

Convex Influences

Anindya De

University of Pennsylvania, Philadelphia, PA, USA

Shivam Nadimpalli

Columbia University, New York, NY, USA

Rocco A. Servedio

Columbia University, New York, NY, USA

Abstract

We introduce a new notion of influence for symmetric convex sets over Gaussian space, which we term “convex influence”. We show that this new notion of influence shares many of the familiar properties of influences of variables for monotone Boolean functions $f : \{\pm 1\}^n \rightarrow \{\pm 1\}$.

Our main results for convex influences give Gaussian space analogues of many important results on influences for monotone Boolean functions. These include (robust) characterizations of extremal functions, the Poincaré inequality, the Kahn-Kalai-Linial theorem [28], a sharp threshold theorem of Kalai [29], a stability version of the Kruskal-Katona theorem due to O’Donnell and Wimmer [44], and some partial results towards a Gaussian space analogue of Friedgut’s junta theorem [24]. The proofs of our results for convex influences use very different techniques than the analogous proofs for Boolean influences over $\{\pm 1\}^n$. Taken as a whole, our results extend the emerging analogy between symmetric convex sets in Gaussian space and monotone Boolean functions from $\{\pm 1\}^n$ to $\{\pm 1\}$.

2012 ACM Subject Classification Mathematics of computing; Mathematics of computing $\rightarrow$ Probability and statistics

Keywords and phrases Fourier analysis of Boolean functions, convex geometry, influences, threshold phenomena

Digital Object Identifier 10.4230/LIPIcs.ITCS.2022.53

Related Version *Full Version:* <https://arxiv.org/abs/2109.03107>

Funding This material is based upon work supported by the National Science Foundation under grant numbers listed above. Any opinions, findings and conclusions or recommendations expressed in this material are those of the authors and do not necessarily reflect the views of the National Science Foundation (NSF).

Anindya De: A.D. is supported by NSF grants CCF-1910534, CCF-1926872, and CCF-2048128.

Shivam Nadimpalli: S.N. is supported by NSF grants CCF-1563155 and by CCF-1763970.

Rocco A. Servedio: R.A.S. is supported by NSF grants CCF-1814873, IIS-1838154, CCF-1563155, and by the Simons Collaboration on Algorithms and Geometry.

Acknowledgements This work was done while A. D. was participating in the "Probability, Geometry, and Computation in High Dimensions" program at the Simons Institute for the Theory of Computing. A. D. wishes to thank the Simons Institute for their support.

1 Introduction

Background: An intriguing analogy. This paper is motivated by an intriguing, but at this point only partially understood, analogy between *monotone Boolean functions over the hypercube* and *symmetric convex sets in Gaussian space*. Perhaps the simplest manifestation of this analogy is the following pair of easy observations: since a Boolean function $f : \{\pm 1\}^n \rightarrow \{\pm 1\}$ is monotone if $f(x) \leq f(y)$ whenever $x_i \leq y_i$ for all i , it is clear that “moving an input up towards 1^n ” by flipping bits from -1 to 1 can never decrease the value of f .



© Anindya De, Shivam Nadimpalli, and Rocco A. Servedio;
licensed under Creative Commons License CC-BY 4.0

13th Innovations in Theoretical Computer Science Conference (ITCS 2022).

Editor: Mark Braverman; Article No. 53; pp. 53:1–53:21

Leibniz International Proceedings in Informatics



LIPICs Schloss Dagstuhl – Leibniz-Zentrum für Informatik, Dagstuhl Publishing, Germany

Similarly, we may view a symmetric¹ convex set $K \subseteq \mathbb{R}^n$ as a 0/1 valued function, and it is clear from symmetry and convexity that “moving an input in towards the origin” can never decrease the value of the function.

The analogy extends far beyond these easy observations to involve many analytic and algorithmic aspects of monotone Boolean functions over $\{\pm 1\}^n$ under the uniform distribution and symmetric convex subsets of $\mathbb{R}^n$ under the Gaussian measure. Below we survey some known points of correspondence (several of which were only recently established) between the two settings:

1. **Density increments.** The well-known Kruskal-Katona theorem [36, 31] gives quantitative information about how rapidly a monotone $f : \{\pm 1\}^n \rightarrow \{\pm 1\}$ increases on average as the input to f is “moved up towards 1^n .” Let $f : \{\pm 1\}^n \rightarrow \{0, 1\}$ be a monotone function and let $\mu_f(j)$ be the fraction of the $\binom{n}{j}$ many weight- j inputs for which f outputs 1; the Kruskal-Katona theorem implies (see e.g. [41]) that if $k = cn$ for some c bounded away from 0 and 1 and $\mu_f(k) \in [0.1, 0.9]$, then $\mu_f(k+1) \geq \mu_f(k) + \Theta(1/n)$. Analogous “density increment” results for symmetric convex sets are known to hold in various forms, where the analogue of moving an input in $\{\pm 1\}^n$ up towards 1^n is now moving an input in $\mathbb{R}^n$ in towards the origin, and the analogue of $\mu_f(j)$ is now $\alpha_r(K)$, which is defined to be the fraction of the origin-centered radius- r sphere $r\mathbb{S}^{n-1}$ that lies in K . For example, Theorem 2 of the recent work [16] shows that if $K \subseteq \mathbb{R}^n$ is a symmetric convex set (which we view as a function $K : \mathbb{R}^n \rightarrow \{0, 1\}$) and $r = \Theta(\sqrt{n})$ satisfies $\alpha_r(K) \in [0.1, 0.9]$, then $\alpha_K(r(1 - 1/n)) \geq \alpha_K(r) + \Theta(1/n)$.
2. **Weak learning from random examples.** Building on the above-described density increment for symmetric convex sets, [16] showed that any symmetric convex set can be learned to accuracy $1/2 + \Omega(1)/\sqrt{n}$ in $\text{poly}(n)$ time given $\text{poly}(n)$ many random examples drawn from $\mathcal{N}(0, 1)^n$. [16] also shows that any $\text{poly}(n)$ -time weak learning algorithm (even if allowed to make membership queries) can achieve accuracy no better than $1/2 + O(\log(n)/\sqrt{n})$. These results are closely analogous to the known (matching) upper and lower bounds for $\text{poly}(n)$ -time weak learning of monotone functions with respect to the uniform distribution over $\{\pm 1\}^n$: Blum et al. [6] showed that $1/2 + \Theta(\log(n)/\sqrt{n})$ is the best possible accuracy for a $\text{poly}(n)$ -time weak learner (even if membership queries are allowed), and O’Donnell and Wimmer [44] gave a $\text{poly}(n)$ time weak learner that achieves this accuracy using random examples only.
3. **Analytic structure and strong learning from random examples.** [11] showed that the Fourier spectrum of any n -variable monotone Boolean function over $\{\pm 1\}^n$ is concentrated in the first $O(\sqrt{n})$ levels. Analogously, [35] showed that the same concentration holds for the first $O(\sqrt{n})$ levels of the Hermite spectrum² of the indicator function of any convex set. In both cases this concentration gives rise to a learning algorithm, using random examples only, running in $n^{O(\sqrt{n})}$ time and learning the relevant class (either monotone Boolean functions over the n -dimensional hypercube or convex sets under Gaussian space) to any constant accuracy.
4. **Qualitative correlation inequalities.** The well-known Harris-Kleitman theorem [26, 34] states that monotone Boolean functions are non-negatively correlated: any monotone $f, g : \{\pm 1\}^n \rightarrow \{0, 1\}$ must satisfy $\mathbf{E}[fg] - \mathbf{E}[f]\mathbf{E}[g] \geq 0$. The Gaussian

¹ A set $K \subseteq \mathbb{R}^n$ is symmetric if $-x \in K$ whenever $x \in K$.

² The Hermite polynomials form an orthonormal basis for the space of square-integrable real-valued functions over Gaussian space; the Hermite spectrum of a function over Gaussian space is analogous to the familiar Fourier spectrum of a function over the Boolean hypercube. See Section 2 for details.

Correlation Inequality [48] gives an exactly analogous statement for symmetric convex sets in Gaussian space: if $K, L \subseteq \mathbb{R}^n$ are any two symmetric convex sets, then $\mathbf{E}[KL] - \mathbf{E}[K] \mathbf{E}[L] \geq 0$, where now expectations are with respect to $\mathcal{N}(0, 1)^n$.

5. **Quantitative correlation inequalities.** Talagrand [50] proved the following *quantitative* version of the Harris–Kleitman inequality: for monotone $f, g : \{\pm 1\}^n \rightarrow \{0, 1\}$,

$$\mathbf{E}[fg] - \mathbf{E}[f] \mathbf{E}[g] \geq \frac{1}{C} \cdot \Psi \left(\sum_{i=1}^n \mathbf{Inf}_i[f] \mathbf{Inf}_i[g] \right). \quad (1)$$

Here $\Psi(x) := x / \log(e/x)$, $C > 0$ is an absolute constant, $\mathbf{Inf}_i[f]$ is the influence of coordinate i on f (see Section 2), and the expectations are with respect to the uniform distribution over $\{\pm 1\}^n$. In a recent work [14] proved a closely analogous quantitative version of the Gaussian Correlation Inequality: for K, L symmetric convex subsets of $\mathbb{R}^n$,

$$\mathbf{E}[KL] - \mathbf{E}[K] \mathbf{E}[L] \geq \frac{1}{C} \cdot \Upsilon \left(\sum_{i=1}^n \tilde{K}(2e_i) \tilde{L}(2e_i) \right), \quad (2)$$

where $\Upsilon : [0, 1] \rightarrow [0, 1]$ is $\Upsilon(x) = \min \left\{ x, \frac{x}{\log^2(1/x)} \right\}$, $C > 0$ is a universal constant, $\tilde{K}(2e_i)$ denotes the degree-2 Hermite coefficient in direction e_i (see Section 2), and expectations are with respect to $\mathcal{N}(0, 1)^n$.

We remark that in many of the above cases the proofs of the two analogous results (Boolean versus Gaussian) are very different from each other even though the statements are quite similar. For example, the Harris-Kleitman theorem has a simple one-paragraph proof by induction on n , whereas the Gaussian Correlation Inequality was a famous conjecture for four decades before Thomas Royen proved it in 2014.

