

Topics in High-Dimensional Probability: Functional Inequalities and Concentration of Measure

Arman Jahangiri
Department of Mathematics
University of Washington

May 2026

Abstract

Functional inequalities are among the fundamental tools of modern probability, analysis, and convex geometry. They quantify how functions fluctuate under a measure and reveal deep connections between geometry, concentration of measure, and the behavior of stochastic processes. This study surveys several central functional inequalities, including the Poincaré inequality, the Brunn–Minkowski inequality, the Prékopa–Leindler inequality, and Cheeger’s inequality. We discuss their geometric and probabilistic interpretations, explore their relationships with concentration and isoperimetry, and present a number of applications in high-dimensional probability and convex geometry. A recurring theme is that geometric and analytic inequalities are often different manifestations of the same underlying structure.

1 Introduction

Functional inequalities play a central role in modern analysis, probability theory, and high-dimensional geometry. They provide powerful tools for understanding concentration of measure, convergence of stochastic processes, and stability properties of probability distributions.

In high-dimensional settings, many classical geometric and analytic intuitions break down. Phenomena such as concentration of measure reveal that functions of many variables often exhibit remarkably rigid behavior, with deviations from their mean becoming exponentially unlikely. Functional inequalities provide a quantitative framework for capturing these effects, linking geometric properties of measures to analytic bounds on functions. They are often manifestations of natural physical phenomena as they often express very general laws of nature formulated in physics, biology, economics and various aspects of engineering. They also form the basis of fundamental mathematical structures such as the calculus of variations [3].

One challenge is to establish such inequalities in high-dimensional spaces, where direct methods are often ineffective. One approach to overcome this challenge is localization, which reduces high-dimensional problems to lower dimensions by decomposing measures into simpler components.

The goal of this work is to study some functional inequalities in high-dimensional spaces, including the Cheeger, Poincaré, and logarithmic Sobolev inequalities. We present their formulations, outline methods of proof and highlight the role of localization as a tool in the analysis when applicable.

This work assumes familiarity with basic concepts from real analysis and probability theory. In particular, we assume that the reader is comfortable with measure-theoretic probability, including probability measures on $\mathbb{R}^d$, expectations, and variances, as well as multivariable calculus notions such as gradients and differentiability. Some exposure to L^p spaces and basic functional analysis will be helpful, though not strictly necessary for understanding the main ideas. We will introduce specialized concepts, such as functional inequalities and log-concavity, as they arise.

In Section 2, we introduce the Poincaré inequality, which provides a fundamental tool for controlling the variance of functions in terms of their gradients. We study its formulation, discuss its significance in high-dimensional settings, and present the Gaussian case as an example. We also highlight its connection to concentration of measure through results such as the Gromov–Milman theorem. In Section 3, we turn to the Prékopa–Leindler inequality, a functional form of the Brunn–Minkowski inequality. We prove it using a dimension reduction technique, and show how it serves as a bridge between convex geometry and functional inequalities. As an application, we derive the Brunn–Minkowski inequality and discuss consequences such as isoperimetric-type results. In Section 4, we study logarithmic Sobolev inequalities, which strengthen the Poincaré inequality by controlling entropy rather than variance. We focus on the role of log-concavity and convexity, and outline how the Prékopa–Leindler inequality leads to logarithmic Sobolev inequalities for strongly log-concave measures. We also discuss their implications for concentration and stability.

We begin by providing a few notations and definitions that we will need to use in the following chapters.

1.1 Notations and relevant definitions

Definition 1.1 (Functional Inequality). Typically, a functional inequality has the form

$$\Phi(f) \leq C \Psi(f),$$

where:

- $f : \mathbb{R}^p \rightarrow \mathbb{R}$ is a function (typically in some function space such as L^p or a Sobolev space),
- Φ and Ψ are *functionals*, i.e., mappings that take a function (and possibly its derivatives) as input and return a real number,
- $C > 0$ is a constant independent of f .

Below, we provide the Cauchy–Schwarz Inequality as a simple example of a functional inequality.

Example 1.2 (Cauchy–Schwarz Inequality). Let (Ω, μ) be a measure space and fix a function $g \in L^2(\mu)$. Define

$$\Phi_g(f) = \left(\int_{\Omega} f(x)g(x) d\mu(x) \right)^2, \quad \Psi(f) = \int_{\Omega} f(x)^2 d\mu(x).$$

Then the Cauchy–Schwarz inequality states that

$$\Phi_g(f) \leq C \cdot \Psi(f),$$

where $C = \left(\int_{\Omega} g(x)^2 d\mu(x) \right)$ is an example of a functional inequality.

Definition 1.3 (Probability measure on $\mathbb{R}^d$). A probability measure μ on $\mathbb{R}^d$ is a nonnegative measure defined on the Borel σ -algebra $\mathcal{B}(\mathbb{R}^d)$ such that

$$\mu(\mathbb{R}^d) = 1.$$

Definition 1.4 (Expectation). Let μ be a probability measure on $\mathbb{R}^d$, and let $f : \mathbb{R}^d \rightarrow \mathbb{R}$ be a measurable function such that $\int |f| d\mu < \infty$. The expectation of f with respect to μ is defined by

$$\mathbb{E}_{\mu}[f] = \int_{\mathbb{R}^d} f(x) d\mu(x).$$

Remark 1.5. Throughout, we may remove the subscript μ from $\mathbb{E}_{\mu}[f]$ for simplicity in cases where it does not lead to any confusion.

Remark 1.6. Throughout this document, we use $X_{i:j}$ to denote the random variables $X_i, \dots, X_j$.

2 Poincaré inequality

The Poincaré inequality that we introduce in Definition 2.1, which was named after Henri Poincaré [12], is an important inequality in modern analysis and the theory of Sobolev spaces. It establishes that the magnitude of a function cannot become arbitrarily large relative to the size of its derivatives. It plays a fundamental role in the analysis of partial differential equations, variational methods, geometric analysis, and probability theory. Beyond its analytic significance, the inequality also reveals a strong connection between the geometry of a domain or probability measure and the behavior of functions defined on it.

Definition 2.1 (Poincaré inequality). Let μ be a probability measure on $\mathbb{R}^d$. We say that μ satisfies a Poincaré inequality with constant $C_P(\mu) > 0$ if, for every function f in the Sobolev space $W_0^{1,p}(\Omega)$ that vanishes on the boundary,

$$\text{Var}_{\mu}[f] := \mathbb{E}_{\mu}[(f - \mathbb{E}_{\mu}[f])^2] \leq C_P(\mu) \mathbb{E}_{\mu}[\|\nabla f\|^2]. \quad (1)$$

The smallest such constant is called the *Poincaré constant* of μ .

Remark 2.2. The Poincaré inequality shows that the variance of a function is controlled by its sensitivity to perturbations of the input, as measured by its gradient. This connection is particularly important in high dimensions, where dimension-free bounds on the Poincaré constant lead to strong concentration of measure phenomena. In particular, for Lipschitz functions, the gradient is uniformly bounded, so the Poincaré inequality yields a uniform variance bound. Such variance estimates can then be strengthened into concentration inequalities. Results such as Theorem 2.14 show that a Poincaré inequality implies exponential concentration for Lipschitz functions, meaning that large deviations from the mean become exponentially unlikely.

The following lemma is an important property of the Poincaré inequality, also known as the tensorization property. The significance of this result is that the Poincaré property is preserved under products of independent random variables. This allows one to deduce

higher-dimensional Poincaré inequalities from one-dimensional ones. In particular, we will prove in Subsection 2.1 (Gaussian Poincaré inequality) that since a standard Gaussian vector in $\mathbb{R}^n$ consists of independent one-dimensional Gaussian coordinates, each having Poincaré constant 1, the tensorization property shows that the standard Gaussian measure in arbitrary dimension also has Poincaré constant 1, since the maximum of all ones is one.

Lemma 2.3 (Tensorization of the Poincaré inequality). *Let $X \in \mathbb{R}^n$ and $Y \in \mathbb{R}^m$ be independent random variables satisfying Poincaré inequalities with constants C_P and C'_P , respectively. Then the pair $(X, Y) \in \mathbb{R}^{n+m}$ satisfies a Poincaré inequality with constant*

$$\max\{C_P, C'_P\}.$$

More precisely, for every $f : \mathbb{R}^{n+m} \rightarrow \mathbb{R}$, in the Sobolev space,

$$\text{Var}[f(X, Y)] \leq \max\{C_P, C'_P\} \mathbb{E}[\|\nabla f(X, Y)\|^2].$$

Proof. By the law of total variance,

$$\text{Var}[f(X, Y)] = \mathbb{E}[\text{Var}[f(X, Y) | Y]] + \text{Var}[\mathbb{E}[f(X, Y) | Y]].$$

We bound each term separately. Fix $Y = y$. Then the function

$$x \mapsto f(x, y)$$

depends only on x . Since X satisfies a Poincaré inequality with constant C_P ,

$$\text{Var}_X[f(X, y)] \leq C_P \mathbb{E}_X[(\partial_x f(X, y))^2].$$

Since this inequality holds for every fixed y , taking expectation over Y gives

$$\mathbb{E}[\text{Var}[f(X, Y) | Y]] \leq C_P \mathbb{E}[(\partial_x f(X, Y))^2].$$

Next, define

$$g(Y) := \mathbb{E}[f(X, Y) | Y].$$

Then g is a function only of Y . Since Y satisfies a Poincaré inequality with constant C'_P ,

$$\text{Var}[g(Y)] \leq C'_P \mathbb{E}[g'(Y)^2].$$

Differentiating under the expectation gives

$$g'(Y) = \mathbb{E}[\partial_y f(X, Y) | Y].$$

Applying Jensen's inequality,

$$g'(Y)^2 = (\mathbb{E}[\partial_y f(X, Y) | Y])^2 \leq \mathbb{E}[(\partial_y f(X, Y))^2 | Y].$$

Taking expectation once more yields

$$\mathbb{E}[g'(Y)^2] \leq \mathbb{E}[(\partial_y f(X, Y))^2].$$

Therefore,

$$\text{Var}[\mathbb{E}[f(X, Y) | Y]] = \text{Var}[g(Y)] \leq C'_p \mathbb{E}[(\partial_y f(X, Y))^2].$$

Combining the bounds for the two terms gives

$$\text{Var}[f(X, Y)] \leq C_p \mathbb{E}[(\partial_x f)^2] + C'_p \mathbb{E}[(\partial_y f)^2].$$

Since

$$\|\nabla f\|^2 = (\partial_x f)^2 + (\partial_y f)^2,$$

we obtain

$$\text{Var}[f(X, Y)] \leq \max\{C_p, C'_p\} \mathbb{E}[\|\nabla f(X, Y)\|^2].$$

Thus (X, Y) satisfies a Poincaré inequality with constant

$$\max\{C_p, C'_p\}.$$

□

2.1 Gaussian Poincaré inequality

We aim to prove the Poincaré inequality for Gaussian measures. The proof follows an approach based on Lemma 2.4 provided below.

Lemma 2.4 (Efron-Stein Inequality). *For $1 \leq i \leq j \leq n$, let*

$$\mathbb{E}_{i:j}[f(X)] = \mathbb{E}[f(X) | X_{i:(j-1)}, X_{(j+1):n}]$$

denote the expectation conditioned on all variables except $X_{i:j}$, and let Var_i denote the conditional variance,

$$\text{Var}_i[f(X)] = \mathbb{E}_{i:i}[(f(X) - \mathbb{E}_{i:i} f(X))^2].$$

Then,

$$\text{Var}[f(X)] \leq \mathbb{E} \sum_{i=1}^n \text{Var}_i f(X).$$

Proof. Note that if $f(x) = \sum_{i=1}^n x_i$, we have an exact equality since $X_i - \mathbb{E}X_i$ are orthogonal in L^2 . We aim to bound $\text{Var} f(X)$ by writing the expression $f(X) - \mathbb{E}f(X)$ as the sum of martingale differences, and using the orthogonality of those differences. Let $Z = f(X)$, $Y = Z - \mathbb{E}Z$, and

$$Y_i = \mathbb{E}_{(i+1):n} Z - \mathbb{E}_{i:n} Z$$

for $i = 1, \dots, n$. Consequently, $\mathbb{E}[Y] = 0$, and

$$\begin{aligned} \sum_{i=1}^n Y_i &= \sum_{i=1}^n \mathbb{E}_{(i+1):n}[Z] - \mathbb{E}_{i:n}[Z] \\ &= \sum_{i=1}^n \mathbb{E}_{(i+1):n}[Z] - \sum_{i=1}^n \mathbb{E}_{i:n}[Z] \\ &= \mathbb{E}[Z | X_{1:n}] - \mathbb{E}[Z] \\ &= Z - \mathbb{E}[Z] \\ &= Y, \end{aligned}$$

The third equality follows from the telescoping structure of the sums, since all intermediate terms cancel, leaving only the final term of the first sum and the initial term of the second sum. Hence, we have

$$\text{Var}[Z] = \text{Var}[Y] = \mathbb{E}[Y^2] = \mathbb{E}\left[\left(\sum_{i=1}^n Y_i\right)^2\right] = \sum_{i=1}^n \mathbb{E}[Y_i^2] + \sum_{i \neq j} \mathbb{E}[Y_i Y_j].$$

The second sum simplifies to zero, because, first

$$\begin{aligned} \mathbb{E}[Y_i | X_{1:(i-1)}] &= \mathbb{E}\left\{\left[\mathbb{E}_{(i+1):n}[Z] - \mathbb{E}_{i:n}[Z]\right] \middle| X_{1:(i-1)}\right\} \\ &= \mathbb{E}\left\{\left[\mathbb{E}[Z | X_{1:i}] - \mathbb{E}[Z | X_{1:(i-1)}]\right] \middle| X_{1:(i-1)}\right\} \\ &= \mathbb{E}\left[\mathbb{E}[Z | X_{1:i}] \middle| X_{1:(i-1)}\right] - \mathbb{E}\left[\mathbb{E}[Z | X_{1:(i-1)}] \middle| X_{1:(i-1)}\right] \\ (*) &= \mathbb{E}[Z | X_{1:(i-1)}] - \mathbb{E}[Z | X_{1:(i-1)}] \\ &= 0, \end{aligned}$$

where (*) follows from the tower property for the first term, and the fact that $\mathbb{E}[U|U] = U$ for any random variable U for the second term. Second, since $\{X_{1:j}\} \supseteq \{X_{1:(i-1)}\}$ for any $i > j$, we also have

$$\mathbb{E}[Y_i | X_{1:j}] = 0,$$

for all $i > j$, and thus,

$$\mathbb{E}[Y_i Y_j] = \mathbb{E}[\mathbb{E}[Y_i Y_j | X_{1:j}]] = \mathbb{E}[Y_j \mathbb{E}[Y_i | X_{1:j}]] = 0,$$

for all $i > j$. Similarly, we can condition on $\{X_1, \dots, X_i\}$ to prove the same equality for $i < j$. Therefore, we have proved that $\sum_{i \neq j} \mathbb{E}[Y_i Y_j] = 0$.

