

Generalized Approach to Matched Filtering using Neural Networks

Jingkai Yan^{1,4}, Mariam Avagyan^{1,4}, Robert E. Colgan^{2,4}, Doğa Veske³,
Imre Bartos⁶, John Wright^{1,4}, Zsuzsa Márka^{4,5}, and Szabolcs Márka^{3,4}

¹*Department of Electrical Engineering, Columbia University in the City of New York, 500 W. 120th St., New York, NY 10027, USA*

²*Department of Computer Science, Columbia University in the City of New York, 500 W. 120th St., New York, NY 10027, USA*

³*Department of Physics, Columbia University in the City of New York, 538 W. 120th St., New York, NY 10027, USA*

⁴*Data Science Institute, Columbia University in the City of New York, 550 W. 120th St., New York, NY 10027, USA*

⁵*Columbia Astrophysics Laboratory, Columbia University in the City of New York, 538 W. 120th St., New York, NY 10027, USA*

⁶*Department of Physics, University of Florida, PO Box 118440, Gainesville, FL 32611-8440, USA*

Gravitational wave science is a pioneering field with rapidly evolving data analysis methodology currently assimilating and inventing deep learning techniques. The bulk of the sophisticated flagship searches of the field rely on the time-tested matched filtering principle within their core. In this paper, we make a key observation on the relationship between the emerging deep learning and the traditional techniques: *matched filtering is formally equivalent to a particular neural network*. This means that a neural network can be constructed analytically to exactly implement matched filtering, and can be further trained on data or boosted with additional complexity for improved performance. Moreover, we show that the proposed neural network architecture can outperform matched filtering, both with or without knowledge of a prior on the parameter distribution. When a prior is given, the proposed neural network can approach the statistically optimal performance. We also propose and investigate two different neural network architectures **MNet-Shallow** and **MNet-Deep**, both of which implement matched filtering at initialization and can be trained on data. **MNet-Shallow** has simpler structure, while **MNet-Deep** is more flexible and can deal with a wider range of distributions. Our theoretical findings are corroborated by experiments using real LIGO data and synthetic injections, where our proposed methods significantly outperform matched filtering at false positive rates above $5 \times 10^{-3}\%$. The fundamental equivalence between matched filtering and neural networks allows us to define a “complexity standard candle” to characterize the relative complexity of the different approaches to gravitational wave signal searches in a common framework. Additionally it also provides a glimpse of an intriguing symmetry that could provide clues on interpretability, namely how neural networks approach the problem of finding signals in overwhelming noise. Finally, our results suggest new perspectives on the role of deep learning in gravitational wave detection.

I. INTRODUCTION

The discovery of cosmic gravitational waves [1], the windfall of binary black-hole (BBH) merger detections [2, 3], and the spectacular insights that multimessenger astrophysics provided [4, 5] revolutionized how we understand the Universe. This leap was due to multiple factors, from instrumental advances to computing breakthroughs. Emerging interferometric gravitational wave detectors, KAGRA [6], GEO600 [7], Virgo [8], and LIGO [9, 10], played a critical role as they provided the technology [11–13] enabling signals to be extracted from ripples in Einstein’s space-time [14, 15]. Of course, as it is not sufficient to have data with faint cosmic signals buried in the noise, the community had to rely on exquisitely sensitive data analysis algorithms to extract transient signals from the noisy data. The bulk of the discoveries were made by two classes of powerful data analysis approaches, *excess power* [16–18] and *matched filtering* [19–25]. The flagship matched filtering methods [26–38] reached unprecedented sophistication and became the workhorse of the field [2, 3]. Insightful work also exist on the extent of optimality, role of intrinsic parameters, and effect of non-Gaussian backgrounds [39–41]. There is more

than historical evidence on their algorithmic power [42], and they are also considered optimal [22] when searching for chirps of known shape [20, 43–45] embedded in well-behaved Gaussian noise. Within the optimality and success lie limitations, as the data is significantly more complex [46, 47] than Gaussian noise and many cosmic signals are not as well known as the BBH models that are being used in searches [48]. Therefore, it is critical that we both seek data analysis methods beyond the horizon of current techniques and rigorously understand the place of current techniques in the broader field of possible methods.

An abundance of prior works has been using deep learning methods for gravitational wave detection. Convolutional neural networks have been shown to be capable of identifying gravitational waves and their parameters from binary black holes and binary neutron stars, with performance approaching the matched filtering search currently used by LIGO, Virgo and KARGA [8, 10, 49–70]. In addition, these machine learning (ML) method can also be applied to glitches and noise transients identification [53, 71–76], signal classification and parameter estimation [77–81], data denoising [82, 83], etc. While these works exhibit neural networks that

could approach the performance of matched filtering, they are still often applied as or considered “black box” models. This makes it challenging to evaluate the statistical evidence provided by neural networks, and to incorporate that evidence in downstream analyses [84].

This paper is motivated by a critical observation, which we substantiate below: *matched filtering with a collection of templates is formally equivalent to a particular neural network*, whose architecture and parameters are dictated by the templates. This observation has precedents in the machine learning literature, where deep neural networks are sometimes viewed as hierarchical template matching methods, with signal-dependent, class-specific templates [85–91]. Here, we delineate a simple and explicit equivalence between matched filtering and particular neural networks, which can be constructed analytically from a set of templates. This equivalence lies in the algorithmic level, and does not depend on specific problem formulations.

In order to study the potential performance gains of using neural networks, we formulate the gravitational wave detection problem abstractly as the detection of a parametric family of signals. Under this framework, we show that the analytically constructed networks can also be used as a principled starting point for learning from data, yielding signal classifiers with better performance than their initialization, namely “standing on the shoulder of giants”. Such learning can be applied to scenarios both with or without a prior distribution on the parameters. In particular, when a prior distribution is given, we show that the learned neural network can (empirically) approach the statistically optimal performance.

We propose and investigate two different neural network architectures for implementing matched filtering, respectively **MNet-Shallow** and **MNet-Deep**. The former has simpler structure, while the latter is more flexible and can deal with a wider range of distributions. These learned classifiers have a number of additional advantages: they do not require prior knowledge of the noise distribution, can be adapted to cope with time-varying noise distributions, and suggest new approaches to computationally efficient signal detection. We conducted experiments using real LIGO data [92] in order to demonstrate the feasibility and power of neural networks in comparison to matched filtering, where we validate our findings empirically that neural networks via training can reach better performance. Finally, interpreting matched filtering and neural networks in a common framework also allows a clear comparison of their computational/storage complexities and statistical strengths, consequently making deep-learning less of a mystery.

The rest of the paper is organized as follows. Section II introduces the problem of parametric signal detection as an abstraction of the gravitational-wave detection problem, and discusses the two formulations of the objective. Section III discusses matched filtering as an approach to solving the parametric detection problem, as well as its limitations. Section IV illustrates how neural net-

work models can be applied in this problem, in a way that exactly implements matched filtering at initialization. Section V discusses the training process of neural network models, and in particular how it is aligned with the parametric signal detection problem. In Section VI we present experimental results on real LIGO data and synthetic injections. We discuss some further implications of this work in Section VII, and conclude in Section VIII.

II. PARAMETRIC SIGNAL DETECTION

The problem of identifying gravitational waves [93] in a single gravitational-wave detector data stream [94] can be formulated as follows: we observe detector strain data $\mathbf{x} \in \mathbb{R}^n$, and wish to determine whether $\mathbf{x}$ consists of astrophysical signal plus noise, or noise alone. We can model possible astrophysical signals as belonging to a parametric family

$$S_\Gamma = \{\mathbf{s}_\gamma \mid \gamma \in \Gamma\} \quad (1)$$

where the parameters γ can represent properties of the objects that generate the gravitational wave, such as masses, orbits and spins. We assume the signals are normalized to have unit power, namely $\|\mathbf{s}_\gamma\|^2 = 1$ for all γ . We model noise as a random vector $\mathbf{z} \in \mathbb{R}^n$, which is assumed to follow distribution ρ_0 and be probabilistically independent of the signal. In this notation, our goal becomes one of solving a hypothesis testing problem:

$$H_0 : \mathbf{x} = \mathbf{z}, \quad (2)$$

$$\text{or } H_1 : \mathbf{x} = \mathbf{s}_\gamma + \mathbf{z} \text{ for some } \gamma \in \Gamma. \quad (3)$$

Note that except for special cases, such as when the hypothesis H_1 is simple, or when the parameters associated with H_1 satisfy certain monotone conditions, we usually do not have a uniformly most powerful test [95].

Our broad goal is to identify decision rules $\delta : \mathbb{R}^n \rightarrow \{0, 1\}$ that (i) have good statistical performance and (ii) can be implemented efficiently. Our approach will start with analytically defined neural networks, which precisely replicate matched filtering, and then train these networks to optimize their statistical performance. We will give training approaches that are compatible with two classical frameworks for formalizing the performance decision rules δ : the *Neyman-Pearson* framework, in which the parameter γ is a random vector with known distribution ν , and the *minimax* framework, in which we control the worst performance over all possible choices of the parameter γ .