Motivation. We feel that the examples presented above motivate a deeper understanding of this “Boolean/Gaussian analogy.” This analogy may be useful in a number of ways; in particular, via this connection known results in one setting may suggest new questions and results for the other setting.³ Thus the overarching goal of this paper is to strengthen the analogy between monotone Boolean functions over $\{\pm 1\}^n$ and symmetric convex sets in Gaussian space. We do this through the study of a new notion of *influence* for symmetric convex sets in Gaussian space.

1.1 This Work: A New Notion of Influence for Symmetric Convex Sets

Before presenting our new notion of influence for symmetric convex sets in Gaussian space, we first briefly recall the usual notion for Boolean functions. For $f : \{\pm 1\}^n \rightarrow \{\pm 1\}^n$, the *influence of coordinate i on f* , denoted $\mathbf{Inf}_i[f]$, is $\Pr[f(\mathbf{x}) \neq f(\mathbf{x}^{\oplus i})]$, where $\mathbf{x}$ is uniform random over $\{\pm 1\}^n$ and $\mathbf{x}^{\oplus i}$ denotes $\mathbf{x}$ with its i -th coordinate flipped. It is a well-known fact (see e.g. Proposition 2.21 of [45]) that for monotone Boolean functions f , we have $\mathbf{Inf}_i[f] = \hat{f}(i)$, the degree-1 Fourier coefficient corresponding to coordinate i .

³ Indeed, the recent Gaussian density increment and weak learning results of [16] were inspired by the Kruskal-Katona theorem and the weak learning algorithms and lower bounds of [6] for monotone Boolean functions. Similarly, the recent quantitative version of the Gaussian Correlation Theorem established in [14] was motivated by the existence of Talagrand’s quantitative correlation inequality for monotone Boolean functions.

Inspired by the relation $\mathbf{Inf}_i[f] = \widehat{f}(i)$ for influence of monotone Boolean functions, and by the close resemblance between Equation (1) and Equation (2), [14] proposed to define the *influence of K along direction v* , for $K : \mathbb{R}^n \rightarrow \{0, 1\}$ a symmetric convex set and $v \in \mathbb{S}^{n-1}$, to be

$$\mathbf{Inf}_v[K] := -\widetilde{K}(2v),$$

the (negated) degree-2 Hermite coefficient⁴ of K in direction v (see Definition 10 for a detailed definition). [14] proved that this quantity is non-negative for any direction v and any symmetric convex K (see Proposition 11). They also defined the *total influence of K* to be

$$\mathbf{I}[f] := \sum_{i=1}^n \mathbf{Inf}_{e_i}[f] \quad (3)$$

and observed that this definition is invariant under different choices of orthonormal basis other than $e_1, \dots, e_n$, but did not explore these definitions further.

The main contribution of the present work is to carry out an in-depth study of this new notion of influence for symmetric convex sets. For conciseness, and to differentiate it from other influence notions (which we discuss later), we will sometimes refer to this new notion as “convex influence.”

Inspired by well known results about influence of monotone Boolean functions, we establish a number of different results about convex influence which show that this notion shares many properties with the familiar Boolean influence notion. Intriguingly, and similar to the Boolean/Gaussian analogy elements discussed earlier, while the statements we prove about convex influence are quite closely analogous to known results about Boolean influences, the proofs and tools that we use (Gaussian isoperimetry, Brascamp-Lieb type inequalities, etc.) are very different from the ingredients that underlie the corresponding results about Boolean influence.

1.2 Results and Organization

We give an overview of our main results below.

1.2.1 Basics, Examples, and Margulis–Russo

We begin in Section 3.1 by working through some basic properties of our new influence notion. After analyzing some simple examples in Section 3.2, we next show in Section 3.3 that the total convex influence for a symmetric convex set is equal to (a scaled version of) the rate of change of the Gaussian volume of the set as the variance of the underlying Gaussian is changed. This gives an alternate characterization of total convex influence, and may be viewed as an analogue of the Margulis–Russo formula for our new influence notion.

1.2.2 Lower Bounds on Total Convex Influence

In Section 4, we give a lower bound on the total convex influence (Equation (3)) for symmetric convex sets, which is closely analogous to the classical KKL Theorem. Our KKL analogue is quadratically weaker than the KKL theorem for Boolean functions; we conjecture that a

⁴ We observe that if K is a symmetric set then since its indicator function is even, the degree-1 Hermite coefficient $\widetilde{K}(v)$ must be 0 for any direction v .

stronger bound in fact holds, which would quantitatively align with the Boolean variant (see Item 1 of Section 1.3). Our proof relies on the Gaussian isoperimetric inequality and differs quite significantly from the proof of the KKL theorem for Boolean functions.

1.2.3 Sharp Thresholds for Sets with All Small Influences

In Section 5, we establish a “sharp threshold” result for symmetric convex sets in Gaussian space, which is analogous to a sharp threshold result for monotone Boolean functions due to Kalai [29]. Building on earlier work of Friedgut and Kalai [25], Kalai [29] showed that if $f : \{\pm 1\}^n \rightarrow \{0, 1\}$ is a monotone Boolean function and $p \in (0, 1)$ is such that (i) all the p -biased influences of f are $o_n(1)$ and (ii) the expectation of f under the p -biased measure is $\Theta(1)$, then f must have a “sharp threshold” in the following sense: the expectation of f under the p_1 -biased measure (p_2 -biased measure, respectively) is $o_n(1)$ ($1 - o_n(1)$, respectively) for some $p_1 < p < p_2$ with $p_2 - p_1 = o_n(1)$. For our sharp threshold result, we prove an analogous statement for symmetric convex sets, where now $\mathcal{N}(0, \sigma^2)$ takes the place of the p -biased distribution over $\{\pm 1\}^n$ and the σ -biased convex influences (see Definition 19) take the place of the p -biased influences. Interestingly, the sharpness of our threshold is quantitatively better than the known analogous result [29] for monotone Boolean functions; see Section 5 for an elaboration of this point.

1.2.4 A Stable Density Increment Result

Finally, in Section 6, we use our new influence notion to give a Gaussian space analogue of a “stability” version of the Kruskal-Katona theorem due to O’Donnell and Wimmer [44]. In [44] it is shown that the $\Omega(1/n)$ density increment of the Kruskal-Katona theorem (see Item 1 at the beginning of this introduction) can be strengthened to $\Omega(\log(n)/n)$ as long as a “low individual influences”-type condition holds. We analogously show that a similar strengthening of the Gaussian space density increment result mentioned in Item 1 earlier can be achieved under the condition that the convex influence in every direction is low.

1.2.5 Additional Results in the Full Version

In the full version of this paper [15], we give a number of additional results about convex influences. These include a convex influence analogue of a consequence of Friedgut’s junta theorem; a convex influence analogue of the Poincaré inequality; a characterization of symmetric convex sets that are extremal with respect to convex influence; and a comparison with previously studied notions of influence over Gaussian space.

1.3 Discussion and Future Work

We believe that much more remains to be discovered about this new notion of influences for symmetric convex sets. We list some natural concrete (and not so concrete) questions for future work:

1. **A stronger KKL-type theorem for convex influences?** We conjecture that the factor of $\sqrt{\log(\text{Var}[K]/\delta)}$ in our KKL analogue, Theorem 22, can be strengthened to $\log(1/\delta)$. As witnessed by Example 18, this would be essentially the strongest possible quantitative result, and would align closely with the original KKL theorem [28].
2. **An analogue of Friedgut’s theorem for convex influences?** As noted earlier, in the full version of this paper, we establish a Gaussian space analogue of a consequence of Friedgut’s Junta Theorem [24] for Boolean functions over $\{\pm 1\}^n$. The following would give a full-fledged Gaussian space analogue of Friedgut’s Junta Theorem:

► **Conjecture 1** (Friedgut’s Junta Theorem for convex influences). *Let $K \subseteq \mathbb{R}^n$ be a convex symmetric set with $\text{Inf}[K] \leq I$. Then there are $J \leq 2^{O(I/\varepsilon)}$ orthonormal directions $v^1, \dots, v^J \in \mathbb{S}^{n-1}$ and a symmetric convex set $L \subseteq \mathbb{R}^n$, such that*

a. $L(x)$ depends only on the values of $v^1 \cdot x, \dots, v^J \cdot x$, and

b. $\Pr_{x \sim \mathcal{N}(0,1)^n}[K(x) \neq L(x)] \leq \varepsilon$.

3. **Are low-influence directions (almost) irrelevant?** Related to the previous question, we note that it seems to be surprisingly difficult to show that low-influence directions “don’t matter much” for convex sets. For example, it is an open question to establish the following, which would give a dimension-free robust version of the last assertion of Proposition 11:

► **Conjecture 2.** *Let $K \subseteq \mathbb{R}^n$ be symmetric and convex, and suppose that $v \in \mathbb{S}^{n-1}$ is such that $\text{Inf}_v[K] \leq \varepsilon$. Then there is a symmetric convex set L such that*

a. $L(x)$ depends only on the projection of x onto the $(n-1)$ dimensional subspace orthogonal to v , and

b. $\Pr_{x \sim \mathcal{N}(0,1)^n}[K(x) \neq L(x)] \leq \tau(\varepsilon)$ for some function τ depending only on ε (in particular, independent of n) and going to 0 as $\varepsilon \rightarrow 0$.

While the corresponding Boolean statement is very easy to establish, natural approaches to Conjecture 2 lead to open (and seemingly challenging) questions regarding dimension-free stable versions of the Ehrhard-Borell inequality [23, 53].

4. **Algorithmic results?** Finally, a broader goal is to further explore the similarities and differences between the theory of convex symmetric sets in Gaussian space and the theory of monotone Boolean functions over $\{\pm 1\}^n$. One topic where the gap in our understanding is particularly wide is the algorithmic problem of *property testing*. The problem of testing monotonicity of functions from $\{\pm 1\}^n$ to $\{\pm 1\}$ is rather well understood, with the current state of the art being an $\tilde{O}(n^{1/2})$ -query upper bound and an $\tilde{\Omega}(n^{1/3})$ -query lower bound [33, 13]. In contrast, the problem of testing whether an unknown region in $\mathbb{R}^n$ is convex (with respect to the standard normal distribution) is essentially wide open, with the best known upper bound being $n^{O(\sqrt{n})}$ queries [12] and no nontrivial lower bounds known.

