

Distribution-free Junta Testing

ZHENGYANG LIU, Shanghai Jiao Tong University, China

XI CHEN, ROCCO A. SERVEDIO, YING SHENG, and JINYU XIE, Columbia University, USA

We study the problem of testing whether an unknown n -variable Boolean function is a k -junta in the *distribution-free* property testing model, where the distance between functions is measured with respect to an arbitrary and unknown probability distribution over $\{0, 1\}^n$. Our first main result is that distribution-free k -junta testing can be performed, with one-sided error, by an adaptive algorithm that uses $\tilde{O}(k^2)/\epsilon$ queries (independent of n). Complementing this, our second main result is a lower bound showing that any *non-adaptive* distribution-free k -junta testing algorithm must make $\Omega(2^{k/3})$ queries even to test to accuracy $\epsilon = 1/3$. These bounds establish that while the optimal query complexity of non-adaptive k -junta testing is $2^{\Theta(k)}$, for adaptive testing it is $\text{poly}(k)$, and thus show that adaptivity provides an exponential improvement in the distribution-free query complexity of testing juntas.

CCS Concepts: • **Theory of computation** → **Streaming, sublinear and near linear time algorithms**;

Additional Key Words and Phrases: Property testing, distribution-free model

ACM Reference format:

Zhengyang Liu, Xi Chen, Rocco A. Servedio, Ying Sheng, and Jinyu Xie. 2018. Distribution-free Junta Testing. *ACM Trans. Algorithms* 15, 1, Article 1 (September 2018), 23 pages.

<https://doi.org/10.1145/3264434>

1 INTRODUCTION

Property testing of Boolean functions was first considered in the seminal works of Blum et al. (1993) and Rubinfeld and Sudan (1996) and has developed into a robust research area at the intersection of sub-linear algorithms and complexity theory. Roughly speaking, a property tester for a class C of functions from $\{0, 1\}^n$ to $\{0, 1\}$ is a randomized algorithm that is given some form of access to the (unknown) input Boolean function f and must with high probability distinguish the case that $f \in C$ versus the case that f is ϵ -far from every function $g \in C$. In the usual (uniform-distribution) property testing scenario, the testing algorithm may access f by making black-box queries on inputs $x \in \{0, 1\}^n$, and the distance between two functions f and g is measured with respect to the uniform distribution on $\{0, 1\}^n$; the goal is to develop algorithms that make as few queries as possible. Many different classes of Boolean functions have been studied from this perspective (Blum et al. 1993; Alon et al. 2005; Bhattacharyya et al. 2010; Parnas et al. 2002; Diakonikolas et al.

This article has been accepted by STOC 2018.

This work is supported by the NSF CCF under Grants No. 1703925, No. 1420349, and No. 1563155.

Authors' addresses: Z. Liu, Shanghai Jiao Tong University, 800 Dongchuan Road, Minhang District, Shanghai, China; email: lzy5118@sjtu.edu.cn; X. Chen, R. A. Servedio, Y. Sheng, and J. Xie, Columbia University, Department of Computer Science, New York, NY, 10027, USA; emails: {xichen, rocco, ys2982, jinyu}@cs.columbia.edu.

Permission to make digital or hard copies of all or part of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyrights for components of this work owned by others than ACM must be honored. Abstracting with credit is permitted. To copy otherwise, or republish, to post on servers or to redistribute to lists, requires prior specific permission and/or a fee. Request permissions from Permissions@acm.org.

© 2018 Association for Computing Machinery.

1549-6325/2018/09-ART1 \$15.00

<https://doi.org/10.1145/3264434>

2007; Goldreich et al. 2000; Fischer et al. 2002; Chakrabarty and Seshadhri 2013; Chen et al. 2014, 2015; Khot et al. 2015; Belovs and Blais 2016; Khot and Shinkar 2016; Chakrabarty and Seshadhri 2016; Baleshazar et al. 2016; Chen et al. 2017b, 2017c; Matulef et al. 2010, 2009; Blais et al. 2011; Blais and Kane 2012; Gopalan et al. 2011), and see other works referenced in the surveys by Ron (2008, 2010) and Goldreich (2010). Among these, the class of k -juntas—Boolean functions that depend only on (an unknown set of) at most k of their n input variables—is one of the best-known and most intensively investigated such classes (Fischer et al. 2004; Chockler and Gutfreund 2004; Blais 2008, 2009; Buhrman et al. 2013; Servedio et al. 2015), with ongoing research on junta testing continuing right up to the present (Chen et al. 2017a).

The query complexity of junta testing in the uniform distribution framework is now well understood. Improving on $\text{poly}(k)/\epsilon$ -query algorithms given in Fischer et al. (2004) (which introduced the junta testing problem), in Blais (2008), Blais gave a non-adaptive algorithm that makes $\tilde{O}(k^{3/2})/\epsilon$ queries, and in Blais (2009), Blais gave an $O(k \log k + k/\epsilon)$ -query adaptive algorithm. On the lower bounds side, Fischer et al. (2004) initially gave an $\Omega(\sqrt{k})$ lower bound for non-adaptively testing k -juntas, which also implies an $\Omega(\log k)$ lower bound for adaptive testing. Chockler and Gutfreund improved the adaptive lower bound to $\Omega(k)$ in Chockler and Gutfreund (2004), and very recently Chen et al. (2017a) gave an $\tilde{\Omega}(k^{3/2})/\epsilon$ non-adaptive lower bound. Thus, in both the adaptive and non-adaptive uniform distribution settings, the query complexity of k -junta testing has now been pinned down to within logarithmic factors.

Distribution-free property testing. This work studies the junta testing problem in the *distribution-free* property testing model that was first introduced in Goldreich et al. (1998). In this model, the distance between Boolean functions is measured with respect to a distribution $\mathcal{D}$ over $\{0, 1\}^n$, which is arbitrary and unknown to the testing algorithm. Since the distribution is unknown, in this model the testing algorithm is allowed (in addition to making black-box queries) to draw random labeled samples $(\mathbf{x}, f(\mathbf{x}))$, where each $\mathbf{x}$ is independently distributed according to $\mathcal{D}$. The query complexity of an algorithm in this framework is the worst-case total number of black-box oracle calls plus random labeled samples that are used, across all possible distributions. (It follows that distribution-free testing of a class C requires at least as many queries as testing C in the standard uniform-distribution model.)

Distribution-free property testing is in the spirit of similar distribution-free models in computational learning theory such as Valiant’s celebrated PAC learning model (Valiant 1984). Such models are attractive because of their minimal assumptions; they are well motivated both because in many natural settings the uniform distribution over $\{0, 1\}^n$ may not be the best way to measure distances, and because they capture the notion of an algorithm dealing with an unknown and arbitrary environment (modeled here by the unknown and arbitrary distribution $\mathcal{D}$ over $\{0, 1\}^n$ and the unknown and arbitrary Boolean function $f : \{0, 1\}^n \rightarrow \{0, 1\}$). Researchers have studied distribution-free testing of a number of Boolean function classes, including monotone functions, low-degree polynomials, dictators (1-juntas) and k -juntas (Halevy and Kushilevitz 2007), disjunctions and conjunctions (monotone and non-monotone), decision lists, and linear threshold functions (Glasner and Servedio 2009; Dolev and Ron 2011; Chen and Xie 2016). Since depending on few variables is an appealingly flexible “real-world” property in comparison with more highly structured syntactically defined properties, we feel that junta testing is a particularly natural task to study in the distribution-free model.

Prior results on distribution-free junta testing. Given how thoroughly junta testing has been studied in the uniform distribution model, surprisingly little was known in the distribution-free setting. The adaptive $\Omega(k)$ and non-adaptive $\tilde{\Omega}(k^{3/2})/\epsilon$ uniform-distribution lower bounds from Chockler and Gutfreund (2004) and Chen et al. (2017a) mentioned earlier trivially extend

to the distribution-free model, but no other lower bounds on distribution-free junta testing were known prior to this work. On the positive side, Halevy and Kushilevitz (2007) showed that any class C that has (i) a one-sided error uniform-distribution testing algorithm and (ii) a self-corrector, has a one-sided error distribution-free testing algorithm. As $\text{poly}(k)/\epsilon$ -query one-sided junta testers were given already in Fischer et al. (2004), and k -juntas have $O(2^k)$ -query self-correctors (Alon and Weinstein 2012), this yields a one-sided non-adaptive distribution-free junta tester with query complexity $O(2^k/\epsilon)$. No other results were known.

Thus, prior to this work there were major gaps in our understanding of distribution-free k -junta testing: Is the query complexity of this problem polynomial in k , exponential in k , or somewhere in between? Does adaptivity confer an exponential advantage, a sub-exponential advantage, or no advantage at all? Our results, described below, answer both these questions.

1.1 Our Results

Our main positive result is a $\text{poly}(k)/\epsilon$ -query one-sided adaptive algorithm for distribution-free k -junta testing:

THEOREM 1.1 (UPPER BOUND). *For any $\epsilon > 0$, there is a one-sided distribution-free adaptive algorithm for ϵ -testing k -juntas with $\tilde{O}(k^2)/\epsilon$ queries.*

Theorem 1.1 shows that k -juntas stand in interesting contrast with many other well-studied classes of Boolean functions in property testing such as conjunctions, decision lists, linear threshold functions, and monotone functions. For each of these classes, distribution-free testing requires dramatically more queries than uniform-distribution testing: for the first three classes the separation is $\text{poly}(1/\epsilon)$ queries in the uniform setting (Parnas et al. 2002; Matulef et al. 2010) versus $n^{\Omega(1)}$ queries in the distribution-free setting (Glasner and Servedio 2009; Chen and Xie 2016); for n -variable monotone functions $\text{poly}(n)$ queries suffice in the uniform setting (Goldreich et al. 2000; Khot et al. 2015), whereas Halevy and Kushilevitz (2007) show that $2^{\Omega(n)}$ queries are required in the distribution-free setting. In contrast, Theorem 1.1 shows that for k -juntas the query complexities of uniform-distribution and distribution-free testing are polynomially related (indeed, within at most a quadratic factor of each other).

Complementing the strong upper bound that Theorem 1.1 gives for adaptive testers, our main negative result is an $\Omega(2^{k/3})$ -query lower bound for non-adaptive testers:

THEOREM 1.2 (LOWER BOUND). *For $k \leq n/200$, any non-adaptive algorithm that distribution-free ϵ -tests k -juntas over $\{0, 1\}^n$, for $\epsilon = 1/3$, must have query complexity $\Omega(2^{k/3})$.*

Theorems 1.1 and 1.2 together show that adaptivity enables an exponential improvement in the distribution-free query complexity of testing juntas. This is in sharp contrast with uniform-distribution junta testing, where the adaptive and non-adaptive query complexities are polynomially related (with an exponent of only $3/2$). To the best of our knowledge, this is the first example of a exponential separation between adaptive and nonadaptive distribution-free testers.

1.2 Ideas and Techniques

The algorithm. As a first step toward our $\tilde{O}(k^2)/\epsilon$ -query algorithm, in Section 3, we first present a simple one-sided adaptive algorithm, which we call **SimpleDJunta**, that distribution-free tests k -juntas using $O((k/\epsilon) + k \log n)$ queries. **SimpleDJunta** uses binary search and is an adaptation to the distribution-free setting of the $O((k/\epsilon) + k \log n)$ -query uniform-distribution algorithm, which is implicit in Blais (2009). The algorithm maintains a set I of *relevant* variables: a string $x \in \{0, 1\}^n$ has been found for each $i \in I$ such that $f(x) \neq f(x^{(i)})$ (we use $x^{(i)}$ to denote the string obtained by flipping the i th bit of x), and the algorithm rejects only when $|I|$ becomes larger than k . In each round, the algorithm samples a string $x \leftarrow \mathcal{D}$ and a subset R of $\bar{I} := [n] \setminus I$ uniformly at random.

A simple lemma, Lemma 3.2, states that if f is far from every k -junta with respect to $\mathcal{D}$, then $f(\mathbf{x}) \neq f(\mathbf{x}^{(\mathbf{R})})$ with at least some moderately large probability as long as $|\mathbf{I}| \leq k$, where we use $\mathbf{x}^{(\mathbf{R})}$ to denote the string obtained from $\mathbf{x}$ by flipping every coordinate in $\mathbf{R}$. With such a pair $(\mathbf{x}, \mathbf{x}^{(\mathbf{R})})$ in hand, it is straightforward to find a new relevant variable using binary search over coordinates in $\mathbf{R}$ (see Figure 1), with at most $\log n$ additional queries.