A. Neyman-Pearson Framework

In this setting, one assumes that γ is a random vector with probability distribution ν . With this distribution ν ,

we can then view H_1 as a simple hypothesis. The false positive rate (FPR) associated with the rule δ is

$$\text{FPR} = \mathbb{P}_{\mathbf{z}} [\delta(\mathbf{z}) = 1] \quad (4)$$

The false negative rate (FNR) at signal $\mathbf{s}_\gamma$ is

$$\text{FNR}_\gamma = \mathbb{P}_{\mathbf{z}} [\delta(\mathbf{s}_\gamma + \mathbf{z}) = 0]. \quad (5)$$

The *overall* false negative rate is

$$\text{FNR} = \int \text{FNR}_\gamma d\nu(\gamma). \quad (6)$$

The Neyman-Pearson criterion seeks the optimal tradeoff between FNR and FPR:

$$\min_{\delta} \text{FNR} \quad \text{subject to} \quad \text{FPR} \leq \alpha, \quad (7)$$

where α is a user-specified significance level.

There is a classical closed form expression for the optimal test under the Neyman-Pearson criterion: if ρ_0 and ρ_1 are the probability densities of the signal $\mathbf{x}$ under hypotheses H_0 and H_1 , respectively, then the optimal test is given by comparing the *likelihood ratio*

$$\lambda(\mathbf{x}) = \frac{\rho_1(\mathbf{x})}{\rho_0(\mathbf{x})} \quad (8)$$

to a threshold τ , which depends on the significance level α . An illustration of an example problem is shown in FIG 1

B. Minimax Framework

When a good prior ν is not available or cannot be assumed, we can instead seek a decision rule that solves

$$\min \text{WFNR} \quad \text{subject to} \quad \text{FPR} \leq \alpha. \quad (9)$$

at a given false positive rate, where WFNR is the *worst false negative rate* defined as

$$\text{WFNR} = \max_{\gamma \in \Gamma} \text{FNR}_\gamma. \quad (10)$$

In contrast to the Neyman-Pearson criterion, there is in general no simple expression for the minimax optimal rule δ [96]. In the next section, we will review matched filtering, a simple, popular approach to detection which is compatible with the minimax framework (albeit suboptimal in terms of [9]), in the sense that it does not require a prior on γ .

III. MATCHED FILTERING FOR PARAMETRIC DETECTION

Matched filtering is a powerful classical approach to signal detection, which applies a linear filter which is chosen to maximize the signal-to-noise ratio (SNR).

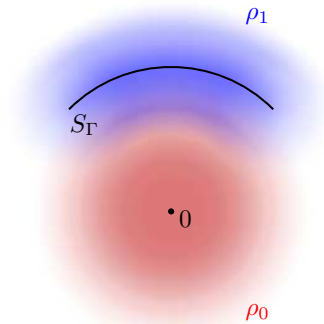


FIG. 1. An example of the parametric signal detection problem with signal space S_Γ . Densities ρ_0 and ρ_1 are shown in red and blue respectively.

A. Optimality for Single Signal Detection

In the simplest possible setting, in which (i) there is only one target signal $\mathbf{s}$, (ii) the observation $\mathbf{x}$ has the same length as $\mathbf{s}$, and (iii) the noise is uncorrelated (i.e., $\mathbb{E}[\mathbf{z}\mathbf{z}^*] = \sigma^2 \mathbf{I}$), matched filtering simply computes the inner product between the target $\mathbf{s}$ and the observation:

$$\delta(\mathbf{x}) = 1 \text{ iff } \langle \mathbf{s}, \mathbf{x} \rangle \geq \tau. \quad (11)$$

When detecting a single signal $\mathbf{s}$ in iid Gaussian noise, this decision rule is optimal in both the Neyman-Pearson and minimax senses: for example, if $\mathbf{z} \sim \mathcal{N}(\mathbf{0}, \sigma^2 \mathbf{I})$, the likelihood ratio

$$\lambda(\mathbf{x}) = \frac{\rho_0(\mathbf{x} - \mathbf{s})}{\rho_0(\mathbf{x})} = \exp \left(\frac{\langle \mathbf{s}, \mathbf{x} \rangle - \|\mathbf{s}\|^2/2}{\sigma^2} \right) \quad (12)$$

is a monotone function of $\langle \mathbf{s}, \mathbf{x} \rangle$, and so matched filtering implements the (optimal) likelihood ratio test. FIG 2 illustrates this optimality geometrically.

The simplicity and optimality in this setting make matched filtering a principled choice for signal detection, and have inspired its application in settings that go far beyond the scope of this rigorous guarantee. In particular, the simplest and most practical extension of this rule to detecting parametric families of signals $\mathbf{s}_\gamma$ is suboptimal in both the Neyman-Pearson and minimax settings. Moreover, there are a number of additional factors which contribute to its suboptimality. These include unknown, non-Gaussian, and possibly time-varying noise distributions as well as density and coverage issues in the template bank, which for complexity reasons may cover only a small portion of the phase space [22]. Nevertheless, we will see how matched filtering can inspire principled approaches to deriving more flexible decision rules which can address many of these challenges.

B. Extensions to Parametric Detection

The simplest extension of the decision rule (11) to *parametric* detection problems, in which there are multiple potential targets $\mathbf{s}_\gamma$, involves taking the maximum over the parameter space:

$$\delta(\mathbf{x}) = 1 \text{ iff } \max_{\gamma \in \Gamma} \langle \mathbf{s}_\gamma, \mathbf{x} \rangle \geq \tau. \quad (13)$$

Here we used the assumption that all templates have unit norm, namely $\|\mathbf{s}_\gamma\|_2^2 = 1, \forall \gamma \in \Gamma$. When this rule (13) is hard to implement in exact form, it can typically be approximated by taking samples $\mathbf{s}_{\gamma_1}, \dots, \mathbf{s}_{\gamma_k}$ and setting

$$\delta(\mathbf{x}) = 1 \text{ iff } \max_{i=1, \dots, k} \langle \mathbf{s}_{\gamma_i}, \mathbf{x} \rangle \geq \tau. \quad (14)$$

When the sampling is sufficiently dense, the sampled matched filter rule (14) accurately approximates the ideal matched filter rule (13) [22]. This rule, while simple, is an important component of many sophisticated data analysis pipelines, including LIGO, Virgo and KARGA's template based searches for compact binary coalescence signals.

Note that the matched filtering decision rule (13) has connections to the (generalized) likelihood ratio test, where H_1 is the composite hypothesis $\mathbf{s}_\gamma \in S_\Gamma$. While this test has nice statistical properties, it is not guaranteed to be the uniformly most powerful test when the hypotheses are composite. For the rest of this paper, the term “likelihood ratio test” will be reserved for the test with a given prior and simple hypotheses, which satisfies the Neyman-Pearson criterion.

In contrast to the single signal setting, the simple extensions (13)-(14) of matched filtering to detecting parametric families of signals are not optimal: in the Neyman-Pearson setting, they do not achieve the minimal FNR for a given FPR, while in the minimax setting, they do not achieve the minimal WFNR for a given FPR.

The suboptimality of (13)-(14) under Neyman-Pearson can be observed by noting that the decision statistic $\max_\gamma \langle \mathbf{s}_\gamma, \mathbf{x} \rangle$ is not a monotone function of the likelihood ratio, which in iid Gaussian noise for example, takes the form

$$\lambda(\mathbf{x}) = \int \exp\left(\frac{\langle \mathbf{s}_\gamma, \mathbf{x} \rangle - \|\mathbf{s}_\gamma\|^2/2}{\sigma^2}\right) d\nu(\gamma). \quad (15)$$

FIG 3 and 4 illustrate such suboptimality for a particular problem configuration in $\mathbb{R}^2$. Note that throughout our paper, we will slightly abuse the term of receiver operating characteristic (ROC) curves by plotting FNR against FPR, instead of the convention of plotting FPR against the true positive rate $\text{TPR} \equiv 1 - \text{FNR}$. This highlights the connection to the notion of error rates in machine learning, and more importantly facilitates demonstration of the curves and axis ranges at very low error rates.

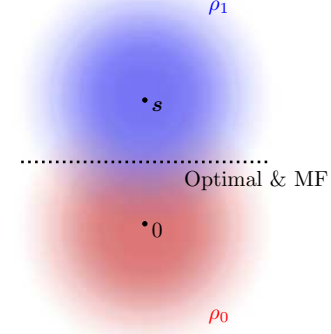


FIG. 2. Optimality of matched filtering in single signal detection.

It is, in a sense, unsurprising that matched filtering is suboptimal in this setting, since the decision rules (13)-(14) do not make use of the prior ν , while the likelihood ratio test assumes (and uses) this prior.

However, the matched filtering rule (13)-(14) is also in general suboptimal in the “prior-free” minimax setting. Consider the scenario in FIG 5 as an example, where the signal space $S_\Gamma \subset \mathbb{R}^2$ consists of only two signals $\mathbf{s}_1 = [1, 0]^T$ and $\mathbf{s}_2 = [0, 1]^T$. Comparing the prior-free matched filtering decision rule δ_{MF} with the optimal decision rule δ_* under the Neyman-Pearson framework with uniform prior over the two signals, we see that δ_{MF} is suboptimal under Neyman-Pearson criterion with uniform prior. Moreover, from symmetry it follows that for symmetric decision rules such as δ_{MF} and δ_* the worst FNR and the overall FNR are equal. This implies that δ_{MF} is also worse than δ_* under the minimax criterion.