2 Preliminaries

In this section we give preliminaries setting notation and recalling useful background on convex geometry, log-concave functions, and Hermite analysis over $\mathcal{N}(0, \sigma^2)^n$.

2.1 Convex Geometry and Log-Concavity

Below we briefly recall some notation, terminology and background from convex geometry and log-concavity. Some of our main results employ relatively sophisticated results from these areas; we will recall these as necessary in the relevant sections and here record only basic facts. For a general and extensive resource we refer the interested reader to [2].

We identify sets $K \subseteq \mathbb{R}^n$ with their indicator functions $K : \mathbb{R}^n \rightarrow \{0, 1\}$, and we say that $K \subseteq \mathbb{R}^n$ is *symmetric* if $K(x) = K(-x)$. We write B_r to denote the origin-centered ball of radius r in $\mathbb{R}^n$. If $K \subseteq \mathbb{R}^n$ is a nonempty symmetric convex set then we let $r_{\text{in}}(K)$ denote $\sup_{r \geq 0} \{r : B_r \subseteq K\}$ and we refer to this as the *in-radius* of K .

Recall that a function $f : \mathbb{R}^n \rightarrow \mathbb{R}_{\geq 0}$ is *log-concave* if its domain is a convex set and it satisfies $f(\theta x + (1-\theta)y) \geq f(x)^\theta f(y)^{1-\theta}$ for all $x, y \in \text{domain}(f)$ and $\theta \in [0, 1]$. In particular, the 0/1-indicator functions of convex sets are log-concave.

Recall that the *marginal* of $f : \mathbb{R}^n \rightarrow \mathbb{R}$ on the set of variables $\{i_1, \dots, i_k\}$ is obtained by integrating out the other variables, i.e. it is the function

$$g(x_{i_1}, \dots, x_{i_k}) = \int_{\mathbb{R}^{n-k}} f(x_1, \dots, x_n) dx_{j_1} \dots dx_{j_{n-k}},$$

where $\{j_1, \dots, j_{n-k}\} = [n] \setminus \{i_1, \dots, i_k\}$. We recall the following fact:

► **Fact 3** ([17, 39, 46, 47] (see Theorem 5.1, [40])). *All marginals of a log-concave function are log-concave.*

The next fact follows easily from the definition of log-concavity:

► **Fact 4** ([27], see e.g. [1]). *A one-dimensional log-concave function is unimodal.*

2.2 Gaussian Random Variables

We write $z \sim \mathcal{N}(0, 1)$ to mean that z is a standard Gaussian random variable, and will use the notation

$$\varphi(z) := \frac{1}{\sqrt{2\pi}} e^{-z^2/2} \quad \text{and} \quad \Phi(z) := \int_{-\infty}^z \varphi(t) dt$$

to denote the pdf and the cdf of this random variable.

Recall that a non-negative random variable r^2 is distributed according to the chi-squared distribution $\chi^2(n)$ if $r^2 = g_1^2 + \dots + g_n^2$ where $g \sim \mathcal{N}(0, 1)^n$, and that a draw from the chi distribution $\chi(n)$ is obtained by making a draw from $\chi^2(n)$ and then taking the square root.

We define the *shell-density function* for K , $\alpha_K : [0, \infty) \rightarrow [0, 1]$, to be

$$\alpha_K(r) := \Pr_{\mathbf{x} \in r\mathbb{S}^{n-1}}[\mathbf{x} \in K], \quad (4)$$

where the probability is with respect to the normalized Haar measure over $r\mathbb{S}^{n-1}$; so $\alpha_K(r)$ equals the fraction of the origin-centered radius- r sphere which lies in K . We observe that if K is convex and symmetric then $\alpha_K(\cdot)$ is a nonincreasing function. A view which will be sometimes useful later is that $\alpha_K(r)$ is the probability that a random Gaussian-distributed point $g \sim \mathcal{N}(0, 1)^n$ lies in K , conditioned on $\|g\| = r$.

2.3 Hermite Analysis over $\mathcal{N}(0, \sigma^2)^n$

Our notation and terminology here follow Chapter 11 of [45]. We say that an n -dimensional *multi-index* is a tuple $\alpha \in \mathbb{N}^n$, and we define

$$|\alpha| := \sum_{i=1}^n \alpha_i. \quad (5)$$

We write $\mathcal{N}(0, \sigma^2)^n$ to denote the n -dimensional Gaussian distribution with mean 0 and variance σ^2 , and denote the corresponding measure by $\gamma_{n, \sigma}(\cdot)$. When the dimension n is clear from context we simply write $\gamma_\sigma(\cdot)$ instead, and sometimes when $\sigma = 1$ we simply write γ for γ_1 . For $n \in \mathbb{N}_{>0}$ and $\sigma > 0$, we write $L^2(\mathbb{R}^n, \gamma_\sigma)$ to denote the space of functions $f : \mathbb{R}^n \rightarrow \mathbb{R}$ that have finite second moment $\|f\|_2^2$ under the Gaussian measure γ_σ , that is:

$$\|f\|_2^2 = \mathbf{E}_{z \sim \mathcal{N}(0, \sigma^2)^n} [f(z)^2]^{1/2} < \infty.$$

We view $L^2(\mathbb{R}^n, \gamma_\sigma)$ as an inner product space with $\langle f, g \rangle := \mathbf{E}_{z \sim \mathcal{N}(0, \sigma^2)^n} [f(z)g(z)]$ for $f, g \in L^2(\mathbb{R}^n, \gamma_\sigma)$. We define “biased Hermite polynomials,” which yield an orthonormal basis for $L^2(\mathbb{R}^n, \gamma_\sigma)$:

► **Definition 5** (Hermite basis). For $\sigma > 0$, the σ -biased Hermite polynomials $(h_{j,\sigma})_{j \in \mathbb{N}}$ are the univariate polynomials defined as

$$h_{j,\sigma}(x) := h_j\left(\frac{x}{\sigma}\right), \quad \text{where} \quad h_j(x) := \frac{(-1)^j}{\sqrt{j!}} \exp\left(\frac{x^2}{2}\right) \cdot \frac{d^j}{dx^j} \exp\left(-\frac{x^2}{2}\right).$$

► **Fact 6** (Easy extension of Proposition 11.33, [45]). For $n \geq 1$ and $\sigma > 0$, the collection of n -variate σ -biased Hermite polynomials given by $(h_{\alpha,\sigma})_{\alpha \in \mathbb{N}^n}$ where

$$h_{\alpha,\sigma}(x) := \prod_{i=1}^n h_{\alpha_i,\sigma}(x_i)$$

forms a complete, orthonormal basis for $L^2(\mathbb{R}^n, \gamma_\sigma)$.

Given a function $f \in L^2(\mathbb{R}^n, \gamma_\sigma)$ and $\alpha \in \mathbb{N}^n$, we define its $(\sigma$ -biased) Hermite coefficient on α to be $\tilde{f}_\sigma(\alpha) := \langle f, h_{\alpha,\sigma} \rangle$. It follows that f is uniquely expressible as $f = \sum_{\alpha \in \mathbb{N}^n} \tilde{f}_\sigma(\alpha) h_{\alpha,\sigma}$ with the equality holding in $L^2(\mathbb{R}^n, \gamma_\sigma)$; we will refer to this expansion as the $(\sigma$ -biased) Hermite expansion of f . When $\sigma = 1$, we will simply write $\tilde{f}(\alpha)$ instead of $\tilde{f}_\sigma(\alpha)$ and h_α instead of $h_{\alpha,1}$. Parseval's and Plancherel's identities hold in this setting:

► **Fact 7.** For $f, g \in L^2(\mathbb{R}^n, \gamma_\sigma)$, we have:

$$\langle f, g \rangle = \mathbf{E}_{z \sim \mathcal{N}(0, \sigma^2)^n} [f(z)g(z)] = \sum_{\alpha \in \mathbb{N}^n} \tilde{f}_\sigma(\alpha) \tilde{g}_\sigma(\alpha), \quad (\text{Plancherel})$$

$$\langle f, f \rangle = \mathbf{E}_{z \sim \mathcal{N}(0, \sigma^2)^n} [f(z)^2] = \sum_{\alpha \in \mathbb{N}^n} \tilde{f}_\sigma(\alpha)^2. \quad (\text{Parseval})$$

The following notation will sometimes come in handy.

► **Definition 8.** Let $v \in \mathbb{S}^{n-1}$ and $f \in L^2(\mathbb{R}^n, \gamma_\sigma)$. We define f 's σ -biased Hermite coefficient of degree k along v , written $\tilde{f}_\sigma(kv)$, to be

$$\tilde{f}_\sigma(kv) := \mathbf{E}_{x \sim \mathcal{N}(0, \sigma^2)^n} [f(x) \cdot h_{k,\sigma}(v \cdot x)]$$

(as usual omitting the subscript when $\sigma = 1$).

► **Notation 9.** We will write $e_i \in \mathbb{N}^n$ to denote the i^{th} standard basis vector for $\mathbb{R}^n$.

In this notation, for example, $\tilde{f}(2e_i) = \mathbf{E}_{x \sim \mathcal{N}(0,1)^n} [f(x) \cdot h_2(x_i)]$. Finally, for a measurable set $K \subseteq \mathbb{R}^n$, it will be convenient for us to write $\gamma(K)$ to denote $\Pr_{x \sim \mathcal{N}(0,1)^n} [x \in K]$, the (standard) Gaussian volume of K .

3 Influences for Symmetric Convex Sets

In this section, we first introduce our new notion of influence for symmetric convex sets over Gaussian space and establish some basic properties. In Section 3.2 we analyze the influences of several natural symmetric convex sets, and in Section 3.3 we give an analogue of the Margulis-Russo formula (characterizing the influences of monotone Boolean functions) which provides an alternative equivalent view of our new notion of influence for symmetric convex sets in terms of the behavior of the sets under dilations.