To achieve a query complexity that is independent of n , clearly one must employ a more efficient approach than binary search over $\Omega(n)$ coordinates (since most likely the set $\mathbf{R}$ has size $\Omega(n)$ for the range of k we are interested in). In the uniform-distribution setting this is accomplished in Blais (2009) by first randomly partitioning the variable space $[n]$ into $s = \text{poly}(k/\epsilon)$ disjoint blocks $B_1, \dots, B_s$ of variables and carrying out binary search over blocks (see Figure 2) rather than over individual coordinates; this reduces the cost of each binary search to $\log(k/\epsilon)$ rather than $\log n$. The algorithm maintains a set of *relevant blocks*: two strings $x, y \in \{0, 1\}^n$ have been found for each such block B that satisfy $f(x) \neq f(y)$ and $y = x^{(S)}$ with $S \subseteq B$, and the algorithm rejects when more than k relevant blocks have been found. In each round the algorithm samples two strings $\mathbf{x}, \mathbf{y}$ uniformly at random conditioned on their agreeing with each other on the relevant blocks that have already been found in previous rounds; if $f(\mathbf{x}) \neq f(\mathbf{y})$, then the binary search over blocks is performed to find a new relevant block. To establish the correctness of this approach, Blais (2009) employs a detailed and technical analytic argument based on the influence of coordinates and the Efron-Stein orthogonal decomposition of functions over product spaces. This machinery is well suited for dealing with product distributions, and indeed the analysis of Blais (2009) goes through for any product distribution over $\{0, 1\}^n$ (and even for more general finite domains and ranges). However, it is far from clear how to extend this machinery to work for the completely unstructured distributions $\mathcal{D}$ that must be handled in the distribution-free model.

Our main distribution-free junta testing algorithm, denoted **MainDJunta**, draws ideas from both **SimpleDJunta** (mainly Lemma 3.2) and the uniform distribution tester of (Blais 2009). To avoid the $\log n$ cost, the algorithm carries out binary search over blocks rather than over individual coordinates, and maintains a set of disjoint relevant blocks $B_1, \dots, B_\ell$, i.e., for each B_j a pair of strings x^j and y^j have been found such that they agree with each other over $\overline{B_j}$ and satisfy $f(x^j) \neq f(y^j)$. Let w^j be the projection of x^j (and y^j) over B_j and let g_j be the Boolean function over $\{0, 1\}^{B_j}$ obtained from f by setting variables in $\overline{B_j}$ to w^j . For clarity, we assume further that every function g_j is very close to a *literal* (i.e., for some $\tau \in \{x_{i_j}, \overline{x_{i_j}}\}$, we have $g_j(x) = \tau$ for all $x \in \{0, 1\}^{B_j}$ for some $i_j \in B_j$) under the *uniform* distribution. (To justify this assumption, we note that if g_j is far from every literal under the uniform distribution, then it is easy to split B_j further into two relevant blocks using the uniform distribution algorithm of (Blais 2009).) Let $I = \{i_j : j \in [\ell]\}$. Even though the algorithm does not know I , there is indeed a way to draw uniformly random subsets $\mathbf{R}$ of $\overline{I}$. First, we draw a partition of B_j into $\mathbf{P}_j$ and $\mathbf{Q}_j$ uniformly at random, for each j . Since g_j is close to a literal, it is not difficult to figure out whether $\mathbf{P}_j$ or $\mathbf{Q}_j$ contains the hidden i_j , say it is $\mathbf{P}_j$ for every j . Then the union of all $\mathbf{Q}_j$'s together with a uniformly random subset of $B_1 \cup \dots \cup B_\ell$, denoted by $\mathbf{R}$, turns out to be a uniformly random subset of $\overline{I}$. With $\mathbf{R}$ in hand, Lemma 3.2 implies that $f(\mathbf{x}) \neq f(\mathbf{x}^{(\mathbf{R})})$ with high probability when $\mathbf{x} \leftarrow \mathcal{D}$, and when this happens, one can carry out binary search over blocks to increase the number of relevant blocks by one. In Section 4.1, we explain the intuition behind the main algorithm in more detail.

The lower bound. As we explain in Section 2, a q -query non-adaptive distribution-free tester is a randomized algorithm A that works as follows. When A is run on an input pair $(\phi, \mathcal{D})$ ¹ it is first

¹For clarity, throughout our discussion of lower bounds, we write ϕ to indicate a function that may be either a “yes-function” or a “no-function,” f to denote a “yes-function” and g to denote a “no-function.”

given the result $(y^1, \phi(y^1)), \dots, (y^q, \phi(y^q))$ of q queries from the sampling oracle. Based on them, it queries the black-box oracle q times on strings $z^1, \dots, z^q$. The z^j 's may depend on the random pairs $(y^i, \phi(y^i))$ received from the sampling oracle, but the j th black-box query string z^j may not depend on the responses $\phi(z^1), \dots, \phi(z^{j-1})$ to any of the $j - 1$ earlier black-box queries.

As is standard in property testing lower bounds, our argument employs a distribution $\mathcal{YES}$ over yes-instances and a distribution $\mathcal{NO}$ over no-instances. Here, $\mathcal{YES}$ is a distribution over (function, distribution) pairs $(f, \mathcal{D})$ in which f is guaranteed to be a k -junta; $\mathcal{NO}$ is a distribution over pairs $(g, \mathcal{D})$ such that with probability $1 - o(1)$, g is $1/3$ -far from every k -junta with respect to $\mathcal{D}$. To prove the desired lower bound against non-adaptive distribution-free testers, it suffices to show that for $q = 2^{k/3}$, any deterministic non-adaptive algorithm A as described above is roughly equally likely to accept whether it is run on an input drawn from $\mathcal{YES}$ or from $\mathcal{NO}$.

Our construction of the $\mathcal{YES}$ and $\mathcal{NO}$ distributions is essentially as follows. In making a draw either from $\mathcal{YES}$ or from $\mathcal{NO}$, first $m = \Theta(2^k \log n)$ strings are selected uniformly at random from $\{0, 1\}^n$ to form a set S , and the distribution $\mathcal{D}$ in both $\mathcal{YES}$ and $\mathcal{NO}$ is set to be the uniform distribution over S . Also, in both $\mathcal{YES}$ and $\mathcal{NO}$, a “background” k -junta h is selected uniformly at random by first picking a set J of k variables at random and then a random truth table for h over the variables in J . We view the variables in J as partitioning $\{0, 1\}^n$ into 2^k disjoint “sections” depending on how they are set.

In the case of a draw from $\mathcal{YES}$, the Boolean function f that goes with the above-described $\mathcal{D}$ is simply the background junta $f = h$. In the case of a draw from $\mathcal{NO}$, the function g that goes with $\mathcal{D}$ is formed by modifying the background junta h in the following way (roughly speaking; see Section 5.1 for precise details): for each $z \in S$, we toss a fair coin $b(z)$ and set the value of all the strings in z 's section that lie within Hamming distance $0.4n$ from z (including z itself) to $b(z)$ (see Figure 7). Note that the value of g at each string in S is a fair coin toss, which is completely independent of the background junta h . Using the choice of m it can be argued (see Section 5.1) that with high probability g is $1/3$ -far from every k -junta with respect to $\mathcal{D}$ as $(g, \mathcal{D}) \leftarrow \mathcal{NO}$.

The rough idea of why a pair $(f, \mathcal{D}) \leftarrow \mathcal{YES}$ is difficult for a $(q = 2^{k/3})$ -query non-adaptive algorithm A to distinguish from a pair $(g, \mathcal{D}) \leftarrow \mathcal{NO}$ is as follows. Intuitively, in order for A to distinguish the no-case from the yes-case, it must obtain two strings x^1, x^2 that belong to the same section but are labeled differently. Since there are 2^k sections but q is only $2^{k/3}$, by the birthday paradox it is very unlikely that A obtains two such strings among the q samples $y^1, \dots, y^q$ that it is given from the distribution $\mathcal{D}$. In fact, in both the yes and no cases, writing $(\phi, \mathcal{D})$ to denote the (function, distribution) pair, the distribution of the q pairs $(y^1, \phi(y^1)), \dots, (y^q, \phi(y^q))$ will be statistically very close to $(x^1, b_1), \dots, (x^q, b_q)$, where each pair (x^j, b_j) is independently drawn uniformly from $\{0, 1\}^n \times \{0, 1\}$. Intuitively, this translates into the examples $(y^i, \phi(y^i))$ from the sampling oracle “having no useful information” about the set J of variables that the background junta depends on.

What about the q strings $z^1, \dots, z^q$ that A feeds to the black-box oracle? It is also unlikely that any two elements of $y^1, \dots, y^q, z^1, \dots, z^q$ belong to the same section but are labeled differently. Fix an $i \in [q]$, we give some intuition here as to why it is very unlikely that there is any j such that z^i lies in the same section as y^j but has $f(z^i) \neq f(y^j)$ (via a union bound, the same intuition handles all $i \in [q]$). Intuitively, since the random examples from the sampling oracle provide no useful information about the set J defining the background junta, the only thing that A can do in selecting z^i is to choose how far it lies, in terms of Hamming distance, from the points in $y^1, \dots, y^q$ (which, recall, are uniform random). Fix $j \in [q]$, if z^i is within Hamming distance $0.4n$ from y^j , then even if z^i lies in the same section as y^j it will be labeled the same way as y^j , whether we are in the yes-case or the no-case. However, if z^i is farther than $0.4n$ in Hamming distance from y^j , then it is

overwhelmingly likely that z^i will lie in a different section from y^j (since it is very unlikely that all $0.4n$ of the flipped coordinates avoid the k -element set J). We prove Theorem 1.2 in Section 5 via a formal argument that proceeds somewhat differently from but is informed by the above intuitions.

Organization. In Section 2, we define the distribution-free testing model and introduce some useful notation. In Section 3 we present **SimpleDJunta** and prove Lemma 3.2 in its analysis. In Section 4, we present our main algorithm **MainDJunta** and prove Theorem 1.1, and in Section 5, we prove Theorem 1.2.

2 PRELIMINARIES

Notation. We use $[n]$ to denote $\{1, \dots, n\}$. We use f and g to denote Boolean functions, which are maps from $\{0, 1\}^n$ to $\{0, 1\}$ for some positive integer n . We use the calligraphic font (e.g., $\mathcal{D}$ and $\mathcal{NO}$) to denote probability distributions, boldface letters such as $\mathbf{x}$ to denote random variables, and write “ $\mathbf{x} \leftarrow \mathcal{D}$ ” to indicate that $\mathbf{x}$ is a random variable drawn from a distribution $\mathcal{D}$. We write $\mathbf{x} \leftarrow \{0, 1\}^n$ to denote that $\mathbf{x}$ is a string drawn uniformly at random. Given $S \subseteq [n]$, we also write $\mathbf{R} \leftarrow S$ to indicate that $\mathbf{R}$ is a subset of S drawn uniformly at random, i.e., each index $i \in S$ is included in $\mathbf{R}$ independently with probability $1/2$.

Given a subset $B \subseteq [n]$, we use $\bar{B}$ to denote its complement with respect to $[n]$, and $\{0, 1\}^B$ to denote the set of all binary strings of length $|B|$ with coordinates indexed by $i \in B$. Given an $x \in \{0, 1\}^n$ and a $B \subseteq [n]$, we write $x_B \in \{0, 1\}^B$ to denote the projection of x over coordinates in B and $x^{(B)} \in \{0, 1\}^n$ to denote the string obtained from x by flipping coordinates in B . Given $x \in \{0, 1\}^B$ and $y \in \{0, 1\}^{\bar{B}}$, we write $x \circ y \in \{0, 1\}^n$ to denote their concatenation, a string that agrees with x over coordinates in B and agrees with y over $\bar{B}$. (As an example of the notation, given $x, y \in \{0, 1\}^n$ and $B \subseteq [n]$, $x_B \circ y_{\bar{B}}$ denotes the string that agrees with x over B and with y over $\bar{B}$.) Given $x, y \in \{0, 1\}^n$, we write $d(x, y)$ to denote the Hamming distance between x and y , and $\text{diff}(x, y) \subseteq [n]$ to denote the set of coordinates $i \in [n]$ with $x_i \neq y_i$.

Given $f, g : \{0, 1\}^n \rightarrow \{0, 1\}$ and a probability distribution $\mathcal{D}$ over $\{0, 1\}^n$, we write

$$\text{dist}_{\mathcal{D}}(f, g) := \Pr_{z \leftarrow \mathcal{D}} [f(z) \neq g(z)]$$

to denote the *distance* between f and g with respect to $\mathcal{D}$. Given a class $\mathbb{C}$ of Boolean functions,

$$\text{dist}_{\mathcal{D}}(f, \mathbb{C}) := \min_{g \in \mathbb{C}} (\text{dist}_{\mathcal{D}}(f, g))$$

denotes the *distance* between f and $\mathbb{C}$ with respect to $\mathcal{D}$, where the minimum is taken over g with the same number of variables as f . We say f is ϵ -far from $\mathbb{C}$ with respect to $\mathcal{D}$ if $\text{dist}_{\mathcal{D}}(f, \mathbb{C}) \geq \epsilon$.

We will often work with restrictions of Boolean functions. Given $f : \{0, 1\}^n \rightarrow \{0, 1\}$, $S \subseteq [n]$ and a string $z \in \{0, 1\}^S$, the *restriction of f over S by z* , denoted by $f|_z$, is the Boolean function $g : \{0, 1\}^{\bar{S}} \rightarrow \{0, 1\}$ defined by $g(x) = f(x \circ z)$ for all $x \in \{0, 1\}^{\bar{S}}$.