For the first sum, note that we have

$$\begin{aligned} Y_i^2 &= (\mathbb{E}[Z | X_{1:i}] - \mathbb{E}[Z | X_{1:(i-1)}])^2 \\ &= (\mathbb{E}[\mathbb{E}[Z | X_{1:n}] | X_{1:i}] - \mathbb{E}[\mathbb{E}[Z | X_{1:(i-1)}, X_{(i+1):n}] | X_{1:i}])^2 \\ &= (\mathbb{E}[\mathbb{E}[Z | X_{1:n}] - \mathbb{E}[Z | X_{1:(i-1)}, X_{(i+1):n}] | X_{1:i}])^2 \\ &= (\mathbb{E}[Z - \mathbb{E}_i Z | X_{1:i}])^2 \\ &\leq \mathbb{E}[(Z - \mathbb{E}_i Z)^2 | X_{1:i}], \end{aligned}$$

where the last inequality follows from Jensen's inequality. Hence, we have

$$\mathbb{E}[Y_i^2] \leq \mathbb{E}[\mathbb{E}[(Z - \mathbb{E}_i Z)^2 | X_{1:i}]] = \mathbb{E}[(Z - \mathbb{E}_i Z)^2] \triangleq \mathbb{E}[\text{Var}_i[Z]].$$

Thus,

$$\text{Var}[Z] = \sum_{i=1}^n \mathbb{E}[Y_i^2] \leq \sum_{i=1}^n \mathbb{E}[\text{Var}_i[Z]] = \mathbb{E}\left[\sum_{i=1}^n \text{Var}_i[Z]\right],$$

which completes the proof. $\square$

We are now ready to prove the Gaussian Poincaré inequality. The proof follows almost the same lines as [1].

Theorem 2.5 (Gaussian Poincaré inequality). *Let $\mu = \mathcal{N}(0, I_d)$ be the standard Gaussian measure on $\mathbb{R}^d$. Then for every $f \in H^1(\mu)$,*

$$\text{Var}_\mu(f) \leq \mathbb{E}_\mu [\|\nabla f\|^2].$$

In particular, the Poincaré constant satisfies

$$C_P(\mathcal{N}(0, I_d)) = 1.$$

Proof. We first prove the result in dimension one for functions $f \in C^2(\mathbb{R})$ with bounded second derivative and sufficient decay, for instance $f \in C_c^2(\mathbb{R})$. Let $\epsilon_1, \dots, \epsilon_n$ be independent Rademacher random variables and define

$$S_n = \frac{1}{\sqrt{n}} \sum_{j=1}^n \epsilon_j.$$

By the central limit theorem, S_n converges in distribution to $X \sim \mathcal{N}(0, 1)$. Since f is smooth and compactly supported, we have

$$\text{Var}(f(S_n)) \rightarrow \text{Var}(f(X)).$$

According to Lemma 2.4 (Efron–Stein inequality),

$$\text{Var}(f(S_n)) \leq \mathbb{E} \left[\sum_{i=1}^n \text{Var}_i(f(S_n)) \right].$$

For each i , conditioning on all variables except ϵ_i , the random variable S_n takes two values differing by $2/\sqrt{n}$. Hence

$$\text{Var}_i(f(S_n)) = \frac{1}{4} \left[f\left(S_n - \frac{\epsilon_i}{\sqrt{n}} + \frac{1}{\sqrt{n}}\right) - f\left(S_n - \frac{\epsilon_i}{\sqrt{n}} - \frac{1}{\sqrt{n}}\right) \right]^2.$$

Therefore,

$$\text{Var}(f(S_n)) \leq \frac{1}{4} \sum_{i=1}^n \mathbb{E} \left[f\left(S_n - \frac{\epsilon_i}{\sqrt{n}} + \frac{1}{\sqrt{n}}\right) - f\left(S_n - \frac{\epsilon_i}{\sqrt{n}} - \frac{1}{\sqrt{n}}\right) \right]^2.$$

Let $K = \sup_x |f''(x)| < \infty$. By Taylor's theorem,

$$\left| f\left(S_n - \frac{\epsilon_i}{\sqrt{n}} + \frac{1}{\sqrt{n}}\right) - f\left(S_n - \frac{\epsilon_i}{\sqrt{n}} - \frac{1}{\sqrt{n}}\right) \right| \leq \frac{2}{\sqrt{n}} |f'(S_n)| + \frac{2K}{n}.$$

Thus,

$$\limsup_{n \rightarrow \infty} \text{Var}(f(S_n)) \leq \mathbb{E}[f'(X)^2].$$

Combining this with $\text{Var}(f(S_n)) \rightarrow \text{Var}(f(X))$, we obtain

$$\text{Var}(f(X)) \leq \mathbb{E}[f'(X)^2].$$

This proves the one-dimensional Gaussian Poincaré inequality for smooth compactly supported functions.

Now let $f : \mathbb{R}^d \rightarrow \mathbb{R}$ be smooth and compactly supported. Applying the one-dimensional inequality conditionally in each coordinate and using the tensorization argument gives

$$\text{Var}_\mu(f) \leq \sum_{i=1}^d \mathbb{E}_\mu \left[\left(\frac{\partial f}{\partial x_i} \right)^2 \right] = \mathbb{E}_\mu [\|\nabla f\|^2].$$

Hence the desired inequality holds for all $f \in C_c^\infty(\mathbb{R}^d)$.

It remains to extend the result to all $f \in H^1(\mu)$. Since $C_c^\infty(\mathbb{R}^d)$ is dense in the Gaussian Sobolev space $H^1(\mu)$, there exists a sequence $f_k \in C_c^\infty(\mathbb{R}^d)$ such that

$$f_k \rightarrow f \quad \text{in } L^2(\mu), \quad \nabla f_k \rightarrow \nabla f \quad \text{in } L^2(\mu).$$

For each k , we have already proved

$$\text{Var}_\mu(f_k) \leq \mathbb{E}_\mu [\|\nabla f_k\|^2].$$

Since $f_k \rightarrow f$ in $L^2(\mu)$, we also have

$$\text{Var}_\mu(f_k) \rightarrow \text{Var}_\mu(f).$$

Moreover,

$$\mathbb{E}_\mu [\|\nabla f_k\|^2] \rightarrow \mathbb{E}_\mu [\|\nabla f\|^2].$$

Passing to the limit in the inequality for f_k gives

$$\text{Var}_\mu(f) \leq \mathbb{E}_\mu [\|\nabla f\|^2].$$

This completes the proof. $\square$

Other examples of distributions that satisfy the Poincaré inequality include log-concave distributions such as the Laplace distribution. We will study this class of distributions in more detail in Section 4.

2.2 Ornstein–Uhlenbeck process

The Gaussian measure possesses additional structure beyond the static formulation of the Poincaré inequality. In particular, many functional inequalities for Gaussian measures can be understood through associated stochastic processes and semigroup methods. One of the most important examples is the Ornstein–Uhlenbeck process, which may be viewed as the canonical diffusion process having the standard Gaussian distribution as its equilibrium measure.

The Ornstein–Uhlenbeck semigroup provides a dynamical perspective on Gaussian functional inequalities and concentration phenomena. In particular, it gives a natural framework for studying the smoothing properties of Gaussian averages and plays a central role in modern proofs of the Gaussian Poincaré and logarithmic Sobolev inequalities. We therefore introduce the process and some of its basic properties.

Definition 2.6 (Ornstein–Uhlenbeck Process). Let $x \in \mathbb{R}^d$. The stochastic process $\{X_t^x\}_{t \geq 0}$ is called an Ornstein–Uhlenbeck process if

$$X_t^x := e^{-t}x + \sqrt{1 - e^{-2t}} G, \quad G \sim \mathcal{N}(0, I). \quad (2)$$

Intuition. At time $t = 0$, we have $X_0^x = x$. As $t \rightarrow \infty$, the process converges to a standard Gaussian random variable G , regardless of the starting point x . Thus, the OU process smoothly interpolates between a deterministic point and a Gaussian random variable. It decomposes into a deterministic part $(\omega_1 x)$ and an independent Gaussian noise term $(\omega_2 G)$ where (ω_1, ω_2) is a unit vector.

Remark 2.7. Note that if the starting point is $G' \sim \mathcal{N}(0, I)$, and if G' is independent of G , then $X_t^{G'} \sim \mathcal{N}(0, I) \equiv G$, meaning that we have a stationary process. This means that if we start the Ornstein–Uhlenbeck process from a Gaussian initial condition, then its distribution does not change over time, since for all $t \geq 0$,

$$X_t^{G'} = e^{-t} G' + \sqrt{1 - e^{-2t}} G \sim \mathcal{N}(0, I),$$

as it is a linear combination of independent zero-mean Gaussians with total variance one: $(e^{-t})^2 + (\sqrt{1 - e^{-2t}})^2 = 1$. Thus, the standard Gaussian distribution is an invariant (equilibrium) distribution for the OU process, and when started from this distribution the process is stationary.

Remark 2.8. The Ornstein–Uhlenbeck process can be defined as the solution to the stochastic differential equation

$$dX_t = -X_t dt + \sqrt{2} dB_t, \quad (3)$$

at time t , where $\{B_t\}_{t \geq 0}$ is a standard Brownian motion. In fact, the expression in (2) is the explicit solution (in distribution) to the stochastic differential equation (3).

Remark 2.9. For all $t > 0$, we have

$$X_t^x \sim \mathcal{N}(e^{-t}x, (1 - e^{-2t})I).$$

For each time t , define an operator P_t acting on functions f by

$$P_t f(x) := \mathbb{E}[f(X_t^x)].$$

So $P_t f(x)$ is the average value of f after evolving the OU process from x for time t . The operators $\{P_t\}_{t \geq 0}$ satisfy the semigroup property $P_{s+t} = P_s P_t$.

Basic properties.

1. $P_0 f(x) = f(x)$.
2. $P_\infty f(x) = \mathbb{E}[f(G)]$.
3. If $G \sim \mathcal{N}(0, I)$, then $\mathbb{E}[P_t f(G)^2] \xrightarrow{t \rightarrow 0} \mathbb{E}[f(G)^2]$.

The intended application is the following. If $\mathbb{E}f(G) = 0$, we have

$$\text{Var}(f(G)) = \mathbb{E}f(G)^2 = \mathbb{E}\left[f(G)(P_0 f(G) - P_\infty f)\right] = \mathbb{E}\left[f(G) \int_0^\infty \partial_t P_t f(G) dt\right]. \quad (4)$$

We look at this last term. Our main claim is that

$$\mathbb{E}[f(G) \cdot \partial_t P_t f(G)] = e^{-t} \mathbb{E}[f'(G)^2].$$

Lemma 2.10 (Gaussian integration by parts). *Let $G \sim \mathcal{N}(0, 1)$, and let $\varphi : \mathbb{R} \rightarrow \mathbb{R}$ be continuously differentiable. Then*

$$\mathbb{E}[G\varphi(G)] = \mathbb{E}[\varphi'(G)]. \quad (5)$$

More generally, if $a, b \in \mathbb{R}$ with $b \neq 0$, then

$$\mathbb{E}[G\varphi(a + bG)] = b \mathbb{E}[\varphi'(a + bG)].$$

Proof. Since the density of G is

$$\gamma(x) = \frac{1}{\sqrt{2\pi}} e^{-x^2/2},$$

we have $\gamma'(x) = -x\gamma(x)$. Hence, assuming the boundary term vanishes,

$$\mathbb{E}[G\varphi(G)] = \int_{\mathbb{R}} x\varphi(x)\gamma(x) dx = - \int_{\mathbb{R}} \varphi(x)\gamma'(x) dx = \int_{\mathbb{R}} \varphi'(x)\gamma(x) dx = \mathbb{E}[\varphi'(G)].$$

Taking $\psi(G) = \varphi(a + bG)$ in (5) gives

$$\mathbb{E}[G\varphi(a + bG)] = \mathbb{E}[\psi'(G)] = b \mathbb{E}[\varphi'(a + bG)],$$

which proves the lemma. □

Claim 2.11. *Define the Ornstein–Uhlenbeck operator in one dimension,*

$$Lf(x) := -xf'(x) + f''(x).$$

Then as operators, $\partial_t P_t = LP_t$.

Proof. Recall that

$$P_t f(x) = \mathbb{E}\left[f\left(e^{-t}x + \sqrt{1 - e^{-2t}} G\right)\right],$$

where $G \sim \mathcal{N}(0, 1)$. Let

$$a_t = e^{-t}, \quad b_t = \sqrt{1 - e^{-2t}}.$$

Then

$$P_t f(x) = \mathbb{E}[f(a_t x + b_t G)].$$

Differentiating with respect to t gives

$$\partial_t P_t f(x) = \mathbb{E}[f'(a_t x + b_t G)(a'_t x + b'_t G)].$$

Since

$$a'_t = -a_t, \quad b'_t = \frac{e^{-2t}}{\sqrt{1 - e^{-2t}}} = \frac{a_t^2}{b_t},$$

we obtain

$$\partial_t P_t f(x) = -a_t x \mathbb{E}[f'(a_t x + b_t G)] + \frac{a_t^2}{b_t} \mathbb{E}[G f'(a_t x + b_t G)].$$

By Gaussian integration by parts (2.10),

$$\mathbb{E}[G f'(a_t x + b_t G)] = b_t \mathbb{E}[f''(a_t x + b_t G)].$$

Therefore,

$$\partial_t P_t f(x) = -a_t x \mathbb{E}[f'(a_t x + b_t G)] + a_t^2 \mathbb{E}[f''(a_t x + b_t G)].$$

On the other hand,

$$\partial_x P_t f(x) = a_t \mathbb{E}[f'(a_t x + b_t G)],$$

and

$$\partial_{xx} P_t f(x) = a_t^2 \mathbb{E}[f''(a_t x + b_t G)].$$

Hence

$$\partial_t P_t f(x) = -x \partial_x P_t f(x) + \partial_{xx} P_t f(x) = L P_t f(x).$$

This proves the claim. $\square$

Claim 2.12. For sufficiently smooth functions h and g ,

$$\mathbb{E}[h(G)Lg(G)] = -\mathbb{E}[h'(G)g'(G)].$$

Proof. By definition,

$$Lg(x) = -xg'(x) + g''(x).$$

Therefore,

$$\mathbb{E}[h(G)Lg(G)] = -\mathbb{E}[h(G)Gg'(G)] + \mathbb{E}[h(G)g''(G)].$$

Using Gaussian integration by parts (2.10) with the function

$$\varphi(x) = h(x)g'(x),$$

we get

$$\mathbb{E}[Gh(G)g'(G)] = \mathbb{E}[\varphi'(G)].$$

By the Leibniz rule,

$$\varphi'(x) = h'(x)g'(x) + h(x)g''(x).$$

Hence

$$\mathbb{E}[Gh(G)g'(G)] = \mathbb{E}[h'(G)g'(G)] + \mathbb{E}[h(G)g''(G)].$$

Substituting this into the expression for $\mathbb{E}[h(G)Lg(G)]$ gives

$$\mathbb{E}[h(G)Lg(G)] = -\mathbb{E}[h'(G)g'(G)].$$

$\square$

Claim 2.13. For every sufficiently smooth function f ,

$$\partial_x P_t f(x) = e^{-t} P_t(f')(x).$$

Proof. Using the definition of the Ornstein–Uhlenbeck semigroup,

$$P_t f(x) = \mathbb{E}\left[f\left(e^{-t}x + \sqrt{1 - e^{-2t}} G\right)\right].$$

Differentiating with respect to x and applying the chain rule gives

$$\partial_x P_t f(x) = \mathbb{E}\left[f'\left(e^{-t}x + \sqrt{1 - e^{-2t}} G\right)e^{-t}\right].$$

Thus,

$$\partial_x P_t f(x) = e^{-t} \mathbb{E} \left[f' \left(e^{-t} x + \sqrt{1 - e^{-2t}} G \right) \right].$$

But the expectation on the right-hand side is exactly $P_t(f')(x)$. Therefore,

$$\partial_x P_t f(x) = e^{-t} P_t(f')(x).$$

□

Finally,

$$\begin{aligned} \mathbb{E}[f(G) \cdot \partial_t P_t f(G)] &= \mathbb{E}[f(G) L P_t f(G)] \\ \text{Claim (2.12)} &= \mathbb{E}[f'(G) \cdot (P_t f)'(G)] \\ \text{Claim (2.13)} &= e^{-t} \mathbb{E}[f'(G) \cdot P_t(f'(G))] \\ (\text{Cauchy-Schwarz inequality}) &\leq e^{-t} (\mathbb{E} f'(G)^2)^{1/2} (\mathbb{E} (P_t f'(G))^2)^{1/2} \\ (\text{third property of } P_t) &\rightarrow e^{-t} \mathbb{E} f'(G)^2. \end{aligned}$$

Integrating over $t \in [0, \infty)$, according to (4), gives

$$\text{Var}(f(G)) = \int_0^\infty \mathbb{E}[f(G) \cdot \partial_t P_t f(G)] dt \leq \int_0^\infty e^{-t} \mathbb{E}[f'(G)^2] dt = \mathbb{E}[f'(G)^2].$$

Thus we obtain the Poincaré inequality for the standard Gaussian, with $C_P = 1$, i.e.,

$$\text{Var}(f(G)) \leq \mathbb{E}[f'(G)^2].$$

2.3 Gromov-Milman theorem

Theorem 2.14 ([4]). *Suppose $X \sim \mu$ on $\mathbb{R}^d$ satisfies a Poincaré inequality with constant $C_P(\mu) = 1$. Let $f : \mathbb{R}^d \rightarrow \mathbb{R}$ be 1-Lipschitz and centered, meaning $\mathbb{E}[f(X)] = 0$. Then there exist constants $C, c > 0$ such that, for all $\lambda > 0$,*

$$\mathbb{P}(|f(X)| > \lambda) \leq C e^{-c\lambda}.$$

Remark 2.15. The assumptions $\mathbb{E}[f(X)] = 0$ and $C_P(\mu) = 1$ is only a normalization. If the Poincaré constant is $C_P(\mu)$, the same proof gives

$$\mathbb{P}(|f(X) - \mathbb{E}f(X)| > \lambda) \leq C \exp \left(-c \frac{\lambda}{\sqrt{C_P(\mu)}} \right).$$

Proof. The idea is to use the Poincaré inequality to control exponential moments of f , and then apply Markov's inequality.

