

PAPER • OPEN ACCESS

Learning from higher-order statistics, efficiently: hypothesis tests, random features, and neural networks^{*}

To cite this article: Eszter Székely *et al* *J. Stat. Mech.* (2025) 084001

View the [article online](#) for updates and enhancements.

You may also like

- [Only strict saddles in the energy landscape of predictive coding networks?](#)
Francesco Innocenti, El Mehdi Achour, Ryan Singh et al.
- [Fixed points and universality classes in coupled Kardar–Parisi–Zhang equations](#)
Dipankar Roy, Abhishek Dhar, Manas Kulkarni et al.
- [How feature learning can improve neural scaling laws](#)
Blake Bordelon, Alexander Atanasov and Cengiz Pehlevan

PAPER: ML 2025

Learning from higher-order statistics, efficiently: hypothesis tests, random features, and neural networks^{*}

Eszter Székely^{1,3}, Lorenzo Bardone^{1,3}, Federica Gerace²
and Sebastian Goldt^{**}

International School of Advanced Studies (SISSA), Trieste, Italy

E-mail: sgoldt@sissa.it, eszekely@sissa.it, lbardone@sissa.it and fgerace@sissa.it

Received 22 May 2025

Accepted for publication 24 June 2025

Published 7 August 2025



Online at stacks.iop.org/JSTAT/2025/084001

<https://doi.org/10.1088/1742-5468/adeb6b>

Abstract. Neural networks excel at discovering statistical patterns in high-dimensional data sets. In practice, higher-order cumulants, which quantify the non-Gaussian correlations between three or more variables, are particularly important for the performance of neural networks. But how efficient are neural networks at extracting features from higher-order cumulants? We study this question in the spiked cumulant model, where the statistician needs to recover a privileged direction or ‘spike’ from the order- $p \geq 4$ cumulants of d -dimensional inputs. We first discuss the fundamental statistical and computational limits of recovering the spike by analysing the number of samples n required to strongly distinguish between inputs from the spiked cumulant model and isotropic Gaussian inputs. Existing literature established the presence of a wide

^{*}This article is an updated version of a paper presented at NEURIPS2024 conference (Szekely E, Bardone L, Gerace F and Goldt S 2024 Learning from higher-order correlations, efficiently: hypothesis tests, random features, and neural networks *Advances in Neural Information Processing Systems* vol 37, ed A Globerson, L Mackey, D Belgrave, A Fan, U Paquet, J Tomczak and C Zhang (Curran Associates, Inc.) pp 78479–522).

¹ Current address: CCFE, Culham Science Centre, Abingdon, Oxon, OX14 3DB, United Kingdom.

² Current address: Dipartimento di Matematica, Università di Bologna, Bologna (BO), Italy.

³ These authors contributed equally.

^{**} Author to whom any correspondence should be addressed.



Original Content from this work may be used under the terms of the [Creative Commons Attribution 4.0 licence](https://creativecommons.org/licenses/by/4.0/). Any further distribution of this work must maintain attribution to the author(s) and the title of the work, journal citation and DOI.

statistical-to-computational gap in this problem. We deepen this line of work by finding an exact formula for the likelihood ratio norm which proves that statistical distinguishability requires $n \gtrsim d$ samples, while distinguishing the two distributions in polynomial time requires $n \gtrsim d^2$ samples for a wide class of algorithms, i.e. those covered by the low-degree conjecture. Numerical experiments show that neural networks do indeed learn to distinguish the two distributions with quadratic sample complexity, while ‘lazy’ methods like random features (RFs) are not better than random guessing in this regime. Our results show that neural networks extract information from higher-order correlations in the spiked cumulant model efficiently, and reveal a large gap in the amount of data required by neural networks and RFs to learn from higher-order cumulants.

Keywords: learning theory, machine learning, statistical inference, deep learning

Contents

1. Introduction	3
1.1. Further related work	5
1.1.1. Detecting spikes in high-dimensional data	5
1.1.2. Non-Gaussian Component Analysis (NGCA), Gaussian pancakes and low-degree polynomials	5
1.1.3. Separation between neural networks and RFs	5
1.1.4. Reproducibility	6
2. The data models	6
2.1. The Gaussian case	6
2.2. The spiked cumulant model	6
3. How many samples do we need to learn?	7
3.1. Statistical distinguishability: LR analysis	7
3.2. Computational distinguishability: LDLR analysis	9
3.3. Statistical-to-computational gaps in the spiked cumulant model	11
4. Learning from HOCs with neural networks and RFs	11
4.1. Spiked Wishart model	12
4.2. Spiked cumulant	13
4.3. Phase transitions and neural network performance in a simple model for images	14
5. Concluding perspectives	16
Acknowledgments	16
Appendix A. Experimental details	16
A.1. Figures 2 and 3	16

A.1.1. Starting from small initial overlap	17
A.2. Figure 4	18
Appendix B. Mathematical details on the hypothesis testing problems...	18
B.1. Notation	18
B.1.1. Multi-index notation	18
B.2. More on LR and LDLR	19
B.2.1. Statistical and computational distinguishability	19
B.2.2. Necessary conditions for distinguishability	20
B.2.3. LDLR analysis for spiked Wishart model.....	21
B.3. Hermite Polynomials	23
B.3.1. Hermite coefficients.....	25
B.3.2. Links with cumulant theory	25
B.4. Details on the spiked cumulant model	26
B.4.1. Prior distribution on the spike	26
B.4.2. Distribution of the non-Gaussianity	26
B.4.3. Computing the whitening matrix	28
B.5. Details of the LR analysis for spiked cumulant model	29
B.5.1. Proof sketch for theorem 2	29
B.5.2. Proof of theorem 2	30
B.5.3. Consequences of theorem 2	32
B.6. LDLR on spiked cumulant model.....	34
B.6.1. Proof sketch	35
B.6.2. Detailed proof	36
B.7. Limitations of our theoretical analysis	44
Appendix C. Details on the replica analysis	44
References.....	47

1. Introduction

Discovering statistical patterns in high-dimensional data sets is the key objective of machine learning. In a classification task, the differences between inputs in different classes arise at different statistical levels of the inputs: two different classes of images will usually have different means, different covariances, and different higher-order cumulants (HOCs), which describe the non-Gaussian part of the correlations between three or more pixels. While differences in the mean and covariance allow for rudimentary classification, Refinetti *et al* [1] recently highlighted the importance of HOCs for the performance of neural networks: when they removed the HOCs per class of the CIFAR10 training set, the test accuracy of various deep neural networks dropped by up to 65 percentage points. The importance of HOCs for classification in general and the performance of

neural networks in particular raises some fundamental questions: what are the fundamental limits of learning from HOCs, i.e. how many samples n (‘sample complexity’) are required to extract information from the HOCs of a data set? How many samples are required when using a *tractable* algorithm? And how do neural networks and other machine learning methods like random features (RFs) compare to those fundamental limits?

In this paper, we study these questions by analysing a series of binary classification tasks. In one class, inputs $x \in \mathbb{R}^d$ are drawn from a normal distribution with zero mean and identity covariance. These inputs are therefore isotropic: the distribution of the high-dimensional points projected along a unit vector in any direction in $\mathbb{R}^d$ is a standard Gaussian distribution. Furthermore, all the HOCs (of order $p \geq 3$) of the inputs are identically zero. Inputs in the second class are also isotropic, except for one special direction $u \in \mathbb{R}^d$ in which their distribution is different. This direction u is often called a ‘spike’, and it can be encoded in different cumulants: for example, we could ‘spike’ the covariance by drawing inputs from a Gaussian with mean zero and covariance $\mathbf{1} + \beta uu^\top$; the signal-to-noise ratio $\beta > 0$ would then control the variance of $\lambda = \langle u, x \rangle$. Likewise, we could spike a HOC of the distribution and ask: what is the minimal number of samples required for a neural network trained with SGD to distinguish between the two input classes? This simple task serves as a proxy for more generic tasks: a neural network cannot be able to *extract* information from a given cumulant if it cannot even recognise that it is different from an isotropic one.

We can obtain the fundamental limits of detecting spikes at different levels of the cumulant hierarchy by considering the hypothesis test between a null hypothesis (the isotropic multivariate normal distribution with zero mean and identity covariance) and an alternative ‘spiked’ distribution. We can then compare the sample complexity of neural networks to the number of samples necessary to distinguish the two distributions using unlimited computational power, or efficiently using algorithms that run in polynomial time. A second natural comparison for neural networks are RFs or kernel methods. Since the discovery of the neural tangent kernel [2], and the practical success of kernels derived from neural networks [3, 4], there has been intense interest in establishing the advantage of neural networks with *learned* feature maps over classical methods with *fixed* feature maps, like kernel machines. The role of higher-order correlations for the relative advantage of these methods has not been studied yet.

In the following, we first introduce some fundamental notions around hypothesis tests and in particular the low-degree method [5–8], which will be a key tool in our analysis, using the classic spiked Wishart model for sample covariance matrices [9, 10]. For spiked cumulants, the existing literature (see section 1.1 for a complete discussion) establishes the presence of a wide statistical-to-computational gap in this model. Our **main contributions** are then as follows:

- We deepen the understanding of the statistical-to-computational gap for learning from higher-order correlations on the statistical side by showing that at *unbounded computational power*, a number of samples *linear* in the input dimension is required to reach *statistical distinguishability*. We prove this by explicitly computing the norm of the *likelihood ratio* (LR), see theorem 2.

- On the algorithmic side, statistical query (SQ) bounds and previous low-degree analyses [11, 12], showed that the sample complexity of learning from HOCs is instead *quadratic* for a wide class of polynomial-time algorithms (section 3.2). Here, we provide a different, more direct proof of such a bound.
- Using these fundamental limits on learning from HOCs as benchmarks, we show numerically that neural networks learn the mixed cumulant model efficiently, while random features do not, revealing a large separation between the two methods (sections 4.1 and 4.2).
- We finally show numerically that the distinguishability in a simple model for images [13] is precipitated by a cross-over in the behaviour of the HOCs of the inputs (section 4.3).

1.1. Further related work

1.1.1. Detecting spikes in high-dimensional data. There is an enormous literature on statistical-to-computational gaps in the detection and estimation of variously structured principal components of high-dimensional data sets. This problem has been studied for the Wigner and Wishart ensembles of random matrices [9, 14–24] and for spiked tensors [25–33]. For comprehensive reviews of the topic, see [34–37]. While the samples in these models are often non-Gaussian, depending on the prior distribution over the spike u , the spike appears already at the level of the covariance of the inputs. Here, we instead study a high-dimensional data model akin to the one used by Wang & Lu [38] to study online independent component analysis (ICA). This data model can be interpreted as having an additional whitening step to pre-process the inputs, which is a common pre-processing step in ICA [39], hence inputs have identity covariance even when cumulants of order $p \geq 3$ are spiked. Wang & Lu [38] proved the existence of a scaling limit for online ICA, but did not consider the sample complexity of recovering the spike/distinguishing the distributions, which is the focus of this paper.

1.1.2. Non-Gaussian Component Analysis (NGCA), Gaussian pancakes and low-degree polynomials. Related models have been introduced under the name of NGCA [40], and studied from the point of view of SQ complexity in a sequence of papers [41–46]. In particular [11] nicknames *Gaussian pancakes* a class of models that includes the *spiked cumulant model* presented here. The SQ bounds found in [11] and refined in the subsequent works (see Diaconikolas *et al* [47] and references therein) could be used, together with SQ-low-degree likelihood ratio (LDLR) equivalence [48], to provide estimates on low-degree polynomials. Finally, [12] proves very general bounds on LDLR norm; our theorem 5 provides an alternative derivation of these bounds, in a setting that is closer to the experiments in section 4.

1.1.3. Separation between neural networks and RFs. The discovery of the neural tangent kernel by Jacot *et al* [2] and the flurry of results on linearised neural networks [2, 3, 49–55] has triggered a new wave of interest in the differences between what neural networks and kernel methods can learn efficiently. While statistical separation results have been well-known for a long time [56], recent work focused on understanding the differences between RFs and neural networks *trained by gradient descent* both with wide

hidden layers [4, 57–61] or even with just a few hidden neurons [62, 63]. At the heart of the data models in all of these theoretical models is a hidden, low-dimensional structure in the task, either in input space (for mixture classification) or in the form of single- or many-index target functions. The impact of higher-order input correlations in mixture classification tasks on the separation between RFs and neural networks has not been directly studied yet.

1.1.4. Reproducibility. We provide code for all of our experiments, including routines to generate the various synthetic data sets we discuss, on GitHub <https://github.com/eszter137/data-driven-separation>.

2. The data models

Throughout the paper, we consider binary classification tasks where high-dimensional inputs $x^\mu = (x_i^\mu) \in \mathbb{R}^d$ have labels $y^\mu = \pm 1$. The total number of training samples is $2n$, i.e. we have n samples per class. For the class $y^\mu = -1$, inputs $x^\mu = z^\mu$, where $z^\mu \sim_{\text{iid}} \mathcal{N}(0, \mathbf{1}_d)$ and $\mathcal{N}$ denotes the normal distribution. For the class $y^\mu = 1$ instead, we consider ‘spiked’ input models as follows.

2.1. The Gaussian case

The simplest spiked model is the spiked Wishart model from random matrix theory [9, 10], in which

$$x^\mu = \sqrt{\frac{\beta}{d}} g^\mu u + z^\mu, \quad g^\mu \sim_{\text{iid}} \mathcal{N}(0, 1), \quad (1)$$

where $u = (u_i)$ is a d -dimensional vector with norm $\|u\| = \sqrt{d}$ whose elements are drawn element-wise i.i.d. according to some probability distribution $\mathcal{P}_u$. In this model, inputs are Gaussian and indistinguishable from white noise except in the direction of the ‘spike’ u , where they have variance $1 + \beta$, where $\beta > 0$ is the *signal-to-noise ratio*. By construction, all HOCs are zero.

2.2. The spiked cumulant model

To study the impact of HOCs, we draw inputs in the ‘spiked’ class from the data model used by Wang & Lu [38] to study online ICA. First, we sample inputs

$$\tilde{x}^\mu = \sqrt{\frac{\beta}{d}} g^\mu u + z^\mu, \quad g^\mu \sim_{\text{i.i.d.}} p_g, \quad (2)$$

as in the Wishart model, but crucially sample the latent variables g^μ from some non-Gaussian distribution $p_g(g^\mu)$, say, the Rademacher distribution; see assumption 1 for a precise statement. For any non-Gaussian p_g , the resulting inputs $\tilde{x}^\mu$ have a non-trivial fourth-order cumulant proportional to $\kappa_4^g u^{\otimes 4}$, where $\kappa_4^g \equiv \mathbb{E}(g^\mu)^4 - 3\mathbb{E}(g^\mu)^2$ is the fourth cumulant of the distribution of g^μ . We fix the mean and variance of p_g to be zero and

one, respectively, so the covariance of inputs has a covariance matrix $\Sigma = \mathbb{1}_d + \beta/duu^\top$. To avoid trivial detection of the spike from the covariance, the key ingredient of the spiked cumulant model is that we whiten the inputs in that class, so that the inputs are finally given by

$$x^\mu = S\tilde{x}^\mu, \quad S = \mathbb{1} - \frac{\beta}{1 + \beta + \sqrt{1 + \beta}} \frac{uu^\top}{d}, \quad (3)$$

with the whitening matrix S (see appendix B.4.3). Hence inputs x^μ are isotropic Gaussians in all directions except u , where they are a weighted sum of g^μ and $\langle z, u \rangle$. The whitening therefore changes the interpretation of β : rather than being a signal-to-noise ratio, as in the spiked Wishart model, here β controls the quadratic interpolation between the distributions of g^μ and z in the direction of u (see equation (55) in the appendix). This leaves us with a data set where inputs in both classes have an average covariance matrix that is the identity, which means that PCA or linear neural networks [64–67] cannot detect any difference between the two classes.

3. How many samples do we need to learn?

Given a data set sampled from the spiked cumulant model, we can now ask: how many samples does a statistician need to reliably detect whether inputs are Gaussian or not, i.e. whether HOCs are spiked or not? This is equivalent to the hypothesis test between n i.i.d. samples of the isotropic normal distribution in $\mathbb{R}^d$ as the null hypothesis $\mathbb{Q}_{n,d}$, and n i.i.d. samples of the spiked cumulant model equation (3) as the alternative hypothesis $\mathbb{P}_{n,d}$. In section 3.1 we will first consider the problem from a statistical point of view, assuming to have *unbounded computational power* and no restrictions on the distinguishing algorithms. Then, in section 3.2 we will use the *low-degree method* to understand how the picture changes when we restrict to algorithms whose running time is at most polynomial in the space dimension d .

3.1. Statistical distinguishability: LR analysis

Recall the notion of *strong asymptotic distinguishability*: two sequences of probability measures are strongly distinguishable if it is possible to design statistical tests that can classify correctly which of the two distributions a sample was drawn from with probabilities of type I and II errors that converge to 0 (see section B.2 for the precise definition). Using this definition of distinguishability, we will ask what is the *statistical sample complexity exponent*, i.e. the minimum θ such that in the high-dimensional limit $d \rightarrow \infty$, if $n \asymp d^\theta$, then $\mathbb{P}_{n,d}$ and $\mathbb{Q}_{n,d}$ are strongly distinguishable (with no constraints on the complexity of statistical tests).

A useful quantity to consider is the LR of probability measures, which is defined as

$$L(x) := \frac{d\mathbb{P}}{d\mathbb{Q}}(x). \quad (4)$$

Computing the LR norm $\|L\|^2 := \mathbb{E}_{\mathbb{Q}}[L^2]$ is an excellent tool to probe for distinguishability: if $(\mathbb{P}_{n,d})$ and $(\mathbb{Q}_{n,d})$ are strongly distinguishable, then $\|L_{n,d}\| \rightarrow \infty$ (this is the well-known *second moment method for distinguishability*, see proposition 6 in the appendix for the precise statement). In the following we will apply this method, finding a formula for the LR norm and then study its limit as a function of θ . Here and throughout, we will denote the data matrix by $\underline{x} = (x^\mu)_{1,\dots,n}^\top$; in general matrices of size $n \times d$ will be denoted with underlined letters; see section B.1 for a complete summary of our notation. We will use Hermite polynomials, denoted by $(h_k)_k$, see appendix B.3 for details. We assume that the spike u is drawn from a known prior $\mathcal{P}(u)$, and that the scalar latent variables $(g^\mu)_{\mu=1,\dots,n}$ are drawn i.i.d. from a distribution p_g with the following properties:

Assumption 1 (assumption on latent variables g^μ). We assume that the one-dimensional probability distribution of the non-Gaussian random variable $p_g(g)$ is an even distribution, $p_g(g = dx) = p_g(-g = -dx)$, with mean 0 and variance 1, and that it satisfies the following requirement on the growth of its Hermite coefficients:

$$\mathbb{E}[h_m(g)] \leq \Lambda^m m! \quad (5)$$

where $\Lambda > 0$ is a positive constant that does not depend on m . Finally, we assume that p_g has tails that cannot go to 0 slower than a standard Gaussian, $\mathbb{E}[\exp(g^2/2)] < +\infty$.

A detailed discussion of these assumptions can be found in appendix B.4, as well as a proof that they are satisfied by a wide class of distributions including all the compactly supported distributions with mean 0 and variance 1 (some concrete examples are $p_g = \text{Rademacher}(1/2)$ or $p_g = \text{Unif}(-\sqrt{3}, \sqrt{3})$).

Under these assumption we can compute the following formula for the LR norm.