3.1 Definitions and Basic Properties

► **Definition 10** (Influence for symmetric log-concave functions). *Let $f \in L^2(\mathbb{R}^n, \gamma)$ be a symmetric (i.e. $f(x) = f(-x)$) log-concave function. Given a unit vector $v \in \mathbb{S}^{n-1}$, we define the influence of direction v on f as being*

$$\mathbf{Inf}_v[f] := -\tilde{f}(2v) = \mathbf{E}_{\mathbf{x} \sim \mathcal{N}(0,1)^n} [-f(\mathbf{x}) h_2(v \cdot \mathbf{x})] = \mathbf{E}_{\mathbf{x} \sim \mathcal{N}(0,1)^n} \left[f(\mathbf{x}) \cdot \left(\frac{1 - (v \cdot \mathbf{x})^2}{\sqrt{2}} \right) \right],$$

the negated “degree-2 Hermite coefficient in the direction v .” Furthermore, we define the total influence of f as

$$\mathbf{I}[f] := \sum_{i=1}^n \mathbf{Inf}_{e_i}[f].$$

Note that the indicator of a symmetric convex set is a symmetric log-concave function, and this is the setting that we will be chiefly interested in. The following proposition (which first appeared in [14], and a proof of which can be found in the full version [15]) shows that these new influences are indeed “influence-like.” An arguably simpler argument for the non-negativity of influences is presented in Section 3.3.

► **Proposition 11** (Influences are non-negative). *If K is a centrally symmetric, convex set, then $\mathbf{Inf}_v[K] \geq 0$ for all $v \in \mathbb{S}^{n-1}$. Furthermore, equality holds if and only if $K(x) = K(y)$ whenever $x_{v^\perp} = y_{v^\perp}$ (i.e. the projection of x orthogonal to v coincides with that of y) almost surely.*

We note that the total influence of a symmetric, convex set K is independent of the choice of basis; indeed, we have

$$\mathbf{I}[K] = \mathbf{E}_{\mathbf{x} \sim \mathcal{N}(0,1)^n} \left[f(\mathbf{x}) \left(\frac{n - \|\mathbf{x}\|^2}{\sqrt{2}} \right) \right] \quad (6)$$

which is invariant under orthogonal transformations. Hence any orthonormal basis $\{v_1, \dots, v_n\}$ could have been used in place of $\{e_1, \dots, e_n\}$ in defining $\mathbf{I}[K]$.

We note that (as is shown in the proof of Proposition 11), the influence of a fixed coordinate is not changed by averaging over some set of other coordinates:

► **Fact 12.** *Let $K \subseteq \mathbb{R}^n$ be a symmetric, convex set, and define the log-concave function $K_{e_i} : \mathbb{R} \rightarrow [0, 1]$ as*

$$K_{e_i}(x) := \mathbf{E}_{\mathbf{x} \sim \mathcal{N}(0,1)^{n-1}} [K(\mathbf{x}_1, \dots, \mathbf{x}_{i-1}, x, \mathbf{x}_{i+1}, \dots, \mathbf{x}_n)]. \quad (7)$$

Then we have

$$\mathbf{Inf}_{e_i}[K] = \mathbf{Inf}_{e_1}[K_{e_i}] = \mathbf{I}[K_{e_i}]. \quad (8)$$

We conclude with the following useful relationship between the in-radius of a symmetric convex set K and its max influence along any direction.

► **Proposition 13.** *Let $K \subseteq \mathbb{R}^n$ be a centrally symmetric convex set with $\gamma(K) \geq \Delta$, and let $r_{\text{in}} = r_{\text{in}}(K)$ be the in-radius of K . Then there is some direction $v \in \mathbb{S}^{n-1}$ such that*

$$\mathbf{Inf}_v[K] \geq \frac{\Delta e^{-r_{\text{in}}^2}}{2^{3/2} \pi}.$$

We will use the following Brascamp–Lieb-type inequality.

► **Lemma 14** (Final assertion of Lemma 4.7 of [51]). *If $g : \mathbb{R} \rightarrow \mathbb{R}_{\geq 0}$ is log-concave and symmetric and supported in $[-c, c]$, then*

$$\frac{\int_{-c}^c x^2 e^{-x^2/2} g(x) dx}{\int_{-c}^c e^{-x^2/2} g(x) dx} \leq 1 - \frac{1}{2\pi} e^{-c^2}.$$

We use this in the proof of the following claim, which will easily yield Proposition 13:

► **Proposition 15.** *Let $K \subseteq \mathbb{R}^n$ be a centrally symmetric convex set with $\gamma(K) \geq \Delta$, and let $v \in \mathbb{S}^{n-1}$ be a unit vector such that $K \subseteq \{x \in \mathbb{R}^n : |v \cdot x| \leq c\}$. Then we have*

$$\mathbf{Inf}_v[K] \geq \frac{\Delta e^{-c^2}}{2^{3/2}\pi}.$$

Proof of Proposition 15. For ease of notation, we take $v = e_1$ and so $K \subseteq \{x \in \mathbb{R}^n : |x_1| \leq c\}$. From Equations (7) and (8), we have that

$$\mathbf{Inf}_v[K] = \mathbf{Inf}_{e_1}[K] = \mathbf{I}[K_{e_1}] = \frac{1}{2\sqrt{\pi}} \int_{\mathbb{R}} K_{e_1}(x)(1 - x^2)e^{-x^2/2} dx \quad (9)$$

where $K_{e_1} : \mathbb{R} \rightarrow [0, 1]$ is the symmetric log-concave function given by

$$K_{e_1}(x) := \mathbf{E}_{\mathbf{x} \sim \mathcal{N}(0,1)^{n-1}} [K(x, \mathbf{x}_2, \dots, \mathbf{x}_n)].$$

As $K(x) = 0$ when $|x_1| > c$ we have that $\text{supp}(g) \subseteq [-c, c]$ and so it follows from Equation (9) that

$$\mathbf{Inf}_v[K] = \frac{1}{2\sqrt{\pi}} \int_{-c}^c K_{e_1}(x)(1 - x^2)e^{-x^2/2} dx. \quad (10)$$

It follows then from Lemma 14 that

$$\mathbf{Inf}_v[K] \geq \frac{1}{2^{3/2}\pi} \left(\frac{e^{-c^2}}{\sqrt{2\pi}} \int_{-c}^c K_{e_1}(x)e^{-x^2/2} dx \right) = \frac{\Delta e^{-c^2}}{2^{3/2}\pi}$$

which completes the proof of Proposition 15. ◀

Proof of Proposition 13. By definition of the in-radius and the supporting hyperplane theorem, there must exist some unit vector $\hat{v} \in \mathbb{R}^n$ such that

$$K \subseteq K_* := \{x \in \mathbb{R}^n : |\hat{v} \cdot x| \leq r_{\text{in}}\},$$

and hence by Proposition 15 we get that

$$\mathbf{Inf}_{\hat{v}}[K] \geq \frac{\gamma(K)e^{-r_{\text{in}}^2}}{2^{3/2}\pi} \geq \frac{\mathbf{Var}[K]e^{-r_{\text{in}}^2}}{2^{3/2}\pi},$$

giving Proposition 13 as claimed. ◀

3.2 Influences of Specific Symmetric Convex Sets

In this subsection we consider some concrete examples by analyzing the influences of a few specific symmetric convex sets, namely “slabs”, balls, and cubes. As we will see, these are closely analogous to well-studied monotone Boolean functions (dictator, Majority, and Tribes, respectively).

► **Example 16** (Analogue of Boolean dictator: a “slab”). Given a vector $w \in \mathbb{R}^n$, define $\text{Dict}_w := \{x \in \mathbb{R}^n : |\langle x, w \rangle| \leq 1\}$. As suggested by the notation, this is the analogue of a single Boolean variable $f(x) = x_i$, i.e. a “dictatorship.” For simplicity, suppose $w := \frac{1}{c} \cdot e_1$ for some $c > 0$, i.e. $\text{Dict}_w = \{x \in \mathbb{R}^n : |x_1| \leq c\}$. We then have

$$\mathbf{Inf}_{e_i}[\text{Dict}_w] = \begin{cases} \Theta(c \cdot \exp(-c^2/2)) & i = 1 \\ 0 & i \neq 1 \end{cases}.$$

Note that while in the setting of the Boolean hypercube there is only one “dictatorship” for each coordinate, in our setting given a particular direction we can have “dictatorships” of varying widths and volumes.

► **Example 17** (Analogue of Boolean Majority: a ball). Let $B_r := \{x \in \mathbb{R}^n : \|x\|_2 \leq r\}$ denote the ball of radius r . Analogous to the Boolean majority function, we argue that for $B = B_{\sqrt{n}}$ we have that $\mathbf{Inf}_{e_i}(B) = \Theta(1/\sqrt{n})$ for all $i \in [n]$.

Recall from Equation (6) that

$$\mathbf{I}[B] = \frac{1}{\sqrt{2}} \mathbf{E}_{\mathbf{x} \sim \mathcal{N}(0,1)^n} [B(\mathbf{x})(n - \|\mathbf{x}\|^2)].$$

By the Berry-Esseen Central Limit Theorem (see [5, 22] or, for example, Section 11.5 of [45]), we have that for $t \in \mathbb{R}$,

$$\left| \mathbf{Pr}_{\mathbf{x} \sim \mathcal{N}(0,1)^n} \left[\frac{\|\mathbf{x}\|^2 - n}{\sqrt{n}} \leq t \right] - \mathbf{Pr}_{\mathbf{y} \sim \mathcal{N}(0,1)} [\mathbf{y} \leq t] \right| \leq \frac{c}{\sqrt{n}}$$

for some absolute constant c . In particular, this implies that

$$\mathbf{Pr}_{\mathbf{x} \sim \mathcal{N}(0,1)^n} [\|\mathbf{x}\|^2 \leq n - \sqrt{n}] \geq \mathbf{Pr}_{\mathbf{y} \sim \mathcal{N}(0,1)} [\mathbf{y} \leq -1] - \frac{c}{\sqrt{n}} \geq 0.15.$$

Since $\mathbf{Pr}_{\mathbf{x} \sim \mathcal{N}(0,1)^n} [B(\mathbf{x}) = 1] = \frac{1}{2} \pm o_n(1)$, and $B(\mathbf{x})(n - \|\mathbf{x}\|^2)$ is never negative, it follows that

$$\mathbf{E}_{\mathbf{x} \sim \mathcal{N}(0,1)^n} [B(\mathbf{x})(n - \|\mathbf{x}\|^2)] \geq \Theta(\sqrt{n})$$

from which symmetry implies that $\mathbf{Inf}_{e_i}[B] \geq \Theta\left(\frac{1}{\sqrt{n}}\right)$ for all $i \in [n]$. The upper bound $\mathbf{Inf}_{e_i}[B] \leq \Theta\left(\frac{1}{\sqrt{n}}\right)$ follows from Parseval’s identity.