We also note that this suboptimality is, in some sense, not because we don't have sufficient templates. In the example shown in FIG 5 the matched filtering model already covers the entire signal set which consists of two signals. Furthermore, we will see in the later discussions that matched filtering has other structural limitations when working with non-Gaussian noise distributions.

IV. FROM MATCHED FILTERING TO NEURAL NETWORKS

Since the matched filtering rule (14) is suboptimal for parametric detection, we will show that (i) the form of this rule suggests approaches to learning optimal rules for parametric detection, and (ii) the resulting classifiers have additional advantages, including greater flexibility and lower computational/storage complexity or cost. Our approach is driven by the observation: the matched filtering rule (14) is equivalent to a feedforward neural network.

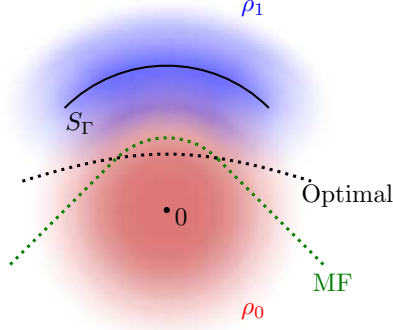


FIG. 3. Suboptimality of matched filtering under the Neyman-Pearson framework.

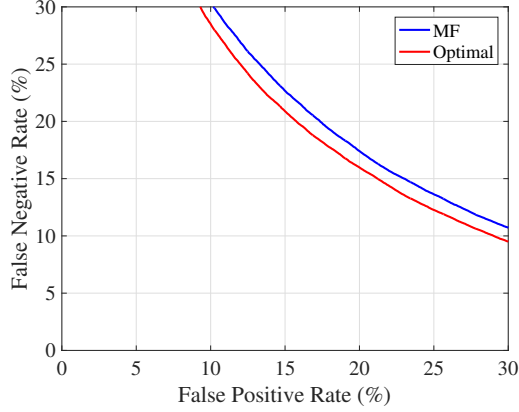


FIG. 4. Comparison of ROC curves of the optimal classifier and matched filtering in the 2-dimensional concept as illustrated in FIG 3

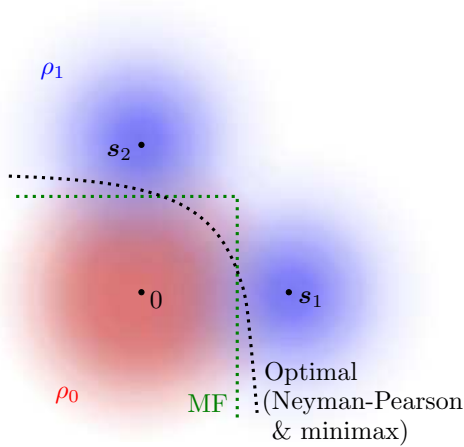


FIG. 5. Suboptimality of matched filtering under the minimax framework.

A. Neural Networks: Notation and Basics

A *neural network* implements a mapping from the signal space $\mathbb{R}^n$ to an output space $\mathbb{R}^d$:

$$f_{\theta} : \mathbb{R}^n \rightarrow \mathbb{R}^d. \quad (16)$$

Here, θ represents the parameters of the network. Specifically, a fully connected neural network can be written as a composition of layers, each of which applies an affine mapping

$$x \mapsto Wx + b \quad (17)$$

followed by an elementwise activation function ϕ :

$$f_{\theta}(x) = W^L \phi \left(W^{L-1} \phi \left(\dots \phi \left(W^1 x + b^1 \right) \dots \right) + b^{L-1} \right) + b^L. \quad (18)$$

With slight abuse of notation, the activation function $\phi : \mathbb{R} \rightarrow \mathbb{R}$ acts elementwise when applied to a vector:

$$\phi([v_1, \dots, v_n]^T) = [\phi(v_1), \dots, \phi(v_n)]^T. \quad (19)$$

The intermediate products

$$\alpha^{\ell}(x) = \phi \left(W^{\ell} \phi \left(\dots \phi \left(W^1 x + b^1 \right) \dots + b^{\ell} \right) \right) \quad (20)$$

are sometimes referred to as *features* [97]. In many situations, it is useful to “pool” features – this is especially useful for data with spatial or temporal structure; combining spatially adjacent features in a nonlinear fashion renders the decision more stable with respect to deformations of the input [98]. For example, *maximum pooling* takes the maximum of adjacent features. In our notation, we can denote this operation by ρ^{ℓ} and write

$$\alpha^{\ell}(x) = \rho^{\ell} \phi \left(W^{\ell} \alpha^{\ell-1}(x) + b^{\ell} \right), \quad (21)$$

where the concise notation ρ^{ℓ} suppresses certain details about which features are combined. For clarity, we summarize this discussion in the following mathematical definition:

Definition 1 (Fully connected neural network). A *fully connected neural network (FCNN)* with **feature dimensions** $n^0, \dots, n^L$, **pre-activation dimensions** $m^1, \dots, m^L$, **parameters**

$$\theta = \left(W^L \in \mathbb{R}^{m^L \times n^{L-1}}, \dots, W^1 \in \mathbb{R}^{m^1 \times n^0}, b^L \in \mathbb{R}^{m^L}, \dots, b^1 \in \mathbb{R}^{m^1} \right), \quad (22)$$

activation function $\phi : \mathbb{R} \rightarrow \mathbb{R}$ (extended to vector inputs by applying it elementwise), and **pooling operations** $\rho^{\ell} : \mathbb{R}^{m^{\ell} \times n^{\ell}}$ given by

$$[\rho^{\ell}]_i(v) = \max_{j \in I_i^{\ell}} v_j, \quad (23)$$

with $I_1^\ell, \dots, I_{n^\ell}^\ell$ being disjoint subsets of $[m^\ell]$, is a mapping $f_\theta : \mathbb{R}^n \rightarrow \mathbb{R}^d$ defined inductively as $f_\theta(\mathbf{x}) = \alpha^L(\mathbf{x})$ by setting $\alpha^0(\mathbf{x}) = \mathbf{x}$, and

$$\alpha^\ell(\mathbf{x}) = \rho^\ell \phi(\mathbf{W}^\ell \alpha^{\ell-1}(\mathbf{x}) + \mathbf{b}^\ell), \quad \ell = 1, \dots, L. \quad (24)$$

When discussing neural networks, it is conventional to distinguish between the *network architecture*, which consists of the choices of feature dimensions n^ℓ, m^ℓ , activation function ϕ , and pooling operators ρ^ℓ , and the *network parameters* θ . Although we have stated a general definition, in specific architectures, the activation function ϕ and/or the pooling operators ρ^ℓ can be chosen to be trivial ($\phi(t) = t$ and/or $\rho^\ell(\mathbf{v}) = \mathbf{v}$).

Architectures. Neural networks are flexible function approximators [99]: universal approximation theorems indicate that *nonlinear* neural networks (with non-polynomial activation ϕ) can accurately approximate any continuous function, as long as the network is sufficiently deep and/or wide [100–102]. There is a growing body of empirical and theoretical evidence showing that (relatively small) neural networks can learn relatively smooth functions over low-dimensional submanifolds of $\mathbb{R}^n$ with a complexity that is proportional to the manifold dimension, which in our problem equals the number of parameters in the parameterization $\gamma \mapsto \mathbf{s}_\gamma$ [103].

Beyond these general considerations, there are scenarios in which the nature of the task dictates specific architectural choices. For example, in the field of inverse problems, neural network architectures can be generated by interpreting various optimization methods as taking on the structure in Definition 1 [104]. Our proposals will have a similar spirit, since they will interpret an existing method (matched filtering) as a particular instance of Definition 1.

Finally, a major architectural choice is whether to enforce additional structure on the matrices $\mathbf{W}^\ell$. When the input $\mathbf{x}$ is a time series, it is natural to structure the linear maps $\alpha \mapsto \mathbf{W}\alpha$ to be time-invariant, i.e., to be convolution operators. To exhibit the equivalence between matched filtering and neural networks in the simplest possible setting, here we train our networks on injections whose starting time is fixed, and focus on fully connected neural networks (not enforcing convolutional structure).

In deployment, the input data is a time series, and astrophysical signals can occur at any time. In this setting, the matched filtering rule is applied in a sliding fashion. Similarly, the neural networks proposed here can be also deployed in a sliding fashion, which effectively converts them to particular convolutional networks. Both the equivalence between matched filtering and particular neural networks and the potential advantages of neural networks carry over to this setting.

Parameters. There are various approaches to choosing the network parameters θ . The dominant approach is to learn these parameters by optimization on data: one chooses initial parameters at random (with appropriate variance to ensure stability), and then iteratively adjusts

them to best fit a given set of “training data”. However, it is also possible in some scenarios to either (i) simply choose the weights at random, or (ii) to generate the weights analytically, either by connecting the network architecture to existing structures/algorithms [104] or from harmonic analysis considerations [105]. There are approaches that lie in between purely data-driven and purely analytical approaches to choosing θ . For example, it is possible generate initial weights analytically, and then tune them on training data. This hybrid approach achieves excellent performance on a number of inverse problems in imaging (super-resolution [106], magnetic resonance image reconstruction [107] etc.).