Fix $\alpha \in (0, 2)$ and define

$$g(x) := e^{\frac{\alpha}{2} f(x)}.$$

By the Poincaré inequality,

$$\text{Var}(g(X)) \leq \mathbb{E} \|\nabla g(X)\|^2.$$

Since f is 1-Lipschitz, we have $\|\nabla f(x)\| \leq 1$ almost everywhere. Hence

$$\|\nabla g(x)\|^2 = \left\| \frac{\alpha}{2} e^{\frac{\alpha}{2} f(x)} \nabla f(x) \right\|^2 = \frac{\alpha^2}{4} e^{\alpha f(x)} \|\nabla f(x)\|^2 \leq \frac{\alpha^2}{4} e^{\alpha f(x)}.$$

Therefore,

$$\text{Var}\left(e^{\frac{\alpha}{2} f(X)}\right) \leq \frac{\alpha^2}{4} \mathbb{E}[e^{\alpha f(X)}].$$

Using

$$\text{Var}(Y) = \mathbb{E}[Y^2] - (\mathbb{E}Y)^2$$

with $Y = e^{\frac{\alpha}{2} f(X)}$, we get

$$\mathbb{E}[e^{\alpha f(X)}] - \left(\mathbb{E}[e^{\frac{\alpha}{2} f(X)}]\right)^2 \leq \frac{\alpha^2}{4} \mathbb{E}[e^{\alpha f(X)}].$$

Rearranging gives

$$\mathbb{E}[e^{\alpha f(X)}] \leq \frac{1}{1 - \frac{\alpha^2}{4}} \left(\mathbb{E}[e^{\frac{\alpha}{2} f(X)}]\right)^2.$$

Iterating this inequality m times gives

$$\mathbb{E}[e^{\alpha f(X)}] \leq \prod_{k=1}^m \left(\frac{1}{1 - \frac{\alpha^2}{4^k}} \right)^{2^{k-1}} \left(\mathbb{E}[e^{\frac{\alpha}{2^m} f(X)}] \right)^{2^m}.$$

Let $u_m = \alpha/2^m$. Since $u_m \rightarrow 0$, we have

$$\mathbb{E}[e^{u_m f(X)}] = 1 + u_m \mathbb{E}[f(X)] + o(u_m).$$

Because $\mathbb{E}[f(X)] = 0$, this becomes

$$\mathbb{E}[e^{u_m f(X)}] = 1 + o(u_m).$$

Since $u_m = \alpha 2^{-m}$, we conclude that

$$\mathbb{E}[e^{\frac{\alpha}{2^m} f(X)}] = 1 + o(2^{-m}).$$

and therefore

$$\left(\mathbb{E}[e^{\frac{\alpha}{2^m} f(X)}] \right)^{2^m} \longrightarrow 1 \quad \text{as } m \rightarrow \infty.$$

Thus

$$\mathbb{E}[e^{\alpha f(X)}] \leq \prod_{k=1}^{\infty} \left(\frac{1}{1 - \frac{\alpha^2}{4^k}} \right)^{2^{k-1}} =: K_{\alpha} < \infty.$$

Now choose any fixed $\alpha \in (0, 2)$, for example $\alpha = 1$. By Markov's inequality,

$$\mathbb{P}(f(X) > \lambda) = \mathbb{P}(e^{\alpha f(X)} > e^{\alpha \lambda}) \leq e^{-\alpha \lambda} \mathbb{E}[e^{\alpha f(X)}] \leq K_{\alpha} e^{-\alpha \lambda}.$$

Applying the same argument to $-f$, which is also centered and 1-Lipschitz, we obtain

$$\mathbb{P}(f(X) < -\lambda) \leq K_{\alpha} e^{-\alpha \lambda}.$$

Therefore,

$$\mathbb{P}(|f(X)| > \lambda) \leq 2K_{\alpha} e^{-\alpha \lambda}.$$

This proves the claim with $C = 2K_{\alpha}$ and $c = \alpha$. $\square$

2.4 KLS conjecture

The Kannan–Lovász–Simonovits (KLS) conjecture is one of the main open problems in modern convex geometry. Broadly speaking, it predicts that log-concave probability distributions on convex bodies have strong “connectivity” properties, meaning that the measure cannot concentrate around narrow “bottlenecks”. We will study this in more depth in this section. This conjecture is also connected to functional inequalities, since good geometric connectivity leads to strong variance bounds, concentration of measure, and rapid mixing of random processes.

To introduce the KLS conjecture, we first need to understand the definition of conductance score.

Definition 2.16 (Conductance score). Let $K \subset \mathbb{R}^n$ be a convex body in $\mathbb{R}^n$ with positive volume. The conductance score of K is defined as

$$\psi_K := \inf_{\substack{S \subseteq K \\ 0 < \text{vol}_n(S) < \text{vol}_n(K)}} \frac{\text{Vol}_{n-1}(\partial S \cap \text{int}(K))}{\min\{\text{vol}_n(S), \text{vol}_n(K \setminus S)\}}.$$

Intuitively, ψ_K measures the *bottleneck* of K : it is the smallest possible boundary “area” (an $(n - 1)$ -dimensional measure) separating K into two parts, normalized by the volume of the smaller part.

Interpretation.

- $\psi_K = 0$ indicates that K is *disconnected* (or can be separated by a cut of zero $(n - 1)$ -dimensional measure).
- Small ψ_K suggests that K has a *thin neck*: one can split K into two relatively large pieces while crossing only a small boundary area.
- Large ψ_K suggests K has *no bottlenecks*: any partition of K into two sets of nontrivial volume must cut through a large boundary area, so K is well-connected in an isoperimetric sense.

A small value of ψ_K indicates the existence of a subset S that can be separated from its complement by a boundary of relatively small $(n - 1)$ -dimensional measure compared to its volume. Such configurations correspond to narrow regions or bottlenecks in K .

Denote the barycenter of K by

$$b(K) = \mathbb{E}[X],$$

where $X \sim \text{Unif}(K)$. Let $A(K)$ be the $n \times n$ inertia matrix of K about its barycenter:

$$A(K) = \mathbb{E}[(X - b(K))(X - b(K))^T], \quad (6)$$

and let $a(K)$ denote the largest eigenvalue of $A(K)$. The p th moment ($p \geq 0$) of a convex body K is defined by

$$M_p(K) = \mathbb{E}[\|X - b(K)\|^p].$$

Theorem 2.17 (Moment bound [5]). *For every convex body K ,*

$$\psi_K \geq \frac{\ln 2}{M_1(K)}.$$

Thus, if the mass of K is concentrated near its center (i.e., $M_1(K)$ is small), the body has large conductance and no significant bottlenecks. Conversely, a large first moment indicates weaker isoperimetry and smaller conductance.

Now, for each $x \in K$, let $\chi(x)$ denote the length of the longest segment contained in K with midpoint x . (Equivalently, $\chi(x)$ is the diameter of $K \cap (2x - K)$.) Define the average chord length

$$\chi(K) = \frac{1}{\text{vol}(K)} \int_K \chi(x) dx.$$

Theorem 2.18 (Chord bound [5]). *For every convex body K ,*

$$\psi(K) \geq \frac{1}{\chi(K)}.$$

This theorem gives a lower bound on the conductance of K in terms of its average chord length. In particular, when the average chord length $\chi(K)$ is small, the lower bound $1/\chi(K)$ is large, forcing the conductance $\psi(K)$ to be large as well. Thus, convex bodies with small average chords cannot have severe bottlenecks in the sense measured by $\psi(K)$. Conversely, large average chords make this lower bound weaker, although they do not by themselves imply small conductance.

Conjecture 2.19 (Kannan–Lovász–Simonovits (KLS)). *For every convex body $K \subset \mathbb{R}^n$,*

$$\psi(K) = \Theta\left(\frac{1}{\sqrt{a(K)}}\right).$$

The KLS conjecture asserts that the conductance is, up to constants, governed by the square root of the largest eigenvalue of the covariance matrix.

Theorem 2.20 (Upper bound [5]). *For every convex body $K \subset \mathbb{R}^n$,*

$$\psi(K) \leq \frac{10}{\sqrt{a(K)}},$$

where $a(K)$ is the largest eigenvalue (spectral norm) of the inertia matrix, defined in (6),

We now specialize the previous results to convex bodies in isotropic position. This normalization simplifies the geometric parameters appearing in the KLS conjecture and allows the bounds to be expressed purely in terms of the dimension. In particular, since isotropic position forces the covariance matrix to be the identity, the quantity $a(K)$ appearing in the general KLS conjecture becomes equal to 1.

Definition 2.21 (Isotropic position). A convex body $K \subset \mathbb{R}^n$ is said to be in *isotropic position* if the uniform probability measure on K has barycenter at the origin and identity covariance matrix. That is, if

$$X \sim \text{Unif}(K),$$

then

$$\mathbb{E}[X] = 0, \quad \mathbb{E}[XX^\top] = I_n.$$

Equivalently,

$$\text{Cov}(X) = I_n.$$

Remark 2.22. Being in Isotropic Position implies that the largest eigenvalue of the covariance matrix satisfies $a(K) = 1$.

The next lemma shows how Theorem 2.17 reduces to a dimension-dependent estimate in isotropic position.

Lemma 2.23. *If $K \subset \mathbb{R}^n$ is isotropic, then*

$$\psi(K) \geq \frac{\ln 2}{\sqrt{n}}.$$

Proof. Note that since $b(K) = 0$, we have

$$M_1(K) = \mathbb{E}\|X - b(K)\|^1 = \mathbb{E}\|X\| \quad (7)$$

Additionally, for any vector $v \in \mathbb{R}^n$, we have Using the identity $v^\top v = \text{tr}(vv^\top)$ for any vector v . Therefore, according to the linearity of expectation and trace,

$$\begin{aligned} \mathbb{E}\|X - \mu\|^2 &= \mathbb{E}[(X - \mu)^\top (X - \mu)] \\ &= \mathbb{E} \text{tr}((X - \mu)(X - \mu)^\top) \\ &= \text{tr}\left(\mathbb{E}[(X - \mu)(X - \mu)^\top]\right) \\ &= \text{tr}(\text{Cov}(X)). \end{aligned}$$

So in our case where $\mathbb{E}[X] = 0$ and $\text{Cov}(X) = I_n$, we have

$$\mathbb{E}\|X\|^2 = \text{tr}(I_n) = n$$

and thus, according to Jensen's inequality,

$$\mathbb{E}\|X\| = \mathbb{E}\left(\sqrt{\|X\|^2}\right) \leq \sqrt{\mathbb{E}\|X\|^2} = \sqrt{n}. \quad (8)$$

Now, applying Theorem 2.17, we have

$$\begin{aligned} \psi(K) &\geq \frac{\ln 2}{M_1(K)} \\ &= \frac{\ln 2}{\mathbb{E}\|X\|} \quad \text{by (7)} \\ &\geq \frac{\ln 2}{\sqrt{n}} \quad \text{by (8)} \end{aligned}$$

□

Similarly, under isotropic normalization, Conjecture 2.19 (KLS) takes the simpler form below.

Lemma 2.24. *If K is isotropic, then*

$$\psi(K) \geq c$$

for some absolute constant $c > 0$ independent of the dimension. Equivalently, since $a(K) = 1$ in isotropic position, the general conjecture

$$\psi(K) = \Theta\left(\frac{1}{\sqrt{a(K)}}\right)$$

predicts that isotropic convex bodies have conductance bounded below by a universal constant.

Equivalently, the general conjecture $\psi(K) = \Theta(1/\sqrt{a(K)})$ predicts $\psi(K) = \Theta(1)$ when $a(K) = 1$.

3 Prékopa–Leindler theorem

The Prékopa–Leindler inequality is an integral inequality that has an important role in analysis, named after András Prékopa and László Leindler. It plays a role in the study of log-concave measures, and has important applications in probability theory, convex geometry, and functional inequalities. For example, it can be used to prove that log-concavity is preserved by marginalization, which we will prove in Section 4.

Additionally, in Section 3.3, as an application of a more geometric nature, we will show how Brunn–Minkowski inequality leads to the classical isoperimetric inequality.

3.1 Prékopa–Leindler theorem

Theorem 3.1 (Prékopa–Leindler). *Let $f, g, h : \mathbb{R}^n \rightarrow [0, \infty)$ be nonnegative measurable functions, and let $t \in (0, 1)$. Assume that for all $x, y \in \mathbb{R}^n$,*

$$h(tx + (1-t)y) \geq f(x)^t g(y)^{1-t}.$$

Then

$$\int_{\mathbb{R}^n} h(x) dx \geq \left(\int_{\mathbb{R}^n} f(x) dx \right)^t \left(\int_{\mathbb{R}^n} g(x) dx \right)^{1-t},$$

i.e.,

$$\|h\|_1 \geq \|f\|_1^{1-\lambda} \|g\|_1^\lambda$$

Proof (induction). To prove the base case ($n=1$), we first consider the normalized case (f and g being density functions). Assume

$$\int_{\mathbb{R}} f = \int_{\mathbb{R}} g = 1.$$

Hence, It suffices to show that $\int_{\mathbb{R}} h \geq 1$.