Theorem 2. Suppose that u has i.i.d. Rademacher prior and that the non-Gaussian distribution p_g satisfies assumption 1. Then the norm of the total LR is given by

$$\|L_{n,d}\|^2 = \sum_{j=0}^d \binom{d}{j} \frac{1}{2^d} f\left(\beta, \frac{2j}{d} - 1\right)^n, \quad (6)$$

where f is defined as the following average over two independent replicas $g_u, g_v \sim g$ of g :

$$f(\beta, \lambda) := \mathbb{E}_{g_u, g_v} \left[\frac{1 + \beta}{\sqrt{(1 + \beta)^2 - \beta^2 \lambda^2}} e^{-\frac{(1+\beta)\left((1+\beta)(g_u^2 + g_v^2) - 2\beta(g_u g_v)\lambda\right) + \frac{g_u^2 + g_v^2}{2}}{2(1+\beta)^2 - 2\beta^2 \lambda^2}} \right]. \quad (7)$$

We prove theorem 2 in appendix B.5. The theorem has key consequences in two directions:

- on the one hand, equation (6) implies that the LR norm is bounded for $\theta < 1$ (see lemma 13 in appendix B.5.3), which confirms that below that threshold *strong distinguishability is impossible*;
- on the other hand, equation (6) implies that whenever there exists $\tilde{\lambda}$ such that $f(\beta, \tilde{\lambda}) > 1$, we find that $\|L_{n,d}\| \geq \frac{f(\beta, \tilde{\lambda})^n}{2^d}$. Thus, the LR norm diverges as soon as n

Learning from higher-order statistics, efficiently: hypothesis tests, random features, and neural networks

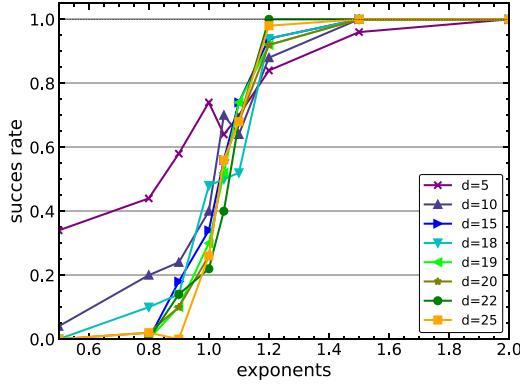


Figure 1. The performance of an exhaustive-search algorithm corroborates the presence of a phase transition for $\theta = 1$, as suggested by theorem 2. Success rate of an exponential-time search algorithm over all the possible spikes in the d -hypercube as a function of the exponent θ that quantifies as $n = d^\theta$ the samples used in the log-likelihood test (8), in the $g \sim \text{Radem}(1/2)$ case.

grows as $n \asymp d^\theta$ with any $\theta > 1$. In appendix B.5.3 we detail as an example the case in which $g \sim \text{Rademacher}(1/2)$ where the norm $\|L_{n,d}\|$ even diverges for $\theta = 1$ and $d \asymp \gamma n$ for some $\gamma > 0$.

So, besides the intrinsic value of providing an exact formula for the LR of the spiked cumulant model, theorem 2 implies the presence of a phase transition at $\theta = 1$ for the *strong statistical distinguishability*.

A complementary approach that substantiates the presence of the statistical phase transition at $\theta = 1$ can be seen in figure 1, where we perform a maximum log-likelihood test along $u \cdot x^\mu$ for *all* the possible spikes u in the d -dimensional hypercube using the formula for the LR conditioned on the spike,

$$\sum_{\mu=1}^n \log \left(\frac{p_x(x^\mu | u)}{p_z(x^\mu)} \right) = \sum_{\mu=1}^n \log \mathbb{E}_g \sqrt{1 + \beta} \exp \left(-\frac{1 + \beta}{2} \left(g - \sqrt{\frac{\beta}{(1 + \beta)d}} x^\mu \cdot u \right)^2 + \frac{g^2}{2} \right), \quad (8)$$

(see equation (61) in section B.5.2 for the derivation of this equation) and output the most likely u . Note that due to the exponential complexity in d of the algorithm, it is unfeasible to reach large values for this parameter. However, even at small d values, the success rate for retrieving the correct spike has a very steep increase at around $\theta = 1$, as predicted by our analysis of the LR norm in theorem 2.

3.2. Computational distinguishability: LDLR analysis

We now compare the statistical threshold of section 3.1 with the *computational sample complexity exponent* that quantifies the sample complexity of detecting non-Gaussianity with an efficient algorithm that runs in a time that is polynomial in the input dimension. The algorithmic sample complexity can be analysed rigorously for a wide class of

algorithms using the *low-degree method* [5–8, 37]. The low-degree method arose from the study of the sum-of-squares hierarchy [5] and rose to prominence when Hopkins & Steurer [7] demonstrated that the method can capture the Kesten–Stigum threshold for community detection in the stochastic block model [68–70]. In the case of hypothesis testing, the key quantity is the *low-degree likelihood ratio* (LDLR) [8, 37].

Definition 3 (LDLR). Let $D \geq 0$ be an integer. The LDLR of degree D is defined as

$$L^{\leq D} := \mathcal{P}^{\leq D} L \quad (9)$$

where $\mathcal{P}^{\leq D}$ projects L onto the space of polynomials of degree up to D , parallel to the Hilbert space structure defined by the scalar product $\langle f, g \rangle_{L^2(\Omega, \mathbb{Q})} := \mathbb{E}_{x \sim \mathbb{Q}}[f(x)g(x)]$.

The idea of this method is that among degree- D polynomials, $L^{\leq D}$ captures optimally the difference between $\mathbb{P}$ and $\mathbb{Q}$, and this difference can be quantified by the norm $\|L^{\leq D}\| = \|L^{\leq D}\|_{L^2(\Omega_n, \mathbb{Q}_n)}$. Hence in analogy to the *second moment method* used in section 3.1, we can expect low-degree polynomials to be able to distinguish $\mathbb{P}$ from $\mathbb{Q}$ only when $\|L_n^{\leq D(n)}\| \xrightarrow{n} \infty$, where $D(n)$ is a monotone sequence diverging with n . Indeed, the following (informal) conjecture from Hopkins [8] states that this criterion is valid not only for polynomial tests, but for all polynomial-time algorithms:

Conjecture 4. For two sequences of measures $\mathbb{Q}_N, \mathbb{P}_N$ indexed by N , suppose that (i) $\mathbb{Q}_N$ is a product measure; (ii) $\mathbb{P}_N$ is *sufficiently* symmetric with respect to permutations of its coordinates; and (iii) $\mathbb{P}_N$ is *robust* with respect to perturbations with small amount of random noise. If $\|L_N^{\leq D}\| = O(1)$ as $N \rightarrow \infty$ and for $D \geq (\log N)^{1+\varepsilon}$, for some $\varepsilon > 0$, then there is no polynomial-time algorithm that strongly distinguishes the distributions $\mathbb{Q}$ and $\mathbb{P}$.

Even though this conjecture is still not proved in general, its empirical confirmation on many benchmark problems has made it an important tool to probe questions of computational complexity, see theorem 8 for a simple example.

We will now compute LDLR norm estimates for the spiked cumulant model, so that the application of conjecture 4 will help to understand the *computational sample complexity* of this model.

Theorem 5 (LDLR for spiked cumulant model). Suppose that $(u_i)_{i=1,\dots,d}$ are drawn i.i.d. from the symmetric Rademacher distribution and that the non-Gaussian distribution p_g satisfies assumption 1. Let $0 < \varepsilon < 1$ and assume $D(n) \asymp \log^{1+\varepsilon}(n)$. Take $n, d \rightarrow \infty$, with the scaling $n \asymp d^\theta$ for $\theta > 0$. The following bounds hold:

$$\left\| L_{n,d}^{\leq D(n)} \right\|^2 \geq \left(\frac{1}{\lfloor D(n)/4 \rfloor} \left(\frac{\beta^2 \kappa_4^g}{(1+\beta)^2} \right)^2 \frac{n}{d^2} \right)^{\lfloor D(n)/4 \rfloor} \quad (10)$$

$$\left\| L_{n,d}^{\leq D(n)} \right\|^2 \leq 1 + \sum_{m=1}^{D(n)} \left(\frac{\Lambda^2 \beta}{1+\beta} \right)^m m^{4m} \left(\frac{n}{d^2} \right)^{m/4}. \quad (11)$$

Taken together, equations (10) and (11) imply the presence of a critical regime for $\theta_c = 2$, and describe the behaviour of $\|L_{n,d}^{\leq D}\|$ for all $\theta \neq \theta_c$

$$\lim_{n,d \rightarrow \infty} \left\| L_{n,d}^{\leq D(n)} \right\| = \begin{cases} 1 & 0 < \theta < 2 \\ +\infty & \theta > 2. \end{cases} \quad (12)$$

This theorem could be derived with different constants from lemma 26 and proposition 8 in Dudeja & Hsu [12]. Here we also provide a different, more direct argument. We sketch the proof of theorem 5 in section B.6.1 and give the complete proof in section B.6.2. We will discuss next the implications of the results presented in sections 3.1 and 3.2

3.3. Statistical-to-computational gaps in the spiked cumulant model

Put together, our results for the statistical and computational sample complexities of detecting non-Gaussianity in theorem 2 and theorem 5 suggest the existence of three different regimes in the spiked cumulant model as we vary the exponent θ , with a statistical-to-computational gap: for $0 \leq \theta < 1$, the problem is *statistically impossible* in the sense that no algorithm is able to strongly distinguish $\mathbb{P}$ and $\mathbb{Q}$ with so few samples, since the LR norm is bounded. For $1 < \theta < 2$, the norm of the LR with $g \sim \text{Rademacher}(1/2)$ diverges (even at $\theta = 1$ for some values of β), the problem could be statistically solvable (as validated by the results of exhaustive-search algorithms in figure 1), but conjecture 4 suggests no polynomial-time algorithm is able to achieve distinguishability in this regime; this is the so-called *hard phase*. If $\theta > 2$, the problem is solvable in polynomial time by evaluating a polynomial function (fourth-order in each sample) and thresholding; this is the *easy phase*.

The spiked cumulant model leads thus to intrinsically harder classification problems than the spiked Wishart model, where the critical regime is at $\theta = 1$. The proof of theorem 5 reveals that this increased difficulty is a direct consequence of the whitening of the inputs in equation (3). Without whitening, degree-2 polynomials would also give contributions to the LDLR (75) which would yield linear sample complexity. The difference in sample complexity of the spiked Wishart and spiked cumulant models mirrors the gap between the sample complexity of the best-known algorithms for matrix factorisation, which require linear sample complexity, and tensor PCA for rank-1 spiked tensors of order k [25, 28, 30], where sophisticated spectral algorithms can match the computational lower bound of $d^{k/2}$ samples.

4. Learning from HOCs with neural networks and RFs

The analysis of the (low-degree) LR has given us a detailed picture of the statistical and computational complexity of extracting information from the covariance or the HOCs of data. We will now use these fundamental limits to benchmark the sample complexity of two-layer neural networks (2LNN) trained with stochastic gradient descent on a binary discrimination task, where inputs in one class are drawn from the normal distribution $\mathcal{N}(0, \mathbf{1}_d)$, while inputs in the other class are drawn from the spiked Wishart or the

spiked cumulant model. In addition, we will also benchmark random feature methods (RF) [71–73] as a finite-dimensional approximation of kernel methods [71–73].

In a nutshell, the idea behind our experiments is to first *validate* both 2LNN and RF on the simpler spiked Wishart task and then to apply both methods to the spiked cumulant model, where inputs are generated in a way that mirrors the spiked Wishart: comparing equation (1) with equations (2) and (3), we see that the only differences are the whitening, and the latent distribution p_g . In our experiments, we choose the latent variables to be standard Gaussian for spiked Wishart, and Rademacher for the spiked cumulant—hence the latent variables have matching first and second moments. However, we will see that the spiked cumulant model exhibits a large gap in the sample complexity required for neural networks or RFs to learn the problem. We relegate details on the experimental setups such as hyper-parameters to appendix A.

4.1. Spiked Wishart model

We trained **two-layer ReLU neural networks** $\phi_\theta(x) = v^\top \max(0, w^T x)$ with $m = 5d$ hidden neurons on the spiked Wishart classification task. We show the early-stopping test accuracy of the networks in the linear and quadratic regimes with $n_{\text{class}} \asymp d, d^2$ samples per class in figures 2(A) and (B), resp. Neural networks are able to solve this task, in the sense that their classification error *per sample* is below 0.5, implying strong distinguishability of the whole data matrix. Indeed, some of the hidden neurons converge to the spike u of the covariance matrix, as can be seen from the maximum value of the normalised overlaps $\max_k w_k^\top u / \sqrt{\|w_k\| \|u\|}$ in figures 2(C) and (D), where w_k is the weight vector of the k th hidden neuron. In the linear regime (C), there is an increasingly clear transition from a random overlap for small data sets to a large overlap at large data sets as we increase the input dimension d ; in the quadratic regime (D), the neural network recovers the spike almost perfectly.

The **relatively large overlap between hidden neurons and spike at small sample complexities** (figures 2(C) and (D)) is due to the fact that we plot the maximum overlaps over a relative large number of hidden neurons $m = 5d$; hence even at initialisation, a few neurons will have large overlaps. We verified that ensuring an initial overlap of only $1/\sqrt{d}$ by explicit orthogonalisation did not change our results on distinguishability, see figure 5. A possible explanation is that the dynamics of the wide network is dominated by the majority of neurons, which do not have a macroscopic overlap.

Meanwhile, we found that **the performance of RFs** tends to random guessing at linear sample complexity, where we let the input dimension $d \rightarrow \infty$ with m/d and n_{class}/d fixed, while at quadratic sample complexity, RFs learn the task, although they perform worse than neural networks. The failure of RF in the linear regime makes sense in light of recent results that suggest that RFs in this scaling regime are limited to learning a linear approximation of the target function [74–77], while the LDLR analysis section B.2.3 shows that the target function, i.e. the LDLR, is a quadratic polynomial. However, these results are, to the best of our knowledge, restricted to the case of Gaussian isotropic inputs.

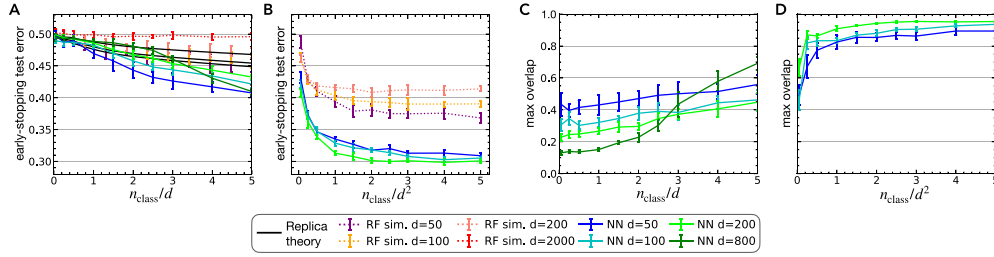


Figure 2. Learning the spiked Wishart task with neural networks and random features. (A) and (B) Test accuracy of random features (RF) and early-stopping test accuracy of two-layer ReLU neural networks (NN) on the spiked Wishart task, equation (1), with linear and quadratic sample complexity ($n_{\text{class}} \propto d$, d^2 , respectively, where d is the input dimension). Predictions for the performance of random features obtained using replicas are shown in black. (C) and (D) Maximum normalised overlaps of the networks' first-layer weights with the spike u , equation (1). Parameters: $\beta = 5$. Neural nets and random features have $m = 5d$ hidden neurons. Full experimental details in appendix A.

To ensure that the performance of RFs does indeed tend to random guessing, we performed a **replica analysis** following Loureiro [78] for mixture classification tasks together with the Gaussian equivalence theorem [79–83] (black lines in figure 2(A), details in appendix C). Replica theory perfectly predicts the performance of RF we obtain in numerical experiments (red-ish dots) for various values of d at linear sample complexity. We thus find a clear separation in the sample complexity required by RFs ($n_{\text{class}} \gtrsim d^2$) and neural networks ($n_{\text{class}} \gtrsim d$) to learn the spiked Wishart task. The replica analysis can be extended to the polynomial regime by a simple rescaling of the free energy [84] on several data models, like the vanilla teacher–student setup [85], the Hidden manifold model [79], and the vanilla Gaussian mixture classification (see figure 8). However, we found that for the spiked Wishart model, the Gaussian equivalence theorem which we need to deal with the RF distribution *fails* at quadratic sample complexity. This might be due to the fact that in this case, the spike induces a dominant direction in the covariance of the RFs, and this direction is also key to solving the task, which can lead to a breakdown of Gaussian equivalence [83]. A similar breakdown of the Gaussian equivalence theorem at quadratic sample complexity has been demonstrated recently in teacher–student setups by Cui *et al* [86] and Camilli *et al* [87].

4.2. Spiked cumulant

Having thus validated the performance of 2LNN and RF on the simpler spiked Wishart model, we turn to the spiked cumulant model and to the question: can neural networks learn higher-order correlations, efficiently? The LDLR analysis predicts that a polynomial-time algorithm requires at least a quadratic number of samples to detect non-Gaussianity, and hence to solve the classification task, and we found indeed that neural networks require at least quadratic sample complexity to solve the task, figures 3(A) and (B). The high values of the maximum overlap between hidden neurons and cumulant spike in the linear regime (compared to $d^{-1/2}$) are again a consequence of

Learning from higher-order statistics, efficiently: hypothesis tests, random features, and neural networks

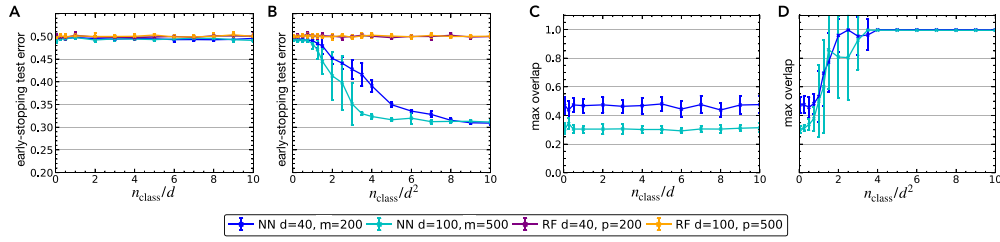


Figure 3. Learning the spiked cumulant task with neural networks and random features. (A), (B) Test accuracy of random features (RF) and early-stopping test accuracy of two-layer ReLU neural networks (NN) on the spiked cumulant task equation (3) with linear and quadratic sample complexity ($n_{\text{class}} \asymp d, d^2$, respectively, where d is the input dimension). (C) and (D) Maximum normalised overlaps of the networks' first-layer weights with the spike u , (3). *Parameters:* $\beta = 10$. Neural nets and random features have $m = 5d$ hidden neurons, same optimisation as in figure 2. Full experimental details in appendix A.

choosing the maximum overlap among $m = 5d$ hidden neurons. RFs cannot solve this task even at quadratic sample complexity, since they are limited to a quadratic approximation of the target function [75–77, 88], but we know from the LDLR analysis that the target function is a fourth-order polynomial. We thus find an even larger separation in the minimal number of samples required for RFs and neural networks to solve tasks that depend on directions which are encoded exclusively in the HOCs of the inputs.

4.3. Phase transitions and neural network performance in a simple model for images

We finally show another example of a separation between the performance of RFs and neural networks in the feature-learning regime on a toy model for images that was introduced recently by Ingrosso & Goldt [13], the non-linear Gaussian process (NLGP). The idea is to generate inputs that are (i) translation-invariant and that (ii) have sharp edges, both of which are hallmarks of natural images [89]. We first sample a vector $z \in \mathbb{R}^d$ from a normal distribution with zero mean and covariance $C_{ij} = \mathbb{E} z_i z_j = \exp(-|i - j|/\xi)$ to ensure translation-invariance of the inputs, with length scale $\xi > 0$. We then introduce edges, i.e. sharp changes in luminosity, by passing z through a saturating non-linearity like the error function, $x_i = \text{erf}(gz_i)/Z(g)$, where $Z(g)$ is a normalisation factor that ensures that the pixel-wise variance $\mathbb{E} x_i^2 = 1$ for all values of the gain $g > 0$. The classification task is to discriminate these ‘images’ from Gaussian inputs with the same mean and covariance, as illustrated in two dimensions in figure 4(A). This task is different from the spiked cumulant model in that the cumulant of the NLGP is not low-rank, so there are many directions that carry a signal about the non-Gaussianity of the inputs.