In the following sections, we will follow this approach: we will give two ways of interpreting the matched filtering decision rule (14) as a fully connected neural network, by making specific (analytical) choices of the architecture and parameters. These analytically chosen parameters can then be used as an initialization for learning on data. We will also see that in addition to this closed-form construction for equivalence, neural network models can be further trained on data to achieve improved performance.

B. Matched Filtering as a Shallow Neural Network

In the language of the previous section, it is not hard to express the decision statistic (14) of matched filtering as a specific fully connected neural network with one layer ($L = 1$). Writing

$$\rho^1(\mathbf{z}) = \max_i z_i, \quad (25)$$

$$\phi(t) = t, \quad (26)$$

$$\mathbf{W}^1 = \begin{bmatrix} \mathbf{s}_{\gamma_1}^* \\ \mathbf{s}_{\gamma_2}^* \\ \vdots \\ \mathbf{s}_{\gamma_k}^* \end{bmatrix} \in \mathbb{R}^{k \times n}, \quad (27)$$

$$\mathbf{b}^1 = \mathbf{0}, \quad (28)$$

(108) we have

$$\max_i \langle \mathbf{s}_{\gamma_i}, \mathbf{x} \rangle = \rho^1 \phi(\mathbf{W}^1 \mathbf{x} + \mathbf{b}^1). \quad (29)$$

In words, the features produced by this neural network correspond to the correlations of the input with the templates $\mathbf{s}_{\gamma_1}, \dots, \mathbf{s}_{\gamma_k}$. FIG 6 illustrates this (simple) architecture, which we label **MNet-Shallow**.

Where needed below, we refer to the input-output relationship implemented by this architecture as

$$f_{\text{MNet-Shallow}, \theta}(\mathbf{x}), \quad (30)$$

where $\theta = (\mathbf{W}^1, \mathbf{b}^1)$ represent the weights and biases. When these are chosen as in (27)–(28), **MNet-Shallow** implements the matched filtering decision rule. We note that these weights can be constructed analytically based on the given templates.

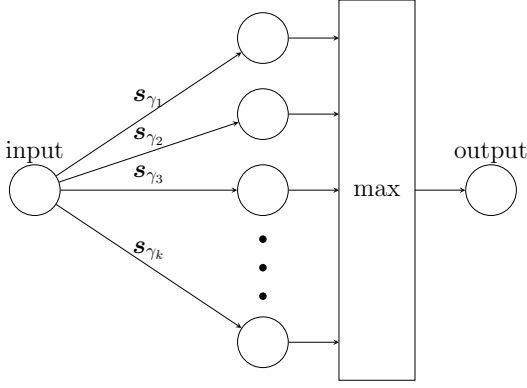


FIG. 6. Illustration of **MNet-Shallow**. Bias terms are omitted in the illustration. (We note that for more complex networks arbitrary pooling operations can replace the “max” box.)

By learning the weights $\mathbf{W}^1$ and biases $\mathbf{b}^1$ from examples, we can further adapt this network to implement a more general family of decision rules, beyond matched filtering [14] with templates $\mathbf{s}_\gamma$. Nevertheless, there are limitations to this architecture. Notice that in **MNet-Shallow** there is only one layer of affine operations, and so this architecture does not satisfy the dictates of the universal approximation theorem [101, 109].

More geometrically, we can notice that the decision rule associated with **MNet-Shallow** is a maximum of affine functions. This means that for any choice of $\mathbf{W}^0$ and $\mathbf{b}^0$, the decision boundary is the boundary of a convex set. This property is also true for matched filtering, which shares exactly the same form. An illustration of this property is shown in FIG 7

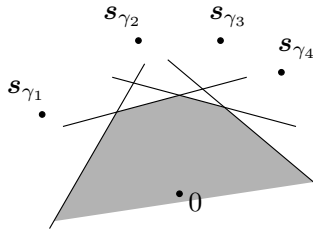


FIG. 7. The set of points classified as noise by matched filtering and **MNet-Shallow** is always a convex set.

How restrictive is this limitation? In the context of parametric detection, this depends largely on the noise distribution. If the noise is Gaussian, the optimal decision boundary is itself the boundary of a convex set:

Proposition 1. *Suppose that the noise $\mathbf{z} \sim \mathcal{N}(\mathbf{0}, \sigma^2 \mathbf{I})$. Then for any significance level α , the optimal (Neyman-Pearson) decision region*

$$\{\mathbf{x} \mid \lambda(\mathbf{x}) \leq \tau\} \quad (31)$$

is a convex subset of $\mathbb{R}^n$, where τ is a constant determined by the significance level α .

Proof. Please see Appendix. $\square$

However, for general non-Gaussian distributions, the optimal decision region is often nonconvex. We illustrate this result in FIG 8. In fact, this suggests an intrinsic structural limitation of matched filtering and similar architectures. Since in reality the noise distribution is not perfectly Gaussian, we cannot expect the optimal decision region to be convex, and hence the matched filtering structure is unable to approach the performance of the likelihood ratio test with arbitrary precision, even if any number of templates (including ones outside the original signal space) are allowed. In such cases, we can benefit from using a more flexible architecture, which we now introduce.

C. Matched Filtering as a Deep Neural Network

We describe an alternative way of expressing template matching as a neural network, which leads to deep, nonlinear architectures that are more flexible than **MNet-Shallow**. We label this structure **MNet-Deep**. In this architecture, we do not compute the maximum in a straightforward way using pooling. Instead, we propose an alternative architecture which is more flexible, and can approximate a wider class of functions. In particular, we will no longer be restricted to implementing decision boundaries that are boundaries of convex sets, allowing us to handle scenarios with non-Gaussian noise. An illustration of this **MNet-Deep** is shown in FIG 9

Our construction is based on the rectified linear unit (ReLU) nonlinearity:

$$\phi(t) = \max(t, 0). \quad (32)$$

This is arguably the most commonly used nonlinearity function in modern deep learning.

The matched filtering decision rule takes the maximum of a family of linear functions $\langle \mathbf{s}_{\gamma_i}, \mathbf{x} \rangle$. Instead of simply “pooling” these functions as in the previous section, we implement the maximum operation using compositions of ReLUs and linear operations. In particular, observe that the maximum of two numbers can be written as a linear combination of 3 ReLU units:

$$\max(a, b) = b + \phi(a - b) = \phi(b) - \phi(-b) + \phi(a - b). \quad (33)$$

The basic idea is to create a hierarchical structure of such 3-ReLU-units, each of which takes a pairwise maximum of its inputs. Our **MNet-Deep** construction will perform convolutions with the templates $\mathbf{s}_{\gamma_i}$, followed by this hierarchical structure for computing the maximum.

FIG 10 illustrates this hierarchical structure for the particular example of four inputs. The network in FIG 10

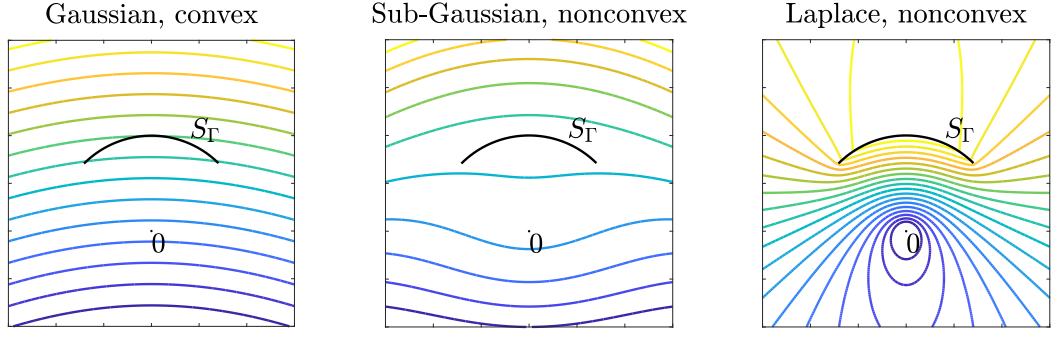


FIG. 8. Contours of log likelihood ratio with various noise distributions, and whether the optimal decision regions with $\delta = 0$ is always convex. Yellow represents larger values and blue represents lower values. From left to right: (1) Gaussian distribution, convex; (2) Sub-Gaussian distribution $\rho_{\text{noise}}(\mathbf{x}) \propto \exp(-C\|\mathbf{x}\|^3)$, not necessarily convex; (3) Laplace distribution, not necessarily convex.

