

Testing high-d distributions with subcube conditioning

Erik Waingarten (Penn CIS)
ewaingar@seas.upenn.edu

A. Levi
(Haifa)

X. Chen
(Columbia)

G. Kamath
(NYU)

C. Canonne
(Sydney)

R. Jayaram
(Google)

S. Ristic
(Penn)

D. Chakrabarty
(Dartmouth)

C. Seshadhri
(UC Santa Cruz)

Brief Outline:

1. Motivation: distribution testing and the curse of dimensionality.
2. Property testing of distributions.
3. New results.

Main Premise:

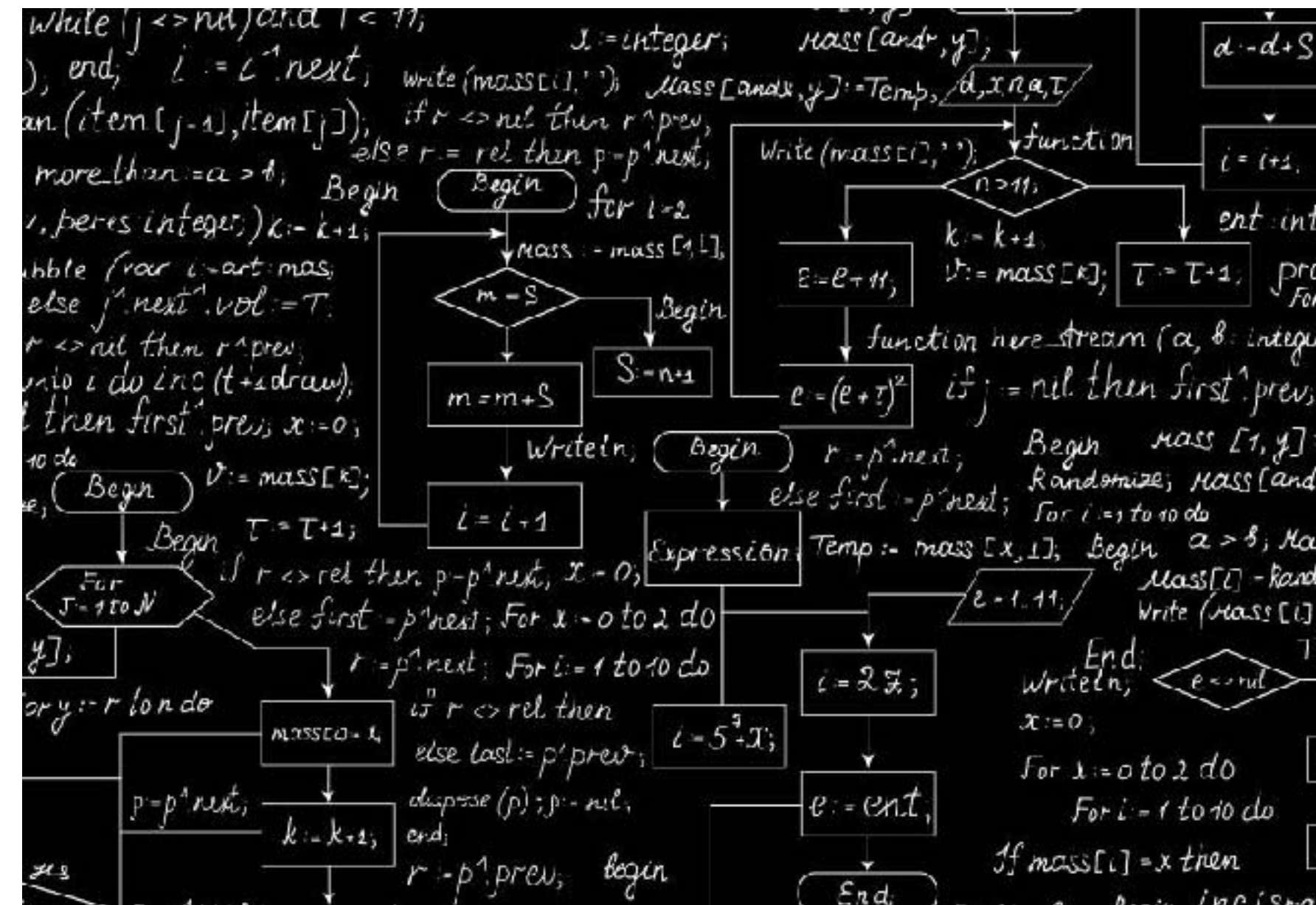
Testing high-d discrete distributions
by *algorithmically* interacting with the distribution.

Types of *Discrete, High-d* distributions

Coming from
Nature



Software or
Generative Model



Discrete Distribution from
Enormous collection

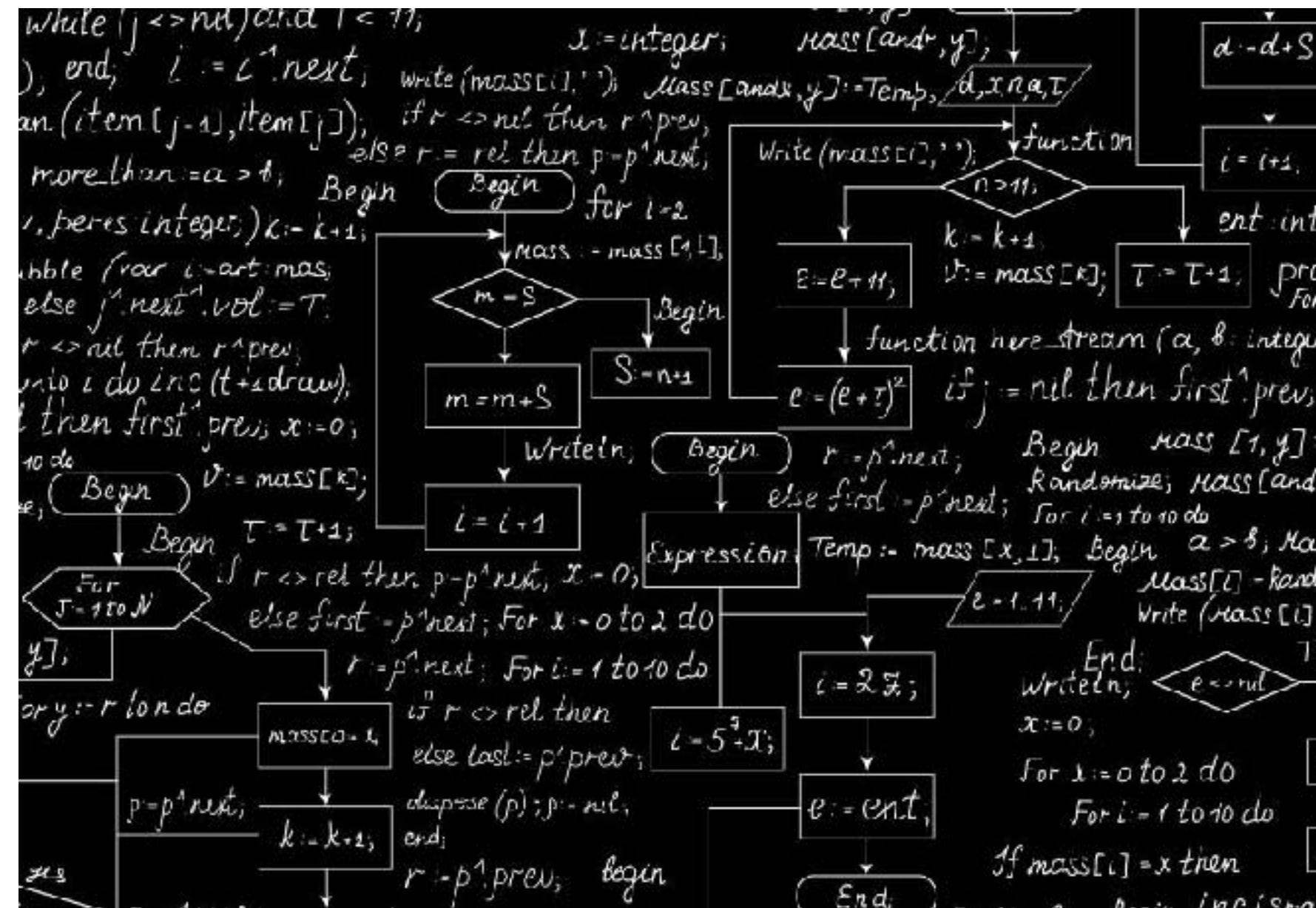


Types of *Discrete, High-d* distributions

Coming from Nature



Software or Generative Model



Discrete Distribution from Enormous collection



Algorithmic question: How to extract information from these high-d distributions?

Which is uniform on $\{0, 1\}^{15}$?

$\mathcal{D}_1$

0	1	0	1	0	1	1	0	0	1	0	1	1	0	0
1	0	1	1	0	1	1	1	1	0	0	1	1	1	1
0	0	0	0	1	1	0	1	1	0	0	1	0	1	0
1	1	1	0	0	1	1	1	0	0	0	0	0	1	1
1	1	0	1	1	1	0	1	1	1	1	0	1	1	1
0	0	0	0	1	1	0	1	1	0	0	1	0	0	0
1	0	1	1	0	1	1	1	1	1	1	1	0	0	1
1	1	1	0	1	0	0	1	0	1	0	0	0	1	0

$\mathcal{D}_2$

0	1	1	1	0	1	0	1	0	0	0	0	1	0	0
1	1	1	0	0	1	0	1	0	1	0	1	0	0	1
1	1	0	1	0	1	0	0	1	0	0	0	1	0	0
1	1	0	1	0	1	1	0	0	1	0	1	1	0	0
0	1	1	0	1	1	0	0	1	1	1	0	0	1	0
1	1	0	1	1	1	0	0	0	1	0	0	1	1	0
0	1	0	0	0	0	1	1	0	1	0	1	1	1	1
1	0	1	0	0	1	0	0	1	1	1	0	0	0	0
0	1	0	0	1	0	0	1	0	1	0	0	1	0	1

Which is uniform on $\{0, 1\}^{15}$?

$\mathcal{D}_1$

0	1	0	1	0	1	1	0	0	1	0	1	1	0	0
1	0	1	1	0	1	1	1	1	0	0	1	1	1	1
0	0	0	0	1	1	0	1	1	0	0	1	0	1	0
1	1	1	0	0	1	1	1	0	0	0	0	0	1	1
1	1	0	1	1	1	0	1	1	1	1	0	1	1	1
0	0	0	0	1	1	0	1	1	0	0	1	0	0	0
1	0	1	1	0	1	1	1	1	1	1	1	0	0	1
1	1	1	0	1	0	0	1	0	1	0	0	0	1	0

$\mathcal{D}_2$

0	1	1	1	0	1	0	1	0	0	0	0	1	0	0
1	1	1	0	0	1	0	1	0	1	0	1	0	0	1
1	1	0	1	0	1	0	0	1	0	0	0	1	0	0
1	1	0	1	0	1	1	0	0	1	0	1	1	0	0
0	1	1	0	1	1	0	0	1	1	1	0	0	1	0
1	1	0	1	1	1	0	0	0	1	0	0	1	1	0
0	1	0	0	0	0	1	1	0	1	0	1	1	1	1
1	0	1	0	0	1	0	0	1	1	1	0	0	0	0
0	1	0	0	1	0	0	1	0	1	0	0	1	0	1

$$x \sim \mathcal{D}_2 : \sum_{i=1}^{15} x_i \text{ even}$$

How to efficiently extract information from a distribution?

Question 1: Is it uniform/the distribution you are expecting? [GR'00, Paninski'08]

Question 2: Are two distributions the same? [BFRSW '00]

Question 3: Estimate properties of distributions ... [VV'11]

How to efficiently extract information from a distribution?

Question 1: Is it uniform/the distribution you are expecting? [GR'00, Paninski'08]

Question 2: Are two distributions the same? [BFRSW '00]

Question 3: Estimate properties of distributions ... [VV'11]

Algorithmic Message: usable information in collision statistics,
often much more efficient than learning.

Birthday Paradox

Theorem (Paninski'08): Testing whether a distribution $\mathcal{D}$ is the uniform dist. (or any fixed reference) over Ω can be done with $\sqrt{|\Omega|}$ samples, and is tight.

Birthday Paradox

Theorem (Paninski'08): Testing whether a distribution $\mathcal{D}$ is the uniform dist. (or any fixed reference) over $\{\pm 1\}^d$ can be done with $2^{d/2}$ samples, and is tight.

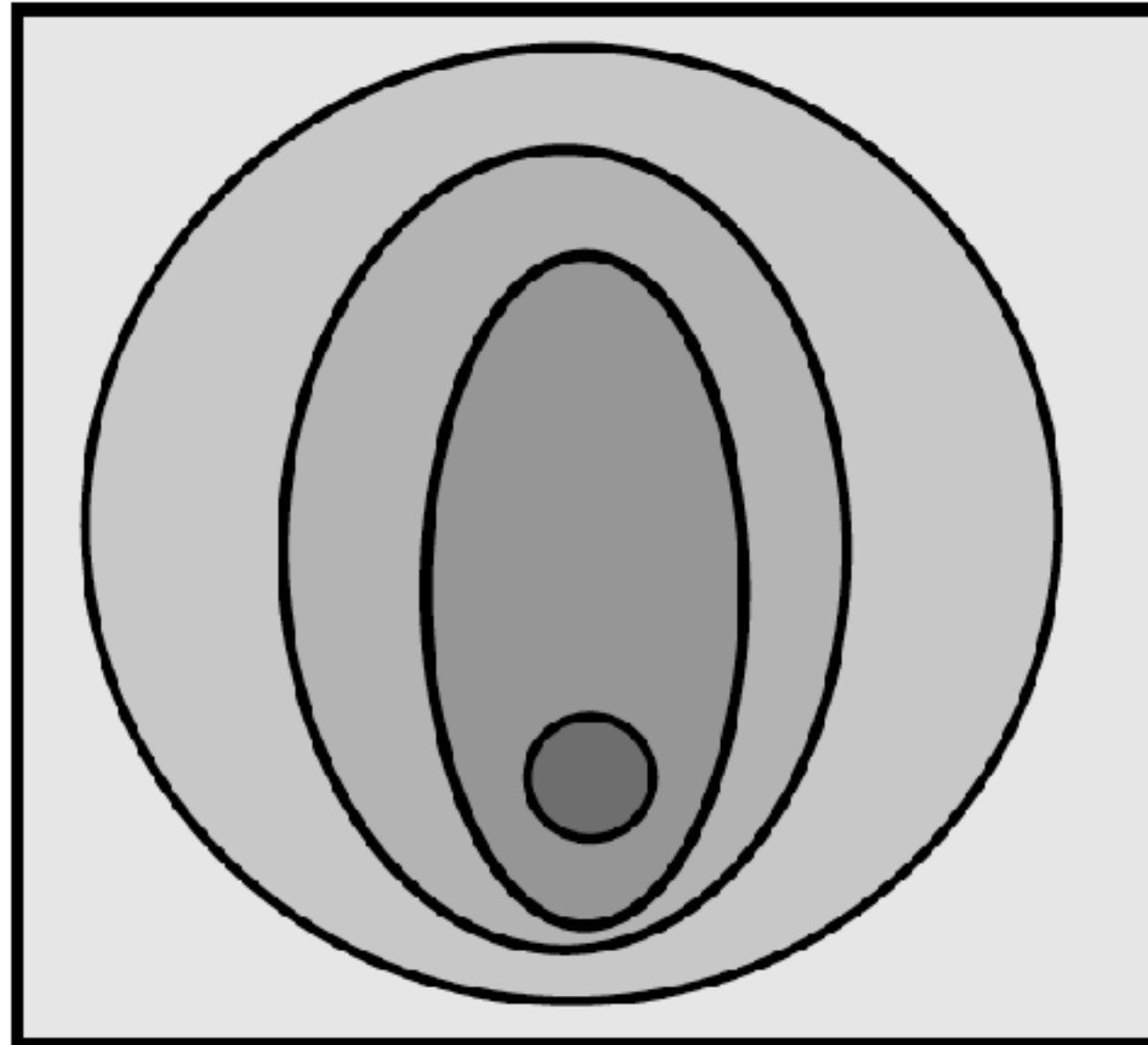
Birthday Paradox

Theorem (Paninski'08): Testing whether a distribution $\mathcal{D}$ is the uniform dist. (or any fixed reference) over $\{\pm 1\}^d$ can be done with $2^{d/2}$ samples, and is tight.

Two Things:

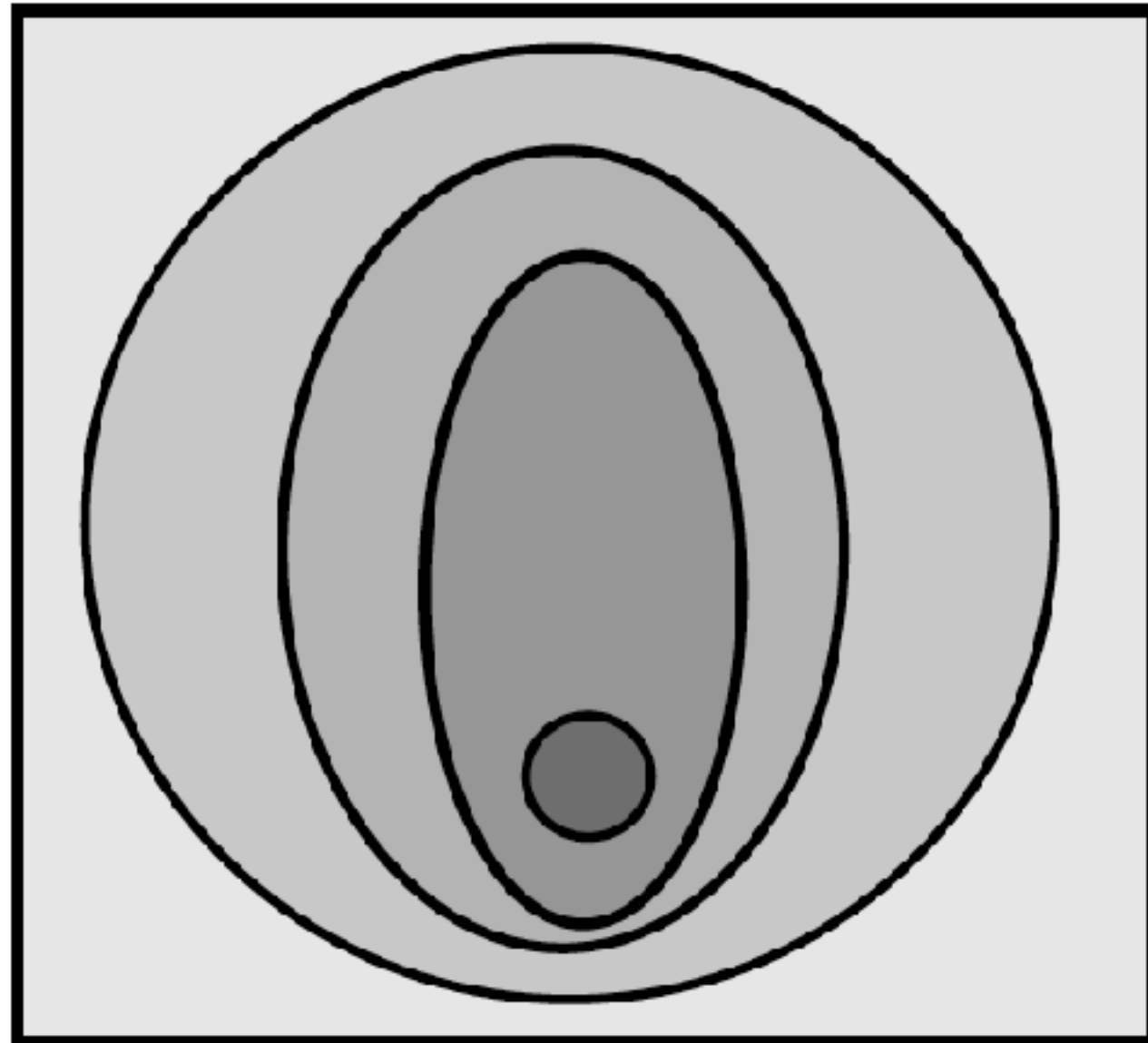
1. Impressive sublinear bound and made it unimpressive.
2. Assigned meaning to coordinates.

Beyond “curse of dimensionality”...

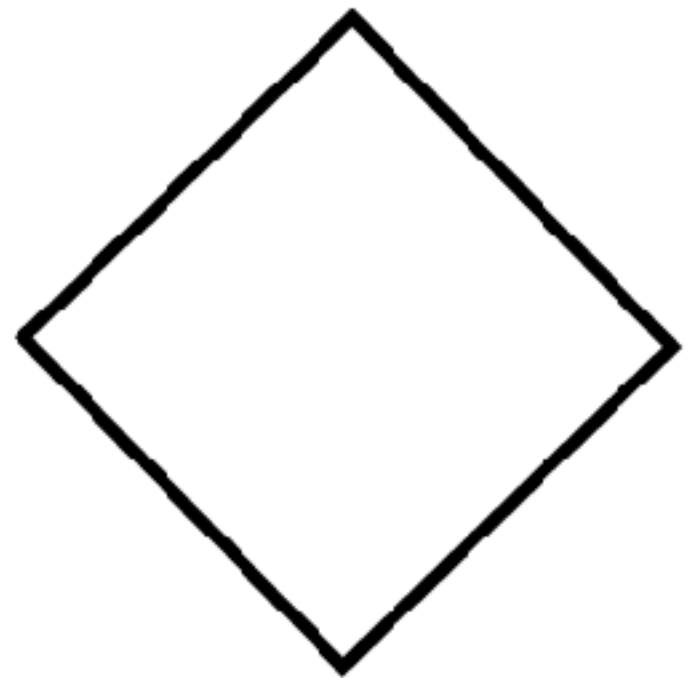


Restricted Inputs
(Ising models, Bayes nets,
small mixtures, monotone dist)

Beyond “curse of dimensionality”...



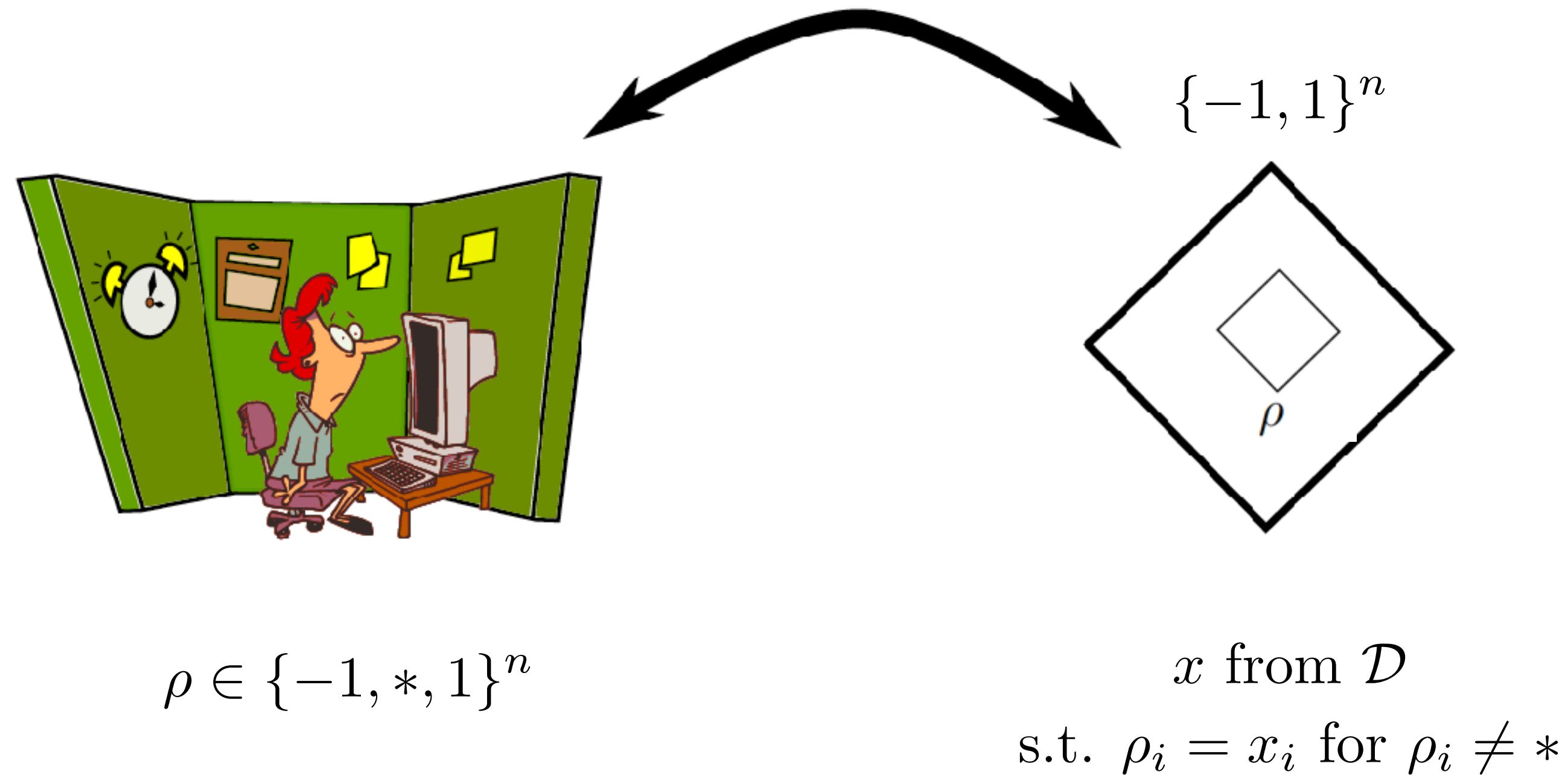
Restricted Inputs
(Ising models, Bayes nets,
small mixtures, monotone dist)



$$\{-1, 1\}^n$$

Stronger Access

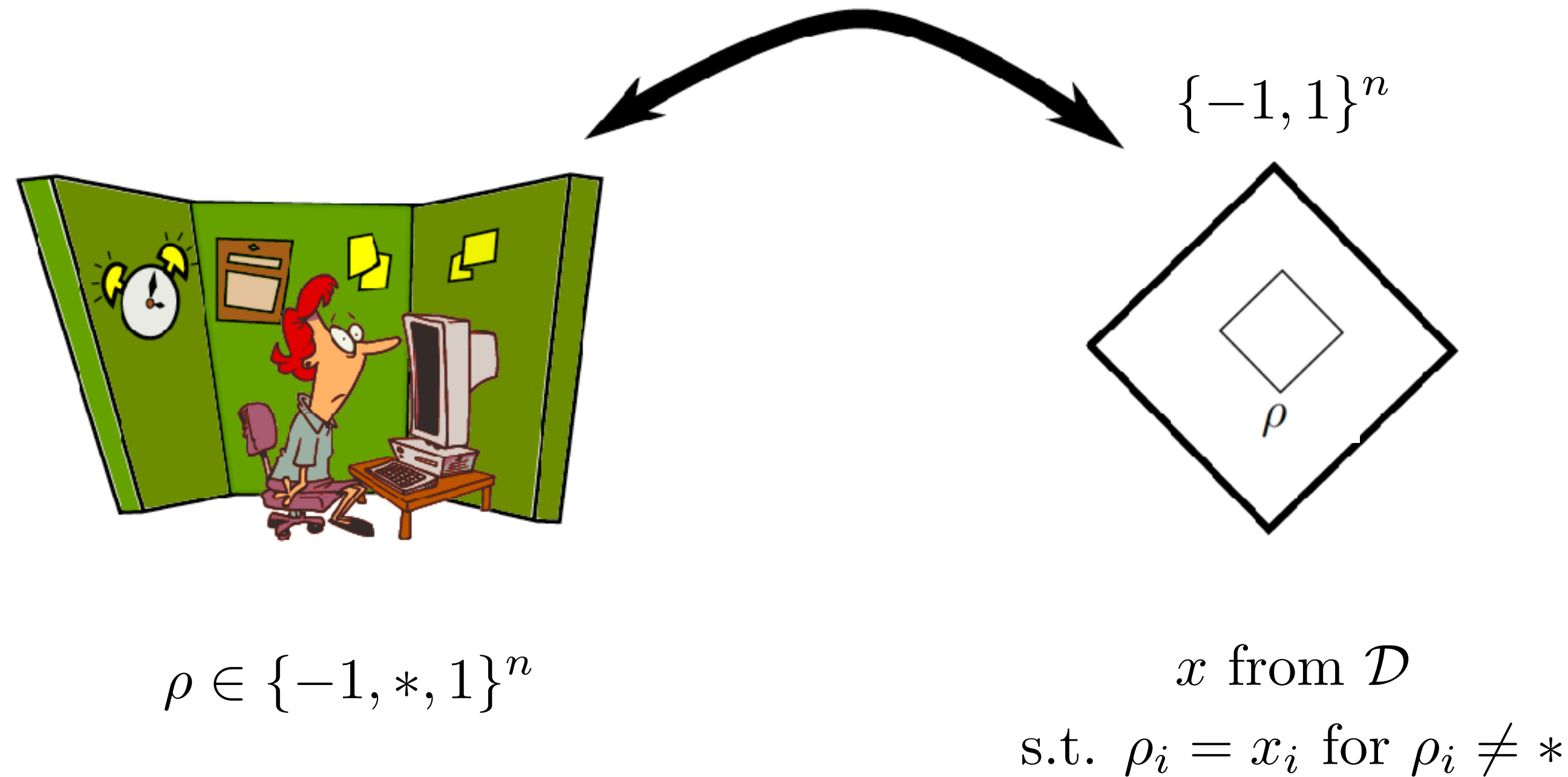
This talk's access: Subcube Conditional Sampling [Bhattacharya, Chakraborty '18]