We trained wide 2LNNs on this task and interpolated between the feature-learning and the ‘lazy’ regimes using the α -renormalisation trick of Chizat *et al* [90]. As we increase α , the networks go from feature-learners ($\alpha = 1$) to an effective RF model and

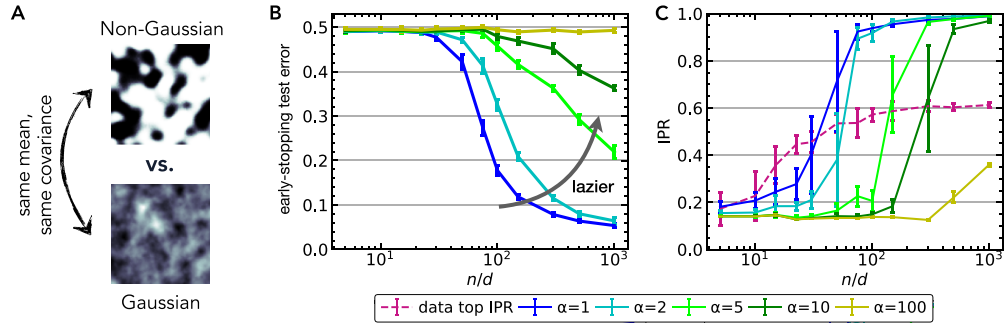


Figure 4. A phase transition in the fourth-order cumulant precedes learning from the fourth cumulant. (A) We train neural networks to discriminate inputs sampled from a simple non-Gaussian model for images introduced by Ingrosso & Goldt [13] (top) from Gaussians with the same mean and covariance (bottom). (B) Test error of two-layer neural networks interpolating between the fully-trained ($\alpha=1$) and lazy regimes (large α) – see section 4.3. (C) The localisation of the leading CP-factor of the non-Gaussian inputs (dashed purple line) and the first-layer weights of the trained networks, as measured by the inverse participation ratio (IPR), equation (13). Large IPR denotes a more localised vector w . Parameters: $g=3$, $\xi=1$, $d=20$, $m=100$. Full details in appendix A.

require an increasing amount of data to solve the task, figure 4(B). There appears to be a sharp transition from random guessing to non-trivial performance as we increase the number of training samples for all values of α . This transition is preceded by a transition in the behaviour of the HOC that was reported by Ingrosso & Goldt [13]. They showed numerically that the CP-factors of the empirical fourth-order cumulant T , defined as the vectors $\hat{u} \in \mathbb{R}^d$ that give the best rank- r approximation $\hat{T} = \sum_{k=1}^r \gamma_k \hat{u}_k^{\otimes 4}$ of T [91], localise in space if the data set from which the empirical cumulant is calculated is large enough. Quantifying the localisation of a weight vector w using the *inverse participation ratio*

$$\text{IPR}(w) = \frac{\sum_{i=1}^d w_i^4}{\left(\sum_{i=1}^d w_i^2\right)^2}, \quad (13)$$

we confirm that the leading CP-factors of the fourth-order cumulant localise (purple dashed line in figure 4(C)). The localisation of the CP-factors occurs with slightly less samples than the best-performing neural network requires to learn ($\alpha=1$). The weights of the neural networks also localise at a sample complexity that is slightly below the sample complexity for solving the task. The laziest network ($\alpha=100$), i.e. the one where the first-layer weights move least and which is hence closest to RFs, does not learn the task even with a training set containing $n=10^3d$ samples when $d=20$, indicating again a large advantage of feature-learners over methods with fixed feature maps, such as RFs.

5. Concluding perspectives

Neural networks crucially rely on the higher-order correlations of their inputs to extract statistical patterns that help them solve their tasks. Here, we have studied the difficulty of learning from higher-order correlations in the spiked cumulant model, where the first non-trivial information in the data set is carried by the input cumulants of order 4 and higher. Our LR analysis of the corresponding hypothesis test confirmed that data sampled from the spiked cumulant model could be statistically distinguishable (in the sense that it passes the *second moment method for distinguishability*) from isotropic Gaussian inputs at linear sample complexity, while the number of samples required to strongly distinguish the two distributions in polynomial time scales as $n \gtrsim d^2$ for the class of algorithms covered by the low-degree conjecture [5–8], suggesting the existence of a large statistical-to-computational gap in this problem. Our experiments with neural networks show that they learn from HOCs efficiently in the sense that they match the sample complexities predicted by the analysis of the hypothesis test, which is in stark contrast to RFs, which require a lot more data. In the future, a key challenge will be extend this framework to null hypotheses that go beyond isotropic Gaussian distributions. It will be intriguing to analyse the *dynamics* of neural networks on spiked cumulant models or the NLGP to understand how neural networks extract information from the HOCs of realistic data sets efficiently [92].

Acknowledgments

We thank Zhou Fan, Yue Lu, Antoine Maillard, Alessandro Pacco, Subhabrata Sen, Gabriele Sicuro, and Ludovic Stephan for stimulating discussions on various aspects of this work. SG acknowledges co-funding from Next Generation EU, in the context of the National Recovery and Resilience Plan, Investment PE1 – Project FAIR ‘Future Artificial Intelligence Research’, and from the European Union—NextGenerationEU, in the framework of the PRIN Project SELFMADE (code 2022E3WYTY—CUP G53D23000780001).

Contributions

ES performed the numerical experiments with neural networks and random features. LB performed the (low-degree) likelihood analysis. FG performed the replica analysis of random features. SG designed research and advised ES and LB. All authors contributed to writing the paper.

Appendix A. Experimental details

A.1. Figures 2 and 3

For the spiked Wishart and spiked cumulant tasks, we trained two-layer ReLU neural networks $\phi_\theta(x) = v^\top \max(0, w^T x)$. The number of hidden neurons $m = 5d$, where d is the input dimension. We train the networks using SGD with a learning rate of 0.002 and

Learning from higher-order statistics, efficiently: hypothesis tests, random features, and neural networks

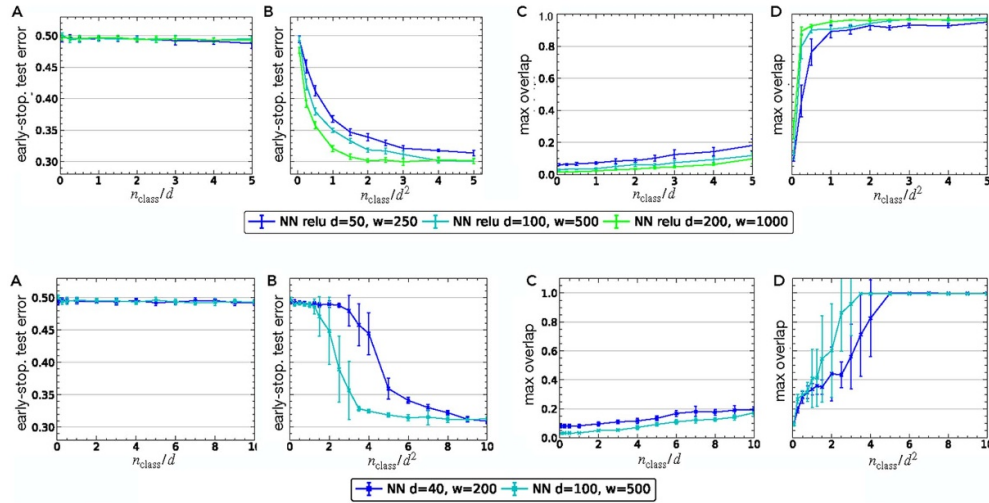


Figure 5. Learning the spiked Wishart and spiked cumulant task, starting from small initial overlaps. We repeat the neural network experiments on the spiked Wishart (top) and spiked cumulant (bottom) task, see figures 2 and 3, while enforcing that all hidden neurons have an overlap of exactly $1/\sqrt{d}$ with the spikes, by simple explicit orthogonalisation. While the maximum overlaps do indeed decrease for small sample complexities, the qualitative behaviour is unchanged. All hyperparameters as in figures 2 and 3, respectively.

a weight-decay of 0.002 for 50 epochs for the spiked Wishart task and for 200 epochs for the spiked cumulant task. The plots show the early-stopping test errors. The results are plotted as averages over 10 random seeds, showing the standard deviation by the errorbars.

The RF models [72, 73] have also a width of $5d$. The ridge regression is performed using scikit-learn [93] with a regularisation of 0.1.

For the spiked datasets, the spikes are from the Rademacher distribution, using a signal-to-noise ratio of 5.0 for the spiked Wishart and 10.0 for the spiked cumulant datasets. For the overlaps between spikes and features overlaps, we plot the highest overlap amongst the incoming weights of the hidden neurons with the spike, calculated as the normalised dot product.

A.1.1. Starting from small initial overlap Since the neural networks have a large number of neurons ($m = 5d$), some of them will have a relatively large overlap with the spike at initialisation, as can be seen in figure 3. To check that this relatively large overlap did not affect our results, we repeated the same set of experiments while enforcing an overlap of all hidden neuron weights with the spike of $1/\sqrt{d}$ by explicit orthogonalisation, as discussed in section 4.1. We show the results in figure 5: while the maximum overlaps do indeed decrease for small sample complexities, the qualitative behaviour is unchanged.

A.2. Figure 4

For the NLGP–GP task, we use the α -scaling trick of Chizat *et al* [90] to interpolate between feature- and lazy-learning. We define the network function as:

$$\phi_{\text{NN}}(x; v, W, v_0, W_0) = \frac{\alpha_{\text{NTK}}}{K} \left[\sum_j^m v_j \sigma \left(\sum_i^d w_i x_i \right) - \sum_j^m v_{0,j} \sigma \left(\sum_i^d w_{0,i} x_i \right) \right] \quad (14)$$

where v_0, w_0 are kept fixed at their initial values. The mean-squared loss is also rescaled by $(1/\alpha_{\text{NTK}}^2)$. Changing α_{NTK} from low to high allows to interpolate between the feature- and the lazy-learning limits as the first-layer weights will move away less from their initial values.

For figure 4, the network has a width of 100 and the optimisation is run by SGD for 200 epochs with a weight-decay of $5 \cdot 10^{-6}$ and learning rate of 0.5. The one-dimensional data vectors have a length of 20; the correlation length is 1 and the signal-to-noise ratio is set to 3. The error shown is early-stopping test error. The localisation of the neural networks' features and the data's fourth moments is shown by the IPR measure. (We used here a lower length for the data vectors so that the calculations for fourth-order cumulants do not have high memory requirements.) For the neural networks, the highest IPR is shown amongst the incoming features of the hidden neurons at the final state of the training. For the data, the highest IPR is the highest amongst the CP-factors of the fourth cumulants of the nlgp-class using a rank-1 PARAFAC decomposition from the tensorly package [94].

Appendix B. Mathematical details on the hypothesis testing problems

In this section, we provide more technical details on Hermite polynomials, LR and LDLR, and present the complete proofs of all theorems stated in the main.

B.1. Notation

We use the convention $0 \in \mathbb{N}$. $n \in \mathbb{N}$ denotes the number of samples, d the dimensionality of each data point x . The letters k, m usually denote free natural parameters. Let $[n] := \{1, \dots, n\}$, $\mu \in [n]$ is an index that ranges on the samples and $i \in [d]$ usually ranges on the data dimension. Underlined letters $\underline{x}, \underline{\mathbb{P}}, \underline{\mathbb{Q}}$ are used for objects that are $n \times d$ dimensional. In proofs the letter C denotes numerical constants whose precise value may change from line to line, and any dependency is denoted as a subscript. $m!$ denotes factorial and $m!!$ denotes the double factorial (product of all the numbers $k \leq m$ with the same parity as m).

B.1.1. Multi-index notation We will need multi-index notation to deal with d -dimensional functions and polynomials. Bold greek letters are used to denote multi-indices and their components, $\boldsymbol{\alpha} = (\alpha_1, \dots, \alpha_d) \in \mathbb{N}^d$. The following conventions are

adopted:

$$|\alpha| := \sum_{i=1}^d \alpha_i, \quad \alpha! := \prod_{i=1}^d \alpha_i!, \quad x^\alpha := \prod_{i=1}^d x_i^{\alpha_i}$$

$$\partial_\alpha f(x_1, \dots, x_d) = \frac{\partial}{\partial x^\alpha} f := \left(\frac{\partial}{\partial x_1^{\alpha_1}} \cdots \frac{\partial}{\partial x_d^{\alpha_d}} \right) f. \quad (15)$$

Since we will have n samples of d -dimensional variables, we will need to consider polynomials and functions in nd variables. To deal with all these variables we introduce multi-indices (denoted by underlined bold Greek letters $\underline{\alpha}, \underline{\beta}, \dots$). They are $n \times d$ matrices with entries in $\mathbb{N}$ (i.e. elements of $\mathbb{N}^{n \times d}$)

$$\underline{\alpha} := (\alpha_i^\mu) = \begin{pmatrix} \alpha_1^1 & \cdots & \alpha_d^1 \\ \vdots & \ddots & \vdots \\ \alpha_1^n & \cdots & \alpha_d^n \end{pmatrix} \quad (16)$$

We denote by α^μ the rows of $\underline{\alpha}$, that are d -dimensional multi-indices.

All the notations (15) generalize to multi-multi-indices in the following way:

$$|\underline{\alpha}| := \sum_{\mu=1}^n |\alpha^\mu| = \sum_{\mu=1}^n \sum_{i=1}^d \alpha_i^\mu, \quad \underline{\alpha}! := \prod_{\mu} \alpha^\mu! = \prod_{\mu} \prod_i \alpha_i^{\mu!},$$

$$\underline{x}^{\underline{\alpha}} := \prod_{\mu=1}^n (x^\mu)^{\alpha^\mu} = \prod_{\mu=1}^n \prod_{i=1}^d (x_i^\mu)^{\alpha_i^\mu} \quad (17)$$

$$\partial_{\underline{\alpha}} f(\underline{x}) := \left(\frac{\partial}{\partial (x^1)^{\alpha^1}} \cdots \frac{\partial}{\partial (x^n)^{\alpha^n}} \right) f(\underline{x}).$$

B.2. More on LR and LDLR

B.2.1. Statistical and computational distinguishability From a statistical point of view, distinguishing two sequences of probability measures $\underline{\mathbb{P}} = (\mathbb{P}_n)_{n \in \mathbb{N}}$ and $\underline{\mathbb{Q}} = (\mathbb{Q}_n)_{n \in \mathbb{N}}$ defined on a sequence of measure spaces $(\Omega_n, \mathcal{F}_n)$ means finding a sequence of tests $f_n: \Omega_n \rightarrow \{0, 1\}$, which are measurable functions that indicate whether a given set of inputs was sampled from $\mathbb{P}$ or $\mathbb{Q}$. We will say that the two measures are **statistically distinguishable** in the strong sense if there exists a statistical test f for which

$$\mathbb{Q}_n(f_n(x) = 0) \xrightarrow{n \rightarrow \infty} 1 \text{ and } \mathbb{P}_n(f_n(x) = 1) \xrightarrow{n \rightarrow \infty} 1. \quad (18)$$

The strong version of statistical distinguishability requires that the probability of success of the statistical test must tend to 1 as $n \rightarrow \infty$, whereas **weak distinguishability** just requires it to be asymptotically greater than $1/2 + \varepsilon$, for any $\varepsilon > 0$. We will call the minimal number of samples required to achieve strong statistical distinguishability the **statistical sample complexity** of the problem. We can obtain a computationally bounded analogue of this definition by restricting the complexity of the statistical test

f_n to functions that are computable in a time that is polynomial in the input dimension d . The **computational statistical complexity** is then the minimal number of samples required to achieve **computational distinguishability**.

B.2.2. Necessary conditions for distinguishability A necessary condition for strong distinguishability is based on the LR of probability measures, which is defined as

$$L_n(x) := \frac{d\mathbb{P}_n}{d\mathbb{Q}_n}(x). \quad (19)$$

The LR provides a necessary condition for strong distinguishability of $\mathbb{P}$ and $\mathbb{Q}$ via the so-called second moment method:

Proposition 6 (second moment method for distinguishability). *Suppose that $\mathbb{P}_n$ is absolutely continuous with respect to $\mathbb{Q}_n$, and let L_n be the corresponding LR. A necessary condition for strong distinguishability of $\mathbb{P}$ from $\mathbb{Q}$ is*

$$\left\| L_n \right\|^2 := \mathbb{E}_{x \sim \mathbb{Q}_n} [L_n(x)^2] \xrightarrow{n \rightarrow \infty} +\infty \quad (20)$$

where $\|\cdot\|$ is the norm with respect to the Hilbert space

$$L^2(\Omega_n, \mathbb{Q}_n) = \{f : \Omega_n \rightarrow \mathbb{R} \mid \mathbb{E}_{\mathbb{Q}_n} [f^2(x)] < \infty\}. \quad (21)$$

Proof. The proof is immediate: if $(\|L_n\|)_n$ were bounded, by Cauchy-Schwartz $\mathbb{Q}(A_n) \rightarrow 0$ would imply $\mathbb{P}(A_n) \rightarrow 0$. But by the definition of strong asymptotic distinguishability there exists a sequence of events $(A_n)_n$ such that $\mathbb{P}(A_n) \rightarrow 1$ and $\mathbb{Q}(A_n) \rightarrow 0$; hence $(\|L_n\|)_n$ must diverge. $\square$

The second moment method has been used to derive statistical thresholds for various high-dimensional inference problems [22, 27, 37, 95]. Note that proposition 6 only states a necessary condition; while it is possible to construct counterexamples, they are usually based on rather artificial constructions (like the following example 1 from [37]), and the second moment method is considered a good proxy for strong statistical distinguishability.

Example 7. Let $\mathbb{P}$ and $\mathbb{Q}$ be two strongly distinguishable sequences of probability measures on $(\mathcal{Y}_n, \mathcal{F}_n)_n$. Define $\mathbb{P}'_n$ on $\mathcal{F}_n$ as a measure that with probability 1/2 follows $\mathbb{P}_n$ and with probability 1/2 follows $\mathbb{Q}_n$. Letting L'_n be the LR of $\mathbb{P}'_n$ with respect to $\mathbb{Q}_n$. We have that $\|L'_n\| \rightarrow \infty$, but $\mathbb{P}'$ and $\mathbb{Q}$ are not strongly distinguishable.

Proof. The density of $\mathbb{P}'_n$ is:

$$p'_n(y) = \frac{p_n(y) + q_n(y)}{2} \quad (22)$$

hence

$$L'_n = \frac{1}{2}(1 + L_n) \quad (23)$$

which clearly has the same asymptotic behaviour as L_n .

On the other hand, it cannot be that $\underline{\mathbb{P}}'$ and $\underline{\mathbb{Q}}'$ are strongly distinguishable since for any event A , $\mathbb{P}'_n(A) \geq \frac{1}{2}\mathbb{Q}_n(A)$. $\square$

B.2.3. LDLR analysis for spiked Wishart model In this subsection we show the application of LDLR method to the spiked Wishart model. We obtain the correct BBP threshold for the signal-to-noise ratio even when we restrict ourselves to the class of polynomials of constant degree:

Theorem 8 (LDLR for spiked Wishart model). *Suppose the prior on u belongs to the following cases:*

- $(u_i)_{i=1,\dots,d}$ are i.i.d. and symmetric Rademacher random variables
- $(u_i)_{i=1,\dots,d}$ are i.i.d. and $u_i \sim \mathcal{N}(0, 1)$.

Let $D \in \mathbb{N}$ and $d, n \rightarrow \infty$, with fixed ratio $\gamma := d/n$, then

$$\lim_{d,n \rightarrow \infty} \|L^{\leq D}\| = \sum_{k=0}^D \frac{(2k-1)!!}{(2k)!!} \frac{\beta^{2k}}{\gamma^k}, \quad (24)$$

which, as D increases, stays bounded for $\beta < \beta_c := \sqrt{\gamma} = \sqrt{d/n}$ and diverges for $\beta > \beta_c$.

The distinguishability threshold that we have recovered here is of course the famous BBP phase transition in the spiked Wishart model [9]: if $\beta < \beta_c = \sqrt{d/n}$, the low-degree LR stays bounded and indeed a set of inputs drawn from **1** is statistically indistinguishable from a set of inputs drawn from $\mathcal{N}(0, \mathbb{1})$. For $\beta > \beta_c$ instead, there is a phase transition for the largest eigenvalue of the empirical covariance of the inputs in the second class, which can be used to differentiate the two classes, and the LDLR diverges.