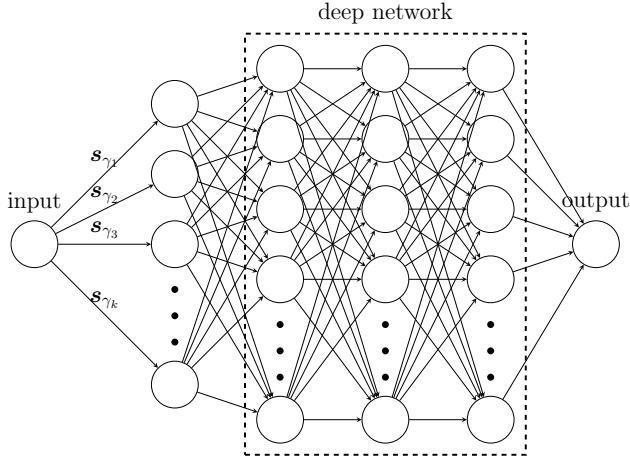


FIG. 9. Illustration of MNet-Deep. Bias terms are omitted in the illustration. This network structure is obtained by replacing the max module in matched filtering (as in FIG 6) with a deep network.

can be expressed as a ReLU network, with sparse weight matrices $\mathbf{W}^\ell$ ($\ell = 0, 1, 2$) for the layers respectively:

$$\mathbf{W}^0 = \begin{bmatrix} 0 & 1 \\ 0 & -1 \\ 1 & -1 \end{bmatrix}, \quad \mathbf{W}^2 = \begin{bmatrix} 1 & -1 & 1 \end{bmatrix}, \quad (34)$$

$$\mathbf{W}^1 = \mathbf{W}^0 \otimes \mathbf{W}^2 = \begin{bmatrix} 0 & 0 & 0 & 1 & -1 & 1 \\ 0 & 0 & 0 & -1 & 1 & -1 \\ 1 & -1 & 1 & -1 & 1 & -1 \end{bmatrix}. \quad (35)$$

Generalizing this construction, we obtain a network that takes the maximum of k numbers, using $\lceil \log_2 k \rceil + 1$ layers.

While the example above delineates a precise form of the ReLU network, this approach can in fact be made flexible. To ensure that the network output is indeed the maximum of the k inputs, we must ensure that at each layer, each feature participates in at least one of

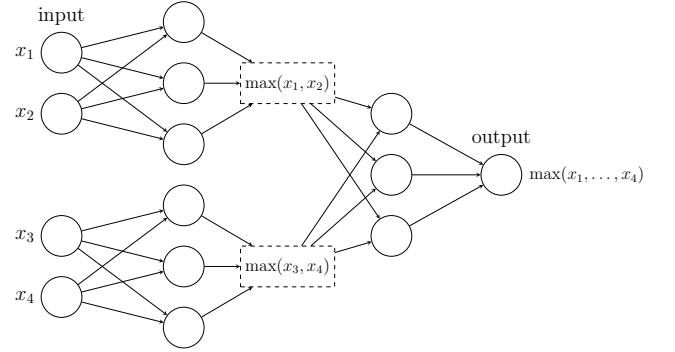


FIG. 10. Illustration of implementing max with a ReLU network. The dashed boxes in the middle are not actual nodes in the network, but “imaginary” nodes to facilitate construction.

the pairwise max operations. This means that at layer ℓ , we must have at least $k/2^\ell$ features. However, we are free to add more intermediate features, with additional (redundant) max operations. This does not change the output of the network, but it affords additional flexibility when we attempt to train the network on data. In particular, this allows the construction of arbitrarily wide or deep ReLU networks, and can therefore approximate any regular continuous function [101, 109].

There is also a degree of freedom in choosing which features participate in each pairwise maximum operation, which could be chosen in various ways. In our implementation we use the following way to pair up the nodes in layer l for pairwise maximum operations that get to layer $l + 1$. Assume layer l contains $2p$ nodes. First pair up the nodes with consecutive indices, namely pair up node $2i - 1$ with node $2i$ for $i = 1, \dots, p$. This ensures that each node is covered by at least one maximum operation. After that, for each leftover node in layer $l + 1$, we establish the corresponding pair in layer l by choosing the nodes at random in layer l . In the following, we label this network MNet-Deep. We emphasize for clarity that the

nodes between consecutive layers are fully connected in the neural network; however, the weights not associated with pairwise maximum operations are all initialized to zero. Below, where needed we refer to the decision rule associated with this network as

$$f_{\text{MNet-Deep}, \boldsymbol{\theta}}(\mathbf{x}), \quad (36)$$

where $\boldsymbol{\theta}$ represent the collection of all weights and biases. The above discussion again gives a recipe for choosing these weights analytically such that the decision rule for MNet-Deep coincides with the matched filtering rule.

In contrast to MNet-Shallow, MNet-Deep is a more flexible architecture. In particular, this architecture satisfies the dictates of the universal approximation theorem. Geometrically, it is not restricted to convex decision regions, which makes it capable of achieving optimal decision boundaries even when the noise is heavy-tailed or has other non-ideal properties.

D. Equivalence of Matched Filtering and Neural Networks

We have demonstrated by construction the following claim:

Given any collection of templates $\mathbf{s}_{\gamma_1}, \dots, \mathbf{s}_{\gamma_k}$ (for any $k \geq 1$), one can analytically determine weights $\boldsymbol{\theta}_s, \boldsymbol{\theta}_d$ such that

$$f_{\text{MNet-Shallow}, \boldsymbol{\theta}_s}(\mathbf{x}) = \max_{i=1 \dots k} \langle \mathbf{s}_{\gamma_i}, \mathbf{x} \rangle \quad (37)$$

$$f_{\text{MNet-Deep}, \boldsymbol{\theta}_d}(\mathbf{x}) = \max_{i=1 \dots k} \langle \mathbf{s}_{\gamma_i}, \mathbf{x} \rangle \quad (38)$$

for all $\mathbf{x} \in \mathbb{R}^n$.

We emphasize the complete generality of this claim: it holds for any number and choice of templates. Moreover, it does not depend on training: the networks can be constructed analytically to implement the matched filtering rule. Nevertheless, we will see in the next section that they can be further adapted based on observed data to strictly outperform matched filtering, in terms of the Neyman-Pearson criterion.

The equivalence between matched filtering and particular neural networks has an additional conceptual advantage: it allows for a clear comparison of the resource complexity of different search methods, in terms of storage and computation. This is valuable because different methods may cut out very different tradeoffs between complexity and accuracy/performance. Neural network implementations of matched filtering can be viewed as “complexity standard candles” against which the performance of more sophisticated networks can be measured. In particular, the complexity of a neural network model may be quantified by the total number of nodes (neurons) in the network, which approximately characterizes

the number of elementary operations performed for evaluating an input instance [110, 111]. We will look for the most appropriate measure of complexity for this problem, and provide detailed analysis in future studies.

V. TRAINING TO APPROACH STATISTICAL OPTIMALITY

In the previous section, we gave two ways of analytically constructing neural networks that reproduce the matched filtering decision rule, and hence exhibit exactly the same performance as matched filtering. The major advantage of this interpretation of matched filtering is that the resulting model can be further trained on sample data to improve its statistical performance or adapt it to handle non-Gaussian noise distributions, or in other words “standing on the shoulder of giants”. In a typical neural network training problem, we have access to labelled samples

$$(\mathbf{x}_1, y_1), \dots, (\mathbf{x}_N, y_N), \quad (39)$$

each of which consists of an observation $\mathbf{x}_i \in \mathbb{R}^n$ and a corresponding label $y_i \in \{0, 1\}$, which indicates whether $\mathbf{x}_i$ contains a noisy signal ($y_i = 1$) or noise only ($y_i = 0$). To date, we have only a moderate number of confirmed gravitational wave detections, and hence have far more negative examples than positive examples. We address this issue by generating our positive training examples by injecting synthetic waveforms into (real) LIGO noise strains. Below, we describe two different training schemes, motivated by the Neyman-Pearson and min-max criteria, which leverage this data to perform training of the neural networks.

Training for Neyman-Pearson. In this setting, we assume that the prior ν is known, and generate positive examples by first sampling $\gamma_i \sim \nu$, and setting $\mathbf{x}_i = \mathbf{s}_{\gamma_i} + \mathbf{z}_i$, where $\mathbf{z}_i$ is observed LIGO noise strain. We solve the following optimization problem:

$$\min_{\boldsymbol{\theta}} \mathcal{R}_N(f_{\boldsymbol{\theta}}) := \frac{1}{N} \sum_{i=1}^N \ell(f_{\boldsymbol{\theta}}(\mathbf{x}_i), y_i). \quad (40)$$

Here, the *loss function* $\ell(\hat{y}, y)$ measures the misfit between the predicted label $\hat{y}$ and the true label y . Typical choices include the square loss $(\hat{y} - y)^2$ and the logistic loss

$$y \log(f_{\text{sigmoid}}(\hat{y})) + (1 - y) \log(1 - f_{\text{sigmoid}}(\hat{y})). \quad (41)$$

where $f_{\text{sigmoid}}(\cdot)$ denotes the logistic/sigmoid function:

$$f_{\text{sigmoid}}(x) = \frac{1}{1 + \exp(-x)} \quad (42)$$

Is this training strategy compatible with the Neyman-Pearson criterion? The following proposition answers this question in the affirmative. Consider the following

setup: training data $(\mathbf{x}_i, y_i)$ are generated independently at random, by setting $y_i = 1$ with probability $p \in (0, 1)$ and choosing $\mathbf{x}_i = \mathbf{s}_{\gamma_i} + \mathbf{z}_i$ when $y_i = 1$ and $\mathbf{x}_i = \mathbf{z}_i$ when $y_i = 0$, with $\gamma_i \sim \nu$, and $\mathbf{z}_i \sim \rho_{\text{noise}}$. Let

$$\mathcal{R}_\infty(f) = \mathbb{E}_{(\mathbf{x}, y)} \ell(f(\mathbf{x}), y). \quad (43)$$

This represents the large-sample limit of $\mathcal{R}_N$: as $N \rightarrow \infty$, $\mathcal{R}_N(f) \rightarrow \mathcal{R}_\infty(f)$. The following proposition shows that the population risk $\mathcal{R}_\infty$ is minimized by (a monotone function of) the likelihood ratio λ :

Proposition 2. *Suppose that for any $y = 0, 1$, the loss $\ell(\hat{y}, y)$ is a strictly convex differentiable function of $\hat{y}$ that is minimized at $\hat{y} = y$. [112] Then the unique optimal solution $f_\star$ to the (functional) optimization problem*

$$\min_f \mathcal{R}_\infty(f) \quad (44)$$

is a strictly increasing function of the likelihood ratio λ :

$$f_\star(\mathbf{x}) = g(\lambda(\mathbf{x})), \quad (45)$$

where g is a strictly increasing function that depends on ℓ .