Distribution-free property testing. Now, we can define distribution-free property testing:

Definition 2.1. We say an algorithm A has *oracle access* to a pair $(f, \mathcal{D})$, where $f : \{0, 1\}^n \rightarrow \{0, 1\}$ is an unknown Boolean function and $\mathcal{D}$ is an unknown probability distribution over $\{0, 1\}^n$, if it can (1) access f via a black-box oracle that returns $f(x)$ when a string $x \in \{0, 1\}^n$ is queried, and (2) access $\mathcal{D}$ via a sampling oracle that, upon each request, returns a pair $(\mathbf{x}, f(\mathbf{x}))$, where $\mathbf{x} \leftarrow \mathcal{D}$ independently.

Let $\mathbb{C}$ be a class of Boolean functions. A *distribution-free testing algorithm A for $\mathbb{C}$* is a randomized algorithm that, given as input a distance parameter $\epsilon > 0$ and oracle access to a pair $(f, \mathcal{D})$, accepts with probability at least $2/3$ if $f \in \mathbb{C}$ and rejects with probability at least $2/3$ if f is ϵ -far from $\mathbb{C}$.

with respect to $\mathcal{D}$. We say A is *one-sided* if it always accepts when $f \in \mathfrak{C}$. The *query complexity* of a distribution-free testing algorithm is the number of queries made on f plus the number of samples drawn from $\mathcal{D}$.

One may assume without loss of generality that a distribution-free testing algorithm consists of two phases: In the first phase, the algorithm draws a certain number of sample pairs $(x, f(x))$ from $\mathcal{D}$; in the second phase, it makes black-box queries to f . In general, a query $x \in \{0, 1\}^n$ made by the algorithm in the second phase may depend on sample pairs it receives in the first phase (e.g. it can choose to query a string that is close to a sample received in the first phase) and results of queries to f made so far. In Section 5, we will prove lower bounds on *non-adaptive* distribution-free testing algorithms. An algorithm is said to be non-adaptive if its black-box queries made in the second phase do not depend on results of previous black-box queries, i.e., all queries during the second phase can be made in a single batch (though we emphasize that they may depend on samples the algorithm received in the first phase).

Juntas and literals. We study the distribution-free testing of the class of k -juntas. Recall that a Boolean function f is a k -junta if it depends on at most k variables. More precisely, f is a k -junta if there exists a list $1 \leq i_1 < \dots < i_k \leq n$ of k indices and a Boolean function $g : \{0, 1\}^k \rightarrow \{0, 1\}$ over k variables such that $f(x_1, \dots, x_n) = g(x_{i_1}, \dots, x_{i_k})$ for all $x \in \{0, 1\}^n$.

We say that a Boolean function f is a *literal* if f depends on exactly one variable, i.e. for some $i \in [n]$, we have that either $f(x) = x_i$ for all x or $f(x) = \overline{x_i}$ for all x . Note that the two constant (all-1 and all-0) functions are one-juntas but are *not* literals.

We often use the term “*block*” to refer to a *nonempty* subset of $[n]$, which should be interpreted as a nonempty subset of the n variables of a Boolean function $f : \{0, 1\}^n \rightarrow \{0, 1\}$. The following definition of distinguishing pairs and relevant blocks will be heavily used in our algorithms.

Definition 2.2 (Distinguishing Pairs and Relevant Blocks). Given $x, y \in \{0, 1\}^n$ and a block $B \subseteq [n]$, we say that (x, y) is a *distinguishing pair* for B if $x_{\overline{B}} = y_{\overline{B}}$ and $f(x) \neq f(y)$. We say B is a *relevant block* of f if such a distinguishing pair exists for B (or equivalently, the influence of B in f is positive).

When $B = \{i\}$ is a relevant block, we simply say that the i th variable is relevant to f .

As will become clear later, all our algorithms reject a function f only when they have found $k + 1$ pairwise disjoint blocks $B_1, \dots, B_{k+1}$ and a distinguishing pair for each B_i . When this occurs, it means that $B_1, \dots, B_{k+1}$ are pairwise disjoint relevant blocks of f , which implies that f cannot be a k -junta. As a result, our algorithms are one-sided. To prove their correctness, it suffices to show that they reject with probability at least $2/3$ when f is ϵ -far from k -juntas with respect to $\mathcal{D}$.

For the standard property testing model under the uniform distribution, Blais (2009) obtained a nearly optimal algorithm:

THEOREM 2.3 (BLAIS 2009). *There exists a one-sided, $O((k/\epsilon) + k \log k)$ -query algorithm **UniformJunta** (f, k, ϵ) that rejects f with probability at least $2/3$ when it is ϵ -far from k -juntas under the uniform distribution. Moreover, it rejects only when it has found $k + 1$ pairwise disjoint blocks and a distinguishing pair of f for each of them.*

Binary Search. The standard binary search procedure (see Figure 1) takes as input two strings $x, y \in \{0, 1\}^n$ with $f(x) \neq f(y)$, makes $O(\log n)$ queries on f , and returns a pair of strings $x', y' \in \{0, 1\}^n$ with $f(x') \neq f(y')$ and $x' = y'^{(i)}$ for some $i \in \text{diff}(x, y)$, i.e., a distinguishing pair for the i th variable for some $i \in \text{diff}(x, y)$.

However, we cannot afford to use the standard binary search procedure directly in our main algorithm due to its query complexity of $O(\log n)$; recall, our goal is to have the query complexity

Procedure BinarySearch(f, x, y)**Input:** Query access to $f: \{0, 1\}^n \rightarrow \{0, 1\}$ and two strings $x, y \in \{0, 1\}^n$ with $f(x) \neq f(y)$.**Output:** Two strings $x', y' \in \{0, 1\}^n$ with $f(x') \neq f(y')$ and $x' = y'^{(i)}$ for some $i \in \text{diff}(x, y)$.

- (1) Let $B \subseteq [n]$ be the set such that $x = y^{(B)}$.
- (2) If $|B| = 1$ return x and y .
- (3) Partition (arbitrarily) B into B_1 and B_2 of size $\lfloor |B|/2 \rfloor$ and $\lceil |B|/2 \rceil$, respectively.
- (4) Query $f(x^{(B_1)})$.
- (5) If $f(x) \neq f(x^{(B_1)})$, return **BinarySearch**($f, x, x^{(B_1)}$).
- (6) Otherwise, return **BinarySearch**($f, x^{(B_1)}, y$).

Fig. 1. Description of the standard binary search procedure.

Procedure BlockBinarySearch($f, x, y, B_1, \dots, B_r$)**Input:** Query access to $f: \{0, 1\}^n \rightarrow \{0, 1\}$, two strings $x, y \in \{0, 1\}^n$ with $f(x) \neq f(y)$, and a sequence of pairwise disjoint blocks $B_1, \dots, B_r$ for some $r \geq 1$ with $\text{diff}(x, y) \subseteq B_1 \cup \dots \cup B_r$.**Output:** Two strings $x', y' \in \{0, 1\}^n$ with $f(x') \neq f(y')$ and $\text{diff}(x, y) \subseteq B_i$ for some $i \in [r]$.

- (1) If $r = 1$ return x and y .
- (2) Let $t = \lfloor r/2 \rfloor$ and B be the intersection of $\text{diff}(x, y)$ and $B_1 \cup \dots \cup B_t$.
- (3) Query $f(x^{(B)})$.
- (4) If $f(x) \neq f(x^{(B)})$, return **BlockBinarySearch**($f, x, x^{(B)}, B_1, \dots, B_t$).
- (5) Otherwise, return **BlockBinarySearch**($f, x^{(B)}, y, B_{t+1}, \dots, B_r$).

Fig. 2. Description of the blockwise version of the binary search procedure.

depend on k only. Instead, we will employ a blockwise version of the binary search procedure, as described in Figure 2. It takes as input two strings $x, y \in \{0, 1\}^n$ with $f(x) \neq f(y)$ and a sequence of pairwise disjoint blocks $B_1, \dots, B_r$ such that

$$\text{diff}(x, y) \subseteq B_1 \cup \dots \cup B_r$$

(i.e., (x, y) is a distinguishing pair for $B_1 \cup \dots \cup B_r$), makes $O(\log r)$ queries on f , and returns two strings $x', y' \in \{0, 1\}^n$ satisfying $f(x') \neq f(y')$ and $\text{diff}(x', y') \subseteq B_i$ for some $i \in [r]$ (i.e., (x', y') is a distinguishing pair for one of the blocks B_i in the input).

3 WARMUP: A TESTER WITH $O((K/\epsilon) + K \log N)$ QUERIES

As a warmup, we present in this section a simple, one-sided distribution-free algorithm for testing k -juntas (**SimpleDJunta**, where the capital letter D is a shorthand for distribution-free). It uses $O((k/\epsilon) + k \log n)$ queries, where n as usual denotes the number of variables of the function being tested. The idea behind **SimpleDJunta** and its analysis (Lemma 3.2 below) will be useful in the next section where we present our main algorithm to remove the dependency on n .

The algorithm **SimpleDJunta** maintains a set $I \subset [n]$, which is such that a distinguishing pair has been found for each $i \in I$ (i.e., I is a set of relevant variables of f discovered so far). The algorithm sets $I = \emptyset$ at the beginning and rejects only when $|I|$ reaches $k + 1$, which implies immediately that the algorithm is one-sided. **SimpleDJunta** proceeds round by round. In each round it draws a pair of random strings x and y with $x_I = y_I$. If $f(x) \neq f(y)$, then the standard binary search procedure is used on x and y to find a distinguishing pair for a new variable $i \in \bar{I}$, which is then added to I .

The description of the algorithm can be found in Figure 3. The following theorem establishes its correctness.

Algorithm SimpleDJunta($f, \mathcal{D}, k, \epsilon$)

Input: Oracle access to a Boolean function $f: \{0, 1\}^n \rightarrow \{0, 1\}$ and a probability distribution $\mathcal{D}$ over $\{0, 1\}^n$, a positive integer k , and a distance parameter $\epsilon > 0$.

Output: Either “accept” or “reject.”

- (1) Set $I = \emptyset$.
- (2) Repeat $8(k+1)/\epsilon$ times:
 - (3) Sample $\mathbf{x} \leftarrow \mathcal{D}$ and a subset $\mathbf{R}$ of $\bar{I}$ uniformly at random. Set $\mathbf{y} = \mathbf{x}^{(\mathbf{R})}$.
 - (4) If $f(\mathbf{x}) \neq f(\mathbf{y})$, then run the standard binary search on $\mathbf{x}, \mathbf{y}$ to find a distinguishing pair for a new relevant variable $i \in \mathbf{R} \subseteq \bar{I}$. Set $I = I \cup \{i\}$.
 - (6) If $|I| > k$, then halt and output “reject.”
- (7) Halt and output “accept.”

Fig. 3. Description of the distribution-free testing algorithm **SimpleDJunta** for k -juntas.

THEOREM 3.1. (i) The algorithm **SimpleDJunta** makes $O((k/\epsilon) + k \log n)$ queries and always accepts when f is a k -junta. (ii) It rejects with probability at least $2/3$ if f is ϵ -far from k -juntas with respect to $\mathcal{D}$.

PROOF. For part (i), note that the algorithm only runs binary search (and spends $O(\log n)$ queries) when $f(\mathbf{x}) \neq f(\mathbf{y})$ and this happens at most $k+1$ times (even though the algorithm has $O(k/\epsilon)$ rounds). The rest of part (i) is immediate from the description of the algorithm.

For part (ii), it suffices to show that when $|I| \leq k$ at the beginning of a round, a new relevant variable is discovered in this round with at least a modestly large probability. For this purpose, we use the following simple but crucial lemma and note the fact that $\mathbf{x}$ and $\mathbf{y}$ on line 3 can be equivalently drawn by first sampling $\mathbf{x} \leftarrow \mathcal{D}$ and $\mathbf{w} \leftarrow \{0, 1\}^n$ and then setting $\mathbf{y} = \mathbf{x}_I \circ \mathbf{w}_{\bar{I}}$ (the way we draw $\mathbf{x}$ and $\mathbf{y}$ in Figure 3 via $\mathbf{R} \leftarrow \bar{I}$ makes it easier to connect with the main algorithm in the next section). $\square$

LEMMA 3.2. If f is ϵ -far from k -juntas with respect to $\mathcal{D}$, then for any $I \subset [n]$ of size at most k , we have

$$\Pr_{\mathbf{x} \leftarrow \mathcal{D}, \mathbf{w} \leftarrow \{0, 1\}^n} [f(\mathbf{x}) \neq f(\mathbf{x}_I \circ \mathbf{w}_{\bar{I}})] \geq \epsilon/2. \quad (1)$$

Before proving Lemma 3.2, we use it to finish the proof of part (ii). Assuming Lemma 3.2 and that f is ϵ -far from k -juntas with respect to $\mathcal{D}$, for each round in which $|I| \leq k$ the algorithm finds a new relevant variable with probability at least $\epsilon/2$. Using a coupling argument, the probability that the algorithm rejects f (i.e., $|I|$ reaches $k+1$ during the $8(k+1)/\epsilon$ rounds) is at least the probability that

$$\sum_{i=1}^{8(k+1)/\epsilon} Z_i \geq k+1,$$

where Z_i 's are i.i.d. $\{0, 1\}$ -variables that are 1 with probability $\epsilon/2$. It follows from the Chernoff bound that the latter probability is at least $2/3$. This finishes the proof of the theorem.

