

LAWS: A Locally Adaptive Weighting and Screening Approach To Spatial Multiple Testing

T. Tony Cai¹, Wenguang Sun², and Yin Xia³

Abstract

Exploiting spatial patterns in large-scale multiple testing promises to improve both power and interpretability of false discovery rate (FDR) analyses. This article develops a new class of locally adaptive weighting and screening (LAWS) rules that directly incorporates useful local patterns into inference. The idea involves constructing robust and structure-adaptive weights according to the estimated local sparsity levels. LAWS provides a unified framework for a broad range of spatial problems and is fully data-driven. It is shown that LAWS controls the FDR asymptotically under mild conditions on dependence. The finite sample performance is investigated using simulated data, which demonstrates that LAWS controls the FDR and outperforms existing methods in power. The efficiency gain is substantial in many settings. We further illustrate the merits of LAWS through applications to the analysis of 2D and 3D images.

Keywords: adjusted p-value; covariate-assisted inference; dependent tests; false discovery rate; structured multiple testing

¹Department of Statistics, The Wharton School, University of Pennsylvania, Philadelphia, PA 19104. The research of Tony Cai was supported in part by NSF Grants DMS-1712735 and DMS-2015259 and NIH grants R01-GM129781 and R01-GM123056.

²Department of Data Sciences and Operations, University of Southern California. The research of Wenguang Sun was supported in part by NSF grants DMS-CAREER 1255406 and DMS-1712983.

³Department of Statistics, School of Management, Fudan University. The research of Yin Xia was supported in part by NSFC Grants 12022103, 11771094 and 11690013.

1 Introduction

1.1 Structural information in spatial multiple testing

Spatial multiple testing arises frequently from a wide variety of applications including functional neuroimaging, environmental studies, disease mapping and astronomical surveys. Intuitively, exploiting spatial structures can help identify signals more accurately and improve the interpretability of scientific findings. There are various ways of incorporating spatial information into the inferential process: the spatial structures and covariates may be utilized to form new hypotheses, define novel error rates, prioritize key tasks and construct new test statistics. For example, in Pacico et al. (2004) and Heller et al. (2006), pre-determined spatial clusters are used to form new hypotheses with clusters as basic inference units. Benjamini and Heller (2007) suggested that aggregating the data from nearby locations can increase the signal-to-noise ratio and reduce the multiplicity. To reflect the relative importance of decision errors, Benjamini and Heller (2007), Sun et al. (2015), and Basu et al. (2018) proposed to take into account spatial covariates such as the size of a cluster when defining the error rates. Moreover, the prior knowledge on the relationship between individual locations and spatial clusters is highly informative and can be utilized to develop new hierarchical testing and selective inference procedures, which promise to improve both the power and interpretability (Yekutieli, 2008; Benjamini and Bogomolov, 2014).

1.2 Challenges of dependence in multiple testing

The localization of sparse signals from massive spatial data often involves conducting thousands and even millions of hypotheses tests. The false discovery rate (FDR; Benjamini and Hochberg, 1995) provides a powerful and practical criterion for multiplicity adjustment in large-scale testing problems. An important line of research is concerned with the impact of dependence on FDR procedures. The Benjamini-Hochberg (BH) method is shown to be valid for FDR control under a range of dependence settings (Benjamini and Yekutieli, 2001; Sarkar, 2002). In particular, Wu (2008) developed conditions under which the BH method controls the FDR for spatially correlated tests in a hidden Markov random field. Meanwhile, Efron (2007) argued that correlation may degrade statistical accuracy and should be

accounted for when conducting simultaneous inference. Optimality under dependence has been investigated in Sun and Cai (2009), which showed, in a hidden Markov model, that incorporating dependence structure into a multiple-testing procedure can greatly improve the efficiency of conventional approaches that ignore dependence. This idea has been further explored in a range of spatial settings, including the Gaussian random field models (Sun et al., 2015), Ising models (Shu et al., 2015), spatial change-point models (Cao and Wu, 2015), and graphical models (Liu et al., 2016a).

Most spatial multiple testing methods have assumed that clusters are known a priori, or the dependence structure can be well estimated from data. However, there are several practical issues. First, the spatial clusters, which are typically formed by aggregating nearby locations according to prior knowledge (e.g. Heller et al., 2006), can be mis-specified. In other works (e.g. Pacico et al., 2004; Sun et al., 2015), the clusters are obtained by inspecting the testing results from a preliminary point-wise analysis, which can be subjective and highly sensitive to the choice of threshold in the tests at individual locations. Second, contiguous spatial clusters may not serve as an appropriate proxy of reality when signals appearing more frequently in a local area but do not form adjoining regions. Hence it is desirable to develop robust and fully data-driven procedures to capture local patterns in spatial data more accurately. Third, although existing spatial FDR methods have good performances when spatial models are estimated well, the commonly used computational algorithms may not produce desired estimates if assumptions on the underlying spatial process are violated, or the model/prior is misspecified. The poor estimates may lead to less powerful and even invalid FDR procedures. Finally, estimating/modeling spatial dependence structures is very challenging in high-dimensional settings, wherein strong regularity conditions and heavy computations greatly limit the scope and applicability of related works.

1.3 The main idea in our approach

The goal of the present paper is to develop simple and robust FDR methods for spatial analysis that are capable of adaptively learning the sparse structure of the underlying spatial process without prior knowledge on clusters, parametric assumptions of the underlying model or intensive computation of posterior distributions. The main idea is to recast spatial

multiple testing in the framework of simultaneous inference with auxiliary information. Under this framework, the p-values play primary roles for assessing the significance, while the spatial locations are viewed as auxiliary variables for providing important structural information to assist inference. We propose a locally adaptive weighting and screening (LAWS) approach that consists of three steps. LAWS first estimates the local sparsity structure using a screening approach, then constructs spatially adaptive weights to reorder the p-values, and finally chooses a threshold to adjust for multiplicity. The proposed method bypasses complicated spatial modeling and directly incorporates useful structures into inference. LAWS is nonparametric and assumption-lean { it only requires that the underlying spatial process is smooth at most locations. By capturing unknown spatial patterns adaptively, LAWS tends to upweight/downweight the p-values in neighborhoods where signals are abundant/sporadic. Our numerical results show that LAWS offers dramatic improvements in power over conventional methods in many settings.

1.4 Connection to existing works and our contributions

Large-scale inference with auxiliary/side information is an important topic that has received much recent attention. There are two lines of research, where the additional information is respectively (i) extracted from the same data set using carefully constructed auxiliary sequences (Liu, 2014; Cai et al., 2019), or (ii) gleaned from secondary data sources such as prior studies and external covariates (Scott et al., 2015; Fortney et al., 2015; Ignatiadis et al., 2016; Basu et al., 2018). Our work departs from these two lines of research in that the side information corresponds to the intrinsic ordering of spatial data. The spatial ordering, which encodes useful patterns such as local clusters and smoothness of the underlying process, is different from conventional auxiliary variables that are either quantitative or qualitative. For example, in the context of inference with side information, the qualitative and quantitative auxiliary variables are often used to create groups to reflect the inhomogeneity among the hypotheses (Efron, 2008; Ferkingstad et al., 2008; Cai and Sun, 2009). The works on multiple testing with groups show that weighted p-values methods can be developed to improve the power of BH (Hu et al., 2010; Liu et al., 2016b; Barber and Ramdas, 2017; Xia et al., 2019). However, the grouping strategy is not suitable for spatial analysis because dividing a region into informative groups requires either good prior knowl-

edge or intensive computation, which becomes infeasible in many scenarios, in particular when hypotheses are located on a two or three dimensional lattice. Moreover, as pointed out by Cai et al. (2019), grouping corresponds to discretizing a continuous variable, which often leads to substantial information loss. By contrast, LAWS directly incorporates the spatial structure into the weights and eliminates the need to define groups.

FDR control via LAWS offers a unified, principled and objective way for exploiting important spatial structures. It has several advantages over recent works on multiple testing with side information such as AdaPT (Lei and Fithian, 2018), SABHA (Li and Barber, 2019) and STAR (Lei et al., 2017). First, LAWS provides a general framework that is capable of handling a broad range of spatial settings. Concretely, SABHA only develops weights for grouped structure and ordered structure along a one-dimensional direction, whereas STAR only works when signals form contiguous clusters with convexity or other shape constraints. By contrast, LAWS is applicable to two or three-dimensional settings, and makes no assumption on the contiguity or convexity of the signal process as required by STAR. Second, LAWS is motivated by the optimality theory in Cai et al. (2019) and built upon solid theoretical foundations. We prove that the oracle LAWS method uniformly dominates BH in ranking and propose data-driven methods that asymptotically emulate the oracle. We present both intuitions and numerical results to demonstrate that the weights in LAWS are in general superior to the weights in SABHA. Finally, in contrast with AdaPT, SABHA and STAR whose performances heavily depend on the quality of prior information or human interactions, LAWS is fully data-driven and provides an objective and principled approach to incorporate side information. This feature is attractive in many scenarios where investigators do not have much flexibility to control the study design or decision-making process. Finally, we develop new theories to prove that LAWS controls the FDR asymptotically under dependence. The theory only requires mild conditions that seem to be substantially weaker than existing results on spatial FDR analysis in the literature.

1.5 Organization

The article is organized as follows. Section 2 introduces the model and problem formulation. Section 3 develops structure-adaptive weights and illustrates its superiority in ranking. In Section 4, we propose the LAWS procedure for spatial multiple testing and study its

theoretical properties. Simulation is conducted in Section 5 to investigate the finite sample performance of LAWS and compare it with existing methods. The merits of LAWS are further illustrated in Section 6 through applications to analyzing 2D and 3D images. The proofs are provided in the Appendix.

2 Model and Problem Formulation

Let $S \subset \mathbb{R}^d$ denote a d -dimensional spatial domain and s a location. We focus on a setting where hypotheses are located on a finite, regular lattice $S = \{s_1, \dots, s_n\}$ and data are observed at every location $s \in S$. We consider the inll{asymptotics framework (Stein, 2012) and assume $n \rightarrow \infty$ in our theoretical analysis. The setup is suitable for analyzing, say, high-frequency linear network data and ne resolution images from satellite monitoring and neuroimaging¹.

Let (s) be a binary variable, with $(s) = 1$ and $(s) = 0$ respectively indicating the presence and absence of a signal of interest at location s . The identification of spatial signals can be formulated as a multiple testing problem:

$$H_0(s) : (s) = 0 \quad \text{versus} \quad H_1(s) : (s) = 1; \quad s \in S: \quad (2.1)$$

Let $f_T(s) : s \in S \rightarrow \mathbb{R}$ be the summary statistic at location s . The common practice in multiple testing is to first convert $T(s)$ to a p -value $p(s)$ and then choose a threshold that corrects for multiplicity. The conditional cumulative distribution functions (CDF) of the p -values are given by

$$P\{p(s) \leq t | (s)\} = 1 - (s)G_1(t | s); \quad (2.2)$$

where $t \in [0, 1]$ and $G_1(t | s)$ is the non-null p -value CDF at s . The corresponding non-null density is denoted by $g_1(t | s)$. Define the sparsity level at location s

$$(s) = P\{p(s) \leq 1\} = 1 - g_1(1 | s); \quad (2.3)$$

¹In other applications such as climate change analysis, one observes incomplete data points at irregular locations (e.g. weather monitoring stations) but needs to make inference at every point in the whole spatial domain. This setting goes beyond the scope of our work; see Sun et al. (2015) for related discussions.

Due to the existence of spatial correlations and external covariates, signals may appear more frequently in certain regions, and the magnitude of non-null effects may also fluctuate across locations. Consequently we allow $\beta(s)$ and $G_1(t|s)$ to vary across the spatial domain to capture important local patterns. A mild condition in our methodological development, characterized precisely in Section 4, is that $\beta(s)$ varies smoothly as a continuous function of s . The smoothness in the sparsity levels provides the key structural information, which can be exploited to integrate information from nearby locations and construct more efficient spatial multiple testing procedures.

We focus on point-wise analysis where testing units are individual locations. The decision at location s is represented by a binary variable $d(s)$, where $d(s) = 1$ if $H_0(s)$ is rejected and $d(s) = 0$ otherwise. The widely used false discovery rate (Benjamini and Hochberg, 1995) is defined as

$$FDR = E \frac{\sum_{s \in S} \beta(s) g(s)}{\max \left\{ \sum_{s \in S} \beta(s); 1 \right\}} : \quad (2.4)$$

The power of an FDR procedure $d(s) : s \in S$ can be evaluated using the expected number of true positives:

$$ETP(d) = E \left(\sum_{s \in S} d(s) \beta(s) \right) : \quad (2.5)$$

It is important to note that although we only consider point-wise tests, the proposed LAWS procedure provides a particularly effective tool for revealing underlying spatial clusters. Hence it may be employed in the preliminary stage of a cluster-wise inference where spatial clusters need to be specified by investigators based on point-wise testing. Moreover, in contrast with existing methods which assume known spatial clusters, LAWS provides a fully data-driven approach to incorporate local structures and does not suffer from possible mis-specifications of the underlying model.

3 Structure{Adaptive Weighting and Its Properties

This section describes a weighted p-value approach to spatial FDR analysis. The key idea is to construct weights by exploiting the local sparsity structure in a spatial domain. A multiple testing procedure involves two steps: ranking and thresholding. It can be represented

by a thresholding rule of the form $(s; t) = \text{If } T(s) \geq t$, where $T(s)$ is the test statistic to order/rank the hypotheses and t is a threshold for adjusting multiplicity. In Section 3.1, we study how to improve the ranking by exploiting the spatial pattern and constructing structure{adaptive weights to adjust the p-values. Further intuition and connections to existing work are discussed in Section 3.2. In Section 3.3, we address the threshold issue and illustrate the superiority of the proposed weighting strategy. Throughout this section we assume that the local sparsity level (s) is known. The setting with unknown sparsity structure is considered in Section 4.