Proof. Please see Appendix. $\square$

This result can be interpreted as saying: “a sufficiently flexible classifier, trained on a sufficiently large dataset will produce the optimal decision rule.” Hence, training to minimize the empirical risk $\mathcal{R}_N(f_\theta)$ is compatible with the Neyman-Pearson criterion.

While this is a promising observation, we should keep in mind a number of remaining issues: How much data is required? What are effective approaches to minimizing the empirical risk $\mathcal{R}_N$? In the next section we investigate these questions experimentally.

Training for Minimax. In this setting, we do not assume any prior, and aim to minimize the worst false negative rate using the formulation in [9]. We convert the constrained problem [9] to an equivalent unconstrained problem,

$$\min_\delta \max_{\gamma \in \Gamma} \text{FNR}_\gamma + c \cdot \text{FPR}, \quad (46)$$

where c is a constant that depends on α . For tractability, we will fix c at a constant value to obtain a concrete optimization objective, and here we fix $c = 1$. In actual deployment where a target significance level α is specified, we can also choose c at the level that corresponds to the specified α . Also, we sample the parameter space Γ at points $\{\gamma_i\}_{i=1}^N$. Since FPR does not depend on γ , it can be moved inside the maximization. Therefore, the minimax optimization problem can be transformed into

$$\min_\delta \max_{i=1, \dots, N} (\text{FNR}_{\gamma_i} + \text{FPR}). \quad (47)$$

This suggests a natural approach to training under the minimax criterion using first-order optimization methods. At each iteration, we estimate FPR and FNR_{γ_i} for

each $i = 1, \dots, N$, and choose $i_\star$ with the highest $\text{FNR}_{\gamma_{i_\star}}$. We then aim to reduce $\text{FNR}_{\gamma_{i_\star}} + \text{FPR}$, which can be estimated by using a sample dataset $\{(\mathbf{x}_i, y_i)\}_{i=1}^N$ as

$$\frac{1}{N} \sum_{i=1}^N \mathbf{1}[f_\theta(\mathbf{x}_i) \neq y_i], \quad (48)$$

where in the dataset all $\mathbf{x}_i$ with corresponding $y_i = 1$ are generated specifically with signal parameter $\gamma_{i_\star}$, and half of data pairs in the dataset have $y_i = 0$. Finally, it is customary in optimization to replace the non-differentiable 0-1 loss with a smooth loss function ℓ , and hence we get the following risk minimization objective:

$$\frac{1}{N} \sum_{i=1}^N \ell(f_\theta(\mathbf{x}_i), y_i). \quad (49)$$

This expression is similar to [40], but the difference is that all positive data in the dataset here are associated with signal parameters $\gamma_{i_\star}$.

VI. SIMULATIONS AND EXPERIMENTS

A. Data Generation

Data-driven methods such as neural networks typically require a large amount of data for training. The question of data sufficiency is especially acute in gravitational wave astronomy: we have only a moderate number of confirmed detections to date. We address this issue by generating our positive training examples by injecting synthetic waveforms into LIGO noise strains [92], which we elaborate below.

For LIGO noise data, we use the L1 strain from LIGO O2 run between August 1 and August 25, 2017, with ANALYSIS_READY segments only. The announced confident detections GW170809, GW170814, GW170817, GW170818 and GW170823 are removed from the strain, such that the data is at least 300 seconds away from these events. We used a total of 338 frame files each of 4096 seconds long, namely a total of 384.57 hours. The strain data is downsampled from the original 4096Hz to 2048Hz for processing efficiency. The downsampled L1 strain data is divided into segments of length 0.6 second, with each successive segment overlapping 50% of the previous segment.

We generate synthetic gravitational wave signals using PyCBC [26, 32], with the following parameters. *Approximant:* IMRPhenomD. *Mass range:* 40 to 50 $M_\odot$, uniformly distributed. *Spin:* 0. *Sampling rate:* 2048Hz. *Low frequency cutoff:* 30Hz. *Coalescence phase:* 0. *Polarization:* plus [113]. With this specified mass range, at least 99.5% of the energy of the signal lies in an interval of length 0.3 second after preprocessing. We note that although the templates are not chosen uniformly in actual LIGO deployment [42, 43, 114–116], we make this

choice here due to simplicity, and also the fact that the large number of templates make up for the possibly sub-optimal choice of templates.

The above data is used to generate training and test datasets of positive and negative labelled data as follows. We divide the collection of downsampled strain segments randomly into training and test sets, ensuring that no training segment overlaps a test segment. Within the training and test sets, we generate both positive and negative examples. The negative examples contain only the strain data. For the positive examples, we inject waveforms into the noise segments by aligning the peak of the waveforms at the 90% location of the center 0.3s, namely at the location of 0.42s within the entire segment of 0.6s. This choice was made as it safely covers the injected waveforms. The amplitude of the injection is set such that after filtering and whitening (to be described below), the resulting signal-to-noise ratio (SNR) is constant. For the experiment, the size of the training and test datasets are respectively 2.62 million and 2 million segments.

We preprocess all training and test data, by applying an FIR bandpass filter with cutoff frequencies 30Hz and 400Hz, whitening using a power spectral density estimated from the L1 strain data, and finally truncating to keep only the center 0.3 second (614 samples).

B. Matched Filtering Configuration

We first need to determine the necessary number of templates to use in matched filtering, given the space of parameters. We set 10, 100, 1000 and 10000 as the candidate numbers of templates. For each candidate number, we independently repeat the following process 30 times: randomly choose the specified number of pairs of parameters uniformly from $[40, 50] \times [40, 50]$, generate waveforms according to these parameters, preprocess (bandpass, whiten and truncate) as described above, and then normalize to equal power. This produces the templates for a matched filtering model. We evaluate the model on the test dataset to obtain an ROC curve. For each candidate number of templates and for each value of FPR, we take the lowest FNR outcome among the 30 independent runs. This is used to approximately represent the best performance achievable with a given number of templates.

The result is shown in FIG 11. We see that the best performance of matched filtering in this setting starts to saturate at approximately 1000 templates, and the best performance with 1000 templates is almost identical to the that with 10000 templates. Therefore, we choose the best performance of matched filtering with 10000 templates, namely the bright blue curve, as the performance curve of the matched filtering method in this setting, against which we will be comparing our neural network method.

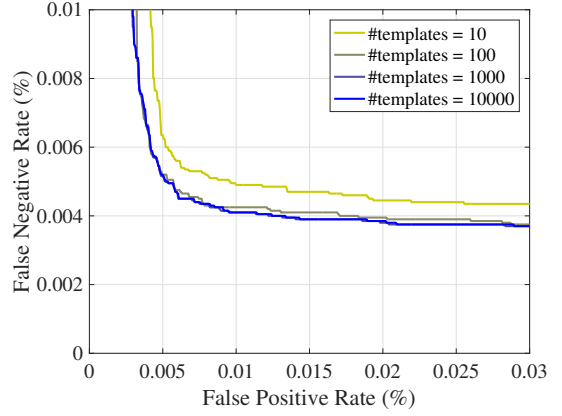


FIG. 11. The best performance of matched filtering with given number of templates across 30 independent runs. The performance starts to saturate above 1000 templates.

C. Neural Network Configuration

To initialize the templates of the neural network models for both **MNet-Shallow** and **MNet-Deep**, we generate 1000 random waveforms from a uniform distribution over the same parameter range, subject to the same preprocessing and normalization process as done in matched filtering.

For the **MNet-Deep** architecture, in addition to the 1000 initialized templates, we also need to specify the number of layers and the feature dimension of each layer. In the experiment we choose $L = 17$ and

$$(n_1, n_2, \dots, n_L) = (1000, 1800, 1200, 720, 480, 300, 180, 120, 90, 60, 36, 24, 18, 12, 6, 3, 1).$$

Here these feature dimensions n_l are chosen arbitrarily so long as they satisfy $n_2 \geq \frac{3}{2}n_1$, $n_l \geq \frac{1}{2}n_{l-1}$ for all $3 \leq l \leq L-1$, $n_{L-1} = 3$, $n_L = 1$, and that $n_2, \dots, n_{L-2}$ are all divisible by 6 (which facilitates construction using our proposed initialization scheme).

For minimax training, in order to search the parameter space for the worst performance, we sample the parameter space $[40, 50] \times [40, 50]$ of (m_1, m_2) using a square grid sampler with interval 0.5. After discarding equivalent samples due to the symmetry between m_1 and m_2 , there are in total 231 samples in the parameter space.