$\rho =$	-1	1	-1	*	*	1	1	*	-1	*	1	1	1	1	-1
$x =$	-1	1	-1	1	-1	1	1	1	-1	-1	1	1	1	1	-1

Prefix [Adar, Fisher, Levi '25]

This talk's access: ~~Subcube~~ **Conditional Sampling** [Bhattacharya, Chakraborty '18]

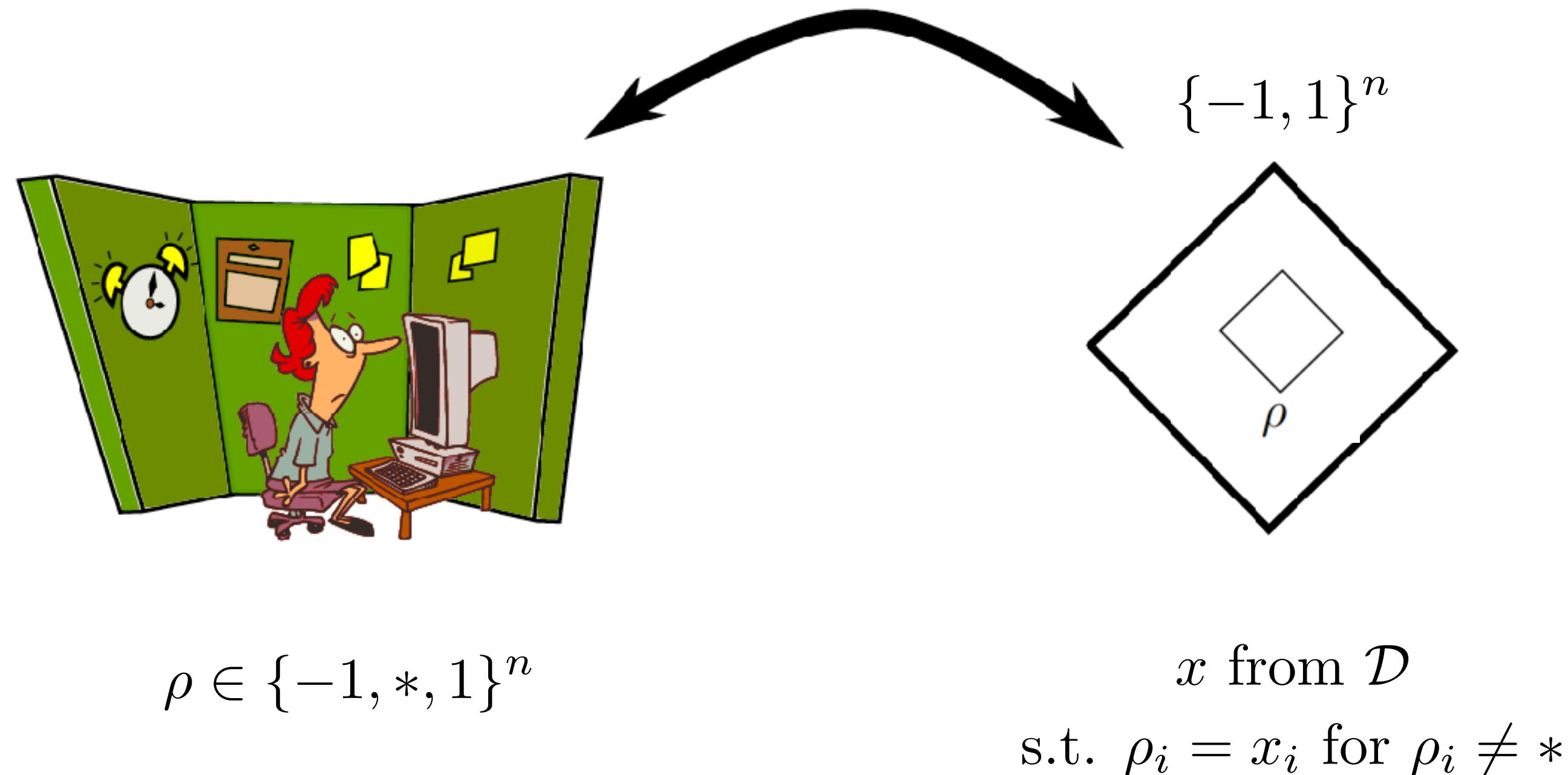


$\rho =$	-1	1	-1	1	*	*	*	*	*	*	*	*	*	*	*
$x =$	-1	1	-1	1	-1	1	1	1	-1	-1	1	1	1	1	-1

Prefix [Adar, Fisher, Levi '25]

This talk's access: Subcube Conditional Sampling [Bhattacharya, Chakraborty '18]

Coordinate [Blanca, Chen, Stefanovic, Vigoda '23]



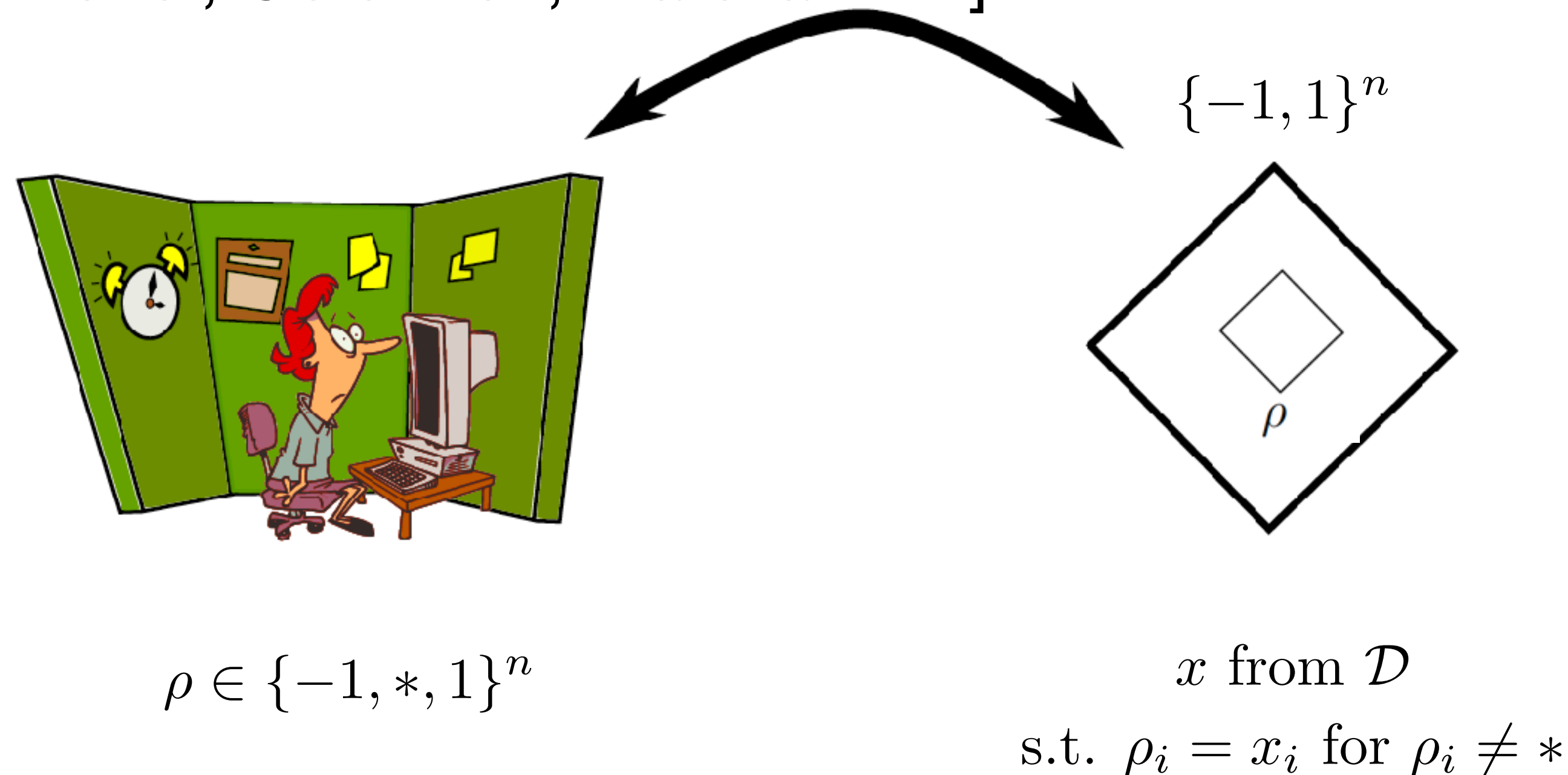
$\rho =$	-1	1	-1	1	*	1	1	1	-1	-1	1	1	1	1	-1
$x =$	-1	1	-1	1	-1	1	1	1	-1	-1	1	1	1	1	-1

Prefix [Adar, Fisher, Levi '25]

This talk's access: ~~Subcube~~ **Conditional Sampling** [Bhattacharya, Chakraborty '18]

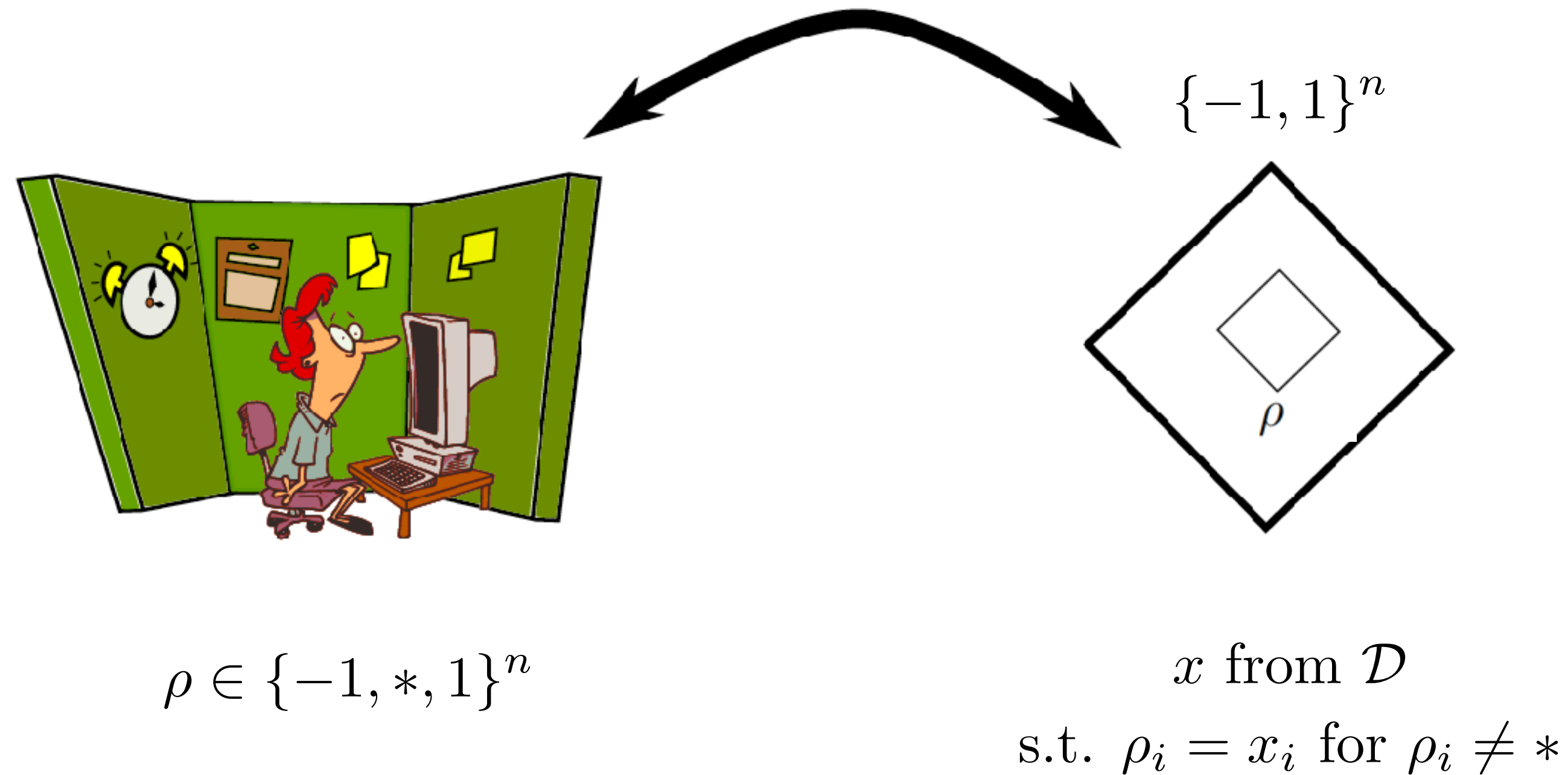
~~Coordinate~~ [Blanca, Chen, Stefanovic, Vigoda '23]

(General) Conditioning [Canonne, Ron, Servedio'12;
Chakraborty, Fisher, Goldhirsh, Matsliah '12]



$\rho =$	-1	1	-1	1	*	1	1	1	-1	-1	1	1	1	1	-1
$x =$	-1	1	-1	1	-1	1	1	1	-1	-1	1	1	1	1	-1

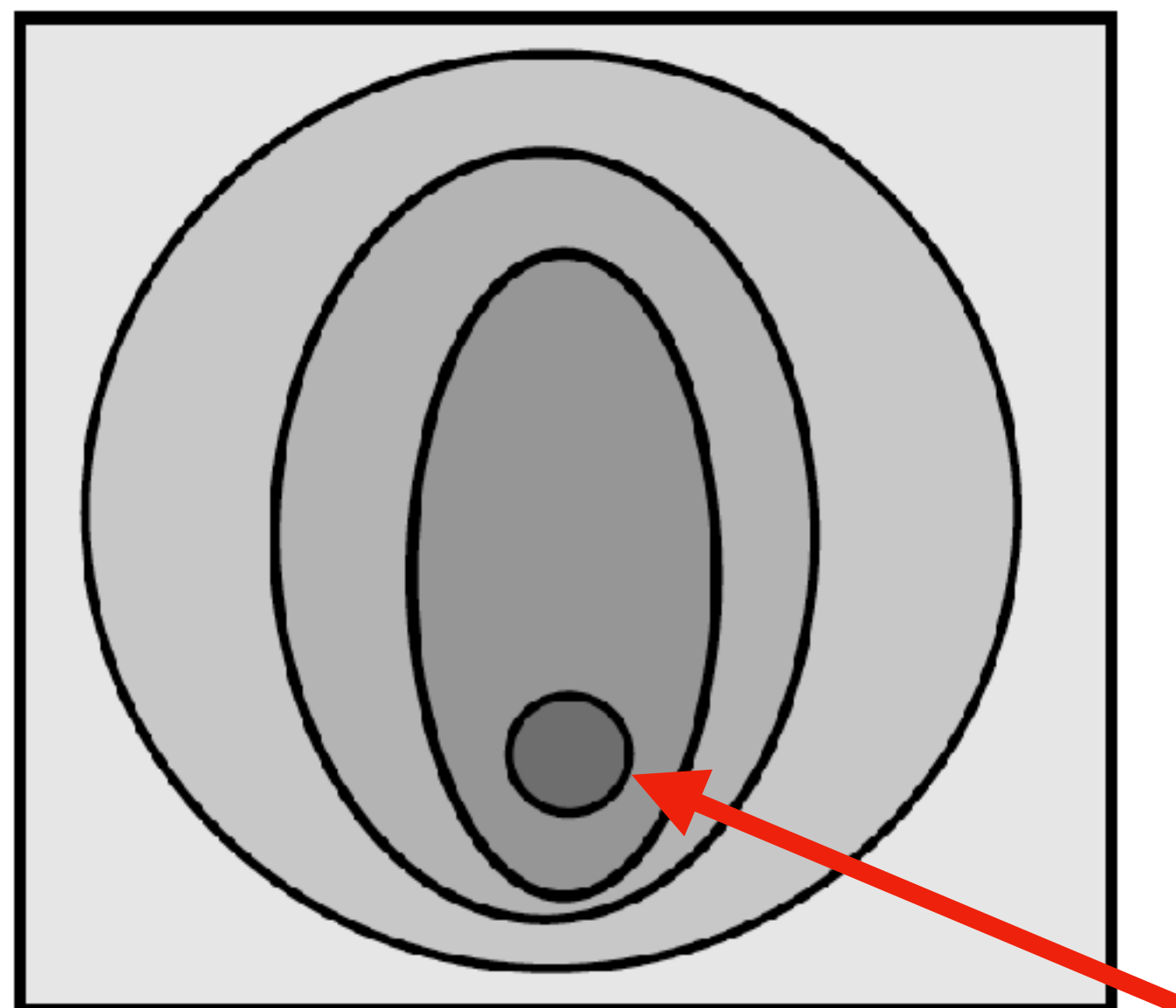
This talk's access: Subcube Conditional Sampling [Bhattacharya, Chakraborty '18]



Algorithmic Message: usable information in marginals of distribution, but only after conditioning on random subcubes.

What we can currently do...

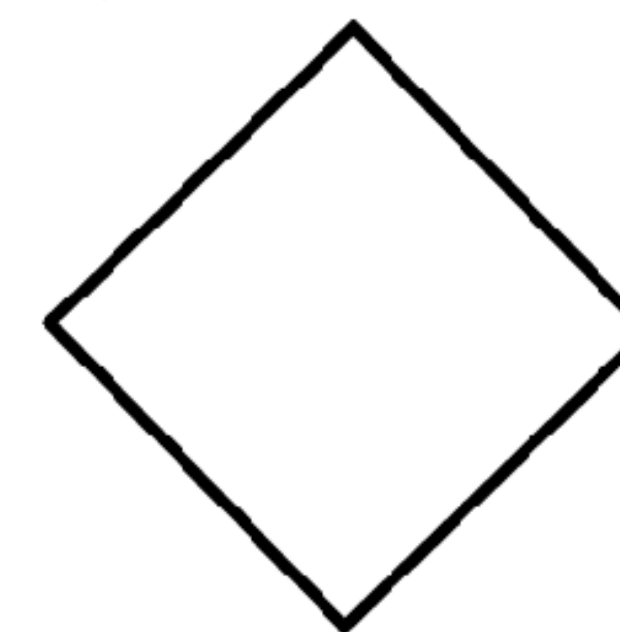
Property	Complexity	$[m]^n$
Uniformity	$\tilde{\Theta}(\sqrt{n}/\epsilon^2)$	$\tilde{O}(\text{poly}(m)\sqrt{n}/\epsilon^2)$ [CM'23]
Junta Testing	$\tilde{\Theta}((k + \sqrt{n})/\epsilon^2)$	
Junta Learning	$\tilde{\Theta}(k/\epsilon^2) \cdot \log n$	
Monotonicity	$\tilde{\Theta}(n/\epsilon^2)$	



Restricted Inputs

Product
distributions

subcube conditioning
=
samples



$$\{-1, 1\}^n$$

Subcube Conditioning

Product Distributions are much simpler.

$\mathcal{D}_1$

1	0	0	0	0	0	1	1	0	1	1	0	1	1	1
0	0	1	0	1	0	0	0	1	1	1	1	1	1	0
0	1	1	0	0	1	0	1	0	1	0	0	1	1	1
0	0	1	1	0	0	0	1	1	0	1	0	1	1	0
0	0	0	0	0	0	0	0	1	1	1	1	1	0	1
0	1	0	0	0	0	1	1	1	0	1	0	1	1	1
0	0	0	0	1	0	0	0	0	1	1	1	0	1	1
0	0	0	0	0	0	0	1	0	1	1	1	1	1	1
0	0	0	0	0	0	0	1	1	1	0	1	1	1	0

$\mathcal{D}_2$

0	1	0	1	0	1	1	0	0	1	0	1	1	0	0
1	0	1	1	0	1	1	1	1	0	0	1	1	1	1
0	0	0	0	1	1	0	1	1	0	0	1	0	1	0
1	1	1	0	0	1	1	1	0	0	0	0	0	1	1
1	1	0	1	1	1	0	1	1	1	1	0	1	1	1
0	0	0	0	1	1	0	1	1	0	0	1	0	0	0
1	0	1	1	0	1	1	1	1	1	1	1	0	0	1
1	1	1	0	1	0	0	1	0	1	0	0	0	1	0

Which is uniform on $\{0, 1\}^{15}$?

Product Distributions are much simpler.

$$\mathcal{D}_1$$

1	0	0	0	0	0	1	1	0	1	1	0	1	1	1
0	0	1	0	1	0	0	0	1	1	1	1	1	1	0
0	1	1	0	0	1	0	1	0	1	0	0	1	1	1
0	0	1	1	0	0	0	1	1	0	1	0	1	1	0
0	0	0	0	0	0	0	0	1	1	1	1	1	0	1
0	1	0	0	0	0	1	1	1	0	1	0	1	1	1
0	0	0	0	1	0	0	0	0	1	1	1	0	1	1
0	0	0	0	0	0	0	1	0	1	1	1	1	1	1
0	0	0	0	0	0	0	1	1	1	0	1	1	1	0

$$\Pr[1] = 0.2$$
$$\Pr[1] = 0.8$$
$$\mathcal{D}_2$$

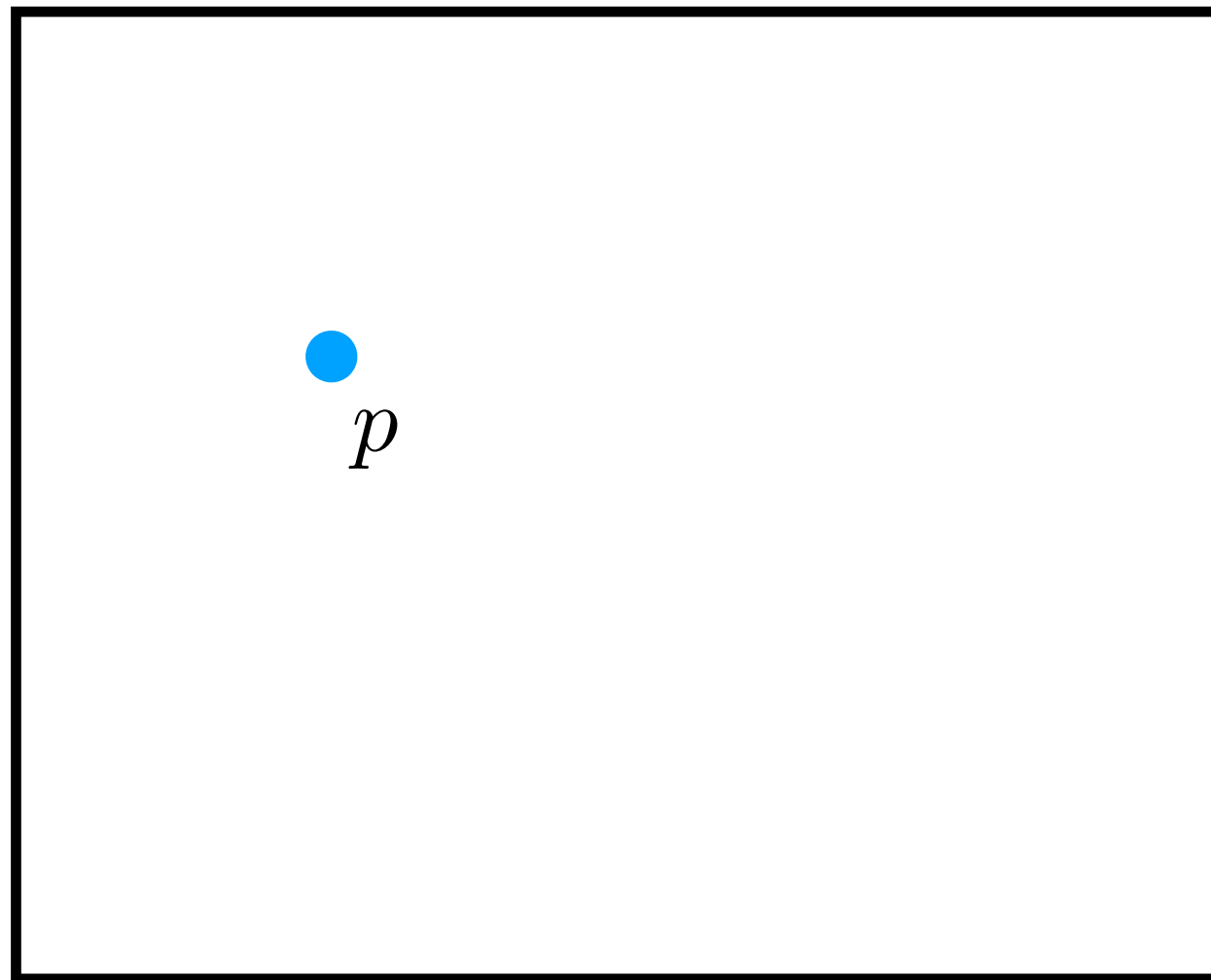
0	1	0	1	0	1	1	0	0	1	0	1	1	0	0
1	0	1	1	0	1	1	1	1	0	0	1	1	1	1
0	0	0	0	1	1	0	1	1	0	0	1	0	1	0
1	1	1	0	0	1	1	1	0	0	0	0	0	1	1
1	1	0	1	1	1	0	1	1	1	1	0	1	1	1
0	0	0	0	1	1	0	1	1	0	0	1	0	0	0
1	0	1	1	0	1	1	1	1	1	1	1	0	0	1
1	1	1	0	1	0	0	1	0	1	0	0	0	1	0

Which is uniform on $\{0, 1\}^{15}$?