Define the distribution functions

$$F(x) := \int_{-\infty}^x f(s) ds, \quad G(y) := \int_{-\infty}^y g(s) ds.$$

Let

$$T(x) := G^{-1}(F(x)).$$

Then T is increasing and differentiable almost everywhere. Furthermore, T transports the density f to the density g , in the sense that the pushforward of the measure μ_f under T equals μ_g . More precisely, for every measurable set $A \subseteq \mathbb{R}$,

$$\mu_f(T^{-1}(A)) = \int_{T^{-1}(A)} f(x) dx = \int_A g(y) dy = \mu_g(A), \quad (9)$$

where μ_f and μ_g denote the probability measures on $\mathbb{R}^n$ with densities f and g , respectively. Using the change of variable $y = T(x)$, we have

$$\int_A g(y) dy = \int_{T^{-1}(A)} g(T(x))T'(x) dx. \quad (10)$$

Hence, since A is arbitrary, (9) and (10) imply that

$$f(x) = g(T(x))T'(x) \quad \text{for a.e. } x.$$

Equivalently,

$$g(T(x)) = \frac{f(x)}{T'(x)}.$$

Now define

$$z(x) := (1 - \lambda)x + \lambda T(x).$$

Since T is increasing, z is also increasing. Therefore, using the change of variables $z = (1 - \lambda)x + \lambda T(x)$, we get

$$\int_{\mathbb{R}} h(z) dz = \int_{\mathbb{R}} h((1 - \lambda)x + \lambda T(x))((1 - \lambda) + \lambda T'(x)) dx.$$

By the hypothesis,

$$h((1 - \lambda)x + \lambda T(x)) \geq f(x)^{1-\lambda} g(T(x))^\lambda.$$

Therefore,

$$\begin{aligned} \int_{\mathbb{R}} h(z) dz &\geq \int_{\mathbb{R}} f(x)^{1-\lambda} g(T(x))^\lambda ((1 - \lambda) + \lambda T'(x)) dx \\ &= \int_{\mathbb{R}} f(x)^{1-\lambda} \left(\frac{f(x)}{T'(x)} \right)^\lambda ((1 - \lambda) + \lambda T'(x)) dx \\ &= \int_{\mathbb{R}} f(x) \frac{(1 - \lambda) + \lambda T'(x)}{(T'(x))^\lambda} dx. \end{aligned}$$

By the weighted AM–GM inequality,

$$(1 - \lambda) + \lambda T'(x) \geq (T'(x))^\lambda.$$

Hence

$$\frac{(1 - \lambda) + \lambda T'(x)}{(T'(x))^\lambda} \geq 1.$$

Thus

$$\int_{\mathbb{R}} h(z) dz \geq \int_{\mathbb{R}} f(x) dx = 1.$$

Finally, for the general case, set

$$A := \int_{\mathbb{R}} f, \quad B := \int_{\mathbb{R}} g.$$

If $A = 0$ or $B = 0$, the result is trivial. Otherwise apply the normalized case to f/A , g/B , and $h/(A^{1-\lambda}B^\lambda)$. This gives

$$\int_{\mathbb{R}} h \geq A^{1-\lambda}B^\lambda = \left(\int_{\mathbb{R}} f \right)^{1-\lambda} \left(\int_{\mathbb{R}} g \right)^\lambda.$$

To prove the inductive step, suppose the Prékopa–Leindler inequality holds in dimension n . We prove it in dimension $n + 1$. Write points in $\mathbb{R}^{n+1}$ as (x, r) , where $x \in \mathbb{R}^n$ and $r \in \mathbb{R}$. Let $f, g, h : \mathbb{R}^{n+1} \rightarrow [0, \infty)$ satisfy

$$h(t(x, r) + (1-t)(y, s)) \geq f(x, r)^t g(y, s)^{1-t}$$

for all $(x, r), (y, s) \in \mathbb{R}^{n+1}$. Fix $x, y \in \mathbb{R}^n$ and set

$$z = tx + (1-t)y.$$

Define one-dimensional functions

$$F_x(r) := f(x, r), \quad G_y(s) := g(y, s), \quad H_z(u) := h(z, u).$$

Then for all $r, s \in \mathbb{R}$,

$$H_z(tr + (1-t)s) = h(z, tr + (1-t)s) \geq f(x, r)^t g(y, s)^{1-t} = F_x(r)^t G_y(s)^{1-t}.$$

By the one-dimensional Prékopa–Leindler inequality,

$$\int_{\mathbb{R}} H_z(u) du \geq \left(\int_{\mathbb{R}} F_x(r) dr \right)^t \left(\int_{\mathbb{R}} G_y(s) ds \right)^{1-t}.$$

Now, define the marginals on $\mathbb{R}^n$ by

$$\widehat{h}(z) := \int_{\mathbb{R}} h(z, u) du, \quad \widehat{f}(x) := \int_{\mathbb{R}} f(x, r) dr, \quad \widehat{g}(y) := \int_{\mathbb{R}} g(y, s) ds.$$

The previous inequality asserts that, for all $x, y \in \mathbb{R}^n$,

$$\widehat{h}(tx + (1-t)y) \geq \widehat{f}(x)^t \widehat{g}(y)^{1-t}.$$

By the induction hypothesis in dimension n ,

$$\int_{\mathbb{R}^n} \widehat{h} \geq \left(\int_{\mathbb{R}^n} \widehat{f} \right)^t \left(\int_{\mathbb{R}^n} \widehat{g} \right)^{1-t}.$$

Using Tonelli's theorem, since all functions are nonnegative,

$$\int_{\mathbb{R}^n} \widehat{h} = \int_{\mathbb{R}^{n+1}} h, \quad \int_{\mathbb{R}^n} \widehat{f} = \int_{\mathbb{R}^{n+1}} f, \quad \int_{\mathbb{R}^n} \widehat{g} = \int_{\mathbb{R}^{n+1}} g.$$

Therefore,

$$\int_{\mathbb{R}^{n+1}} h \geq \left(\int_{\mathbb{R}^{n+1}} f \right)^t \left(\int_{\mathbb{R}^{n+1}} g \right)^{1-t}.$$

This proves the result in dimension $n + 1$, and hence, by induction, in all dimensions. $\square$

Remark 3.2. The Prékopa–Leindler inequality can be viewed as a functional inequality when reduced to special cases. Indeed, by applying Prékopa–Leindler inequality to indicator functions of measurable sets, one recovers a special case, which is also known as the Brunn–Minkowski inequality. We provide a proof of this connection below.

3.2 Brunn–Minkowski inequality

Theorem 3.3 (Brunn–Minkowski inequality). *Let $S, T \subset \mathbb{R}^n$ and their Minkowski sum*

$$S + T = \{x + y : x \in S, y \in T\}$$

be measurable sets. Then

$$\text{Vol}(S + T)^{1/n} \geq \text{Vol}(S)^{1/n} + \text{Vol}(T)^{1/n}.$$

Proof. Let $A, B \subset \mathbb{R}^n$ be measurable sets and fix $t \in (0, 1)$. Define

$$f = \mathbf{1}_A, \quad g = \mathbf{1}_B, \quad h = \mathbf{1}_{tA + (1-t)B}.$$

We claim that for all $x, y \in \mathbb{R}^n$,

$$h(tx + (1-t)y) \geq f(x)^t g(y)^{1-t}.$$

Indeed, if $x \in A$ and $y \in B$, then $tx + (1-t)y \in tA + (1-t)B$, so both sides equal 1. Otherwise, at least one of $f(x)$ or $g(y)$ is zero, hence the right-hand side is zero and the inequality still holds. Therefore, the assumptions of the Prékopa–Leindler inequality are satisfied, and we obtain

$$\text{Vol}(tA + (1-t)B) \geq \text{Vol}(A)^t \text{Vol}(B)^{1-t}.$$

Now, using the scaling property of Lebesgue measure,

$$\text{Vol}(\lambda E) = \lambda^n \text{Vol}(E), \quad \lambda > 0,$$

we write

$$tA + (1-t)B = t \left(A + \frac{1-t}{t} B \right),$$

and hence

$$\text{Vol}(tA + (1-t)B) = t^n \text{Vol} \left(A + \frac{1-t}{t} B \right).$$

Substituting into the previous inequality gives

$$\text{Vol}\left(A + \frac{1-t}{t}B\right) \geq t^{-n} \text{Vol}(A)^t \text{Vol}(B)^{1-t}.$$

Now choose

$$t = \frac{\text{Vol}(A)^{1/n}}{\text{Vol}(A)^{1/n} + \text{Vol}(B)^{1/n}}.$$

A straightforward computation then yields

$$\text{Vol}(A + B)^{1/n} \geq \text{Vol}(A)^{1/n} + \text{Vol}(B)^{1/n},$$

which completes the proof. $\square$

Remark 3.4. Brun-Minkowski inequality asserts that for Minkowski combinations

$$A_t = (1-t)A + tB,$$

the function

$$t \rightarrow \text{Vol}(A_t)^{1/n}$$

is concave.

One of the consequences of the Prékopa–Leindler inequality is its connection to log-concavity. In particular, it implies that marginals of log-concave functions are log-concave. Moreover, the inequality serves as a cornerstone for many important functional inequalities, including logarithmic Sobolev and transportation inequalities.

3.3 Isoperimetric inequality

Let $A \subset \mathbb{R}^n$ be measurable and let $\bar{B}$ be a Euclidean ball with

$$\text{Vol}(A) = \text{Vol}(\bar{B}).$$

Define

$$A_\varepsilon := A + \varepsilon B_n,$$

where B_n is the unit ball in $\mathbb{R}^n$. By Brunn–Minkowski inequality,

$$\text{Vol}(A_\varepsilon)^{1/n} \geq \text{Vol}(A)^{1/n} + \text{Vol}(\varepsilon B_n)^{1/n} = \text{Vol}(\bar{B})^{1/n} + \text{Vol}(\varepsilon B_n)^{1/n} = \text{Vol}(\bar{B}_\varepsilon)^{1/n}.$$

where the last equality follows from the fact that for Euclidean balls, Minkowski addition corresponds to adding radii. Therefore, we have proved that

$$\text{Vol}(A_\varepsilon) \geq \text{Vol}(\bar{B}_\varepsilon), \tag{11}$$

Thus, among all sets of fixed volume, Euclidean balls minimize the growth of volume under Minkowski enlargement. This is an equivalent formulation of the classical isoperimetric inequality, which states that balls minimize surface area.

3.4 Log-concavity

A probabilistic analogue of convexity is provided by the notion of *log-concavity*. Just as convexity imposes rigid geometric structure on sets, log-concavity imposes strong regularity on probability measures and their densities. Log-concave measures form a stable class under natural operations and they imply powerful concentration phenomena. More importantly, Conjecture 2.19 (KLS) asks whether every isotropic log-concave measure satisfies a dimension-free Poincaré inequality, thereby placing log-concavity at the heart of one of the open problems in the field. We now introduce this fundamental class of probability measures.

Definition 3.5 (Log-concave function). A function $f : \mathbb{R}^n \rightarrow \mathbb{R}^+$ is called *log-concave* if there exists a concave function

$$V : \mathbb{R}^n \rightarrow \mathbb{R}$$

such that

$$f(x) = e^{V(x)}.$$

Equivalently, f is log-concave if $\log f$ is a concave function on the domain

$$\{x \in \mathbb{R}^n : f(x) > 0\}.$$

That is, for all $x, y \in \mathbb{R}^n$ and all $\lambda \in [0, 1]$,

$$f(\lambda x + (1 - \lambda)y) \geq f(x)^\lambda f(y)^{1-\lambda}.$$

Claim 3.6. Let $V : \mathbb{R}^n \rightarrow \mathbb{R}$ be a concave function. Define the probability measure

$$\mu = e^{V(x)} dx.$$

Then, μ is a log-concave measure, i.e., for all $t \in [0, 1]$ and measurable subsets A, B ,

$$\mu(tA + (1 - t)B) \geq \mu(A)^t \mu(B)^{1-t}. \quad (12)$$

Proof. For measurable sets $A, B \subset \mathbb{R}^n$ define

$$f(x) = e^{V(x)} \mathbf{1}_A(x), \quad g(x) = e^{V(x)} \mathbf{1}_B(x), \quad h(x) = e^{V(x)} \mathbf{1}_{tA + (1-t)B}(x).$$

Then for all x, y ,

$$h(tx + (1 - t)y) \geq f(x)^t g(y)^{1-t}.$$

Applying Prékopa–Leindler inequality yields

$$\int_{\mathbb{R}^n} h \geq \left(\int_{\mathbb{R}^n} f \right)^t \left(\int_{\mathbb{R}^n} g \right)^{1-t},$$

i.e.

$$\mu(tA + (1 - t)B) \geq \mu(A)^t \mu(B)^{1-t}.$$

So μ is a log-concave measure. □

Examples of log-concave measures

The class of log-concave measures on $\mathbb{R}^n$ is very broad. Important examples include:

1. Gaussian measures.

Any Gaussian measure on $\mathbb{R}^n$ is log-concave. Indeed, if $X \sim N(m, \Sigma)$, then its density is

$$f(x) = \frac{1}{(2\pi)^{n/2} \det(\Sigma)^{1/2}} \exp\left(-\frac{1}{2}(x-m)^T \Sigma^{-1}(x-m)\right).$$

Therefore

$$\log f(x) = C - \frac{1}{2}(x-m)^T \Sigma^{-1}(x-m),$$

where C is a constant. Since Σ^{-1} is positive semidefinite, the quadratic form $(x-m)^T \Sigma^{-1}(x-m)$ is convex, and hence $\log f$ is concave. Thus f is log-concave.

2. Uniform measures on convex bodies.

If $K \subseteq \mathbb{R}^n$ is a convex Borel set, then

$$d\mu(x) = \frac{1}{|K|} \mathbf{1}_K(x) dx.$$

Since the indicator of a convex set satisfies $\mathbf{1}_K((1-\lambda)x + \lambda y) \geq \mathbf{1}_K(x)^{1-\lambda} \mathbf{1}_K(y)^\lambda$, the density is log-concave.

3. Laplace distributions.

The one-dimensional density

$$f(x) = \frac{1}{2} e^{-|x|}$$

is log-concave since $\log f(x) = C - |x|$, and $-|x|$ is concave on $\mathbb{R}$.

4. Logistic distributions.

The density

$$f(x) = \frac{e^{-x}}{(1 + e^{-x})^2}$$

is log-concave. Indeed, $\log f(x) = -x - 2 \log(1 + e^{-x})$, whose second derivative is nonpositive.

5. Gamma distributions with shape parameter $\alpha \geq 1$.

The density

$$f(x) = \frac{1}{\Gamma(\alpha)\theta^\alpha} x^{\alpha-1} e^{-x/\theta} \mathbf{1}_{(0,\infty)}(x)$$

is log-concave whenever $\alpha \geq 1$. Indeed, $\log f(x) = C + (\alpha-1) \log x - \frac{x}{\theta}$, which is concave on $(0, \infty)$ when $\alpha \geq 1$.

6. Beta distributions with parameters $\alpha, \beta \geq 1$.

The density

$$f(x) = \frac{\Gamma(\alpha + \beta)}{\Gamma(\alpha)\Gamma(\beta)} x^{\alpha-1} (1-x)^{\beta-1} \mathbf{1}_{(0,1)}(x)$$

is log-concave whenever $\alpha, \beta \geq 1$. Indeed, $\log f(x) = C + (\alpha - 1) \log x + (\beta - 1) \log(1 - x)$, a concave function on $(0, 1)$ when $\alpha, \beta \geq 1$.

Moreover, the class of log-concave measures is closed under several natural operations. We will prove the closure under marginalization in Theorem 3.7 below and the rest, such as affine transformations and independent products, can be found in the survey of Saumard and Wellner [13]

Theorem 3.7. *Marginalization preserves log-concavity. In other words, if (X, Y) is a random vector on $\mathbb{R}^n \times \mathbb{R}^m$ with log-concave joint density $f_{X,Y}$, then the marginal density*

$$f_X(x) := \int_{\mathbb{R}^m} f_{X,Y}(x, y) dy$$

is log-concave.