For the optimization parameters of the neural network, we train the network using logistic loss, the Adam optimizer [117], and a constant learning rate of 10^{-5} .

D. Simulation Results

Performance under minimax. In this experiment we perform injections such that SNR is 5, and only for the **MNet-Shallow** model. While this SNR value is smaller than the range of meaningful observed events, we choose

this value for the simplicity of exposition and reduction of training time, since the training procedure for minimax criterion is rather computationally heavy. Similar results should hold at higher SNR values. FIG 12 plots the ROC curves for both matched filtering and MNet-Shallow trained for minimax, measured in terms of both worst performance and the average performance over a uniform prior. We see that the trained neural network achieves better performance than matched filtering under minimax, while achieving approximately identical performance as matched filtering under Neyman-Pearson with a uniform prior. This is not surprising since the training process is designed to only optimize for the minimax criterion, and not the Neyman-Pearson criterion with uniform prior.

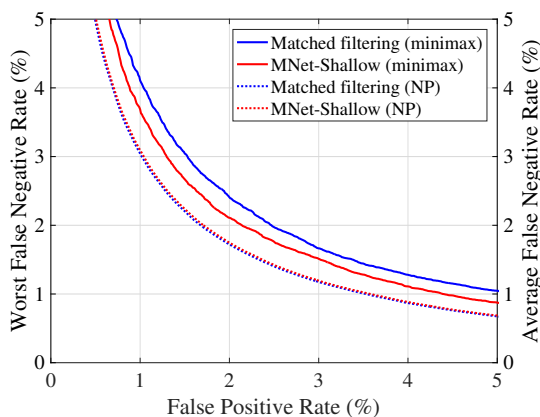


FIG. 12. ROC curves of the trained shallow neural network and matched filtering. The solid curves correspond to the vertical axis on the left, and the dotted curves correspond to the vertical axis on the right. For both models we show both the worst (minimax) performance and the average performance under Neyman-Pearson (NP) setting with a uniform prior. The neural network with minimax training outperforms matched filtering in terms of the minimax criterion. The performance of the two models under NP is similar, which is reasonable since our optimization for the neural network was aimed for the minimax criterion only.

Performance under a uniform prior. In this experiment we perform injections such that SNR is 9. Figure 13 plots the ROC curves for both formulations MNet-Shallow and MNet-Deep trained for Neyman-Pearson, as well as that of matched filtering. As expected, the neural network models strictly improves over matched filtering. Moreover, the MNet-Deep architecture has a slight performance advantage over MNet-Shallow. The performance improvement of the trained models over matched filtering is especially remarkable with low FNR values, which is arguably the more important scenario for gravitational wave detection, since we can hardly afford to miss actual astrophysical events which are quite scarce.

VII. DISCUSSION

Our experiments demonstrate the potential of neural networks to outperform matched filtering, especially at low false negative rates. The flexibility of neural networks also enables this architecture to implement more general variations of matched filtering, such as with weights or aggregation functions different from the maximum. Neural networks have additional potential advantages: deep networks can adapt to unknown and/or non-Gaussian noise distributions. In addition, architectural ideas in deep networks such as pooling help to convey invariances that may be helpful in detecting some “unknown unknowns” that lie outside of the span of a pre-specified family of templates. This should be investigated in the future.

The proposed architectures can be adapted to time-varying noise distributions, by pre-training on very large collections of (synthetic) Gaussian noise and then adapting the pre-trained network using a smaller number of on-line examples. This kind of pre-training may also be helpful in deploying our methods across larger mass ranges, which require more training data.

We note that it is, in some sense, unsurprising that deep networks can exhibit advantages over matched filtering, since the former can be made arbitrarily complex, and can approximate essentially arbitrary functions. An important direction for future work is to study architectures that not only approach optimal statistical performance, but exhibit good *complexity-performance* tradeoffs. There are a number of concrete directions for achieving this – in particular, the weight matrices learned by our Neyman-Pearson networks exhibit particular types of low-dimensional (low-rank and sparse) structure, which can be leveraged to reduce complexity. Interpreting matched filtering as a particular neural network facilitates the study of complexity-performance tradeoffs, since it allows these distinct methods to be studied in a unified framework. Another avenue for complexity reduction is to define and train very large (over-parameterized) networks and then prune them to produce much smaller subnetworks with good performance. MNet-Deep is particularly promising in this regard, since this construction yields networks of arbitrary depth.

One future possibility of the approach is to go beyond the fixed template banks that constrain the limited set of parameters taken into account. For example, to limit the size of the template bank, BH spins that are misaligned from the orbital angular momentum are not widely used yet. Also, due to the lack of available template banks, some astrophysically feasible scenarios receive relatively little attention, including eccentric binary merger template banks where every new template requires a computationally very expensive general-relativity simulation. Therefore generalized matched filtering needs to be investigated in this context, to measure its performance on signal classes that current templates don’t cover. Additionally, training it with a sample of eccentric wave-

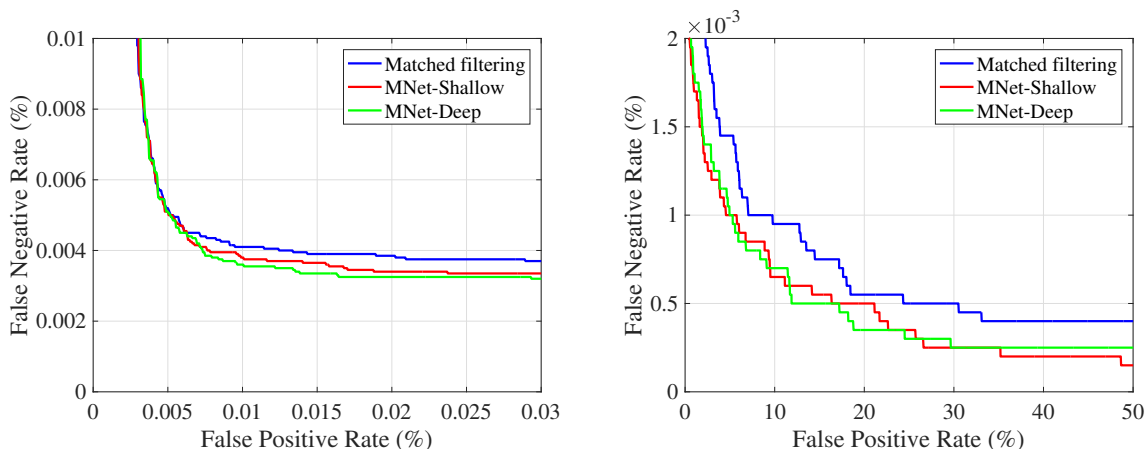


FIG. 13. ROC curves of the trained **MNet-Shallow** and **MNet-Deep** models compared with matched filtering. Left and right panels plot the same curves, but have different axis ranges to better show the contrast between the curves.

forms could enable the detection of other eccentric BBHs even with properties not covered by the limited simulation used for training. Exploring these scenarios are very important experiments for the future.

Another desirable goal is to allow matched filtering algorithms to run "coherently", treating the GW detectors worldwide as a single detector and analyzing data from multiple GW detectors together as a single data stream. The main difficulty is that the sky direction of the cosmic source is unknown, therefore there are many unknown time shifts among the detectors' data. Searching a large number of different combinations can be cost prohibitive with current approaches. It is important to experimentally investigate the ML extensions to matched filtering to measure the increased sensitivity due to the coherent framework.

Furthermore, experiments on the natural generalization of the approach where one does not aim to find the best matching waveform, but instead aims to estimate the parameters of the BBH system are needed. For example, instead of having the maximum reported, one could report the probability distribution over parameters. The difficulty here is that searches usually have much fewer parameters than what is used for parameter estimation. The performance of the ML framework in parameter estimation should be quantified in the future, even if it comes at the price of precision and is therefore only used as a first estimate.

VIII. CONCLUSION

In this paper, we highlighted the idea that matched filtering currently applied by LIGO is formally equivalent to a particular neural network, which can be defined analytically in closed form. We also modeled the LIGO gravitational wave search as the parametric signal detection problem, and illustrated the suboptimality of matched filtering regardless of whether a prior distribution on the

parameter space is given. On the other hand, we proposed neural network architectures **MNet-Shallow** and **MNet-Deep**, which are initialized to implement matched filtering exactly, and then trained on data for improved performance. In particular, we showed that when the prior distribution is known, the training process is aligned with the statistically optimal decision rule. Between the two proposed architectures, the former more closely resembles the architecture of matched filtering, while the latter has a more flexible architecture capable of dealing with a wider range of distributions. We conducted experiments using LIGO strain data from O2 and synthetic waveform injections, and showed that our trained network can achieve uniformly better performance than matched filtering both with or without a known prior, especially in scenarios where false negative rate is low.

Through this work, we seek to bridge the gap between data-driven methods such as deep learning and those detection methods currently in use in LIGO, and explore the possibility of incorporating them into the gravitational wave search of LIGO, as well as broader areas of scientific discovery. In the future work, we aim to explore the potentials of efficiency gains of neural networks over matched filtering, and also establish an end-to-end guarantee for the performance of the proposed framework.