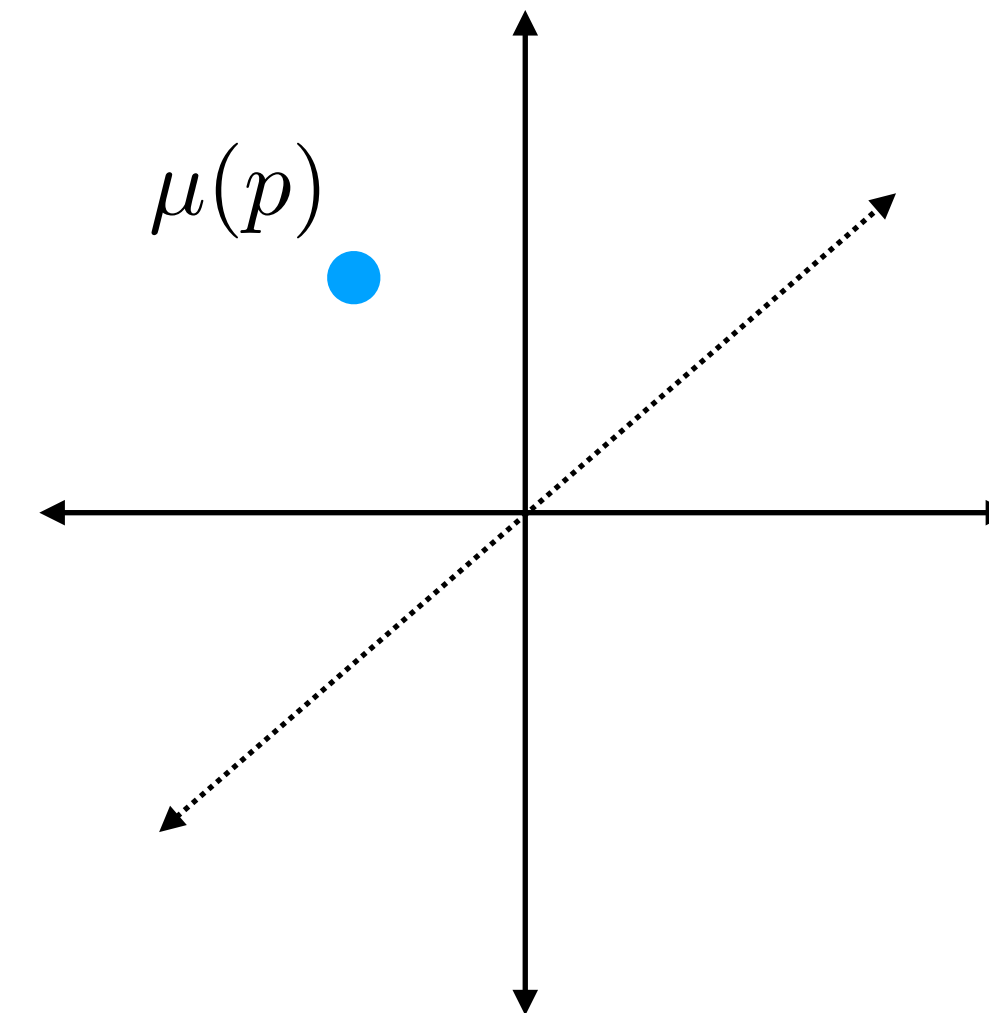
Product Distributions on $\{-1, 1\}^n$

$$\mu(p) = (\mu(p)_1, \mu(p)_2, \dots, \mu(p)_n) = \mathbb{E}_{\mathbf{x} \sim p} [\mathbf{x}] \in [-1, 1]^n$$

↖
bias of each
coordinate



Space of distributions in $\{-1, 1\}^n$

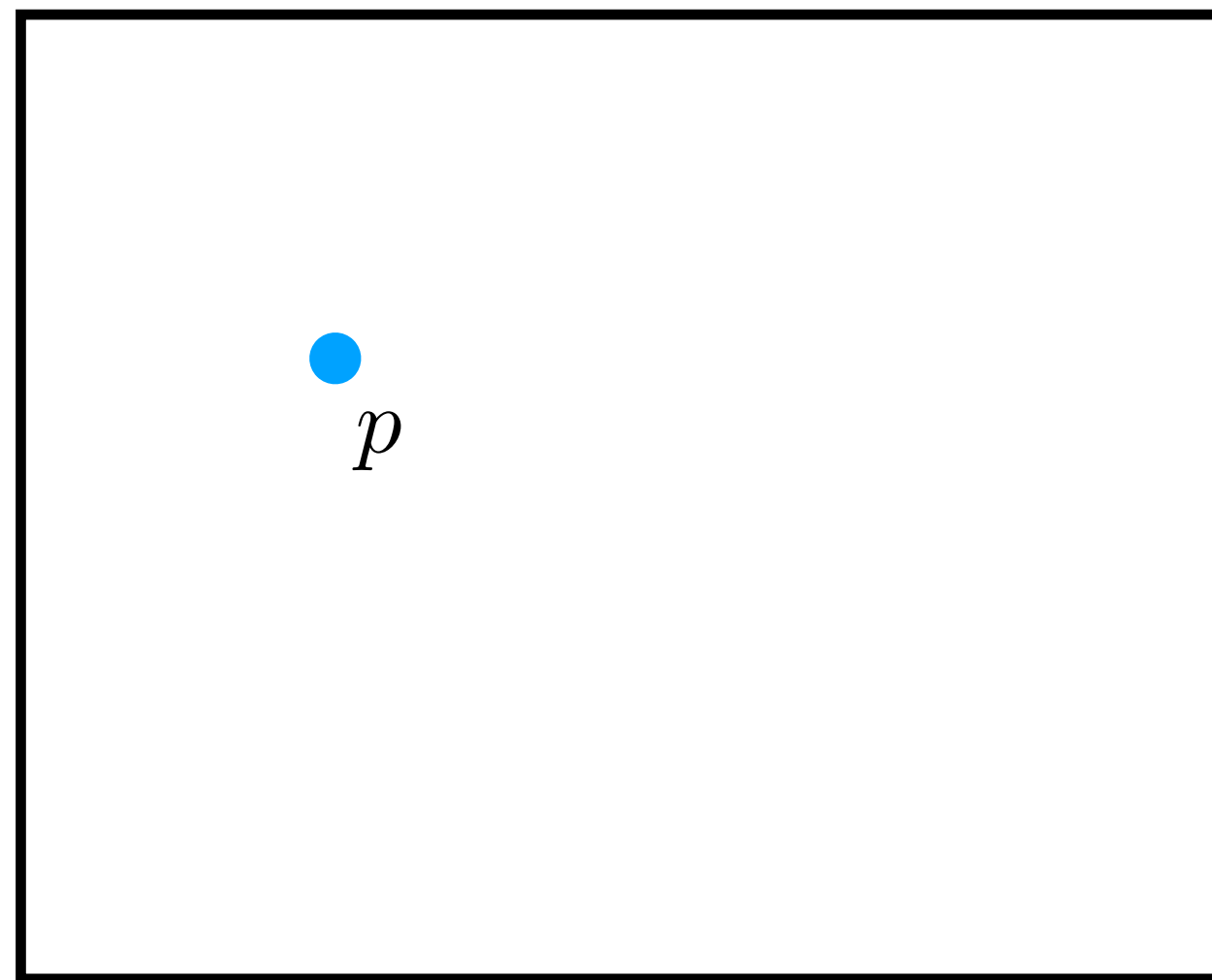


Space of $\mu(p) \in [-1, 1]^n$

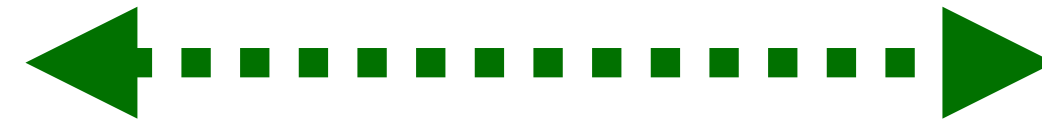
Product Distributions on $\{-1, 1\}^n$

$$\mu(p) = (\mu(p)_1, \mu(p)_2, \dots, \mu(p)_n) = \mathbb{E}_{\mathbf{x} \sim p} [\mathbf{x}] \in [-1, 1]^n$$

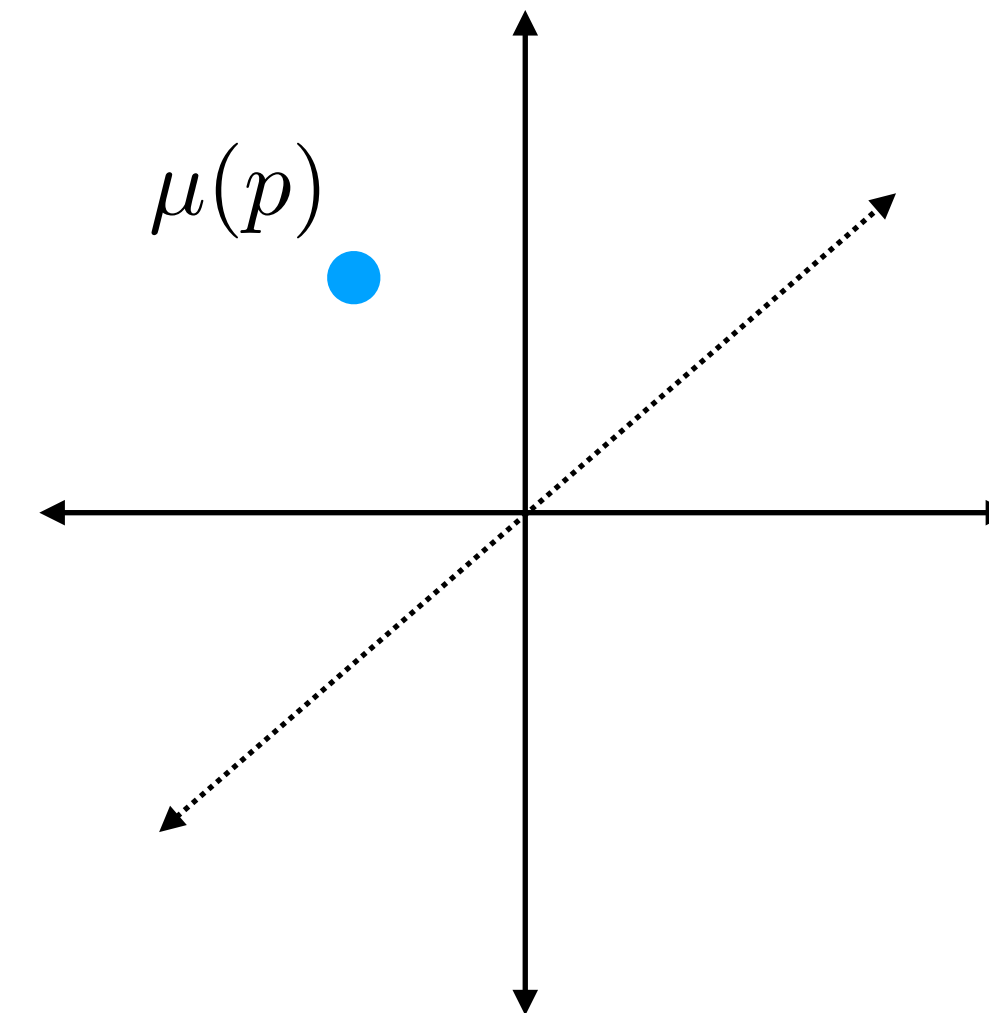
↖
bias of each
coordinate



Space of distributions in $\{-1, 1\}^n$



Huge reduction
in complexity!

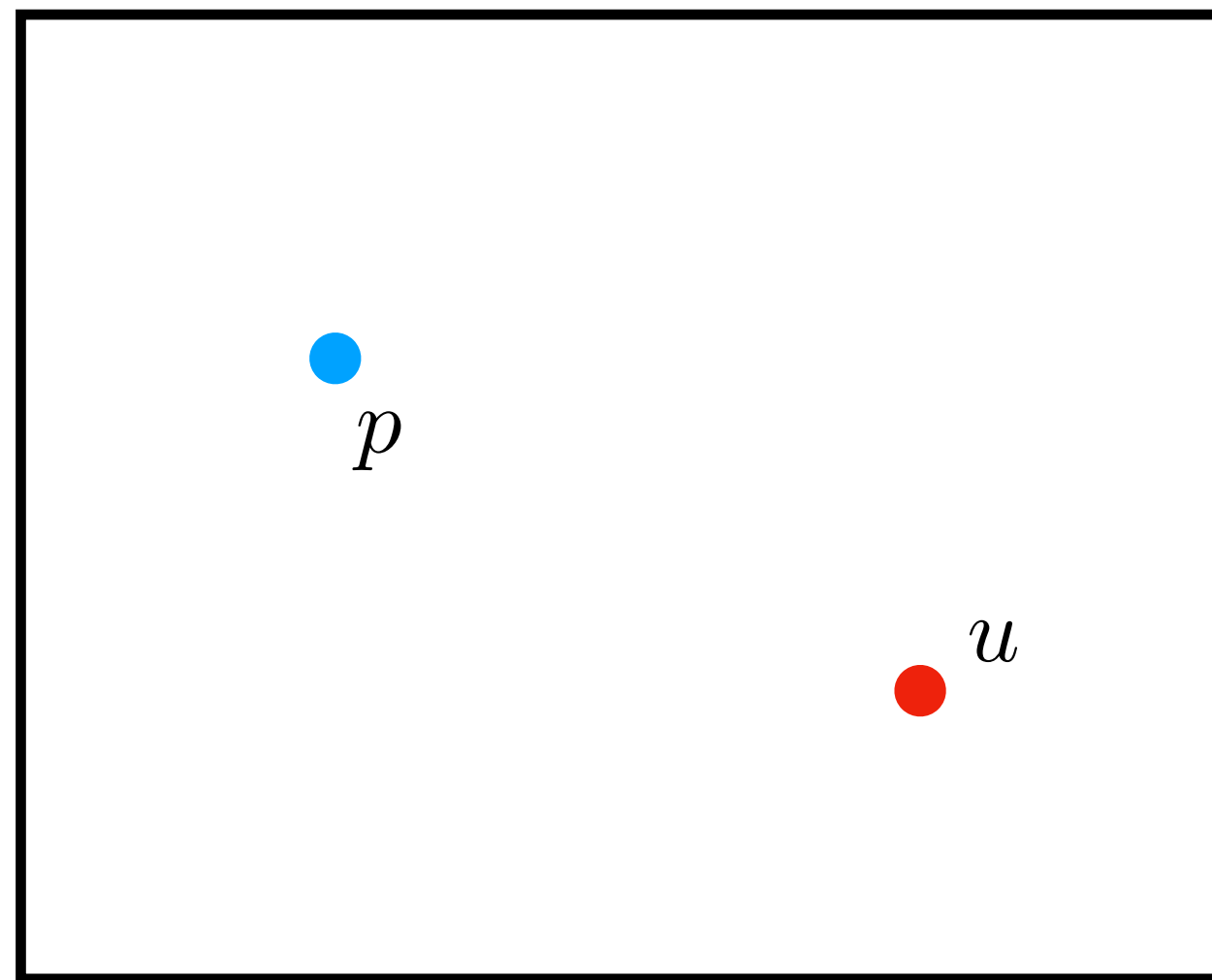


Space of $\mu(p) \in [-1, 1]^n$

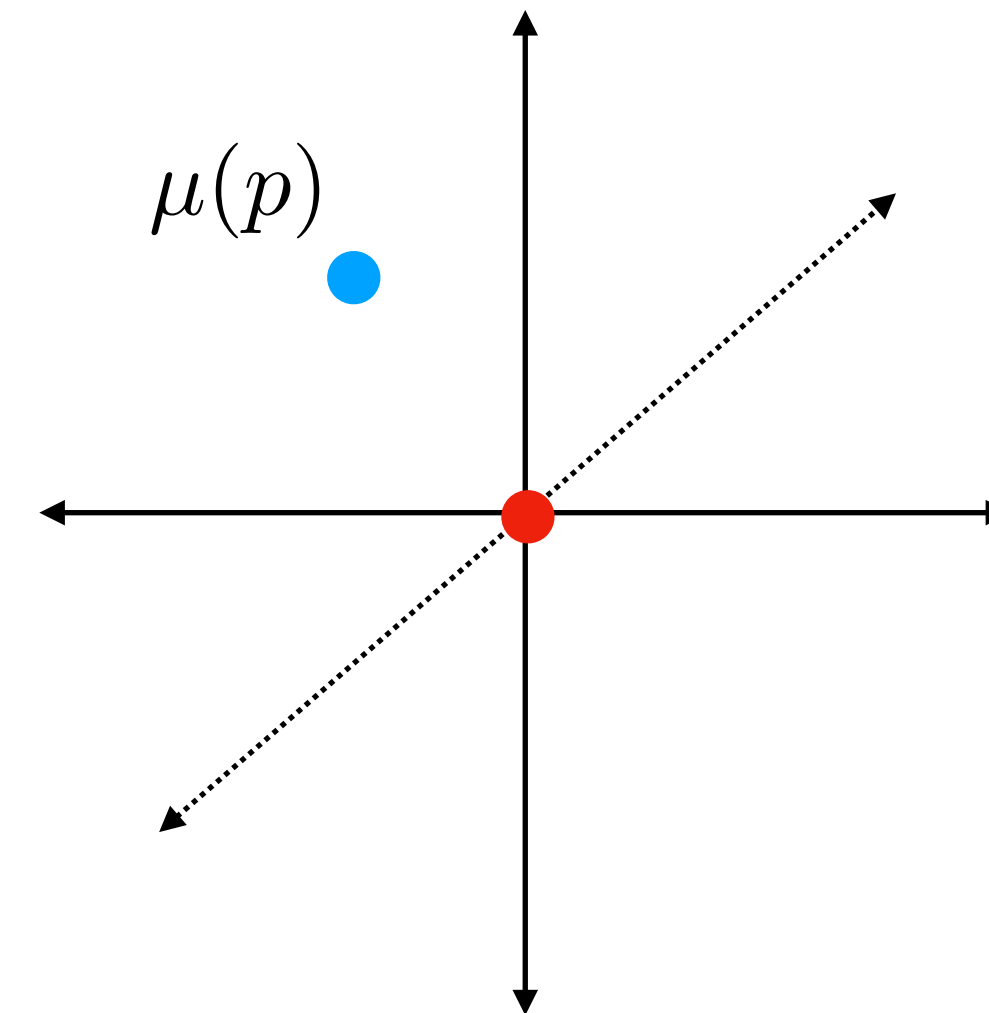
Product Distributions on $\{-1, 1\}^n$

$$\mu(p) = (\mu(p)_1, \mu(p)_2, \dots, \mu(p)_n) = \mathbb{E}_{\mathbf{x} \sim p} [\mathbf{x}] \in [-1, 1]^n$$

bias of each
coordinate



Space of distributions in $\{-1, 1\}^n$

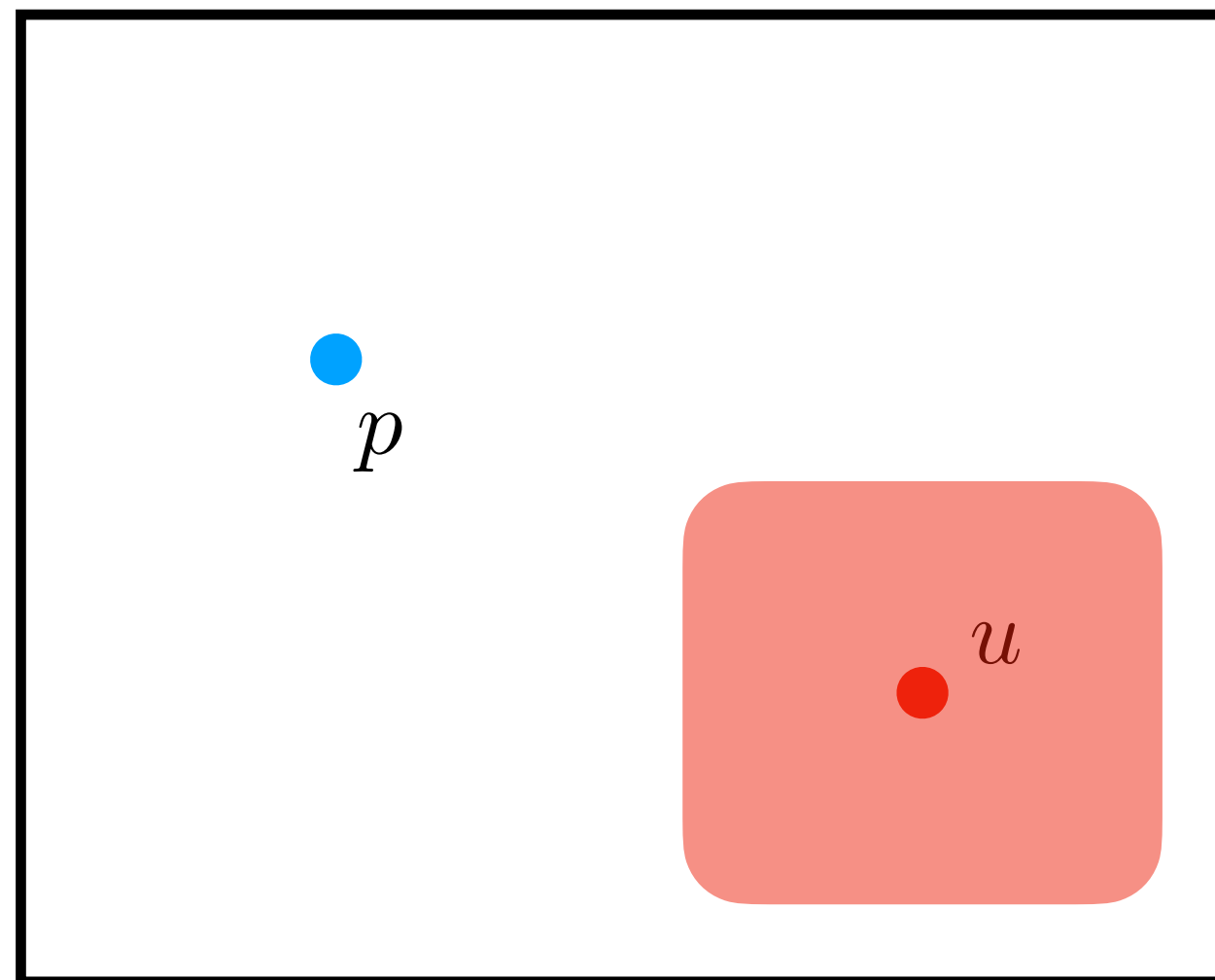


Space of $\mu(p) \in [-1, 1]^n$

Product Distributions on $\{-1, 1\}^n$

$$\mu(p) = (\mu(p)_1, \mu(p)_2, \dots, \mu(p)_n) = \mathbb{E}_{\mathbf{x} \sim p} [\mathbf{x}] \in [-1, 1]^n$$

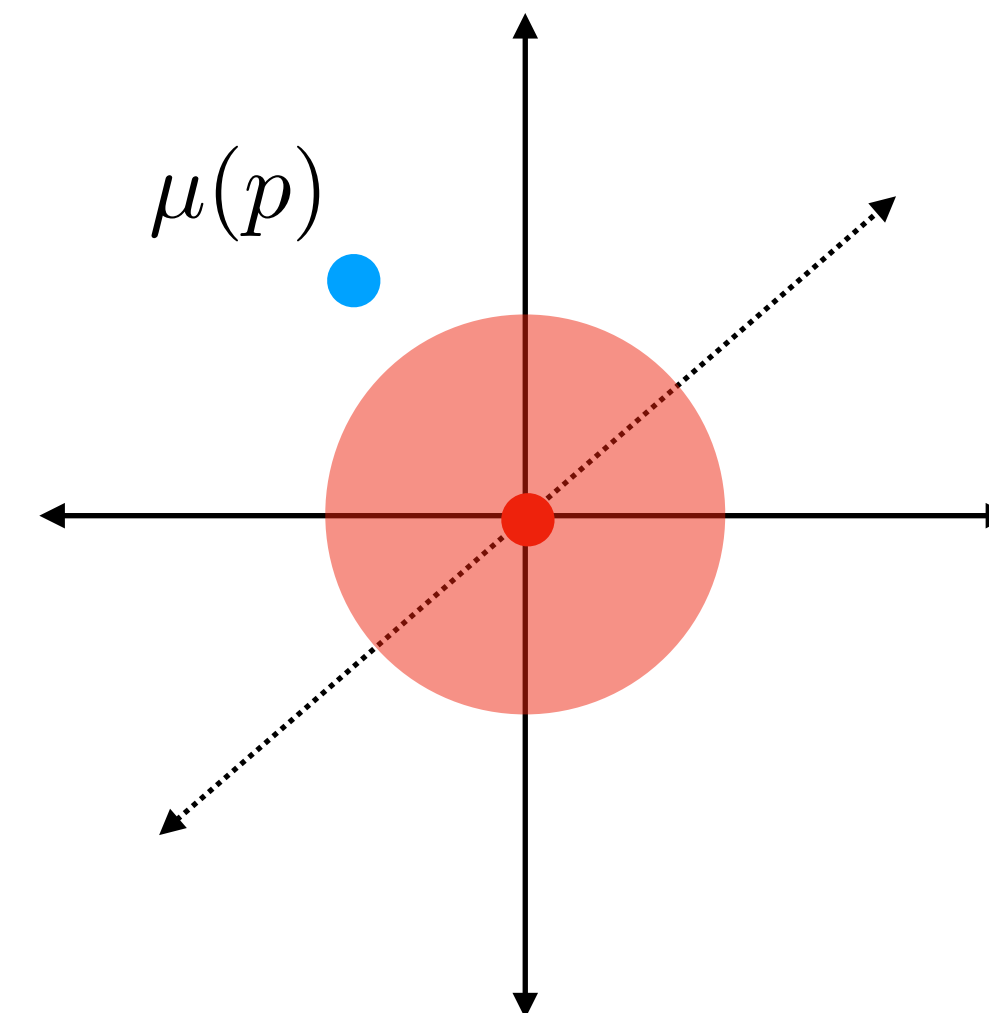
↖
bias of each
coordinate



Space of distributions in $\{-1, 1\}^n$

Simple Claim:

$$\|\mu(p)\|_2 \gtrsim d_{\text{TV}}(p, u)$$



Space of $\mu(p) \in [-1, 1]^n$

Product Distributions on $\{-1, 1\}^n$

$$\mu(p) = (\mu(p)_1, \mu(p)_2, \dots, \mu(p)_n) = \mathbb{E}_{\mathbf{x} \sim p} [\mathbf{x}] \in [-1, 1]^n$$

↖
bias of each
coordinate

Theorem: There is an algorithm using $O(\sqrt{n}/\epsilon^2)$ samples from a product dist. which can test whether it is uniform or far-from uniform, and is tight.