Proof. Let (X, Y) be a random vector on $\mathbb{R}^n \times \mathbb{R}^m$ with joint density

$$f_{X,Y} : \mathbb{R}^{n+m} \rightarrow [0, \infty).$$

Assume that $f_{X,Y}$ is log-concave. Define the marginal density of X by

$$f_X(x) := \int_{\mathbb{R}^m} f_{X,Y}(x, y) dy.$$

We show that f_X is log-concave. Fix $x_0, x_1 \in \mathbb{R}^n$ and $\lambda \in [0, 1]$, and set

$$x_\lambda := \lambda x_0 + (1 - \lambda)x_1.$$

For $y_0, y_1 \in \mathbb{R}^m$, log-concavity of $f_{X,Y}$ gives

$$f_{X,Y}(x_\lambda, \lambda y_0 + (1 - \lambda)y_1) \geq f_{X,Y}(x_0, y_0)^\lambda f_{X,Y}(x_1, y_1)^{1-\lambda}. \quad (13)$$

Define

$$f(y) := f_{X,Y}(x_0, y), \quad g(y) := f_{X,Y}(x_1, y), \quad h(y) := f_{X,Y}(x_\lambda, y).$$

Then, (13) is equivalent to

$$h(\lambda y_0 + (1 - \lambda)y_1) \geq f(y_0)^\lambda g(y_1)^{1-\lambda}.$$

By the Prékopa–Leindler inequality,

$$\int_{\mathbb{R}^m} h(y) dy \geq \left(\int_{\mathbb{R}^m} f(y) dy \right)^\lambda \left(\int_{\mathbb{R}^m} g(y) dy \right)^{1-\lambda}.$$

Therefore

$$f_X(x_\lambda) \geq f_X(x_0)^\lambda f_X(x_1)^{1-\lambda}.$$

Hence

$$f_X(\lambda x_0 + (1 - \lambda)x_1) \geq f_X(x_0)^\lambda f_X(x_1)^{1-\lambda}.$$

Thus f_X is log-concave. Same argument can be used to prove f_Y is log-concave. $\square$

Another important stability property is preservation under convolution. In probabilistic terms, this means that the sum of independent log-concave random variables remains log-concave.

Definition 3.8 (Convolution). Let $f, g : \mathbb{R}^n \rightarrow [0, \infty)$ be integrable functions. Their *convolution* is the function $f * g : \mathbb{R}^n \rightarrow [0, \infty)$ defined by

$$(f * g)(z) := \int_{\mathbb{R}^n} f(x) g(z - x) dx.$$

If f and g are probability densities, then $f * g$ is the density of the sum of two independent random variables with densities f and g .

Claim 3.9 (Convolution preserves log-concavity). *Let $f, g : \mathbb{R}^n \rightarrow [0, \infty)$ be log-concave integrable functions. Then their convolution*

$$(f * g)(z) := \int_{\mathbb{R}^n} f(x) g(z - x) dx$$

is log-concave on $\mathbb{R}^n$.

Proof. Fix $z_0, z_1 \in \mathbb{R}^n$ and $\lambda \in [0, 1]$, and set

$$z_\lambda := (1 - \lambda)z_0 + \lambda z_1.$$

Define

$$u(x) := f(x)g(z_0 - x), \quad v(y) := f(y)g(z_1 - y),$$

and

$$w(t) := (f * g)(t) = \int_{\mathbb{R}^n} f(x)g(t - x) dx.$$

We claim that for all $x, y \in \mathbb{R}^n$,

$$w((1 - \lambda)x + \lambda y) \geq u(x)^{1-\lambda} v(y)^\lambda.$$

Indeed, by log-concavity of f and g , we have

$$\begin{aligned} f((1 - \lambda)x + \lambda y) &\geq f(x)^{1-\lambda} f(y)^\lambda, \\ g(z_\lambda - ((1 - \lambda)x + \lambda y)) &= g((1 - \lambda)(z_0 - x) + \lambda(z_1 - y)) \\ &\geq g(z_0 - x)^{1-\lambda} g(z_1 - y)^\lambda. \end{aligned}$$

Multiplying these inequalities yields

$$\begin{aligned} u(x)^{1-\lambda} v(y)^\lambda &= (f(x)g(z_0 - x))^{1-\lambda} (f(y)g(z_1 - y))^\lambda \\ &\leq f((1 - \lambda)x + \lambda y) g(z_\lambda - ((1 - \lambda)x + \lambda y)). \end{aligned}$$

Thus, if we define

$$h(\xi) := f(\xi) g(z_\lambda - \xi),$$

then for all $x, y \in \mathbb{R}^n$,

$$h((1 - \lambda)x + \lambda y) \geq u(x)^{1-\lambda} v(y)^\lambda.$$

We may therefore apply the Prékopa–Leindler inequality to the functions u , v , and h , obtaining

$$\int_{\mathbb{R}^n} h(\xi) d\xi \geq \left(\int_{\mathbb{R}^n} u(x) dx \right)^{1-\lambda} \left(\int_{\mathbb{R}^n} v(y) dy \right)^\lambda.$$

By definition of convolution,

$$\int_{\mathbb{R}^n} h(\xi) d\xi = (f * g)(z_\lambda), \quad \int_{\mathbb{R}^n} u(x) dx = (f * g)(z_0), \quad \int_{\mathbb{R}^n} v(y) dy = (f * g)(z_1).$$

Hence,

$$(f * g)(z_\lambda) \geq (f * g)(z_0)^{1-\lambda} (f * g)(z_1)^\lambda.$$

This shows that $f * g$ is log-concave. $\square$

4 Logarithmic Sobolev inequality

The logarithmic Sobolev inequality provides a refinement of the Poincaré inequality by controlling not only the variance of a function, but a stronger quantity known as entropy. In particular, while the Poincaré inequality bounds fluctuations of a function around its mean, the logarithmic Sobolev inequality captures the rate at which a function concentrates and spreads under nonlinear transformations.

A key feature of the logarithmic Sobolev inequality is that it links three fundamental objects:

- entropy (a global, nonlinear measure of dispersion),
- gradients (local regularity of functions), and
- the underlying geometry of the measure.

The significance of Chapter 3 extends beyond Brunn–Minkowski and isoperimetric inequalities. One of the consequences of the Prékopa–Leindler theorem is that it provides a way to prove log-concavity properties. Since strongly log-concave measures satisfy logarithmic Sobolev inequalities, Prékopa–Leindler becomes a bridge between convexity and concentration phenomena.

4.1 Log–Sobolev inequality for strongly log-concave measures

We now show how the Prékopa–Leindler inequality leads to a logarithmic Sobolev inequality for strongly log-concave measures. This result strengthens the Poincaré inequality discussed earlier: instead of controlling variance, it controls entropy.

Definition 4.1 (Entropy). Let μ be a probability measure and let $h \geq 0$ be integrable. The entropy of h with respect to μ is defined by

$$\text{Ent}_\mu(h) := \int h \log h d\mu - \left(\int h d\mu \right) \log \left(\int h d\mu \right).$$

Definition 4.2 (Logarithmic Sobolev inequality). Let μ be a probability measure on $\mathbb{R}^n$. We say that μ satisfies a logarithmic Sobolev inequality with constant $\rho > 0$ if, for every sufficiently smooth function $h : \mathbb{R}^n \rightarrow \mathbb{R}$,

$$\text{Ent}_\mu(h^2) \leq \frac{2}{\rho} \int \|\nabla h\|^2 d\mu.$$

While the Poincaré inequality controls variance, the logarithmic Sobolev inequality controls entropy. This stronger control leads to Gaussian-type concentration estimates for Lipschitz functions, as shown in Theorem 4.3 provided below.

Theorem 4.3. Let μ be a probability measure on $\mathbb{R}^n$ satisfying the logarithmic Sobolev inequality, i.e.,

$$\text{Ent}_\mu(h^2) \leq \frac{2}{\rho} \int \|\nabla h\|^2 d\mu$$

for every sufficiently smooth function h . Let $f : \mathbb{R}^n \rightarrow \mathbb{R}$ be 1-Lipschitz and assume that

$$\int f d\mu = 0.$$

Then, for every $\lambda > 0$,

$$\int e^{\lambda f} d\mu \leq e^{\lambda^2/(2\rho)}.$$

Consequently, if $X \sim \mu$, then for every $t > 0$,

$$\mathbb{P}(|f(X)| \geq t) \leq 2e^{-\rho t^2/2}.$$

Proof. Define

$$F(\lambda) := \log \left(\int e^{\lambda f} d\mu \right).$$

Then

$$F'(\lambda) = \frac{\int f e^{\lambda f} d\mu}{\int e^{\lambda f} d\mu}.$$

Applying the logarithmic Sobolev inequality to

$$h = e^{\lambda f/2}$$

gives

$$\text{Ent}_\mu(e^{\lambda f}) \leq \frac{2}{\rho} \int \|\nabla e^{\lambda f/2}\|^2 d\mu.$$

Since

$$\nabla e^{\lambda f/2} = \frac{\lambda}{2} e^{\lambda f/2} \nabla f,$$

we get

$$\text{Ent}_\mu(e^{\lambda f}) \leq \frac{\lambda^2}{2\rho} \int \|\nabla f\|^2 e^{\lambda f} d\mu.$$

Because f is 1-Lipschitz, $\|\nabla f\| \leq 1$ almost everywhere. Hence

$$\text{Ent}_\mu(e^{\lambda f}) \leq \frac{\lambda^2}{2\rho} \int e^{\lambda f} d\mu. \quad (14)$$

Now write

$$Z(\lambda) := \int e^{\lambda f} d\mu.$$

Since $F(\lambda) = \log Z(\lambda)$, we have

$$F'(\lambda) = \frac{Z'(\lambda)}{Z(\lambda)}.$$

Using the definition of entropy,

$$\text{Ent}_\mu(e^{\lambda f}) = \int \lambda f e^{\lambda f} d\mu - Z(\lambda) \log Z(\lambda).$$

Thus

$$\text{Ent}_\mu(e^{\lambda f}) = \lambda Z'(\lambda) - Z(\lambda) F(\lambda).$$

Dividing by $Z(\lambda)$ gives

$$\frac{\text{Ent}_\mu(e^{\lambda f})}{Z(\lambda)} = \lambda F'(\lambda) - F(\lambda).$$

Combining this with (14)

$$\lambda F'(\lambda) - F(\lambda) \leq \frac{\lambda^2}{2\rho}.$$

Equivalently,

$$\left(\frac{F(\lambda)}{\lambda} \right)' = \frac{\lambda F'(\lambda) - F(\lambda)}{\lambda^2} \leq \frac{1}{2\rho}. \quad (15)$$

Since $\int f d\mu = 0$, we have

$$\lim_{\lambda \downarrow 0} \frac{F(\lambda)}{\lambda} = \int f d\mu = 0.$$

Integrating (15) from 0 to λ gives

$$\frac{F(\lambda)}{\lambda} \leq \frac{\lambda}{2\rho}.$$

Therefore,

$$\int e^{\lambda f} d\mu = e^{F(\lambda)} \leq e^{\lambda^2/(2\rho)}.$$

Finally, by Markov's inequality, for every $\lambda > 0$,

$$\mathbb{P}(f(X) \geq t) = \mathbb{P}(e^{\lambda f(X)} \geq e^{\lambda t}) \leq e^{-\lambda t} \int e^{\lambda f} d\mu \leq \exp\left(-\lambda t + \frac{\lambda^2}{2\rho}\right).$$

Optimizing over $\lambda > 0$ gives

$$\lambda = \rho t.$$

Therefore,

$$\mathbb{P}(f(X) \geq t) \leq e^{-\rho t^2/2}.$$

Applying the same argument to $-f$ gives

$$\mathbb{P}(f(X) \leq -t) \leq e^{-\rho t^2/2}.$$

Thus

$$\mathbb{P}(|f(X)| \geq t) \leq 2e^{-\rho t^2/2}.$$

□

We now explain how logarithmic Sobolev inequality arise from convexity assumptions on the underlying measure. In particular, a strongly log-concave measure satisfies a logarithmic Sobolev inequality.

Theorem 4.4. *Let μ be a probability measure on $\mathbb{R}^d$ of the form*

$$d\mu(x) = e^{-V(x)} dx, \quad \int_{\mathbb{R}^d} e^{-V(x)} dx = 1,$$

where $V : \mathbb{R}^d \rightarrow \mathbb{R}$ is sufficiently smooth. Assume that there exists $\rho > 0$ such that

$$\nabla^2 V(x) \geq \rho I_d \quad \text{for all } x \in \mathbb{R}^d.$$

Equivalently,

$$V(tx + (1-t)y) \leq tV(x) + (1-t)V(y) - \frac{\rho}{2}t(1-t)\|x - y\|^2$$

for all $x, y \in \mathbb{R}^d$ and $t \in [0, 1]$. Then, for every sufficiently smooth function f ,

$$\text{Ent}_\mu(f^2) \leq \frac{2}{\rho} \int \|\nabla f\|^2 d\mu.$$

Proof. To simplify the argument, we first consider functions of the form $f^2 = e^g$. This is not a restriction: every strictly positive function $F = f^2$ can be written uniquely as $F = e^g$ with $g = \log F$. Functions that may vanish can be treated by a standard approximation argument, replacing f^2 by $f^2 + \varepsilon$ and letting $\varepsilon \downarrow 0$. Thus it suffices to establish the inequality in the parametrization $f^2 = e^g$. Since

$$\nabla f = \frac{1}{2}e^{g/2}\nabla g,$$

we have

$$\|\nabla f\|^2 = \frac{1}{4}e^g\|\nabla g\|^2.$$

Thus the desired inequality is equivalent to

$$\text{Ent}_\mu(e^g) \leq \frac{1}{2\rho} \int \|\nabla g\|^2 e^g d\mu.$$

Now, fix $t \in (0, 1)$ and set $s = 1 - t$. Define

$$g_t(z) := \sup_{z=tx+sy} \{g(x) - tV(x) - sV(y) + V(tx + sy)\}.$$

Consider the functions

$$h(z) := \exp(g_t(z) - V(z)), \quad f_t(x) := \exp\left(\frac{g(x)}{t} - V(x)\right), \quad \varphi(y) := \exp(-V(y)).$$

By the definition of g_t , whenever $z = tx + sy$, we have

$$g_t(z) - V(z) \geq g(x) - tV(x) - sV(y).$$

Therefore,

$$h(tx + sy) \geq f_t(x)^t \varphi(y)^s.$$

Applying Theorem 3.1 (Prékopa–Leindler inequality) gives

$$\int_{\mathbb{R}^d} e^{g_t(z) - V(z)} dz \geq \left(\int_{\mathbb{R}^d} e^{g(x)/t - V(x)} dx \right)^t \left(\int_{\mathbb{R}^d} e^{-V(y)} dy \right)^s.$$

Since μ is a probability measure, $\int e^{-V} = 1$. Hence

$$\int e^{g_t} d\mu \geq \left(\int e^{g/t} d\mu \right)^t. \quad (16)$$

Define

$$B(t) := \int e^{g/t} d\mu, \quad \text{and} \quad \Phi(t) := (B(t))^t.$$

Then, we have $\log \Phi(t) = t \log B(t)$. Differentiating it, we have

$$\frac{\Phi'(t)}{\Phi(t)} = \log B(t) + t \frac{B'(t)}{B(t)}. \quad (17)$$

Since

$$B'(t) = -\frac{1}{t^2} \int g e^{g/t} d\mu,$$

equation (17) implies

$$\Phi'(1) = \left(\int e^g d\mu \right) \log \left(\int e^g d\mu \right) - \int g e^g d\mu = -\text{Ent}_\mu(e^g).$$