As mentioned in the main text, a thorough LDLR analysis of a more general model that encompasses the spiked Wishart model was performed by Bandeira *et al* [96]. Their theorem 3.2 is a more general version of our theorem 8, that generalizes it in two directions: it works for a class of subgaussian priors on u that satisfy some concentration property, and also allows for negative SNR (the requirement is $\beta > -1$). For completeness, here we show that this result can be obtained as a straightforward application of a Gaussian additive model.

Proof. We note that the problem belongs the class of *Additive Gaussian Noise models* which is studied in depth by Kunisky *et al* [37]. In those models the two hypotheses have to be expressed as an additive perturbation of white noise:

- $\mathbb{P}_n$: $\underline{x}_n = y_n + z_n$,
- $\mathbb{Q}_n$: $\underline{x}_n = z_n$.

The spiked Wishart model that we consider belongs to this class. It can be seen by defining $\mathbb{R}^{nd} \ni y_n = \left(\sqrt{\frac{\beta}{d}} g^1 u, \dots, \sqrt{\frac{\beta}{d}} g^n u \right)^\top$. So we can apply theorem 2.6 from Kunisky *et al* [37], that computes the norm of the LDLR by using two independent replicas of the variable y . Denoting the two replicas by $\hat{y}$ and $\tilde{y}$, we get

$$\begin{aligned} \|L_n^{\leq D}\|^2 &= \mathbb{E} \left[\sum_{m=0}^D \frac{1}{m!} (\hat{y} \cdot \tilde{y})^m \right] \\ &= \mathbb{E} \left[\sum_{m=0}^D \frac{\beta^m}{m! d^m} \left(\sum_{\mu=1}^n \hat{g}^\mu \tilde{g}^\mu \hat{u} \cdot \tilde{u} \right)^m \right] \\ &= \sum_{m=0}^D \frac{\beta^m}{m! d^m} \mathbb{E} \left[(\hat{u} \cdot \tilde{u})^m \left(\sum_{\mu=1}^n \hat{g}^\mu \tilde{g}^\mu \right)^m \right] \\ &= \sum_{m=0}^D \frac{\beta^m}{m! d^m} \mathbb{E} \left[\left(\sum_{\mu=1}^n \hat{g}^\mu \tilde{g}^\mu \right)^m \right] \mathbb{E} [(\hat{u} \cdot \tilde{u})^m]. \end{aligned} \quad (25)$$

Note now that $\sum_{\mu=1}^n \hat{g}^\mu \tilde{g}^\mu$ has the same distribution as $-\sum_{\mu=1}^n \hat{g}^\mu \tilde{g}^\mu$, so its distribution is even and all the odd moments are 0. This means that we can reduce to the case $m = 2k$, so we need to study:

$$\|L_n^{\leq D}\|^2 = \sum_{k=0}^{\lfloor D/2 \rfloor} \frac{\beta^{2k}}{(2k)! d^{2k}} \underbrace{\mathbb{E} \left[\left(\sum_{\mu=1}^n \hat{g}^\mu \tilde{g}^\mu \right)^{2k} \right]}_{T_1} \underbrace{\mathbb{E} [(\hat{u} \cdot \tilde{u})^{2k}]}_{T_2}. \quad (26)$$

Let us consider the term T_1 . Call $Y^\mu := \hat{g}^\mu \tilde{g}^\mu$. The distribution of each of the Y^μ is not Gaussian, but we have that $\mathbb{E}[Y^\mu] = 0$ and $\text{Var}(Y^\mu) = 1$. So by the central limit theorem $S_n := \frac{1}{\sqrt{n}} \sum_{\mu=1}^n Y^\mu$ converges to a standard normal in distribution, as $n \rightarrow \infty$. Note that the cumulants of S_n can be computed thanks to linearity and additivity of cumulants and are $\kappa_{2k}^{(S_n)} = n^{1-k} \kappa_{2k}^Y$. So they all go to 0 except from the variance. Since the moments can be written as a function of the cumulants up to that order, it follows that $\lim_n \mathbb{E}[S_n^{2k}]$ will be the $2k$ th moment of the standard normal distribution, which means that we have the following:

$$\lim_{n \rightarrow +\infty} \mathbb{E} \left[\left(\frac{1}{\sqrt{n}} \sum_{\mu=1}^n \hat{g}^\mu \tilde{g}^\mu \right)^{2k} \right] = (2k-1)!! \quad (27)$$

We turn now to T_2 and do the same reasoning. Define $v_i := \hat{u}_i \tilde{u}_i$, both in Rademacher and in Gaussian prior case, we have that the $(v_i)_{i=1, \dots, d}$ is an independent family of random variables that have 0 mean and variance equal to 1. So we can again apply the

central limit theorem to get that:

$$\lim_{d \rightarrow +\infty} \mathbb{E} \left[\left(\frac{\hat{u} \cdot \tilde{u}}{\sqrt{d}} \right)^{2k} \right] = (2k-1)!! \quad (28)$$

Taking the limit for $n, d \rightarrow +\infty$ with the constraint $\gamma = d/n$, we have that:

$$\begin{aligned} \lim_{n, d \rightarrow \infty} \|L_n^{\leq D}\|^2 &= \lim_{n, d \rightarrow \infty} \sum_{k=0}^{\lfloor D/2 \rfloor} \frac{\beta^{2k} n^k}{(2k)! d^k} \mathbb{E} \left[\left(\frac{1}{\sqrt{n}} \sum_{\mu=1}^n \hat{g}^\mu \tilde{g}^\mu \right)^{2k} \right] \mathbb{E} \left[\left(\frac{\hat{u} \cdot \tilde{u}}{\sqrt{d}} \right)^{2k} \right] \\ &= \sum_{k=0}^D \frac{(\beta^k (2k-1)!!)^2}{(2k)! \gamma^k} = \sum_{k=0}^D \frac{(2k-1)!!}{(2k)!} \frac{\beta^{2k}}{\gamma^k} \end{aligned} \quad (29)$$

which is what we wanted to prove.

As a final note we remark that the basic ideas of the arguments in [96] coincide with what exposed above. However, the increased generality of the statement requires the use of the abstract theory of Umbral calculus to generalize the notion of Hermite polynomials to negative SNR cases, as well as more technical work to achieve the bounds on the LDLR projections. $\square$

B.3. Hermite Polynomials

We recall here the definitions and key properties of the Hermite polynomials.

Definition 9. The Hermite polynomial of degree m is defined as

$$h_m(x) := (-1)^m e^{\frac{x^2}{2}} \frac{d^m}{dx^m} \left(e^{-\frac{x^2}{2}} \right). \quad (30)$$

Here is a list of the first 5 Hermite polynomials:

$$\begin{aligned} h_0(x) &= 1 \\ h_1(x) &= x \\ h_2(x) &= x^2 - 1 \\ h_3(x) &= x^3 - 3x \\ h_4(x) &= x^4 - 6x^2 + 3. \end{aligned} \quad (31)$$

The Hermite polynomials enjoy the following properties that we will use in the subsequent proofs (for details see section 5.4 in McCullagh [97], Szegő [98] and Abramowitz & Stegun [99]):

- they are orthogonal with respect to the L^2 product weighted with the density of the Normal distribution:

$$\frac{1}{\sqrt{2\pi}} \int_{-\infty}^{\infty} h_n(x) h_m(x) e^{-\frac{x^2}{2}} dx = n! \delta_{m,n}; \quad (32)$$

- h_m is a monic polynomial of degree m , hence $(h_m)_{m \in \{1, \dots, N\}}$ generates the space of polynomials of degree $\leq N$;
- the previous two properties imply that the family of Hermite polynomials is an orthogonal basis for the Hilbert space $L^2(\mathbb{R}, \mathbb{Q})$ where $\mathbb{Q}$ is the normal distribution;
- they enjoy the following recurring relationship

$$h_{m+1}(x) = xh_m(x) - h'_m(x), \quad (33)$$

which can also be expressed as a relationship between coefficients as follows. If $h_m(x) = \sum_{k=0}^m a_{m,k} x^k$, then

$$a_{m+1,k} = \begin{cases} -a_{m,1} & k=0, \\ a_{m,k-1} - (k+1)a_{m,k+1} & k>0. \end{cases} \quad (34)$$

They also satisfy identities of binomial type, like:

$$h_m(x+y) = \sum_{k=0}^m \binom{m}{k} x^{m-k} h_k(y) \quad (35)$$

$$h_m(\gamma x) = \sum_{j=0}^{\lfloor m/2 \rfloor} \gamma^{m-2j} (\gamma^2 - 1)^j \binom{m}{2j} (2j-1)!! h_{m-2j}(x). \quad (36)$$

Multivariate case. In the multivariate m -dimensional case we can consider Hermite tensors $(H_\alpha)_{\alpha \in \mathbb{N}^m}$ defined as:

$$H_\alpha(x_1, \dots, x_m) = \prod_{i=1}^m h_{\alpha_i}(x_i) \quad (37)$$

all the properties of the one-dimensional Hermite polynomials generalize to this case, in particular they form an orthogonal basis of $L^2(\mathbb{R}^m, \mathbb{Q})$ where $\mathbb{Q}$ is a multivariate normal distribution. If $\langle \cdot, \cdot \rangle$ is the inner product of that Hilbert space, we have that:

$$\langle H_\alpha, H_\beta \rangle = \alpha! \delta_{\alpha, \beta}. \quad (38)$$

Of course all this is valid in the case $m = nd$, where we can also use multi-multi-indices to get the following identity:

$$H_{\underline{\alpha}}(x_1^1, \dots, x_d^n) = \prod_{\mu=1}^n H_{\alpha^\mu}(x^\mu) = \prod_{\mu=1}^n \prod_{i=1}^d h_{\alpha_i^\mu}(x_i^\mu). \quad (39)$$

We are now ready to see the proof of 75.

Lemma 10. We consider a hypothesis testing problem in $\mathbb{R}^d$ where the null hypothesis $\mathbb{Q}$ is the multivariate Gaussian distribution $\mathcal{N}(0, \mathbb{1}_{d \times d})$, while the alternative hypothesis $\mathbb{P}$ is absolutely continuous with respect to it. Then:

$$L^{\leq D} = \sum_{|\alpha| \leq D} \frac{\langle L, H_\alpha \rangle H_\alpha}{\alpha!} = \sum_{|\alpha| \leq D} \frac{\mathbb{E}_{x \sim \mathbb{P}} [H_\alpha(x)] H_\alpha}{\alpha!} \quad (40)$$

which implies

$$\|L^{\leq D}\|^2 = \sum_{|\alpha| \leq D} \frac{\langle L, H_\alpha \rangle^2}{\alpha!} = \sum_{|\alpha| \leq D} \frac{\mathbb{E}_{x \sim \mathbb{P}} [H_\alpha(x)]^2}{\alpha!}. \quad (41)$$

Proof. First note that

$$\langle L, H_\alpha \rangle = \mathbb{E}_{x \sim \mathbb{P}} [H_\alpha(x)] \quad (42)$$

due to the definition of LR and a change of variable in the expectation. Then we can just use the fact that $(H_\alpha)_\alpha \in \mathbb{N}^d$ are an orthogonal basis for $L^2(\mathbb{R}^d, \mathbb{Q})$, and if we consider the Hermite polynomials up to degree D , they are also a basis of the space of polynomials in which we want to project L to get $L^{\leq D}$. Hence the formulas follow by just computing the projection using this basis. $\square$

Note that we set the lemma in $\mathbb{R}^d$, but of course it holds also in $\mathbb{R}^{nd}$ just switching to multi-multi-index notation.

B.3.1. Hermite coefficients Lemma 10 translates the problem of computing the norm of the LDLR to the problem of computing the projections $\langle L, H_\alpha \rangle$. Note that this quantity is equal to $\mathbb{E}_{\mathbb{P}}[H_\alpha(x)]$, which we will call α th *Hermite coefficient* of the distribution $\mathbb{P}$.

The following lemma from [96] provides a version of the integration by parts technique that is tailored for Hermite polynomials.

Lemma 11. Let $f : \mathbb{R}^d \rightarrow \mathbb{R}^d$ be a function that is continuously differentiable k times. Assume that f and all of its partial derivatives up to order k are bounded by $O(\exp(|y|^\lambda))$ for some $\lambda \in (0, 2)$, then for any $\alpha \in \mathbb{N}^d$ such that $|\alpha| \leq k$

$$\langle f, H_\alpha \rangle = \mathbb{E}_{y \sim \mathcal{N}(0, \mathbb{1})} [H_\alpha(y) f(y)] = \mathbb{E}_{y \sim \mathcal{N}(0, \mathbb{1})} [\partial_\alpha f(y)]. \quad (43)$$

Proof. 43 can be proved by doing induction on k using 33, see [96] for details. $\square$

B.3.2. Links with cumulant theory The cumulants $(\kappa_\alpha)_{\alpha \in \mathbb{N}^d}$ of a random variable $x \in \mathbb{R}^d$, can be defined as the Taylor coefficients of the expansion in ξ of the *cumulant generating function*:

$$K_x(\xi) := \log(\mathbb{E}[e^{\xi \cdot x}]). \quad (44)$$

Order-one cumulant is the mean, order-two is the variance, and higher-order cumulants encode more complex correlations among the variables. The Gaussian distribution is the only non constant distribution to have a polynomial cumulant generating function (as proved in theorem 7.3.5 in [100]), if $z \sim \mathcal{N}(\mu, \Sigma)$, then:

$$K_z(\xi) = \mu\xi + \frac{1}{2}\xi^\top \Sigma \xi.$$

Hence cumulants with order higher than three can also be seen as a measure of how much a distribution deviates from Gaussianity (i.e. how much it deviates from its best Gaussian approximation).

On this point the similarity with Hermite coefficients is evident. Indeed, up to order-five, on whitened distributions, cumulants and Hermite coefficients coincide. But from sixth-order onward, they start diverging. They are still linked by deterministic relationship, but its combinatorial complexity increases swiftly and it is not easy to translate formulas involving Hermite coefficients into cumulants and vice versa. For this reason, our low-degree analysis of the LR leading to 5 keeps the formalism that naturally arises from the computations, which is based on Hermite coefficients. A detailed discussion of the relation between Hermite polynomials and cumulants in the context of asymptotic expansions of distributions like the Gram–Charlier and Edgeworth series can be found in chapter 5 of McCullagh [97].

B.4. Details on the spiked cumulant model

Here we will expand on the mathematical details of the spiked cumulant model.

B.4.1. Prior distribution on the spike For the prior distribution on u , $\mathcal{P}$, its role is analogous to the spiked Wishart model, so all the choices commonly used for that model can be considered applicable to this case. Namely symmetric distributions so that $\| \frac{u}{\sqrt{d}} \| \approx 1$ as $d \rightarrow \infty$. In the following we will make the computations assuming u_i i.i.d. and with Rademacher prior:

$$u_i \sim \text{Rademacher}(1/2). \quad (45)$$

It helps to have i.i.d. components and constant norm $\| \frac{u}{\sqrt{d}} \| \equiv 1$. However all the results should hold, with more involved computations, also with the following priors:

- $(u_i)_{i=1,\dots,d}$ are i.i.d. and $u_i \sim \mathcal{N}(0, 1)$
- $u \in \text{Unif}(\partial B(0, \sqrt{d}))$.

B.4.2. Distribution of the non-Gaussianity As detailed in 1, we need the non-Gaussianity to satisfy specific conditions. Some of the requirements are fundamental and cannot be avoided, whereas others could likely be removed or modified with only technical repercussion that do not change the essence of the results.

The most vital assumptions are the first and the last: $\mathbb{E}[g] \neq 0$ would introduce signal in the first cumulant, changing completely the model. It is important to control the tails of the distribution with

$$\mathbb{E}[\exp(g^2/2)] < +\infty, \quad (46)$$

a fat-tailed g would make the LR-LDLR technique pointless due to the fact that $L \notin \mathcal{L}^2(\mathbb{R}^d, \mathbb{Q})$ and $\|L\| = \infty$ for any n, d . For example it is not possible to use Laplace distribution for g .

On the opposite side, the least important assumptions are that $\text{Var}(g) = 1$ and equation (5); removing the first would just change the formula for whitening matrix S , while equation (5) is a very weak requirement and it is needed just to estimate easily the Hermite coefficients' contribution and reach 11, it may be even possible to remove it and try to derive a similar estimate from the assumption on the tails 46.

Finally, the requirement of symmetry of the distribution has been chosen arbitrarily thinking about applications: the idea is that the magnitude of the fourth-order cumulant, the kurtosis, is sometimes used as a test for non-Gaussianity, so it is interesting to isolate the contribution given by the kurtosis, and higher-order, even degree, cumulants, by cancelling the contribution of odd-order cumulants. However, it would be interesting to extend the model in the case of a centered distribution with non-zero third-order cumulant. It is likely that the same techniques can be applied, but the final thresholds on θ may differ.

The following lemma gives a criterion of admissibility that ensures Radem(1/2) and Unif($-\sqrt{3}, \sqrt{3}$) (together with many other compactly supported distributions) satisfy 1.

Lemma 12. *Suppose p_g is a probability distribution, compactly supported in $[-\Lambda, \Lambda]$, $\Lambda \geq 1$, then*

$$\mathbb{E}_{g \sim p_g} [h_m(g)] \leq \Lambda^m m!.$$

Proof. Use the notation $h_m(x) = \sum_{k=0}^m a_{m,k} x^k$ and $S_m := \sum_{k=0}^m |a_{m,k}|$. Then we have that

$$\begin{aligned} \mathbb{E}_{g \sim p_g} [h_m(g)] &= \mathbb{E}_{g \sim p_g} \left[\sum_{k=0}^m a_{m,k} g^k \right] \\ &\leq \mathbb{E}_{g \sim p_g} \left[\sum_{k=0}^m |a_{m,k}| |g|^k \right] \\ &\leq \Lambda^m S_m \end{aligned} \quad (47)$$

So we just need to prove that $S_m \leq m!$, which can be done by induction using 34. Suppose it true for m , we prove it for $m+1$. By 34 we have that:

$$|a_{m+1,k}| \leq \begin{cases} |a_{m,1}| & k=0 \\ |a_{m,k-1}| + (k+1)|a_{m,k+1}| & k>0. \end{cases} \quad (48)$$

Summing on both sides (and using $a_{m,k} = 0$ when $k > m$), we get:

$$\begin{aligned}
 S_{m+1} &\leq |a_{m,1}| + \sum_{k=1}^{m+1} |a_{m,k-1}| + (k+1) |a_{m,k+1}| \\
 S_{m+1} &\leq S_m + \sum_{j=0}^m j |a_{m,j}| \\
 S_{m+1} &\leq S_m + m \sum_{j=0}^m |a_{m,j}| \\
 S_{m+1} &\leq (m+1) S_m.
 \end{aligned} \tag{49}$$

Hence by application of the inductive hypothesis we get $S_{m+1} \leq (m+1)!$, completing the proof. $\square$

B.4.3. Computing the whitening matrix In this paragraph all the expectations are made assuming u fixed and it will be best to work with its normalized version $\bar{u} = u/\sqrt{d}$. Note that $\mathbb{E}[x] = 0$ and we want also that

$$\begin{aligned}
 \mathbb{1}_{d \times d} &= \mathbb{E}[xx^\top] = \mathbb{E}\left[S\left(\sqrt{\beta}g\bar{u} + z\right)\left(\sqrt{\beta}g\bar{u}^\top + z^\top\right)S^\top\right] \\
 &= S\left(\mathbb{1} + \beta\bar{u}\bar{u}^\top\right)S^\top.
 \end{aligned} \tag{50}$$

Hence we need

$$S^2 = SS^\top = (\mathbb{1} + \beta\bar{u}\bar{u}^\top)^{-1} = \mathbb{1} - \frac{\beta}{1+\beta}\bar{u}\bar{u}^\top. \tag{51}$$

So we look for γ such that:

$$(\mathbb{1} + \gamma\bar{u}\bar{u}^\top)^2 = \mathbb{1} - \frac{\beta}{1+\beta}\bar{u}\bar{u}^\top. \tag{52}$$

By solving the second-order equation we get

$$\gamma_{\pm} = -\left(1 \pm \frac{1}{\sqrt{1+\beta}}\right). \tag{53}$$

We choose the solution with $-$, so that S is positive definite:

$$S = \mathbb{1} - \left(1 - \frac{1}{\sqrt{1+\beta}}\right)\bar{u}\bar{u}^\top = \mathbb{1} - \frac{\beta}{1+\beta+\sqrt{1+\beta}}\frac{uu^\top}{d}. \tag{54}$$

Hence we can compute also the explicit expression for x :

$$\begin{aligned}
 x &= z - \left(1 - \frac{1}{\sqrt{1+\beta}}\right) \bar{u}^\top z \bar{u} + \sqrt{\frac{\beta}{1+\beta}} g \bar{u} \\
 x &= \underbrace{z - \bar{u}^\top z \bar{u}}_{z_{\perp u}} + \left(\sqrt{\frac{1}{1+\beta}} \bar{u}^\top z + \underbrace{\sqrt{\frac{\beta}{1+\beta}} g}_{\eta} \right) \bar{u} \\
 &= z_{\perp u} + \left(\sqrt{1-\eta^2} \bar{u}^\top z + \eta g \right) \bar{u}.
 \end{aligned} \tag{55}$$

So x is standard Gaussian in the directions orthogonal to u , whereas in u direction, it is a weighted sum between g and z . Note that it is a quadratic interpolation: the sum of the square of the weights is 1.