Our last example is analogous to the “Tribes CNF” function introduced by Ben-Or and Linial [4] (alternatively, see Definition 2.7 of [45]):

► **Example 18** (Analogue of Boolean Tribes: a cube). Let $C_r := \{x \in \mathbb{R}^n : |x_i| \leq r \text{ for all } i \in [n]\}$ denote the axis-aligned cube of side-length $2r$ and $\gamma(C_r) = \frac{1}{2}$, i.e. let $r > 0$ be the unique value such that

$$\mathbf{Pr}_{\mathbf{g} \sim \mathcal{N}(0,1)} [|\mathbf{g}| \leq r] = \left(\frac{1}{2}\right)^{1/n} = 1 - \frac{\Theta(1)}{n}. \quad (11)$$

By standard tail bounds on the Gaussian distribution, we have that $r = \Theta(\sqrt{\log n})$. Because of the symmetry of C_r , we have $\mathbf{Inf}_{e_i}[C_r] = \mathbf{Inf}_{e_j}[C_r]$ for all $i, j \in [n]$. Note, however, that we can write

$$C_r(x) = \prod_{i=1}^n \text{Dict}_{1/r}(x_i)$$

where $\text{Dict}_{1/r} : \mathbb{R} \rightarrow \{0, 1\}$ is as defined in Example 16. By considering the Hermite representation of $C_r(x)$, it is easy to see that

$$\mathbf{Inf}_{e_i}[C_r] = \mathbf{E}_{\mathbf{g} \sim \mathcal{N}(0,1)} [\text{Dict}_{1/r}(\mathbf{g})]^{n-1} \mathbf{I}[\text{Dict}_{1/r}].$$

By our choice of r above, we have $\mathbf{E}[\text{Dict}_{1/r}] = \sqrt[n]{1/2}$ and so

$$\mathbf{E}_{\mathbf{g} \sim \mathcal{N}(0,1)} [\text{Dict}_{1/r}(\mathbf{g})]^{n-1} = \Theta(1).$$

From Example 16, we know $\mathbf{I}[\text{Dict}_{1/r}] = \Theta(re^{-r^2/2})$, and so we have

$$\mathbf{Inf}_{e_i}[C_r] = \Theta(re^{-r^2/2}). \quad (12)$$

We now recall the following tail bound on the normal distribution (see Theorem 1.2.6 of [18] or Equation 2.58 of [52]):

$$\varphi(r) \left(\frac{1}{r} - \frac{1}{r^3} \right) \leq \mathbf{Pr}_{\mathbf{g} \sim \mathcal{N}(0,1)} [\mathbf{g} \geq r] \leq \varphi(r) \left(\frac{1}{r} - \frac{1}{r^3} + \frac{3}{r^5} \right), \quad (13)$$

where $\varphi(r) = \frac{1}{\sqrt{2\pi}} e^{-r^2/2}$ is the density function of $N(0, 1)$. Combining Equation (11), Equation (12) and Equation (13) we get that $\mathbf{Inf}_{e_i}[C_r] = \Theta(r^2) \cdot \mathbf{Pr}_{\mathbf{g} \sim \mathcal{N}(0,1)} [\mathbf{g} \geq r] = \Theta(\log(n)) \cdot \Theta(1/n)$, which corresponds to the influence of each individual variable on the Boolean “tribes” function.

3.3 Margulis-Russo for Convex Influences: An Alternative Characterization of Influences via Dilations

In this subsection we give an alternative view of the notion of influence defined above, in terms of the behavior of the Gaussian measure of the set as the variance of the underlying Gaussian is changed.⁵ This is closely analogous to the Margulis-Russo formula for monotone Boolean functions on $\{\pm 1\}^n$ (see [49, 42] or Equation (8.9) in [45]), which relates the derivative (with respect to p) of the p -biased measure of a monotone function f to the p -biased total influence of f .

We start by defining σ -biased convex influences, which are analogous to p -biased influences from Boolean function analysis (see Section 8.4 of [45]).

► **Definition 19** (σ -biased influence). *Given a centrally symmetric convex set $K \subseteq \mathbb{R}^n$, we define the σ -biased influence of direction v on K as being*

$$\mathbf{Inf}_v^{(\sigma)}[K] := -\tilde{f}_\sigma(2v) = \mathbf{E}_{\mathbf{x} \sim \mathcal{N}(0,1)^n} [-f(\mathbf{x}) h_{2,\sigma}(v \cdot \mathbf{x})],$$

⁵ Since $\gamma_\sigma(K) = \gamma(K/\sigma)$, decreasing (respectively increasing) the variance of the underlying Gaussian measure is equivalent to dilating (respectively shrinking) the set.

the negated degree-2 σ -biased Hermite coefficient in the direction v . We further define the σ -biased total influence of K as

$$\mathbf{I}^{(\sigma)}[K] := \sum_{i=1}^n \mathbf{Inf}_{e_i}^{(\sigma)}[K].$$

The proof of the following proposition, which asserts that the rate of the change of the Gaussian measure of a symmetric convex set K with respect to σ^2 is (up to scaling) equal to the σ -biased total influence of K , is deferred to the full version [15]. We note that this relation was essentially known to experts (see e.g. [37]), though we are not aware of a specific place where it appears explicitly in the literature.

► **Proposition 20** (Margulis-Russo for symmetric convex sets). *Let $K \subseteq \mathbb{R}^n$ be a centrally symmetric convex set. Then for any $\sigma > 0$ we have*

$$\frac{d}{d\sigma^2} \mathbf{E}_{\mathbf{x} \sim \mathcal{N}(0, \sigma^2)^n} [K(\mathbf{x})] = \frac{-\mathbf{I}^{(\sigma)}[K]}{\sigma^2 \sqrt{2}} = \frac{-1}{\sigma^2 \sqrt{2}} \sum_{i=1}^n \mathbf{Inf}_{e_i}^{(\sigma)}[K].$$

Note that decreasing (respectively increasing) the variance of the background Gaussian measure is equivalent to dilating (respectively shrinking) the symmetric convex set while keeping the background measure fixed; this lets us write

$$\mathbf{I}[K] = \frac{1}{\sqrt{2}} \lim_{\delta \rightarrow 0} \frac{\gamma_n(K) - \gamma_n((1 - \delta)K)}{\delta} \quad (14)$$

for a symmetric convex $K \subseteq \mathbb{R}^n$. We also note that Proposition 20 easily extends to the following coordinate-by-coordinate version (which also admits a similar description in terms of dilations):

► **Proposition 21** (Coordinate-wise Margulis-Russo). *Let $K \subseteq \mathbb{R}^n$ be a centrally symmetric convex set. Then for any $\sigma > 0$, we have*

$$\frac{d}{d\sigma_i^2} \mathbf{E}_{\substack{\mathbf{x}_i \sim \mathcal{N}(0, \sigma_i^2) \\ j \neq i : \mathbf{x}_j \sim \mathcal{N}(0, \sigma^2)}} [K(\mathbf{x})] \Big|_{\sigma_i^2 = \sigma^2} = \frac{-1}{\sigma^2 \sqrt{2}} \mathbf{Inf}_{e_i}^{(\sigma)}[K].$$

In particular, we have

$$\mathbf{Inf}_{e_i}[K] = -\sqrt{2} \frac{d}{d\sigma^2} \mathbf{E}_{\substack{\mathbf{x}_i \sim \mathcal{N}(0, \sigma^2) \\ j \neq i : \mathbf{x}_j \sim \mathcal{N}(0, 1)}} [K(\mathbf{x})] \Big|_{\sigma^2 = 1}.$$

Note that decreasing the variance of the underlying Gaussian measure along a coordinate direction cannot cause the volume of the set to decrease. It follows then that $\mathbf{Inf}_{e_i}[K] \geq 0$ for all e_i .

4 A KKL Analogue for Symmetric Convex Sets

A fundamental result on the influence of variables for Boolean functions $f : \{\pm 1\}^n \rightarrow \{\pm 1\}$ is the celebrated “KKL Theorem” of Kahn, Kalai, and Linial [28], which gives a lower bound on total influence. The KKL theorem shows (roughly speaking) that if all influences are small then the total influence must be somewhat large. Several proofs of the KKL Theorem are now known, using a range of different techniques such as the famous hypercontractive inequality [7, 3] (the original approach), methods of stochastic calculus [21], and the Log-Sobolev inequality [32]. In this section we prove our convex influence analogue of the KKL theorem using the Gaussian Isoperimetric Theorem [8].

Recall that the KKL theorem for Boolean functions over $\{\pm 1\}^n$ states that if no coordinate influence is allowed to be large (each is at most δ), then the total influence must be large (at least $\Omega(\text{Var}[f] \cdot \log(1/\delta))$). We now prove an analogous result for convex influences, though we only achieve a quadratically weaker bound in terms of the max influence:

► **Theorem 22** (KKL for symmetric convex sets). *Let $K \subseteq \mathbb{R}^n$ be a symmetric convex set with $\text{Inf}_v[K] \leq \delta \leq \text{Var}[K]/10$ for all $v \in \mathbb{S}^{n-1}$. Then*

$$\mathbf{I}[K] \geq \Omega\left(\text{Var}[K] \sqrt{\log\left(\frac{\text{Var}[K]}{\delta}\right)}\right). \quad (15)$$

Our proof of Theorem 22 is inspired by the approach of [38]. The main technical ingredient we use is the *Gaussian isoperimetric inequality*:

► **Proposition 23** (Gaussian isoperimetric inequality, [8]). *Given any Borel set $A \subseteq \mathbb{R}^n$, we have*

$$\Phi^{-1}(\gamma_n(A_t)) \geq \Phi^{-1}(\gamma_n(A)) + t$$

where $A_t := A + B_t$ is the t -enlargement of A .

We remark that it is easy to obtain Proposition 23 from the Ehrhard-Borell inequality [20, 9, 10]. We will also require the following easy estimate on the *Gaussian isoperimetric function* $\varphi \circ \Phi^{-1}(\cdot)$.