3.1 Incorporating sparsity structure by adjusting the p-values

To motivate our weighting strategy, consider the following covariate{adjusted mixture model under the independence assumption

$$X(s) \stackrel{\text{iid}}{\sim} f(x|s) = (1 - (s))f_0(x|s) + (s)f_1(x|s); \quad (3.1)$$

where the covariate s encodes useful side information, $f_0(x|s)$ and $f_1(x|s)$ are the null and non-null densities, (s) is the sparsity level and $f(x|s)$ is the mixture density. Ignoring the inhomogeneity captured by the covariate s , Model (3.1) reduces to the widely used random mixture model (Efron et al., 2001; Newton et al., 2004; Sun and Cai, 2007)

$$X(s) \stackrel{\text{iid}}{\sim} f(x) = (1 - (s))f_0(x) + (s)f_1(x); \quad (3.2)$$

Define the conditional (or covariate{adjusted) local false discovery rate

$$\text{CLfdr}(x|s) = \mathbb{P}(f(s) = 0|x; sg) = \frac{f(1 - (s))f_0(x|s)}{f(x|s)}. \quad (3.3)$$

It follows from the optimality theory in Cai et al. (2019) [Section 4.1] that under Model (3.1), the CLfdr thresholding rule is optimal in the sense that it maximizes the ETP subject to the constraint on FDR.

However, CLfdr cannot handle dependent tests. Under the spatial setting, we aim to

develop weighted p-values to approximate the optimal ranking by CLfdr. Let

$$w(s) = \frac{1}{f_1(s)} \frac{f_0(xjs)}{f_1(xjs)}. \quad (3.4)$$

Then $CLfdr = \alpha/(1 + \alpha)$ is monotone in α . The inspection of (3.4) reveals that whether $\alpha(s) = 1$ should be decided based on two factors: (a) the information of the sparsity structure that reflects how frequently signals appear in the neighborhood, i.e. $\frac{1}{f_1(s)}$; (b) the information exhibited by the data itself that indicates the strength of evidence against the null, i.e. $\frac{f_0(xjs)}{f_1(xjs)}$. The term $\frac{f_0(xjs)}{f_1(xjs)}$ is extremely difficult to model and calculate, we propose to replace it by the p-value, which also captures the evidence against the null in the data. Combining the above concerns, we define the weighted p-values:

$$p^w(s) = \min\left(\frac{1}{w(s)} p(s); 1\right) = \min\left(p(s); \frac{1}{w(s)}\right); \quad (3.5)$$

where $w(s) = \frac{1}{f_1(s)}$. Similar to (3.4), the weighted p-values (3.5) combines the structural information in the neighborhood and evidence of the signal at a specific location s .

3.2 Intuitions and connections to existing weighting methods

Weighting is a widely used strategy for incorporating side information into FDR analyses (Benjamini and Hochberg, 1997; Genovese et al., 2006; Roquain and Van De Wiel, 2009; Basu et al., 2018). Unlike other methods where the side information is acquired externally through domain knowledge or prior data, our inference aims to utilize the spatial information, which encodes the intrinsic structure of the collected data. The spatial structure is effectively incorporated into inference via $w(s) = \frac{1}{f_1(s)}$. The key structural assumption, which is suitable for a wide range of applications, is that the local sparsity level (s) varies smoothly in s . Our proposed LAWS procedure employs a kernel screening method to estimate s by pooling information from points close to s . It effectively takes into account important local patterns such as spatial clusters in a data-adaptive fashion. For example, suppose there are many signals in the neighborhood of s , then LAWS tends to produce a large estimate of s , thereby upweighting the p-values in the neighborhood.

The SABHA algorithm by Li and Barber (2019) adopts a different set of weights

$w^0(s) = \frac{1}{1 + \bar{p}(s)}$. Under Model (3.2), SABHA reduces to the methods in Benjamini and Hochberg (2000); Genovese and Wasserman (2002); Storey (2002), who suggested applying BH procedure to adjusted p-values $(1 - \bar{p}(s))p(s)$. These works showed that exploiting the global sparsity structure can improve the power of BH by raising the FDR from $(1 - \bar{p}(s))$ to the nominal level. The ideas in SABHA and LAWS further illustrate that exploiting the local patterns can improve the efficiency even more dramatically; the idea is formalized in our theoretical analysis in Section 3.3. Compared to the SABHA weight $w^0(s) = \frac{1}{1 + \bar{p}(s)}$, our weight $w(s) = \frac{f_0(x|s)}{f_1(x|s)}$ can separate clustered non-null p-values more effectively; this is intuitively justified by the connection to the optimality theory (3.3) and confirmed by our simulation studies. Moreover, the motivation, interpretation and estimation of our weighted p-values are all fundamentally different from the weights in Hu et al. (2010); Xia et al. (2019), which are developed under the group setting.

Finally, we stress that $w(s)$ only captures the sparsity structure, and the amplitude and variance structures of the underlying spatial process, which is subsumed in the ratio $\frac{f_0(x|s)}{f_1(x|s)}$, has been intentionally discarded when constructing our weights. This leads to a much simpler and theoretically sound methodology. It remains an open question regarding the information loss when suppressing other structural information in the proposed weights. The heterogeneity issue and the derivation of optimal weighting functions are highly non-trivial (Peña et al., 2011; Ignatiadis et al., 2016; Habiger, 2017; Habiger et al., 2017). Note that existing methods are already very complicated for the independent tests and it would require substantial efforts to extend these methods to the spatial setting.

3.3 A theoretical analysis of ranking

This section demonstrates the benefit of weighting. Let $\psi(t) = \{f^\psi(s; t) : s \in S\}$ denote a class of testing rules where $f^\psi(s; t) = \mathbb{I}\{p^\psi(s) \leq t\}$, and $p^\psi(s) = \min_{v(s)} \frac{p(s)}{v(s)}$ with $v(s)$ being the pre-specified weight. Consider the covariate-adjusted p-value mixture model (2.2). It is shown in Proposition 2 of Appendix C that, under mild conditions, the FDR of $\psi(t)$ can be written as

$$\text{FDR}_{f^\psi(t)} = Q^\psi(t) + o(1) = \mathbb{P} \left[\frac{\sum_{s \in S} f_1(s) g v(s) t}{\sum_{s \in S} f_1(s) g v(s) t + \sum_{s \in S} G f v(s) t j s g} \right] + o(1); \quad (3.6)$$

The power of $\psi(t)$ is evaluated using the ETP

$$f^\psi(t)g = \sum_{s \in S} (s)G_1 f^\psi(s)tjsg:$$

To focus on the main idea, we derive the oracle FDR procedure under an asymptotic setting, which uses the leading term $Q^\psi(t)$ in (3.6) to approximate the actual FDR. Define the oracle threshold $t_{OR} \asymp \sup\{t : Q^\psi(t) \leq g\}$. Then the oracle procedure is

$$\psi_{OR}^\psi(t_{OR}) \triangleq \{i : f^\psi(s_i) \leq t_{OR}g : s_i \in S\}:$$

Next we demonstrate that the weighted p-values $p^w(s)$ defined in (3.5) produces better ranking than the unweighted p-values. Our basic strategy is to show that at the same FDR level, thresholding (oracle) weighted p-value always yields larger ETP than unweighted p-values. Consider two sets of weights $f^\psi(s) = 1 : s \in S_g$ and $f^\psi(s) = w(s) : s \in S_g$. The asymptotic FDR and ETP of $\psi^1(t)$ and $\psi^w(t)$ are denoted by $Q^1(t)$, $Q^w(t)$, $\psi^1(t)$ and $\psi^w(t)$. The corresponding oracle procedures are defined as $\psi_{OR}^1(t_{OR}^1)$ and $\psi_{OR}^w(t_{OR}^w)$. The next theorem shows that ψ_{OR}^w uniformly dominates ψ_{OR}^1 .

Theorem 1 Assume that $\frac{p_{s \in S}^1(s)}{f^1(s)g} \leq 1$. For each $s \in S$, if the function $t \mapsto G_1(tjs)$ is concave and the function $x \mapsto G_1(xjs)$ is convex for $\min_{s \in S} w^{-1}(s) \leq x \leq \max_{s \in S} w^{-1}(s)$, then we have

$$(a) Q^w(t_{OR}^1) \leq Q^1(t_{OR}^1); \quad (b) \psi^w(t_{OR}^w) \geq \psi^1(t_{OR}^1):$$

Condition $\frac{p_{s \in S}^1(s)}{f^1(s)g} \leq 1$ in Theorem 1 is mild. It only requires that the expected number of alternative hypotheses is smaller than or equal to the expected number of null hypotheses. The condition corresponds to the notion of sparsity that holds trivially in most practical situations. By Theorem 1, we conclude that the superiority of ψ_{OR}^w over ψ_{OR}^1 is due to the improved ranking via weighted p-values since with the same threshold t_{OR}^1 , we simultaneously have $Q^w(t_{OR}^1) \leq Q^1(t_{OR}^1)$ and $\psi^w(t_{OR}^1) \geq \psi^1(t_{OR}^1)$.

4 Spatial Multiple Testing by LAWS

This section discusses a locally adaptive weighting and screening (LAWS) approach to spatial multiple testing. To emulate the oracle procedure $\mathbf{w}_{\text{OR}}$, we need to estimate two unknown quantities: the sparsity level (s) and threshold t_{OR}^w . We first develop a nonparametric screening approach for estimating (s) in Section 4.1, then propose a data-driven procedure to approximate t_{OR}^w in Section 4.2, and finally establish the theoretical properties of LAWS in Section 4.3.

4.1 Sparsity Estimation via Screening

The direct estimation of (s) is very difficult. We instead introduce an intermediate quantity to approximate (s) :

$$(s) = 1 - \frac{\mathbb{P}(\text{fp}(s) > g)}{1}, \quad 0 < g < 1: \quad (4.1)$$

We first present some intuitions to explain why (s) provides a good approximation to (s) , then describe a screening approach to estimate (s) and finally establish the theoretical properties of the proposed estimator.

The relative bias of the approximation can be calculated as

$$\frac{(s) - (s)}{1} = \frac{G_1(j(s)) - (s)}{1}:$$

This result has two implications. First, the bias is always negative, which desirably leads to conservative FDR control as we show in Theorem 2. Second, as g becomes larger, we expect that the null p-values will become increasingly dominant in the right tail area $[g, 1)$ compared to the non-null p-values, making the ratio $\frac{1 - G_1(j(s))}{1}$ very small. Hence (s) provides a good approximation to (s) with a suitably chosen g .

We now describe two key steps in estimating (s) : smoothing and screening. In the smoothing step, we exploit the structural assumption that (s) [thus (s)] varies as a smooth function of spatial location s . In reality we only have one observation at location s . To pool information from nearby locations, we use a kernel function to assign weights to observations according to their distances to s . Specifically, for any given grid S on $S \subset \mathbb{R}^d$,

let $K : \mathbb{R}^d \rightarrow \mathbb{R}$ be a positive, bounded and symmetric kernel function satisfying

$$\int_{\mathbb{R}^d} K(t) dt = 1; \quad \int_{\mathbb{R}^d} t K(t) dt = 0; \quad \int_{\mathbb{R}^d} t^T t K(t) dt < 1 :$$

Denote by $K_h(t) = h^{-1}K(t/h)$, where h is the bandwidth. At location s , define

$$v_h(s; s^0) = \frac{K_h(s - s^0)}{K_h(0)}; \quad (4.2)$$

for all $s_0 \in S$. Under the spatial setting, $K_h(s - s^0)$ is computed as a function of the Euclidean distance $\|s - s^0\|$ and $h > 0$ is a scalar. Now consider the quantity $m_s = \sum_{s^0 \in S} v_h(s; s^0)$. We can conceptualize m_s as the "total mass" (or "total number of observations") at location s . This is a key quantity in our methodological development. Thus, the smoothing step utilizes the spatial structure to calculate m_s by borrowing strength from points close to s while placing little weight on points far apart from s .

Next we explain the screening step. Motivated by (4.1), we first apply a screening procedure with threshold α to obtain a subset $T(\alpha) = \{s \in S : p(s) > \alpha\}$. Suppose we are interested in counting how many p -values from the null are greater than α among the m_s "observations" at s . The empirical count, which assumes that the majority cases in $T(\alpha)$ come from the null, is given by

$$\sum_{s^0 \in T(\alpha)} v_h(s; s^0); \quad (4.3)$$

By contrast, the expected count can be calculated theoretically as

$$\sum_{s^0 \in S} v_h(s; s^0) \mathbb{I}(p(s^0) > \alpha); \quad (4.4)$$

Setting Equations (4.3) and (4.4) equal, we obtain the following estimate

$$\hat{\alpha}(s) = \frac{\sum_{s^0 \in T(\alpha)} v_h(s; s^0)}{\sum_{s^0 \in S} v_h(s; s^0)}; \quad (4.5)$$

Next we justify the estimator (4.5) by showing that $\hat{\alpha}(s)$ converges to $\alpha(s)$ for every $s \in S$ as $|S| \rightarrow \infty$ by appealing to the in-sample asymptotics framework (Stein, 2012), where the grid S becomes denser and denser in a fixed and finite domain $S \subset \mathbb{R}^d$. For each $s \in S$, let $\lambda_{\min}(s)$ and $\lambda_{\max}(s)$ respectively be the smallest and largest eigenvalues of the Hessian

matrix $P^{(2)}(p(s) > \gamma) \in \mathbb{R}^{d \times d}$. We introduce the following technical assumptions.

(A1) Assume that $p(s)$ has continuous first and second partial derivatives and there exists a constant $C > 0$ that $C \min(s) \leq \max(s) \leq C$ uniformly for all $s \in S$.

(A2) Assume that $\text{Var}(\sum_{s \in S}^P \text{Ifp}(s) > \gamma) \leq C^0 \sum_{s \in S}^P \text{Var}(\text{Ifp}(s) > \gamma)$ for some constant $C^0 > 1$.

Proposition 1 Under (A1) and (A2), if $h \propto |S|^{-1}$, we have, uniformly for all $s \in S$,

$$E f^{\wedge}(s) - (s)g^2 \leq 0; \text{ as } |S| \rightarrow \infty:$$

Remark 1 Assumption (A1) is a mild regularity condition on the alternative CDF $G_1(j|s)$. (A2) assumes that most of the p-values are weakly correlated and it can be further relaxed with a larger choice of the bandwidth. For example, with the common choice of $h \propto |S|^{-1/5}$, by the proof of Proposition 1, we can relax (A2) to $\text{Var}(\sum_{s \in S}^P \text{Ifp}(s) > \gamma) \leq C^0 |S|^c \sum_{s \in S}^P \text{Var}(\text{Ifp}(s) > \gamma)$ for some constant $c < 4/5$, which allows the p-values to be highly correlated.