B.5. Details of the LR analysis for spiked cumulant model

B.5.1. Proof sketch for theorem 2 Since the samples are independent, the total LR factorises as

$$L(\underline{y}) = \mathbb{E}_u \left[\prod_{\mu=1}^n l(y^\mu | u) \right], \tag{56}$$

where the sample-wise LR is

$$l(y|u) = \frac{p_x(y|u)}{p_z(y)} = \mathbb{E}_{g \sim p_g} \left[\sqrt{1+\beta} \exp \left(-\frac{1+\beta}{2} \left(g - \sqrt{\frac{\beta}{(1+\beta)d}} y \cdot u \right)^2 + \frac{g^2}{2} \right) \right]. \tag{57}$$

To compute the norm of the LR, we consider two independent replicas of the spike, u and v , to get

$$\|L_{n,d}\|^2 = \mathbb{E}_{\underline{y} \sim \underline{\mathbb{Q}}} \left[\mathbb{E}_u \left[\prod_{\mu=1}^n l(y^\mu | u) \right] \mathbb{E}_v \left[\prod_{\mu=1}^n l(y^\mu | v) \right] \right]. \tag{58}$$

The hard part now is to simplify the high-dimensional integrals in this expression, it can be done because the integrand is almost completely symmetric, and the only asymmetries lie on the subspace spanned by u and v

$$\|L_{n,d}\|^2 = \mathbb{E}_{u,v} \left[\mathbb{E}_{g_u, g_v} \left[\frac{1+\beta}{\sqrt{(1+\beta)^2 - \beta^2 \left(\frac{u \cdot v}{d}\right)^2}} e^{-\frac{(1+\beta) \left((1+\beta) (g_u^2 + g_v^2) - 2\beta (g_u g_v) \left(\frac{u \cdot v}{d}\right) \right) + \frac{g_u^2 + g_v^2}{2}}}{2(1+\beta)^2 - 2\beta^2 \left(\frac{u \cdot v}{d}\right)^2} + \frac{g_u^2 + g_v^2}{2} \right] \right]^n \right].$$

Note that the integrand depends on u and v only through $\frac{u \cdot v}{d} =: \lambda$. So, using that the prior on u, v is i.i.d. Rademacher, the outer expectation can be transformed in a one dimensional expectation over λ , leading to (6).

B.5.2. Proof of theorem 2 $\tilde{x}^\mu = \sqrt{\frac{\beta}{d}}g^\mu u + z^\mu$, to find the marginal density of x^μ we integrate over the possible values of g

$$p_{\tilde{x}}(y|u) = \mathbb{P}(\tilde{x}^\mu \in dy|u) = \mathbb{E}_g \left[p_z \left(y - \sqrt{\frac{\beta}{d}}gu \right) \right] \quad (59)$$

where p_z is the density of a standard normal d -dimensional variable $z \sim \mathcal{N}(0, \mathbb{1}_d)$. But we are interested in the density of the whitened variable x , $p_x(\cdot|u)$, which can be seen as the push forward of the density of $\tilde{x}$ with respect to the linear transformation S . So

$$p_x(y|u) = p_{\tilde{x}}(S^{-1}y|u) |\det S^{-1}|. \quad (60)$$

It is easy to see from 3 that $|\det S^{-1}| = \sqrt{1+\beta}$, so we can plug it into 59 and after expanding the computations we get to

$$p_x(y|u) = p_z(y) \mathbb{E}_g \left[\sqrt{1+\beta} \exp \left(-\frac{1+\beta}{2} \left(g - \sqrt{\frac{\beta}{(1+\beta)d}} y \cdot u \right)^2 + \frac{g^2}{2} \right) \right]. \quad (61)$$

So we have found the LR for a single sample, conditioned on the spike:

$$l(y|u) = \frac{p_x(y|u)}{p_z(y)} = \mathbb{E}_g \left[\sqrt{1+\beta} \exp \left(-\frac{1+\beta}{2} \left(g - \sqrt{\frac{\beta}{(1+\beta)d}} y \cdot u \right)^2 + \frac{g^2}{2} \right) \right]. \quad (62)$$

Note that conditioning on u the samples are independent, so we have the following formula:

$$L(y) = \mathbb{E}_u \left[\prod_{\mu=1}^n l(y^\mu|u) \right]. \quad (63)$$

So to compute the norm we consider two independent replicas of the spike, u and v , then we switch the order of integration to get

$$\begin{aligned} \|L_{n,d}\|^2 &= \mathbb{E}_{\underline{y} \sim \underline{\mathbb{Q}}} \left[\mathbb{E}_u \left[\prod_{\mu=1}^n l(y^\mu|u) \right] \mathbb{E}_v \left[\prod_{\mu=1}^n l(y^\mu|v) \right] \right] \\ &= \mathbb{E}_{u,v} \left[\mathbb{E}_{\underline{y} \sim \underline{\mathbb{Q}}} \left[\prod_{\mu=1}^n l(y^\mu|u) l(y^\mu|v) \right] \right] \\ &= \mathbb{E}_{u,v} \left[\mathbb{E}_{y \sim \mathbb{Q}} [l(y|u) l(y|v)]^n \right] \\ &= \mathbb{E}_{u,v} \left[\mathbb{E}_{y \sim \mathbb{Q}} \left(\mathbb{E}_{g_u} \left[\sqrt{1+\beta} \exp \left(-\frac{1+\beta}{2} \left(g_u - \sqrt{\frac{\beta}{(1+\beta)d}} y \cdot u \right)^2 + \frac{g_u^2}{2} \right) \right] \right. \right. \\ &\quad \cdot \left. \mathbb{E}_{g_v} \left[\sqrt{1+\beta} \exp \left(-\frac{1+\beta}{2} \left(g_v - \sqrt{\frac{\beta}{(1+\beta)d}} y \cdot v \right)^2 + \frac{g_v^2}{2} \right) \right] \right)^n \right]. \end{aligned} \quad (64)$$

Switching the integral over y inside the new integrals over g_u and g_v we can isolate the following integral over y :

$$I := \mathbb{E}_{y \sim \mathbb{Q}} \left[\exp \left(-\frac{1+\beta}{2} \left(g_u - \sqrt{\frac{\beta}{(1+\beta)d}} y \cdot u \right)^2 - \frac{1+\beta}{2} \left(g_v - \sqrt{\frac{\beta}{(1+\beta)d}} y \cdot v \right)^2 \right) \right]. \quad (65)$$

It can be computed by noting the subspace orthogonal to $\{u, v\}$, we just have the integral of a standard normal, and the remaining 2-dimensional integral can be computed explicitly. Since the Rademacher prior implies that $\|u\| = \|v\| = \sqrt{d}$, the result depends only on their overlap λ (i.e. $\lambda = \frac{u \cdot v}{d}$), leading to:

$$I = \frac{1}{\sqrt{(1+\beta)^2 - \beta^2 \lambda^2}} \exp \left(-\frac{1+\beta}{2(1+\beta)^2 - 2\beta^2 \lambda^2} ((1+\beta)(g_u^2 + g_v^2) - 2\beta(g_u g_v) \lambda) \right). \quad (66)$$

If we plug this formula inside 64 and rearrange the terms we get an expression that can be written in terms of the density of two bi-dimensional centered Gaussians

$$\begin{aligned} \|L_{n,d}\|^2 &= \mathbb{E}_\lambda \left[\mathbb{E}_{g_u, g_v} \left[\frac{1+\beta}{\sqrt{(1+\beta)^2 - \beta^2 \lambda^2}} e^{-\frac{(1+\beta)((1+\beta)(g_u^2 + g_v^2) - 2\beta(g_u g_v) \lambda)}{2(1+\beta)^2 - 2\beta^2 \lambda^2} + \frac{g_u^2 + g_v^2}{2}} \right]^n \right] \\ &= \mathbb{E}_\lambda \left[\mathbb{E}_{g_u, g_v} \left[\mathcal{N}((g_u, g_v); \Sigma) \mathcal{N}((g_u, g_v); \mathbb{1}_{2 \times 2})^{-1} \right]^n \right] \end{aligned} \quad (67)$$

where

$$\Sigma^{-1} = \frac{1+\beta}{(1+\beta)^2 - \beta^2 \lambda^2} \begin{pmatrix} 1+\beta & -\beta \lambda \\ -\beta \lambda & 1+\beta \end{pmatrix}, \quad \lambda = \frac{u \cdot v}{d}. \quad (68)$$

Now we turn to compute the expectation over λ . Note that since u, v are independent and their components are Rademacher distributed, product of Rademacher is still Rademacher, hence $u \cdot v$ is the sum of d independent Rademacher random variables. Using moreover that the Rademacher distribution is just a linear transformation of the Bernoulli, we can link the distribution of the overlap λ to a linear transformation of a binomial distribution: $\frac{d}{2}(\lambda + 1) \sim \text{Binom}(d, 1/2)$.

We therefore define the auxiliary function

$$f(\beta, \lambda) := \mathbb{E}_{g_u, g_v} \left[\frac{1+\beta}{\sqrt{(1+\beta)^2 - \beta^2 \lambda^2}} e^{-\frac{(1+\beta)((1+\beta)(g_u^2 + g_v^2) - 2\beta(g_u g_v) \lambda)}{2(1+\beta)^2 - 2\beta^2 \lambda^2} + \frac{g_u^2 + g_v^2}{2}} \right] \quad (69)$$

so that the LR norm can be rewritten as

$$\|L_{n,d}\|^2 = \sum_{j=0}^d \binom{d}{j} \frac{1}{2^d} f\left(\beta, \frac{2j}{d} - 1\right)^n \quad (70)$$

which is what we needed to prove.

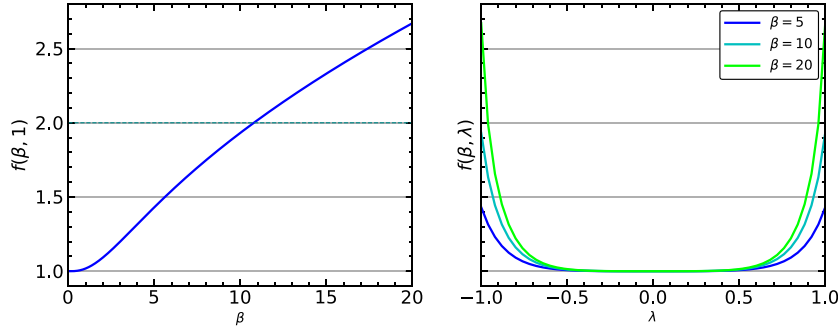


Figure 6. Graphs of f , defined in (7), when $g \sim \text{Rademacher}(1/2)$.

B.5.3. Consequences of theorem 2 The precise value of f depends on the choice of distribution for the non-Gaussianity g , however, it is possible to prove the following

Lemma 13. *If p_g satisfies assumptions 1 and $\theta < 1$, then there exists C such that:*

$$\|L_{n,d}\| \leq C \quad \forall n, d.$$

Proof. We first prove that, thanks to the sub-Gaussianity of g , we have that

$$f(\beta, \lambda) \leq 1 + C_\beta \lambda^2, \quad \forall \lambda \in [-1, 1]. \quad (71)$$

To do it, first note that $f(\beta, 0) = 1$, and $f(\beta, \cdot)$ is bounded on $[-1, 1]$ thanks to the sub-gaussianity of g :

$$\begin{aligned} |f(\beta, \lambda)| &\leq \mathbb{E}_{g_u, g_v} \left[e^{\frac{1}{2}(g_u^2 + g_v^2)} \right] \sup_{g_u, g_v} \left(\frac{1 + \beta}{\sqrt{(1 + \beta)^2 - \beta^2 \lambda^2}} e^{-\frac{(1 + \beta) \left((1 + \beta)(g_u^2 + g_v^2) - 2\beta(g_u g_v) \lambda \right)}{2(1 + \beta)^2 - 2\beta^2 \lambda^2}} \right) \\ &\leq C \frac{1 + \beta}{\sqrt{(1 + \beta)^2 - \beta^2 \lambda^2}} \\ &\leq C_1 (1 + C_2 \lambda^2). \end{aligned}$$

So we just need to prove that, up to changes on C_2 , we can take $C_1 = 1$. To do that it is sufficient to study the behaviour around $\lambda = 0$:

$$\frac{\partial}{\partial \lambda} f(\beta, 0) = \mathbb{E} \left[\frac{\beta}{1 + \beta} g_u g_v \right] = 0.$$

So there is a neighborhood of 0 such that $f(\beta, 0) = 1 + C\lambda^2 + o(\lambda^2)$. Hence we can take a suitable constant C_β so that equation (71) holds.

Now we can apply equation (71) to (6), to get

$$\begin{aligned} \|L_{n,d}\|^2 &\leq \sum_{j=0}^d \binom{d}{j} \frac{1}{2^d} \left(1 + C_\beta \left(\frac{2j}{d} - 1 \right)^2 \right)^n \\ &= \sum_{j=0}^d \binom{d}{j} \frac{1}{2^d} \sum_{k=0}^n \binom{n}{k} C_\beta^k \left(\frac{2j}{d} - 1 \right)^{2k} \\ &= \sum_{k=0}^n \binom{n}{k} \left(\frac{C_\beta}{d^2} \right)^k \mathbb{E} [Y_d^{2k}] \end{aligned}$$

where Y_d is a random variable distributed as the sum of d independent Rademacher with parameter $1/2$. So we can apply the central limit theorem on $\frac{Y_d}{\sqrt{d}}$ to get that $\frac{\mathbb{E}[Y_d^{2k}]}{d} \rightarrow (2k-1)!!$ hence we get:

$$\begin{aligned} \|L_{n,d}\|^2 &\leq C \sum_{k=0}^n \binom{n}{k} \left(\frac{C_\beta}{d} \right)^k (2k-1)!! \\ &\leq C_1 \sum_{k=0}^n \left(\frac{C_2 n}{d} \right)^k \end{aligned}$$

where C_1 and C_2 are constants independent of n, d . So now we can substitute that $n \asymp d^\theta$, and since $\theta < 1$ the series converges for all d , hence the LR norm is bounded. $\square$

We will analyze in more detail the case in which $g \sim \text{Rademacher}(1/2)$. It is a case that is particularly interesting for the point of view of the applications because it amounts to comparing a standard Gaussian with a Gaussian mixture with same mean and covariance. Moreover, with this choice of non-Gaussianity, the technical work simplifies because the troublesome integral over g_u, g_v in 7 becomes just a simple sum over 4 possibilities. In this case f can be computed exactly and it is displayed in 6. The maximum of $f(\beta, \cdot)$ is attained at $\lambda = \pm 1$ and $f(\cdot, 1)$ is monotonically increasing.

Assume that $n \asymp d^\theta$, then a sufficient condition for LR to diverge is $\frac{f(\beta, 1)^{d^\theta}}{2^d} \rightarrow \infty$, which holds as soon as $\theta > 1$. Moreover, even at linear sample complexity, it is possible to find regimes that ensure divergence of the LR norm, similar to BBP phase transition in spiked Wishart model. Assume that samples and dimensions scale at the same rate as in spiked Wishart model: $n \asymp \frac{d}{\gamma}$, then a sufficient condition for divergence is that

$$\frac{f(\beta, 1)^{d/\gamma}}{2^d} \rightarrow \infty$$

which holds if and only if $f(\beta, 1)^{1/\gamma} > 2$. Hence given β , you can always find

$$\gamma_\beta := \frac{\log(f(\beta, 1))}{\log 2}. \quad (72)$$

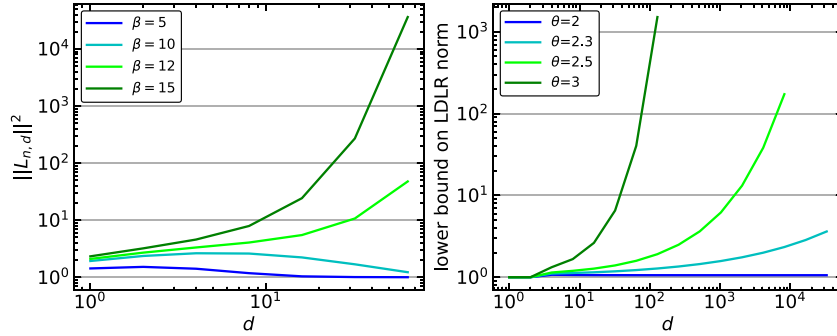


Figure 7. On the left, LR norm when $g \sim \text{Rademacher}(1/2)$ in the regime $n = \gamma d$ with $\gamma = 1$. When $\beta < \beta_\gamma \approx 10.7$ the LR remains bounded, whereas it goes to $+\infty$ for $\beta > \beta_\gamma$. On the right, the lower bound on $\|L_{n,d}^{\leq D(n)}\|$ given by (73) goes to $+\infty$ for $\theta > 2$. Parameters for the plot: $g \sim \text{Radem}(1/2)$, $\beta = 10$, $D(n) = \log^{3/2}(n)$

And for $\gamma > \gamma_\beta$ there is guarantee that $\|L_{n,d}\| \rightarrow \infty$. Vice-versa, we could also fix γ and define β_γ as only value of β that makes (72) true.

It is spontaneous to ask whether this condition for divergence of LR norm is also necessary, making the SNR threshold β_γ the analogous in this model of the threshold $\beta_c = \sqrt{\gamma}$ in spiked Wishart model (8). Although a rigorous proof of this is still missing, numerical evidence (7) suggests that for $\beta \leq \beta_\gamma$ the LR norm stays indeed bounded.

B.6. LDLR on spiked cumulant model

To discuss the proof of theorem 5, it is best to state it in two separate parts

Theorem 14 (LDLR for spiked cumulant model). Suppose that $(u_i)_{i=1,\dots,d}$ are drawn i.i.d. from the symmetric Rademacher distribution. If the non-Gaussian distribution p_g satisfies assumption 1, then the following lower and upper bounds hold:

- Let $D \in \mathbb{N}$ such that $D/4 \leq n$, then:

$$\left\| L_{n,d}^{\leq D} \right\|^2 \geq \sum_{m=0}^{\lfloor D/4 \rfloor} \binom{n}{m} \binom{d+1}{2}^m \left(\frac{\beta^2 \kappa_4^g}{\sqrt{4!} d^2 (1+\beta)^2} \right)^{2m} \quad (73)$$

where κ_4^g is the fourth-order cumulant of g .

- Conversely, for any D, n, d :

$$\left\| L_{n,d}^{\leq D} \right\|^2 \leq 1 + \sum_{m=1}^D \frac{C_m}{d^m} \sup_{k \leq m} \left(\mathbb{E}[h_k(g)]^{2m/k} \right) \binom{n}{\lfloor m/4 \rfloor} \binom{d}{\lfloor m/2 \rfloor} \quad (74)$$

where $C_m := \left(\frac{\beta}{1+\beta} \right)^m \binom{\lfloor m/4 \rfloor \lfloor m/2 \rfloor + m - 1}{m}$.