► **Proposition 24.** *Let $\Phi : \mathbb{R} \rightarrow [0, 1]$ denote the cumulative distribution function of the standard one-dimensional Gaussian distribution, and let $\varphi := \Phi'$ denote its density. Then for all $\alpha \in (0, 1)$, we have*

$$\varphi \circ \Phi^{-1}(\alpha) \geq \sqrt{\frac{2}{\pi}} \min(\alpha, 1 - \alpha).$$

Proof. By symmetry, it suffices to show that $\varphi \circ \Phi^{-1}(\alpha) \geq \sqrt{\frac{2}{\pi}} \alpha$ for $\alpha \in [0, \frac{1}{2}]$. This is immediate from the fact that

$$\varphi \circ \Phi^{-1}(0) = 0 \quad \text{and} \quad \varphi \circ \Phi^{-1}\left(\frac{1}{2}\right) = \frac{1}{\sqrt{2\pi}},$$

and the concavity of $\varphi \circ \Phi^{-1}$ (see, for example, Exercise 5.43 of [45]). ◀

Proof of Theorem 22. Let r_{in} denote the in-radius of K . We will show that

$$\mathbf{I}[K] \geq \frac{1}{\sqrt{\pi}} \text{Var}[K] \cdot r_{\text{in}} \quad (16)$$

and that

$$r_{\text{in}} \geq \Omega(\sqrt{\ln(\text{Var}[K]/\delta)}) \quad (17)$$

from which the desired result follows.

For Equation (17), by Proposition 13 we have that for some direction $v \in \mathbb{S}^{n-1}$,

$$\text{Inf}_v[K] \geq \frac{\gamma(K)e^{-r_{\text{in}}^2}}{2^{3/2}\pi} \geq \frac{\text{Var}[K]e^{-r_{\text{in}}^2}}{2^{3/2}\pi}.$$

Combining this with $\mathbf{Inf}_\delta[K] \leq \delta$ and recalling that $\delta \leq \mathbf{Var}[K]/10$, we get Equation (17).

We turn now to establishing Equation (16). Recall from Equation (14) of Section 3.3 (our Margulis-Russo formula) that

$$\mathbf{I}[K] = \frac{1}{\sqrt{2}} \lim_{\delta \rightarrow 0} \frac{\gamma(K) - \gamma((1-\delta)K)}{\delta}. \quad (18)$$

We proceed to upper-bound $\gamma((1-\delta)K)$ in terms of $\gamma(K)$. Since r_{in} is the in-radius of K , for all $0 < \delta \leq 1$, we have that

$$(1-\delta)K + \delta r_{\text{in}} B_1 = (1-\delta)K + B_{\delta r_{\text{in}}} \subseteq K. \quad (19)$$

Let $K^c := \mathbb{R} \setminus K$, and let $(K^c)_{\delta r_{\text{in}}} := K^c + B_{\delta r_{\text{in}}}$ be the δr_{in} -enlargement of K^c . It follows from Equation (19) that $(1-\delta)K \cap (K^c)_{\delta r_{\text{in}}} = \emptyset$, which in turn implies that

$$\gamma((1-\delta)K) + \gamma((K^c)_{\delta r_{\text{in}}}) \leq 1, \quad \text{and so} \quad \gamma((1-\delta)K) \leq 1 - \gamma((K^c)_{\delta r_{\text{in}}}). \quad (20)$$

However, from the Gaussian isoperimetric inequality (Proposition 23), we know that

$$\gamma((K^c)_{\delta r_{\text{in}}}) \geq \Phi(\Phi^{-1}(\gamma(K^c)) + \delta r_{\text{in}}). \quad (21)$$

Let $\alpha = \gamma(K^c)$, so $\gamma(K) = 1 - \alpha$. Putting Equations (18), (20), and (21) together, we get

$$\begin{aligned} \mathbf{I}[K] &\geq \frac{1}{\sqrt{2}} \lim_{\delta \rightarrow 0} \frac{\Phi(\Phi^{-1}(\alpha) + \delta r_{\text{in}}) - \alpha}{\delta} \\ &= \frac{1}{\sqrt{2}} r_{\text{in}} \left(\lim_{\varepsilon \rightarrow 0} \frac{\Phi(\Phi^{-1}(\alpha) + \varepsilon) - \Phi(\Phi^{-1}(\alpha))}{\varepsilon} \right) \\ &= \frac{1}{\sqrt{2}} r_{\text{in}} \cdot \Phi'(\Phi^{-1}(\alpha)) \\ &= \frac{1}{\sqrt{2}} r_{\text{in}} \cdot \varphi \circ \Phi^{-1}(\alpha) \end{aligned}$$

by making the change of variables $\varepsilon := \delta r_{\text{in}}$ and using the fact that $\varphi = \Phi'$. It follows then from Proposition 24 that

$$\mathbf{I}[K] \geq \frac{1}{\sqrt{2}} r_{\text{in}} \cdot \left(\sqrt{\frac{2}{\pi}} \min(\alpha, 1-\alpha) \right) \geq \frac{1}{\sqrt{\pi}} \mathbf{Var}[K] \cdot r_{\text{in}}$$

which completes the proof. $\blacktriangleleft$

As discussed in Item 1 of Section 1.3, we conjecture that the RHS of Equation (15) can be strengthened to $\Omega(\mathbf{Var}[K] \log(\frac{1}{\delta}))$, which would be the best possible bound by Example 18.

5 Sharp Threshold Results for Symmetric Convex Sets

For any symmetric convex set $K \subseteq \mathbb{R}^n$, we have that $\gamma_\sigma(K) = \gamma(K/\sigma)$, and hence the map $\Psi_K : \sigma \mapsto \gamma_\sigma(K)$ is a non-increasing function of σ (since $K/\sigma_1 \subseteq K/\sigma_2$ whenever $\sigma_1 \geq \sigma_2$). Given this, it is natural to study the rate of decay of Ψ_K for different symmetric convex sets $K \subseteq \mathbb{R}^n$.

The S -inequality of [37] can be interpreted as saying that the slowest rate of decay across all symmetric convex sets of a given volume is achieved by a symmetric strip. Let K_* be such a strip, i.e. we may take $K_* = \{x \in \mathbb{R}^n : |x_1| \leq c_*\}$ where $c_* = \Theta(\sqrt{\ln(1/\varepsilon)})$ is chosen

so that $\Psi_{K_*}(1) = 1 - \varepsilon$ (and hence $\gamma(K_*) = 1 - \varepsilon$). With this choice of c_* , it follows that $\Psi_{K_*}(\sigma) = \varepsilon$ for $\sigma = \tilde{\Theta}(1/\varepsilon)$. Hence, for the volume of K_* to shrink from $1 - \varepsilon$ to ε , the variance of the underlying Gaussian has to increase very dramatically, by a factor of $\tilde{O}(1/\varepsilon^2)$. Taking, for example, $\varepsilon = 0.01$, we see that in order for the symmetric strip K_* to have its Gaussian volume change from $\gamma_1(K_*) = 0.99$ to $\gamma_\sigma(K_*) = 0.01$, the parameter σ must vary over an interval of size $\Theta(1)$, so the strip K_* does not exhibit a “sharp threshold.”

Of course, as we have seen before, the symmetric strip K_* has an extremely large (constant) convex influence in the direction e_1 . We now show that large individual influences are the only obstacle to sharp thresholds, i.e. any symmetric convex set in which no direction has large convex influence must exhibit a sharp threshold:

► **Theorem 25** (Sharp thresholds for symmetric convex sets with small max influence). *Let $K \subseteq \mathbb{R}^n$ be a centrally symmetric convex set. Suppose $\varepsilon, \delta > 0$ where $\delta \leq \varepsilon^{-10 \log(1/\varepsilon)}$ and $\varepsilon > 0$ is sufficiently small (at most some fixed absolute constant). Suppose that $\gamma(K) \leq 1 - \varepsilon$ and $\max_{v \in \mathbb{S}^{n-1}} [\mathbf{Inf}_v(K)] \leq \delta$. Then, for $\sigma = 1 + \Theta\left(\frac{\ln(1/\varepsilon)}{\sqrt{\ln(\varepsilon/\delta)}}\right)$, we have $\gamma_\sigma(K) \leq \varepsilon$.*

Setting $\varepsilon = 0.01$ and $\delta = o(1)$, the above theorem implies that for K a symmetric convex set K with $\max_{v \in \mathbb{S}^{n-1}} [\mathbf{Inf}_v(K)] = o(1)$, it must be the case that $\gamma_\sigma(K)$ changes from 0.99 to 0.01 as the underlying σ changes from 1 to $1 + o(1)$.

Discussion. Theorem 25 can be seen as a convex influence analogue of a “sharp threshold” result due to Kalai [29]. Building on [25], Kalai [29] shows that if $f : \{\pm 1\}^n \rightarrow \{0, 1\}$ is monotone and its max influence is $o(1)$, then $\mu_p(f)$ must have a sharp threshold (where $\mu_p(f)$ is the expectation of f under the p -biased measure) (see also Theorem 3.8 of [30]). This is closely analogous to Theorem 25, which establishes a sharp threshold for $\gamma_\sigma(K)$ under the assumption that the max convex influence of K is $o(1)$. We note an interesting quantitative distinction between Theorem 25 and the result of [29]: if the max influence of a monotone $f : \{\pm 1\}^n \rightarrow \{0, 1\}$ function is δ , then [29] shows that $\mu_p(f)$ goes from 0.01 to 0.99 in an interval of width $\approx 1/\text{poly}(\log \log(1/\delta))$ (see the discussion following Theorem 3.8 of [30]). In contrast, Theorem 25 shows that $\gamma_\sigma(K)$ goes from 0.01 to 0.99 in an interval of width $\approx 1/\sqrt{\log(1/\delta)}$, thus establishing an exponentially “sharper threshold” in the convex setting.⁶

Proof of Theorem 25. We may assume that $\gamma(K) \geq \varepsilon$, since otherwise, there is nothing to prove. Let $r_{\text{in}} = r_{\text{in}}(K)$ be the in-radius of K . By Proposition 13 we get that

$$r_{\text{in}} \geq \sqrt{\ln\left(\frac{\gamma(K)}{2^{3/2}\pi\delta}\right)} \geq \sqrt{\ln\left(\frac{\varepsilon}{2^{3/2}\pi\delta}\right)} \quad (22)$$

(note that our assumptions on δ and ε imply that the right hand side of (22) is well-defined). Next, we observe that a *mutatis mutandis* modification of the proof of Equation (16) gives that

$$\mathbf{I}^{(\sigma)}[K] \geq \frac{1}{\sqrt{\pi}} \cdot r_{\text{in}} \cdot \mathbf{Var}_\sigma[K]. \quad (23)$$

⁶ Roughly speaking, the extra exponential factor in [29] arises because of Friedgut’s junta theorem; our proof takes a different path and does not incur this quantitative factor.