4.2 Data-driven procedure

This section describes the proposed LAWS procedure for FDR control. Define the locally adaptive weights

$$w(s) = \frac{\hat{\lambda}(s)}{1 + \hat{\lambda}(s)}; \quad s \in S; \quad (4.6)$$

where $\hat{\lambda}(s)$ is estimated by the screening approach (4.5) [the tuning parameter γ has been suppressed in the expression]. To increase the stability of the algorithm, we take $\hat{\lambda}(s) = (1 + \gamma)$ if $\hat{\lambda}(s) > 1 + \gamma$ and take $\hat{\lambda}(s) = \gamma$ if $\hat{\lambda}(s) < \gamma$ with $\gamma = 10^{-5}$. Next we order the weighted p-values from the smallest to largest. If $\hat{\lambda}(s)$ is known and the threshold is given by t_w , then the expected number of false positives (EFP) can be calculated as

$$EFP = \sum_{s \in S} P \{ p^w(s) \leq t_w | s \} = \sum_{s \in S} P \{ f(s) \leq t_w \} = \sum_{s \in S} \hat{\lambda}(s) t_w; \quad (4.7)$$

It follows that if j hypotheses are rejected along the ranking, then we expect that $\sum_{s \in S}^P \hat{\lambda}(s) p_{(j)}^w$ rejections are likely to be false positives. It follows that $\sum_{s \in S}^P \hat{\lambda}(s) p_{(j)}^w$ provides a good

estimate of the false discovery proportion (FDP). The following step-wise algorithm selects a threshold to maximize the number of rejections subject to the FDP constraint.

Algorithm 1 The LAWS Procedure

- 1: Order the weighted p-values from the smallest to largest $p_{(1)}^{(w)}; \dots; p_{(m)}^{(w)}$ and denote corresponding null hypotheses $H_{(1)}; \dots; H_{(m)}$.
 - 2: Let $k^w = \max_{j: j \leq \frac{1}{\alpha} \sum_{s \in S} \hat{\lambda}(s) p_{(j)}^{(w)}} j$; 3: Reject $H_{(1)}; \dots; H_{(k^w)}$.
-

Consider the special case where $\hat{\lambda}(s) = \hat{\lambda}$ for all $s \in S$. Then LAWS coincides with the SABHA (Li and Barber, 2019), and both recover the methods in Benjamini and Hochberg (2000); Storey (2002); Genovese and Wasserman (2002), which are essentially equivalent to applying the BH algorithm to the adjusted p-values $(1 - \hat{\lambda})p(s)$. However, the ranking by LAWS is substantially different from SABHA when $\hat{\lambda}(s)$ are heterogeneous. Our simulations show that LAWS is more powerful than SABHA and the power gain can be substantial in many settings. Moreover, SABHA does not provide a systematic way to estimate $\hat{\lambda}(s)$. It also requires preordering or grouping of the hypotheses, which is not suitable for handling higher-dimensional spatial settings.

4.3 Theoretical properties

This section studies the theoretical properties of the LAWS procedure. Define the z-values by $z(s) = \frac{1}{\alpha} (1 - \hat{\lambda}(s) p(s))$, for $s \in S$, and let $m = |S|$. Arrange $s \in S$ in any pre-specified order $s_1, \dots, s_m$ ² and denote the corresponding z-values $Z = (z_1, \dots, z_m)^T$. We collect below several regularity conditions for the asymptotic error rates control. In spatial data analysis with a latent process $f(s) : s \in S$, the dependence among p-values may come from two possible sources: the correlations among p-values when $\hat{\lambda}(s)$ are given and the correlations among $\hat{\lambda}(s)$. Our conditions on these two types of correlations are respectively specified in (A3) and (A4).

(A3) Define $(r_{i,j})_{m \times m} = R = \text{Corr}(Z)$. Assume $\max_{1 \leq i < j \leq m} |r_{i,j}| \leq r < 1$ for some

²The actual order would not affect the methodology or theory as the weights are fully determined by the spatial structure. We only need an ordering for characterizing the dependence structure between all pairs of p-values.

constant $r > 0$. Moreover, there exists $\epsilon > 0$ such that $\max_{f_i: (s_i)=0} \sum_{j: |i-j| \leq r} f_j(s_j) = o(m)$ for some constant $0 < \epsilon < 1+r$, where $f_j(s_j) = f_j : 1 \leq j \leq m; j \in \mathbb{Z}^d$ $(\log m)^{-2} g$.

(A4) Under Model (2.3), there exists a sufficiently small constant $\epsilon > 0$, such that $\sum_{s \in S} f(s) = 0$ $g = O(m^{1+\epsilon})$ for some constant $0 < \epsilon < 1$.

(A5) Define $S = \{i : 1 \leq i \leq m; \sum_{j \in \mathbb{Z}^d} f_j(s_j) > (\log m)^{(1+\epsilon)/2}\}$; where $f_i = E(z_i)$. For some $\epsilon > 0$ and some $\epsilon > 0$, $|S| \geq (1-\epsilon)m + \epsilon(\log m)^{1/2}$, where 3.14 is a math constant.

Remark 2 Condition (A3) assumes that most of the null p-values [i.e. given that $(s) = 0$] are weakly correlated. The condition can be fulfilled by a wide class of correlation structures because (i) it still allows each p-value to be highly correlated with polynomially growing number of other p-values under the null and (ii) we do not impose any conditions on the correlation structures of the p-values under the alternative. Condition (A4) only assumes that the latent variables $f(s) : s \in S$ g are not perfectly correlated. It allows highly correlated (s) so is a rather weak condition. In the case where (s) are mutually independent, (A4) is satisfied trivially with $\epsilon = 0$. Condition (A5) is mild, as it only requires that there exist a few spatial locations with mean effects of z-values exceeding $(\log m)^{(1+\epsilon)/2}$ for some $\epsilon > 0$.

Our theoretical analysis is divided into two steps. We first consider the setup where (s) is known (Theorem 2) and then turn to the case where (s) must be estimated (Theorem 3). Define the FDP of a decision rule $\psi(t)$ by

$$FDP_{\psi}(t) = \frac{\sum_{s \in S} f(s) \psi(s; t)}{\max_{s \in S} \sum_{s \in S} \psi(s; t) g(s)}.$$

We first take $\psi(s)$ as $w(s) = \frac{f(s)}{\sum_{s \in S} f(s)}$ with known (s) . Then similar to Algorithm 1, we order the weighted p-values from the smallest to largest $p_{(1)}^w; \dots; p_{(m)}^w$, and calculate $k^w = \max_{j: \sum_{i=1}^j p_{(i)}^w \leq \alpha \sum_{i=1}^m p_{(i)}^w}$. The corresponding decision rule, denoted ψ^w , is to reject $H_0(s)$ with $p^w(s) \leq p_{(k^w)}^w$.

Let H_0 be the set of null hypotheses and H_1 be the set of alternatives. Without loss of generality, we assume that $m_0 = jH_0j \leq cm$ for some $c > 0$. (Otherwise we could simply reject all the hypotheses, and the FDR would tend to zero.) The next theorem shows that ψ^w controls both the FDP and FDR at the nominal level asymptotically under dependency.

Theorem 2 Under Conditions (A3) - (A5), we have for any $\alpha > 0$

$$\lim_{m \rightarrow \infty} \overline{\text{FDR}}(\mathbf{w}) \leq \alpha; \text{ and } \lim_{m \rightarrow \infty} \mathbb{P}(\text{FDP}(\mathbf{w}) \geq 1 - \alpha) = 1:$$

Remark 3 The decision rule $\mathbf{w}$ is defined based on the weight $w(s) = \frac{1 - \frac{(\hat{s})}{s}}{1 + \frac{(\hat{s})}{s}}$, where $\hat{s}$ is a conservative approximation of s as explained in Section 4.1. Theorem 2 shows that the use of $\hat{s}$ instead of s leads to conservative error rates control.

The next theorem establishes the theoretical properties of the data-driven LAWS procedure (Algorithm 1, with decision rule denoted by $\mathbf{w} = \mathbf{w}(\hat{\mathbf{p}}_{\mathbf{w}})$, which utilizes the estimated weights via (4.5).

Theorem 3 Under the conditions in Proposition 1 and Theorem 2, we have for any $\alpha > 0$

$$\lim_{s \rightarrow \infty} \overline{\text{FDR}}(\mathbf{w}) \leq \alpha; \text{ and } \lim_{s \rightarrow \infty} \mathbb{P}(\text{FDP}(\mathbf{w}) \geq 1 - \alpha) = 1:$$

5 Simulation

This section conducts simulation studies to compare the proposed LAWS procedure with several competing methods. The implementation details are first described in Section 5.1. Sections 5.2 and 5.3 respectively consider linear block and triangle block patterns. The applications to higher dimensional settings (2D and 3D) for identifying more complicated spatial patterns are illustrated in Section 6.

5.1 Estimating the conditional proportions

The proposed estimator (4.5) captures the sparsity structure and plays a key role in constructing the weights. This section first discusses its implementation and illustrates its effectiveness. To create the screening subset T , we choose α as the p-value threshold of the BH procedure at $\alpha = 0.9$. This ensures that the null cases are dominant in T . See Appendix B for a more detailed discussion on the bias-variance tradeoff when calibrating α . The bandwidth h is set using the "h.cv" option in the R package kedd.

Next we investigate the performance of $\hat{\mathbf{w}}$ using simulated data. We generate $m = 5,000$

hypotheses from the following normal mixture model:

$$X_{ij|i} \stackrel{\text{ind}}{\sim} \epsilon_i N(0; 1) + \gamma_i N(\mu; 1); \quad \gamma_i \sim \text{Bernoulli}(\pi_i); \quad (5.1)$$

We consider two setups under which the signals appear with elevated frequencies in the following blocks [1001; 1200]; [2001; 2200]; [3001; 3200]; [4001; 4200]. The patterns of (s) , which are piecewise constants and triangle blocks, are shown in the top and bottom rows in Figure 1 (solid red lines), respectively. We can see that the varying sparsity structure of the spatial data can be reasonably captured by the estimated $\hat{\pi}_s$ (dashed blue lines). As predicted by theory, our estimated $\hat{\pi}_s$ tend to be smaller than true π_s within the blocks where signals are observed with elevated frequencies. The underestimation of (s) leads to conservative FDR levels. This is conrmed by the simulation in the next section.

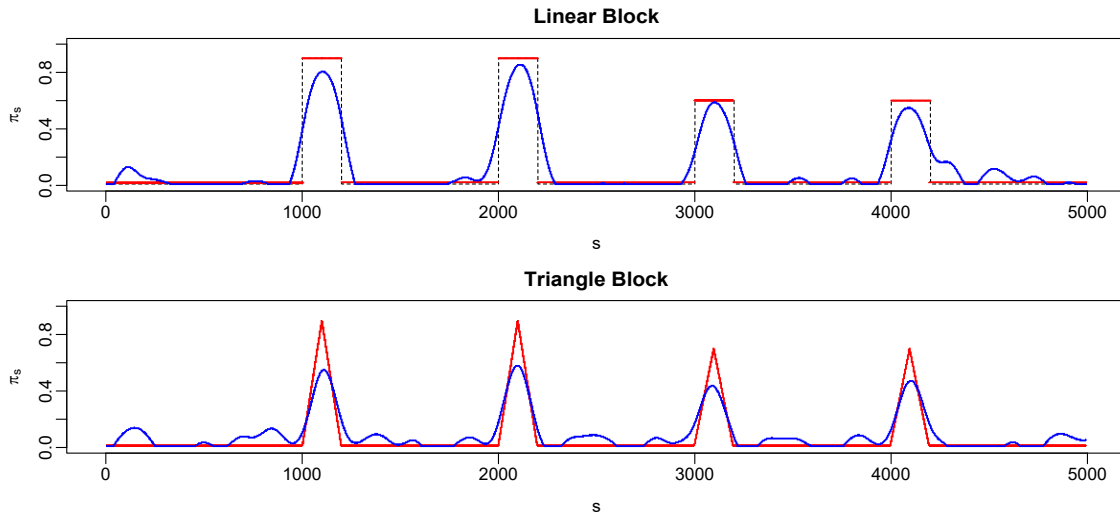


Figure 1: True π_s (solid lines) vs estimated $\hat{\pi}_s$ (dashed lines). Top row: piecewise constants; bottom row: triangle blocks.

5.2 The blockwise 1D setting with piece-wise constants

This section compares LAWS with competitive methods. Similar to the previous section, we generate data from (5.1) under the setup where (s) is a piecewise constant function (top row of Figure 1). The following methods are applied to the simulated data:

- Benjamini-Hochberg procedure (BH);

- SABHA with known (s) (SABHA.OR);
- data-driven SABHA with estimated $\hat{(s)}$ (SABHA.DD);
- LAWS with known (s) (LAWS.OR); and
- data-driven LAWS with estimated $\hat{(s)}$ (LAWS.DD).

We stress that our proposed estimator (4.5) has been used to implement SABHA. The SABHA paper does not provide an estimator of (s) with proven theoretical properties. The inclusion of SABHA is to illustrate the superiority of the LAWS weight $w(s) = (s)f_1(s)g$ over the SABHA weight $1=f_1(s)g$. The FDR and average power [denoted as $E f_{s \geq s}^P(s) \mid s = s_{2s}^P(s)g]$ of different methods are computed by averaging over 200 replications, and the nominal level is chosen at $\alpha = 0.05$. The simulation results are summarized in Figure 2.

In the top row, the signals appear with elevated frequencies in the following blocks:

$$(s) = 0.9 \text{ for } s \in [1001; 1200] \cup [2001; 2200]; (s) = 0.6 \text{ for } s \in [3001; 3200] \cup [4001; 4200];$$

For rest of the locations, we have $(s) = 0.01$. We vary ρ from 2 to 4 to investigate the impact of the signal strength. In the bottom row, we $x = 2.5$. We let $(s) = \rho$ in the above specified blocks and $(s) = 0.01$ elsewhere. Then ρ is varied from 0.3 to 0.9 to investigate the impact of sparsity structure.