Corollary 15 (asymptotics of LDLR bounds). Assume the hypotheses of theorem 5. Let $0 < \varepsilon < 1$ and assume $D(n) \asymp \log^{1+\varepsilon}(n)$. Take $n, d \rightarrow \infty$, with the scaling $n \asymp d^\theta$ for

$\theta > 0$. Estimate (73) implies that for n, d large enough, the following lower bound holds:

$$\left\| L_{n,d}^{\leq D(n)} \right\|^2 \geq \left(\frac{1}{\lfloor D(n)/4 \rfloor} \left(\frac{\beta^2 \kappa_4^g}{(1+\beta)^2} \right)^2 \frac{n}{d^2} \right)^{\lfloor D(n)/4 \rfloor}.$$

Conversely, (74) leads to:

$$\left\| L_{n,d}^{\leq D(n)} \right\|^2 \leq 1 + \sum_{m=1}^{D(n)} \left(\frac{\Lambda^2 \beta}{1+\beta} \right)^m m^{4m} \left(\frac{n}{d^2} \right)^{m/4}.$$

Taken together, equations (10) and (11) imply the presence of a critical regime for $\theta_c = 2$, and describe the behaviour of $\|L_{n,d}^{\leq D}\|$ for all $\theta \neq \theta_c$

$$\lim_{n,d \rightarrow \infty} \left\| L_{n,d}^{\leq D(n)} \right\| = \begin{cases} 1 & 0 < \theta < 2 \\ +\infty & \theta > 2. \end{cases}$$

In the following we will present the proofs of 14 and 15. We first give a sketch of the proof in B.6.1 before giving the detailed proof in B.6.2.

B.6.1. Proof sketch The starting point of the argument is the observation that since the null hypothesis is white Gaussian noise, the space $L^2(\mathbb{R}^{nd}, \langle \cdot, \cdot \rangle)$ has an explicit orthogonal basis in the set of multivariate Hermite polynomials (H_{α}). The multi-index $\alpha \in \mathbb{N}^{nd}$ denotes the degree of the Hermite polynomial which is computed for each entry of the data matrix; see B.3 for a detailed explanation. We can then expand the LDLR norm in this basis to write:

$$\|L_{n,d}^{\leq D}\|^2 = \sum_{|\alpha| \leq D} \frac{\langle L, H_{\alpha} \rangle^2}{\alpha!} = \sum_{|\alpha| \leq D} \frac{1}{\alpha!} \mathbb{E}_{x \sim \mathbb{P}} [H_{\alpha}(x)]^2. \quad (75)$$

From now on the two bounds are treated separately.

Lower bound The idea is to bound the LDLR from below by computing the sum 75 only over a restricted set of addends $\mathcal{I}_m$ which we can compute explicitly. In particular, we consider the set $\mathcal{I}_m$ of all the polynomials with degree $4m$ in the data matrix, which are at most of order 4 in each individual sample x^{μ} . Then we use that the expectation of such Hermite polynomials conditioned on u can be split into m integrals of degree-4 Hermite polynomials in d variables. In this way we can use our knowledge of the fourth-order cumulant of x to compute those expectations (see (93)), so we find that

$$\|L_{n,d}^{\leq D}\|^2 \geq \sum_m \sum_{|\alpha| \in \mathcal{I}_m} \frac{1}{\alpha!} \mathbb{E}_{x \sim \mathbb{P}} [H_{\alpha}(x)]^2 \geq \sum_m \sum_{\alpha \in \mathcal{I}_m} \left(\frac{\beta^2 \kappa_4^g}{\sqrt{4!} d^2 (1+\beta)^2} \right)^{2m}. \quad (76)$$

Thanks to this manipulation, the bound does not depend on α explicitly, so we can complete the bound by estimating the cardinality of $\mathcal{I}_m$. Thanks to the fact that we

have n i.i.d. copies of variable x at disposal, the class of polynomials $\mathcal{I}_m$ has size that grows at least as $\binom{n}{m} \binom{d+1}{2}^m$, allowing to reach the lower bound in the statement.

Upper bound We can show that the expectation of Hermite polynomials for a single sample x^μ over the distribution $\mathbb{P}(\cdot|u)$ can be written as $\mathbb{E}_{x \sim \mathbb{P}(\cdot|u)}[H_\alpha(x)] = T_{|\alpha|,g}/d^{|\alpha|/2} u^\alpha$, where $T_{|\alpha|,g} = \left(\frac{\beta}{1+\beta}\right)^{|\alpha|/2} \mathbb{E}[h_{|\alpha|}(g)]$, cf lemma 16. Using the fact that inputs are sampled i.i.d. from $\mathbb{P}(\cdot|u)$ and substituting this result into (75), we find that

$$\|L_{n,d}^{\leq D}\|^2 = \sum_{|\underline{\alpha}| \leq D} \frac{\left(\prod_{\mu=1}^n T_{|\alpha^\mu|,g}\right)^2}{\underline{\alpha}! d^{|\underline{\alpha}|}} \mathbb{E}_{u \sim \mathcal{P}(u)}[u^{\underline{\alpha}}]^2. \quad (77)$$

To obtain an upper bound, we will first show that many addends in this sum are equal to zero, then estimate the remainder.

On the one hand, since we chose a Rademacher prior over the elements of the spike u_i , the expectation over the prior yields either 1 or 0 depending on whether the exponent of each u_i is even or odd. On the other hand, the whitening of the data means that $T_{|\alpha|,g} = 0$ for $0 < |\alpha| < 4$ (as proved in lemma 16).

For each $m \in \mathbb{N}$, we can denote $\mathcal{A}_m$ the set of multi-indices $|\underline{\alpha}| = m$ that give non-zero contributions in 77, so that we can write:

$$\|L_{n,d}^{\leq D}\|^2 = \sum_{m=0}^D \sum_{\underline{\alpha} \in \mathcal{A}_m} \frac{\left(\prod_{\mu=1}^n T_{|\alpha^\mu|,g}\right)^2}{\underline{\alpha}! d^m}. \quad (78)$$

Now we proved that we can bound the inner terms so that they depend on $\underline{\alpha}$ only through the norm $|\underline{\alpha}| = m$ (cf (82) and (110)):

$$\begin{aligned} \|L_{n,d}^{\leq D}\|^2 &\leq \sum_{m=0}^D \sum_{\underline{\alpha} \in \mathcal{A}_m} \frac{1}{d^m} \left(\frac{\beta}{1+\beta}\right)^m \sup_{k \leq m} \left(\mathbb{E}[h_k(g)]^{2m/k}\right) \\ &= \sum_{m=0}^D \frac{1}{d^m} \left(\frac{\beta}{1+\beta}\right)^m \sup_{k \leq m} \left(\mathbb{E}[h_k(g)]^{2m/k}\right) \#\mathcal{A}_m. \end{aligned} \quad (79)$$

Finally, the cardinality of $\mathcal{A}_m$ is 1 in case $m=0$ (which leads to the addend 1 in (74)) and if $m>0$ it can be bounded by (for details see (111) in the appendix)

$$\binom{\lfloor m/4 \rfloor \lfloor m/2 \rfloor + m - 1}{m} \binom{n}{\lfloor m/4 \rfloor} \binom{d}{\lfloor m/2 \rfloor}, \quad (80)$$

leading to (74).

B.6.2. Detailed proof First we prove a lemma that provides formulas and estimates for projections of the sample-wise LR $l(\cdot|u)$ on the Hermite polynomials.

Lemma 16. Let $x = S\left(\sqrt{\beta/d}gu + z\right)$ be a spiked cumulant random variable. Then for any $\alpha \in \mathbb{N}^d$, with $|\alpha| = m$, we have that:

$$\langle l(\cdot|u), H_\alpha \rangle = \mathbb{E}_{x \sim \mathbb{P}(\cdot|u)} [H_\alpha(x)] = \frac{T_{m,g}}{d^{m/2}} u^\alpha \quad (81)$$

where $T_{m,g}$ is a coefficient defined as:

$$T_{m,g} = \left(\frac{\beta}{1+\beta}\right)^{m/2} \mathbb{E}[h_m(g)]. \quad (82)$$

Proof. Recall that by (57)

$$l(y|u) = \mathbb{E}_g \left[\sqrt{1+\beta} \exp \left(-\frac{1}{2} \left(\sqrt{1+\beta}g - \sqrt{\frac{\beta}{d}}y \cdot u \right)^2 + \frac{g^2}{2} \right) \right]. \quad (83)$$

Note that this expression, thanks to (46) that bounds the integral, is differentiable infinitely many times in the y variable. It can be quickly proven by induction on $|\alpha| = m$ (using the recursive definition of Hermite polynomials (33)) that:

$$\begin{aligned} \partial_\alpha l(y|u) &= \left(\frac{\beta}{d}\right)^{m/2} u^\alpha \\ &\cdot \mathbb{E}_g \left[\sqrt{1+\beta} h_m \left(\sqrt{1+\beta}g - \sqrt{\frac{\beta}{d}}y \cdot u \right) \exp \left(-\frac{1}{2} \left(\sqrt{1+\beta}g - \sqrt{\frac{\beta}{d}}y \cdot u \right)^2 + \frac{g^2}{2} \right) \right]. \end{aligned} \quad (84)$$

Hence we can again use 46 together with the fact that

$$\sup_g \left| h_m \left(\sqrt{1+\beta}g - \sqrt{\frac{\beta}{d}}y \cdot u \right) \exp \left(-\frac{1}{2} \left(\sqrt{1+\beta}g - \sqrt{\frac{\beta}{d}}y \cdot u \right)^2 \right) \right| < \infty$$

to deduce that the hypothesis of lemma 11 are met for any $\alpha \in \mathbb{N}^d$, leading to:

$$\langle l(\cdot|u), H_\alpha \rangle = \mathbb{E}_{x \sim \mathbb{P}(\cdot|u)} [H_\alpha(x)] = \mathbb{E}_{y \sim \mathbb{Q}} [\partial_\alpha l(y|u)]. \quad (85)$$

This, and (84) already prove (81). Now to compute the exact value of $T_{m,g}$ we need to take the expectation with respect to $y \sim \mathbb{Q} = \mathcal{N}(0, \mathbb{1}_{d \times d})$. Note that by the choice of Rademacher prior on u , we know that $\|u\| = \sqrt{d}$. So, conditioning on u , $y \cdot \frac{u}{\sqrt{d}} \sim \mathcal{N}(0, 1)$.

Hence switching the expectations in (84), we get:

$$\begin{aligned}
 T_{m,g} &= \beta^{m/2} \mathbb{E}_g \left[\int_{-\infty}^{\infty} \frac{dz}{\sqrt{2\pi}} \sqrt{1+\beta} h_m \left(\sqrt{1+\beta}g - \sqrt{\beta}z \right) e^{-\frac{1}{2} \left(\sqrt{1+\beta}g - \sqrt{\beta}z \right)^2 + \frac{g^2 - z^2}{2}} \right] \\
 &= \beta^{m/2} \mathbb{E}_g \left[\int_{-\infty}^{\infty} \frac{dz}{\sqrt{2\pi}} \sqrt{1+\beta} h_m \left(\sqrt{1+\beta}g - \sqrt{\beta}z \right) \exp \left(-\frac{1}{2} \left(\sqrt{\beta}g - \sqrt{1+\beta}z \right)^2 \right) \right] \\
 &\stackrel{\tilde{z}=\sqrt{1+\beta}z}{=} \beta^{m/2} \mathbb{E}_g \left[\int_{-\infty}^{\infty} \frac{d\tilde{z}}{\sqrt{2\pi}} h_m \left(\sqrt{1+\beta}g - \sqrt{\frac{\beta}{1+\beta}}\tilde{z} \right) \exp \left(-\frac{1}{2} \left(\sqrt{\beta}g - \tilde{z} \right)^2 \right) \right].
 \end{aligned} \tag{86}$$

Now we use (35) on

$$\begin{aligned}
 x &= \frac{\beta}{\sqrt{1+\beta}}g - \sqrt{\frac{\beta}{1+\beta}}\tilde{z} \\
 y &= \sqrt{1+\beta}g - \frac{\beta}{\sqrt{1+\beta}}g = \frac{g}{\sqrt{1+\beta}}
 \end{aligned} \tag{87}$$

applying also the translation change of variable $\hat{z} = \tilde{z} - \sqrt{\beta}g$ we get:

$$\begin{aligned}
 &\mathbb{E}_{y \sim \mathbb{Q}} [\partial_{\alpha} l(y|u)] \\
 &= \left(\frac{\beta}{d} \right)^{m/2} u^{\alpha} \mathbb{E}_g \left[\sum_{k=0}^m \binom{m}{k} h_k \left(\frac{g}{\sqrt{1+\beta}} \right) \left(-\sqrt{\frac{\beta}{1+\beta}} \right)^{m-k} \mathbb{E}_{\hat{z} \sim \mathcal{N}(0,1)} [\hat{z}^{m-k}] \right].
 \end{aligned} \tag{88}$$

Now recall the formula for the moments of the standard Gaussian (see for instance section 3.9 in [97]):

$$\mathbb{E}_{\hat{z} \sim \mathcal{N}(0,1)} [\hat{z}^{m-k}] = \begin{cases} (m-k-1)!! & \text{if } m-k \text{ is even} \\ 0 & \text{if } m-k \text{ is odd} \end{cases} \tag{89}$$

Plugging this formula inside and changing the summation index $2j := m-k$ we get

$$T_{m,g} = \beta^{m/2} \sum_{j=0}^{\lfloor m/2 \rfloor} \binom{m}{2j} \left(\frac{\beta}{1+\beta} \right)^j (2j-1)!! \mathbb{E} \left[h_{m-2j} \left(\frac{g}{\sqrt{1+\beta}} \right) \right]. \tag{90}$$

Note that (36) with $x = \frac{g}{\sqrt{1+\beta}}$ and $\gamma = \sqrt{1+\beta}$ gives the following rewriting of $h_m(g)$:

$$h_m(g) = \sum_{j=0}^{\lfloor m/2 \rfloor} \left(\sqrt{1-\beta} \right)^{m-2j} (\beta)^j \binom{m}{2j} (2j-1)!! h_{m-2j} \left(\frac{g}{\sqrt{1+\beta}} \right) \tag{91}$$

which is almost the same expression as in (90). This allows simplify everything, leading to (82). $\square$

Note that lemma 16 together with 1 on g imply:

$$m = 2 \text{ or } m \text{ odd} \implies T_{m,g} = 0. \quad (92)$$

So, apart from $m=0$ which gives the trivial contribution $+1$, the first non zero contributions to the LDLR norm is at degree $m=4$:

$$\mathbb{E}_{x \sim \mathbb{P}(\cdot|u)} [H_{\alpha}(x)] = \frac{\beta^2}{(1+\beta)^2} \kappa_4^g \frac{u^{\alpha}}{d^2}. \quad (93)$$

From now on we will consider separately the lower and the upper bounds.

Lower bound The idea of the proof is to start from

$$\|L_{n,d}^{\leq D}\|^2 = \sum_{\substack{\alpha \in \mathbb{N}^{nd} \\ |\alpha| \leq D}} \frac{\langle L, H_{\alpha} \rangle^2}{\alpha!} \quad (94)$$

and to estimate it from below by picking only few terms in the sum. So we restrict to particular sets of multi-multi-indices for which we can exploit (93). Let $m \in \mathbb{N}$ such that $4m \leq D$, define

$$\mathcal{I}_m = \left\{ \alpha \in (\mathbb{N}^d)^n \mid |\alpha| = 4m, |\alpha^{\mu}| \in \{0, 4\} \ \forall 1 \leq \mu \leq n \right\}. \quad (95)$$

$\mathcal{I}_m$ is non empty since $m < n$ and for each $\alpha \in \mathcal{I}_m$ we can enumerate all the indices $\mu_1, \dots, \mu_m$ such that $\alpha^{\mu_i} \neq 0$.

Now we go on to compute the term $\langle L, H_{\alpha} \rangle$ for $\alpha \in \mathcal{I}_m$:

$$\begin{aligned} \langle L, H_{\alpha} \rangle &= \int_{\mathcal{U}} \mathbb{E}_{\underline{x} \sim \mathbb{P}(\cdot|u)^{\otimes n}} [H_{\alpha}(\underline{x})] d\mathcal{P}(u) \\ &= \int_{\mathcal{U}} \mathbb{E}_{\underline{x} \sim \mathbb{P}(\cdot|u)^{\otimes n}} \left[\prod_{i=1}^m H_{\alpha^{\mu_i}}(x^{\mu_i}) \right] d\mathcal{P}(u). \end{aligned} \quad (96)$$

Since the samples are independent conditionally on u , we can split the inner expectation along the m contributing directions:

$$\langle L, H_{\alpha} \rangle = \int_{\mathcal{U}} \prod_{i=1}^m \mathbb{E}_{x \sim \mathbb{P}(\cdot|u)} [H_{\alpha^{\mu_i}}(x^{\mu_i})] d\mathcal{P}(u). \quad (97)$$

Now we need to compute the inner d -dimensional expectation. For that we use that $|\alpha^{\mu_i}| = 4$, so, recalling the notation $\eta = \sqrt{\frac{\beta}{1+\beta}}$, we can apply (93) to get for each i

$$\mathbb{E}_{x \sim \mathbb{P}(\cdot|u)} [H_{\alpha^{\mu_i}}(x^{\mu_i})] = \frac{\eta^4}{d^2} \kappa_4^g u^{\alpha^{\mu_i}}. \quad (98)$$

So the resulting integral can be written in the following way:

$$\langle L, H_{\underline{\alpha}} \rangle = \left(\frac{\eta^4 \kappa_4^g}{d^2} \right)^m \int_{\mathcal{U}} \prod_{j=1}^d u_j^{\gamma_j} d\mathcal{P}(u) \quad (99)$$

where for each $j \in \{1, \dots, d\}$, $\gamma_j := \sum_{\mu} \alpha_j^{\mu}$. So we also have that $\sum_j \gamma_j = 4m$.

Now we take the expectation with respect to $\mathcal{P}(u)$, and use the fact that the components are i.i.d. Rademacher so the result depends on the parity of the γ_j in the following way:

$$\langle L, H_{\underline{\alpha}} \rangle = \begin{cases} \left(\frac{\eta^4 \kappa_4^g}{d^2} \right)^m & \text{if all } (\gamma_j)_{j=1, \dots, n} \text{ are even} \\ 0 & \text{if there is at least one } \gamma_j \text{ which is an odd number} \end{cases} \quad (100)$$

Hence, if we restrict to the set:

$$\tilde{\mathcal{I}}_m = \left\{ \underline{\alpha} \in \mathcal{I}_m \mid \forall j \in \{1, \dots, d\} \sum_{\mu=1}^n \alpha_j^{\mu} \text{ is even} \right\}. \quad (101)$$

We have that all the indices belonging to $\tilde{\mathcal{I}}_m$ give the same contribution:

$$\langle L, H_{\underline{\alpha}} \rangle^2 = \left(\frac{\eta^4}{d^2} \kappa_4^g \right)^{2m}. \quad (102)$$

Also, note that inside $\mathcal{I}_m$, $\alpha! \leq (4!)^m$, so get the following estimates

$$\begin{aligned} \|L_{n,d}^{\leq D}\|^2 &= \sum_{|\underline{\alpha}| \leq D} \frac{\langle L, H_{\underline{\alpha}} \rangle^2}{\underline{\alpha}!} \\ &\geq \sum_{m=0}^{\lfloor D/4 \rfloor} \sum_{\underline{\alpha} \in \tilde{\mathcal{I}}_m} \frac{\langle L, H_{\underline{\alpha}} \rangle^2}{\underline{\alpha}!} \\ &\geq \sum_{m=0}^{\lfloor D/4 \rfloor} \sum_{\underline{\alpha} \in \tilde{\mathcal{I}}_m} \left(\frac{\eta^4 \kappa_4^g}{\sqrt{4!} d^2} \right)^{2m} \\ &\geq \sum_{m=0}^{\lfloor D/4 \rfloor} \# \tilde{\mathcal{I}}_m \left(\frac{\eta^4 \kappa_4^g}{\sqrt{4!} d^2} \right)^{2m}. \end{aligned} \quad (103)$$

Now we just need to estimate the cardinality of $\tilde{\mathcal{I}}_m$. A lower bound can be provided by considering

$$\hat{\mathcal{I}}_m = \{ \underline{\alpha} \in \mathcal{I}_m \mid \forall \mu \in [n] \ \forall j \in [d] \ \alpha_j^{\mu} \text{ is even} \}. \quad (104)$$

Clearly $\hat{\mathcal{I}}_m \subseteq \tilde{\mathcal{I}}_m$ and $\# \hat{\mathcal{I}}_m = \binom{n}{m} \binom{d+1}{2}^m$ because first we can pick the $n-m$ rows of $\underline{\alpha}$ that will be $0 \in \mathbb{N}^d$, which can be done in $\binom{n}{m}$. Then, for each non-zero row, we need to

Learning from higher-order statistics, efficiently: hypothesis tests, random features, and neural networks

pick two columns (with repetitions) in which to place a 2 and leave all the other entries as 0, that can be done in $\binom{d+1}{2}^m$ ways.