- It implies $\Omega(\sqrt{n}/\epsilon^2)$ subcube queries are necessary.

Simple Claim:

$$\|\mu(p)\|_2 \gtrsim d_{\text{TV}}(p, u)$$

General Distributions on $\{-1, 1\}^n$

$$\mu(p) = (\mu(p)_1, \mu(p)_2, \dots, \mu(p)_n) = \mathbb{E}_{\mathbf{x} \sim p} [\mathbf{x}] \in [-1, 1]^n$$

↖
bias of each
coordinate

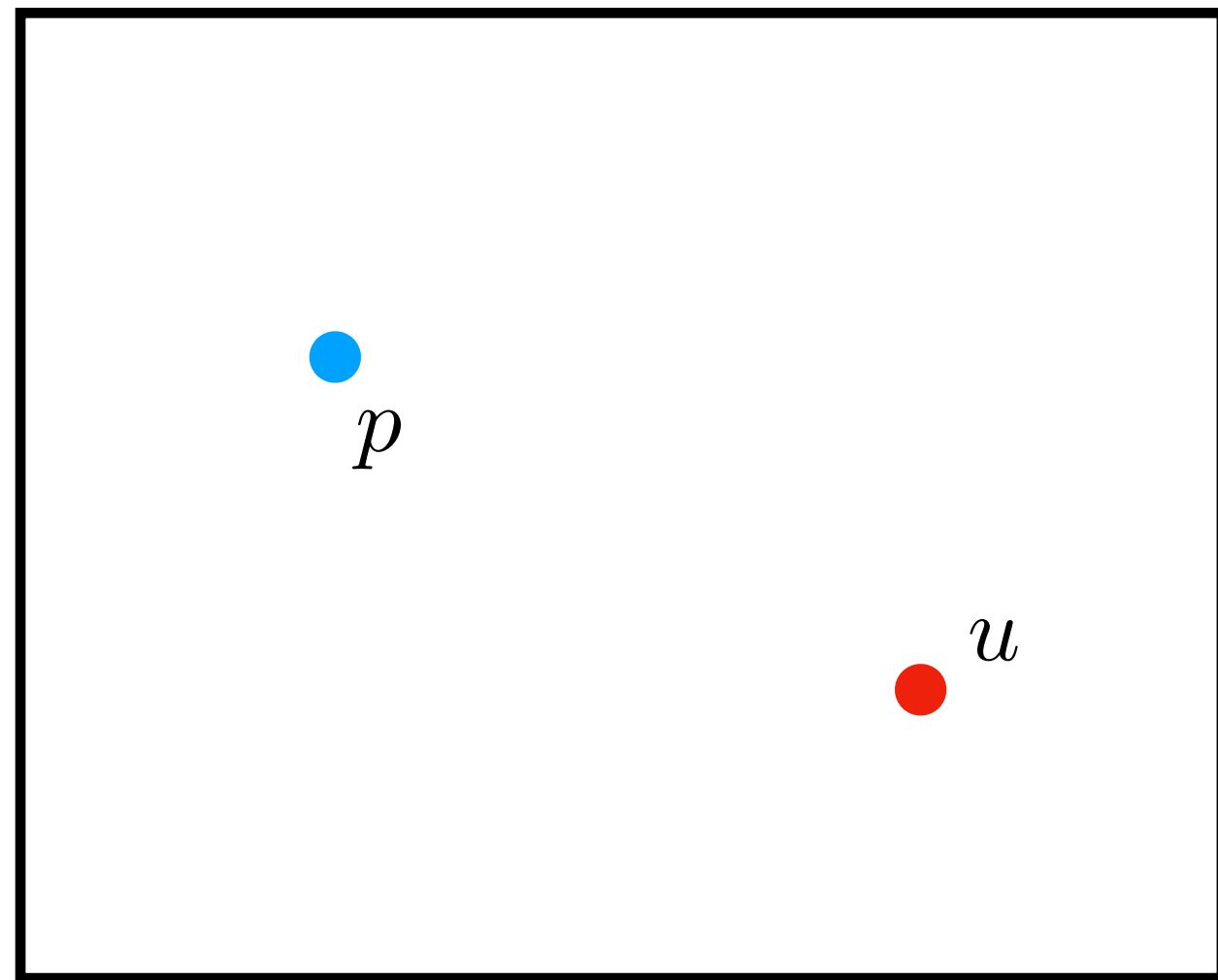
Theorem: There is an algorithm using $\tilde{O}(\sqrt{n}/\epsilon^2)$ ^{subcube queries} ~~samples~~ from a ^{general} ~~product~~ dist. which can test whether it is uniform or far-from uniform, and is tight.

- It implies $\Omega(\sqrt{n}/\epsilon^2)$ subcube queries are necessary.

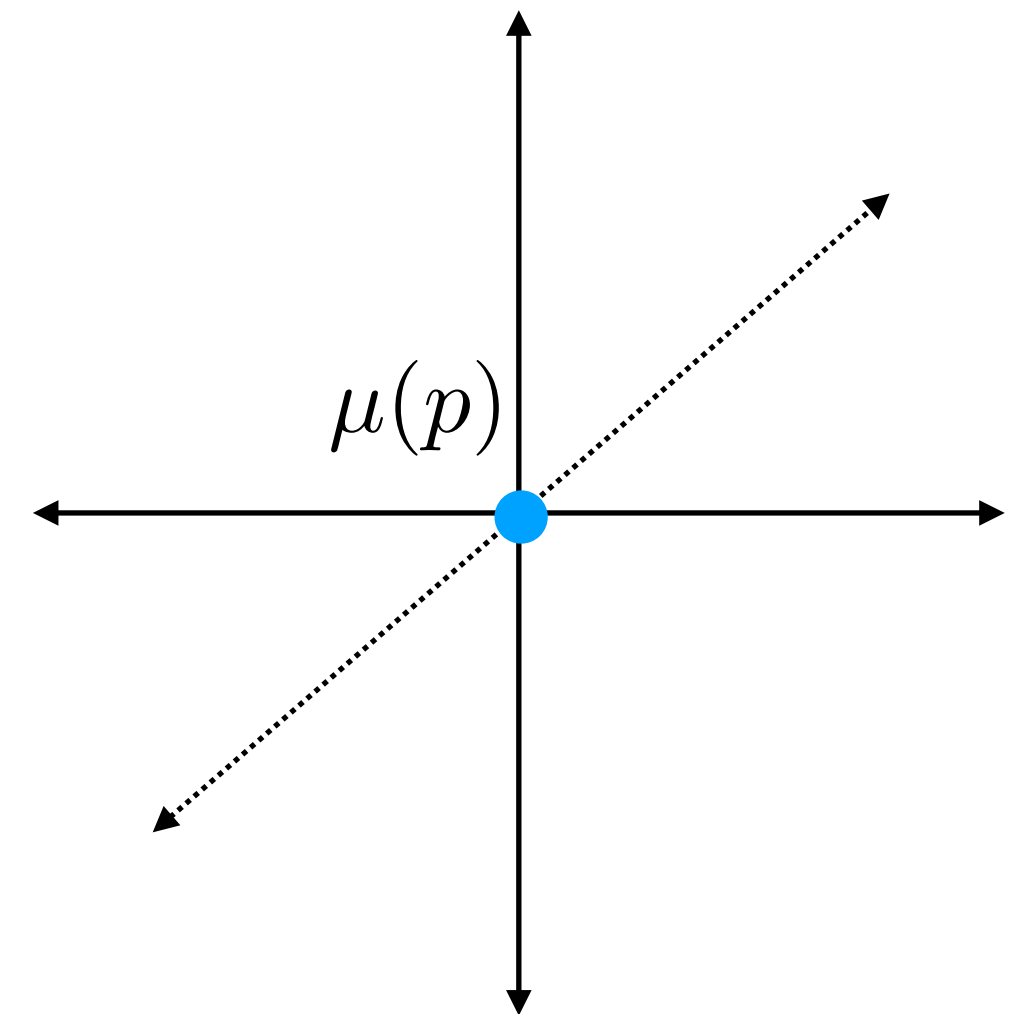
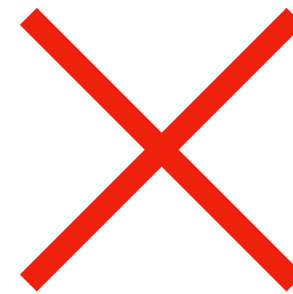
Simple Claim:

$$\|\mu(p)\|_2 \gtrsim d_{\text{TV}}(p, u)$$

General Distributions over $\{-1, 1\}^n$: mean vectors alone are not good



Space of distributions in $\{-1, 1\}^n$

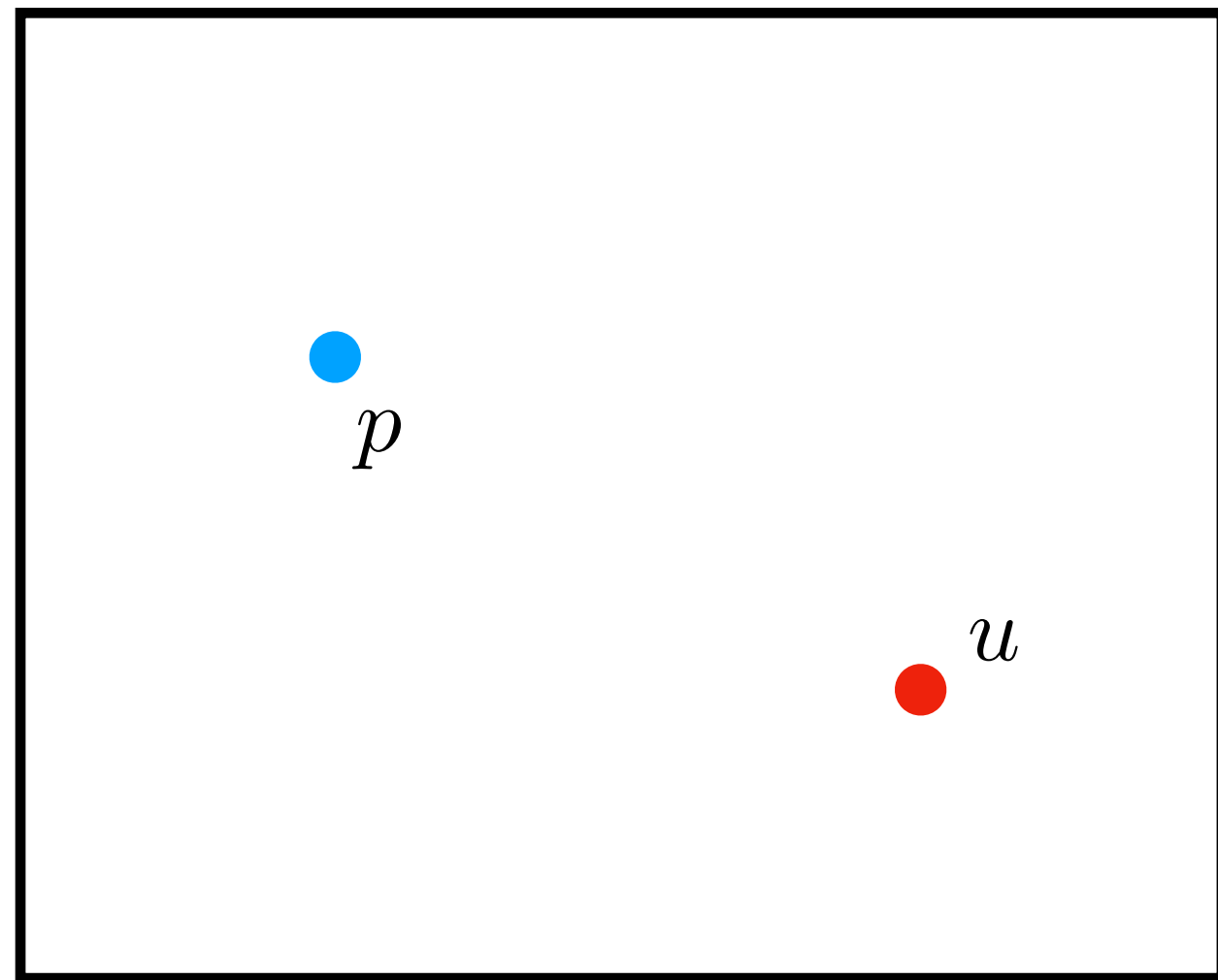


Space of $\mu(p) \in [-1, 1]^n$

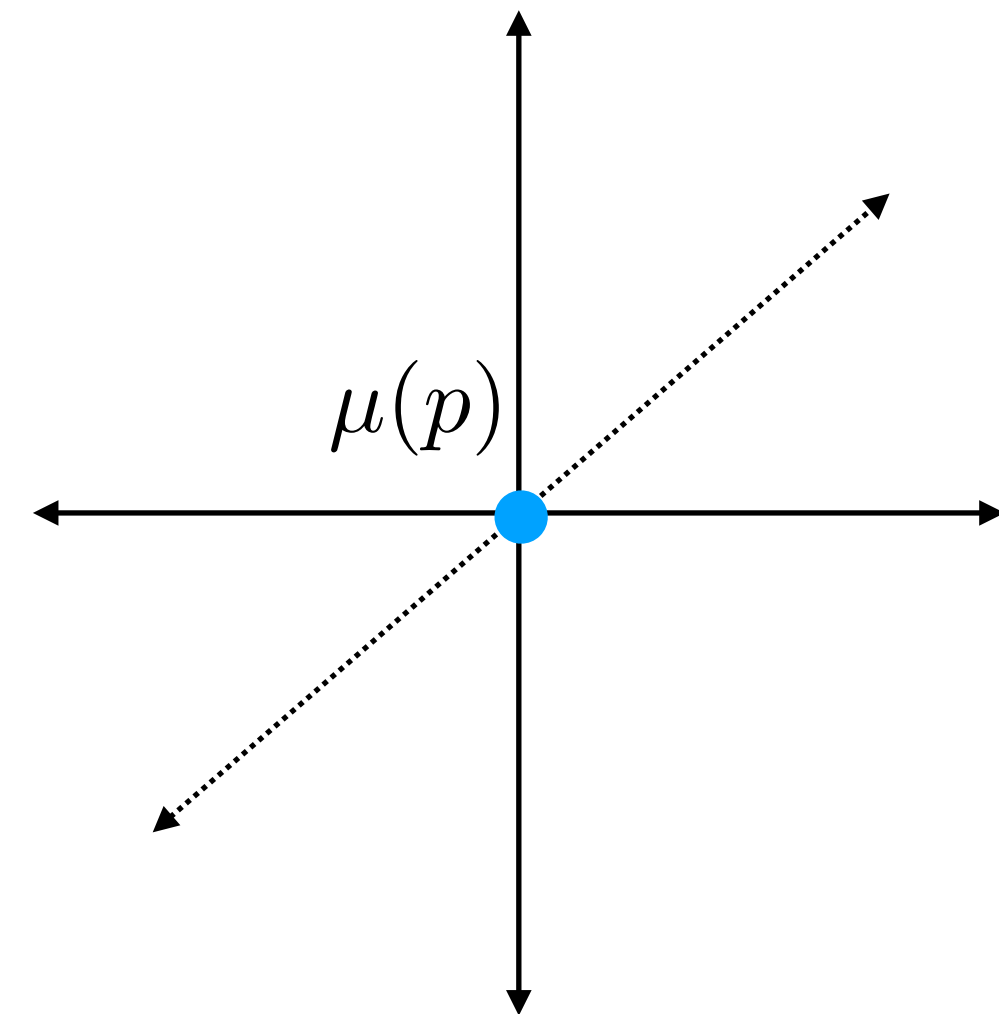
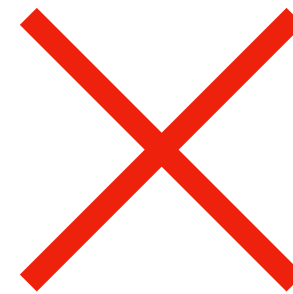
~~Simple Claim:~~

$$\|\mu(p)\|_2 \gtrsim d_{\text{TV}}(p, u)$$

General Distributions over $\{-1, 1\}^n$: mean vectors alone are not good



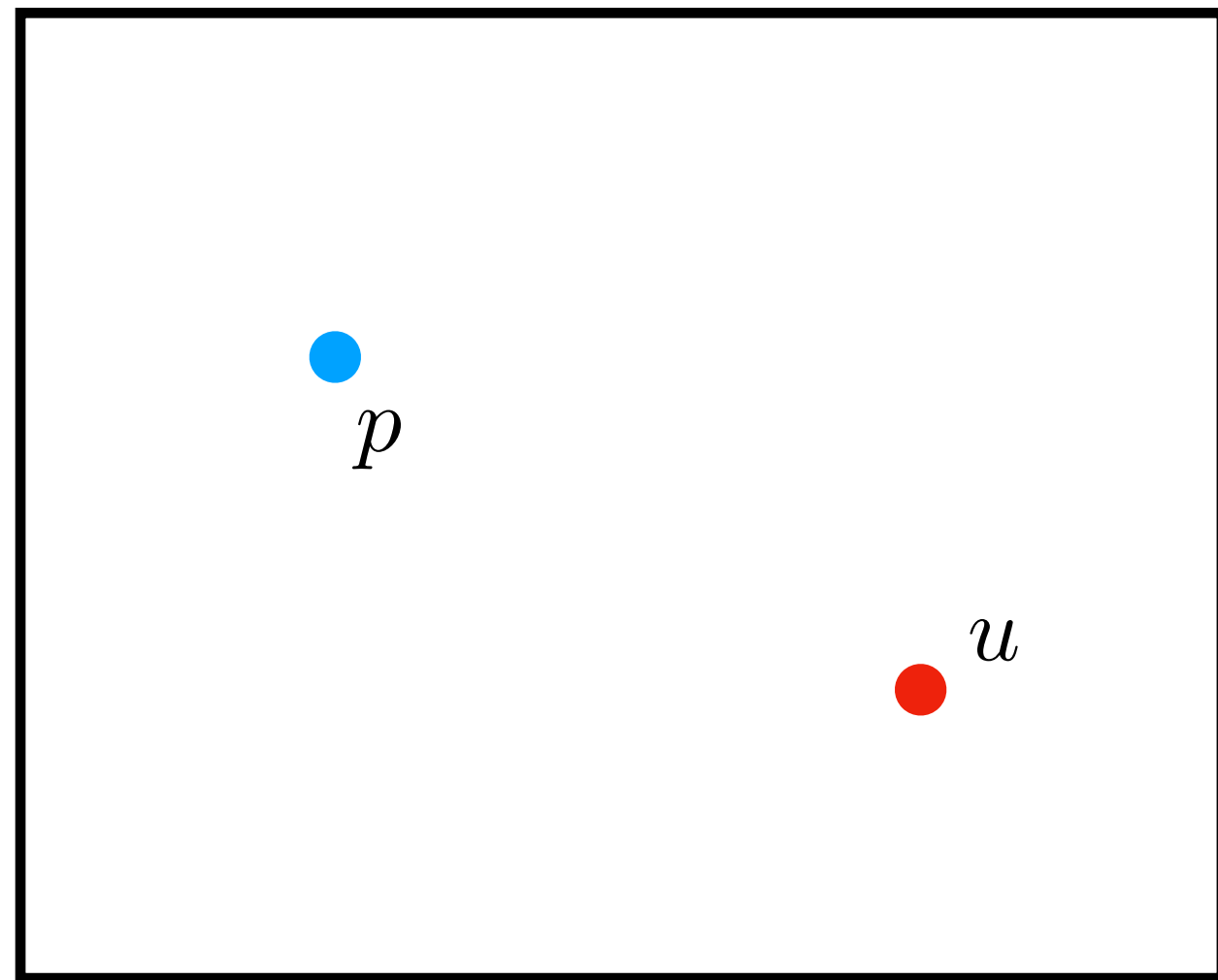
Space of distributions in $\{-1, 1\}^n$



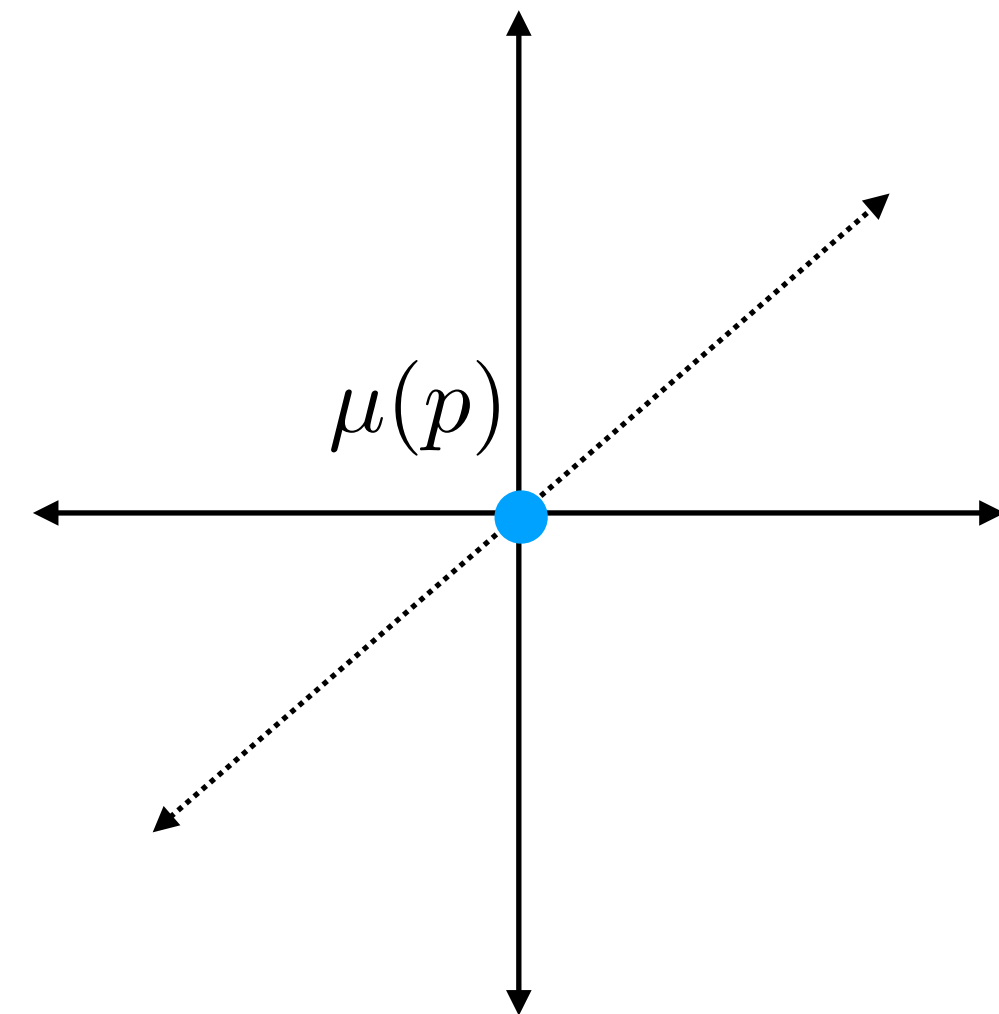
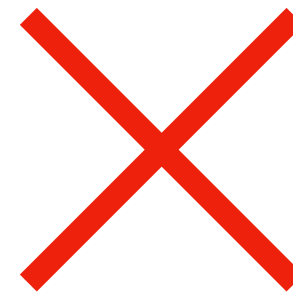
Space of $\mu(p) \in [-1, 1]^n$

$$x \sim p \quad : \quad \begin{cases} (1, 1, \dots, 1) & \text{w.p } 1/2 \\ (-1, -1, \dots, -1) & \text{w.p } 1/2 \end{cases}$$

General Distributions over $\{-1, 1\}^n$: mean vectors alone are not good



Space of distributions in $\{-1, 1\}^n$



Space of $\mu(p) \in [-1, 1]^n$

$$x \sim p \quad : \quad \begin{cases} (1, 1, \dots, 1) & \text{w.p } 1/2 \\ (-1, -1, \dots, -1) & \text{w.p } 1/2 \end{cases}$$

General Distributions over $\{-1, 1\}^n$: mean vectors on random subcubes

Theorem: For any distribution p supported on $\{-1, 1\}^n$,

1. Sample ‘star-scale’ $i \sim [\log_2 n]$, i.e., $\approx 2^i$ stars
2. Subcube $\rho \sim \mathcal{R}_i(p)$,

General Distributions over $\{-1, 1\}^n$: mean vectors on random subcubes

Theorem: For any distribution p supported on $\{-1, 1\}^n$,

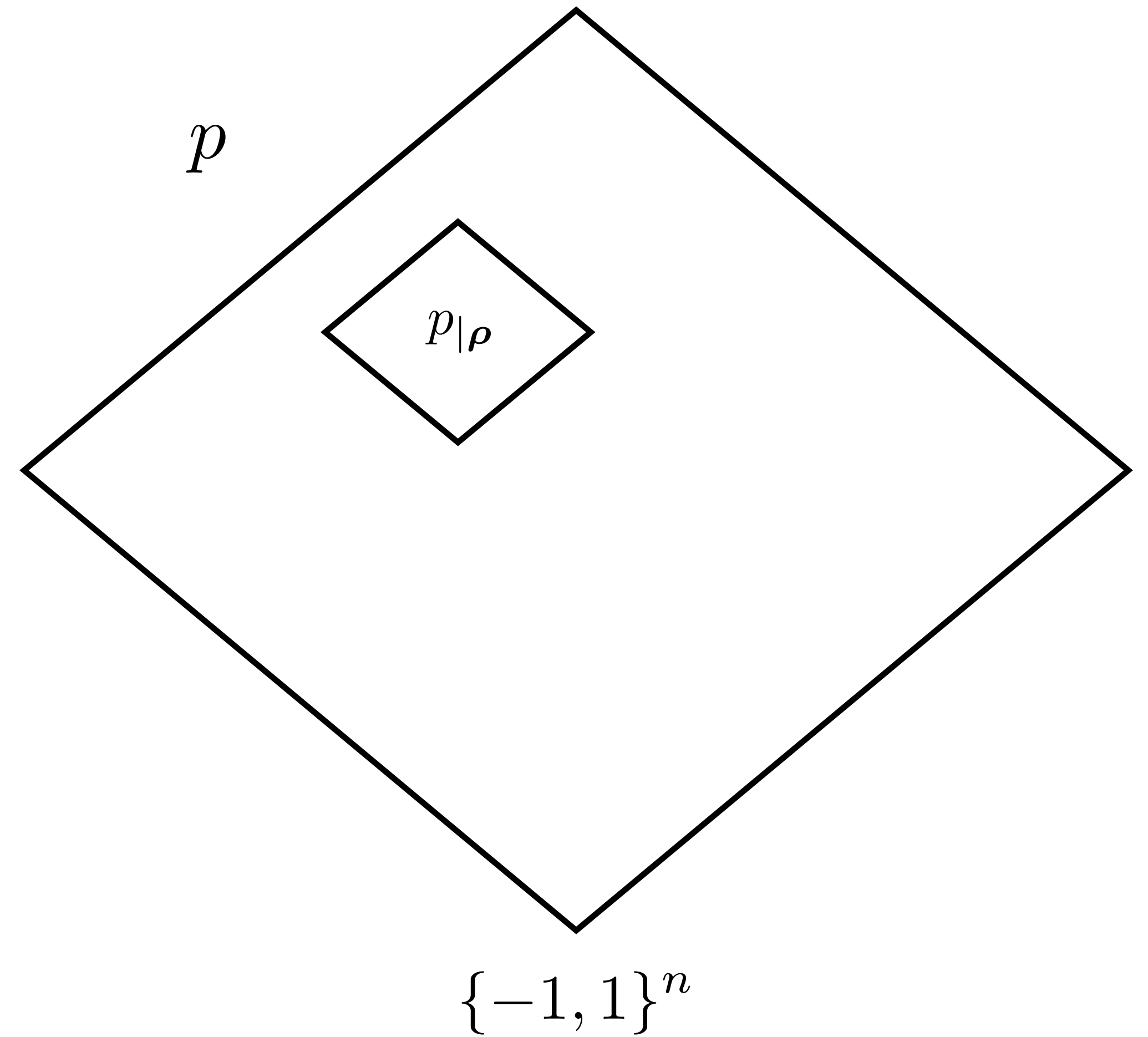
1. Sample ‘star-scale’ $i \sim [\log_2 n]$, i.e., $\approx 2^i$ stars
2. Subcube $\rho \sim \mathcal{R}_i(p)$,

$$\mathbb{E}_{\substack{i \sim [\log n] \\ \rho \sim \mathcal{R}_i(p)}} [\|\mu(p|_{\rho})\|_2] \gtrsim \frac{1}{\text{polylog}(n)} \cdot d_{\text{TV}}(p, u)$$