We can see from Figure 2 that all methods control the FDR at the nominal level, with LAWS.DD being conservative due to the underestimation of (s) in the linear blocks (see also Figure 1). LAWS.OR substantially outperforms SABHA.OR, showing the superiority of the LAWS weight. Similarly, LAWS.DD outperforms SABHA.DD. Both LAWS and SABHA, which exploit the varying sparsity structure, outperform BH. This illustrates the benefits of incorporating side information into inference.

Finally, we can see that the efficiency gain of LAWS over competing methods is more pronounced when the signals are relatively weak (top row of Figure 2). This shows the advantage of LAWS, which integrates information from nearby locations via the weighted kernel. Moreover, the power improvement by LAWS is greater when the signals are more concentrated in the designated blocks (bottom row of Figure 2). This is consistent with our intuition since larger ρ

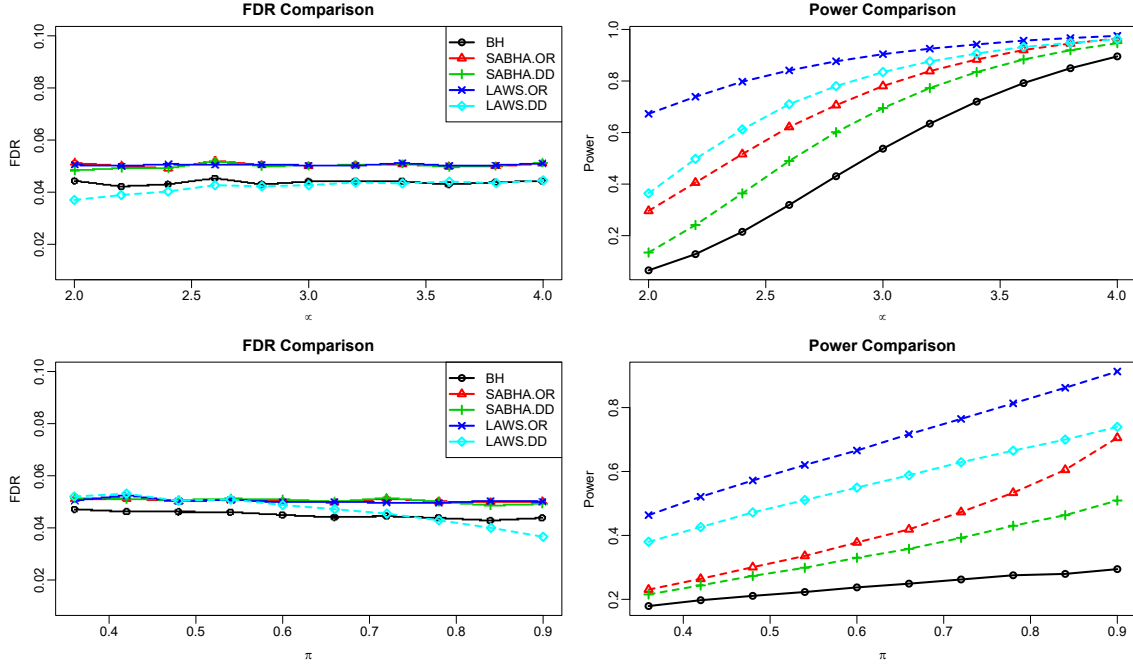


Figure 2: FDR and Power comparisons: the linear block pattern.

indicates greater disparity among spatial locations (and hence more informative spatial structure).

5.3 The blockwise 1D setting with triangular patterns

We generate data from (5.1) under the setup where (s) follows a triangular block pattern; see the bottom row of Figure 1 for an illustration. We apply BH, SABHA.OR, SABHA.DD, LAWS.OR and LAWS.DD to the simulated data and summarize the results in Figure 3. Similar as before, in the top and bottom rows we respectively vary the signal strength and sparsity levels. We can see that the power of BH is improved by SABHA, which is further improved by LAWS. The proposed method is in particular useful when the signals are weak and the structural information is strong in the spatial data.

5.4 Simulation in 2D setting

This section presents simulation results in the 2D setting. We did not compare with SABHA and STAR because (i) the original SABHA algorithm cannot be implemented since it is unclear how to order the hypotheses as a xed sequence or divide them into groups; (ii) the

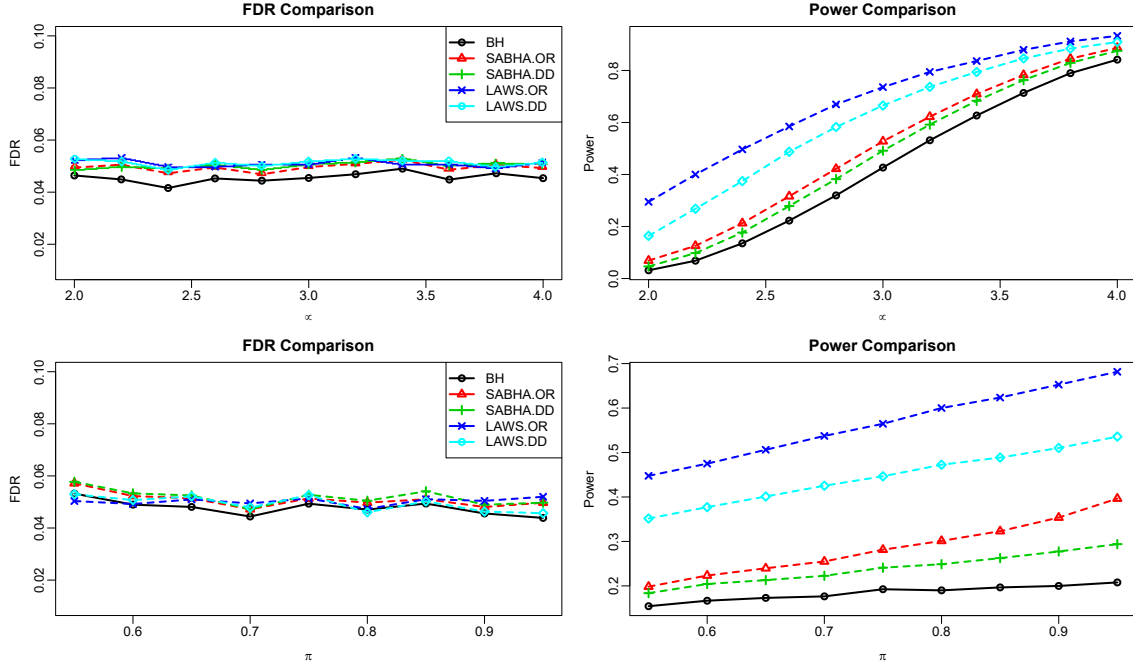


Figure 3: FDR and Power comparisons: the triangular block pattern.

STAR algorithm does not work in one of our settings where the underlying shape is not a convex region. Instead, we compare with the FDR smoothing method proposed in Tansey et al. (2018) ("Tansey" in short) for exploiting spatial structure in large multiple-testing problems.

We generate the data by Model (5.1) on a 200 200 lattice, where the signals are more likely to be located on a double-triangle or a rectangle shape as shown in Figure 4. We let $(s) = 0:9$ for the left triangle and left half of the rectangle respectively, $(s) = 0:6$ for the right triangle and right half of the rectangle and let $(s) = 0:01$ for the rest of the locations. Similarly as in the 1D setting, we first vary α from 2.5 to 4 to investigate the impact of the signal strength. We then $\alpha = 3$, let $(s) = 0$ in the triangle and rectangle patterns, $(s) = 0:01$ for the rest, and vary π from 0.6 to 0.9 to illustrate the impact of sparsity structure. The empirical FDR and power are computed over 200 replications with nominal level $\alpha = 0:05$.

We can see from Figures 5 and 6 that, all methods except Tansey controls the empirical FDR well and LAWS.DD is slightly more conservative than LAWS.OR due to the negative bias of α as explained in Section 4.1. By successfully incorporating the spatial information,

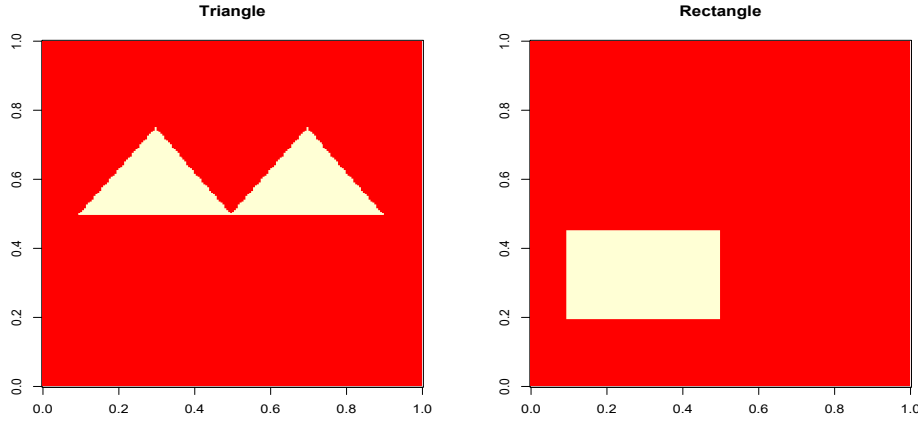


Figure 4: 2D triangle and rectangle pattern.

the empirical power of BH has been significantly improved by LAWS.OR and LAWS.DD for varying signal strengths and sparsity levels. The improvement is more significant when the signals are weaker or the sparse structure is more informative. Note that, due to the seriously inflated empirical FDRs, Tansey has higher empirical power than the competing methods.

5.5 Simulation in 3D setting

We compare in this section the numerical performance of LAWS.DD and LAWS.OR with Tansey and the BH method in 3D spatial settings. The data are generated by Model (5.1) on a 3D 20 25 30 lattice, where the signals are located on a cubic with dimension 10 10 15. We let $(s) = 0.8$ within the cubic and let $(s) = 0.01$ for the rest of the locations. To show the impact of the signal strength and the impact of sparsity structure, we vary α and (s) ($\alpha = 3.5$ for the latter) in the same way as the 2D settings. The empirical FDR and power are computed over 200 replications with nominal level $\alpha = 0.05$.

We see from Figure 7 that, similarly as 1D and 2D settings, all methods except Tansey control the FDR and LAWS.DD is slightly more conservative than LAWS.OR; the empirical power of BH is significantly improved by LAWS.OR and LAWS.DD for different signal strengths and sparsity levels.

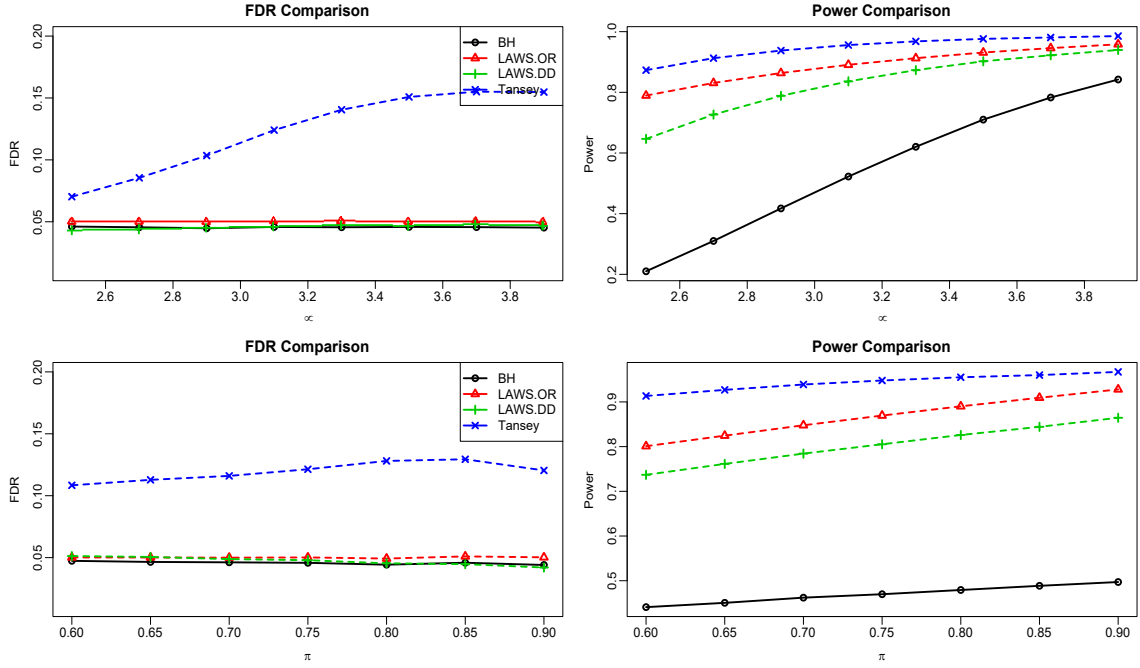


Figure 5: FDR and Power comparisons: the 2D triangle pattern.