PROOF OF LEMMA 3.2. Let I be a subset of $[n]$ of size at most k . To prove Equation (1) for I , we use I to define the following Boolean function $h: \{0, 1\}^n \rightarrow \{0, 1\}$ over n variables: for each $\mathbf{x} \in \{0, 1\}^n$, we set

$$h(\mathbf{x}) := \arg \max_{b \in \{0, 1\}} \left\{ \Pr_{\mathbf{w} \leftarrow \{0, 1\}^n} [f(\mathbf{x}_I \circ \mathbf{w}_{\bar{I}}) = b] \right\},$$

where we break ties arbitrarily. Then for any $\mathbf{x} \in \{0, 1\}^n$, we have

$$\Pr_{\mathbf{w} \leftarrow \{0, 1\}^n} [f(\mathbf{x}_I \circ \mathbf{w}_{\bar{I}}) = h(\mathbf{x})] \geq 1/2. \quad (2)$$

Furthermore, we have

$$\begin{aligned}
& \Pr_{\mathbf{x} \leftarrow \mathcal{D}, \mathbf{w} \leftarrow \{0,1\}^n} [f(\mathbf{x}) \neq f(\mathbf{x}_I \circ \mathbf{w}_{\bar{I}})] \\
&= \sum_{z \in \{0,1\}^n} \Pr_{\mathbf{x} \leftarrow \mathcal{D}} [\mathbf{x} = z] \cdot \Pr_{\mathbf{w} \leftarrow \{0,1\}^n} [f(z) \neq f(z_I \circ \mathbf{w}_{\bar{I}})] \\
&\geq \sum_{z \in \{0,1\}^n} \Pr_{\mathbf{x} \leftarrow \mathcal{D}} [\mathbf{x} = z] \cdot ((1/2) \cdot \mathbf{1}[f(z) \neq h(z)]) \\
&= (1/2) \cdot \Pr_{\mathbf{x} \leftarrow \mathcal{D}} [f(\mathbf{x}) \neq h(\mathbf{x})] \geq \epsilon/2,
\end{aligned}$$

where the first inequality follows from Equation (2) and the second inequality follows from the assumption that f is ϵ -far from every k -junta with respect to $\mathcal{D}$ and the fact that h is a k -junta (since it only depends on variables in I and $|I| \leq k$). This finishes the proof of the lemma. $\square$

4 PROOF OF THEOREM 1.1: A TESTER WITH $\tilde{O}(K^2)/\epsilon$ QUERIES

In this section, we present our main $\tilde{O}(k^2)/\epsilon$ -query algorithm for the distribution-free testing of k -juntas. We start with some intuition behind the algorithm.

4.1 Intuition

Recall that the factor of $\log n$ in the query complexity of **SimpleDJunta** from the previous section is due to the use of the standard binary search procedure. To avoid it, one could choose to terminate each call to binary search early but this ends up giving us relevant *blocks* of variables instead of relevant variables. To highlight the challenge, imagine that the algorithm has found so far $\ell \leq k$ many pairwise disjoint relevant blocks B_j , $j \in [\ell]$; i.e., it has found a distinguishing pair for each block B_j . By definition, each B_j must contain at least one relevant variable $i_j \in B_j$. However, we do not know exactly which variable in B_j is i_j , and thus it is not clear how to draw a set $\mathbf{R}$ from $\bar{I}$ uniformly at random, where $I = \{i_j : j \in [\ell]\}$, as on line 3 of **SimpleDJunta**, to apply Lemma 3.2 to discover a new relevant block. It seems that we are facing a dilemma when trying to improve **SimpleDJunta** to remove the $\log n$ factor: unless we pin down a set of relevant variables, it is not clear how to draw a random set from their complement, but pinning down a single relevant variable using the standard binary search procedure would already cost $\log n$ queries.

To explain the main idea behind our $\tilde{O}(k^2)/\epsilon$ -query algorithm, let us assume again that $\ell \leq k$ many disjoint relevant blocks B_j have been found so far, with a distinguishing pair $(x^{[j]}, y^{[j]})$ for each B_j (satisfying that $\text{diff}(x^{[j]}, y^{[j]}) \subseteq B_j$ and $f(x^{[j]}) \neq f(y^{[j]})$ by definition). Let

$$w^{[j]} = (x^{[j]})_{\overline{B_j}} = (y^{[j]})_{\overline{B_j}} \in \{0,1\}^{\overline{B_j}}.$$

Next let us assume further that the function $g_j := f|_{w^{[j]}}$, for each $j \in [\ell]$, is a *literal*, i.e., either $g_j(z) = z_{i_j}$ for all $z \in \{0,1\}^{B_j}$ or $g_j(z) = \bar{z}_{i_j}$ for all $z \in \{0,1\}^{B_j}$, for some variable $i_j \in B_j$, but the variable i_j is of course unknown to the algorithm. (While this may seem very implausible, we make this assumption for now and explain below why it is not too far from real situations.)

To make progress, we draw a random two-way partition of each B_j into $\mathbf{P}_j$ and $\mathbf{Q}_j$, i.e., each $i \in B_j$ is added to $\mathbf{P}_j$ or $\mathbf{Q}_j$ with probability $1/2$ (so they are disjoint and $B_j = \mathbf{P}_j \cup \mathbf{Q}_j$). We make three simple but crucial observations to increase the number of disjoint relevant blocks by one.

- (1) Since g_j is assumed to be a literal on the i_j th variable (and by the definition of g_j we have query access to g_j), it is easy to tell whether $i_j \in \mathbf{P}_j$ or $i_j \in \mathbf{Q}_j$, simply by picking an arbitrary string $x \in \{0,1\}^{B_j}$ and comparing $g_j(x)$ with $g_j(x^{(\mathbf{P}_j)})$. Below, we assume that the

algorithm correctly determines whether i_j is in $\mathbf{P}_j$ or $\mathbf{Q}_j$ for all $j \in [\ell]$. We let $\mathbf{S}_j$ denote the element of $\{\mathbf{P}_j, \mathbf{Q}_j\}$ that contains i_j and let $\mathbf{T}_j$ denote the other one. We also assume below that the algorithm has obtained a distinguishing pair of g_j for each block $\mathbf{S}_j$.

- (2) Next, we draw a subset $\mathbf{T}$ of $\overline{B_1 \cup \dots \cup B_\ell}$ uniformly at random. Crucially, the way that $\mathbf{P}_j$ and $\mathbf{Q}_j$ were drawn, and the above assumption that $\mathbf{S}_j$ contains i_j , implies that

$$\mathbf{R} := \mathbf{T} \cup \mathbf{T}_1 \cup \dots \cup \mathbf{T}_\ell$$

is indeed a subset of $\bar{I}$ drawn uniformly at random (recall that $I = \{i_j : j \in [\ell]\}$), since other than those in I , each variable is included in $\mathbf{R}$ independently with probability $1/2$. If we draw a random string $\mathbf{x} \leftarrow \mathcal{D}$, then Lemma 3.2 implies that $f(\mathbf{x}) \neq f(\mathbf{y})$, where $\mathbf{y} = \mathbf{x}^{(\mathbf{R})}$, with probability at least $\epsilon/2$.

- (3) Finally, assuming that $f(\mathbf{x}) \neq f(\mathbf{y})$ (with $\text{diff}(\mathbf{x}, \mathbf{y}) = \mathbf{R}$), running the blockwise binary search on $\mathbf{x}, \mathbf{y}$ and blocks $\mathbf{T}, \mathbf{T}_1, \dots, \mathbf{T}_\ell$ will lead to a distinguishing pair for one of these blocks and will only require $O(\log \ell) \leq O(\log k)$ queries. If it is a distinguishing pair for $\mathbf{T}$, then we can add $\mathbf{T}$ to the list of relevant blocks $B_1, \dots, B_\ell$ and they remain pairwise disjoint. If it is $\mathbf{T}_j$ for some $j \in [\ell]$, then we can replace B_j in the list by $\mathbf{S}_j$ and $\mathbf{T}_j$, for each of which we have found a distinguishing pair (recall that a distinguishing pair has already been found for each $\mathbf{S}_j$ in the first step). In either case, we have that the number of pairwise disjoint relevant blocks grows by one.

Coming back to the assumption we made earlier, although g_j is very unlikely to be a literal, it must fall into one of the following three cases: (1) close to a literal; (2) close to a (all-0 or all-1) constant function; or (3) far from 1-juntas. Here in all cases “close” and “far” means with respect to the *uniform distribution* over $\{0, 1\}^{B_j}$. As we discuss in more detail in the rest of the section, with some more careful probability analysis the above arguments generalize to the case in which every g_j is only close to (rather than exactly) a literal. However, if one of the blocks B_j is in case (2) or (3), then (using the fact that we have a distinguishing pair for B_j) it is easy to split B_j into two blocks and find a distinguishing pair for each of them. (For example, for case (3) this can be done by running Blais’s uniform distribution junta testing algorithm.) As a result, we can always make progress by increasing the number of pairwise disjoint relevant blocks by one. Our algorithm basically keep repeating these steps until the number of such blocks reaches $k + 1$.

4.2 Description of the Main Algorithm and the Proof of Correctness

Our algorithm **MainDJunta**($f, \mathcal{D}, k, \epsilon$) is described in Figure 4. It maintains two collections of blocks $V = \{B_1, \dots, B_v\}$ (V for “verified”) and $U = \{C_1, \dots, C_u\}$ (U for “unverified”) for some non-negative integers v and u . They are set to be $\emptyset$ at initialization and always satisfy the following properties:

- (A) $B_1, \dots, B_v, C_1, \dots, C_u \subseteq [n]$ are pairwise disjoint (nonempty) blocks of variables;
- (B) A distinguishing pair has been found for each of these blocks. For notational convenience, we use $(x^{[j]}, y^{[j]})$ to denote the distinguishing pair for each B_j and $(x^{[C]}, y^{[C]})$ to denote the distinguishing pair for each block $C \in U$. We also use the notation

$$w^{[j]} := (x^{[j]})_{\overline{B_j}} = (y^{[j]})_{\overline{B_j}} \in \{0, 1\}^{\overline{B_j}} \quad \text{and} \quad w^{[C]} := (x^{[C]})_{\overline{C}} = (y^{[C]})_{\overline{C}} \in \{0, 1\}^{\overline{C}},$$

and we let $g_j := f|_{w^{[j]}}$ and $g_C := f|_{w^{[C]}}$, Boolean functions over $\{0, 1\}^{B_j}$ and $\{0, 1\}^C$, respectively.

The algorithm rejects only when the total number of blocks $v + u \geq k + 1$ so it is one-sided.

Algorithm MainDJunta($f, \mathcal{D}, k, \epsilon$) with the same input / output as **SimpleDJunta** in Figure 3.

- (1) Initialization: Set $V = U = \emptyset$, $r_1 = 64k/\epsilon$ and $r_2 = 3(k+1)$.
- (2) While $r_1 > 0$ and $r_2 > 0$ do (letting $V = \{B_1, \dots, B_v\}$ and $U = \{C_1, \dots, C_u\}$)
- (3) If $u = 0$, then
- (4) Set r_1 to be $r_1 - 1$.
- (5) For $j = 1$ to v do $((x^{[j]}, y^{[j]}):$ distinguishing pair for B_j , $w^{[j]} = (x^{[j]})_{\overline{B_j}}$, $g_j = f \upharpoonright_{w^{[j]}}$)
- (6) Draw a random partition P_j, Q_j of B_j and run **WhereIsTheLiteral**(g_j, P_j, Q_j).
- (7) If it returns a distinguishing pair of g_j for P_j , set $S_j = P_j$ and $T_j = Q_j$;
- (8) Else if it returns a distinguishing pair of g_j for Q_j , set $S_j = Q_j$ and $T_j = P_j$;
- (9) Else (it returns “fail”), skip this round and go back to line 2.
- (10) Draw $x \leftarrow \mathcal{D}$ and a subset T of $B_1 \cup \dots \cup B_v$ uniformly at random.
- (11) If $f(x) \neq f(y)$, where $y = x^{(R)}$ with $R = T \cup T_1 \cup \dots \cup T_v$, then
- (12) Run the blockwise binary search on x and y with blocks $T, T_1, \dots, T_v$.
- (13) If a distinguishing pair of f for T is found, add T to U .
- (14) Else (a distinguishing pair of f for T_{j^*} , for some $j^* \in [v]$, is found)
- (15) Concatenate $w^{[j^*]}$ to the distinguishing pair of g_{j^*} for S_{j^*} found on line 7-8.
- (16) This gives us a distinguishing pair of f for S_{j^*} .
- (17) Remove B_{j^*} from V and add both S_{j^*} and T_{j^*} to U .
- (18) Else (i.e., $u > 0$)
- (19) Set r_2 to be $r_2 - 1$.
- (20) Pick a $C \in U$ arbitrarily; let (x, y) be its distinguishing pair, $w = x_{\overline{C}}$ and $g = f \upharpoonright_w$.
- (21) If **Literal**(g) returns “true,” remove C from U and add it to V .
- (22) Else (it returns disjoint subsets C', C^* of C and each a distinguishing pair of g_C)
- (23) Concatenate w to obtain a distinguishing pair of f for each of C' and C^*
- (24) Remove C from U and add both C' and C^* to U .
- (25) If $|V| + |U| \geq k + 1$, then halt and output “reject.”
- (26) Halt and output “accept.”