The Taylor expansion of $\Phi(t)$ around $t = 1$ implies

$$\Phi(t) = \Phi(1) + (t - 1)\Phi'(1) + O((t - 1)^2).$$

Letting $t = 1 - s$, we have $s \downarrow 0$, and

$$\Phi(1 - s) = \Phi(1) - s\Phi'(1) + O(s^2) = \int e^g d\mu + s \text{Ent}_\mu(e^g) + O(s^2).$$

Recall that $\Phi(1 - s) = \Phi(t) = B(t)^t = \left(\int e^{g/t} d\mu \right)^t$. Therefore,

$$\left(\int e^{g/t} d\mu \right)^t = \int e^g d\mu + s \text{Ent}_\mu(e^g) + O(s^2). \quad (18)$$

Next we estimate g_t . Let $z = tx + sy$ and set

$$w = z - y = t(x - y), \quad r := \frac{s}{t}.$$

Using the uniform convexity assumption on V , we have

$$tV(x) + sV(y) \geq V(tx + sy) + \frac{\rho}{2}st\|x - y\|^2.$$

Since $w = t(x - y)$, this implies

$$g_t(z) \leq \sup_w \left\{ g(z + rw) - \frac{\rho r}{2} \|w\|^2 \right\}.$$

For small r , Taylor expanding $g(z + rw)$ gives

$$g(z + rw) = g(z) + r\langle \nabla g(z), w \rangle + o(r).$$

Therefore,

$$g_t(z) \leq g(z) + \sup_w \left\{ r\langle \nabla g(z), w \rangle - \frac{\rho r}{2} \|w\|^2 \right\} + o(r).$$

The supremum is attained at

$$w = \frac{1}{\rho} \nabla g(z),$$

and its value is

$$\frac{r}{2\rho} \|\nabla g(z)\|^2.$$

Thus

$$g_t(z) \leq g(z) + \frac{r}{2\rho} \|\nabla g(z)\|^2 + o(r).$$

Consequently,

$$\int e^{g_t} d\mu \leq \int e^g \left(1 + \frac{r}{2\rho} \|\nabla g\|^2 + o(r) \right) d\mu. \quad (19)$$

Combining (16), (18), and (19), and using $r = s/t \sim s$ as $s \downarrow 0$, we obtain

$$\text{Ent}_\mu(e^g) \leq \frac{1}{2\rho} \int \|\nabla g\|^2 e^g d\mu.$$

Returning to the parametrization $f^2 = e^g$, this becomes

$$\text{Ent}_\mu(f^2) \leq \frac{2}{\rho} \int \|\nabla f\|^2 d\mu,$$

which proves the logarithmic Sobolev inequality. $\square$

Corollary 4.5 (Gaussian logarithmic Sobolev inequality). *Let γ denote the standard Gaussian measure on $\mathbb{R}^d$. Then, for every sufficiently smooth function f ,*

$$\text{Ent}_\gamma(f^2) \leq 2 \int \|\nabla f\|^2 d\gamma.$$

Proof. For the standard Gaussian measure,

$$d\gamma(x) = e^{-V(x)} dx, \quad V(x) = \frac{\|x\|^2}{2} + \frac{d}{2} \log(2\pi).$$

Hence

$$\nabla^2 V(x) = I_d.$$

Thus the previous theorem applies with $\rho = 1$, giving

$$\text{Ent}_\gamma(f^2) \leq 2 \int \|\nabla f\|^2 d\gamma.$$

□

5 Cheeger's inequality

In the previous chapters we studied functional inequalities such as the Poincaré and logarithmic Sobolev inequalities, which quantify concentration through the behavior of functions and their gradients. We now turn to a geometric perspective based on the geometric properties of sets. While the Poincaré inequality bounds fluctuations of functions using gradients, Cheeger's inequality bounds the measure of sets using their geometric expansion. Remarkably, these viewpoints are closely linked—in fact, Cheeger's inequality is equivalent to an L^1 -version of the Poincaré inequality, which we will prove in Theorem 5.2.

Definition 5.1 (Cheeger inequality). Let μ be a probability measure on $\mathbb{R}^d$, and let $A \subset \mathbb{R}^d$ be a Borel set. Let $B := \{x \in \mathbb{R}^d : \|x\|_2 \leq 1\}$ be the unit Euclidean ball. For $\varepsilon > 0$, define the ε -enlargement of A by

$$A_\varepsilon := A + \varepsilon B = \{x + y : x \in A, \|y\|_2 \leq \varepsilon\}.$$

The (outer) boundary measure of A with respect to μ is defined as

$$\mu^+(A) := \liminf_{\varepsilon \rightarrow 0^+} \frac{\mu(A_\varepsilon) - \mu(A)}{\varepsilon}.$$

We say that μ satisfies a *Cheeger isoperimetric inequality* if there exists a constant $C_I(\mu) > 0$ such that for all Borel sets $A \subset \mathbb{R}^d$,

$$\mu^+(A) \geq C_I(\mu) \cdot \min\{\mu(A), 1 - \mu(A)\}.$$

The largest such constant $C_I(\mu)$ is called the *Cheeger constant* of μ .

Geometric and functional inequalities are often two manifestations of the same underlying phenomenon. The main goal of this section is to make this connection precise by showing that Cheeger's inequality is equivalent to an L^1 -Poincaré inequality and, in particular, implies the usual L^2 -Poincaré inequality.

Theorem 5.2. *Let μ be a probability measure on $\mathbb{R}^d$. The following are equivalent.*

1. (Cheeger's inequality) There exists $C_I(\mu) > 0$ such that for all measurable sets A ,

$$\mu^+(A) \geq C_I(\mu) \min\{\mu(A), 1 - \mu(A)\}.$$

2. (L^1 -Poincaré inequality) For every Lipschitz function $f : \mathbb{R}^d \rightarrow \mathbb{R}$,

$$\mathbb{E}_\mu |f - \text{med}_\mu(f)| \leq \frac{1}{C_I(\mu)} \mathbb{E}_\mu \|\nabla f\|,$$

where

$$\text{med}_\mu(f) := \inf \left\{ x \in \mathbb{R} : \mu(f \leq x) \geq \frac{1}{2} \right\}$$

denotes a median of f under μ .

Proof. We first prove that the Cheeger's inequality implies the L^1 -Poincaré inequality. Let $\tilde{f} = f - \text{med}_\mu(f)$. Then $\text{med}_\mu(\tilde{f}) = 0$ and $\nabla \tilde{f} = \nabla f$, so it suffices to assume that $\text{med}_\mu(f) = 0$. By the coarea formula,

$$\mathbb{E}_\mu \|\nabla f\| = \int_{-\infty}^{\infty} \mu^+(\{f > t\}) dt.$$

Applying the Cheeger's inequality to the level sets $\{f > t\}$,

$$\mu^+(\{f > t\}) \geq C_I(\mu) \min(\mu(f > t), 1 - \mu(f > t)).$$

Integrating,

$$\mathbb{E}_\mu \|\nabla f\| \geq C_I(\mu) \int_{-\infty}^{\infty} \min(\mu(f > t), 1 - \mu(f > t)) dt.$$

Since 0 is a median,

$$\mu(f > 0) \leq \frac{1}{2}, \quad \mu(f < 0) \leq \frac{1}{2}.$$

Hence

$$\min(\mu(f > t), 1 - \mu(f > t)) = \begin{cases} \mu(f > t), & t > 0, \\ 1 - \mu(f > t), & t < 0. \end{cases}$$

Therefore,

$$\int_{-\infty}^{\infty} \min(\mu(f > t), 1 - \mu(f > t)) dt = \int_0^{\infty} \mu(f > t) dt + \int_{-\infty}^0 \mu(f \leq t) dt.$$

By the layer-cake representation,

$$\int_0^{\infty} \mu(f > t) dt = \mathbb{E}_\mu f_+, \quad \int_{-\infty}^0 \mu(f \leq t) dt = \mathbb{E}_\mu f_-,$$

so

$$\int_{-\infty}^{\infty} \min(\mu(f > t), 1 - \mu(f > t)) dt = \mathbb{E}_\mu |f|.$$

Thus

$$\mathbb{E}_\mu \|\nabla f\| \geq C_I(\mu) \mathbb{E}_\mu |f|,$$

which gives

$$\mathbb{E}_\mu |f - \text{med}_\mu(f)| \leq \frac{1}{C_I(\mu)} \mathbb{E}_\mu \|\nabla f\|.$$

For the reverse implication, assume that the L^1 -Poincaré inequality holds. Let $A \subseteq \mathbb{R}^d$ be a measurable set. We want to prove

$$\mu^+(A) \geq C_I(\mu) \min\{\mu(A), 1 - \mu(A)\}.$$

The idea is to apply the L^1 -Poincaré inequality to the indicator function χ_A . However, χ_A is not Lipschitz in general, so we instead approximate it by Lipschitz functions. For $\varepsilon > 0$, define

$$f_\varepsilon(x) := \max\left\{1 - \frac{\text{dist}(x, A)}{\varepsilon}, 0\right\}.$$

Then f_ε is ε^{-1} -Lipschitz, satisfies

$$0 \leq f_\varepsilon \leq 1, \quad f_\varepsilon = 1 \text{ on } A, \quad f_\varepsilon = 0 \text{ outside } A_\varepsilon,$$

and $f_\varepsilon \rightarrow \chi_A$ pointwise as $\varepsilon \downarrow 0$. Moreover, by the coarea formula,

$$\int \|\nabla f_\varepsilon\| d\mu = \int_0^1 \mu^+(\{f_\varepsilon > t\}) dt.$$

The level sets $\{f_\varepsilon > t\}$ are contained in A_ε and shrink to A as $\varepsilon \downarrow 0$. Thus, using the definition of the outer boundary measure,

$$\liminf_{\varepsilon \downarrow 0} \int \|\nabla f_\varepsilon\| d\mu \leq \mu^+(A).$$

Applying the L^1 -Poincaré inequality to f_ε gives

$$\mathbb{E}_\mu |f_\varepsilon - \text{med}_\mu(f_\varepsilon)| \leq \frac{1}{C_I(\mu)} \int \|\nabla f_\varepsilon\| d\mu.$$

Letting $\varepsilon \downarrow 0$, the left-hand side converges to

$$\mathbb{E}_\mu |\chi_A - \text{med}_\mu(\chi_A)| = \min\{\mu(A), 1 - \mu(A)\}.$$

Therefore,

$$\min\{\mu(A), 1 - \mu(A)\} \leq \frac{1}{C_I(\mu)} \mu^+(A).$$

Rearranging gives

$$\mu^+(A) \geq C_I(\mu) \min\{\mu(A), 1 - \mu(A)\},$$

which is precisely the Cheeger's inequality. □

Corollary 5.3. *Assume that μ satisfies the Cheeger's inequality. Then,*

$$C_P(\mu) \leq \frac{4}{C_I(\mu)^2}. \tag{20}$$

Proof. Let m be a median of f , and replace f by $f - m$. Since subtracting a constant does not change the gradient and

$$\text{Var}_\mu(f) \leq \int (f - m)^2 d\mu,$$

(as $u^* = \mathbb{E}_\mu[f]$ minimizes the function $h(u) = \int (f - u)^2 d\mu$), it is enough to prove that if $\text{med}(f) = 0$, then

$$\int f^2 d\mu \leq \frac{4}{C_I(\mu)^2} \int |\nabla f|^2 d\mu.$$

Assume therefore that $\text{med}(f) = 0$, and define

$$f_+ := \max(f, 0), \quad f_- := \max(-f, 0).$$

Then

$$f = f_+ - f_-, \quad |f| = f_+ + f_-,$$

and since $f_+(x)f_-(x) = 0$ for every x , we also have

$$f^2 = f_+^2 + f_-^2. \tag{21}$$

Therefore,

$$\int f^2 d\mu = \int f_+^2 d\mu + \int f_-^2 d\mu.$$

Now since $\text{med}(f) = 0$, we have

$$\mu(f \leq 0) \geq \frac{1}{2} \quad \text{and} \quad \mu(f \geq 0) \geq \frac{1}{2}.$$

According to the definitions of f_+ and f_- ,

$$\mu(f_+^2 = 0) = \mu(f_+ = 0) = \mu(f \leq 0) \geq \frac{1}{2}, \quad \mu(f_-^2 = 0) = \mu(f_- = 0) = \mu(f \geq 0) \geq \frac{1}{2},$$

so $\text{med}(f_+) = \text{med}(f_-) = 0$. Applying Theorem 5.2 to f_+^2 and f_-^2 , we have

$$\int f_+^2 d\mu \leq \frac{1}{C_I(\mu)} \int |\nabla(f_+^2)| d\mu, \quad \int f_-^2 d\mu \leq \frac{1}{C_I(\mu)} \int |\nabla(f_-^2)| d\mu.$$

Summing,

$$\int f^2 d\mu \leq \frac{1}{C_I(\mu)} \left(\int |\nabla(f_+^2)| d\mu + \int |\nabla(f_-^2)| d\mu \right).$$

Using the chain rule,

$$|\nabla(f_+^2)| = 2f_+|\nabla f_+|, \quad |\nabla(f_-^2)| = 2f_-|\nabla f_-|.$$

So

$$\int f^2 d\mu \leq \frac{2}{C_I(\mu)} \left(\int f_+|\nabla f_+| d\mu + \int f_-|\nabla f_-| d\mu \right).$$

Now apply the Cauchy–Schwarz to each term. We have

$$\int f_+|\nabla f_+| d\mu \leq \left(\int f_+^2 d\mu \right)^{1/2} \left(\int |\nabla f_+|^2 d\mu \right)^{1/2},$$

$$\int f_- |\nabla f_-| d\mu \leq \left(\int f_-^2 d\mu \right)^{1/2} \left(\int |\nabla f_-|^2 d\mu \right)^{1/2}.$$

Thus

$$\int f^2 d\mu \leq \frac{2}{C_I(\mu)} \left[\left(\int f_+^2 d\mu \right)^{1/2} \left(\int |\nabla f_+|^2 d\mu \right)^{1/2} + \left(\int f_-^2 d\mu \right)^{1/2} \left(\int |\nabla f_-|^2 d\mu \right)^{1/2} \right].$$

Using the inequality

$$a_1 b_1 + a_2 b_2 \leq \sqrt{a_1^2 + a_2^2} \sqrt{b_1^2 + b_2^2}$$

with

$$\begin{aligned} a_1 &= \left(\int f_+^2 d\mu \right)^{1/2}, & a_2 &= \left(\int f_-^2 d\mu \right)^{1/2}, \\ b_1 &= \left(\int |\nabla f_+|^2 d\mu \right)^{1/2}, & b_2 &= \left(\int |\nabla f_-|^2 d\mu \right)^{1/2}, \end{aligned}$$

we obtain

$$\int f^2 d\mu \leq \frac{2}{C_I(\mu)} \left(\int (f_+^2 + f_-^2) d\mu \right)^{1/2} \left(\int (|\nabla f_+|^2 + |\nabla f_-|^2) d\mu \right)^{1/2}.$$

According to equation (21), for any function f we also have

$$|\nabla f_+|^2 + |\nabla f_-|^2 = |\nabla f|^2 \quad \text{a.e.},$$

which implies

$$\int f^2 d\mu \leq \frac{2}{C_I(\mu)} \left(\int f^2 d\mu \right)^{1/2} \left(\int |\nabla f|^2 d\mu \right)^{1/2}.$$

If $\int f^2 d\mu = 0$, there is nothing to prove. Otherwise divide by

$$\left(\int f^2 d\mu \right)^{1/2}$$

to get

$$\left(\int f^2 d\mu \right)^{1/2} \leq \frac{2}{C_I(\mu)} \left(\int |\nabla f|^2 d\mu \right)^{1/2}.$$

Squaring both sides yields

$$\int f^2 d\mu \leq \frac{4}{C_I(\mu)^2} \int |\nabla f|^2 d\mu.$$

as desired. This, together with the definition of Poincaré constant, implies

$$C_P(\mu) \leq \frac{4}{C_I(\mu)^2}.$$

□

Remark 5.4. The previous corollary shows that a lower bound on the Cheeger's constant implies an upper bound on the Poincaré constant:

$$C_P(\mu) \leq \frac{4}{C_I(\mu)^2}.$$

For general probability measures, the converse implication is false. However, under suitable convexity assumptions, such as for absolutely continuous log-concave probability measures, the reverse inequality holds up to a universal constant. More precisely, in 2009, Emanuel Milman [9] showed that under certain convexity assumptions the Cheeger's and Poincaré inequalities are equivalent up to universal constants. In our notation, this implies the converse direction of equation (20)

$$C_P(\mu) \asymp \frac{1}{C_I(\mu)^2}.$$

Thus, under certain assumptions, the geometric information captured by the isoperimetric profile and the functional information captured by the Poincaré inequality are quantitatively equivalent.