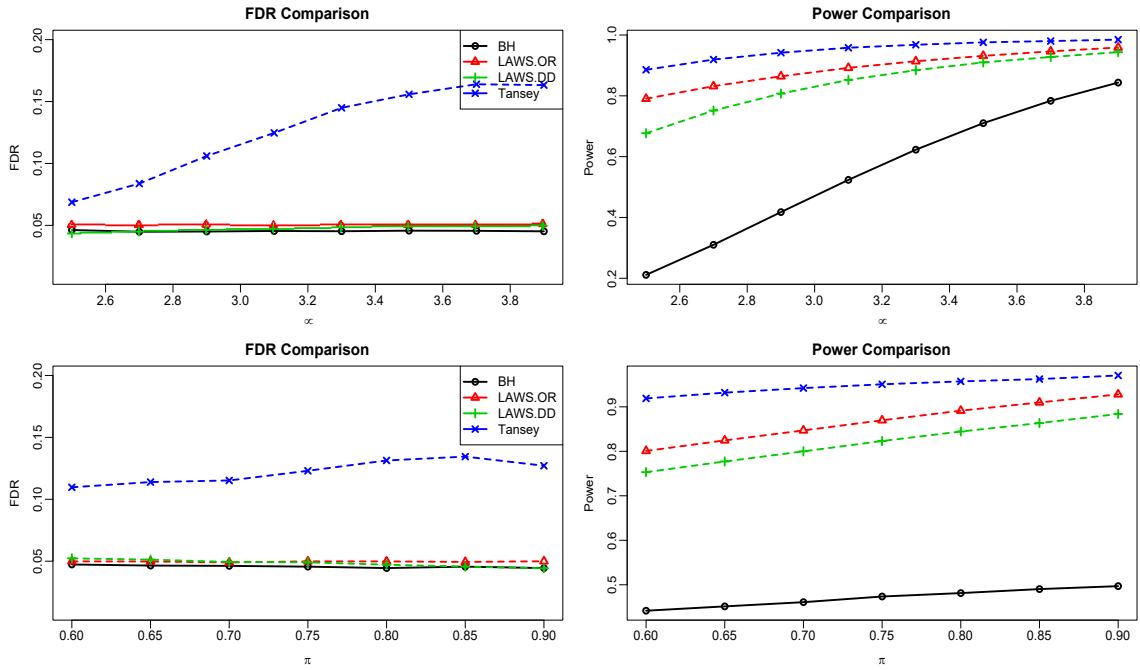


Figure 6: FDR and Power comparisons: the 2D rectangle pattern.

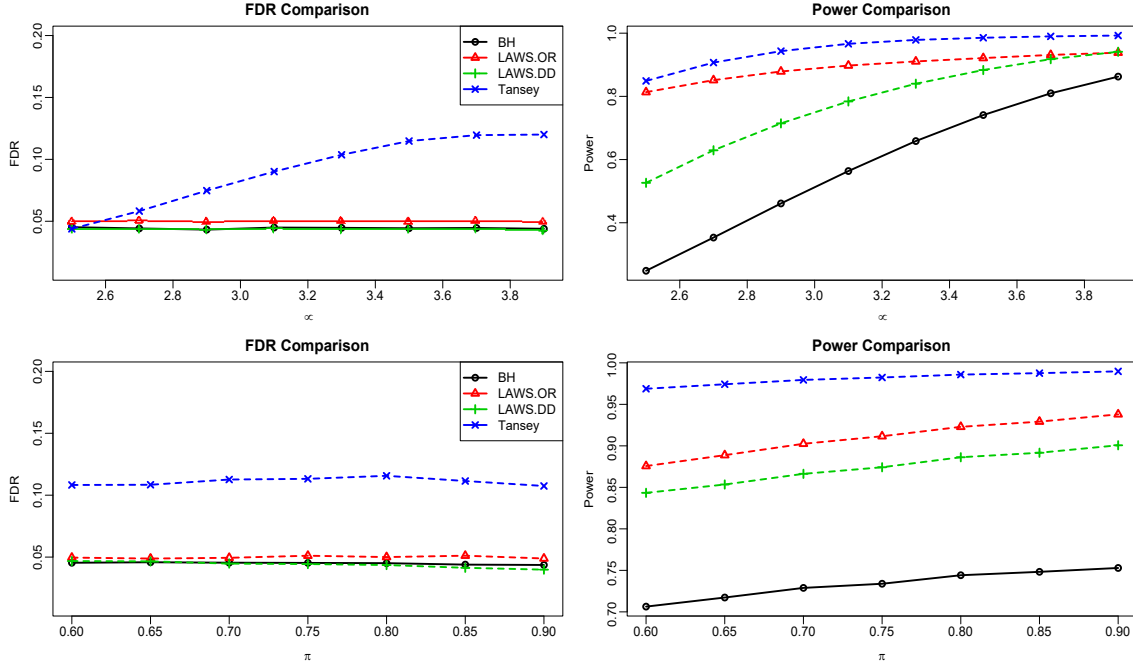


Figure 7: FDR and Power comparisons: the 3D cubic pattern.

6 Applications

This section applies LAWS to identify 2D spatial clusters (Section 6.1) and signal patterns in 3D image data (Section 6.2). LAWS has several advantages over existing structure-adaptive testing methods. For example, SABHA requires that the hypotheses can be divided into groups or should be ordered as a fixed sequence, which are not always feasible in 2D and 3D spatial applications. By contrast, LAWS constructs weights based on the distance between spatial locations (4.2) and can easily handle higher dimensional spatial settings. Unlike the STAR procedure (Lei et al., 2017) which requires that the spatial region must be contiguous and convex, LAWS is applicable to a wider types of settings where the local sparsity patterns are heterogeneous. We present two examples to show that LAWS is more accurate and effective in identifying and recovering specific patterns of interest in analysis of 2D and 3D image data.

6.1 The 2D setting with spatial clusters

We simulate data on a 200 200 lattice. The signals of interest form two spatial clusters respectively with donut and square shapes. The observations follow the random mixture

model (5.1), where $(s) = 1$ if s is within the donut or square and $(s) = 0$ otherwise. We first obtain $\hat{\lambda}(s)$ using (4.5), where $\|s - s_0^k\|$ is calculated as the usual Euclidean distance. We then obtain two-sided p-values and finally apply both BH and LAWS to the simulated data set. From the first to last row, we vary the signal strength from 2.0 to 3.0. The true states, and the rejected locations by BH and LAWS are respectively displayed from Column 1 to Column 3.

We can see that LAWS is more powerful than BH in uncovering the underlying truth. Both spatial patterns, namely the donut and square, can be more easily identified based on the results of LAWS. The key idea is that $\hat{\lambda}(s)$ tend to be very large in the neighborhood of clustered signals (yellow areas) due to the strong spatial correlations. Therefore the p-values in these neighborhood are upweighted via data-driven weights. We conclude that by exploiting the local sparsity structures, LAWS is more effective in rejecting the hypotheses in regions where signals appear in clusters. This property is in particular attractive in spatial data analysis.

6.2 The 3D setting: application to fMRI data

We further illustrate the LAWS procedure through a magnetic resonance imaging (MRI) data for a study of attention deficit hyperactivity disorder (ADHD). The dataset is available at <http://neurobureau.projects.nitrc.org/ADHD200/Data.html>. The images were produced by the ADHD-200 Sample Initiative, then preprocessed by the Neuro Bureau.

We first reduce the resolution of MRI images from $256 \times 198 \times 256$ to $30 \times 36 \times 30$ (Li and Zhang, 2017) by aggregating the corresponding pixels into blocks. This helps the analysis in several ways. First, the aggregation of pixels not only increases the signal to noise ratio but also effectively avoids misalignments of brain regions. Second, the p-values, which are calculated based on normal approximations, should satisfy the required accuracy needed in the large p small n paradigm (Fan et al., 2007; Liu and Shao, 2010; Chernozhukov et al., 2017). The downsizing helps to increase the precision of the approximations. Finally, the aggregation can effectively eliminate noises and make it easier to visualize interesting spatial patterns.

The dataset consists of 931 subjects, among whom 356 are combined ADHD subjects and 575 are normal controls. We conduct two-sample t-tests to compare the two groups

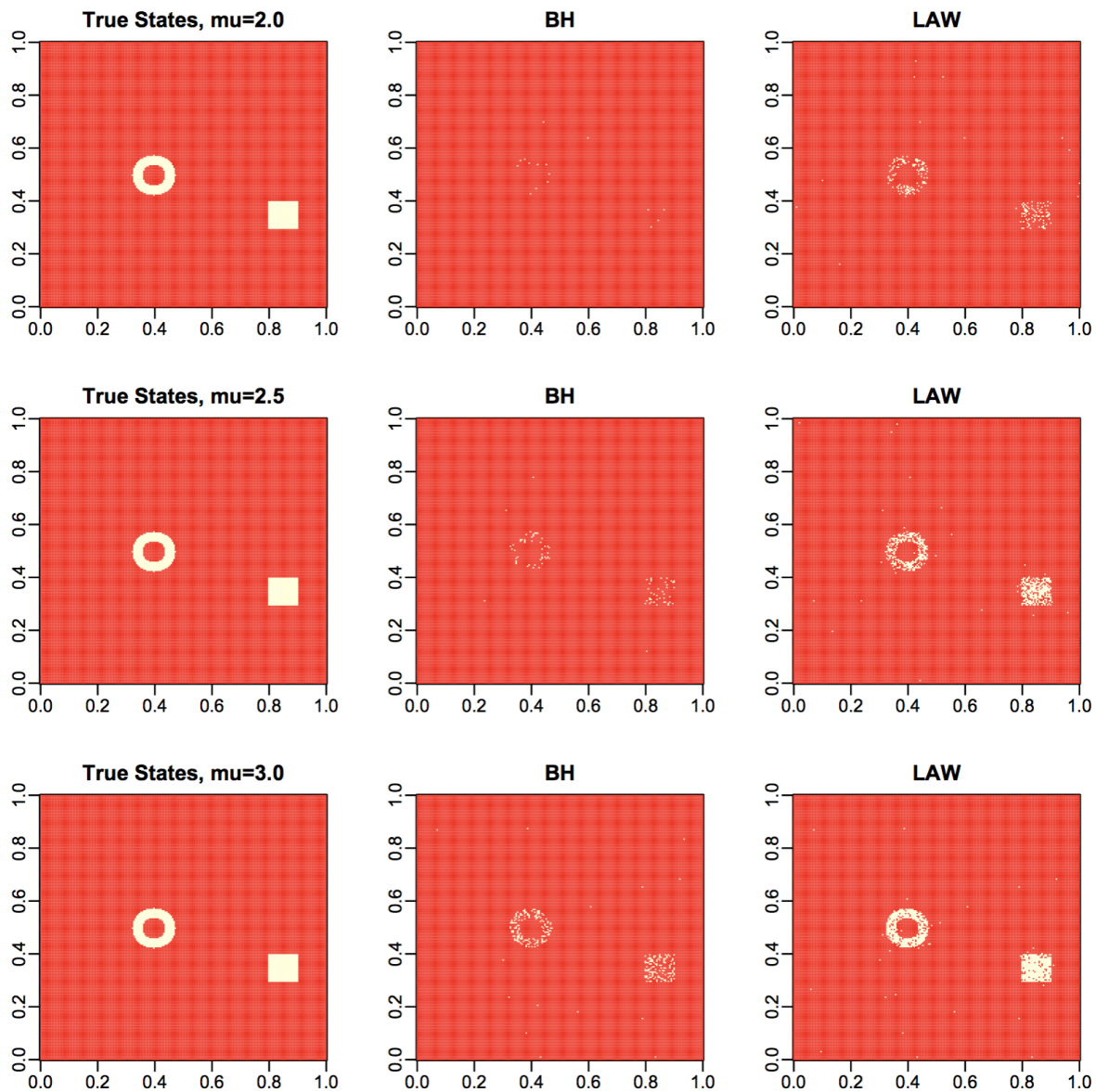


Figure 8: Spatial FDR Analysis in 2D setting. LAWS is more effective in revealing the donut and square shapes by up-weighting the p-values in the regions where signal appear in clusters.

and use normal approximation to obtain the p-values. Finally, we apply LAWS and BH procedures to identify brain regions that exhibit significant differences between subjects with and without ADHD.

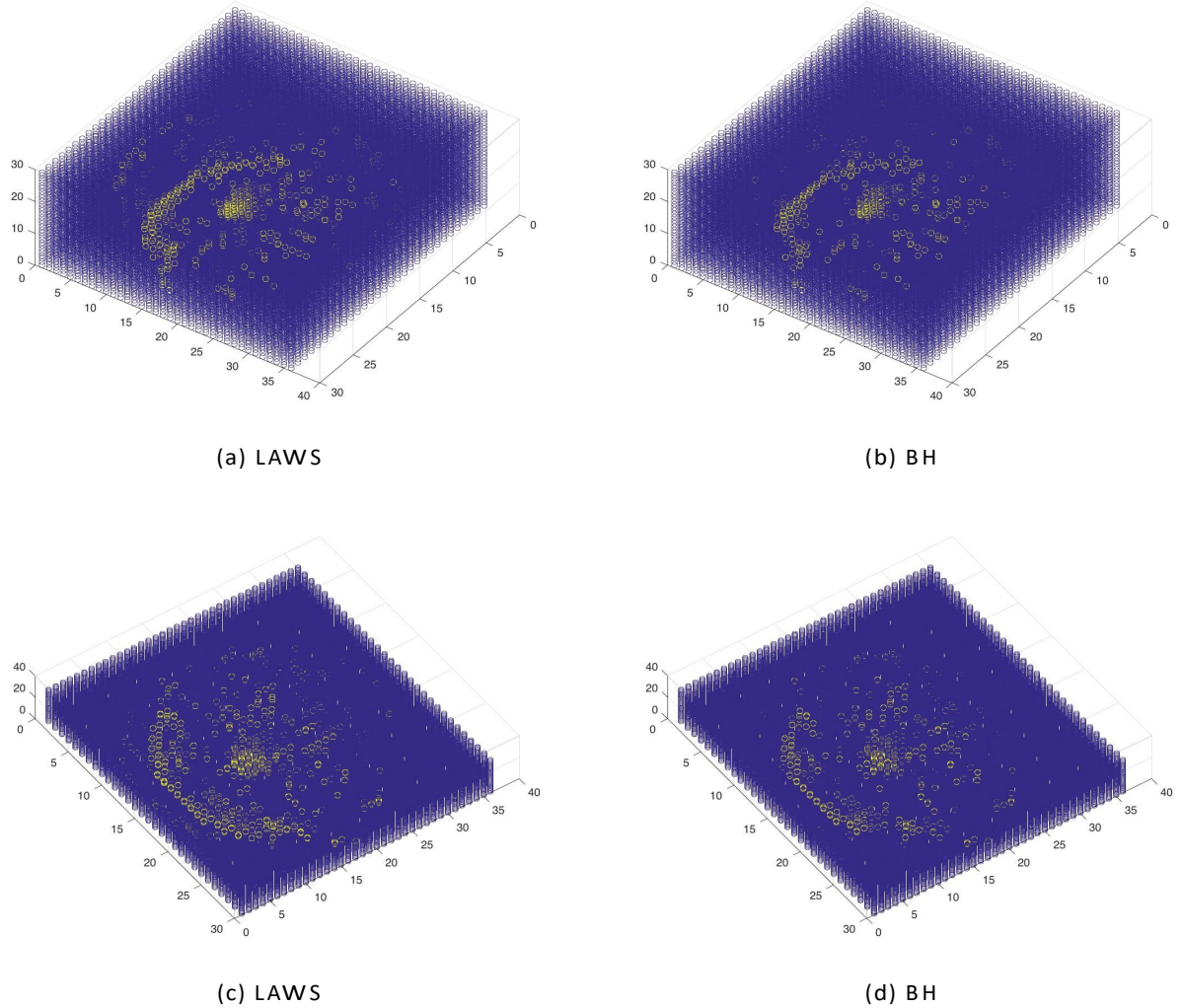


Figure 9: Significant brain regions (yellow) after applying LAWS (left) and BH (right), view with azimuth and elevation angles (35; 65) on top row and (35; 80) on bottom row. FDR level = 0.05.