Fig. 4. Description of the distribution-free testing algorithm **MainDJunta** for k -juntas.

Throughout the algorithm and its analysis, we set a key parameter $\gamma := 1/(8k)$. Blocks in V are intended to be those that have been “verified” to satisfy the condition that g_j is γ -close to a literal (for some unknown variable $i_j \in B_j$) under the uniform distribution, while blocks in U have not been verified yet so they may or may not satisfy the condition. More formally, at any point in the execution of the algorithm we say that the algorithm is in *good condition* if its current collections V and U satisfy conditions (A), (B), and

(C) Every $g_j, j \in [v]$, is γ -close to a literal under the uniform distribution over $\{0, 1\}^{B_j}$.

The algorithm **MainDJunta**(f, k, ϵ) starts with $V = U = \emptyset$ and proceeds round by round. For each round, we consider two different types that the round may have: type 1 is that $u = 0$, and type 2 is that $u > 0$. In a type-1 round (with $u = 0$), we follow the idea sketched in Section 4.1 to increase the total number of disjoint relevant blocks by one. We prove the following lemma for this case in Section 4.3.

LEMMA 4.1. *Assume that **MainDJunta** is in good condition at the beginning of a round, with $u = 0$ and $v \leq k$. Then it must remain in good condition at the end of this round. Moreover, letting V' and U' be the two collections of blocks at the end of this round, we have either $|V'| = v$ and $|U'| = 1$, or $|V'| = v - 1$ and $|U'| = 2$ with probability at least $\epsilon/4$.*

In a type-2 round (with $u \geq 1$), we pick an arbitrary block C from U and check whether g_C is close to a literal under the uniform distribution. The following lemma, which we prove in Section 4.4, shows that with high probability, either C is moved to collection V and the algorithm

remains in good condition, or the algorithm finds two disjoint subsets of C and a distinguishing pair for each of them so that V stays the same but $|U|$ goes up by one (we add these two blocks to U , since they have not yet been verified).

LEMMA 4.2. *Assume that **MainDJunta** is in good condition at the beginning of a round, with $u > 0$ and $v + u \leq k$. Then with probability at least $1 - 1/(64k)$, one of the following two events occurs at the end of this round (letting V' and U' be the two collections of blocks at the end of this round):*

- (1) *The algorithm remains in good condition with $|V'| = |V| + 1$ and $|U'| = |U| - 1$; or*
- (2) *The algorithm remains in good condition with $V' = V$ and $|U'| = |U| + 1$.*

Assuming Lemmas 4.1 and 4.2, we are ready to prove the correctness of **MainDJunta**.

THEOREM 4.3. (i) *The algorithm **MainDJunta** makes $\tilde{O}(k^2)/\epsilon$ queries and always accepts f when it is a k -junta. (ii) *It rejects with probability at least $2/3$ when f is ϵ -far from every k -junta with respect to $\mathcal{D}$.**

PROOF OF THEOREM 4.3 ASSUMING LEMMAS 4.1 AND 4.2. **MainDJunta** is one-sided, since it rejects f only when it has found $k + 1$ pairwise disjoint relevant blocks of f . Its query complexity is $(\# \text{ type-1 rounds}) \cdot (\# \text{ queries per type-1 round}) + (\# \text{ type-2 rounds}) \cdot (\# \text{ queries per type-2 round})$
 $= O(k/\epsilon) \cdot (O(k) + O(\log k)) + O(k) \cdot O(\log k) \cdot O(k) = O(k^2/\epsilon) + O(k^2 \log k) = O(k^2 \log k)/\epsilon$.

In the rest of the proof we show that it rejects f with probability at least $2/3$ when f is ϵ -far from every k -junta with respect to $\mathcal{D}$.

For this purpose, we introduce a simple potential function F to measure the progress:

$$F(V, U) := 3|V| + 2|U|.$$

Each round of the algorithm is either of type-1 (when $|U| = 0$) or of type-2 (when $|U| > 0$). By Lemma 4.1, if the algorithm is in good condition at the beginning of a type-1 round, then the algorithm ends the round in good condition and the potential function F goes up by at least one with probability at least $\epsilon/4$ (in which case, we say that the algorithm succeeds in this type-1 round). By Lemma 4.2, if the algorithm is in good condition at the beginning of a type-2 round, then the algorithm ends the round in good condition and F goes up by at least one with probability at least $1 - 1/(32k)$ (in which case, we say it succeeds in this type-2 round).

Note that F is 0 at the beginning ($V = U = \emptyset$) and that we must have $|U| + |V| \geq k + 1$ (and thus, the algorithm rejects) when the potential function F reaches $3(k + 1)$ or above. As a result, a necessary condition for the algorithm to accept is that one of the following two events occurs:

- E_1 : At least one of the (no more than $3(k + 1)$ many) type-2 rounds fails.
- E_2 : E_1 does not occur (so the algorithm ends every round in good condition, and the reason that the algorithm accepts cannot be that it uses up all the $3(k + 1)$ many type-2 rounds), and the algorithm uses up all the $64k/\epsilon$ many type-1 rounds but at most $3k + 2$ of them succeed.

By a union bound, the probability of E_1 is at most

$$3(k + 1) \cdot 1/(64k) \leq 6k \cdot 1/(64k) < 1/8.$$

As the algorithm ends every round in good condition, it follows from Lemma 4.1 from a coupling argument that the probability of E_2 is at most the probability that

$$\sum_{i=1}^{64k/\epsilon} Z_i \leq 3k + 2,$$

Subroutine WhereIsTheLiteral(g, P, Q)**Input:** Oracle access to a Boolean function g over $\{0, 1\}^B$ with P, Q being a partition of B .**Output:** Either a distinguishing pair for P , a distinguishing pair for Q , or “fail.”

- (1) Draw $\mathbf{w} \leftarrow \{0, 1\}^Q$ and $\mathbf{z} \leftarrow \{0, 1\}^P$ independently and uniformly at random.
- (2) If $g(\mathbf{w} \circ \mathbf{z}) \neq g(\mathbf{w} \circ \mathbf{z}^{(P)})$, return $(\mathbf{w} \circ \mathbf{z}, \mathbf{w} \circ \mathbf{z}^{(P)})$ as a distinguishing pair for P .
- (3) Draw $\mathbf{w}' \leftarrow \{0, 1\}^P$ and $\mathbf{z}' \leftarrow \{0, 1\}^Q$ independently and uniformly at random.
- (4) If $g(\mathbf{w}' \circ \mathbf{z}') \neq g(\mathbf{w}' \circ \mathbf{z}'^{(Q)})$, return $(\mathbf{w}' \circ \mathbf{z}', \mathbf{w}' \circ \mathbf{z}'^{(Q)})$ as a distinguishing pair for Q .
- (5) Return “fail.”

Fig. 5. Description of the subroutine **WhereIsTheLiteral**.

where Z_i 's are i.i.d. $\{0, 1\}$ -valued random variables that take 1 with probability $\epsilon/4$. It follows from the Chernoff bound the probability is at most (using $3k + 2 \leq 5k$)

$$\exp\left(-\left(\frac{11}{16}\right)^2 \cdot \frac{16k}{2}\right) = \exp\left(-\frac{121k}{32}\right) < \exp(-3) < 1/8.$$

Finally, it follows from a union bound that the algorithm rejects with probability at least $3/4$. $\square$

4.3 Proof of Lemma 4.1

We start with a lemma for the subroutine **WhereIsTheLiteral**, which is described in Figure 5.

LEMMA 4.4. *Assume that $g : \{0, 1\}^B \rightarrow \{0, 1\}$ is γ -close (with respect to the uniform distribution) to a literal x_i or $\bar{x}_i$ for some $i \in B$. If $i \in P$, then **WhereIsTheLiteral**(g, P, Q) returns a distinguishing pair of g for P with probability at least $1 - 4\gamma$; If $i \in Q$, then it returns a distinguishing pair of g for Q with probability at least $1 - 4\gamma$.*

PROOF. Let K be the set of strings $x \in \{0, 1\}^B$ such that $g(x)$ disagrees with the literal to which it is γ -close (so $|K| \leq \gamma \cdot 2^{|B|}$). We work on the case when $i \in Q$; the case when $i \in P$ is similar.

By the description of **WhereIsTheLiteral**, it returns a distinguishing pair for Q if

$$g(\mathbf{w} \circ \mathbf{z}) = g(\mathbf{w} \circ \mathbf{z}^{(P)}) \quad \text{and} \quad g(\mathbf{w}' \circ \mathbf{z}') \neq g(\mathbf{w}' \circ \mathbf{z}'^{(Q)}).$$

Note that this holds if all four strings fall outside of K and thus, the probability that it does not hold is at most the probability that at least one of these four strings falls inside K . The latter by a union bound is at most 4γ , since each of these four strings is drawn uniformly at random from $\{0, 1\}^B$ by itself. This finishes the proof of the lemma. $\square$

We are now ready to prove Lemma 4.1.

PROOF OF LEMMA 4.1. First, it is easy to verify that if the algorithm starts a round in good condition, then it ends it in good condition. This is because whenever a block is added to U , it is disjoint from other blocks, and we have found a distinguishing pair for it.

Next it follows directly from Lemma 4.4 and a union bound that, for any sequence of partitions P_j and Q_j of B_j picked on line 6, the probability that the for-loop correctly sets S_j to be the one that contains the special variable i_j for all $j \in [v]$ is at least (recalling that $\gamma = 1/(8k)$)

$$1 - 4\gamma \cdot v \geq 1 - 4\gamma \cdot k = 1/2.$$

Now, we can view the process equivalently as follows. First, we draw $\mathbf{x} \leftarrow \mathcal{D}$, $\mathbf{T} \leftarrow \overline{B_1 \cup \dots \cup B_v}$, and random partitions P_j, Q_j of each B_j . If we let $\mathbf{T}_j^*$ be the set in P_j, Q_j that does not contain the special variable, then $\mathbf{R}^* = \mathbf{T} \cup \mathbf{T}_1^* \cup \dots \cup \mathbf{T}_v^*$ is a set drawn uniformly at random from $\bar{I}$, where $I = \{i_j : j \in [v]\}$ consists of the special variables. Therefore, it follows from Lemma 3.2 that

Subroutine Literal(g)**Input:** Oracle access to a Boolean function g over $\{0, 1\}^C$ where C has a distinguishing pair.**Output:** “True” or disjoint nonempty subsets C', C^* of C and a distinguishing pair for each.

- (1) Repeat $\log k + 6$ times:
 - (2) If **UniformJunta**($g, 1, \gamma$) rejects, then
 - (3) Return the two disjoint blocks it has found and a distinguishing pair for each.
- (4) Let (x, y) be the distinguishing pair for C .
- (5) Repeat $\log k + 3$ times:
 - (6) Draw a random partition C', C^* of C and query $g(x^{(C')}), g(x^{(C^*)}), g(y^{(C')}), g(y^{(C^*)})$.
 - (7) If $g(x^{(C')}) = g(x^{(C^*)}) \neq g(x)$, then
 - (8) Return C', C^* and $(x, x^{(C')})$ and $(x, x^{(C^*)})$ as their distinguishing pairs.
 - (9) If $g(y^{(C')}) = g(y^{(C^*)}) \neq g(y)$, then
 - (10) Return C', C^* and $(y, y^{(C')})$ and $(y, y^{(C^*)})$ as their distinguishing pairs.
- (11) Return “true.”

Fig. 6. Description of the subroutine **Literal**.

$f(\mathbf{x}) \neq f(\mathbf{x}^{(\mathbf{R}^*)})$ with probability at least $\epsilon/2$. Since with probability at least $1/2$, the set $\mathbf{R}$ on line 11 agrees with $\mathbf{R}^*$, we have that the algorithm reaches line 12 with $f(\mathbf{x}) \neq f(\mathbf{y})$ with probability at least $\epsilon/4$. Given this, the lemma is immediate by inspection of lines 12-17 of the algorithm. $\square$

4.4 Proof of Lemma 4.2

First it follows from the description of the subroutine **Literal**(g) in Figure 6 that it either returns “true” or a pair of nonempty disjoint subsets C', C^* of C and a distinguishing pair of g for each of them (see Theorem 2.3). Next, let $C \in V$ be the block picked in line 20. If g is γ -close to a literal, then it is easy to verify that one of the two events described in Lemma 4.2 must hold (using the property of **Literal**(g) above). So, we focus on the other two cases in the rest of the proof: g is γ -far from 1-juntas or g is γ -close to a (all-1 or all-0) constant function. In both cases, we show below that the second event happens with high probability.