6 Payne–Weinberger theorem

In the previous sections, we studied several functional inequalities and their consequences in high-dimensional settings. A recurring difficulty in proving such inequalities is that direct arguments in high dimensions are often complicated and difficult to control geometrically. One strategy for overcoming this difficulty is to reduce the problem to lower-dimensional statements, where the geometry becomes more transparent and sharper tools are available. This idea can be used in the proof of the Payne–Weinberger theorem. The localization method provides a systematic framework for carrying out this dimension reduction. Roughly speaking, localization allows one to decompose a high-dimensional measure into a family of one-dimensional components, often supported on intervals, while preserving enough structure to transfer inequalities from the one-dimensional setting back to the original space. In this section, we introduce the localization lemma and then apply it in the proof of the Payne–Weinberger theorem.

6.1 Localization lemma

The heart of the method is the ability to reduce certain integral inequalities in $\mathbb{R}^n$ to one-dimensional inequalities along *needles* (segments) with a specific log-concave weight.

Lemma 6.1 (Localization lemma / needle decomposition [8, 6]). *Let μ be a log-concave measure on $\mathbb{R}^n$, and let $f, g : \mathbb{R}^n \rightarrow \mathbb{R}$ be continuous and integrable with respect to μ . Assume*

$$\int f \, d\mu = 0 \quad \text{and} \quad \int g \, d\mu \geq 0.$$

Then there exist two points $a, b \in \mathbb{R}^n$ and a log-concave function $\psi : [0, 1] \rightarrow \mathbb{R}_+$ such that, writing the segment parametrization

$$x(t) := (1 - t)a + tb \quad (t \in [0, 1]),$$

we have the one-dimensional reductions

$$\int_0^1 f(x(t)) \psi(t) dt = 0 \quad \text{and} \quad \int_0^1 g(x(t)) \psi(t) dt \geq 0.$$

Moreover, when K is a convex body and μ is the uniform measure on K , one may take $\psi(t)$ to be a power of an affine function (heuristically $\psi(t) \propto L(t)^{n-1}$ for an affine L coming from slicing geometry).

Theorem 6.2 (Payne–Weinberger, 1960 [11]). *Let $K \subset \mathbb{R}^n$ be convex, and let μ_K be the uniform probability measure on K . Then for every locally Lipschitz $f : K \rightarrow \mathbb{R}$,*

$$\text{Var}_{\mu_K}(f) \leq \frac{\text{diam}(K)^2}{\pi^2} \int_K \|\nabla f(x)\|_2^2 d\mu_K(x),$$

where $\text{diam}(K) := \sup\{\|x - y\|_2 : x, y \in K\}$. Equivalently,

$$C_P(\mu_K) \leq \frac{\text{diam}(K)^2}{\pi^2}.$$

Moreover, the constant $\text{diam}(K)^2/\pi^2$ is sharp already in dimension 1 (on an interval), and thus cannot be improved if one only controls the diameter.

Proof (using localization). We outline the standard contradiction and localization argument that reduces the theorem to a one-dimensional weighted Poincaré inequality. Set $D := \text{diam}(K)$. Suppose the conclusion fails. Then there exists a locally Lipschitz function $f : K \rightarrow \mathbb{R}$ with mean zero $\int_K f d\mu_K = 0$ such that

$$\int_K f(x)^2 d\mu_K(x) > \frac{D^2}{\pi^2} \int_K \|\nabla f(x)\|_2^2 d\mu_K(x). \quad (22)$$

Equivalently,

$$\int_K \left(f(x)^2 - \frac{D^2}{\pi^2} \|\nabla f(x)\|_2^2 \right) d\mu_K(x) > 0.$$

Now define

$$F(x) := f(x), \quad G(x) := f(x)^2 - \frac{D^2}{\pi^2} \|\nabla f(x)\|_2^2.$$

Then $\int F d\mu_K = 0$ and $\int G d\mu_K > 0$. Apply the localization lemma (Lemma 6.1) to (F, G) under μ_K . We obtain a segment $x(t) = (1-t)a + tb$ and a log-concave weight $\psi(t)$ on $[0, 1]$ such that

$$\int_0^1 f(x(t)) \psi(t) dt = 0 \quad \text{and} \quad \int_0^1 \left(f(x(t))^2 - \frac{D^2}{\pi^2} \|\nabla f(x(t))\|_2^2 \right) \psi(t) dt > 0.$$

At this point we *compare gradients*. Along the segment, define the one-variable function

$$u(t) := f(x(t)).$$

Then $u'(t) = \langle \nabla f(x(t)), b - a \rangle$, so by Cauchy–Schwarz,

$$|u'(t)|^2 \leq \|b - a\|_2^2 \|\nabla f(x(t))\|_2^2 \leq D^2 \|\nabla f(x(t))\|_2^2,$$

since $\|b - a\|_2 \leq D$. Thus

$$\|\nabla f(x(t))\|_2^2 \geq \frac{|u'(t)|^2}{D^2}.$$

Plugging into the strict positivity inequality on the needle yields

$$\int_0^1 \left(u(t)^2 - \frac{D^2}{\pi^2} \cdot \|\nabla f(x(t))\|_2^2 \right) \psi(t) dt \leq \int_0^1 \left(u(t)^2 - \frac{1}{\pi^2} |u'(t)|^2 \right) \psi(t) dt.$$

Hence we deduce a one-dimensional strict violation:

$$\int_0^1 u(t) \psi(t) dt = 0 \quad \text{and} \quad \int_0^1 u(t)^2 \psi(t) dt > \frac{1}{\pi^2} \int_0^1 |u'(t)|^2 \psi(t) dt. \quad (23)$$

So the n -dimensional counterexample implies a 1D counterexample for a log-concave weight ψ . It therefore remains to rule out the one-dimensional violation (23). The key point is that for a log-concave density ψ on an interval $I = [a, b]$, the sharp one-dimensional Poincaré inequality holds. More precisely, if $|I| = \ell$ and u is absolutely continuous with

$$\int_I u \psi = 0,$$

then

$$\int_I u^2 \psi \leq \frac{\ell^2}{\pi^2} \int_I |u'|^2 \psi.$$

Since the interval I corresponds to a segment contained in K , its length satisfies

$$\ell \leq \text{diam}(K) = D.$$

Consequently,

$$\frac{\ell^2}{\pi^2} \leq \frac{D^2}{\pi^2}.$$

A convenient way to understand this inequality is to compare with the uniform measure on an interval. For the uniform measure on $[0, \ell]$, the sharp Poincaré constant is ℓ^2/π^2 , attained by the first nontrivial Neumann eigenfunction

$$u(t) = \cos\left(\frac{\pi t}{\ell}\right).$$

The Payne–Weinberger theorem shows that introducing a log-concave weight cannot increase the optimal constant. Equivalently, the first nonzero eigenvalue of the associated weighted Sturm–Liouville operator is at least π^2/ℓ^2 ; see [11].

Applying the one-dimensional weighted Poincaré inequality to the localized needle, whose underlying segment has length $\ell \leq D$, yields

$$\int_0^1 u(t)^2 \psi(t) dt \leq \frac{1}{\pi^2} \int_0^1 |u'(t)|^2 \psi(t) dt,$$

contradicting (23). This contradiction shows that the assumed counterexample cannot exist. Therefore,

$$\mathrm{Var}_{\mu_K}(f) \leq \frac{D^2}{\pi^2} \int_K \|\nabla f(x)\|_2^2 d\mu_K(x),$$

for every locally Lipschitz f . Hence

$$C_P(\mu_K) \leq \frac{D^2}{\pi^2}.$$

□

To check the sharpness of the constant, consider $K = [-R, R] \subset \mathbb{R}$ equipped with the uniform measure. Then $\mathrm{diam}(K) = 2R$, and the classical one-dimensional Poincaré inequality gives

$$C_P(\mu_K) = \frac{(2R)^2}{\pi^2} = \frac{4R^2}{\pi^2}.$$

Hence the Payne–Weinberger bound is attained in dimension one, showing that the constant $1/\pi^2$ is optimal.

Recall that A random vector $X \sim \mu$ is called *isotropic* if $\mathbb{E}X = 0$ and $\mathrm{Cov}(X) = I_n$. In particular,

$$\mathbb{E}\|X\|_2^2 = \mathrm{tr}(\mathrm{Cov}(X)) = n.$$

Thus, a typical Euclidean radius scale for an isotropic distribution is $\sqrt{n}$. For an isotropic convex body K (or isotropic log-concave measure), one indeed has strong concentration of $\|X\|_2$ around $\sqrt{n}$; for example, in [10], Paouris shows that if X is isotropic and log-concave, then

$$\mathbb{P}(\|X\|_2 \geq c\sqrt{n}t) \leq \exp(-c'\sqrt{n}t), \quad t \geq 1,$$

for universal constants $c, c' > 0$. However, Theorem 6.2 (Payne–Weinberger) only yields

$$C_P(\mu_K) \leq \frac{\mathrm{diam}(K)^2}{\pi^2}.$$

For an isotropic convex body, the diameter (or enclosing radius) may be of order $\sqrt{n}$, leading to the estimate

$$C_P(\mu_K) \lesssim n,$$

which grows linearly with the dimension. In contrast, Conjecture 2.19 (KLS) predicts that for *isotropic* log-concave measures one should have

$$C_P(\mu) \leq C$$

for a universal constant C , which remains open. A major breakthrough was obtained by Klartag [7], who has provided the best bound, proving that it holds true up to a factor of $\log n$.

7 Applications

The functional and geometric inequalities developed in the previous sections have important algorithmic consequences. One application arises in the problem of sampling and volume estimation for high-dimensional convex bodies. The analysis of random-walk sampling algorithms led Kannan, Lovász, and Simonovits to formulate the Kannan–Lovász–Simonovits (KLS) conjecture, which connects mixing times, conductance, and Poincaré inequalities.

7.1 Volume Estimation (convex bodies)

Definition 7.1 (Convex Body). The subset $K \subseteq \mathbb{R}^n$ is called an n -dimensional *convex body* if it is convex, compact, and has nonempty interior.

Let $K \subseteq \mathbb{R}^n$ be a convex body, and let $xB^n = \{y \in \mathbb{R}^n : \|y\|_2 \leq x\}$, i.e., a ball of radius x . Similarly, $1B^n = B^n$ is the unit ball. Furthermore, assume

$$rB^n \subset K \subset RB^n,$$

A problem in computational geometry is to compute the volume of a convex body. For a convex body $K \subseteq \mathbb{R}^n$, a membership procedure is specified by $R \geq r > 0$ and a point $x_0 \in K$ such that

$$x_0 + rB^n \subseteq K \subseteq RB^n$$

The procedure decides whether a queried point $x \in \mathbb{R}^n$ belongs to K . Estimate

$$\text{Vol}(K) \approx \frac{\#\{x_i \in K\}}{m} \cdot \text{Vol}(RB^n).$$

Indeed,

$$\frac{\#\{x_i \in K\}}{m} = \frac{1}{m} \sum_{i=1}^m \mathbb{1}\{x_i \in K\}.$$

Since the points x_i are sampled independently and uniformly from RB^n , the random variables $\mathbb{1}\{x_i \in K\}$ are i.i.d. Bernoulli with mean

$$\mathbb{E}[\mathbb{1}\{x_i \in K\}] = \mathbb{P}(x_i \in K) = \frac{\text{Vol}(K)}{\text{Vol}(RB^n)}.$$

Therefore, by the Law of Large Numbers,

$$\frac{\#\{x_i \in K\}}{m} \longrightarrow \frac{\text{Vol}(K)}{\text{Vol}(RB^n)} \quad \text{as } m \rightarrow \infty.$$

But this method suffers from the *curse of dimensionality*, because the probability of success, i.e.,

$$p := \mathbb{P}_{X \sim \text{Unif}(RB^n)}(X \in K) = \frac{\text{Vol}(K)}{\text{Vol}(RB^n)},$$

can be exponentially small in n in the worst case, forcing an exponentially large sample size to obtain a reliable estimate. Indeed, if

$$p = e^{-cn}$$

for some constant $c > 0$, then the expected number of sampled points that fall inside K is

$$\mathbb{E}[\#\{x_i \in K\}] = mp.$$

Therefore, to obtain even a constant expected number of “successful” samples, one requires

$$m \asymp \frac{1}{p} = e^{cn},$$

which grows exponentially with the dimension. Consequently, this method, also known as the Monte Carlo estimator, may require exponentially many samples before it can accurately estimate $\text{Vol}(K)$.

The following remark can illustrate how exponentially small volume ratios naturally arise in high dimensions.

Remark 7.2 (Boundary concentration for the Euclidean ball). Let $B^n = \{x \in \mathbb{R}^n : \|x\|_2 \leq 1\}$ be the unit Euclidean ball, and fix $\varepsilon \in (0, 1)$. Then

$$\frac{\text{Vol}((1 - \varepsilon)B^n)}{\text{Vol}(B^n)} = (1 - \varepsilon)^n.$$

Since

$$(1 - \varepsilon)^n \leq e^{-n\varepsilon},$$

the fraction of the volume contained in the smaller concentric ball decays exponentially fast with n .

Equivalently, if $X \sim \text{Unif}(B^n)$, then

$$\mathbb{P}(\|X\|_2 \leq 1 - \varepsilon) = (1 - \varepsilon)^n \leq e^{-n\varepsilon}.$$

Thus, for large n , almost all of the volume of B^n lies in the thin outer shell

$$\{x \in \mathbb{R}^n : 1 - \varepsilon \leq \|x\|_2 \leq 1\}.$$

This demonstrates that exponentially small volume ratios are common in high dimensions, which explains why the naive Monte Carlo approach can become computationally infeasible.

The example above suggests that directly estimating the volume of a convex body through a single Monte Carlo experiment can be expensive in high dimensions. A better strategy is to avoid estimating one extremely small probability and instead decompose the volume computation into a sequence of easier estimation problems. The key idea is to compute $\text{Vol}(K)$ via a product of *moderate* ratios rather than one tiny probability. Dyer, Frieze, and Kannan [2] introduced a randomized polynomial-time algorithm based on the following successive approximation of $\text{Vol}(K)$.