General Distributions over $\{-1, 1\}^n$: mean vectors on random subcubes

$$\mathbb{E}_{\substack{i \sim [\log n] \\ \rho \sim \mathcal{R}_i(p)}} [\|\mu(p|_{\rho})\|_2] \gtrsim \frac{1}{\text{polylog}(n)} \cdot d_{\text{TV}}(p, u)$$

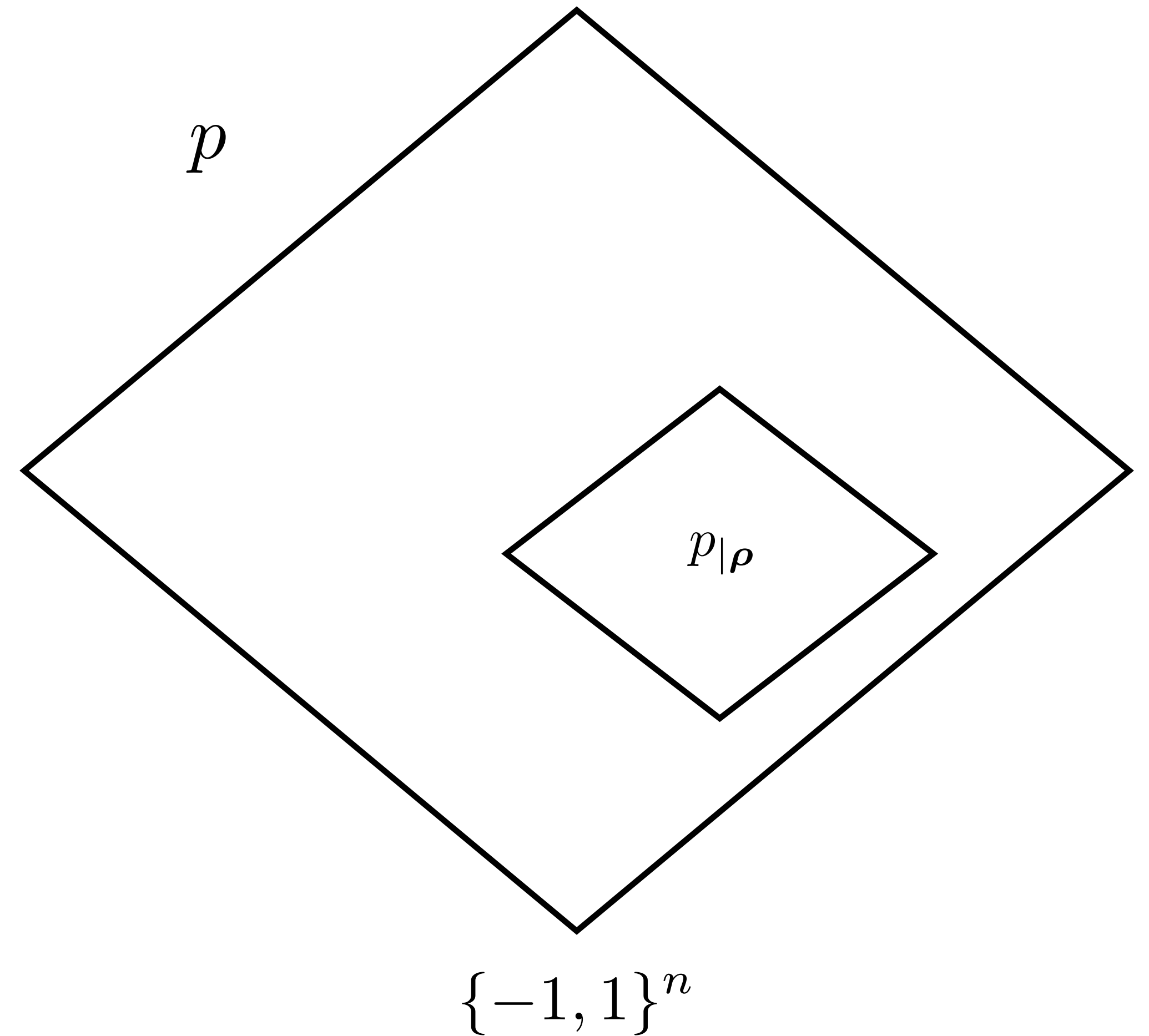
“If all subcubes have unbiased marginals,
close to uniform.”



General Distributions over $\{-1, 1\}^n$: mean vectors on random subcubes

$$\mathbb{E}_{\substack{i \sim [\log n] \\ \rho \sim \mathcal{R}_i(p)}} [\|\mu(p|_{\rho})\|_2] \gtrsim \frac{1}{\text{polylog}(n)} \cdot d_{\text{TV}}(p, u)$$

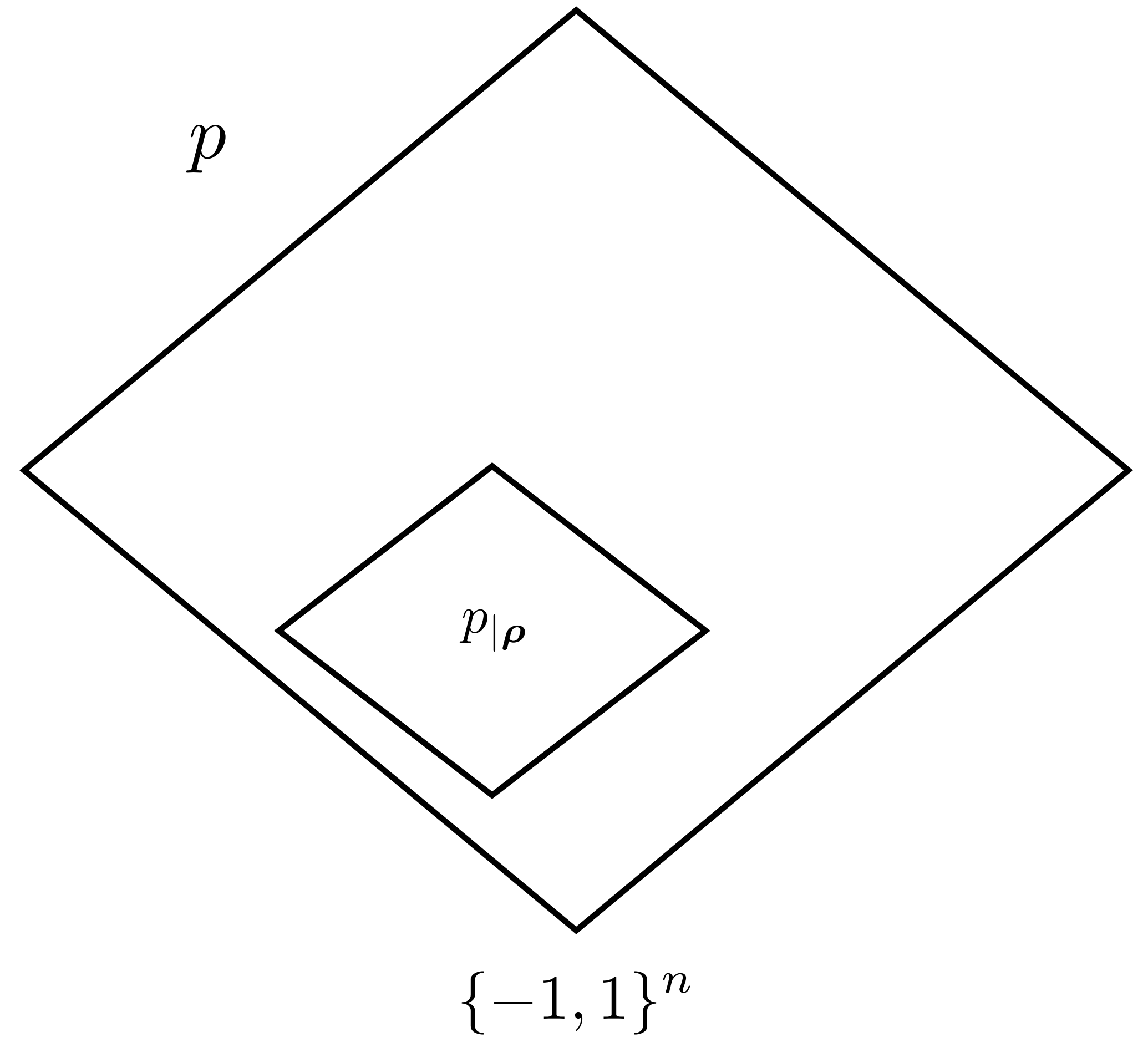
“If all subcubes have unbiased marginals,
close to uniform.”



General Distributions over $\{-1, 1\}^n$: mean vectors on random subcubes

$$\mathbb{E}_{\substack{i \sim [\log n] \\ \rho \sim \mathcal{R}_i(p)}} [\|\mu(p|_{\rho})\|_2] \gtrsim \frac{1}{\text{polylog}(n)} \cdot d_{\text{TV}}(p, u)$$

“If all subcubes have unbiased marginals,
close to uniform.”

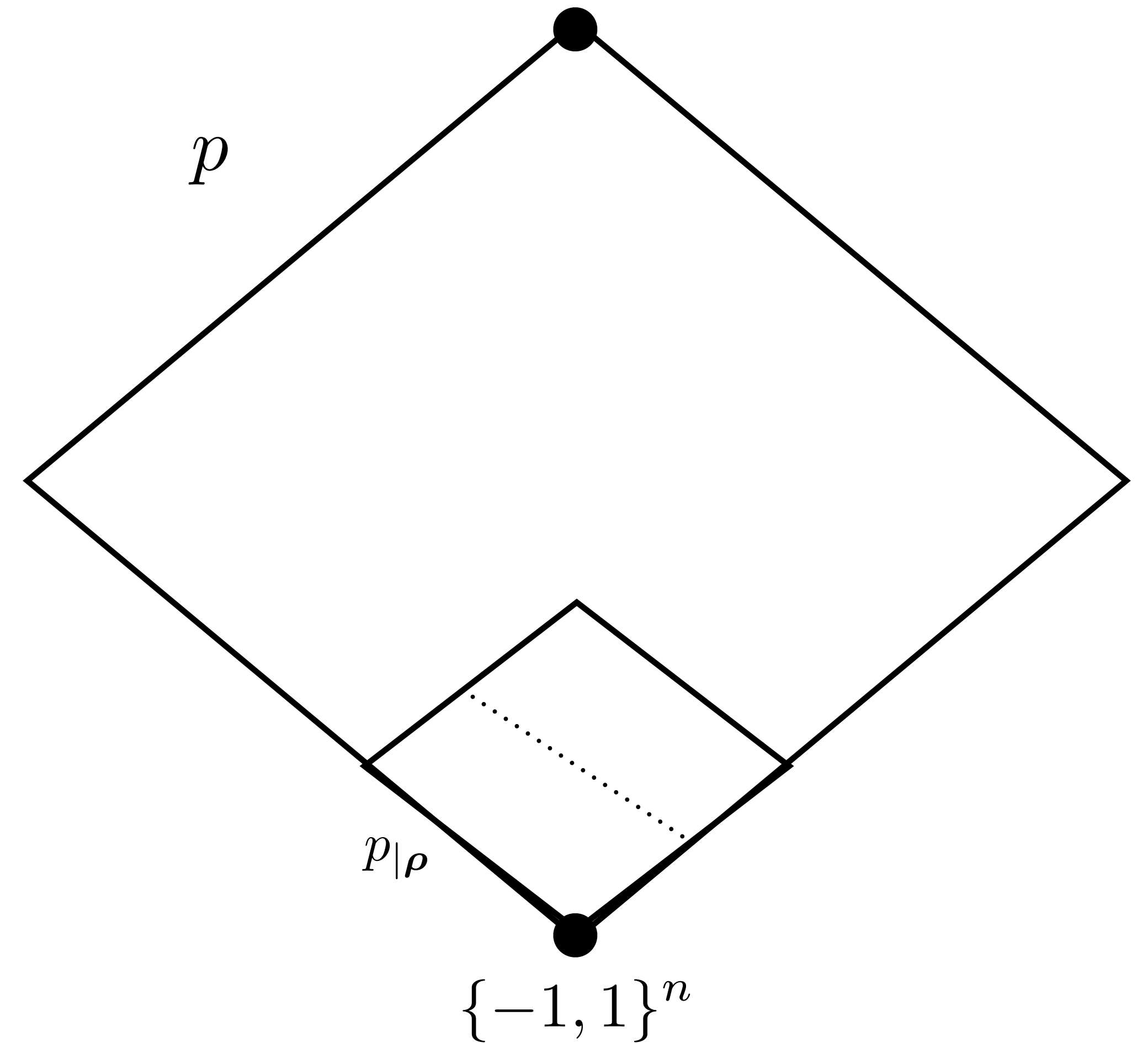


General Distributions over $\{-1, 1\}^n$: mean vectors on random subcubes

Examples:

$$\mathbf{x} \sim p \quad : \quad \begin{cases} (1, 1, \dots, 1) & \text{w.p } 1/2 \\ (-1, -1, \dots, -1) & \text{w.p } 1/2 \end{cases}$$

$$\mathbb{E}_{\substack{i \sim [\log n] \\ \rho \sim \mathcal{R}_i(p)}} [\|\mu(p|_{\rho})\|_2] \approx \sqrt{n}/\log(n) \quad \text{very loose}$$

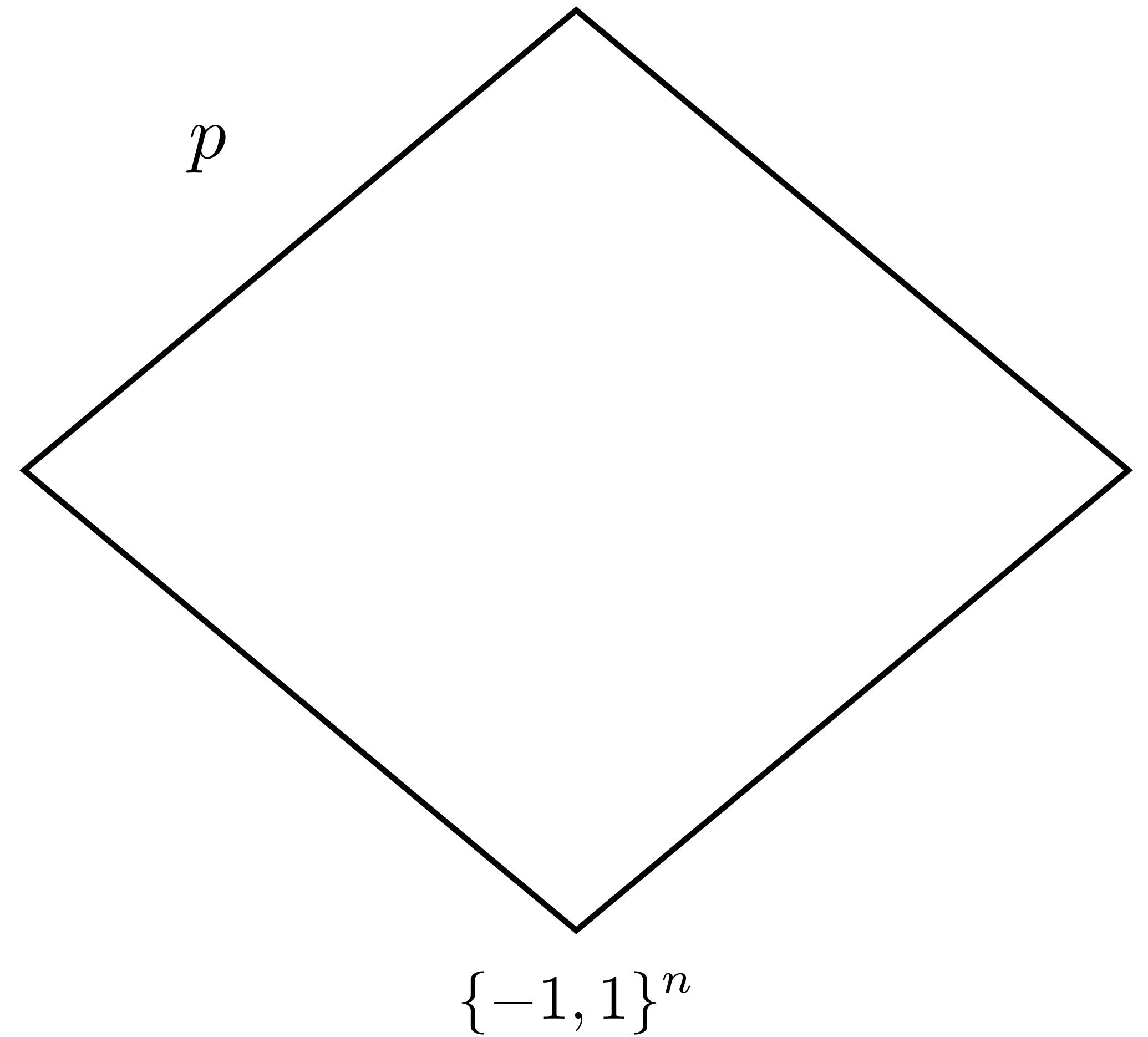


General Distributions over $\{-1, 1\}^n$: mean vectors on random subcubes

Examples:

p : most coords uniform, hidden set

$$S \subset [n] \text{ s.t. } x \sim p, \prod_{j \in S} x_j = 1$$



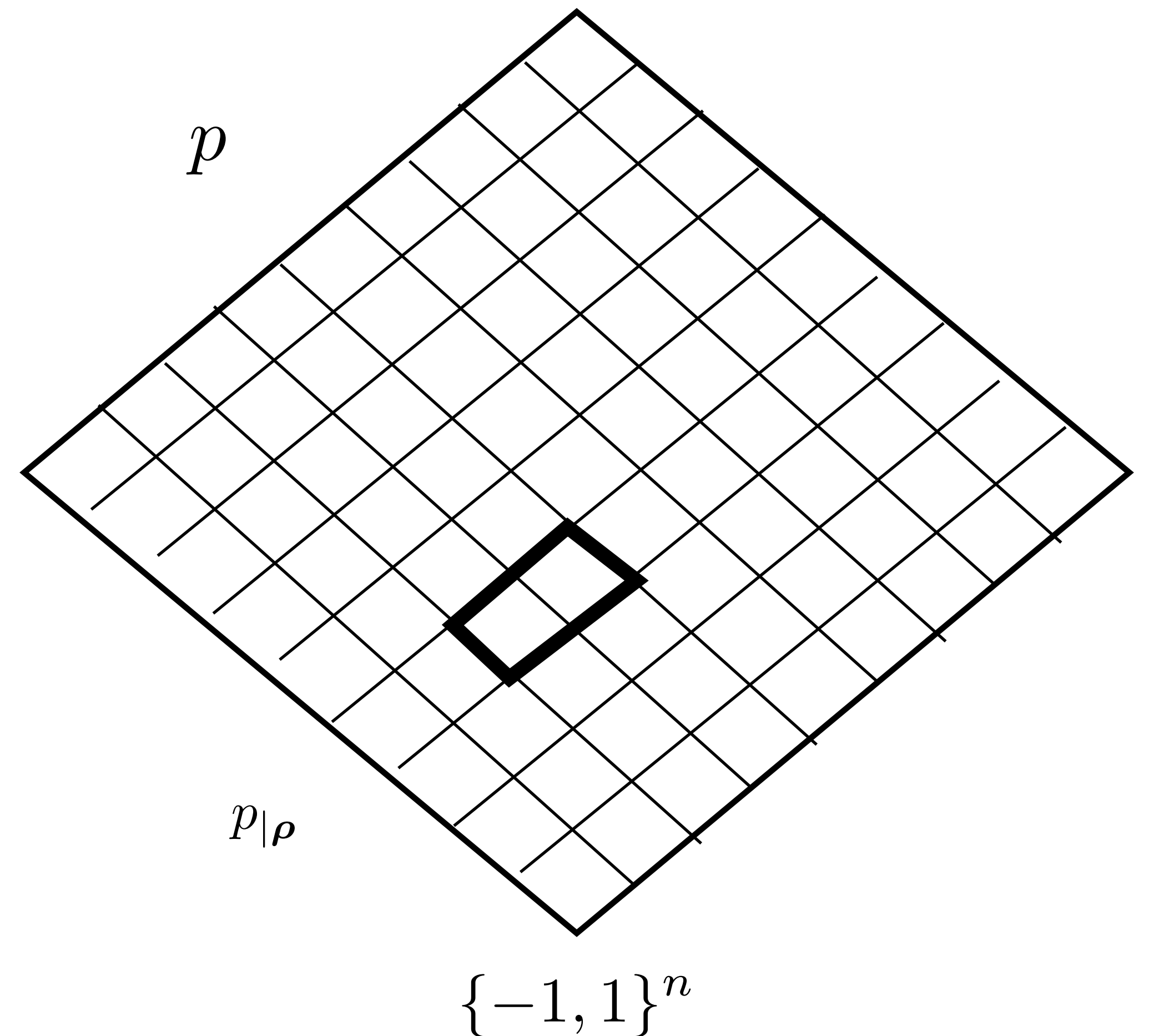
General Distributions over $\{-1, 1\}^n$: mean vectors on random subcubes

Examples:

p : most coords uniform, hidden set

$$S \subset [n] \text{ s.t. } x \sim p, \prod_{j \in S} x_j = 1$$

exists scale $i \in [\log n]$ s.t. $\rho \sim \mathcal{R}_i(p)$
has exactly one 'star' in S .



General Distributions over $\{-1, 1\}^n$: mean vectors on random subcubes

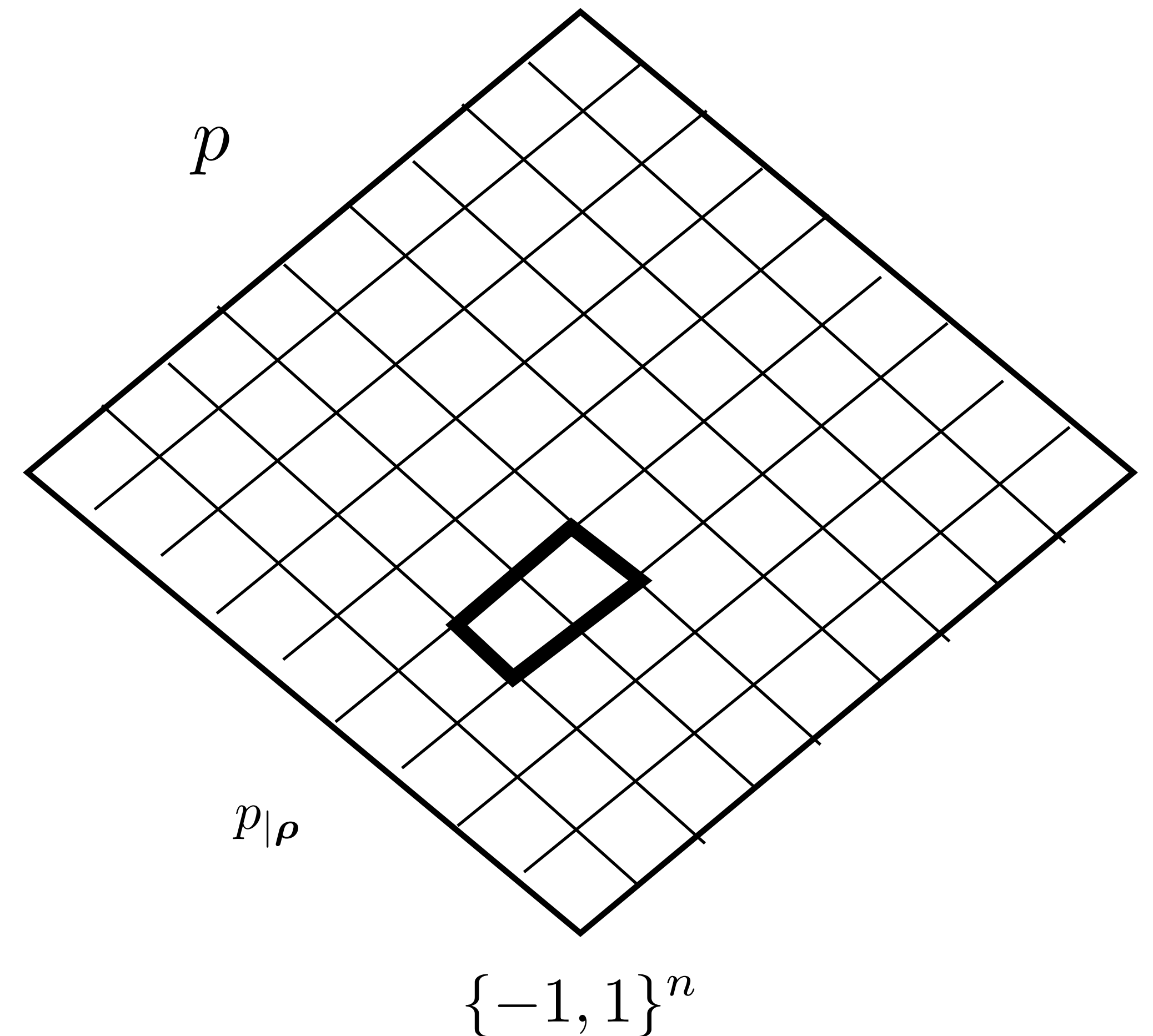
Examples:

p : most coords uniform, hidden set

$$S \subset [n] \text{ s.t. } x \sim p, \prod_{j \in S} x_j = 1$$

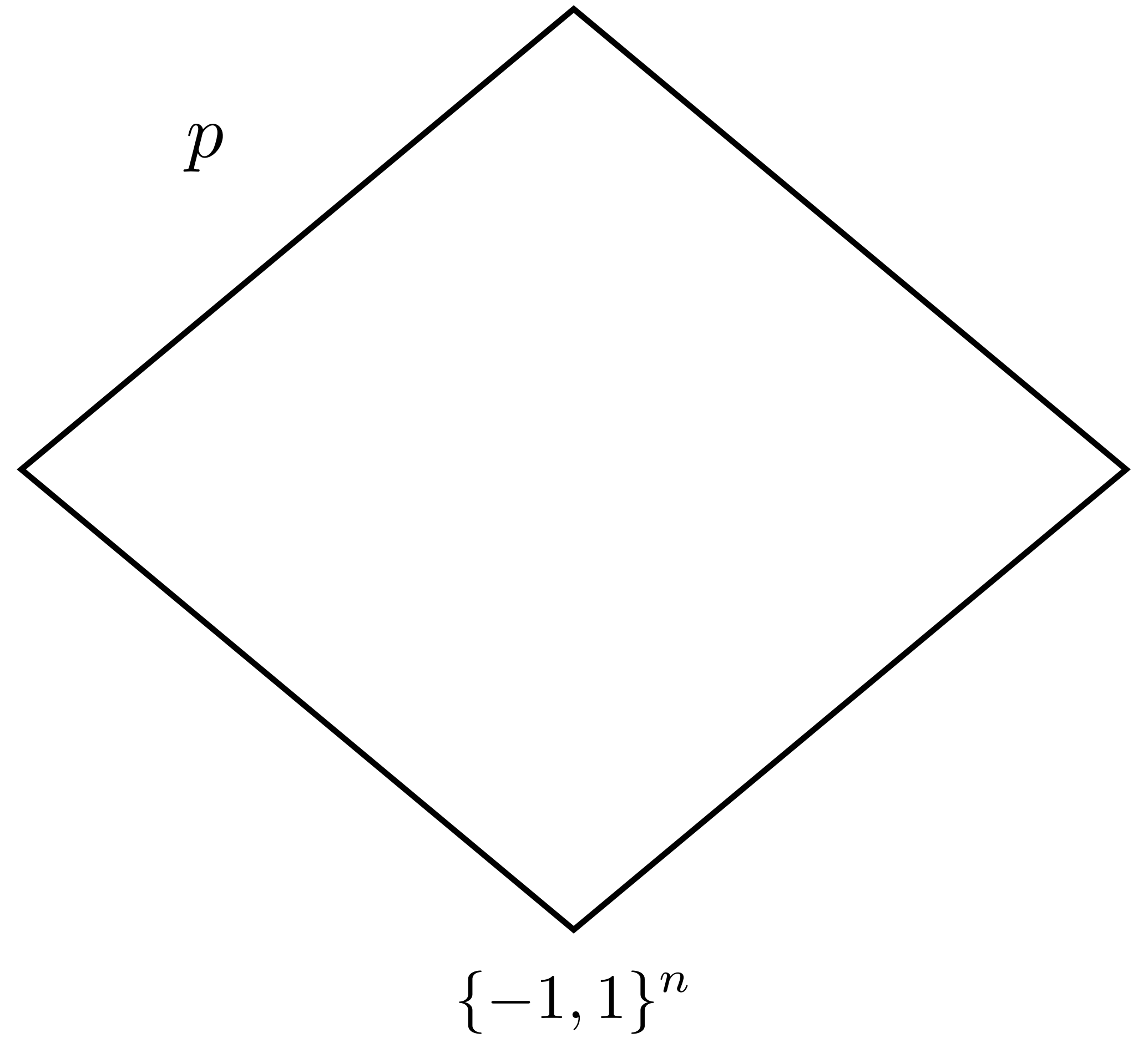
exists scale $i \in [\log n]$ s.t. $\rho \sim \mathcal{R}_i(p)$
has exactly one 'star' in S .

$$\mathbb{E}_{\substack{i \sim [\log n] \\ \rho \sim \mathcal{R}_i(p)}} [\|\mu(p|_{\rho})\|_2] \approx 1/\log(n)$$



Why should something like this be true?