When g is γ -far from 1-juntas under the uniform distribution, we have that one of the $\log k + 6$ calls to **UniformJunta** in **Literal** rejects with probability at least

$$1 - (1/3)^{\log k + 6} > 1 - 1/(64k).$$

The second event in Lemma 4.2 occurs when this happens.

When g is γ -close to a constant function (say the all-1 function), we have that either string x or y in the distinguishing pair for C disagrees with the function (say $g(x) = 0$, since $g(x) \neq g(y)$). Let K be the set of strings in $\{0, 1\}^C$ that disagree with the all-1 function. Then line 7 of **Literal**(g) does not hold only when one of $x^{(C')}$ or $x^{(C^*)}$ lies in K . As both strings are distributed uniformly over $\{0, 1\}^C$ by themselves, this happens with probability at most 2γ by a union bound. Therefore, the probability that line 7 holds at least once is at least

$$1 - (2\gamma)^{\log k + 3} = 1 - (1/(4k))^{\log k + 3} > 1 - (1/4)^{\log k + 3} = 1 - 1/(64k^2).$$

As a result, the second event in Lemma 4.2 occurs with probability at least $1 - 1/(64k^2)$.

This finishes the proof of Lemma 4.2.

5 PROOF OF THEOREM 1.2: AN $\Omega(2^{K/3})$ -QUERY NON-ADAPTIVE LOWER BOUND

In this section, we prove the $\Omega(2^{k/3})$ lower bound for the non-adaptive distribution-free testing of k -juntas that was stated as Theorem 1.2. We start with some notation. Given a sequence $Y = (y^i : i \in [q])$ of q strings in $\{0, 1\}^n$ and a Boolean function $\phi : \{0, 1\}^n \rightarrow \{0, 1\}$, we write $\phi(Y)$ to denote

the q -bit string α with $\alpha_i = \phi(y^i)$ for each $i \in [q]$. We also write $Y = (y^i : i \in [q]) \leftarrow \mathcal{D}^q$ to denote a sequence of q independent draws from the same probability distribution $\mathcal{D}$.

Let k and n be two positive integers that satisfy $k \leq n/200$. We may further assume that k is at least some absolute constant C (to be specified later), since otherwise, the claimed $\Omega(2^{k/3})$ lower bound on query complexity holds trivially due to the constant hidden behind the Ω . Let $q = 2^{k/3}$. For convenience, we refer to an algorithm as a q -query algorithm if it makes q sample queries and q black-box queries *each*. Such algorithms are clearly at least as powerful as those that make q queries in total. Our goal is then to show that there exists no q -query non-adaptive (randomized) algorithm for the distribution-free testing of k -juntas over Boolean functions of n variables, even when the distance parameter ϵ is $1/3$.

By Yao's minimax principle, we focus on q -query non-adaptive *deterministic* algorithms. Such an algorithm A (which consists of two deterministic maps A_1 and A_2 as discussed below) works as follows. Upon an input pair $(\phi, \mathcal{D})$, where $\phi : \{0, 1\}^n \rightarrow \{0, 1\}$ and $\mathcal{D}$ is a probability distribution over $\{0, 1\}^n$, the algorithm receives in the first phase a sequence $Y = (y^i : i \in [q])$ of q strings (which should be thought of as samples from $\mathcal{D}$) and a binary string $\alpha = \phi(Y)$ of length q . In the second phase, the algorithm A uses the first map A_1 to obtain a sequence of q strings $Z = (z^i : i \in [q]) = A_1(Y, \alpha)$ and feeds them to the black-box oracle. Once the query results $\beta = \phi(Z)$ are back, $A_2(Y, \alpha, \beta)$ returns either 0 or 1 (notice that we do not need to include Z as an input of A_2 , since it is determined by Y and α), in which cases the algorithm A either rejects or accepts, respectively. A randomized algorithm T works similarly and consists of two similar maps T_1 and T_2 but both are randomized.

Given the description above, unlike typical deterministic algorithms, whether A accepts or not depends on not only $(\phi, \mathcal{D})$ but also the sample strings $Y \leftarrow \mathcal{D}^q$ it draws. Formally, we have

$$\Pr[A \text{ accepts } (\phi, \mathcal{D})] = \Pr_{Y \leftarrow \mathcal{D}^q} [A_2(Y, \phi(Y), \phi(A_1(Y, \phi(Y)))) = 1].$$

The plan of the rest of the section is as follows. We define in Section 5.1 a pair of probability distributions $\mathcal{YES}$ and $\mathcal{NO}$ over pairs $(\phi, \mathcal{D})$, where ϕ is a Boolean function over n variables and $\mathcal{D}$ is a distribution over $\{0, 1\}^n$. For clarity, we use $(f, \mathcal{D})$ to denote pairs in the support of $\mathcal{YES}$ and $(g, \mathcal{D})$ to denote pairs in the support of $\mathcal{NO}$. We show that (1) Every $(f, \mathcal{D})$ in the support of $\mathcal{YES}$ satisfies that f is a k -junta (Lemma 5.2); (2) With probability $1 - o_k(1)$, $(g, \mathcal{D}) \leftarrow \mathcal{NO}$ satisfies that g is $1/3$ -far from every k -junta with respect to $\mathcal{D}$ (Lemma 5.3). To obtain Theorem 1.2, it suffices to prove the following main technical lemma, which informally says that any q -query non-adaptive deterministic algorithm must behave similarly when it is run on $(f, \mathcal{D}) \leftarrow \mathcal{YES}$ versus $(g, \mathcal{D}) \leftarrow \mathcal{NO}$:

LEMMA 5.1. *Any q -query deterministic algorithm A satisfies*

$$| \mathbb{E}_{(f, \mathcal{D}) \leftarrow \mathcal{YES}} [\Pr [A \text{ accepts } (f, \mathcal{D})]] - \mathbb{E}_{(g, \mathcal{D}) \leftarrow \mathcal{NO}} [\Pr [A \text{ accepts } (g, \mathcal{D})]] | \leq 1/4. \quad (3)$$

PROOF OF THEOREM 1.2 ASSUMING LEMMAS 5.1, 5.2, AND 5.3. Assume for a contradiction that there exists a q -query non-adaptive randomized algorithm T for the distribution-free testing of k -juntas over n -variable Boolean functions when $\epsilon = 1/3$. Then it follows from Lemmas 5.2 and 5.3 that

$$\mathbb{E}_{(f, \mathcal{D}) \leftarrow \mathcal{YES}} [\Pr [T \text{ accepts } (f, \mathcal{D})]] - \mathbb{E}_{(g, \mathcal{D}) \leftarrow \mathcal{NO}} [\Pr [T \text{ accepts } (g, \mathcal{D})]] \geq 1/3 - o_k(1),$$

since the first expectation is at least $2/3$ and the second is at most

$$1/3(1 - o_k(1)) + o_k(1) \leq 1/3 + o_k(1).$$

As T is a probability distribution over deterministic algorithms, there must exist a q -query non-adaptive deterministic algorithm A that satisfies

$$\mathbb{E}_{(f, \mathcal{D}) \leftarrow \mathcal{YES}}[\Pr [A \text{ accepts } (f, \mathcal{D})]] - \mathbb{E}_{(g, \mathcal{D}) \leftarrow \mathcal{NO}}[\Pr [A \text{ accepts } (g, \mathcal{D})]] \geq 1/3 - o_k(1),$$

a contradiction with Lemma 5.1 when k is sufficiently large. $\square$

5.1 The $\mathcal{YES}$ and $\mathcal{NO}$ Distributions

Given $J \subseteq [n]$, we partition $\{0, 1\}^n$ into *sections* (with respect to J) where the z -section, $z \in \{0, 1\}^J$, consists of those $x \in \{0, 1\}^n$ that have $x_J = z$. We write $\mathcal{JUNT} \mathcal{A}_J$ to denote the uniform distribution over all juntas over J . More precisely, a Boolean function $\mathbf{h} : \{0, 1\}^n \rightarrow \{0, 1\}$ drawn from $\mathcal{JUNT} \mathcal{A}_J$ is generated as follows: For each $z \in \{0, 1\}^J$, a bit $\mathbf{b}(z)$ is chosen independently and uniformly at random, and for each $x \in \{0, 1\}^n$ the value of $\mathbf{h}(x)$ is set to $\mathbf{b}(x_J)$. Let

$$m := 36 \cdot 2^k \ln n.$$

We start with $\mathcal{YES}$. A pair $(f, \mathcal{D})$ drawn from $\mathcal{YES}$ is generated as follows:

- (1) First, we draw independently a subset J of $[n]$ of size k uniformly at random and a subset S of $\{0, 1\}^n$ of size m uniformly at random.
- (2) Next, we draw $f \leftarrow \mathcal{JUNT} \mathcal{A}_J$ and set $\mathcal{D}$ to be the uniform distribution over S .

For technical reasons that will become clear in Section 5.2 we use $\mathcal{YES}^*$ to denote the probability distribution supported over triples $(f, \mathcal{D}, J)$, with $(f, \mathcal{D}, J) \leftarrow \mathcal{YES}^*$ being generated by the same two steps above (so, the only difference is that we include J in elements of $\mathcal{YES}^*$).

The following observation is straightforward from the definition of $\mathcal{YES}$.

LEMMA 5.2. *The function f is a k -junta for every pair $(f, \mathcal{D})$ in the support of $\mathcal{YES}$.*

We now describe $\mathcal{NO}$. A pair $(g, \mathcal{D})$ drawn from $\mathcal{NO}$ is generated as follows:

- (1) We draw J and S in the same way as the first step of $\mathcal{YES}$.
- (2) Next, we draw $\mathbf{h} \leftarrow \mathcal{JUNT} \mathcal{A}_J$ and a map $\gamma : S \rightarrow \{0, 1\}$ uniformly at random by choosing a bit independently and uniformly at random for each string in S . We usually refer to $\mathbf{h}$ as the “background junta.”
- (3) The distribution $\mathcal{D}$ is set to be the uniform distribution over S , which is the same as $\mathcal{YES}$. The function $g : \{0, 1\}^n \rightarrow \{0, 1\}$ is defined using $\mathbf{h}$, S and γ as follows:
 - (a) For each string $y \in S$, set $g(y) = \gamma(y)$;
 - (b) For each string $x \notin S$, if there exists no $y \in S$ with $y_J = x_J$ and $d(x, y) \leq 0.4n$, set $g(x) = \mathbf{h}(x)$; otherwise, we set $g(x) = 1$ if there exists such a $y \in S$ with $\gamma(y) = 1$, and set $g(x) = 0$ if every such $y \in S$ has $\gamma(y) = 0$. (The choice of the tie-breaking rule here is not important; we just pick one to make sure that g is well defined in all cases.)

Similarly, we let $\mathcal{NO}^*$ denote the distribution supported on triples $(g, \mathcal{D}, J)$ as generated above.

See Figure 7 for an illustration of a function drawn from $\mathcal{NO}$. To gain some intuition, we first note that about half of the strings $z \in S$ have $g(z)$ disagree with the value of the background junta on the section it lies in. With such a string z in hand (from one of the samples received in the first phase), an algorithm may attempt to find a string w that lies in the same section as z but satisfies $g(z) \neq g(w)$. If such a string is found, then the algorithm knows for sure that $(g, \mathcal{D})$ is from the $\mathcal{NO}$ distribution. However, finding such a w is not easy, because one must flip more than $0.4n$ bits of z , but without knowing the variables in J it is hard to keep w in the same section as z after flipping this many bits.

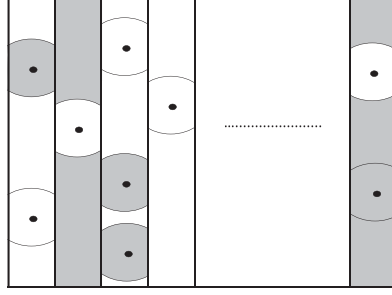


Fig. 7. A schematic depiction of how $\{0,1\}^n$ is labeled by a function g from $\mathcal{NO}$. The domain $\{0,1\}^n$ is partitioned into 2^k sections corresponding to different settings of the variables in J ; each section is a vertical strip in the figure. Shaded regions correspond to strings where g evaluates to 1 and unshaded regions to strings where g evaluates to 0. Each string in S is a black dot and the value of g on each such string is chosen uniformly at random. Since in this figure the truncated circles are disjoint, the tie-breaking rule does not come into effect, and for each $z \in S$, all strings in its section within distance at most $0.4n$ (the points in the truncated circle around z) have the same value as z . The value of g on other points is determined by the background junta h , which assigns a uniform random bit to each section.