Figure 9 displays the testing results from two different angles of the 3D image, with FDR level equal to 0.05. The significant brain regions identified by BH are a subset of those identified by LAWS. To be more specific, the LAWS procedure identifies 538 regions, while BH recovers 349. The graph further shows that LAWS has superior power performance

over BH in identifying spatial signals.

7 Discussions

This paper develops a new locally-adaptive weighting approach that incorporates the spatial structure into statistical inference. It provides a unified framework for a broad range of spatial multiple testing problems and is fully data-driven. We show that LAWS controls the FDR asymptotically under dependence and outperforms existing methods in power.

LAWS is powerful yet simple. It is capable of adaptively learning the sparse structure of the underlying spatial process without prior knowledge. The spatial locations are viewed as auxiliary variables for providing important structural information to assist inference. However, as explained in Section 3.1 of the paper, there are two pieces of information that could potentially be useful in spatial setting: the varying sparsity structure that we have considered, and the varying distributional information that we have replaced by individual p-values. Such replacement may lead to certain information loss which requires further investigation. The estimation of the distributional information is challenging and computationally intensive. Substantial work is needed to extend existing methods to the spatial setting; such analysis is beyond the scope of the current paper. The development of more powerful weighting strategies to incorporate other types of side information is an interesting and important direction for future research.

References

- Barber, R. F. and Ramdas, A. (2017). The p-lter: multilayer false discovery rate control for grouped hypotheses. *J. Roy. Statist. Soc. B*, 79(4):1247{1268.
- Basu, P., Cai, T. T., Das, K., and Sun, W. (2018). Weighted false discovery rate control in large-scale multiple testing. *J. Am. Statist. Assoc.*, 113(523):1172{1183.
- Benjamini, Y. and Bogomolov, M. (2014). Selective inference on multiple families of hypotheses. *J. Roy. Statist. Soc. B*, 76(1):297{318.

- Benjamini, Y. and Heller, R. (2007). False discovery rates for spatial signals. *J. Amer. Statist. Assoc.*, 102(480):1272{1281.
- Benjamini, Y. and Hochberg, Y. (1995). Controlling the false discovery rate: a practical and powerful approach to multiple testing. *J. Roy. Statist. Soc. B*, 57(1):289{300.
- Benjamini, Y. and Hochberg, Y. (1997). Multiple hypotheses testing with weights. *Scand. J. Stat.*, 24(3):407{418.
- Benjamini, Y. and Hochberg, Y. (2000). On the adaptive control of the false discovery rate in multiple testing with independent statistics. *J. Educ. Behav. Stat.*, 25(1):60{83.
- Benjamini, Y. and Yekutieli, D. (2001). The control of the false discovery rate in multiple testing under dependency. *Ann. Statist.*, 29(4):1165{1188.
- Cai, T. T. and Sun, W. (2009). Simultaneous testing of grouped hypotheses: Finding needles in multiple haystacks. *J. Amer. Statist. Assoc.*, 104(488):1467{1481.
- Cai, T. T., Sun, W., and Wang, W. (2019). CARS: Covariate assisted ranking and screening for large-scale two-sample inference (with discussion). *J. Roy. Statist. Soc. B*, 81(2):187{234.
- Cao, H. and Wu, W.-B. (2015). Changepoint estimation: another look at multiple testing problems. *Biometrika*, 102(4):974{980.
- Chernozhukov, V., Chetverikov, D., Kato, K., et al. (2017). Central limit theorems and bootstrap in high dimensions. *Ann. Probab.*, 45(4):2309{2352.
- Efron, B. (2007). Correlation and large-scale simultaneous significance testing. *J. Amer. Statist. Assoc.*, 102(477):93{103.
- Efron, B. (2008). Microarrays, empirical Bayes and the two-groups model. *Statist. Sci.*, 23:1{22.
- Efron, B., Tibshirani, R., Storey, J. D., and Tusher, V. (2001). Empirical Bayes analysis of a microarray experiment. *J. Amer. Statist. Assoc.*, 96(456):1151{1160.

- Fan, J., Hall, P., and Yao, Q. (2007). To how many simultaneous hypothesis tests can normal, student's t or bootstrap calibration be applied? *J. Amer. Statist. Assoc.*, 102(480):1282{1288.
- Ferkingstad, E., Frigessi, A., Rue, H., Thorleifsson, G., and Kong, A. (2008). Unsupervised empirical bayesian multiple testing with external covariates. *Ann. Appl. Stat.*, 2(2):714{735.
- Fortney, K., Dobriban, E., Garagnani, P., Pirazzini, C., Monti, D., Mari, D., Atzmon, G., Barzilai, N., Franceschi, C., Owen, A. B., et al. (2015). Genome-wide scan informed by age-related disease identities loci for exceptional human longevity. *PLoS genetics*, 11(12):e1005728.
- Genovese, C. and Wasserman, L. (2002). Operating characteristics and extensions of the false discovery rate procedure. *J. R. Stat. Soc. B*, 64(3):499{517.
- Genovese, C. R., Roeder, K., and Wasserman, L. (2006). False discovery control with p-value weighting. *Biometrika*, 93(3):509{524.
- Habiger, J. (2017). Adaptive false discovery rate control for heterogeneous data. *Stat. Sin.*, pages 1731{1756.
- Habiger, J., Watts, D., and Anderson, M. (2017). Multiple testing with heterogeneous multinomial distributions. *Biometrics*, 73(2):562{570.
- Heller, R., Stanley, D., Yekutieli, D., Rubin, N., and Benjamini, Y. (2006). Cluster-based analysis of fmri data. *NeuroImage*, 33(2):599{608.
- Hu, J. X., Zhao, H., and Zhou, H. H. (2010). False discovery rate control with groups. *J. Amer. Statist. Assoc.*, 105(491):1215{1227.
- Ignatiadis, N., Klaus, B., Zaugg, J. B., and Huber, W. (2016). Data-driven hypothesis weighting increases detection power in genome-scale multiple testing. *Nat. Methods*, 13(7):577.
- Lei, L. and Fithian, W. (2018). Adapt: an interactive procedure for multiple testing with side information. *J. Roy. Statist. Soc. B*, 80(4):649{679.

- Lei, L., Ramdas, A., and Fithian, W. (2017). Star: A general interactive framework for fdr control under structural constraints. *arXiv preprint arXiv:1710.02776*.
- Li, A. and Barber, R. F. (2019). Multiple testing with the structure-adaptive benjamini-hochberg algorithm. *J. Roy. Statist. Soc. B*, 81(1):45{74.
- Li, L. and Zhang, X. (2017). Parsimonious tensor response regression. *J. Amer. Statist. Assoc.*, 112(519):1131{1146.
- Liu, J., Zhang, C., Page, D., et al. (2016a). Multiple testing under dependence via graphical models. *Ann. Appl. Stat.*, 10(3):1699{1724.
- Liu, W. (2014). Incorporation of sparsity information in large-scale multiple two-sample t tests. *arXiv preprint arXiv:1410.4282*.
- Liu, W. and Shao, Q.-M. (2010). Cramer-type moderate deviation for the maximum of the periodogram with application to simultaneous tests in gene expression time series. *Ann. Statist.*, 38(3):1913{1935.
- Liu, Y., Sarkar, S. K., and Zhao, Z. (2016b). A new approach to multiple testing of grouped hypotheses. *J. Stat. Plan. Inference*, 179:1{14.
- Newton, M. A., Noueiry, A., Sarkar, D., and Ahlquist, P. (2004). Detecting differential gene expression with a semiparametric hierarchical mixture method. *Biostatistics*, 5(2):155{176.
- Pacico, M., Genovese, C., Verdinelli, I., and Wasserman, L. (2004). False discovery control for random elds. *J. Amer. Statist. Assoc.*, 99(468):1002{1014.
- Peña, E. A., Habiger, J. D., and Wu, W. (2011). Power-enhanced multiple decision functions controlling family-wise error and false discovery rates. *Ann. Statist.*, 39(1):556.
- Roquain, E. and Van De Wiel, M. A. (2009). Optimal weighting for false discovery rate control. *Electron. J. Stat.*, 3:678{711.
- Sarkar, S. K. (2002). Some results on false discovery rate in stepwise multiple testing procedures. *Ann. Statist.*, 30:239{257.

- Scott, J. G., Kelly, R. C., Smith, M. A., Zhou, P., and Kass, R. E. (2015). False discovery rate regression: an application to neural synchrony detection in primary visual cortex. *J. Amer. Statist. Assoc.*, 110(510):459{471.
- Shu, H., Nan, B., and Koeppel, R. (2015). Multiple testing for neuroimaging via hidden markov random eld. *Biometrics*, 71(3):741{750.
- Stein, M. L. (2012). *Interpolation of spatial data: some theory for kriging*. Springer Science & Business Media.
- Storey, J. D. (2002). A direct approach to false discovery rates. *J. Roy. Statist. Soc. B*, 64(3):479{498.
- Sun, W. and Cai, T. T. (2007). Oracle and adaptive compound decision rules for false discovery rate control. *J. Amer. Statist. Assoc.*, 102(479):901{912.
- Sun, W. and Cai, T. T. (2009). Large-scale multiple testing under dependence. *J. R. Stat. Soc. B*, 71(2):393{424.
- Sun, W., Reich, B. J., Cai, T. T., Guindani, M., and Schwartzman, A. (2015). False discovery control in large-scale spatial multiple testing. *J. Roy. Statist. Soc. B*, 77(1):59{83.
- Tansey, W., Koyejo, O., Poldrack, R. A., and Scott, J. G. (2018). False discovery rate smoothing. *Journal of the American Statistical Association*, 113(523):1156{1171.
- Wu, W. B. (2008). On false discovery control under dependence. *Ann. Statist.*, 36(1):364{380.
- Xia, Y., Cai, T. T., and Sun, W. (2019). GAP: A General Framework for Information Pooling in Two-Sample Sparse Inference. *J. Amer. Statist. Assoc.*, to appear.
- Yekutieli, D. (2008). Hierarchical false discovery rate{controlling methodology. *J. Amer. Statist. Assoc.*, 103(481):309{316.

Supplementary Material for “LAWS: A Locally Adaptive Weighting and Screening Approach To Spatial Multiple Testing”

We first prove the main results in Section A. Additional explanations for the variance-bias tradeoff in choosing $\hat{\tau}$ and Equation (3.6) are provided in Sections B and C, respectively.

A Proofs of Main Results

A.1 Proof of Theorem 1

Proof: Note that

$$Q_w(t) = \frac{\sum_{s \in S} \frac{1 - \hat{\tau}(s)}{t + \frac{\sum_{s \in S} (1 - \hat{\tau}(s))}{\sum_{s \in S} \hat{\tau}(s)}}}{\sum_{s \in S} (1 - \hat{\tau}(s))t + \frac{\sum_{s \in S} (1 - \hat{\tau}(s))}{\sum_{s \in S} \hat{\tau}(s)}}, \quad \frac{\sum_{s \in S} \hat{\tau}(s) G_1\{w(s)t | s\}}{\sum_{s \in S} [\hat{\tau}(s) G_1\{t | s\}]},$$

and

$$Q_1(t) = \frac{\sum_{s \in S} \frac{1 - \hat{\tau}(s)}{t + \frac{\sum_{s \in S} (1 - \hat{\tau}(s))}{\sum_{s \in S} \hat{\tau}(s)}}}{\sum_{s \in S} (1 - \hat{\tau}(s))t + \frac{\sum_{s \in S} (1 - \hat{\tau}(s))}{\sum_{s \in S} \hat{\tau}(s)}} \frac{\sum_{s \in S} \hat{\tau}(s) G_1\{t | s\}}{\sum_{s \in S} [\hat{\tau}(s) G_1\{t | s\}]}$$

Because $t \mapsto G_1(t | s)$ is concave and $x \mapsto G_1(t/x | s)$ is convex for $\min_{s \in S} w^{-1}(s) \leq x \leq \max_{s \in S} w^{-1}(s)$, together with the condition that $\frac{\sum_{s \in S} (1 - \hat{\tau}(s))}{\sum_{s \in S} \hat{\tau}(s)} \leq 1$, we have

$$\begin{aligned} \frac{\sum_{s \in S} \frac{1 - \hat{\tau}(s)}{\sum_{s \in S} \hat{\tau}(s)} \sum_{s \in S} \hat{\tau}(s) G_1\{w(s)t | s\}}{\sum_{s \in S} (1 - \hat{\tau}(s))t + \frac{\sum_{s \in S} (1 - \hat{\tau}(s))}{\sum_{s \in S} \hat{\tau}(s)}} &= \frac{\sum_{s \in S} \frac{1 - \hat{\tau}(s)}{\sum_{s \in S} \hat{\tau}(s)} \sum_{s \in S} \hat{\tau}(s) G_1\{t/w^{-1}(s) | s\}}{\sum_{s \in S} (1 - \hat{\tau}(s))G_1\left\{\frac{t}{\frac{\sum_{s \in S} (1 - \hat{\tau}(s))}{\sum_{s \in S} \hat{\tau}(s)} w^{-1}(s)}\right\}} \\ &= \sum_{s \in S} (1 - \hat{\tau}(s)) G_1\left\{\frac{t}{\frac{\sum_{s \in S} (1 - \hat{\tau}(s))}{\sum_{s \in S} \hat{\tau}(s)} w^{-1}(s)}\right\} \\ &= \sum_{s \in S} (1 - \hat{\tau}(s)) G_1\left\{\frac{\sum_{s \in S} \hat{\tau}(s)}{\sum_{s \in S} (1 - \hat{\tau}(s))} t\right\} \\ &= \sum_{s \in S} \hat{\tau}(s) G_1(t | s). \end{aligned}$$