So plugging the lower bound in the previous estimate, we reach the inequality that we wanted to prove:

$$\begin{aligned} \|L_{n,d}^{\leq D}\|^2 &\geq \sum_{m=0}^{\lfloor D/4 \rfloor} \binom{n}{m} \binom{d+1}{2}^m \left(\frac{\eta^4 \kappa_4^g}{\sqrt{4!} d^2} \right)^{2m} \\ &= \sum_{m=0}^{\lfloor D/4 \rfloor} \binom{n}{m} \binom{d+1}{2}^m \left(\frac{\beta^2 \kappa_4^g}{\sqrt{4!} d^2 (1+\beta)^2} \right)^{2m}. \end{aligned} \quad (105)$$

Upper bound We start from (75) using the following rewriting:

$$\begin{aligned} \|L_{n,d}^{\leq D}\|^2 &= \sum_{|\underline{\alpha}| \leq D} \frac{\langle L, H_{\underline{\alpha}} \rangle^2}{\underline{\alpha}!} \\ &= \sum_{|\underline{\alpha}| \leq D} \frac{\mathbb{E}_{\underline{x} \sim \mathbb{P}} [H_{\underline{\alpha}}(\underline{x})]^2}{\underline{\alpha}!} \\ &= \sum_{|\underline{\alpha}| \leq D} \frac{1}{\underline{\alpha}!} \mathbb{E}_{u \sim \mathcal{P}(u)} \left[\prod_{\mu=1}^n \mathbb{E}_{x^\mu \sim \mathbb{P}(\cdot|u)} [H_{\underline{\alpha}^\mu}(x^\mu)] \right]^2. \end{aligned} \quad (106)$$

Now use 16 and plug in the formulas for the inner expectations.

$$\begin{aligned} \|L_{n,d}^{\leq D}\|^2 &= \sum_{|\underline{\alpha}| \leq D} \frac{1}{\underline{\alpha}!} \mathbb{E}_{u \sim \mathcal{P}(u)} \left[\prod_{\mu=1}^n \frac{T_{|\underline{\alpha}^\mu|,g}}{d^{|\underline{\alpha}^\mu|/2}} u^{\underline{\alpha}^\mu} \right]^2 \\ &= \sum_{|\underline{\alpha}| \leq D} \frac{\left(\prod_{\mu=1}^n T_{|\underline{\alpha}^\mu|,g} \right)^2}{\underline{\alpha}! d^{|\underline{\alpha}|}} \mathbb{E}_{u \sim \mathcal{P}(u)} [u^{\underline{\alpha}}]^2 \end{aligned} \quad (107)$$

Now we use our prior assumption that $u_i \stackrel{\text{i.i.d.}}{\sim} \text{Rad}(1/2)$. Note that odd moments of u_i are equal to 0 and even moments are equal to 1, so (denoting by $\chi_A(\cdot)$ the indicator function of set A):

$$\mathbb{E}_{u \sim \mathcal{P}(u)} [u^{\underline{\alpha}}] = \prod_{i=1}^d \mathbb{E}_{u_i \sim \text{Rademacher}(1/2)} \left[u_i^{\sum_{\mu=1}^n \alpha_i^\mu} \right] = \chi_{\{\sum_{\mu=1}^n \alpha_i^\mu \text{ is even for all } i\}}(\underline{\alpha}). \quad (108)$$

Now the key point of the proof: this last formula, together with (92), implies that most of the addends in the sum in (107) are zero. Set $|\underline{\alpha}| = m$, then the set of multi-indices that could give a non-zero contribution is

$$\mathcal{A}_m = \left\{ \underline{\alpha} \in \mathbb{N}^{n \times d} \mid |\underline{\alpha}| = m, \forall \mu \in [n], \underline{\alpha}^\mu = 0 \text{ or } |\underline{\alpha}^\mu| > 4, \text{ and } \forall i \in [d] \sum_{\mu=1}^n \alpha_i^\mu \text{ is even} \right\}. \quad (109)$$

Using this fact together with (82) we get:

$$\begin{aligned}
 \|L_{n,d}^{\leq D}\|^2 &= \sum_{m=0}^D \sum_{\underline{\alpha} \in \mathcal{A}_m} \frac{\left(\prod_{\mu=1}^n T_{|\alpha^\mu|,g}\right)^2}{\underline{\alpha}! d^{|\underline{\alpha}|}} \\
 &\leq \sum_{m=0}^D \sum_{\underline{\alpha} \in \mathcal{A}_m} \frac{1}{d^m} \prod_{\mu=1}^n \left(\left(\frac{\beta}{1+\beta} \right)^{|\alpha^\mu|} \mathbb{E}[h_{|\alpha^\mu|}(g)]^2 \right) \\
 &\leq \sum_{m=0}^D \sum_{\underline{\alpha} \in \mathcal{A}_m} \frac{1}{d^m} \left(\frac{\beta}{1+\beta} \right)^m \sup_{k \leq m} \left(\mathbb{E}[h_k(g)]^{2m/k} \right) \\
 &= \sum_{m=0}^D \frac{1}{d^m} \left(\frac{\beta}{1+\beta} \right)^m \sup_{k \leq m} \left(\mathbb{E}[h_k(g)]^{2m/k} \right) \# \mathcal{A}_m.
 \end{aligned} \tag{110}$$

In the third step we used the following inequality

$$\prod_{\mu} \mathbb{E}[h_{|\alpha^\mu|}(g)]^2 \leq \prod_{\mu} \sup_{k \leq m} \left(\mathbb{E}[h_k(g)]^{1/k} \right)^{2|\alpha^\mu|} = \sup_{k \leq m} \mathbb{E}[h_k(g)]^{2m/k}$$

That holds since $\underline{\alpha} \in \mathcal{A}_m$ hence $\sum_{\mu} |\alpha^\mu| = m$.

It is hard to compute exactly the cardinality of $\mathcal{A}_m$, but we can estimate it by considering the inclusion $\mathcal{A}_m \subseteq \tilde{\mathcal{A}}_m$ defined as

$$\begin{aligned}
 \tilde{\mathcal{A}}_m &= \left\{ \underline{\alpha} \in \mathbb{N}^{n \times d} \mid |\underline{\alpha}| = m, \text{ there are at most } \left\lfloor \frac{m}{4} \right\rfloor \text{ non-zero rows and} \right. \\
 &\quad \left. \times \left\lfloor \frac{m}{2} \right\rfloor \text{ non-zero columns} \right\}.
 \end{aligned} \tag{111}$$

Assume now $m > 0$ (the case $m = 0$ can be treated separately, leading to the addend +1 in (74)). To compute the cardinality of $\tilde{\mathcal{A}}_m$ we just have to multiply the different contributions

- There are $\binom{n}{\lfloor m/4 \rfloor}$ possibilities for the non-zero rows
- There are $\binom{d}{\lfloor m/2 \rfloor}$ possibilities for the non-zero columns
- once restricted to a $\lfloor m/4 \rfloor \times \lfloor m/2 \rfloor$ we have to place the units to get to norm m . It is the counting problem of placing m units inside the $\lfloor m/4 \rfloor \lfloor m/2 \rfloor$ matrix entries. The possibilities are

$$\binom{\lfloor m/4 \rfloor \lfloor m/2 \rfloor + m - 1}{m}. \tag{112}$$

So we get the following estimate for the cardinality of $\mathcal{A}_m$

$$\# \mathcal{A}_m \leq \binom{n}{\lfloor m/4 \rfloor} \binom{d}{\lfloor m/2 \rfloor} \binom{\lfloor m/4 \rfloor \lfloor m/2 \rfloor + m - 1}{m}. \tag{113}$$

Plugging into (110) we reach the final formula (74), which completes the proof.

Proof of corollary 15. Now we turn to the proof of 15 on the asymptotic behavior of the bound

$$\binom{n}{m} = \frac{n^m}{m!} + O(n^{m-1}) \quad (114)$$

$$\binom{d+1}{2}^m \geq \frac{d^{2m}}{2^m}. \quad (115)$$

So we have that:

$$\begin{aligned} \|L_{n,d}^{\leq D(n)}\|^2 &\geq \sum_{m=0}^{\lfloor D/4 \rfloor} \left(\frac{n^m}{m!} + O(n^{m-1}) \right) \left(\frac{d^{2m}}{2^m} + O(1) \right) \left(\frac{\beta^2 \kappa_4^g}{\sqrt{4!} d^2 (1+\beta)^2} \right)^{2m} \\ &= \left(\sum_{m=0}^{\lfloor D/4 \rfloor} \frac{1}{m!} \left(\frac{\beta^2 \kappa_4^g}{\sqrt{2} \sqrt{4!} (1+\beta)^2} \right)^{2m} \left(\frac{n}{d^2} \right)^m \right) + O \left(\sum_m \frac{n^{m-1}}{m! d^{2m}} \left(\frac{\beta^2 \kappa_4^g}{\sqrt{2} \sqrt{4!} (1+\beta)^2} \right)^{2m} \right) \end{aligned} \quad (116)$$

Then we lower-bound the sum by considering just the last term

$$\|L_{n,d}^{\leq D(n)}\|^2 \geq \frac{1}{\lfloor D/4 \rfloor!} \left(\frac{\beta^2 \kappa_4^g}{\sqrt{2} \sqrt{4!} (1+\beta)^2} \right)^{2\lfloor D/4 \rfloor} \left(\frac{n}{d^2} \right)^{\lfloor D/4 \rfloor} + o \left(\frac{1}{\lfloor D/4 \rfloor!} \left(\frac{Cn}{d^2} \right)^{\lfloor D/4 \rfloor} \right). \quad (117)$$

So by using $k! \leq k^k$ on $\lfloor D/4 \rfloor$ and picking n, d large enough so that numerical constants and the $o(\dots)$ become negligible for the estimate, we get (10).

Plugging in the scaling $n \asymp d^\theta$ and $D(n) \asymp \log^{1+\varepsilon}(n)$ it is clear that what decides the behaviour of the sequence is the term $\frac{n}{d^2}$.

- If $0 < \theta \leq 2$, $\frac{n}{d^2} \rightarrow 0$ and so (10) does not provide information on the divergence of the LDLR norm.
- If $\theta > 2$ it is $\frac{n}{d^2} \rightarrow \infty$ faster than the logarithmic term at denominator, so the right hand side in (10) diverges at ∞ , proving the second regime in (12).

Now let us turn to proving (11). We are in the regime $m \leq D \ll \min(n, d)$, hence we can use the following estimates of binomial coefficients and factorial

$$\begin{aligned} \binom{n}{\lfloor m/4 \rfloor} &\leq n^{m/4} \\ \binom{d}{\lfloor m/2 \rfloor} &\leq d^{m/2} \\ \binom{\lfloor m/4 \rfloor \lfloor m/2 \rfloor + m - 1}{m} &\leq (m^2/8)^m \\ m! &\leq m^m \end{aligned} \quad (118)$$

on (74) to get

$$\|L_{n,d}^{\leq D}\|^2 \leq 1 + \sum_{m=1}^D \left(\frac{\beta}{1+\beta}\right)^m m^{2m} \sup_{k \leq m} \left(\mathbb{E}[h_k(g)]^{2m/k}\right) \left(\frac{n}{d^2}\right)^{m/4}. \quad (119)$$

Now we use estimate the term depending on the Hermite coefficients of g , using the assumption 5:

$$\begin{aligned} \sup_{k \leq m} \left(\mathbb{E}[h_k(g)]^{2m/k}\right) &\leq \sup_{k \leq m} \left((\Lambda^k k!)^{2m/k}\right) \\ &\leq \Lambda^{2m} \sup_{k \leq m} (k^{2m}) \\ &\leq \Lambda^{2m} m^{2m}. \end{aligned} \quad (120)$$

Plugging into the estimate for the LDLR we get (11)

$$\|L_{n,d}^{\leq D}\|^2 \leq 1 + \sum_{m=1}^D \left(\frac{\Lambda^2 \beta}{1+\beta}\right)^m m^{4m} \left(\frac{n}{d^2}\right)^{m/4}. \quad (121)$$

By letting $n \asymp d^\theta$, for $\theta > 0$ and $D(n) \asymp \log^{1+\varepsilon}(n)$, we can see that the bound goes to 1 when $\theta < 2$, proving the first regime in (12).

B.7. Limitations of our theoretical analysis

The main limitation of the theoretical portion of our work is that it relies on the assumption that the null hypothesis is standard i.i.d. Gaussian noise. Since this assumption is central to the large body of work analysing hypothesis tests [37], a very interesting future direction is to develop tools for analysing the case of a different null hypothesis. On the technical side, while the assumption of Rademacher prior on u is not essential and could be easily generalized to other isotropic distributions, assumption 1 is a key requirement to carry out the proofs. We discuss the limitations of the RF analysis, and in particular the Gaussian equivalence theorem [79–83], at the end of section 4.1.

As is generally the case in high-dimensional statistics, our results do not constitute a complete proof of the existence of a statistical-to-computational gap: our argument relies on the low-degree conjecture 4, and the divergence of the norm of the LR is only a necessary condition for *strong distinguishability*. In fact, there are no techniques at the moment that can prove that average-case problems require super-polynomial time in the hard phase, even if we assume $P \neq NP$ [37]. So our work should be seen as providing rigorous evidence for a statistical-to-computational gap in the spiked cumulant model.

Appendix C. Details on the replica analysis

We analytically compute the generalisation error of a RF model trained on the Gaussian mixture classification task of B.2.3 by exploiting the Gaussian equivalence theorem for dealing with structured data distributions [79–83]. In particular, we use equations

Learning from higher-order statistics, efficiently: hypothesis tests, random features, and neural networks

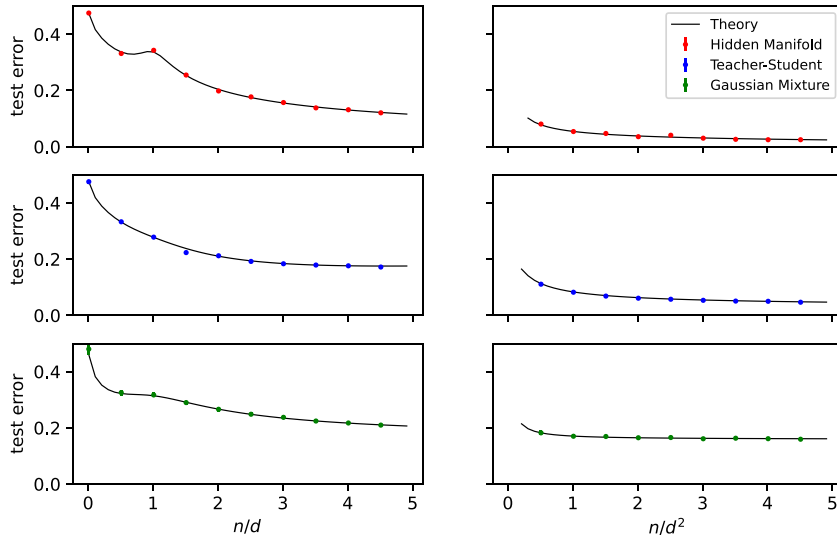


Figure 8. Linear and quadratic sample regimes for different synthetic data models. (*Right*) Generalization error of the hidden manifold model (top), the teacher-student setup (center) and a mixture of two Gaussians as a function of the ratio of the number of samples and the input dimension. (*Left*) Same except that the number of samples scales with the square of the input dimension. The solid black line corresponds to the replica theory prediction while the coloured dots display the outcome of the numerical simulations averaged over 10 different seeds. In all the panels, $d = 1000$ and $d = 20$ for linear and quadratic sample regimes respectively, $\lambda = 0.01$ for both the teacher-student setup and the hidden manifold model while $\lambda = 0.1$ for Gaussian mixtures. In the Gaussian mixture case, $\mu_{\pm} = (\pm 1, 0, \dots, 0) \in \mathbb{R}^d$ while the covariance matrices are both isotropic and, in particular, both equal to the identity matrix: $\Sigma_{\pm} = \mathbb{I}$.

derived by Loureiro [78] that describe the generalisation error of RFs on Gaussian mixture classification using the replica method of statistical physics. The equations describe the high-dimensional limit where n, d , and the size of the hidden layer $p \rightarrow \infty$ while their ratios stay finite, $\alpha = n/d$ and $\gamma = d/p \sim O(1)$. In this regime, the generalisation error is a function of only scalar quantities, i.e.

$$\epsilon_g = 1 - \rho_+ \mathbb{E}_{\xi} \left[\text{sign} \left(m_+^* + \sqrt{q_+^*} \xi + b^* \right) - \right] + \rho_- \mathbb{E}_{\xi} \left[\text{sign} \left(m_-^* + \sqrt{q_-^*} \xi + b^* \right) \right], \quad (122)$$

where $\xi \sim \mathcal{N}(0, 1)$, $\rho_{\pm}$ defines the fraction of data points belonging to class $y = \pm 1$ while $m_{\pm}^*$, $q_{\pm}^*$ and b^* are the so-called overlap parameters and bias term respectively. Their value for every size of the training set can be determined by solving the following optimisation problem [78]:

$$f_{\beta} = \underset{\{q_k, m_k, V_k, \hat{q}_k, \hat{m}_k, \hat{V}_k, b\}_{k=+, -}}{\text{extr}} \left[\sum_{k=+, -} -\frac{1}{2} \left(\hat{q}_k V_k - q_k \hat{V}_k \right) + m_k \hat{m}_k + \lim_{p \rightarrow \infty} \frac{1}{p} \Psi_s + \alpha \gamma \Psi_e \right] \quad (123)$$

where the entropic potential Ψ_s can be expressed as a function of the means $\boldsymbol{\mu}_\pm \in \mathbb{R}^p$ and covariances $\Sigma_\pm \in \mathbb{R}^{p \times p}$ of the hidden-layer activations of the RFs model, while the energetic potential Ψ_e depends on the specific choice of the loss function $\ell(\cdot)$.