Next, we prove that with high probability, $(g, \mathcal{D}) \leftarrow \mathcal{NO}$ satisfies that g is $1/3$ -far from every k -junta with respect to $\mathcal{D}$:

LEMMA 5.3. *With probability at least $1 - o_k(1)$, $(g, \mathcal{D}) \leftarrow \mathcal{NO}$ is such that g is $1/3$ -far from every k -junta with respect to the distribution $\mathcal{D}$.*

PROOF. Fix a k -junta h , i.e., any set $I \subset [n]$ with $|I| = k$ and any 2^k -bit truth table over variables in I . We have that $\text{dist}_{\mathcal{D}}(g, h)$ is precisely the fraction of strings $z \in S$ such that $\gamma(z) \neq h(z)$. Since each bit $\gamma(z)$ is drawn independently and uniformly at random, we have that

$$\Pr_{(g, \mathcal{D}) \leftarrow \mathcal{NO}} [\text{dist}_{\mathcal{D}}(g, h) \leq 1/3] = \Pr_{j \leftarrow \text{Bin}(m, 1/2)} [j \leq m/3],$$

which, recalling that $m = 36 \cdot 2^k \ln n$, by a standard Chernoff bound is at most $e^{-m/36} = n^{-2^k}$. The result follows by a union bound over all (at most)

$$\binom{n}{k} \cdot 2^{2^k} \leq n^k \cdot 2^{2^k} = o_k(n^{2^k})$$

possible k -juntas h over n variables. This finishes the proof of the lemma. $\square$

Given Lemmas 5.2 and 5.3, to prove Theorem 1.2 it remains only to prove Lemma 5.1.

5.2 Proof of Lemma 5.1

The following definitions will be useful. Let $Y = (y^i : i \in [q])$ be a sequence of q strings in $\{0,1\}^n$, α be a q -bit string, and $J \subset [n]$ be a set of variables of size k . We say that (Y, α, J) is *consistent* if

$$\alpha_i = \alpha_j \quad \text{for all } i, j \in [q] \text{ with } y_J^i = y_J^j. \quad (4)$$

Given a consistent triple (Y, α, J) , we write $\mathcal{JUNTA}_{Y, \alpha, J}$ to denote the uniform distribution over all juntas h over J that are consistent with (Y, α) . More precisely, a draw of $h \leftarrow \mathcal{JUNTA}_{Y, \alpha, J}$ is generated as follows: For each $z \in \{0,1\}^J$, if there exists a y^i such that $y_J^i = z$, then $h(x)$ is set to α_i for all $x \in \{0,1\}^n$ with $x_J = z$; if no such y^i exists, then a uniform random bit $b(z)$ is chosen independently and $h(x)$ is set to $b(z)$ for all x with $x_J = z$.

To prove Lemma 5.1, we first derive from A a new *randomized* algorithm A' that works on triples $(\phi, \mathcal{D}, J)$ from the support of either $\mathcal{YES}^*$ or $\mathcal{NO}^*$. Again for clarity we use ϕ to denote a function from the support of $\mathcal{YES}/\mathcal{YES}^*$ or $\mathcal{NO}/\mathcal{NO}^*$, f to denote a function from $\mathcal{YES}/\mathcal{YES}^*$ and g to denote a function from $\mathcal{NO}/\mathcal{NO}^*$.

In addition to being randomized, A' differs from A in two important ways:

- (1) Like A , A' receives samples $Y \leftarrow \mathcal{D}^q$ and $\phi(Y)$, but unlike A , A' also receives J for free.
- (2) Unlike A , A' does not make any black-box queries but simply runs on the triple $(Y, \phi(Y), J)$ it receives at the beginning. So formally A' is a randomized algorithm that runs on triples (Y, α, J) , where $Y = (y^i : i \in [q])$ is a sequence of q strings, α is a q -bit string, and $J \subset [n]$ is a set of variables of size k , and outputs “accept” or “reject.”

A detailed description of the randomized algorithm A' running on (Y, α, J) is as follows:

- (1) First, if (Y, α, J) is not consistent, then A' immediately halts and rejects (simply because this can never occur if (Y, α, J) is obtained from a triple $(f, \mathcal{D}, J)$ in the support of $\mathcal{YES}^*$). Otherwise, A' applies A_1 on (Y, α) to obtain a sequence $Z = (Z^i : i \in [q])$ of q strings.
- (2) Next, A' draws $\mathbf{h}' \leftarrow \mathcal{JUNT} \mathcal{A}_{Y, \alpha, J}$. (This is the only part of A' that is randomized.)
- (3) Finally, A' runs $A_2(Y, \alpha, \mathbf{h}'(Z))$ and outputs the same result (accept or reject).

From the description of A' above, whether it accepts a triple $(\phi, \mathcal{D}, J)$ or not depends on both the randomness of Y and $\mathbf{h}'$. Formally, we have

$$\Pr [A' \text{ accepts } (\phi, \mathcal{D}, J)] = \Pr_{Y, \mathbf{h}'} [(Y, \alpha, J) \text{ is consistent and } A_2(Y, \alpha, \mathbf{h}'(Z)) = 1].$$

Lemma 5.1 follows immediately from the following three lemmas (note that the marginal distribution of $(f, \mathcal{D})$ in $\mathcal{YES}^*$ (or $(g, \mathcal{D})$ in $\mathcal{NO}^*$) is the same as $\mathcal{YES}$ (or $\mathcal{NO}$)). In all three lemmas, we assume that A is a q -query non-adaptive deterministic algorithm while A' is the randomized algorithm derived from A as described above.

LEMMA 5.4 (A' BEHAVES SIMILARLY ON $\mathcal{YES}^*$ AND $\mathcal{NO}^*$). *We have*

$$| \mathbb{E}_{(f, \mathcal{D}, J) \leftarrow \mathcal{YES}^*} [\Pr [A' \text{ accepts } (f, \mathcal{D}, J)]] - \mathbb{E}_{(g, \mathcal{D}, J) \leftarrow \mathcal{NO}^*} [\Pr [A' \text{ accepts } (g, \mathcal{D}, J)]] | \leq 1/8.$$

LEMMA 5.5 (A AND A' BEHAVE IDENTICALLY ON $\mathcal{YES}$ AND $\mathcal{YES}^*$, RESPECTIVELY). *We have*

$$\mathbb{E}_{(f, \mathcal{D}, J) \leftarrow \mathcal{YES}^*} [\Pr [A \text{ accepts } (f, \mathcal{D})]] = \mathbb{E}_{(f, \mathcal{D}, J) \leftarrow \mathcal{YES}^*} [\Pr [A' \text{ accepts } (f, \mathcal{D}, J)]]. \quad (5)$$

LEMMA 5.6 (A AND A' BEHAVE SIMILARLY ON $\mathcal{NO}$ AND $\mathcal{NO}^*$, RESPECTIVELY). *We have*

$$| \mathbb{E}_{(g, \mathcal{D}, J) \leftarrow \mathcal{NO}^*} [\Pr [A \text{ accepts } (g, \mathcal{D})]] - \mathbb{E}_{(g, \mathcal{D}, J) \leftarrow \mathcal{NO}^*} [\Pr [A' \text{ accepts } (g, \mathcal{D}, J)]] | \leq 1/8.$$

We start with the proof of Lemma 5.4, which says that a limited algorithm such as A' cannot effectively distinguish between a draw from $\mathcal{YES}^*$ versus $\mathcal{NO}^*$:

PROOF OF LEMMA 5.4. Since A' runs on (Y, α, J) , it suffices to show that the distributions of (Y, α, J) induced from $\mathcal{YES}^*$ and $\mathcal{NO}^*$ have small total variation distance. For this purpose, we first note that the distributions of (Y, J) induced from $\mathcal{YES}^*$ and $\mathcal{NO}^*$ are identical: In both cases, Y and J are independent; J is a random subset of $[n]$ of size k ; Y is obtained by first sampling a subset S of $\{0, 1\}^n$ of size m and then drawing a sequence of q strings from S with replacement.

Fix a pair (Y, J) in the support of (Y, J) . We say Y is *scattered* by J if $y_j^i \neq y_j^j$ for all $i \neq j \in [q]$. In particular, this implies that no string appears more than once in Y . The following claim, whose proof we defer, shows that Y is scattered by J with high probability.

CLAIM 5.7. *We have that Y is scattered by J with probability at least $1 - O(2^{-k/3})$.*

Fix any (Y, J) in the support of (Y, J) such that Y is scattered by J . We claim that the distributions of α conditioning on $(Y, J) = (Y, J)$ in the $\mathcal{YES}^*$ case and the $\mathcal{NO}^*$ case are identical, from which it follows that the total variation distance between the distributions of (Y, α, J) in the two cases is at most $O(2^{-k/3}) \leq 1/8$ when k is sufficiently large. Indeed, α is uniform over strings of length q in both cases. This is trivial for $\mathcal{NO}^*$. For $\mathcal{YES}^*$ note that α is determined by the random k -junta $f \leftarrow \mathcal{JUNT}\mathcal{A}_J$; the claim follows from the assumption that Y is scattered by J .

PROOF OF CLAIM 5.7. We fix J and show that Y is scattered by J with high probability. As strings of Y are drawn one by one, the probability of y^i colliding with one of the previous samples is at most $(i-1)/m \leq q/m$. By a union bound, all strings in Y are distinct with probability at least

$$1 - q \cdot (q/m) = 1 - q^2/m = 1 - O(2^{-k/3}).$$

Conditioning on this event, $Y = (y^i)$ is distributed precisely as a uniform random sequence from $\{0, 1\}^n$ with no repetition and thus each pair (y^i, y^j) is distributed uniformly over pairs of distinct strings in $\{0, 1\}^n$. As a result, we have

$$\Pr[y_J^i = y_J^j] = \frac{2^{n-k} - 1}{2^n - 1} \leq \frac{1}{2^k}.$$

By a union bound over $\binom{q}{2}$ pairs, we have that the probability of Y being scattered by J is at least

$$(1 - O(2^{-k/3})) \cdot \left(1 - \binom{q}{2} \cdot 2^{-k}\right) \geq 1 - O(2^{-k/3}).$$

This finishes the proof of the claim. $\square$

Next, we prove Lemma 5.5.

PROOF OF LEMMA 5.5. The first expectation in Equation (5) is equal to the probability that

$$A_2(Y, \alpha, f(A_1(Y, \alpha))) = 1,$$

where $(f, \mathcal{D}, J) \leftarrow \mathcal{YES}^*$, $Y \leftarrow \mathcal{D}^q$ and $\alpha = f(Y)$. For the second expectation, since the triple on which we run A' is always consistent, we can rewrite it as the probability that

$$A_2(Y, \alpha, h'(A_1(Y, \alpha))) = 1,$$

where $(f, \mathcal{D}, J) \leftarrow \mathcal{YES}^*$, $Y \leftarrow \mathcal{D}^q$, $\alpha = f(Y)$ and $h' \leftarrow \mathcal{JUNT}\mathcal{A}_{Y, \alpha, J}$.

To show that these two probabilities are equal, we first note that the distributions of (Y, α, J) are identical. Fixing any triple (Y, α, J) in the support of (Y, α, J) , which must be consistent, we claim that the distribution of f conditioning on $(Y, \alpha, J) = (Y, \alpha, J)$ is exactly $\mathcal{JUNT}\mathcal{A}_{Y, \alpha, J}$. This is because, for each $z \in \{0, 1\}^J$, if $y_J^j = z$ for some y^j in Y , then we have $f(x) = \alpha_i$ for all strings x with $x_J = z$; otherwise, we have $f(x) = b(z)$ for all x with $x_J = z$, where $b(z)$ is an independent and uniform bit. This is the same as how $h' \leftarrow \mathcal{JUNT}\mathcal{A}_{Y, \alpha, J}$ is generated. It follows directly from this claim that the two probabilities are the same. This finishes the proof of the lemma. $\square$

Finally, we prove Lemma 5.6, the most difficult among the three lemmas:

PROOF OF LEMMA 5.6. Similar to the proof of Lemma 5.5, the first expectation is the probability of

$$A_2(Y, \alpha, g(A_1(Y, \alpha))) = 1,$$

where $(g, \mathcal{D}, J) \leftarrow \mathcal{NO}^*$ and $\alpha = g(Y)$, while the second expectation is the probability of

$$(Y, \alpha, J) \text{ is consistent and } A_2(Y, \alpha, h'(A_1(Y, \alpha))) = 1,$$

where $(g, \mathcal{D}, J) \leftarrow \mathcal{NO}^*$, $Y \leftarrow \mathcal{D}^q$, $\alpha = g(Y)$, and $h' \leftarrow \mathcal{JUNT}\mathcal{A}_{Y, \alpha, J}$. We note that the distributions of $(Y, \alpha, J, \mathcal{D})$ in the two cases are identical.

The following definition is crucial. We say a tuple $(Y, \alpha, J, \mathcal{D})$ in the support of $(Y, \alpha, J, \mathcal{D})$ is *good* if it satisfies the following three conditions (S below is the support of $\mathcal{D}$):

E_0 : Y is scattered by J .

