perturbation analysis of optimization problems*. Springer Science & Business Media, 2013.
- Lawrence D Brown. In-season prediction of batting averages: A field test of empirical bayes and bayes methodologies. *The Annals of Applied Statistics*, pages 113–152, 2008.
- Lawrence D Brown and Eitan Greenshtein. Nonparametric empirical bayes and compound decision approaches to estimation of a high-dimensional vector of normal means. *The Annals of Statistics*, pages 1685–1704, 2009.
- Lawrence D Brown, Eitan Greenshtein, and Ya’acov Ritov. The poisson compound decision problem revisited. *Journal of the American Statistical Association*, 108(502):741–749, 2013.
- Kacper Chwialkowski, Heiko Strathmann, and Arthur Gretton. A kernel test of goodness of fit. *JMLR: Workshop and Conference Proceedings*, 2016.
- M Lawrence Clevenson and James V Zidek. Simultaneous estimation of the means of independent poisson laws. *Journal of the American Statistical Association*, 70(351a): 698–705, 1975.
- Anna Piil Damm and Christian Dustmann. Does growing up in a high crime neighborhood affect youth criminal behavior? *American Economic Review*, 104(6):1806–32, 2014.
- Bradley Efron. Tweedie’s formula and selection bias. *Journal of the American Statistical Association*, 106(496):1602–1614, 2011.
- Bradley Efron. *Large-scale inference: empirical Bayes methods for estimation, testing, and prediction*, volume 1. Cambridge University Press, 2012.
- Bradley Efron. Two modeling strategies for empirical bayes estimation. *Statistical science: a review journal of the Institute of Mathematical Statistics*, 29(2):285, 2014.
- Bradley Efron. Empirical bayes deconvolution estimates. *Biometrika*, 103(1):1–20, 2016.

- Bradley Efron. Bayes, oracle bayes and empirical bayes. *Statistical Science*, 34(2):177–201, 2019.
- Dominique Fourdrinier and Christian P Robert. Intrinsic losses for empirical bayes estimation: A note on normal and poisson cases. *Statistics & probability letters*, 23(1):35–44, 1995.
- Dominique Fourdrinier, William E Strawderman, and Martin T Wells. *Shrinkage Estimation*. Springer, 2018.
- Anqi Fu, Balasubramanian Narasimhan, and Stephen Boyd. Cvxr: An r package for disciplined convex optimization. *arXiv preprint arXiv:1711.07582*, 2017.
- Jiaying Gu and Roger Koenker. Empirical bayesball remixed: Empirical bayes methods for longitudinal data. *Journal of Applied Econometrics*, 32(3):575–599, 2017.
- Nikolaos Ignatiadis and Stefan Wager. Bias-aware confidence intervals for empirical bayes analysis. *arXiv preprint arXiv:1902.02774*, 2019.
- William James and Charles Stein. Estimation with quadratic loss. In *Proceedings of the fourth Berkeley symposium on mathematical statistics and probability*, volume 1, pages 361–379, 1961.
- Wenhua Jiang and Cun-Hui Zhang. General maximum likelihood empirical bayes estimation of normal means. *The Annals of Statistics*, 37(4):1647–1684, 2009.
- Jack Kiefer and Jacob Wolfowitz. Consistency of the maximum likelihood estimator in the presence of infinitely many incidental parameters. *The Annals of Mathematical Statistics*, pages 887–906, 1956.
- Achim Klenke. *Probability Theory: A Comprehensive Course*, pages 331–349. Springer London, 2014.
- Roger Koenker and Gilbert Bassett Jr. Regression quantiles. *Econometrica: journal of the Econometric Society*, pages 33–50, 1978.
- Roger Koenker and Jiaying Gu. Rebayes: An r package for empirical bayes mixture methods. *Journal of Statistical Software*, 82(1):1–26, 2017.
- Roger Koenker and Ivan Mizera. Convex optimization, shape constraints, compound decisions, and empirical bayes rules. *Journal of the American Statistical Association*, 109(506):674–685, 2014.
- Susan V Koski, David Bowers, and SE Costanza. State and institutional correlates of reported victimization and consensual sexual activity in juvenile correctional facilities. *Child and Adolescent Social Work Journal*, 35(3):243–255, 2018.
- Nan Laird. Nonparametric maximum likelihood estimation of a mixing distribution. *Journal of the American Statistical Association*, 73(364):805–811, 1978.

- Qiang Liu and Dilin Wang. Stein variational gradient descent: A general purpose bayesian inference algorithm. In *Advances In Neural Information Processing Systems*, pages 2378–2386, 2016.
- Qiang Liu, Jason D Lee, and Michael I Jordan. A kernelized stein discrepancy for goodness-of-fit tests. In *Proceedings of the International Conference on Machine Learning (ICML)*, 2016.
- Nuno Moniz and Luís Torgo. Multi-source social feedback of online news feeds. *arXiv preprint arXiv:1801.07055*, 2018.
- Albert Noack. A class of random variables with discrete distributions. *The Annals of Mathematical Statistics*, 21(1):127–132, 1950.
- Chris J Oates, Mark Girolami, and Nicolas Chopin. Control functionals for monte carlo integration. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)*, 79(3):695–718, 2017.
- David Pollard. A few good inequalities. available at: <http://www.stat.yale.edu/~pollard/Books/Mini/Basic.pdf>, 2015.
- Herbert Robbins. An empirical bayes approach to statistics. Technical report, COLUMBIA UNIVERSITY New York City United States, 1956.
- Ken-iti Sato and Sato Ken-Iti. *Lévy processes and infinitely divisible distributions*. Cambridge university press, 1999.
- Robert J Serfling. *Approximation theorems of mathematical statistics*, volume 162. John Wiley & Sons, 2009.
- Galit Shmueli, Thomas P Minka, Joseph B Kadane, Sharad Borle, and Peter Boatwright. A useful distribution for fitting discrete data: revival of the conway–maxwell–poisson distribution. *Journal of the Royal Statistical Society: Series C (Applied Statistics)*, 54(1):127–142, 2005.
- US Department of Justice and Federal Bureau of Investigation. *Uniform Crime Reporting Program Data: County-Level Detailed Arrest and Offense Data, United States, 2012*. Inter-University Consortium for Political and Social Research Ann Arbor, MI, 2014. doi: <https://doi.org/10.3886/ICPSR35019.v1>.
- Hal R Varian. A bayesian approach to real estate assessment. *Studies in Bayesian econometric and statistics in Honor of Leonard J. Savage*, pages 195–208, 1975.
- Asaf Weinstein, Zhuang Ma, Lawrence D Brown, and Cun-Hui Zhang. Group-linear empirical bayes estimates for a heteroscedastic normal mean. *Journal of the American Statistical Association*, pages 1–13, 2018.
- Xianchao Xie, SC Kou, and Lawrence D Brown. Sure estimates for a heteroscedastic hierarchical model. *Journal of the American Statistical Association*, 107(500):1465–1479, 2012.

Jiasen Yang, Qiang Liu, Vinayak Rao, and Jennifer Neville. Goodness-of-fit testing for discrete distributions via stein discrepancy. In International Conference on Machine Learning, pages 5561–5570, 2018.