Thus, take $t = t_{OR}^1$, we have

$$Q_w(t_{OR}^1) \geq Q_1(t_{OR}^1) \geq \alpha.$$

Hence, we have $t_{OR}^w \rightarrow t_{OR}^1$, which yields that

$$w(t_{OR}^w) \rightarrow w(t_{OR}^1) = 1(t_{OR}^1).$$

This concludes the proof of Theorem 1. ■

A.2 Proof of Proposition 1

Proof: Note that, by the definition of $v_h(s, s^0)$,

$$1 - \hat{\pi}(s) = \frac{\int_{s^0 \in S} v_h(s, s^0) dP(s^0)}{\int_{s^0 \in S} v_h(s, s^0) dP(s^0)} = \frac{\int_{s^0 \in S} K_h(s - s^0) dP(s^0)}{\int_{s^0 \in S} K_h(s - s^0) dP(s^0)}.$$

Thus, uniformly for all $s \in S$, we have

$$\begin{aligned} E_{s^0 \in S} [K_h(s - s^0)] &= E_{s^0 \in S} [K_h(s - s^0) I\{p(s^0) > \pi\}] \\ &= \int_{s^0 \in S} [K_h(s - s^0) P\{p(s^0) > \pi\}] dP(s^0). \end{aligned}$$

Under the conditions of the proposition, the eigenvalues of the Hessian matrix $P^{(2)}(p(s) > \pi)$ $\mathbb{R}^{d \times d}$ are bounded from above and below, namely, we have $C \geq \lambda_{\min}(s) \geq \lambda_{\max}(s) \geq C$, for all $s \in S$. Furthermore, $\hat{\pi}(s)$ is bounded, and thus the entries of the first derivative vector are also bounded. Hence, by multivariate Taylor expansion, we have

$$\begin{aligned} &\int_S K_h(s - s^0) P(p(s^0) > \pi) ds^0 \\ &= P(p(s) > \pi) \int_S K_h(s - s^0) ds^0 + v^T \int_S \frac{s - s^0}{h} K\left(\frac{s - s^0}{h}\right) ds^0 \\ &\quad + \frac{h^2}{2} O(1) \int_{\mathbb{R}^d} x^T x K(x) dx + o(h^2), \end{aligned}$$

where $v = (v_1, \dots, v_d)^T$ satisfies $v_j = O(1)$ for $j = 1, \dots, d$. It follows that, as $S \rightarrow S$ in the summation $\int_{s^0 \in S}$,

$$E(1 - \hat{\pi}(s)) \rightarrow \frac{\int_S K_h(s - s^0) P(p(s^0) > \pi) ds^0}{\int_S K_h(s - s^0) ds^0}.$$

Thus, we have, uniformly for all $s \in S$, there exists some constant $c > 0$,

$$[E\{\hat{\tau}^{\alpha}(s)\} - \tau^{\alpha}(s)]^2 \leq c \int_S \int_S \frac{K_h(s-s^0)}{h} K_h(s-s^0) ds^0 ds + ch^4 \int_S x^T x K(x) dx \int_S K_h(s-s^0) ds^0. \quad (A.2)$$

We next show the variance term. Note that, for some constant $c_0 > 1$,

$$\text{Var} \left(\frac{1}{|S|} \sum_{s^0 \in S} K_h(s-s^0) \right) = \text{Var} \left(\frac{1}{|S|} \sum_{s^0 \in S} [K_h(s-s^0) I\{p(s^0) > \alpha\}] \right) \\ \leq \frac{1}{|S|} \sum_{s^0 \in S} [K_h^2(s-s^0) P\{p(s^0) > \alpha\} (1 - P\{p(s^0) > \alpha\})].$$

Thus, as $|S| \rightarrow \infty$ in the summation $\sum_{s^0 \in S}$, we have,

$$\text{Var} \left(\frac{1}{|S|} \sum_{s^0 \in S} K_h(s-s^0) \right) \leq \frac{1}{|S|} \int_S \int_S \frac{K_h^2(s-s^0) P\{p(s^0) > \alpha\} (1 - P\{p(s^0) > \alpha\}) ds^0}{\left[\int_S K_h(s-s^0) ds^0 \right]^2} \\ \leq c^0 |S|^{-1} (1 - \alpha)^{-2} \int_S K^2(x) dx \int_S K_h(s-s^0) ds^0 \\ \leq c^{\infty} |S|^{-1} (1 - \alpha)^{-2} \int_S K_h(s-s^0) ds^0, \quad (A.3)$$

for some constants $c^0, c^{\infty} > 0$, where the last inequality comes from the fact that $K(\cdot)$ is positive and bounded. Combining (A.2) and (A.3) and letting $|S| \rightarrow \infty$ and thus $h \rightarrow 0$, by the conditions that the domain S is finite, $\int_{R^d} K(t) dt = 1$, $\int_{R^d} t K(t) dt = 0$ and $\int_{R^d} t^T t K(t) dt < 1$, Proposition 1 is proved. ■

A.3 Proof of Theorem 2

Proof: We first introduce the following procedure in order to prove the theorem.

Procedure 1 Calculate the weights as

$$\tilde{w}(s) = \left(\sum_{s^0 \in S} \frac{\hat{\tau}^{\alpha}(s^0)}{1 - \hat{\tau}^{\alpha}(s^0)} \right)^{-1} \frac{m \hat{\tau}^{\alpha}(s)}{1 - \hat{\tau}^{\alpha}(s)}, \quad s \in S \quad (A.4)$$

For hypothesis $H_0(s) : \tau(s) = 0$, define adjusted p-values as $\tilde{p}^w(s) = \min\{p(s)/\tilde{w}(s), 1\}$. Apply the BH procedure at level α to all adjusted p-values.

The following lemma develops the theoretical property of Procedure 1 for each realiza-

tion of $\{\sqrt{\cdot}(s), s \in S\}$.

Lemma 1 Under Conditions (A3) and (A5), and assume that, there exists a sufficiently small constant $\epsilon > 0$, such that $\hat{\pi}^{\otimes}(s) \in [\epsilon, 1 - \epsilon]$, then we have

$$\lim_{m \rightarrow \infty} \overline{\text{FDR}}_{\text{Procedure 1}} \leq \epsilon, \text{ and } \lim_{m \rightarrow \infty} P(\text{FDP}_{\text{Procedure 1}} \leq \epsilon + \epsilon) = 1.$$

for any $\epsilon > 0$.

We remark here that, from the proof of Lemma 1, by replacing $\hat{\pi}^{\otimes}(s)$ by $\hat{\pi}(s)$ in Procedure 1, Lemma 1 still holds under the conditions of Proposition 1.

Now we prove Theorem 2. We first note that, by the proof of Lemma 3 in Xia et al. (2019), Algorithm 1 is equivalent to the following algorithm.

Algorithm 2 An equivalent LAWS Procedure

1: Estimate the FDP by

$$\hat{\text{FDP}}_w(t) = \frac{\sum_{s \in S} \hat{\pi}(s)t}{\max\{\sum_{s \in S} \hat{p}^w(s) \mathbb{I}(p^w(s) \leq t), 1\}} \quad (\text{A.5})$$

2: Obtain the data-driven threshold $\hat{t}_w = \sup_t \{t : \hat{\text{FDP}}_w(t) \leq \epsilon\}$.

3: Reject $H_0(s)$ if $p^w(s) \leq \hat{t}_w$.

Thus, in the oracle case when $\{\hat{\pi}(s), s \in S\}$ are known, the corresponding oracle FDP and its conservative estimator can be equivalently written by

$$\text{FDP}_w(t) = \frac{\sum_{s \in S} \mathbb{I}(p^w(s) \leq t)}{\max\{R_w(t), 1\}}, \quad \hat{\text{FDP}}_w(t) = \frac{\sum_{s \in S} \hat{\pi}(s)t}{\max\{R_w(t), 1\}}, \quad (\text{A.6})$$

where $R_w(t) = \sum_{s \in S} \mathbb{I}(p^w(s) \leq t)$ denotes the oracle total number of rejections with the threshold t , and $w(s) = \frac{\hat{\pi}(s)}{1 - \hat{\pi}(s)}$. Define the oracle threshold $t_w = \sup_t \{t : \text{FDP}_w(t) \leq \epsilon\}$. The oracle decision rule $\hat{w}(t)$ is then equivalent to reject $H_0(s)$ if $p^w(s) \leq t_w$, and we have that $\text{FDP}_w(t_w) = \text{FDP}[\hat{w}\{p^w\}]$.

Based on the definition of t_w , to prove Theorem 2, it is enough to show that, uniformly for all $t \leq t_w$,

$$\frac{\sum_{s \in S} [\mathbb{I}(p^w(s) \leq t, \sqrt{\cdot}(s) = 0) - c\hat{\pi}(s)t]}{\sum_{s \in S} \hat{\pi}(s)t} \leq 0$$

in probability, for some constant $0 < c \leq 1$.

Note that the weights in the oracle procedure and the weights in Procedure 1 are proportional. Thus the adjusted p-values $\{p^w(s), i = 1, \dots, m\}$ and $\{\tilde{p}^w(s), i = 1, \dots, m\}$ have exactly the same order, with $\tilde{p}^w(s) = p^w(s) m^{-1} \prod_{s \in S} \frac{\uparrow^w(s)}{1 - \uparrow^w(s)}$. Also note that Procedure 1 is equivalent to find

$$\tilde{t}_w = \sup_t \left\{ t : \frac{m t}{\max_{s \in S} \{l(\tilde{p}^w(s) \cdot t), 1\}} \leq 1 \right\},$$

and reject the hypothesis $H_0(s)$ if $\tilde{p}^w(s) < \tilde{t}_w$. By letting $\tilde{t} = m^{-1} \prod_{s \in S} \frac{\uparrow^w(s)}{1 - \uparrow^w(s)} t$, we have that

$$\frac{m t}{\max_{s \in S} \{l(\tilde{p}^w(s) \cdot t), 1\}} = \frac{m \tilde{t}}{\max_{s \in S} \{l(p^w(s) \cdot \tilde{t}), 1\}},$$

and thus the threshold for $\tilde{p}^w(s)$ in Procedure 1 is

$$\tilde{t}_w = \sup_{\tilde{t}} \left\{ \tilde{t} : \frac{m \tilde{t}}{\max_{s \in S} \{l(p^w(s) \cdot \tilde{t}), 1\}} \leq 1 \right\} \cdot m^{-1} \prod_{s \in S} \frac{\uparrow^w(s)}{1 - \uparrow^w(s)}.$$

Hence the threshold for $p^w(s)$ in Procedure 1 is

$$t_w^1 = \sup_t \left\{ t : \frac{t}{\max_{s \in S} \{l(p^w(s) \cdot t), 1\}} \leq 1 \right\}.$$

Comparing the definition of t_w and t_w^1 , it is easily to see that $t_w^1 \leq t_w$. By the proofs of Lemma 1, we have that, the threshold for the z-values corresponding to threshold of p-values: t_w^1 is no larger than t_m as defined in the proof of Lemma 1. We further learn from the proofs of Lemma 1 that, for every realization of $\{\sqrt{\cdot}(s), s \in S\}$, uniformly for all $t \leq t_w^1$, as $m \rightarrow \infty$,

$$\frac{\prod_{\sqrt{\cdot}(s)=0} P(p^w(s) \cdot t) - \prod_{\sqrt{\cdot}(s)=0} P(p^w(s) \cdot t | \sqrt{\cdot}(s) = 0)}{\prod_{\sqrt{\cdot}(s)=0} P(p^w(s) \cdot t | \sqrt{\cdot}(s) = 0)} \rightarrow 0$$

in probability. Note that, by Condition (A4),

$$\frac{\prod_{\sqrt{\cdot}(s)=0} P(p^w(s) \cdot t | \sqrt{\cdot}(s) = 0)}{\prod_{s \in S} P(p^w(s) \cdot t, \sqrt{\cdot}(s) = 0)} \rightarrow 0$$

in probability, where $\{\check{\nu}(s), s \in S\}$ in the above equation represent random variables as modeled in (2.3). Thus we have

$$\frac{P_{s \in S} [I(p^w(s) \leq t, \check{\nu}(s) = 0) - P(p^w(s) \leq t, \check{\nu}(s) = 0)]}{P_{s \in S} (p^w(s) \leq t, \check{\nu}(s) = 0)} \xrightarrow{P} 0$$

in probability. Note that

$$\begin{aligned} \sum_{s \in S} P(p^w(s) \leq t, \check{\nu}(s) = 0) &= \sum_{s \in S} P(p^w(s) \leq t | \check{\nu}(s) = 0) P(\check{\nu}(s) = 0) \\ &= \sum_{s \in S} \frac{\hat{\nu}^w(s)}{1 - \hat{\nu}^w(s)} t (1 - \hat{\nu}^w(s)) \sum_{s \in S} \hat{\nu}^w(s) t. \end{aligned}$$

Thus we have that, uniformly for all $t \leq t_w^1$, there exists a constant $0 < c \leq 1$, as $m \rightarrow \infty$,

$$\frac{P_{s \in S} [I(p^w(s) \leq t, \check{\nu}(s) = 0) - c \hat{\nu}^w(s) t]}{P_{s \in S} \hat{\nu}^w(s) t} \xrightarrow{P} 0$$

in probability. This concludes the proof of Theorem 2. ■

A.4 Proof of Theorem 3

Proof: As shown in the proof of Theorem 2, the data-driven decision rule $\hat{w}(t)$ is equivalent to reject $H_0(s)$ if $p^w(s) \leq t_w$, and we have that $FDP_{\hat{w}}(t_w) = FDP[\hat{w}\{p^w_{(\frac{1}{k})}\}]$.