E_1 : Let $Z = A_1(Y, \alpha)$. Then every z in Z and every x in $S \setminus Y$ have $d(x, z) > 0.4n$. (In $S \setminus Y$, we abuse notation and use Y as a set that contains all strings in the sequence Y .)

E_2 : If a string z in Z satisfies $z_J = y_J$ for some y in Y , then we have $d(y, z) \leq 0.4n$. $\square$

We delay the proof of the following claim to the end.

CLAIM 5.8. *We have that $(Y, \alpha, J, \mathcal{D})$ is good with probability at least $7/8$.*

Fix any good $(Y, \alpha, J, \mathcal{D})$ in the support and let $Z = A_1(Y, \alpha)$. We finish the proof by showing that the distribution of $\mathbf{g}(Z)$, a binary string of length q , conditioning on $(Y, \alpha, J, \mathcal{D}) = (Y, \alpha, J, \mathcal{D})$ is the same as that of $\mathbf{h}'(Z)$ with $\mathbf{h}' \leftarrow \mathcal{JUNT} \mathcal{A}_{Y, \alpha, J}$. This combined with Claim 5.8 implies that the difference of the two probabilities has absolute value at most $1/8$.

To see this is the case, we partition strings of Z into Z_w , where each Z_w is a nonempty set that contains all z in Z with $z_J = w \in \{0, 1\}^J$. For each Z_w , we consider the following two cases:

- (1) If there exists no string y in Y with $y_J = w$, then by E_1 strings in Z_w are all far from strings of S (i.e., the support of $\mathcal{D}$) in this section and thus, $\mathbf{g}(z) = \mathbf{b}(w)$ for some independent and uniform bit $\mathbf{b}(w)$, for all strings $z \in Z_w$.
- (2) If there exists a y in Y with $y_J = w$ (which must be unique by E_0), say y^i , then by E_1 and E_2 strings in Z_w are all close to y and far from other strings of S in this section. As a result, we have $\mathbf{g}(z) = \alpha_i$ for all strings $z \in Z_w$.

So the conditional distribution of $\mathbf{g}(Z)$ is identical to that of $\mathbf{h}'(Z)$ with $\mathbf{h}' \leftarrow \mathcal{JUNT} \mathcal{A}_{Y, \alpha, J}$. This finishes the proof of the lemma.

PROOF OF CLAIM 5.8. We bound the probabilities of $(Y, \alpha, J, \mathcal{D})$ violating each of the three conditions E_0 , E_1 and E_2 and apply a union bound. By Claim 5.7, E_0 is violated with probability $O(2^{-k/3})$.

For E_1 , we fix a pair (Y, α) in the support and let $\ell \leq q$ be the number of distinct strings in Y and $Z = A_1(Y, \alpha)$. Conditioning on $Y = Y$, $S \setminus Y$ is a uniformly random subset of $\{0, 1\}^n \setminus Y$ of size $m - \ell$. Instead of working with $S \setminus Y$, we let T denote a set obtained by making $m - \ell$ draws from $\{0, 1\}^n$ uniformly at random (with replacements).

On the one hand, the total variation distance between $S \setminus Y$ and T is exactly the probability that either (1) $T \cap Y$ is nonempty or (2) $|T| < m - \ell$. By two union bounds, (1) happens with probability at most $(m - \ell) \cdot (\ell/2^n) \leq mq/2^n$ and (2) happens with probability at most $(m/2^n) \cdot m$. As a result, the total variation distance is at most $(mq + m^2)/2^n$. On the other hand, the probability that one of the strings of T has distance at most $0.4n$ with one of the strings of Z is at most $mq \cdot \exp(-n/100)$ by a Chernoff bound followed by a union bound. Thus, the probability of violating E_1 is at most (using the assumption that $k \leq n/200$)

$$(mq + m^2)/2^n + mq \cdot \exp(-n/100) = O(2^{-n/300} \ln n).$$

For E_2 , we fix a pair (Y, α) in the support and let $Z = A_2(Y, \alpha)$. Because J is independent from (Y, α) , it remains a subset of $[n]$ of size k drawn uniformly at random. For each pair (y, z) with y from Y and z from Z that satisfy $d(y, z) > 0.4n$, the probability of $y_J = z_J$ is at most

$$\frac{\binom{0.6n}{k}}{\binom{n}{k}} \leq (0.6)^k.$$

Since there are at most q^2 many such pairs, it follows from a union bound that the probability of violating E_2 is at most $q^2 \cdot (0.6)^k \leq e^{-0.07k}$.

Finally, the lemma follows from a union bound when k (and thus, n) is sufficiently large. $\square$

REFERENCES

- N. Alon, T. Kaufman, M. Krivelevich, S. Litsyn, and D. Ron. 2005. Testing Reed-Muller codes. *IEEE Trans. Info. Theory* 51, 11 (2005), 4032–4039.
- Noga Alon and Amit Weinstein. 2012. Local correction of juntas. *Inf. Process. Lett.* 112, 6 (2012), 223–226.
- Roksana Baleshzar, Meiram Murzabulatov, Ramesh Krishnan S. Pallavoor, and Sofya Raskhodnikova. 2016. Testing unateness of real-valued functions. *CoRR abs/1608.07652*.
- Aleksandrs Belovs and Eric Blais. 2016. A polynomial lower bound for testing monotonicity. In *Proceedings of the 48th Annual ACM Symposium on Theory of Computing (STOC'16)*. ACM, New York, NY, 1021–1032. DOI: <https://doi.org/10.1145/2897518.2897567>
- Arnab Bhattacharyya, Swastik Kopparty, Grant Schoenebeck, Madhu Sudan, and David Zuckerman. 2010. Optimal testing of Reed-Muller codes. In *Proceedings of the 51st Annual IEEE Symposium on Foundations of Computer Science (FOCS'10)*. 488–497.
- Eric Blais. 2008. Improved bounds for testing juntas. In *Proceedings of the International Conference on Randomization and Computation (RANDOM'08)*. 317–330.
- Eric Blais. 2009. Testing juntas nearly optimally. In *Proceedings of the 41st Annual ACM Symposium on Theory of Computing (STOC'09)*. 151–158.
- Eric Blais, Joshua Brody, and Kevin Matulef. 2011. Property testing lower bounds via communication complexity. In *Proceedings of the Computational Complexity Conference (CCC'11)*. 210–220.
- Eric Blais and Daniel M. Kane. 2012. Tight bounds for testing k -linearity. In *Proceedings of the International Conference on Randomization and Computation (RANDOM'12)*. 435–446.
- M. Blum, M. Luby, and R. Rubinfeld. 1993. Self-testing/correcting with applications to numerical problems. *J. Comput. Syst. Sci.* 47 (1993), 549–595.
- Harry Buhrman, David García-Soriano, Arie Matsliah, and Ronald de Wolf. 2013. The non-adaptive query complexity of testing K -parities. *Chicago J. Theoret. Comput. Sci.* 6 (2013), 1–11.
- Deeparnab Chakrabarty and C. Seshadhri. 2013. An $o(n)$ monotonicity tester for Boolean functions over the hypercube. In *Proceedings of the 45th ACM Symposium on Theory of Computing*. 411–418.
- Deeparnab Chakrabarty and C. Seshadhri. 2016. A $\tilde{O}(n)$ non-adaptive tester for unateness. *CoRR abs/1608.06980*.
- Xi Chen, Anindya De, Rocco A. Servedio, and Li-Yang Tan. 2015. Boolean function monotonicity testing requires (almost) $n^{1/2}$ non-adaptive queries. In *Proceedings of the 47th ACM Symposium on Theory of Computing*. 519–528.
- X. Chen, R. A. Servedio, and L.-Y. Tan. 2014. New algorithms and lower bounds for monotonicity testing. In *Proceedings of the 55th IEEE Symposium on Foundations of Computer Science (FOCS'14)*. 286–295.
- Xi Chen, Rocco A. Servedio, Li-Yang Tan, Erik Waingarten, and Jinyu Xie. 2017a. Settling the query complexity of non-adaptive junta testing. In *Proceedings of the 32nd Computational Complexity Conference (CCC'17)*. 26:1–26:19.
- Xi Chen, Erik Waingarten, and Jinyu Xie. 2017b. Beyond talagrand functions: New lower bounds for testing monotonicity and unateness. In *Proceedings of the 49th Annual ACM SIGACT Symposium on Theory of Computing (STOC'17)*. 523–536.
- Xi Chen, Erik Waingarten, and Jinyu Xie. 2017c. Boolean unateness testing with $\tilde{O}(n^{3/4})$ adaptive queries. In *Proceedings of the 58th Annual IEEE Symposium on Foundations of Computer Science (FOCS'17)*. 868–879.
- Xi Chen and Jinyu Xie. 2016. Tight bounds for the distribution-free testing of monotone conjunctions. In *Proceedings of the 27th Annual ACM-SIAM Symposium on Discrete Algorithms (SODA'16)*. 54–71.
- H. Chockler and D. Gutfreund. 2004. A lower bound for testing juntas. *Inform. Process. Lett.* 90, 6 (2004), 301–305.
- I. Diakonikolas, H. Lee, K. Matulef, K. Onak, R. Rubinfeld, R. Servedio, and A. Wan. 2007. Testing for concise representations. In *Proceedings of the 48th Annual Symposium on Computer Science (FOCS'07)*. 549–558.
- E. Dolev and D. Ron. 2011. Distribution-free testing for monomials with a sublinear number of queries. *Theory Comput.* 7, 1 (2011), 155–176.
- E. Fischer, G. Kindler, D. Ron, S. Safra, and A. Samorodnitsky. 2004. Testing juntas. *J. Comput. Syst. Sci.* 68, 4 (2004), 753–787.
- E. Fischer, E. Lehman, I. Newman, S. Raskhodnikova, R. Rubinfeld, and A. Samorodnitsky. 2002. Monotonicity testing over general poset domains. In *Proceedings of the 34th Annual ACM Symposium on the Theory of Computing*. 474–483.
- Dana Glasner and Rocco A. Servedio. 2009. Distribution-free testing lower bound for basic boolean functions. *Theory Comput.* 5, 1 (2009), 191–216.
- O. Goldreich (Ed.). 2010. *Property Testing: Current Research and Surveys*. Springer.
- O. Goldreich, S. Goldwasser, E. Lehman, D. Ron, and A. Samordnitsky. 2000. Testing monotonicity. *Combinatorica* 20, 3 (2000), 301–337.

- O. Goldreich, S. Goldwasser, and D. Ron. 1998. Property testing and its connection to learning and approximation. *J. ACM* 45 (1998), 653–750.
- P. Gopalan, R. O'Donnell, R. Servedio, A. Shpilka, and K. Wimmer. 2011. Testing fourier dimensionality and sparsity. *SIAM J. Comput.* 40, 4 (2011), 1075–1100.
- S. Halevy and E. Kushilevitz. 2007. Distribution-free property testing. *SIAM J. Comput.* 37, 4 (2007), 1107–1138.
- Subhash Khot, Dor Minzer, and Muli Safra. 2015. On monotonicity testing and Boolean isoperimetric type theorems. In *Proceedings of the 56th Annual Symposium on Foundations of Computer Science*. 52–58.
- Subhash Khot and Igor Shinkar. 2016. An $O(n)$ queries adaptive tester for unateness. In *Approximation, Randomization, and Combinatorial Optimization Algorithms and Techniques*. Schloss Dagstuhl–Leibniz-Zentrum fuer Informatik.
- K. Matulef, R. O'Donnell, R. Rubinfeld, and R. Servedio. 2010. Testing halfspaces. *SIAM J. Comput.* 39, 5 (2010), 2004–2047.
- Kevin Matulef, Ryan O'Donnell, Ronitt Rubinfeld, and Rocco A. Servedio. 2009. Testing ± 1 -weight halfspace. In *Proceedings of the International Conference on Approximation Algorithms for Combinatorial Optimization Problems and the International Conference on Randomization and Computation (APPROX-RANDOM'09)*. 646–657.
- M. Parnas, D. Ron, and A. Samorodnitsky. 2002. Testing basic boolean formulae. *SIAM J. Disc. Math.* 16 (2002), 20–46. Retrieved from citeseer.ifi.unizh.ch/parnas02testing.html.
- D. Ron. 2008. Property testing: A learning theory perspective. *Found. Trends Mach. Learn.* 1, 3 (2008), 307–402.
- D. Ron. 2010. Algorithmic and analysis techniques in property testing. *Found. Trends Theoret. Comput. Sci.* 5, 2 (2010), 73–205.
- Ronitt Rubinfeld and Madhu Sudan. 1996. Robust characterizations of polynomials with applications to program testing. *SIAM J. Comput.* 25, 2 (1996), 252–271. DOI: Retrieved from <https://doi.org/10.1137/S0097539793255151>.
- Rocco Servedio, Li-Yang Tan, and John Wright. 2015. Adaptivity helps for testing juntas. In *Proceedings of the 30th IEEE Conference on Computational Complexity*, vol. 33 of LIPIcs, 264–279.
- L. Valiant. 1984. A theory of the learnable. *Commun. ACM* 27, 11 (1984), 1134–1142.

Received February 2018; accepted July 2018