We further recall that by our Margulis-Russo formula for symmetric convex sets (Proposition 20), we have

$$\frac{d\gamma_\sigma(K)}{d\sigma^2} = -\frac{1}{\sigma^2\sqrt{2}}\mathbf{I}^{(\sigma)}[K]. \quad (24)$$

Combining (22), (23) and (24), we get that

$$\frac{d\gamma_\sigma(K)}{d\sigma^2} \leq -\frac{1}{\sqrt{2\pi}\sigma^2} \cdot \mathbf{Var}_\sigma[K] \cdot \sqrt{\ln\left(\frac{\varepsilon}{2^{3/2}\pi\delta}\right)}.$$

Expressing $\mathbf{Var}_\sigma[K]$ as $\gamma_\sigma(K) \cdot (1 - \gamma_\sigma(K))$ and “solving” the above differential equation for $\gamma_\sigma(K)$, we get that

$$\ln\left(\frac{\gamma_\sigma(K)}{1 - \gamma_\sigma(K)}\right) - \ln\left(\frac{\gamma(K)}{1 - \gamma(K)}\right) \leq \frac{-1}{\sqrt{2\pi}} \cdot \sqrt{\ln\left(\frac{\varepsilon}{2^{3/2}\pi\delta}\right)} \cdot 2 \ln \sigma. \quad (25)$$

Using the assumption that $\gamma(K) \leq 1 - \varepsilon$, it follows that for $\sigma \geq 1$, we have

$$\ln\left(\frac{\gamma_\sigma(K)}{1 - \gamma_\sigma(K)}\right) \leq \ln(1/\varepsilon) + \frac{-1}{\sqrt{2\pi}} \cdot \sqrt{\ln\left(\frac{\varepsilon}{2^{3/2}\pi\delta}\right)} \cdot 2 \ln \sigma.$$

Recalling the assumption that $\delta \leq \varepsilon^{-10 \log(1/\varepsilon)}$, we see that choosing

$$\sigma = 1 + \Theta\left(\frac{\ln(1/\varepsilon)}{\sqrt{\ln(\varepsilon/\delta)}}\right),$$

we get $\gamma_\sigma(K) \leq \varepsilon$ as claimed. $\blacktriangleleft$

► **Remark 26.** We close this section by noting that the type of threshold phenomenon studied here has previously been considered in geometric functional analysis. In particular, the seminal work of Milman [43], using concentration of measure to establish Dvoretzky’s theorem [19] on almost Euclidean sections of symmetric convex sets, implies a type of threshold phenomenon for symmetric convex sets. Milman’s result can be shown to imply that if the “Dvoretzky number” of a symmetric convex set is $\omega_n(1)$, then the set must exhibit a type of sharp threshold behavior. Indeed, Milman’s theorem can be used to give an alternate proof of a result that is similar to Theorem 25.

6 A Robust Kruskal-Katona Analogue for Symmetric Convex Sets

Recall from Equation (4) that for a symmetric convex set $K \subseteq \mathbb{R}^n$, the shell density function $\alpha_K : [0, \infty) \rightarrow [0, 1]$ is defined to be $\alpha_K(r) := \Pr_{\mathbf{x} \in \mathbb{S}_r^{n-1}}[\mathbf{x} \in K]$, and that $\alpha_K(\cdot)$ is non-increasing. In [16], De and Servedio established the following *quantitative* lower bound on the rate of decay of $\alpha_K(\cdot)$:

► **Theorem 27** (Theorem 12 of [16]). *Let $K \subseteq \mathbb{R}^n$ be a convex body that has in-radius $r_{\text{in}} > 0$. Then for $r > r_{\text{in}}$ such that $\min\{\alpha_K(r), (1 - \alpha_K(r))\} \geq e^{-n/4}$, as $\Delta r \rightarrow 0^+$ we have that*

$$\alpha_K(r - \Delta r) - \alpha_K(r) \geq \Omega\left(\frac{r_{\text{in}} \cdot \sqrt{n} \cdot \Delta r}{r^2}\right) \alpha_K(r) (1 - \alpha_K(r)).$$

As alluded to in Item 1 of Section 1, the above result can be used to obtain a Kruskal-Katona type theorem for centrally symmetric convex sets. In particular, we have the following corollary:

► **Corollary 28.** *Let $K \subseteq \mathbb{R}^n$ be a symmetric convex set and $r = \Theta(\sqrt{n})$ be such that $\alpha_K(r) \in [1/10, 9/10]$. Then, as $\varepsilon \rightarrow 0^+$, we have that*

$$\alpha_K(r(1 - \varepsilon)) - \alpha_K(r) = \Omega(\varepsilon).$$

Proof. Let r_{in} denote the in-radius of K , so for any $\zeta > 0$, there is a point z_* such that $z_* \notin K$ and $\|z_*\|_2 = r_{\text{in}} + \zeta$. By the separating hyperplane theorem, it follows that there is a unit vector $\hat{v} \in \mathbb{R}^n$ such that

$$K \subseteq K_* := \{x \in \mathbb{R}^n : |\hat{v} \cdot x| \leq r_{\text{in}} + \zeta\}. \quad (26)$$

We next upper bound $\alpha_{K_*}(r)$. For this, without loss of generality, we may assume that $\hat{v} = e_1$. We have

$$\alpha_{K_*}(r) = \Pr_{y \in \mathbb{S}_r^{n-1}}[|y_1| \leq r_{\text{in}} + \zeta] \leq O\left(\frac{(r_{\text{in}} + \zeta) \cdot \sqrt{n}}{r}\right),$$

where the upper bound is an easy consequence of well-known concentration of measure results for the n -dimensional unit sphere. Now, using (26) and letting $\zeta \rightarrow 0$, we have

$$\alpha_K(r) \leq \alpha_{K_*}(r) \leq O\left(\frac{r_{\text{in}} \cdot \sqrt{n}}{r}\right).$$

Since $\alpha_K(r) \geq 0.1$ by assumption, it follows that $r_{\text{in}} = \Omega(1)$. Corollary 28 now follows from Theorem 27. ◀

A Robust Analogue of Kruskal-Katona. The lower bound given by Corollary 28 cannot be improved in general; for example, the convex set $K = \text{Dict}_{e_1} := \{x : |x_1| \leq 1\}$ satisfies the conditions of Corollary 28 and has

$$\alpha_{\text{Dict}_{e_1}}(r(1 - \varepsilon)) - \alpha_{\text{Dict}_{e_1}}(r) = \Theta(\varepsilon)$$

for $r = \Theta(\sqrt{n})$. This is closely analogous to how the $\Omega(1/n)$ density increment of the original Kruskal-Katona theorem for monotone Boolean functions (recall Item 1 of Section 1) cannot be improved in general because of functions like the Boolean dictator function $f(x) = x_1$. However, if “large single-coordinate influences” are disallowed then stronger forms of the Kruskal-Katona theorem, with larger density increments, hold for monotone Boolean functions. In particular, O’Donnell and Wimmer proved the following “robust” version of the Kruskal-Katona theorem:

► **Theorem 29** (Theorem 1.3 of [44]). *Let $f : \{\pm 1\}^n \rightarrow \{0, 1\}$ be a monotone function and let $n/10 \leq j \leq 9n/10$. If $1/10 \leq \mu_j(f) \leq 9/10$ and it holds for all $i \in [n]$ that*

$$\left| \Pr_{x \sim \binom{[n]}{j}}[f(x) = 1 | x_i = 1] - \Pr_{x \sim \binom{[n]}{j}}[f(x) = 1 | x_i = -1] \right| \leq \frac{1}{n^{1/10}}, \quad (27)$$

then

$$\mu_{j+1}(f) \geq \mu_j(f) + \Omega\left(\frac{\log n}{n}\right).$$

In words, under condition Equation (27) (which is akin to saying that each variable x_i has “low influence on f ”), the much larger density increment $\Omega(\log(n)/n)$ must hold.

Using our notion of convex influences, we now establish a robust version of Corollary 28 which is similar in spirit to the Boolean “robust Kruskal-Katona” result given by Theorem 29. Intuitively, our result says that if all convex influences are small, then we get a stronger density increment than Corollary 28:

► **Theorem 30.** *Let $K \subseteq \mathbb{R}^n$ be a centrally symmetric convex set and $\sqrt{n} \leq r = \Theta(\sqrt{n})$ be such that $\alpha_K(r) \in [1/10, 9/10]$. If $\mathbf{Inf}_v[K] \leq \delta$ for all $v \in \mathbb{S}^{n-1}$ then as $\varepsilon \rightarrow 0^+$ we have that*

$$\alpha_K(r(1 - \varepsilon)) - \alpha_K(r) = \Omega(\varepsilon \sqrt{\ln(1/\delta)}).$$

Proof of Theorem 30. We begin by proving that $\gamma(K) = \Theta(1)$. Note that

$$\gamma(K) = \int_{r=0}^{\infty} \alpha_K(r) \cdot \chi_n(r) dr \geq \int_{r=0}^{\sqrt{n}} \alpha_K(r) \cdot \chi_n(r) dr,$$

where $\chi_n(\cdot)$ is the pdf of the χ -distribution with n -degrees of freedom. Now, since $\alpha_K(\cdot)$ is non-increasing and $\int_{r=0}^{\sqrt{n}} \chi_n(r) dr = \Omega(1)$, it must be the case that

$$\gamma(K) \geq \alpha_K(\sqrt{n}) \cdot \int_{r=0}^{\sqrt{n}} \chi_n(r) dr = \Theta(1), \quad (28)$$

where the last equality uses the fact that $r \geq \sqrt{n}$ and $\alpha_K(r) \geq 1/10$.