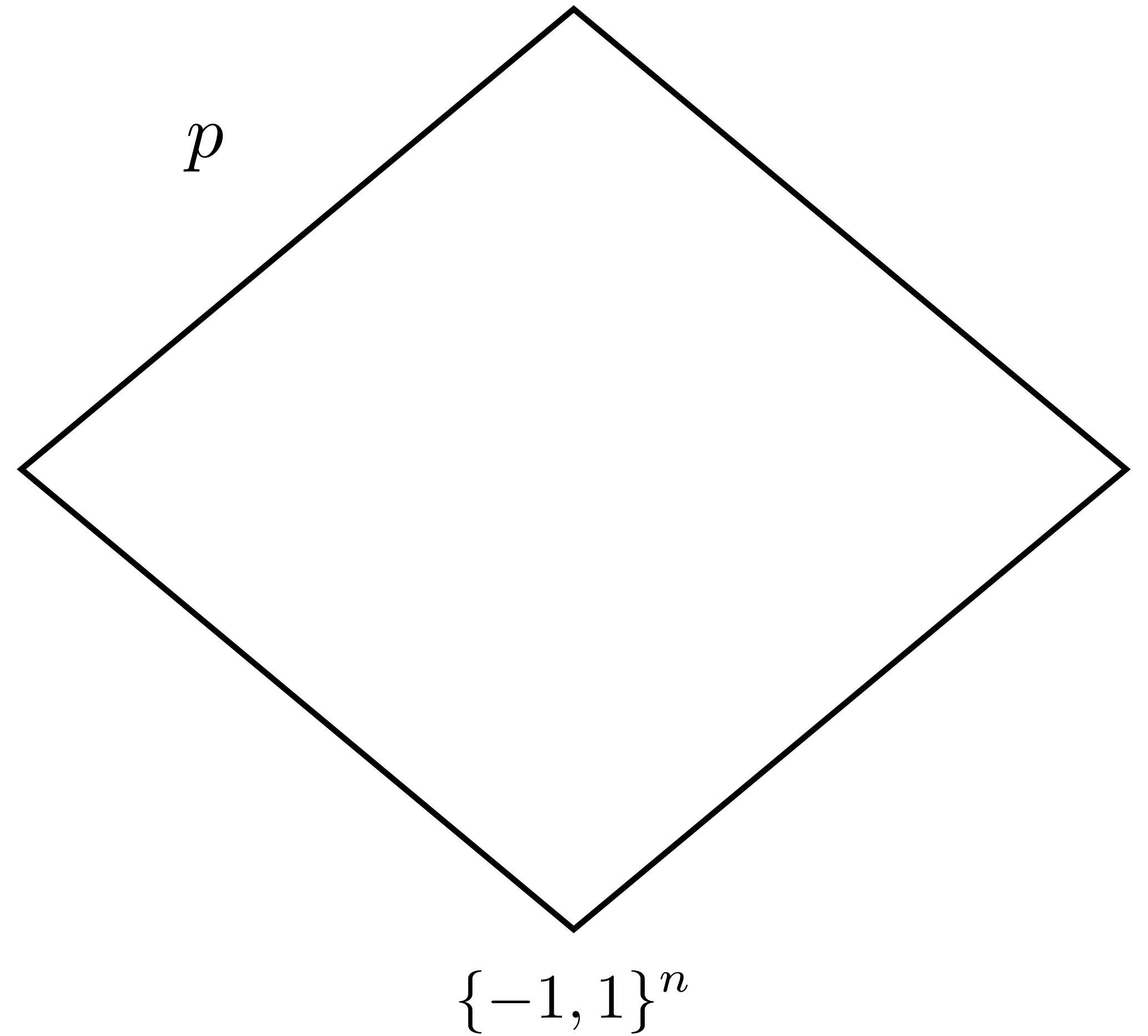
A distribution is uniform iff:



Why should something like this be true?

A distribution is uniform iff:

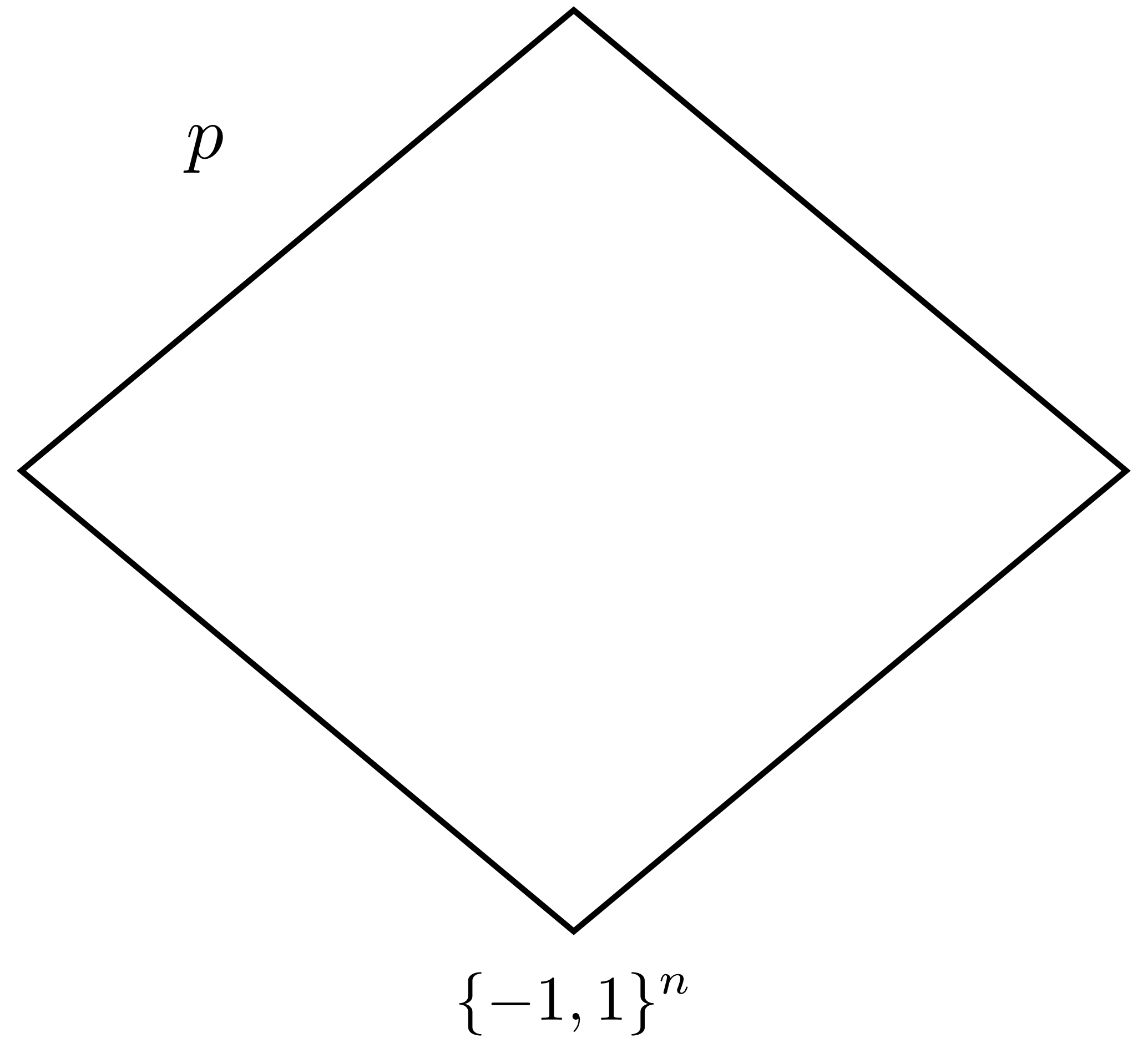
- $p(x) = 1/2^{-n}$ for all $x \in \{-1, 1\}^n$.



Why should something like this be true?

A distribution is uniform iff:

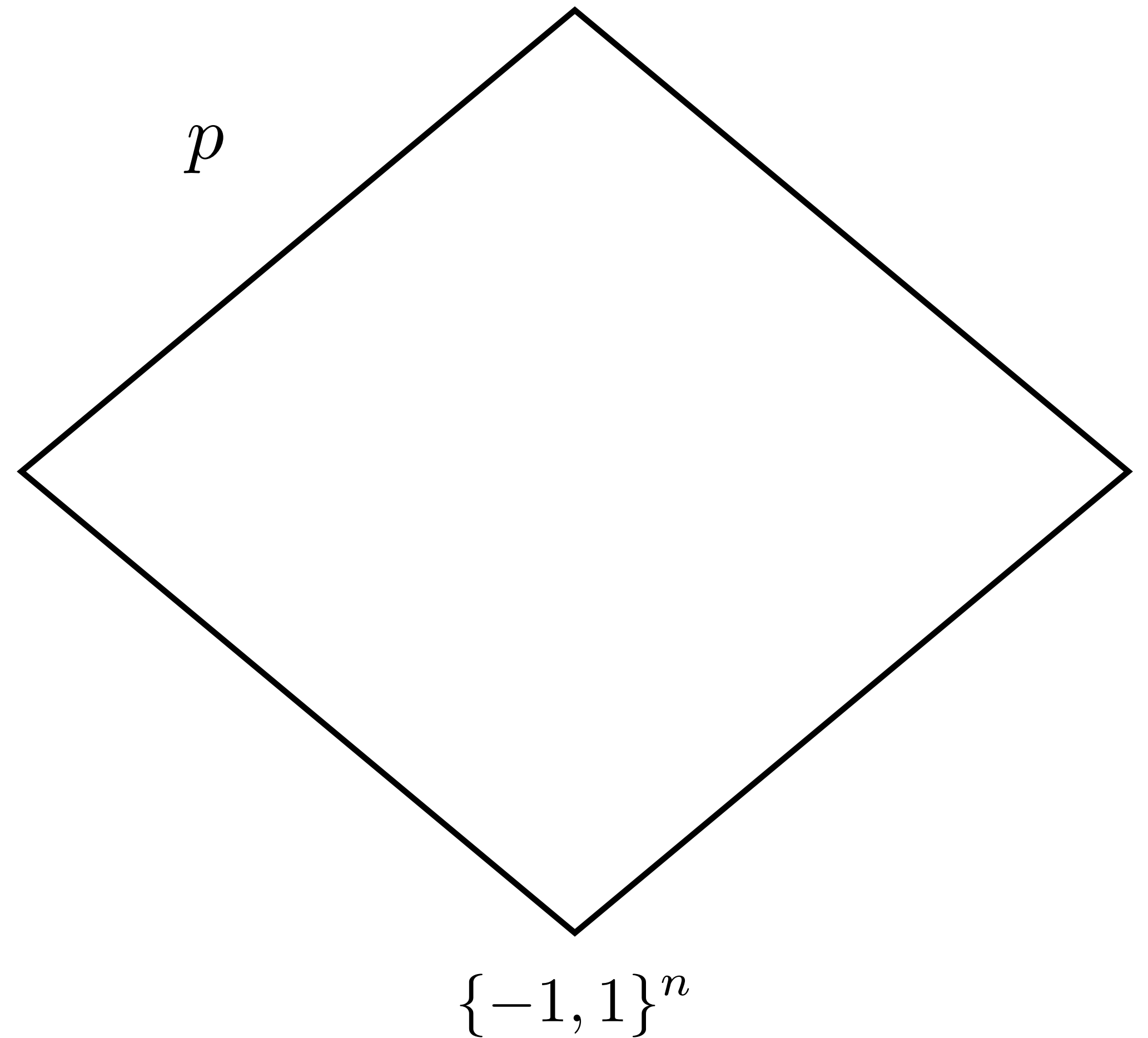
- $p(x) = 1/2^{-n}$ for all $x \in \{-1, 1\}^n$.
- $\|p\|_2^2$ minimal



Why should something like this be true?

A distribution is uniform iff:

- $p(x) = 1/2^{-n}$ for all $x \in \{-1, 1\}^n$.
- $\|p\|_2^2$ minimal
- $p(x) = p(x^{(i)}) \quad \forall x \in \{-1, 1\}^n, \forall i \in [n]$



New Theorem:

$$\mathbb{E}_{\substack{i \sim [\log n] \\ \rho \sim \mathcal{R}_i(p)}} [\|\mu(p|_{\rho})\|_2] \gtrsim \frac{1}{\text{polylog}(n)} \cdot d_{\text{TV}}(p, u)$$

Intuition of the Proof — 3 Conceptual Steps

New Theorem: $\frac{1}{\text{polylog}(n)} \cdot d_{\text{TV}}(p, u) \lesssim \mathbb{E}_{\substack{i \sim [\log n] \\ \rho \sim \mathcal{R}_i(p)}} [\|\mu(p|_{\rho})\|_2]$

Intuition of the Proof — 3 Conceptual Steps

New Theorem: $\frac{1}{\text{polylog}(n)} \cdot d_{\text{TV}}(p, u) \lesssim \mathbb{E}_{\substack{i \sim [\log n] \\ \rho \sim \mathcal{R}_i(p)}} [\|\mu(p|_{\rho})\|_2]$

Intuition of the Proof — 3 Conceptual Steps



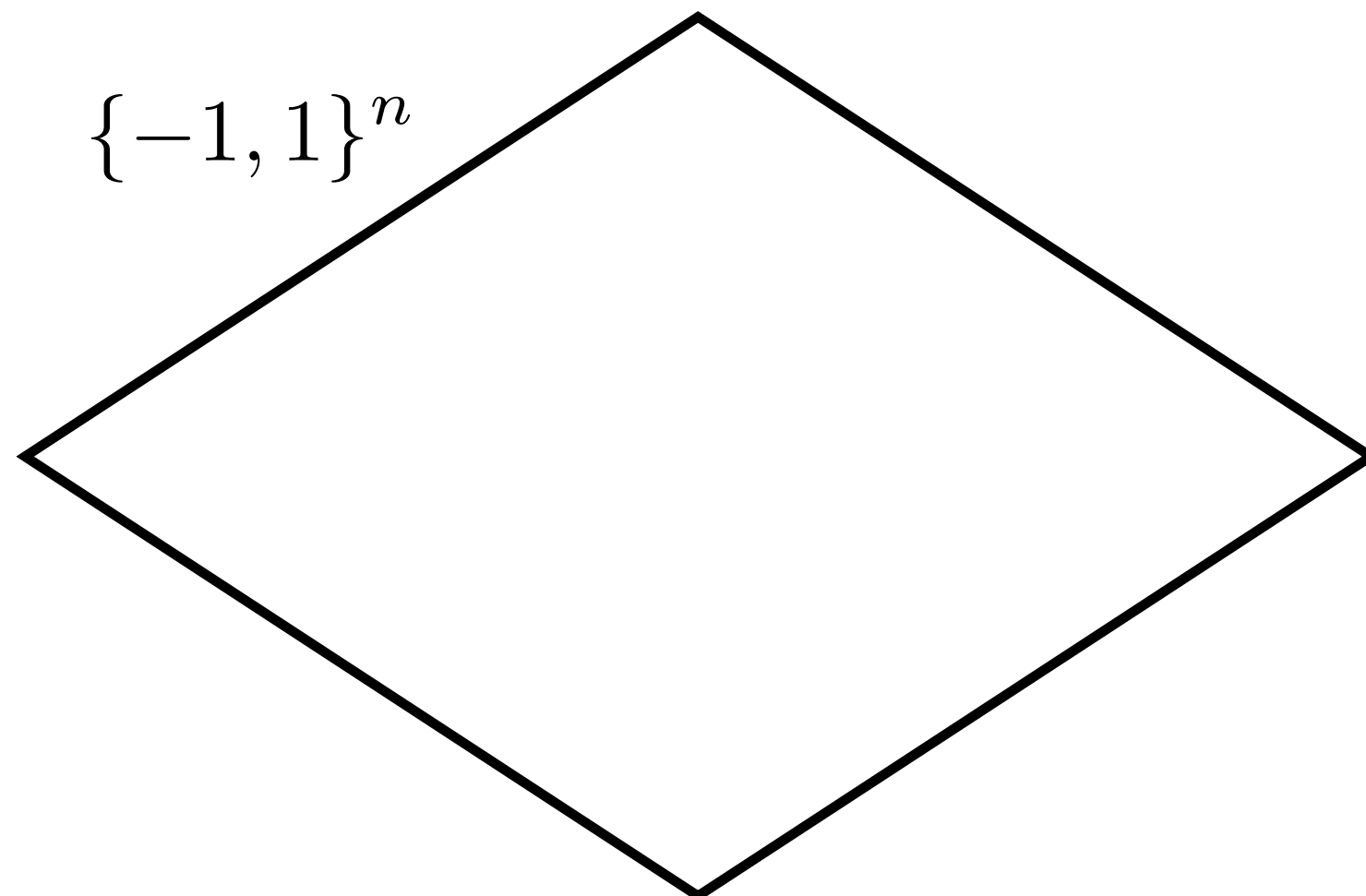
New Theorem: $\frac{1}{\text{polylog}(n)} \cdot d_{\text{TV}}(p, u) \lesssim \mathbb{E}_{\substack{i \sim [\log n] \\ \rho \sim \mathcal{R}_i(p)}} [\|\mu(p|_{\rho})\|_2]$

Intuition of the Proof — 3 Conceptual Steps

**Isoperimetric
Inequalities**



**Symmetry in the
Hypercube**



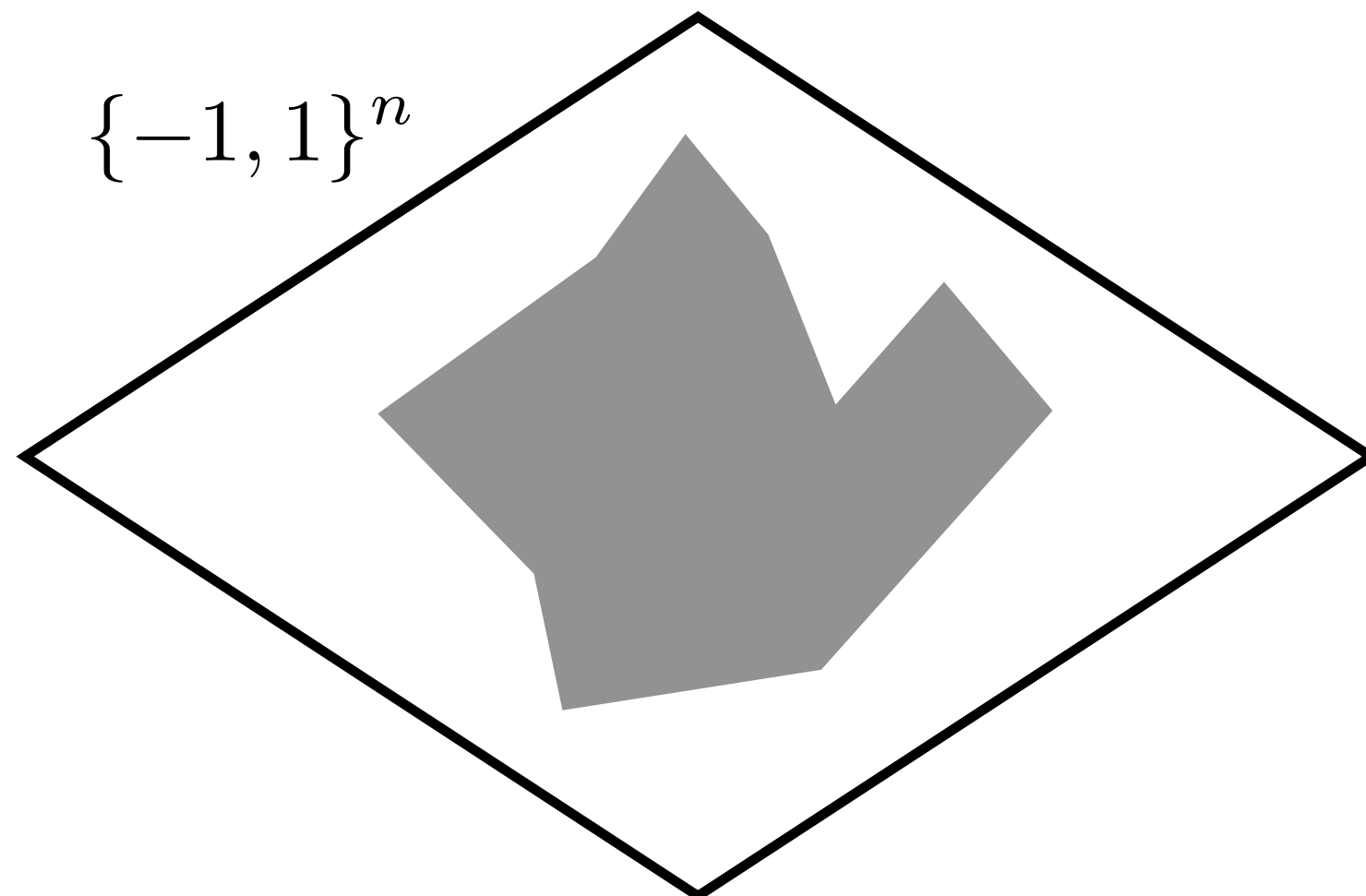
New Theorem: $\frac{1}{\text{polylog}(n)} \cdot d_{\text{TV}}(p, u) \lesssim \mathbb{E}_{\substack{i \sim [\log n] \\ \rho \sim \mathcal{R}_i(p)}} [\|\mu(p|_{\rho})\|_2]$

Intuition of the Proof — 3 Conceptual Steps

**Isoperimetric
Inequalities**



**Symmetry in the
Hypercube**



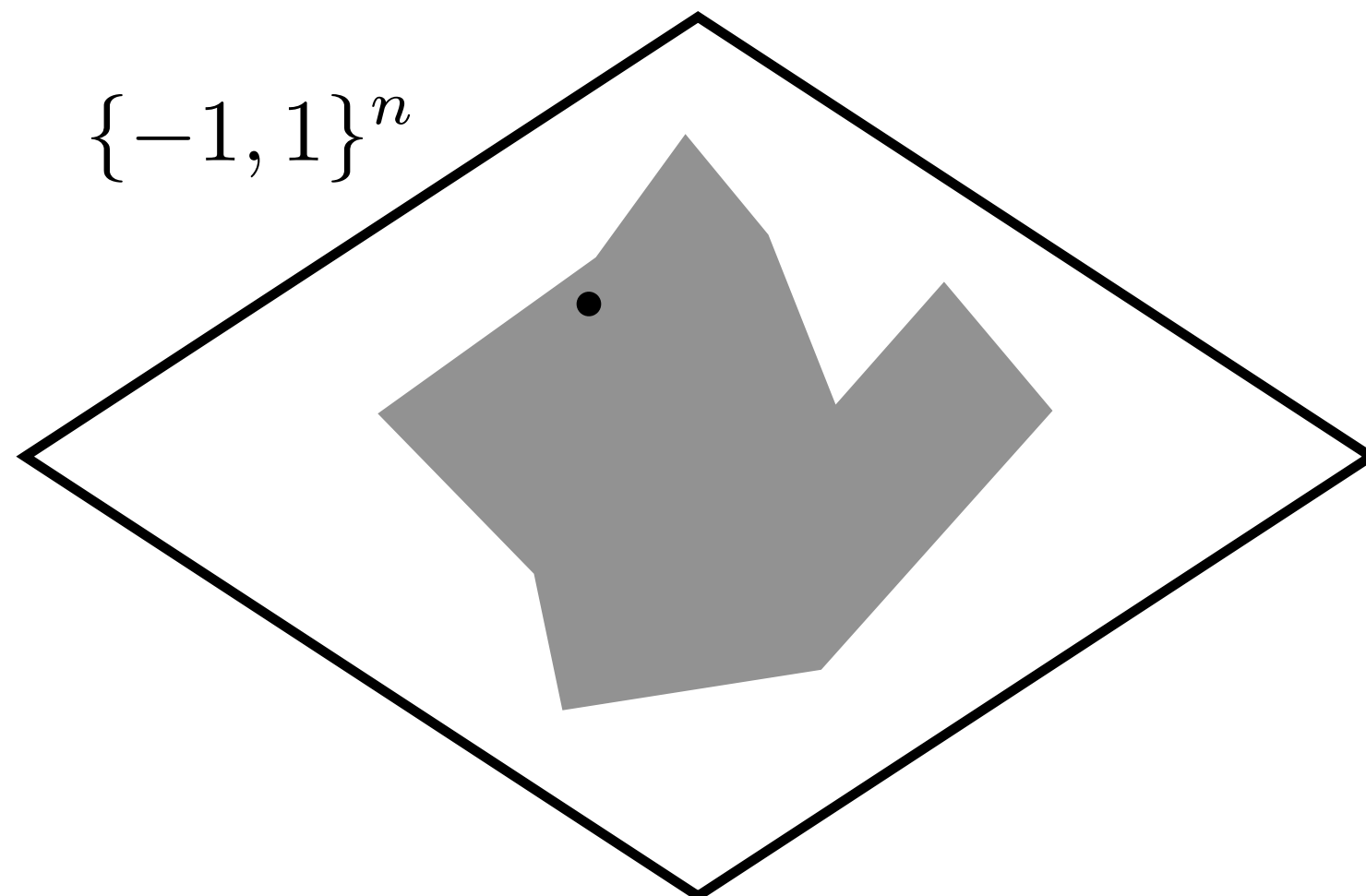
New Theorem: $\frac{1}{\text{polylog}(n)} \cdot d_{\text{TV}}(p, u) \lesssim \mathbb{E}_{\substack{i \sim [\log n] \\ \rho \sim \mathcal{R}_i(p)}} [\|\mu(p|_{\rho})\|_2]$