Let $\{r_i \mid i = 0, \dots, m\} \subseteq [r, R]$ be an increasing sequence of radii, i.e.,

$$r = r_0 < r_1 < r_2 < \dots < r_m = R,$$

and define the nested bodies

$$K_i = K \cap r_i B, \quad (\text{so } rB = K_0 \subseteq \dots \subseteq K_m = K)$$

Then

$$\prod_{i=1}^m \frac{\text{Vol}(K_i)}{\text{Vol}(K_{i-1})} = \frac{\text{Vol}(K_m)}{\text{Vol}(K_0)} = \frac{\text{Vol}(K)}{r^n \text{Vol}(B)}$$

For example, we can use a small multiplicative expansion per phase, i.e.,

$$r_i = \left(1 + \frac{1}{n}\right) \cdot r_{i-1}.$$

so

$$\text{Vol}(r_i B) = \left(1 + \frac{1}{n}\right)^n \cdot \text{Vol}(r_{i-1} B)$$

For the multiphase estimator to be practical, we must be able to estimate each ratio

$$\frac{\text{Vol}(K_{i-1})}{\text{Vol}(K_i)}.$$

This can be done probabilistically. If $X \sim \text{Unif}(K_i)$, then

$$\frac{\text{Vol}(K_{i-1})}{\text{Vol}(K_i)} = \Pr(X \in K_{i-1}).$$

Thus, by drawing many approximately uniform samples from K_i , we can estimate this ratio by the empirical fraction of samples that fall in K_{i-1} .

This reduces the volume estimation problem to a sampling problem: how can we efficiently generate approximately uniform samples from a high-dimensional convex body? Random walks provide a powerful approach. The idea is to construct a Markov chain whose stationary distribution is the uniform measure on K_i . If the chain converges rapidly to this stationary distribution, then the samples produced by the walk can be used to estimate the volume ratios above.

A simple but important example is the ball walk. It was introduced and analyzed by Lovász and Simonovits [8] as a Markov-chain method for sampling from convex bodies.

- Start with an arbitrary point in the convex body K ($x_0 \in K$).
- Given the current state $x_t \in K$, sample x_{t+1} according to the following rule.

$$\tilde{x}_{t+1} = x_t + y_t,$$

where $y_t \sim \text{Unif}(\delta B^n)$ (uniform on a ball of tiny radius δ).

- Define

$$x_{t+1} = \begin{cases} \tilde{x}_{t+1}, & \text{if } \tilde{x}_{t+1} \in K, \\ x_t, & \text{otherwise.} \end{cases}$$

Question. How fast does this work? What can go wrong?

Ball walk is *local*: it moves only within distance δ at each step. Consequently, the geometry of the convex body strongly influences how rapidly the walk can explore the state space. Recall from Section 2.4 that conductance measures how easily probability mass can move from a set to its complement. For a subset $S \subseteq K$, the conductance is determined by the size of the boundary

separating S from $K \setminus S$ relative to the volume of these sets. Small conductance corresponds to a narrow bottleneck, while large conductance indicates that every sufficiently large set has substantial boundary. A theorem of Lovász and Simonovits [8] shows that such bottlenecks are essentially the only obstruction to rapid mixing. Thus, in order to understand the mixing behavior of the ball walk, it is natural to study how boundary size can be quantified geometrically. This motivates the study of linear cuts of convex bodies and their boundary volume.

7.2 Linear cuts and boundary volume

To understand why conductance governs the mixing behavior of the ball walk, it is useful to examine how boundary size is measured. For convex bodies, a simple and geometrically transparent class of cuts is given by hyperplane cuts. For simplicity, we illustrate the main ideas in the two-dimensional case. The higher-dimensional analogue is similar. The goal is to show how the boundary length of a linear cut can be recovered from the volume of a thin tube around the cut, and then interpreted probabilistically through a one-dimensional projection.

Suppose we have an isotropic convex body $K \subseteq \mathbb{R}^2$ with the uniform measure over it, and let $X = (X_1, X_2)$ denote a random vector distributed uniformly on K . For simplicity, instead of considering arbitrary measurable cuts $S \subseteq K$, we restrict attention to *linear cuts*, that is, cuts induced by a hyperplane (a line in $\mathbb{R}^2$). Fix a line that intersects the interior of K ,

$$H = \{x \in \mathbb{R}^2 : u^\top x = c\},$$

where $\|u\| = 1$ is a unit normal vector to the 1-dimensional hyperplane (line). This line partitions K into two measurable subsets,

$$S = K \cap \{x \in \mathbb{R}^2 : u^\top x \leq c\}, \quad K \setminus S = K \cap \{x \in \mathbb{R}^2 : u^\top x \geq c\}.$$

Geometrically, S and $K \setminus S$ are the two regions obtained by cutting the convex body K with the line H , as illustrated in Figure 1.

The boundary between these two parts inside K is the line segment

$$\partial S \cap \text{int}(K) = H \cap K,$$

which is a one-dimensional set. Since the conductance involves the $(n-1)$ -dimensional “volume” of this boundary (which, in two dimensions, is its length), it is convenient to express this quantity in terms of ordinary two-dimensional area by a limiting procedure.

Approximation of the boundary. For $\varepsilon > 0$, define the ε -neighborhood of the hyperplane H by

$$T_\varepsilon = \{x \in \mathbb{R}^2 : \text{dist}(x, H) \leq \varepsilon\}.$$

Since $\|u\|_2 = 1$, we have $\text{dist}(x, H) = |u^\top x - c|$, so equivalently

$$T_\varepsilon(H) = \{x \in \mathbb{R}^2 : |u^\top x - c| \leq \varepsilon\}.$$

The set $K \cap T_\varepsilon(H)$ is a thin strip of width 2ε around the chord $H \cap K$. Its area satisfies

$$\text{vol}_2(K \cap T_\varepsilon(H)) = 2\varepsilon \text{length}(H \cap K) + o(\varepsilon) \quad \text{as } \varepsilon \rightarrow 0.$$

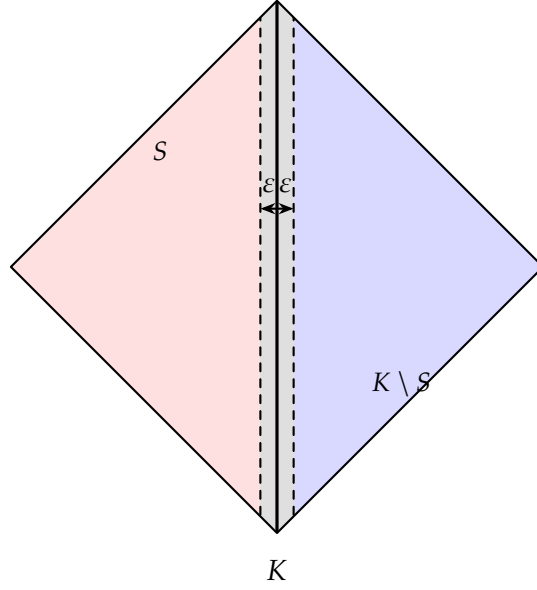


Figure 1: A linear cut of a convex body K by the line $H = \{u^\top x = c\}$, producing the two sets S and $K \setminus S$. The shaded strip represents the ε -tube T_ε around H .

Consequently, the one-dimensional volume (length) of the boundary segment $H \cap K$ can be computed as the limit

$$\text{length}(H \cap K) = \lim_{\varepsilon \rightarrow 0} \frac{1}{2\varepsilon} \text{vol}_2(K \cap T_\varepsilon(H)). \quad (24)$$

Now if $Z \sim \text{Unif}(K)$, then

$$\mathbb{P}(Z \in K \cap T_\varepsilon(H)) = \frac{\text{vol}_2(K \cap T_\varepsilon(H))}{\text{vol}_2(K)}.$$

Combining with (24),

$$\lim_{\varepsilon \rightarrow 0} \frac{1}{2\varepsilon} \mathbb{P}(Z \in K \cap T_\varepsilon) = \frac{\text{length}_1(H \cap K)}{\text{vol}_2(K)}.$$

Now define the *one-dimensional signed distance to the hyperplane*

$$X := u^\top Z - c \in \mathbb{R}.$$

Then, by the characterization $T_\varepsilon = \{x : |u^\top x - c| \leq \varepsilon\}$,

$$\mathbb{P}(|X| \leq \varepsilon) = \mathbb{P}(Z \in T_\varepsilon(H)),$$

and hence, provided X admits a density f_X that is continuous at 0,

$$\lim_{\varepsilon \rightarrow 0} \frac{1}{2\varepsilon} \mathbb{P}(|X| \leq \varepsilon) = \lim_{\varepsilon \rightarrow 0} \frac{1}{2\varepsilon} \mathbb{P}(Z \in T_\varepsilon(H)) = f_X(0).$$

In this way, the boundary volume can be interpreted probabilistically as the density at 0 of the signed-distance variable X (equivalently, the density of the projection $u^\top Z$ evaluated at c).

Connection to conductance. We can now return to conductance. For a linear cut induced by

$$H = \{u^\top x = c\},$$

the conductance numerator is the boundary volume

$$\text{Vol}_1(\partial S \cap \text{int}(K)) = \text{length}(H \cap K).$$

The calculation above shows that this boundary term can be computed geometrically through the tube limit

$$\text{length}(H \cap K) = \lim_{\varepsilon \rightarrow 0} \frac{1}{2\varepsilon} \text{Vol}_2(K \cap T_\varepsilon(H)),$$

or probabilistically through the small-ball probability

$$\lim_{\varepsilon \rightarrow 0} \frac{1}{2\varepsilon} \mathbb{P}(|X| \leq \varepsilon), \quad X = u^\top Z - c,$$

where $Z \sim \text{Unif}(K)$.

Thus, for hyperplane cuts, the geometric boundary term appearing in conductance can be interpreted as the density at zero of the signed-distance variable $X = u^\top Z - c$, or equivalently as the density of the one-dimensional projection $u^\top Z$ evaluated at c . This explains why one-dimensional projections naturally appear when studying conductance and mixing for convex bodies.

The discussion above connects the sampling problem back to the functional and isoperimetric inequalities studied earlier. In Section 5, we saw that Cheeger-type inequalities quantify the expansion of sets, and that for log-concave measures the Cheeger and Poincaré viewpoints are equivalent up to universal constants. Thus, understanding conductance for convex bodies is closely related to understanding Poincaré inequalities for the corresponding uniform measures. This leads naturally to the following question.

Question. If $K \subset \mathbb{R}^n$ is a convex body and μ_K denotes the uniform probability measure on K , does μ_K satisfy a Poincaré inequality with a dimension-free constant? More generally, does every isotropic log-concave probability measure on $\mathbb{R}^n$ satisfy such an inequality?

This is precisely the content of the Kannan–Lovász–Simonovits conjecture introduced in Section 2.4. The Payne–Weinberger theorem proved in Section 6 gives a general bound in terms of the diameter of the support, namely

$$C_P(\mu_K) \leq \frac{\text{diam}(K)^2}{\pi^2}.$$

If $K \subseteq RB_2^n$, this yields

$$C_P(\mu_K) \leq \frac{4R^2}{\pi^2}.$$

However, in isotropic scaling one typically has R^2 of order n , so this bound grows with the dimension. The KLS conjecture predicts a much stronger dimension-free bound for isotropic log-concave measures. Thus, the study of volume algorithms, random walks, conductance, and Poincaré inequalities all lead to the same central geometric question.

This interplay between algorithms, geometry, and functional inequalities is one of the main reasons the KLS conjecture occupies a central position in high-dimensional probability.

8 Conclusion

Functional inequalities provide a framework for understanding a wide range of phenomena in probability theory, analysis, and convex geometry. Although the inequalities studied in this work arise in different settings and possess distinct formulations, they are all manifestations of a common principle: geometric structure imposes strong constraints on the behavior of functions and measures.

We began with the Poincaré inequality, which quantifies how the variance of a function can be controlled by its gradient and serves as a tool in the study of concentration of measure. We then turned to the Prékopa–Leindler inequality, a functional analogue of the Brunn–Minkowski inequality, which illustrates connections between convexity, volume, and log-concavity. Building on these ideas, we studied logarithmic Sobolev inequalities, which strengthen variance bounds by controlling entropy and lead to powerful concentration results. Finally, we examined Cheeger-type inequalities and the Kannan–Lovász–Simonovits conjecture, showing the relationship between geometric connectivity, isoperimetry, and functional inequalities.

A recurring theme throughout this work is that high-dimensional problems often become tractable only after identifying the appropriate geometric structure. Convexity, log-concavity, and isoperimetric principles repeatedly emerge as the mechanisms that govern concentration phenomena and the behavior of probability measures in large dimensions.

Another theme is the role of dimensional reduction. Techniques such as localization make it possible to reduce high-dimensional questions to one-dimensional problems, where the structure becomes more transparent. This perspective has proved remarkably successful in modern convex geometry and continues to play an important role in current research.

Despite substantial progress the Kannan–Lovász–Simonovits conjecture continues to be one of the central open problems in high-dimensional geometry and probability. Its resolution would significantly deepen the understanding of the relationship between geometry, concentration, and mixing phenomena for log-concave measures.

The study of functional inequalities remains an active and developing area of research. Beyond their mathematical interest, these inequalities continue to find applications in fields ranging from stochastic processes and partial differential equations to statistics, machine learning, and data science. As a result, they provide not only a collection of powerful technical tools, but also a unifying language through which many different areas of modern mathematics can be understood.

References

- [1] Stéphane Boucheron, Gábor Lugosi, and Olivier Bousquet. Concentration inequalities. In *Advanced Lectures on Machine Learning*, pages 208–240. Springer, 2004.
- [2] Martin E. Dyer, Alan M. Frieze, and Ravi Kannan. A random polynomial time algorithm for approximating the volume of convex bodies. In *Symposium on the Theory of Computing*, 1989.
- [3] Nassif Ghoussoub and Amir Moradifam. *Functional Inequalities: New Perspectives and New Applications*, volume 187 of *Mathematical Surveys and Monographs*. American Mathematical Society, 2013.

-
- [4] Mikhail Gromov and Vitali D. Milman. Generalization of the spherical isoperimetric inequality to uniformly convex banach spaces. *Compositio Mathematica*, 62(3):263–282, 1987.
 - [5] Ravi Kannan, László Lovász, and Miklós Simonovits. Isoperimetric problems for convex bodies and a localization lemma. *Discrete & Computational Geometry*, 13(3):541–559, 1995.
 - [6] Ravindran Kannan, László Lovász, and Miklós Simonovits. Isoperimetric problems for convex bodies and a localization lemma. *Discrete & Computational Geometry*, 13(3–4):541–559, 1995.
 - [7] Bo’az Klartag. Logarithmic bounds for isoperimetry and slices of convex sets. *Ars Inveniendi Analytica*, (4):1–17, 2023.
 - [8] László Lovász and Miklós Simonovits. Random walks in a convex body and an improved volume algorithm. *Random Structures & Algorithms*, 4(4):359–412, 1993.
 - [9] Emanuel Milman. On the role of convexity in functional and isoperimetric inequalities. *Proceedings of the London Mathematical Society*, 99(1):32–66, 2009.
 - [10] Grigoris Paouris. Concentration of mass on convex bodies. *Geometric and Functional Analysis*, 16(5):1021–1049, 2006.
 - [11] L. E. Payne and H. F. Weinberger. An optimal Poincaré inequality for convex domains. *Archive for Rational Mechanics and Analysis*, 5(1):286–292, 1960.
 - [12] H. Poincaré. Sur les equations aux dérivées partielles de la physique mathématique. *American Journal of Mathematics*, 12(3):211–294, 1890.
 - [13] Adrien Saumard and Jon A. Wellner. Log-concavity and strong log-concavity: A review. *Statistics Surveys*, 8:45–114, 2014.