Solving the optimization problem in 123, leads to a set of coupled saddle-point equations for the overlap parameters and the bias term. This set of equations is nothing but a special case of the equations already derived in [78], that is high-dimensional classification of a mixture of two Gaussians with RFs, except for the typo in the self-consistency equation of the bias term

$$b = \left(\sum_{k=\pm} \frac{\rho_k}{V_k} \right)^{-1} \sum_{k=\pm} \rho_k \mathbb{E}_\xi [\eta_k - m_k], \quad (124)$$

with η_k being the extremizer

$$\eta_k = \underset{\lambda}{\operatorname{argmin}} \left[\frac{(\lambda - \sqrt{q_k} \xi - m_k - b)^2}{2V_+} + \ell(y_k, \lambda) \right] \quad (125)$$

and $\ell(\cdot)$ being any convex loss function. Analogously to the more general setting of [78], the Gaussian Equivalence Theorem allows to express the entropic potential Ψ_s as a function of the means $\boldsymbol{\mu}_k \in \mathbb{R}^p$ and covariances $\Sigma_k \in \mathbb{R}^{p \times p}$ of the RFs

$$\begin{aligned} \Psi_s = & \lim_{p \rightarrow \infty} \frac{1}{p} (\hat{q}_+ \boldsymbol{\mu}_+ + \hat{q}_- \boldsymbol{\mu}_-)^t \left(\lambda \mathbb{I}_p + \hat{V}_+ \Sigma_+ + \hat{V}_- \Sigma_- \right)^{-1} (\hat{q}_+ \boldsymbol{\mu}_+ + \hat{q}_- \boldsymbol{\mu}_-) \\ & + \lim_{p \rightarrow \infty} \frac{1}{p} \operatorname{Tr} \left((\hat{q}_+ \Sigma_+ + \hat{q}_- \Sigma_-) \left(\lambda \mathbb{I}_p + \hat{V}_+ \Sigma_+ + \hat{V}_- \Sigma_- \right)^{-1} \right); \end{aligned} \quad (126)$$

while the energetic potential Ψ_e depends on the specific choice of the loss function $\ell(\cdot)$

$$\Psi_e = -\frac{1}{2} \mathbb{E}_\xi \left[\sum_{k=\pm} \rho_k \frac{(\eta_k - \sqrt{q_k} \xi - m_k - b)^2}{2V_k} + \ell(y_k, \eta_k) \right]. \quad (127)$$

As discussed in 4.1, the Gaussian Equivalence Theorem breaks in the quadratic sample regime if the data are sampled from the spiked Wishart model. Interestingly, this is not the case for the Hidden Manifold model. Indeed, as shown in figure 8, we still get a perfect match between the Gaussian theory and the numerical simulations even in the quadratic sample regime. The theoretical prediction can be easily derived by rescaling the free-energy associated to the Hidden Manifold model as in [80] by a factor $1/d^2$. This is the same trick already proposed in [84] for teacher–student settings and displayed in the middle panel of figure 8.

References

- [1] Refinetti M, Ingrosso A and Goldt S 2023 Neural networks trained with SGD learn distributions of increasing complexity *Int. Conf. on Machine Learning* pp 28843–63
- [2] Jacot A, Gabriel F and Hongler C 2018 Neural tangent kernel: convergence and generalization in neural networks *Advances in Neural Information Processing Systems* vol 31 (Curran Associates, Inc.)
- [3] Arora S, Du S S, Hu W, Li Z, Salakhutdinov R R and Wang R 2019 On exact computation with an infinitely wide neural net *Advances in Neural Information Processing Systems* vol 32
- [4] Geiger M, Spigler S, Jacot A and Wyart M 2020 Disentangling feature and lazy training in deep neural networks *J. Stat. Mech.* **113301**
- [5] Barak B, Hopkins S, Kelner J, Kothari P K, Moitra A and Potechin A 2019 A nearly tight sum-of-squares lower bound for the planted clique problem *SIAM J. Comput.* **48** 687–735
- [6] Hopkins S B, Kothari P K, Potechin A, Raghavendra P, Schramm T and Steurer D 2017 The power of sum-of-squares for detecting hidden structures *IEEE 58th Annual Symp. on Foundations of Computer Science (FOCS)* pp 720–31
- [7] Hopkins S B and Steurer D 2017 Bayesian estimation from few samples: community detection and related problems (arXiv: [1710.00264](https://arxiv.org/abs/1710.00264))
- [8] Hopkins S 2018 Statistical inference and the sum of squares method *PhD Thesis*
- [9] Baik J, Ben Arous G and Pécché S 2005 Phase transition of the largest eigenvalue for nonnull complex sample covariance matrices *Ann. Probab.* **33** 1643–97
- [10] Potters M and Bouchaud J-P 2020 *A First Course in Random Matrix Theory: for Physicists, Engineers and Data Scientists* ed S Bengio H Wallach H Larochelle K Grauman N Cesa-Bianchi R Garnett (Cambridge University Press)
- [11] Diakonikolas I, Kane D M and Stewart A 2017 Statistical query lower bounds for robust estimation of high-dimensional Gaussians and Gaussian mixtures *2017 IEEE 58th Annual Symp. on Foundations of Computer Science (FOCS)* pp 73–84
- [12] Dudeja R and Hsu D 2024 Statistical-computational trade-offs in tensor PCA and related problems via communication complexity *Ann. Stat.* **52** 131–56
- [13] Ingrosso A and Goldt S 2022 Data-driven emergence of convolutional structure in neural networks *Proc. Natl Acad. Sci.* **119** e2201854119
- [14] Miolane L 2018 Phase transitions in spiked matrix estimation: information-theoretic analysis (arXiv: [1806.04343](https://arxiv.org/abs/1806.04343))
- [15] Lelarge M and Miolane L 2019 Fundamental limits of symmetric low-rank matrix estimation *Probab. Theory Relat. Fields* **173** 859–929
- [16] El Alaoui A, Krzakala F and Jordan M 2020 Fundamental limits of detection in the spiked Wigner model *Ann. Stat.* **48** 863–85
- [17] Paul D 2007 Asymptotics of sample eigenstructure for a large dimensional spiked covariance model *Stat. Sin.* **17** 1617–42
- [18] Berthet Q and Rigollet P 2012 Optimal detection of sparse principal components in high dimension *Ann. Stat.* **41** 1780–815
- [19] Berthet Q and Rigollet P 2013 Complexity theoretic lower bounds for sparse principal component detection *Conf. on Learning Theory (COLT)*
- [20] Lesieur T, Krzakala F and Zdeborová L 2015 Phase transitions in sparse PCA *2015 IEEE Int. Symp. on Information Theory (ISIT)* pp 1635–9
- [21] Lesieur T, Krzakala F and Zdeborová L 2015 MMSE of probabilistic low-rank matrix estimation: Universality with respect to the output channel *2015 53rd Annual Allerton Conf. on Communication, Control and Computing (Allerton)* pp 680–7
- [22] Perry A, Wein A S, Bandeira A S and Moitra A 2016 Optimality and sub-optimality of PCA for spiked random matrices and synchronization (arXiv: [1609.05573](https://arxiv.org/abs/1609.05573))
- [23] Krzakala F, Xu J and Zdeborová L 2016 Mutual information in rank-one matrix estimation *2016 IEEE Information Theory Workshop (ITW)* pp 71–75
- [24] Dia M, Macris N, Krzakala F, Lesieur T and Zdeborová L 2016 Mutual information for symmetric rank-one matrix estimation: a proof of the replica formula *Advances in Neural Information Processing Systems* vol 29
- [25] Richard E and Montanari A 2014 A statistical model for tensor PCA *Advances in Neural Information Processing Systems* vol 27

- [26] Hopkins S B, Shi J and Steurer D 2015 Tensor principal component analysis via sum-of-square proofs *Conf. on Learning Theory* pp 956–1006
- [27] Montanari A, Reichman D and Zeitouni O 2015 On the limitation of spectral methods: from the Gaussian hidden clique problem to rank-one perturbations of Gaussian tensors *Advances in Neural Information Processing Systems* vol 28
- [28] Perry A, Wein A S and Bandeira A S 2016 Statistical limits of spiked tensor models (arXiv: [1612.07728](#))
- [29] Kim C, Bandeira A S and Goemans M X 2017 Community detection in hypergraphs, spiked tensor models and sum-of-squares *2017 Int. Conf. on Sampling Theory and Applications (SampTA)* pp 124–8
- [30] Lesieur T, Miolane L, Lelarge M, Krzakala F and Zdeborová L 2017 Statistical and computational phase transitions in spiked tensor estimation *2017 IEEE Int. Symp. on Information Theory (ISIT)* pp 511–5
- [31] Arous G B, Gheissari R and Jagannath A 2020 Algorithmic thresholds for tensor PCA *Ann. Probab.* **48** 2052–87
- [32] Jagannath A, Lopatto P and Miolane L 2020 Statistical thresholds for tensor PCA *Ann. Appl. Probab.* **30** 1910–33
- [33] Niles-Weed J and Rigollet P 2022 Estimation of wasserstein distances in the spiked transport model *Bernoulli* **28** 2663–88
- [34] Zdeborová L and Krzakala F 2016 Statistical physics of inference: thresholds and algorithms *Adv. Phys.* **65** 453–552
- [35] Lesieur T, Krzakala F and Zdeborová L 2017 Constrained low-rank matrix estimation: phase transitions, approximate message passing and applications *J. Stat. Mech.* **073403**
- [36] Bandeira A S, Perry A and Wein A S 2018 Notes on computational-to-statistical gaps: predictions using statistical physics *Port. Math.* **75** 159–86
- [37] Kunisky D, Wein A S and Bandeira A S 2019 Notes on computational hardness of hypothesis testing: predictions using the low-degree likelihood ratio *ISAAC Congress (Int. Society for Analysis, its Applications and Computation)* pp 1–50
- [38] Wang C and Lu Y M 2017 The scaling limit of high-dimensional online independent component analysis *Advances in Neural Information Processing Systems* vol 31
- [39] Hyvärinen A and Oja E 2000 Independent component analysis: algorithms and applications *Neural Netw.* **13** 411–30
- [40] Blanchard G, Kawanabe M, Sugiyama M, Spokoiny V and Müller K -R 2006 In Search of Non-Gaussian components of a high-dimensional distribution *J. Mach. Learn. Res.* **7** 247–82
- [41] Bean D M 2014 *Non-Gaussian Component Analysis* (University of California, Berkeley)
- [42] Vempala S S and Xiao Y 2012 Structure from local optima: learning subspace juntas via higher order PCA (arXiv: [1108.3329](#) [cs.CC])
- [43] Tan Y S and Vershynin R 2018 Polynomial time and sample complexity for Non-Gaussian component analysis: spectral methods *Proc. 31st Conf. On Learning Theory (Proc. of Machine Learning Research* vol 75) pp 498–534
- [44] Goyal N and Shetty A 2018 Non-Gaussian component analysis using entropy methods *Proc. 51st Annual ACM SIGACT Symp. on Theory of Computing*
- [45] Mao C and Wein A S 2022 Optimal spectral recovery of a planted vector in a subspace (arXiv: [2105.15081](#) [math.ST])
- [46] Damian A, Pillaud-Vivien L, Lee J D and Bruna J 2024 Computational-statistical gaps in Gaussian single-index models (arXiv: [2403.05529](#) [cs.LG])
- [47] Diakonikolas I, Kane D, Ren L and Sun Y 2023 SQ lower bounds for Non-Gaussian component analysis with weaker assumptions *37th Conf. on Neural Information Processing Systems*
- [48] Brennan M S, Bresler G, Hopkins S, Li J and Schramm T 2021 Statistical query algorithms and low degree tests are almost equivalent *Proc. 34th Conf. on Learning Theory (Proc. of Machine Learning Research* vol 134) ed M Belkin and S Kpotufe pp 774–774
- [49] Li Y and Yuan Y 2017 Convergence analysis of two-layer neural networks with relu activation *Advances in Neural Information Processing Systems* vol 30
- [50] Du S S, Zhai X, Póczos B and Singh A 2019 Gradient descent provably optimizes over-parameterized neural networks *Int. Conf. on Learning Representations*
- [51] Li Y and Liang Y 2018 Learning overparameterized neural networks via stochastic gradient descent on structured data *Advances in Neural Information Processing Systems* vol 31
- [52] Allen-Zhu Z, Li Y and Song Z 2019 A convergence theory for deep learning via over-parameterization *Int. Conf. on Machine Learning* pp 242–52
- [53] Bordelon B, Canatar A and Pehlevan C 2020 Spectrum dependent learning curves in kernel regression and wide neural networks *Int. Conf. on Machine Learning* pp 1024–34

- [54] Canatar A, Bordelon B and Pehlevan C 2021 Spectral bias and task-model alignment explain generalization in kernel regression and infinitely wide neural networks *Nat. Commun.* **12** 2914
- [55] Nguyen Q, Mondelli M and Montufar G F 2021 Tight bounds on the smallest eigenvalue of the neural tangent kernel for deep relu networks *Int. Conf. on Machine Learning* pp 8119–29
- [56] Bach F 2017 Breaking the curse of dimensionality with convex neural networks *J. Mach. Learn. Res.* **18** 629–81
- [57] Ghorbani B, Mei S, Misiakiewicz T and Montanari A 2019 Limitations of lazy training of two-layers neural network *Advances in Neural Information Processing Systems* vol 32 pp 9111–21
- [58] Ghorbani B, Mei S, Misiakiewicz T and Montanari A 2020 When do neural networks outperform kernel methods? *Advances in Neural Information Processing Systems* vol 33
- [59] Chizat L and Bach F 2020 Implicit bias of gradient descent for wide two-layer neural networks trained with the logistic loss *Conf. on Learning Theory* pp 1305–38
- [60] Daniely A and Malach E 2020 Learning parities with neural networks *Advances in Neural Information Processing Systems* vol 33 pp 20356–65
- [61] Paccolat J, Petrini L, Geiger M, Tyloo K and Wyart M 2021 Geometric compression of invariant manifolds in neural networks *J. Stat. Mech.* 044001
- [62] Yehudai G and Shamir O 2019 On the power and limitations of random features for understanding neural networks *Advances in Neural Information Processing Systems* vol 32 pp 6598–608
- [63] Refinetti M, Goldt S, Krzakala F and Zdeborová L 2021 Classifying high-dimensional Gaussian mixtures: where kernel methods fail and neural networks succeed *Int. Conf. on Machine Learning* pp 8936–47
- [64] Baldi P and Chauvin Y 1991 Temporal evolution of generalization during learning in linear networks *Neural Comput.* **3** 589–603
- [65] Le Cun Y, Kanter I and Solla S A 1991 Eigenvalues of covariance matrices: application to neural-network learning *Phys. Rev. Lett.* **66** 2396
- [66] Saxe A M, McClelland J L and Ganguli S 2014 Exact solutions to the nonlinear dynamics of learning in deep linear neural networks *ICLR*
- [67] Advani M S, Saxe A M and Sompolinsky H 2020 High-dimensional dynamics of generalization error in neural networks *Neural Netw.* **132** 428–46
- [68] Kesten H and Stigum B P 1966 Additional limit theorems for indecomposable multidimensional Galton-Watson processes *Ann. Math. Stat.* **37** 1463–81
- [69] Decelle A, Krzakala F, Moore C and Zdeborová L 2011 Inference and phase transitions in the detection of modules in sparse networks *Phys. Rev. Lett.* **107** 065701
- [70] Decelle A, Krzakala F, Moore C and Zdeborová L 2011 Asymptotic analysis of the stochastic block model for modular networks and its algorithmic applications *Phys. Rev. E* **84** 066106
- [71] Balcan M -F, Blum A and Vempala S 2006 Kernels as features: on kernels, margins and low-dimensional mappings *Mach. Learn.* **65** 79–94
- [72] Rahimi A and Recht B 2008 Random features for large-scale kernel machines *Advances in Neural Information Processing Systems* pp 1177–84
- [73] Rahimi A and Recht B 2009 Weighted sums of random kitchen sinks: replacing minimization with randomization in learning *Advances in Neural Information Processing Systems* pp 1313–20
- [74] Ghorbani S, Mei T and Misiakiewicz A 2021 Linearized two-layers neural networks in high dimension *Ann. Stat.* **49** 1029–54
- [75] Xiao L, Hu H, Misiakiewicz T, Lu Y and Pennington J 2022 Precise learning curves and higher-order scalings for dot-product kernel regression *Advances in Neural Information Processing Systems* vol 35 pp 4558–70
- [76] Mei S, Misiakiewicz T and Montanari A 2022 Generalization error of random feature and kernel methods: hypercontractivity and kernel matrix concentration *Appl. Comput. Harmon. Anal.* **59** 3–84
- [77] Misiakiewicz T and Montanari A 2023 Six lectures on linearized neural networks (arXiv: 2308.13431)
- [78] Loureiro B, Sicuro G, Gerbelot C, Pocco A, Krzakala F and Zdeborová L 2021 Learning Gaussian mixtures with generalized linear models: precise asymptotics in high-dimensions *Advances in Neural Information Processing Systems* vol 34 pp 10144–57
- [79] Goldt S, Mézard M, Krzakala F and Zdeborová L 2020 Modeling the influence of data structure on learning in neural networks: the hidden manifold model *Phys. Rev. X* **10** 041044
- [80] Gerace F, Loureiro B, Krzakala F, Mézard M and Zdeborová L 2020 Generalisation error in learning with random features and the hidden manifold model *Int. Conf. on Machine Learning* pp 3452–62
- [81] Hu H and Lu Y M 2022 Universality laws for high-dimensional learning with random features *IEEE Trans. Inf. Theory* **69** 1932–64
- [82] Mei S and Montanari A 2022 The generalization error of random features regression: precise asymptotics and the double descent curve *Commun. Pure Appl. Math.* **75** 667–766

- [83] Goldt S, Loureiro B, Reeves G, Krzakala F, Mézard M and Zdeborová L 2022 The Gaussian equivalence of generative models for learning with shallow neural networks *Mathematical and Scientific Machine Learning* pp 426–71
- [84] Dietrich R, Oppen M and Sompolinsky H 1999 Statistical mechanics of support vector networks *Phys. Rev. Lett.* **82** 2975
- [85] Gardner E and Derrida B 1989 Three unfinished works on the optimal storage capacity of networks *J. Phys. A: Math. Gen.* **22** 1983
- [86] Cui H, Krzakala F and Zdeborova L 2023 Bayes-optimal learning of deep random networks of extensive-width *Int. Conf. on Machine Learning* pp 6468–521
- [87] Camilli F, Tieplová D and Barbier J 2023 Fundamental limits of overparametrized shallow neural networks for supervised learning (arXiv: [2307.05635](https://arxiv.org/abs/2307.05635))
- [88] Hu H and Lu Y M 2022 Sharp asymptotics of kernel ridge regression beyond the linear regime (arXiv: [2205.06798](https://arxiv.org/abs/2205.06798))
- [89] Bell A and Sejnowski T J 1996 Edges are the ‘Independent Components’ of Natural Scenes *Advances in Neural Information Processing Systems* vol 9, ed M C Mozer M Jordan T Petsche (MIT Press)
- [90] Chizat L, Oyallon E and Bach F 2019 On lazy training in differentiable programming *Advances in Neural Information Processing Systems* vol 32, ed H Wallach H Larochelle A Beygelzimer F d’Alché-Buc E Fox R Garnett (Curran Associates, Inc.)
- [91] Kolda T G and Bader B W 2009 Tensor decompositions and applications *SIAM Rev.* **51** 455–500
- [92] Bardone L and Goldt S 2024 Sliding down the stairs: how correlated latent variables accelerate learning with neural networks *Proc. 41st Int. Conf. on Machine Learning (Proc. of Machine Learning Research* vol 235) ed R Salakhutdinov Z Kolter K Heller A Weller N Oliver J Scarlett F Berkenkamp (PMLR) pp 3024–45
- [93] Pedregosa F *et al* 2011 Scikit-learn: machine learning in python *J. Mach. Learn. Res.* **12** 2825–30
- [94] Kossaifi J, Panagakis Y, Anandkumar A and Pantic M 2019 TensorLy: tensor learning in python *J. Mach. Learn. Res.* **20** 26
- [95] Banks J, Moore C, Neeman J and Netrapalli P 2016 Information-theoretic thresholds for community detection in sparse networks *Conf. on Learning Theory* pp 383–416
- [96] Bandeira A S, Kunisky D and Wein A S 2020 Computational Hardness of certifying bounds on constrained PCA problems *11th Innovations in Theoretical Computer Science Conf. (ITCS 2020) (Leibniz Int. Proc. in Informatics (LIPIcs))* vol 151, ed T Vidick (Schloss Dagstuhl–Leibniz-Zentrum fuer Informatik) pp 78:1–78:29
- [97] McCullagh P 2018 *Tensor Methods in Statistics* S Bubeck V Perchet P Rigollet (Courier Dover Publications)
- [98] Szegő G 1939 *Orthogonal Polynomials* (American mathematical society)
- [99] Abramowitz M and Stegun I A 1964 *Handbook of Mathematical Functions With Formulas, Graphs and Mathematical tables (Applied Mathematics Series)* vol 55 (National Bureau of Standards) p xiv+1046
- [100] Lukacs E 1972 A survey of the theory of characteristic functions *Adv. Appl. Probab.* **4** 1–38