Intuition of the Proof — 3 Conceptual Steps

**Isoperimetric
Inequalities**



**Symmetry in the
Hypercube**



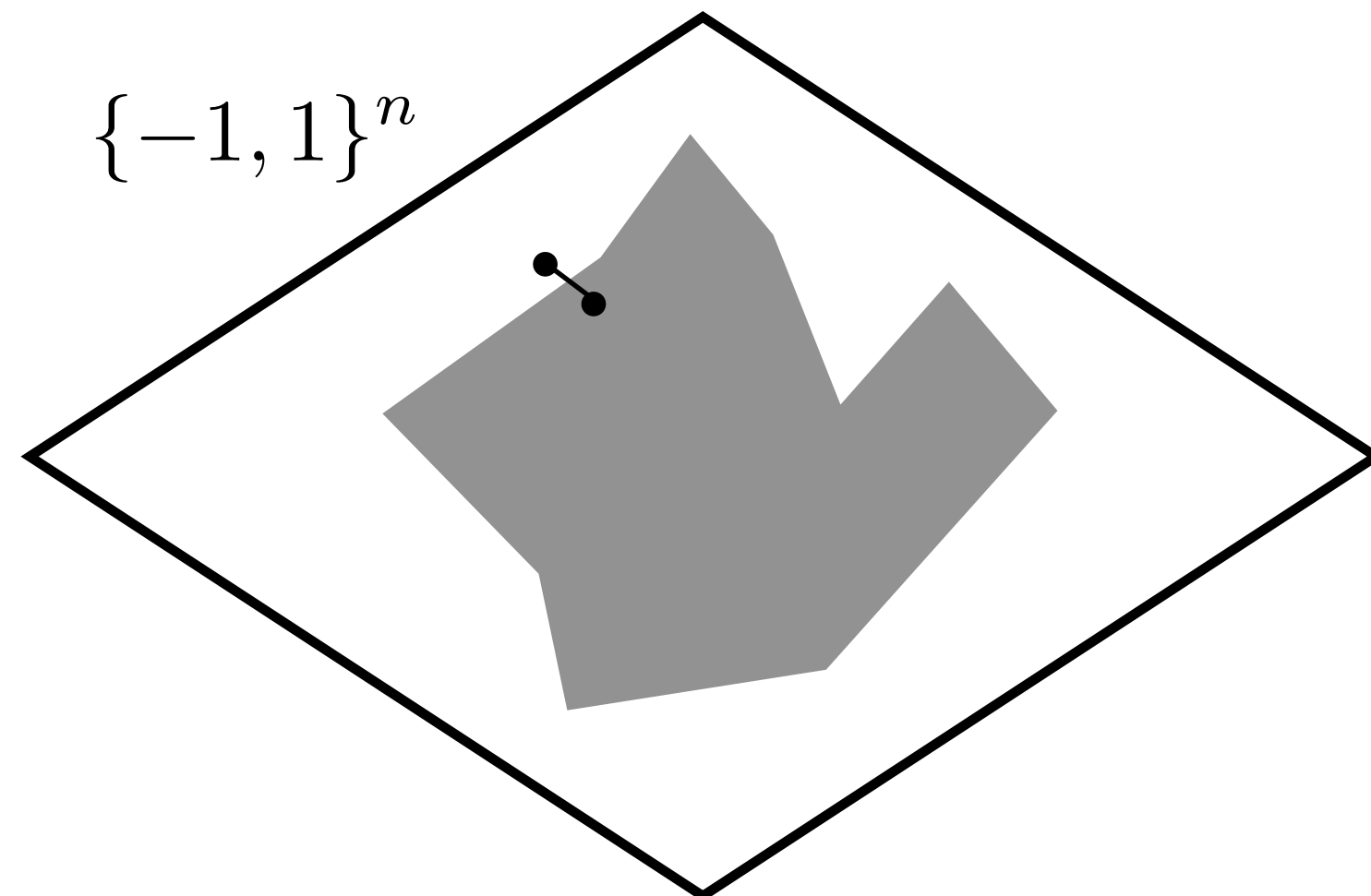
New Theorem:
$$\frac{1}{\text{polylog}(n)} \cdot d_{\text{TV}}(p, u) \lesssim \mathbb{E}_{\substack{i \sim [\log n] \\ \rho \sim \mathcal{R}_i(p)}} [\|\mu(p|_{\rho})\|_2]$$

Intuition of the Proof — 3 Conceptual Steps

**Isoperimetric
Inequalities**



**Symmetry in the
Hypercube**



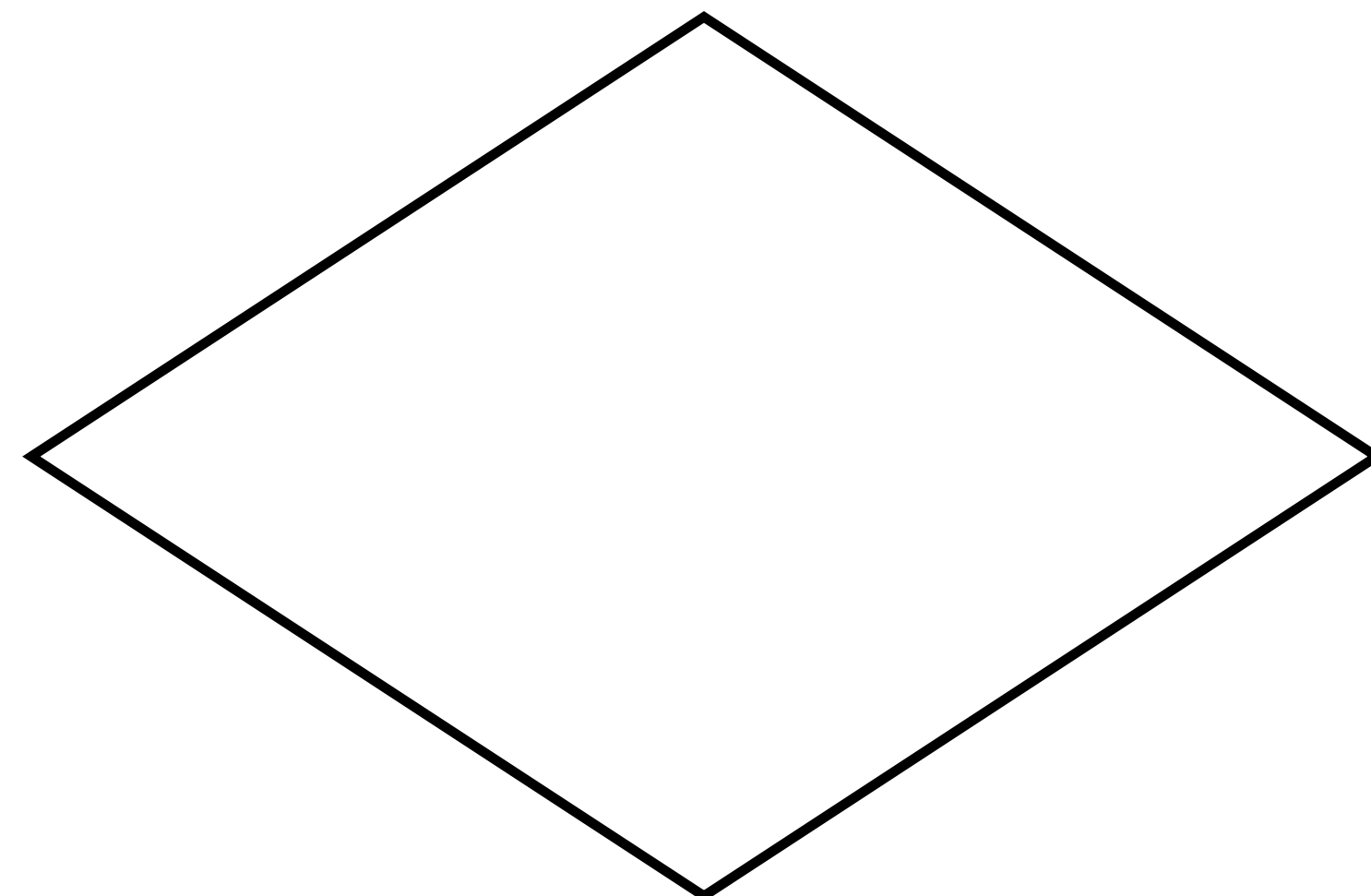
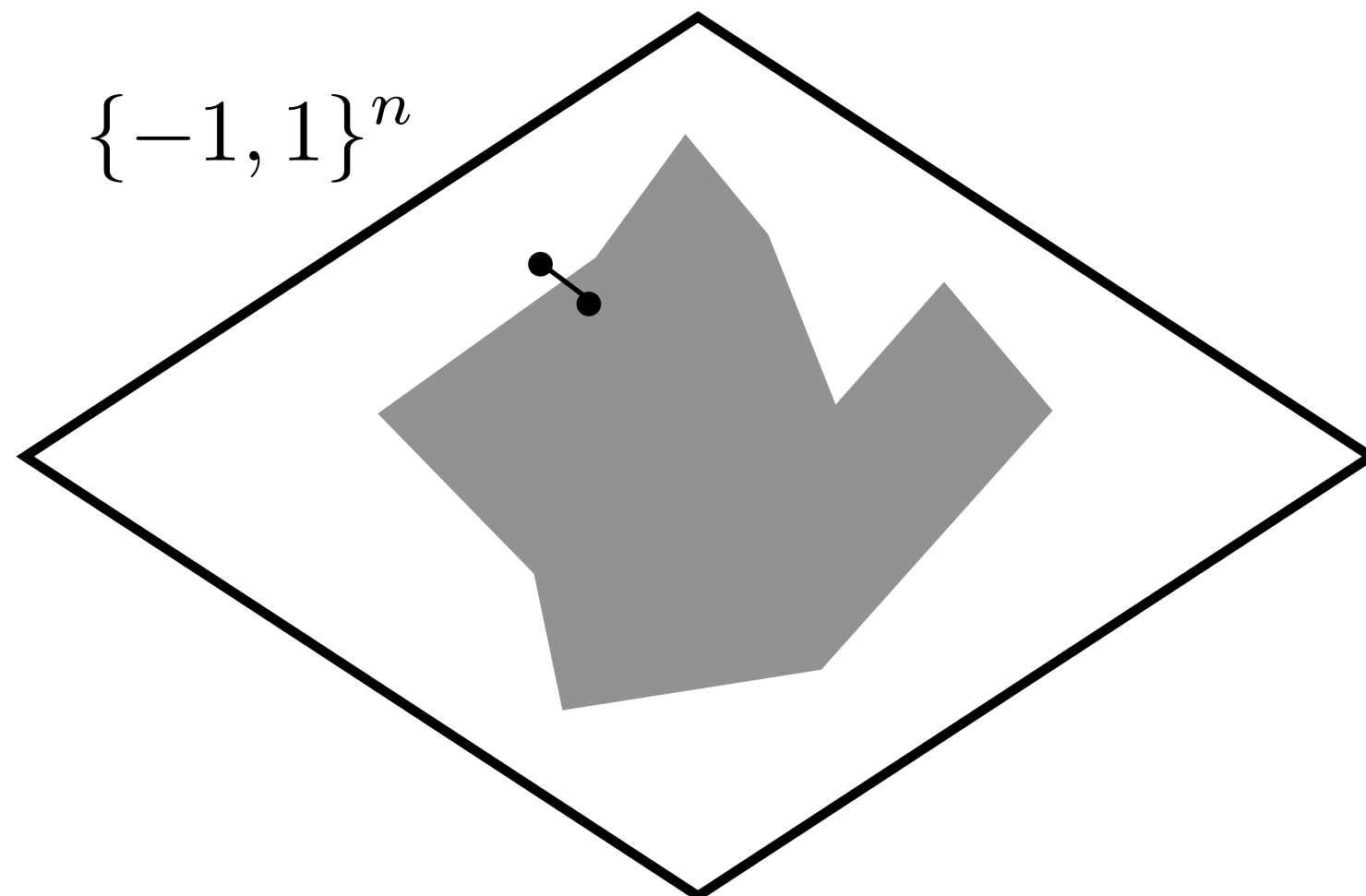
New Theorem:
$$\frac{1}{\text{polylog}(n)} \cdot d_{\text{TV}}(p, u) \lesssim \mathbb{E}_{\substack{i \sim [\log n] \\ \rho \sim \mathcal{R}_i(p)}} [\|\mu(p|_{\rho})\|_2]$$

Intuition of the Proof — 3 Conceptual Steps

**Isoperimetric
Inequalities**



**Symmetry in the
Hypercube**



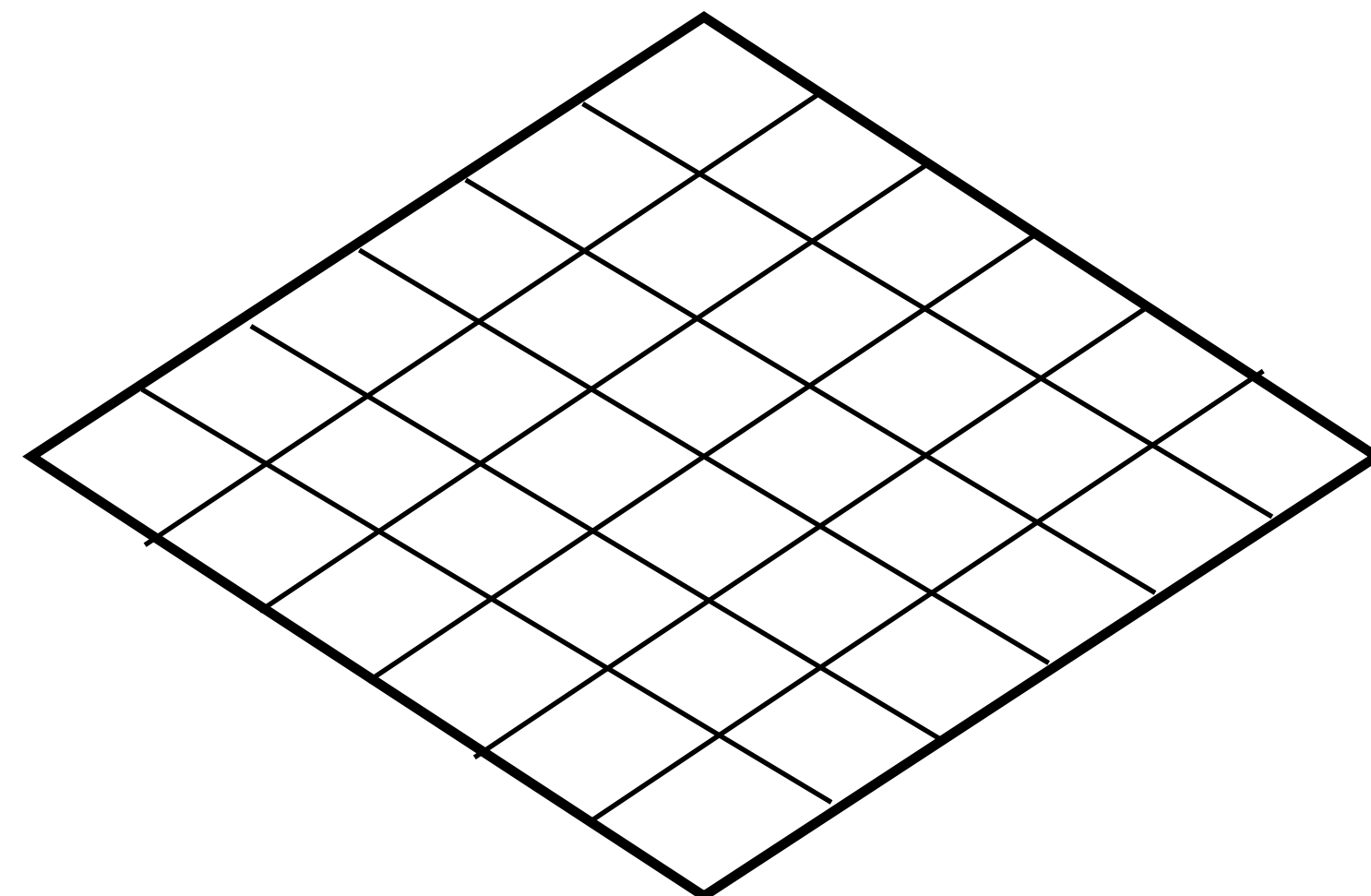
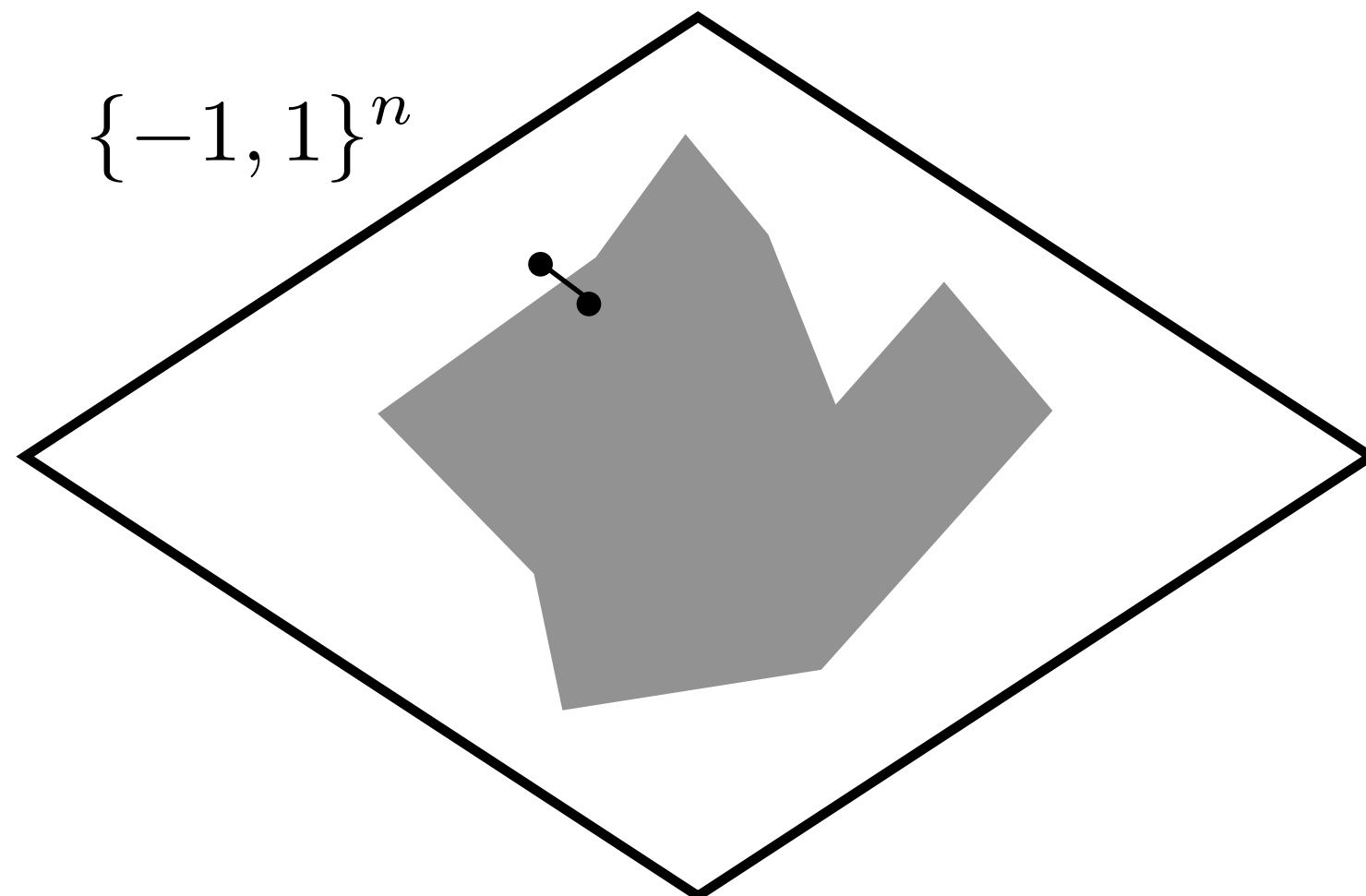
New Theorem:
$$\frac{1}{\text{polylog}(n)} \cdot d_{\text{TV}}(p, u) \lesssim \mathbb{E}_{\substack{i \sim [\log n] \\ \rho \sim \mathcal{R}_i(p)}} [\|\mu(p|_{\rho})\|_2]$$

Intuition of the Proof — 3 Conceptual Steps

**Isoperimetric
Inequalities**



**Symmetry in the
Hypercube**



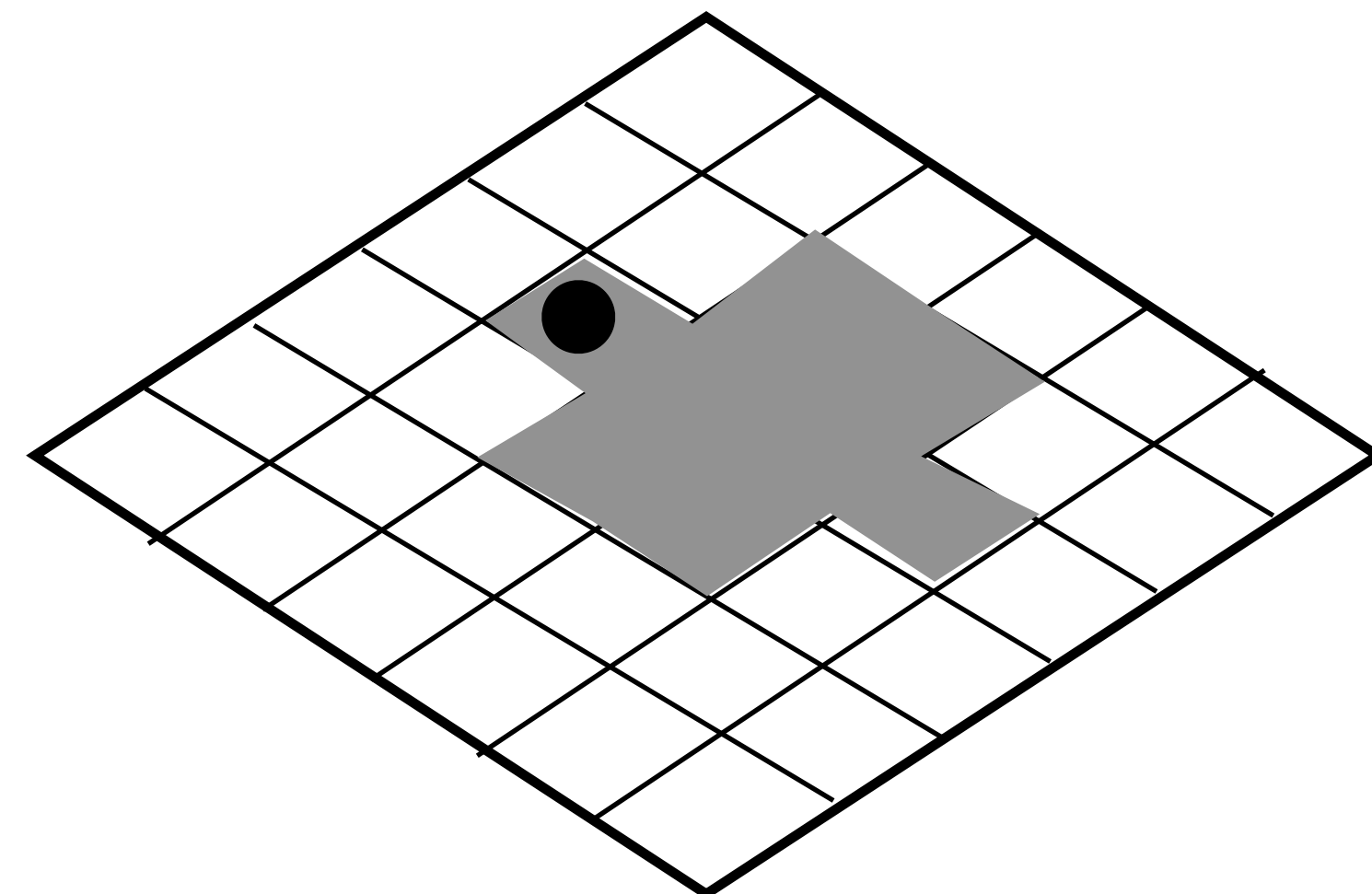
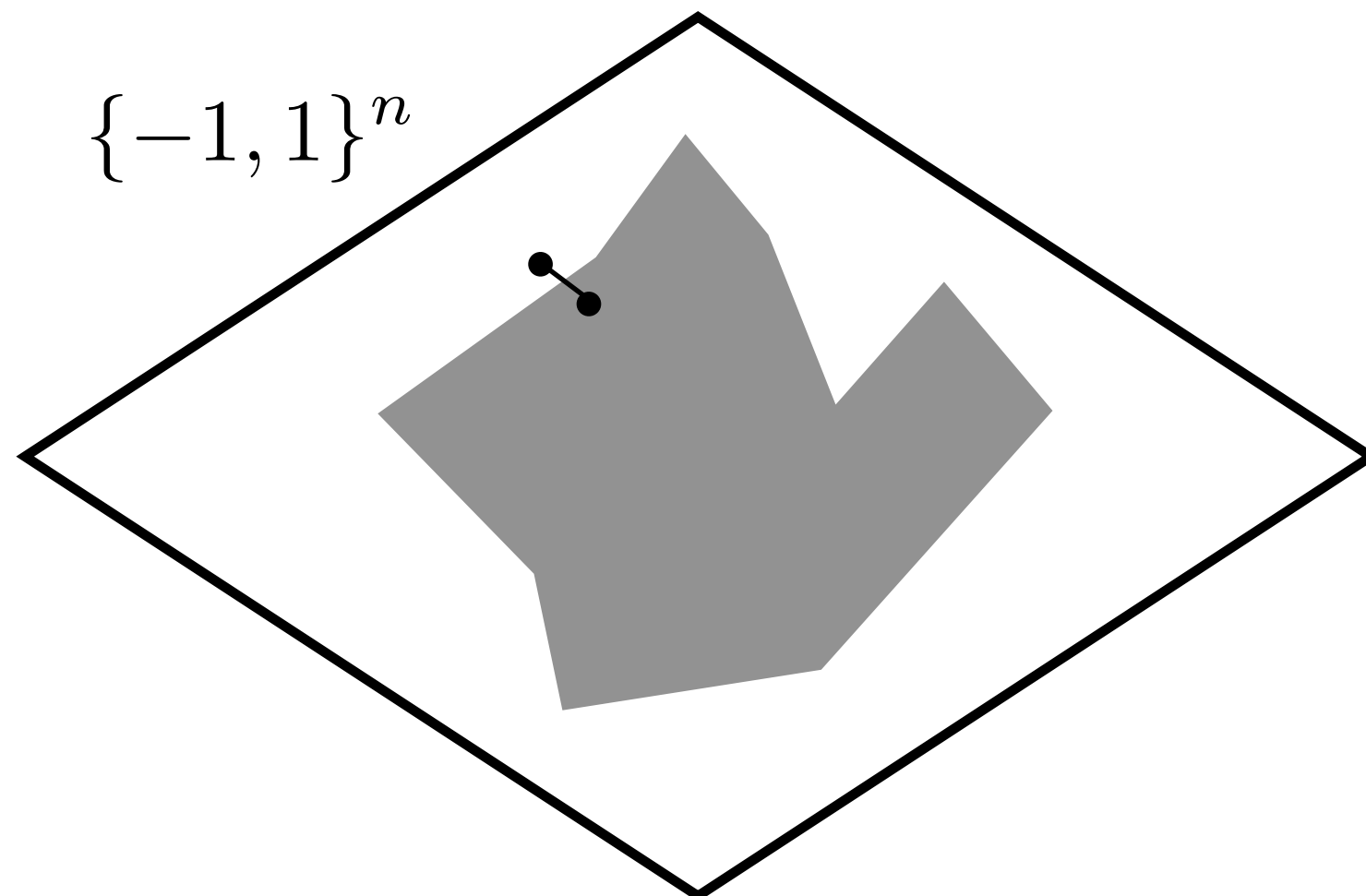
New Theorem: $\frac{1}{\text{polylog}(n)} \cdot d_{\text{TV}}(p, u) \lesssim \mathbb{E}_{\substack{i \sim [\log n] \\ \rho \sim \mathcal{R}_i(p)}} [\|\mu(p|_{\rho})\|_2]$