Similarly as the proof of Theorem 2, based on the proofs of Lemma 1 and Condition (A4), we have that

$$\frac{P_{s \in S} (I(p^{\hat{w}}(s) \leq t, \check{\nu}(s) = 0) - P(p^{\hat{w}}(s) \leq t, \check{\nu}(s) = 0))}{P_{s \in S} (p^{\hat{w}}(s) \leq t, \check{\nu}(s) = 0)} \xrightarrow{P} 0$$

in probability, uniformly for all $t \leq t_w^1$, where t_w^1 is defined as in the proof of Theorem 2, by replacing $\hat{\nu}^w(s)$ by $\hat{\nu}(s)$ in Procedure 1. Recall that $\hat{\nu}^w(s) = 1 - \frac{P(p(s) > \frac{1}{k})}{1 - \frac{1}{k}}$, and that the bias of $1 - \hat{\nu}^w(s)$ is always non-negative. By Proposition 1 and the fact that $\hat{\nu}^w(s) \in [\frac{1}{k}, 1 - \frac{1}{k}]$ for some sufficiently small constant $\frac{1}{k} > 0$,

$$\sum_{s \in S} P(p^{\hat{w}}(s) \leq t, \check{\nu}(s) = 0) = \sum_{s \in S} P(p^{\hat{w}}(s) \leq t | \check{\nu}(s) = 0) P(\check{\nu}(s) = 0)$$

$$\begin{aligned}
&= (1 + o(1)) \prod_{s \in S} \frac{\hat{p}(s)}{1 - \hat{p}(s)} t(1 - \hat{p}(s)) \\
&\stackrel{s \in S}{\rightarrow} (1 + o(1)) \prod_{s \in S} \frac{\hat{p}(s)}{1 - \hat{p}(s)} t(1 - \hat{p}(s)) \\
&= (1 + o(1)) \prod_{s \in S} \hat{p}(s) t,
\end{aligned}$$

where $o(1)$ in the above equations are in the limit of $S \rightarrow \infty$. Thus, based on Proposition 1, we have that, uniformly for all $t \leq t_w^1$, there exists a constant $0 < c \leq 1$ such that

$$\frac{\mathbb{P}_{s \in S}(\prod_{s \in S} (\hat{p}(s) - t) \sqrt{s} = 0) - c \hat{p}(s) t}{\prod_{s \in S} \hat{p}(s) t} \rightarrow 0$$

in probability. Namely, the data-driven procedure provides a more conservative FDR and FDP control asymptotically. This concludes the proof of Theorem 3.

A.5 Proof of Lemma 1

We now prove Lemma 1 below. We let $p^w(s) = \tilde{p}^w(s)$ in the proof of this lemma for notation simplicity. We arrange $\{s \in S\}$ in any order $\{s_1, \dots, s_m\}$ and note that $\mathbb{P}_{\sqrt{(s_i)=0}}(z_i^w \leq t)$ is equivalent to $\mathbb{P}_{\sqrt{(s_i)=0}}(z_i \leq t | \sqrt{(s_i)=0})$ and $\prod_{i=1, \dots, m} \mathbb{P}(z_i \leq t, \sqrt{(s_i)=0})$ for a given realization of the hypotheses, where $z^w = (1 - p^w(s_i)/2)$, for $i = 1, \dots, m$. We first show

that by applying BH to the weighted p-values controls FDR exactly under the independence of z_i^w . We then prove that, in the dependent case, its performance asymptotically the same as case when z_i^w are independent.

Let $t_m = (2 \log m - 2 \log \log m)^{1/2}$. By Condition (A5), we have

$$\prod_{\sqrt{(s_i)=1}} \mathbb{P}(|z_i| \leq (c \log m)^{1/2 + o(1/4)}) \leq \{1/(c^{1/2} + 1)\} (\log m)^{1/2},$$

with probability going to one, for some constant $c > 0$. Recall that we have $w(s) = \prod_{s \in S} \frac{\hat{p}(s)}{1 - \hat{p}(s)} \leq 1$ and $\hat{p}(s) \in [0, 1]$. Thus, for those indices $i \in H_1$ (equivalently $\sqrt{(s_i)=1}$) such that $|z_i| \leq (c \log m)^{1/2 + o(1/4)}$, we have

$$p^w(s_i) = p(s_i)/w(s_i) \leq (1 - ((c \log m)^{1/2 + o(1/4)})/w(s_i) = o(m^{-M}),$$

for any constant $M > 0$. Thus we have

$$\sum_{i=1}^m \mathbb{P}\{z_i^w \leq (2 \log m)^{1/2} - \{1/(\hat{\tau}^{1/2} \epsilon) + \epsilon\}(\log m)^{1/2}\},$$

with probability going to one. Hence, with probability tending to one, we have

$$\mathbb{P}\left\{\frac{\sum_{i=1}^m \mathbb{P}\{z_i^w \leq (2 \log m)^{1/2} - \{1/(\hat{\tau}^{1/2} \epsilon) + \epsilon\}(\log m)^{1/2}\}}{m} \geq 2m\{1/(\hat{\tau}^{1/2} \epsilon) + \epsilon\}^{-1}(\log m)^{-1/2}\right\}.$$

Because $1 - \beta(t_m) \leftarrow 1/\{(2\hat{\tau})^{1/2}t_m\} \exp(-t_m^2/2)$, it suffices to show that, uniformly in $0 \leq t \leq t_m$, there exists a constant $0 < c \leq 1$, such that

$$\mathbb{P}\left\{\frac{\sum_{\sqrt{(s_i)}=0} \mathbb{P}(z_i^w \leq t) - cm_0 G(t)}{cm_0 G(t)} \leq 0\right\}, \quad (\text{A.7})$$

in probability, where $G(t) = 2(1 - \beta(t))$.

We first consider the case when z_i^w are independent with each other. For $i = 1, \dots, m$, the ideal choice of the threshold t^0 for z_i^w in order to control the FDR is that

$$t^0 = \inf\{t \geq 0, \mathbb{P}\left\{\frac{\sum_{\sqrt{(s_i)}=0} \mathbb{P}(z_i^w \leq t)}{\sum_{i=1}^m \mathbb{P}(z_i^w \leq t)} \geq \epsilon\right\} \leq \alpha\}.$$

It is easy to show that, under independence of z_i^w ,

$$\frac{\mathbb{P}\left\{\sum_{\sqrt{(s_i)}=0} \mathbb{P}(z_i^w \leq t)\right\}}{\sum_{\sqrt{(s_i)}=0} \mathbb{P}(z_i^w \leq t)} \geq \frac{\mathbb{P}\left\{\sum_{\sqrt{(s_i)}=0} \mathbb{P}(z_i^w \leq t)\right\}}{\sum_{i=1}^m \mathbb{P}(z_i^w \leq t)} \geq 0, \text{ as } m \rightarrow \infty.$$

Thus a good estimate of t^0 can be written by

$$\hat{t}^0 = \inf\{t \geq 0, \mathbb{P}\left\{\frac{\sum_{\sqrt{(s_i)}=0} \mathbb{P}(z_i^w \leq t)}{\sum_{i=1}^m \mathbb{P}(z_i^w \leq t)} \geq \epsilon\right\} \leq \alpha\}. \quad (\text{A.8})$$

Since the spatial locations $\{s \in S\}$ does not change the null distribution of p-values, according to Theorem 1 in Genovese et al. (2006), and by the fact that the original p-values of the null hypotheses are uniformly distributed, the procedure by applying BH procedure

on the weighted p-values controls the FDR at level $\alpha m_0/m$. That is, if

$$k = \max\{i : p_{(i)}^w \leq i\alpha/m\},$$

where $p_{(1)}^w \leq \dots \leq p_{(m)}^w$ are the ordered weighted p-values, and we reject all k hypotheses associated with $p_{(1)}^w, \dots, p_{(k)}^w$, then we have

$$E \frac{P_{\sqrt{(s_i)=0}} \sum_{i=1}^k (p^w(s_i) - p_{(k)}^w)}{\max\{\sum_{i=1}^k (p^w(s_i) - p_{(k)}^w), 1\}} \leq \alpha m_0/m. \quad (\text{A.9})$$

By the definition of z_i^w , it is equivalent to reject all k hypotheses with

$$\frac{2m(1 - (z_{(i)}^w))}{i} \leq \alpha.$$

That is to find

$$\hat{t} = \inf\{t \geq 0, P_{\sqrt{(s_i)=0}} \sum_{i=1}^k \frac{2m(1 - (z_i^w - t))}{i} \leq \alpha\}, \quad (\text{A.10})$$

and reject all hypotheses with $z_i^w \leq \hat{t}$. This yields that

$$E \frac{P_{\sqrt{(s_i)=0}} \sum_{i=1}^k (z_i^w - \hat{t})}{\sum_{i=1}^k (z_i^w - \hat{t})} \leq \alpha m_0/m.$$

Hence, this procedure is more conservative than rejecting all hypotheses with $z_i^w \leq t^0$ as defined in (A.8). Thus, there exists a constant $0 < c \leq 1$, such that, uniformly in $0 \leq t \leq t_m$,

$$\frac{P_{\sqrt{(s_i)=0}} \sum_{i=1}^k (z_i^w - t) - cm_0 G(t)}{cm_0 G(t)} \leq 0. \quad (\text{A.11})$$

Under Assumption (A1), that is, when the weighted z-values are weakly dependent with each other, by the proof of Theorem 1 in Xia et al. (2019) and the fact that

$$\frac{1}{n} \{1 - [1 - \{(c_1 \log m + c_2 \log \log m)^{1/2}\}/w(s_i)]\} = c_1 \log m + c_2 \log \log m + c_3,$$

for some constant c_1, c_2, c_3 , we have,

$$\frac{P_{\sqrt{(s_i)=0}}(I(z_i^w \leq t) - P(z_i^w \leq t))}{P_{\sqrt{(s_i)=0}}(z_i^w \leq t)} \leq 0, \quad (\text{A.12})$$

in probability, uniformly in $0 \leq t \leq t_m$. By (A.11) and (A.12), (A.7) is proved and thus Procedure 1 controls FDR and FDP asymptotically under dependency. This concludes the proof of Lemma 1. ■

B Explanation of the bias-variance tradeoff in choosing η

There is a bias-variance tradeoff in the choice of η in the proposed estimator $\hat{\eta}^\eta(s)$. It is easy to see that a large η will simultaneously reduce the bias (desirable) and decrease the sample size (undesirable).

To reduce the bias in the proposed estimator $\hat{\eta}^\eta(s)$, one needs to choose a relatively large η to ensure the “purity” of $T(\eta) = \{s \in S : p(s) > \eta\}$, i.e. we wish to have a screening set where majority of the cases come from the null. Although the common choice of $\eta = 0.5$ succeeds in many situations, we propose a new scheme to carefully calibrate η that is adaptive to the observed data. Let $\hat{\eta}$ be the threshold determined by the BH algorithm with $\alpha = 0.9$. Then roughly speaking, in the subset $T(\hat{\eta}) \approx \{s \in S : p(s) \geq \hat{\eta}\}$, 90% of the cases come from the null (e.g. the expected proportion of false positives made by BH). It is anticipated that in the remaining set $T(\eta) = \{s \in S : p(s) > \eta\}$, which is used to construct our estimator, an overwhelming proportion (more than 90%) of the cases should come from the null. In the simulation (say the 2D rectangle case), $\hat{\eta}$ roughly ranges from 0.38 to 0.47 when we run BH algorithm at $\alpha = 0.9$. This data-driven scheme ensures the purity of the screening set while maintaining a larger sample size compared to the standard choice of $\eta = 0.5$.

C Explanation of Equation (3.6) in Section 3.3

In this section we justify the approximation in Equation (3.6). Denote $\mathcal{V}(t) = \{v(s, t) : s \in S\}$ a class of testing rules, where $v(s, t) = I\{p^v(s) \geq t\}$, $v(s)$ is the pre-specified weight and $p^v(s) = \min \left\{ \frac{p(s)}{v(s)}, 1 \right\}$. Denote by $V(t) = \sum_{s \in H_0} v(s, t)$ and $R(t) = \sum_{s \in S} v(s, t)$. We

show that $Q^v(t)$ provides a good approximation to the actual FDR level under the following condition:

$$\text{Var} [R(t)/E\{R(t)\}] = o(1). \quad (\text{C.13})$$

We remark here that for a fixed threshold $t > 0$ and $v(s) \in [\epsilon, 1 - \epsilon]$ for a sufficiently small constant $\epsilon > 0$, Condition (C.13) is satisfied if $\text{Var} \sum_{s \in S} I\{p(s) \leq t\}/m = o(1)$. This is a weaker condition compared to (A2) in Section 4.1. Hence the condition can be fulfilled by both BH and LAWS under the general class of dependence structures being considered in this article.

Proposition 2 Assume that Condition (C.13) holds, then we have

$$\text{FDR}\{v(t)\} = Q^v(t) + o(1) = \frac{E \sum_{s \in H_0} v(s, t)^{2H}}{E \sum_{s \in H_0} v(s, t) + E \sum_{s \in H_1} v(s, t)} + o(1).$$

Proof of Proposition 2: Note that $R(t) = 0$ implies $V(t) = 0$, hence we have

$$\text{FDR}\{v(t)\} = E \frac{V(t)}{R(t)} I\{R(t) > 0\}, \text{ and } Q^v(t) = \frac{[V(t) I\{R(t) > 0\}]}{E\{R(t)\}}.$$

It follows that

$$\begin{aligned} | \text{FDR}\{v(t)\} - Q^v(t) | &\leq E \left| \frac{V(t)}{R(t)} - \frac{V(t)}{E\{R(t)\}} I\{R(t) > 0\} \right| \\ &= E \left| \frac{V(t) R(t) - E\{R(t)\} V(t)}{R(t) E\{R(t)\}} \right| I\{R(t) > 0\}. \end{aligned}$$

Note that $V(t)$ is always no larger than $R(t)$, we have

$$| \text{FDR}\{v(t)\} - Q^v(t) | \leq \text{Var}^{1/2} [R(t)/E\{R(t)\}] = o(1),$$

which proves the proposition.