Let r_{in} denote the in-radius of K . Exactly as reasoned in the proof of Corollary 28, there exists a unit vector $v \in \mathbb{S}^{n-1}$ such that $K \subseteq \{x \in \mathbb{R}^n : |v \cdot x| \leq r_{\text{in}} + \zeta\}$ for any $\zeta > 0$. Since $\gamma(K) = \Omega(1)$, it now follows from Proposition 15 that there is a direction v such that

$$\mathbf{Inf}_v[K] = \Omega(e^{-(r_{\text{in}} + \zeta)^2}).$$

By our hypothesis, we have that $\mathbf{Inf}_v[K] \leq \delta$, so taking $\zeta \rightarrow 0$ we get that $r_{\text{in}} = \Omega(\sqrt{\ln(1/\delta)})$ (note that we may assume δ is at most some sufficiently small constant, since otherwise the claimed result is given by Corollary 28). We now apply Theorem 27 to obtain that

$$\alpha_K(r(1 - \varepsilon)) - \alpha_K(r) = \Omega(\varepsilon \sqrt{\ln(1/\delta)}),$$

thus proving Theorem 30. ◀

References

- 1 M. Y. An. Log-concave probability distributions: Theory and statistical testing. Technical Report Economics Working Paper Archive at WUSTL, Washington University at St. Louis, 1995.
- 2 S. Artstein-Avidan, A. Giannopoulos, and V.D. Milman. *Asymptotic Geometric Analysis, Part I*, volume 202 of *Mathematical Surveys and Monographs*. American Mathematical Society, 2015.
- 3 W. Beckner. Inequalities in Fourier analysis. *Annals of Mathematics*, 102:159–182, 1975.
- 4 M. Ben-Or and N. Linial. Collective coin flipping. In *Proc. 26th Annual Symposium on Foundations of Computer Science (FOCS)*, pages 408–416, 1985.
- 5 Andrew C. Berry. The Accuracy of the Gaussian Approximation to the Sum of Independent Variates. *Transactions of the American Mathematical Society*, 49(1):122–136, 1941.
- 6 A. Blum, C. Burch, and J. Langford. On learning monotone boolean functions. In *Proceedings of the Thirty-Ninth Annual Symposium on Foundations of Computer Science*, pages 408–415, 1998.

- 7 A. Bonami. Etude des coefficients Fourier des fonctions de $L^p(G)$. *Ann. Inst. Fourier (Grenoble)*, 20(2):335–402, 1970.
- 8 C. Borell. The Brunn-Minkowski inequality in Gauss space. *Invent. Math.*, 30:207–216, 1975.
- 9 C. Borell. The Ehrhard inequality. *C. R. Math. Acad. Sci. Paris*, 337:663–666, 2003.
- 10 C. Borell. Inequalities of the Brunn-Minkowski type for Gaussian measures. *Probab. Theory Related Fields*, 140:195–205, 2008.
- 11 Nader Bshouty and Christino Tamon. On the Fourier spectrum of monotone functions. *Journal of the ACM*, 43(4):747–770, 1996.
- 12 Xi Chen, Adam Freilich, Rocco A Servedio, and Timothy Sun. Sample-Based High-Dimensional Convexity Testing. In *APPROX/RANDOM 2017*, 2017.
- 13 Xi Chen, Erik Waingarten, and Jinyu Xie. 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 2017*, pages 523–536, 2017.
- 14 A. De, S. Nadimpalli, and R. A. Servedio. Quantitative Correlation Inequalities via Semigroup Interpolation. In *12th Innovations in Theoretical Computer Science Conference (ITCS 2021)*, 2021.
- 15 Anindya De, Shivam Nadimpalli, and Rocco A. Servedio. Convex influences. *arXiv preprint*, 2021. [arXiv:2109.03107](https://arxiv.org/abs/2109.03107).
- 16 Anindya De and Rocco A. Servedio. Weak learning convex sets under normal distributions. In *Conference on Learning Theory, COLT 2021*, volume 134 of *Proceedings of Machine Learning Research*, pages 1399–1428. PMLR, 2021.
- 17 A. Dinghas. Über eine Klasse superadditiver Mengenfunktionale von Brunn Minkowski Lusternik-schem Typus. *Math. Zeitschr.*, 68:111–125, 1957.
- 18 Rick Durrett. *Probability: Theory and Examples*. Cambridge Series in Statistical and Probabilistic Mathematics. Cambridge University Press, 5 edition, 2019. doi:10.1017/9781108591034.
- 19 A. Dvoretzky. Some results on convex bodies and Banach spaces. In *Proc. Internat. Sympos. Linear Spaces (Jerusalem, 1960)*, pages 123–160. Jerusalem Academic Press, Jerusalem; Pergamon, Oxford, 1961.
- 20 A. Ehrhard. Symétrisation dans l’espace de gauss. *Math Scand.*, 53:281–301, 1983.
- 21 Ronen Eldan and Renan Gross. Concentration on the Boolean hypercube via pathwise stochastic analysis. In *Proceedings of the 52nd Annual ACM SIGACT Symposium on Theory of Computing, STOC 2020*, pages 208–221. ACM, 2020.
- 22 Carl-Gustav Esseen. On the Liapunoff limit of error in the theory of probability. *Arkiv för matematik, astronomi och fysik*, A:1–19, 1942.
- 23 Alessio Figalli. Personal communication, 2020.
- 24 E. Friedgut. Boolean functions with low average sensitivity depend on few coordinates. *Combinatorica*, 18(1):474–483, 1998.
- 25 E. Friedgut and G. Kalai. Every monotone graph property has a sharp threshold. *Proc. Amer. Math. Soc.*, 124:2993–3002, 1996.
- 26 T.E. Harris. A lower bound for the critical probability in a certain percolation process. *Proc. Camb. Phil. Soc.*, 56:13–20, 1960.
- 27 I.A. Ibragimov. On the composition of unimodal distributions. *Theory Prob. Appl.*, 1, 1956.
- 28 J. Kahn, G. Kalai, and N. Linial. The influence of variables on boolean functions. In *Proc. 29th Annual Symposium on Foundations of Computer Science (FOCS)*, pages 68–80, 1988.
- 29 G. Kalai. Social indeterminacy. *Econometrica*, 72(5):1565–1581, 2004.
- 30 Gil Kalai and Shmuel Safra. Threshold Phenomena and Influence with Some Perspectives from Mathematics, Computer Science, and Economics. In *Computational Complexity and Statistical Physics*, pages 25–60. Oxford University Press, 2005.
- 31 G. Katona. A theorem of finite sets. *Journal of Graph Theory - JGT*, pages 381–401, June 2010. doi:10.1007/978-0-8176-4842-8_27.

- 32 Esty Kelman, Subhash Khot, Guy Kindler, Dor Minzer, and Muli Safra. Theorems of KKL, Friedgut, and Talagrand via Random Restrictions and Log-Sobolev Inequality. In James R. Lee, editor, *12th Innovations in Theoretical Computer Science Conference, ITCS 2021*, volume 185 of *LIPICs*, pages 26:1–26:17, 2021.
- 33 Subhash Khot, Dor Minzer, and Muli Safra. On monotonicity testing and boolean isoperimetric type theorems. In *Proceedings of the 56th Annual Symposium on Foundations of Computer Science*, pages 52–58, 2015.
- 34 Daniel J Kleitman. Families of non-disjoint subsets. *Journal of Combinatorial Theory*, 1(1):153–155, 1966.
- 35 A. Klivans, R. O’Donnell, and R. Servedio. Learning geometric concepts via Gaussian surface area. In *Proc. 49th IEEE Symposium on Foundations of Computer Science (FOCS)*, pages 541–550, 2008.
- 36 J. Kruskal. The number of simplices in a complex. In R. Bellman, editor, *Mathematical Optimization Techniques*, pages 251–278. University of California, Press, Berkeley, 1963.
- 37 Rafał Łatała and Krzysztof Oleszkiewicz. Gaussian measures of dilatations of convex symmetric sets. *Ann. Probab.*, 27(4):1922–1938, October 1999. doi:10.1214/aop/1022874821.
- 38 Rafał Łatała and Krzysztof Oleszkiewicz. Small ball probability estimates in terms of width. *Studia Mathematica*, 169(3):305–314, 2005. doi:10.4064/sm169-3-6.
- 39 L. Leindler. On a certain converse of Hölder’s Inequality II. *Acta Sci. Math. Szeged*, 33:217–223, 1972.
- 40 L. Lovász and S. Vempala. The geometry of logconcave functions and sampling algorithms. *Random Structures and Algorithms*, 30(3):307–358, 2007.
- 41 László Lovász. *Combinatorial problems and exercises*, volume 361. American Mathematical Soc., 1981.
- 42 G. Margulis. Probabilistic characteristics of graphs with large connectivity. *Prob. Peredachi Inform.*, 10:101–108, 1974.
- 43 V.D. Milman. A new proof of A. Dvoretzky’s theorem on cross-sections of convex bodies. *Funkcional. Anal. i Priložen*, 5(4):28–37, 1971.
- 44 R. O’Donnell and K. Wimmer. KKL, Kruskal-Katona, and monotone nets. In *Proc. 50th IEEE Symposium on Foundations of Computer Science (FOCS)*, 2009.
- 45 Ryan O’Donnell. *Analysis of Boolean Functions*. Cambridge University Press, 2014.
- 46 A. Prekopa. Logarithmic concave measures and functions. *Acta Sci. Math. Szeged*, 34:335–343, 1973.
- 47 A. Prekopa. On logarithmic concave measures with applications to stochastic programming. *Acta Sci. Math. Szeged*, 32:301–316, 1973.
- 48 Thomas Royen. A simple proof of the Gaussian correlation conjecture extended to multivariate gamma distributions. *arXiv preprint*, 2014. arXiv:1408.1028.
- 49 L. Russo. On the critical percolation probabilities. *Z. Wahrsch. verw. Gebiete*, 56(2):229–237, 1981.
- 50 M. Talagrand. How much are increasing sets positively correlated? *Combinatorica*, 16(2):243–258, 1996.
- 51 Santosh S. Vempala. Learning convex concepts from gaussian distributions with PCA. In *51th Annual IEEE Symposium on Foundations of Computer Science, FOCS*, pages 124–130. IEEE Computer Society, 2010.
- 52 M. Wainwright. Basic tail and concentration bounds, 2015. URL: www.stat.berkeley.edu/~mhwain/stat210b/Chap2_TailBounds_Jan22_2015.pdf.
- 53 Artem Zvavitch. Personal communication, 2020.