Intuition of the Proof — 3 Conceptual Steps

Isoperimetric
Inequalities



Symmetry in the
Hypercube



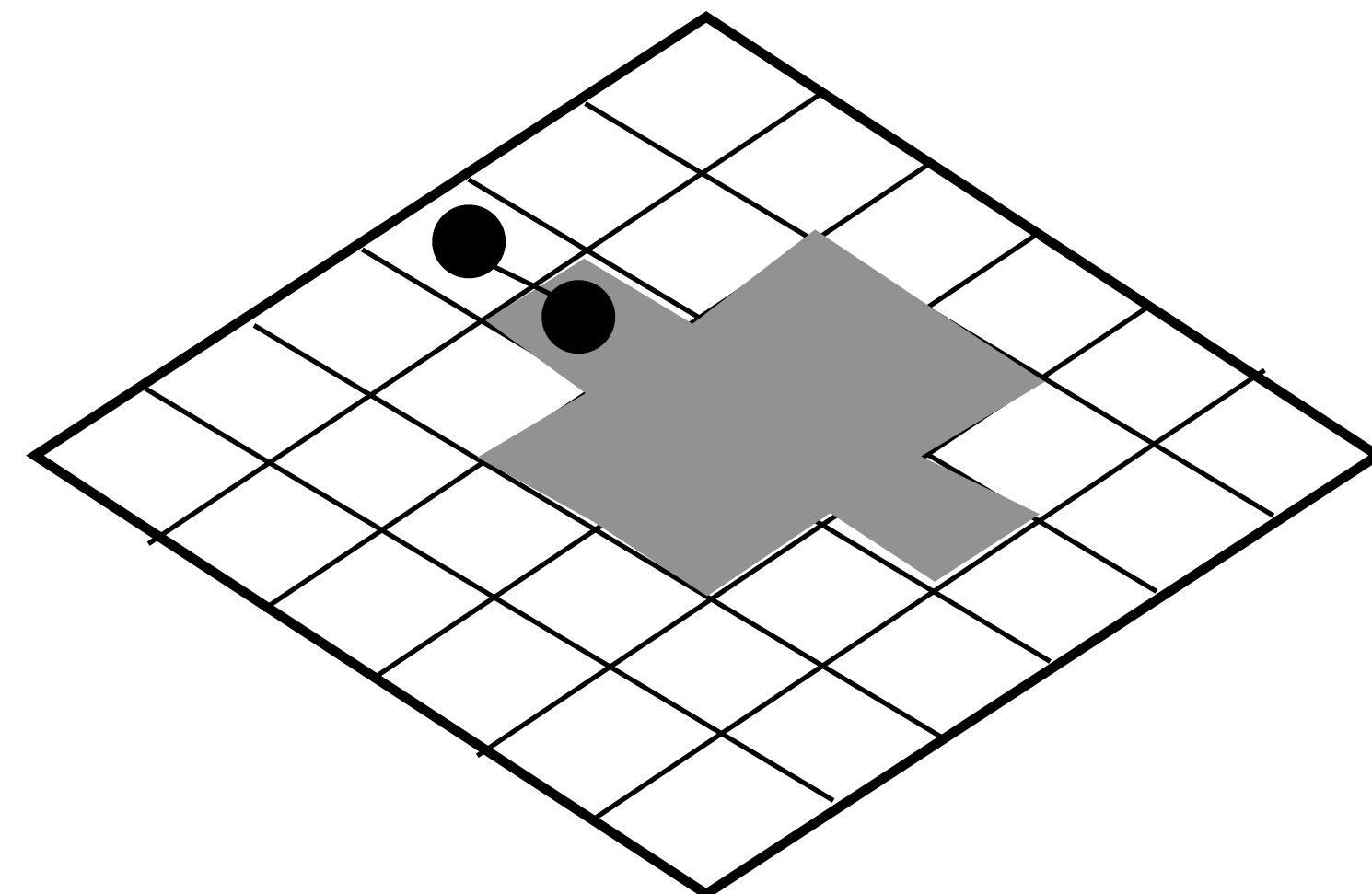
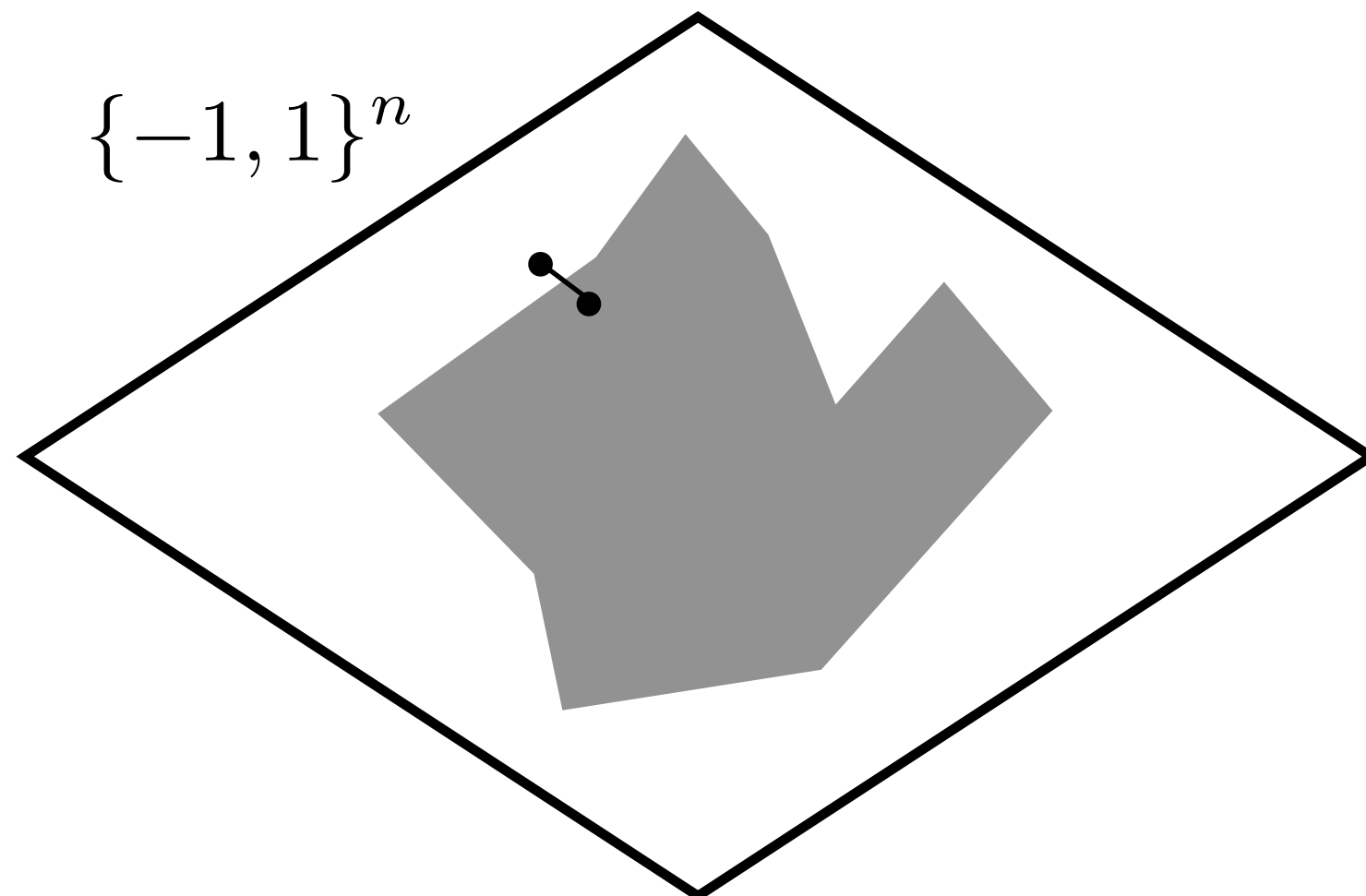
New Theorem: $\frac{1}{\text{polylog}(n)} \cdot d_{\text{TV}}(p, u) \lesssim \mathbb{E}_{\substack{i \sim [\log n] \\ \rho \sim \mathcal{R}_i(p)}} [\|\mu(p|_{\rho})\|_2]$

Intuition of the Proof — 3 Conceptual Steps

Isoperimetric
Inequalities



Symmetry in the
Hypercube



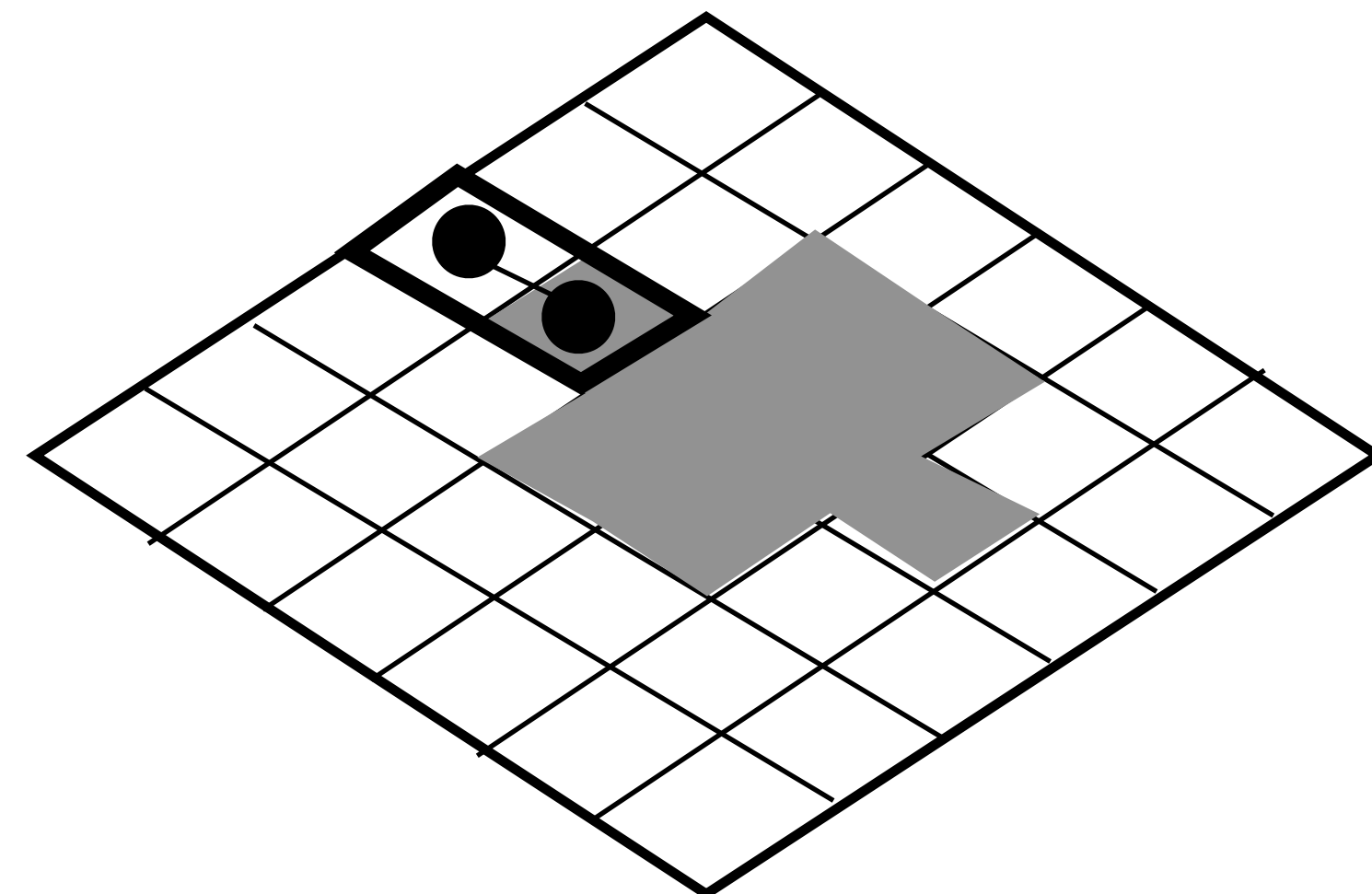
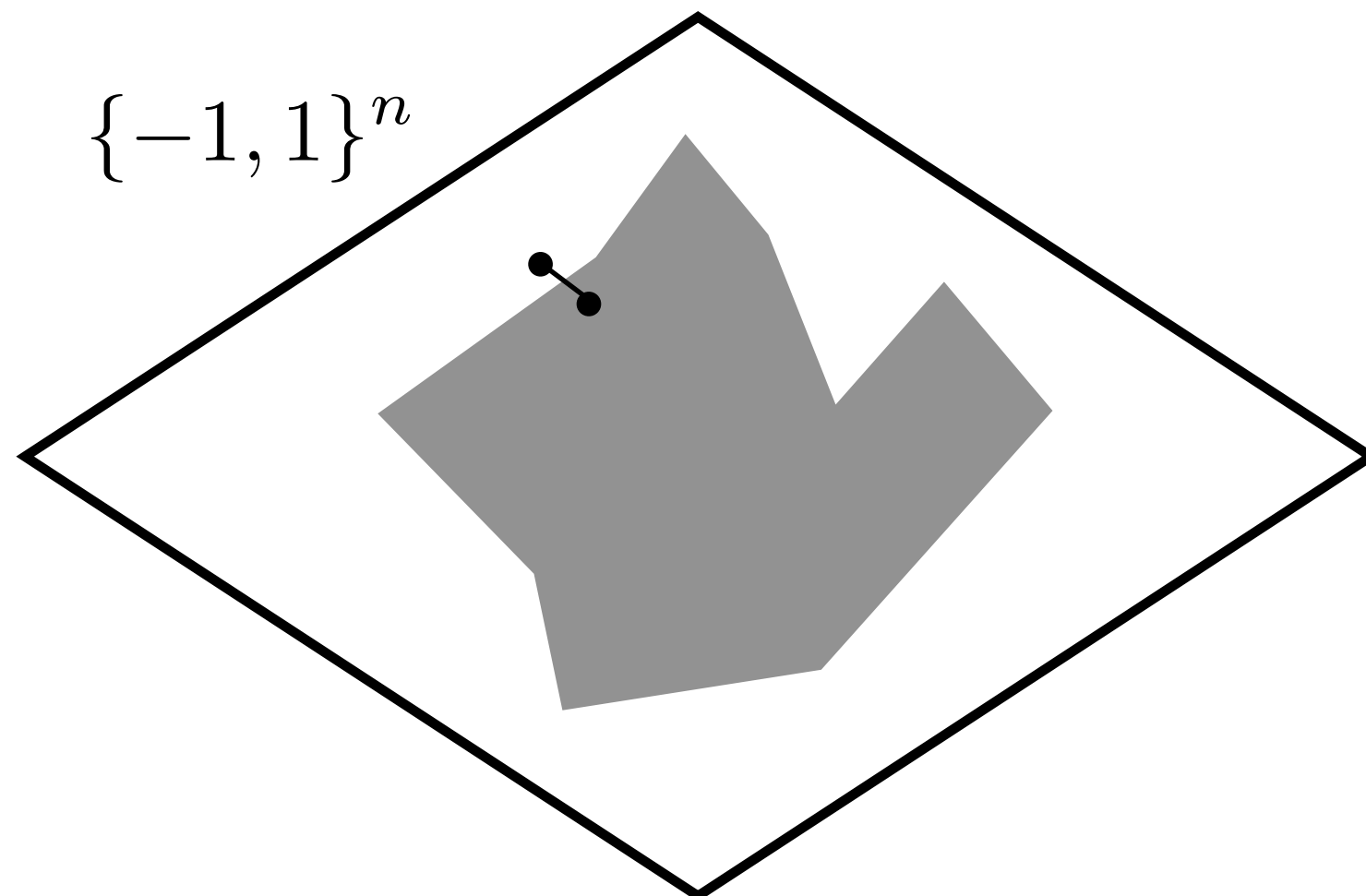
New Theorem:
$$\frac{1}{\text{polylog}(n)} \cdot d_{\text{TV}}(p, u) \lesssim \mathbb{E}_{\substack{i \sim [\log n] \\ \rho \sim \mathcal{R}_i(p)}} [\|\mu(p|_{\rho})\|_2]$$

Intuition of the Proof — 3 Conceptual Steps

Isoperimetric
Inequalities



Symmetry in the
Hypercube



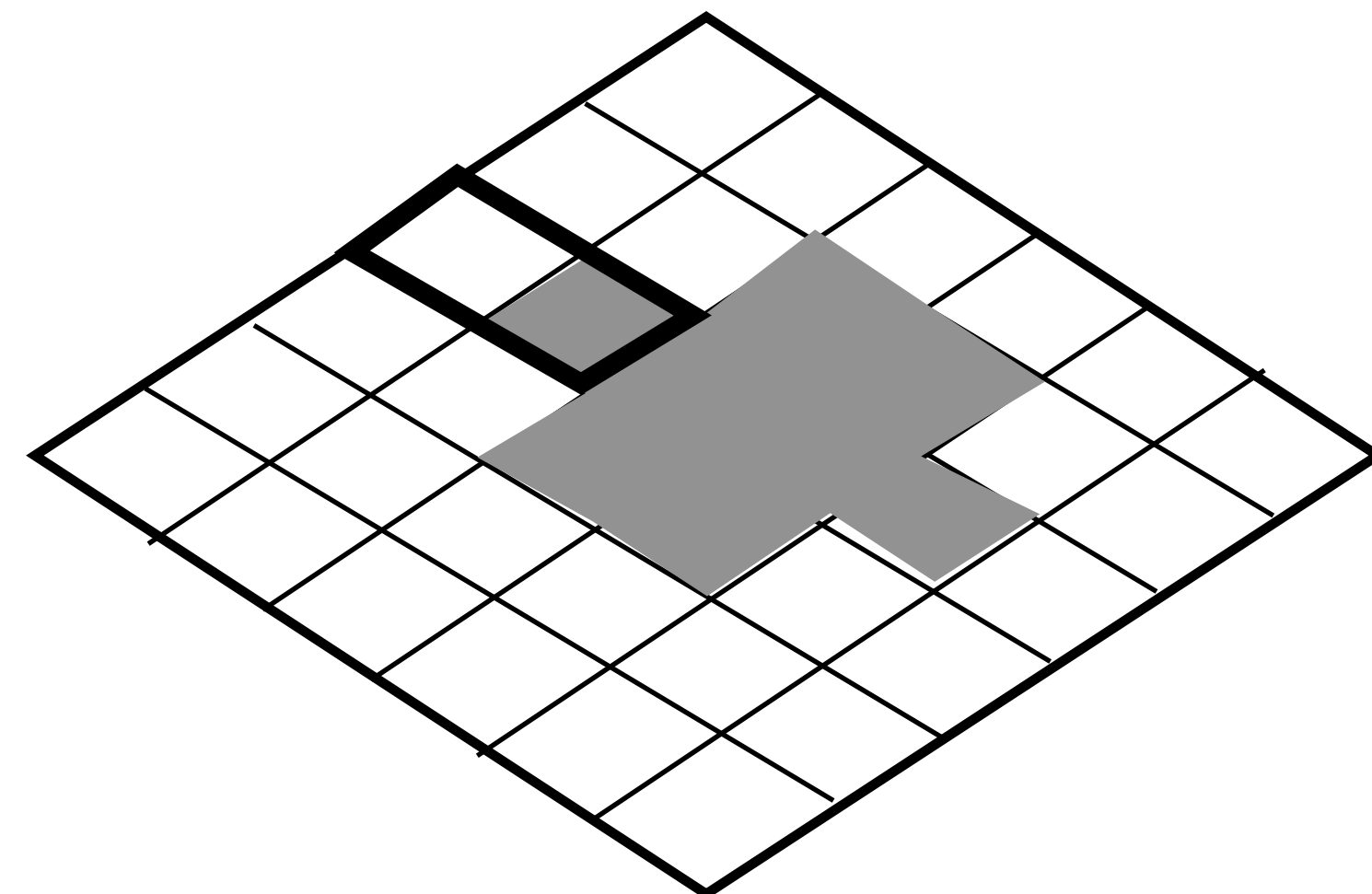
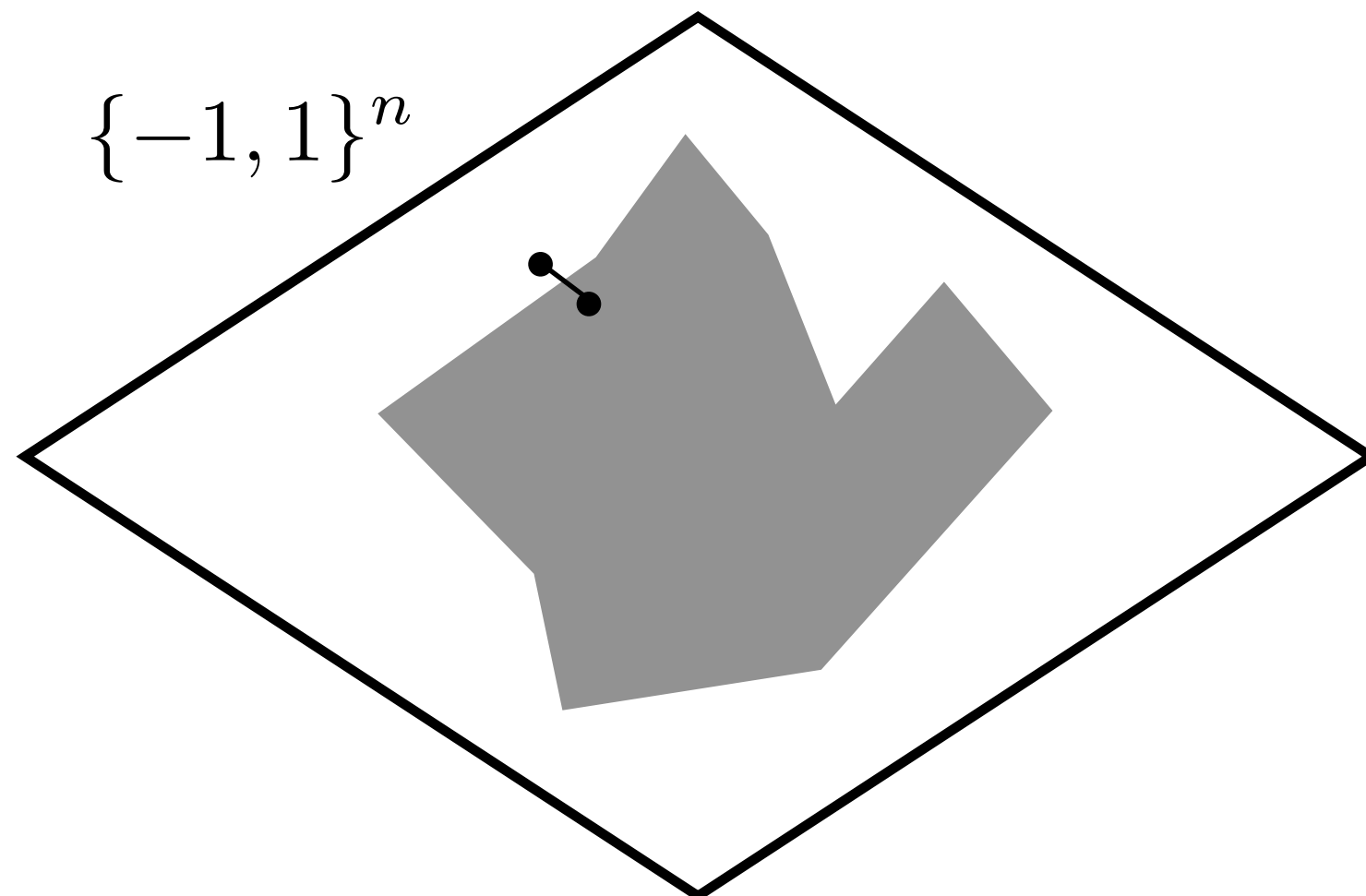
New Theorem: $\frac{1}{\text{polylog}(n)} \cdot d_{\text{TV}}(p, u) \lesssim \mathbb{E}_{\substack{i \sim [\log n] \\ \rho \sim \mathcal{R}_i(p)}} [\|\mu(p|_{\rho})\|_2]$

Intuition of the Proof — 3 Conceptual Steps

Isoperimetric
Inequalities



Symmetry in the
Hypercube



How to efficiently extract information from a distribution?

Question 1: Is it uniform/the distribution you are expecting?

Question 2: Are two distributions the same?

Question 3: Estimate properties of distributions ...

Thanks!