ACKNOWLEDGMENTS

We acknowledge computing resources from Columbia University's Shared Research Computing Facility project, which is supported by NIH Research Facility Improvement Grant 1G20RR030893-01, and associated funds from the New York State Empire State Development, Division of Science Technology and Innovation (NYSTAR) Contract C090171, both awarded April 15, 2010.

The authors are grateful for the LIGO Scientific Collaboration for the careful review of the paper. This pa-

per is assigned a LIGO DCC number of LIGO-P2100086. The authors acknowledge the LIGO Laboratory and Scientific Collaboration for the detectors, data, and the game changing computing resources (National Science Foundation Grants PHY-0757058 and PHY-0823459). The authors would like to thank colleagues of the LIGO Scientific Collaboration and the Virgo Collaboration and Columbia University for their help and useful comments, in particular the CBC group, Andrew Williamson, Stefan Countryman, William Tse, Nicolas Beltran, Asif Malik, Sireesh Gururaja, and Thomas Dent which we hereby gratefully acknowledge. SM thanks David Spergel, Rainer Weiss, Rana Adhikari, and Kipp Canon for the motivating general discussions related to the role of machine learning and data analysis.

This research has made use of data, software and/or web tools obtained from the Gravitational Wave Open Science Center [92, 118] (<https://www.gwopenscience.org/>), a service of LIGO Laboratory, the LIGO Scientific Collaboration and the Virgo Collaboration. LIGO Laboratory and Advanced LIGO are funded by the United States National Science Foundation (NSF) as well as the Science and Technology Fa-

cilities Council (STFC) of the United Kingdom, the Max-Planck-Society (MPS), and the State of Niedersachsen/Germany for support of the construction of Advanced LIGO and construction and operation of the GEO600 detector. Additional support for Advanced LIGO was provided by the Australian Research Council. Virgo is funded, through the European Gravitational Observatory (EGO), by the French Centre National de Recherche Scientifique (CNRS), the Italian Istituto Nazionale di Fisica Nucleare (INFN) and the Dutch Nikhef, with contributions by institutions from Belgium, Germany, Greece, Hungary, Ireland, Japan, Monaco, Poland, Portugal, Spain.

The authors thank the University of Florida and Columbia University in the City of New York for their generous support. The authors are grateful for the generous support of the National Science Foundation under grant CCF-1740391. The authors thank Sharon Sputz of Columbia University for her effort in facilitating this collaboration.

I.B. acknowledges the support of the National Science Foundation under grant PHY-1911796 and the Alfred P. Sloan Foundation.

-
- [1] LIGO and Virgo Collaborations, *Phys. Rev. Lett.* **116**, 061102 (2016)
 - [2] B. P. Abbott and et al., *Physical Review X* **9** (2019), 10.1103/physrevx.9.031040
 - [3] V. LSC and KAGRA, arXiv e-prints, arXiv:2010.14527 (2020), arXiv:2010.14527 [gr-qc]
 - [4] B. P. Abbott and et al., *Phys. Rev. Lett.* **119**, 161101 (2017)
 - [5] B. P. Abbott and et al., *Astrophysical Journal Letters* **848**, L12 (2017), arXiv:1710.05833 [astro-ph.HE]
 - [6] T. Akutsu *et al.*, *Progress of Theoretical and Experimental Physics* (2020), 10.1093/ptep/ptaa125, ptaa125, <https://academic.oup.com/ptep/advance-article-pdf/doi/10.1093/ptep/ptaa125/34386189/ptaa125.pdf>
 - [7] K. L. Dooley, J. R. Leong, T. Adams, C. Affeldt, A. Bisht, C. Bogan, J. Degallaix, C. Gräf, S. Hild, J. Hough, A. Khalaidovski, N. Lastzka, J. Lough, H. Lück, D. Macleod, L. Nuttall, M. Prijatelj, R. Schnabel, E. Schreiber, J. Slutsky, B. Sorazu, K. A. Strain, H. Vahlbruch, M. Was, B. Willke, H. Wittel, K. Danzmann, and H. Grote, *Classical and Quantum Gravity* **33**, 075009 (2016)
 - [8] F. Acernese *et al.*, *Classical and Quantum Gravity* **32**, 024001 (2015), arXiv:1408.3978 [gr-qc]
 - [9] A. Abramovici, W. E. Althouse, R. W. P. Drever, Y. Gursel, S. Kawamura, F. J. Raab, D. Shoemaker, L. Sievers, R. E. Spero, K. S. Thorne, R. E. Vogt, R. Weiss, S. E. Whitcomb, and M. E. Zucker, *Science* **256**, 325 (1992)
 - [10] B. P. Abbott *et al.*, *Classical and Quantum Gravity* **32**, 074001 (2015), arXiv:1411.4547 [gr-qc]
 - [11] C. Affeldt, K. Danzmann, K. L. Dooley, H. Grote, M. Hewitson, S. Hild, J. Hough, J. Leong, H. Lück, M. Prijatelj, S. Rowan, A. Rüdiger, R. Schilling, R. Schnabel, E. Schreiber, B. Sorazu, K. A. Strain, H. Vahlbruch, B. Willke, W. Winkler, and H. Wittel, *Classical and Quantum Gravity* **31**, 224002 (2014)
 - [12] F. Acernese, M. Agathos, L. Aiello, A. Allocca, A. Amato, S. Ansoldi, S. Antier, M. Arène, N. Arnaud, S. Ascenzi, P. Astone, F. Aubin, S. Babak, P. Bacon, F. Badaracco, M. K. M. Bader, J. Baird, F. Baldaccini, G. Ballardín, G. Baltus, C. Barbieri, P. Barneo, F. Barone, M. Barsuglia, D. Barta, A. Basti, M. Bawaj, M. Bazzan, M. Bejger, I. Belahcene, S. Bernuzzi, D. Bersanetti, A. Bertolini, M. Bischì, M. Bitossi, M. A. Bizouard, F. Bobba, M. Boer, G. Bogaert, F. Bondu, R. Bonnand, B. A. Boom, V. Boschi, Y. Bouffanais, A. Bozzi, C. Bradaschia, M. Branchesi, M. Breschi, T. Briant, F. Brighenti, A. Brillet, J. Brooks, G. Bruno, T. Bulik, H. J. Bulten, D. Buskulic, G. Cagnoli, E. Calloni, M. Canepa, G. Carapella, F. Carbognani, G. Carullo, J. Casanueva Diaz, C. Casentini, J. Castañeda, S. Caudill, F. Cavalier, R. Cavalieri, G. Cella, P. Cerdá-Durán, E. Cesarini, O. Chaibi, E. Chassande-Mottin, F. Chiadini, R. Chierici, A. Chincarini, A. Chiummo, N. Christensen, S. Chua, G. Ciani, P. Ciecìelag, M. Cieřlar, R. Cioffi, F. Cipriano, A. Cirone, S. Clesse, F. Cleva, E. Coccia, P.-F. Cohadon, D. Cohen, M. Colpi, L. Conti, I. Cordero-Carrión, S. Corezzi, D. Corre, S. Cortese, J.-P. Coulon, M. Croquette, J.-R. Cudell, E. Cuoco, M. Curylo, B. D’Angelo, S. D’Antonio, V. Dattilo, M. Davier, J. Degallaix, M. De Laurentis, S. Deléglise, W. Del Pozzo, R. De Pietri, R. De Rosa, C. De Rossi, T. Dietrich, L. Di Fiore, C. Di Giorgio, F. Di Giovanni, M. Di Giovanni, T. Di Girolamo, A. Di Lieto, S. Di Pace, I. Di Palma, F. Di Renzo, M. Drago, J.-G. Ducoin, O. Durante, D. D’Urso, M. Eisenmann, L. Errico, D. Es-

- tevez, V. Fafone, S. Farinon, F. Feng, E. Fenyvesi, I. Ferrante, F. Fidecaro, I. Fiori, D. Fiorucci, R. Fitipaldi, V. Fiumara, R. Flaminio, J. A. Font, J.-D. Fournier, S. Frasca, F. Frasconi, V. Frey, G. Fronzè, F. Garufi, G. Gemme, E. Genin, A. Gennai, A. Ghosh, B. Giacomazzo, M. Gosselin, R. Gouaty, A. Grado, M. Granata, G. Greco, G. Grignani, A. Grimaldi, S. J. Grimm, P. Gruning, G. M. Guidi, G. Guixé, Y. Guo, P. Gupta, O. Halim, T. Harder, J. Harms, A. Heidmann, H. Heitmann, P. Hello, G. Hemming, E. Hennes, T. Hinderer, D. Hofman, D. Huet, V. Hui, B. Idzkowski, A. Iess, G. Intini, J.-M. Isac, T. Jacqmin, P. Jaranowski, R. J. G. Jonker, S. Katsanevas, F. Kéfélian, I. Khan, N. Khetan, G. Koekoek, S. Koley, A. Królak, A. Kutyina, D. Laghi, A. Lamberts, I. La Rosa, A. Lartaux-Vollard, C. Lazzaro, P. Leaci, N. Leroy, N. Letendre, F. Linde, M. Llorens-Monteaudo, A. Longo, M. Lorenzini, V. Lorette, G. Losurdo, D. Lumaca, A. Macquet, E. Majorana, I. Maksimovic, N. Man, V. Mangano, M. Mantovani, M. Mapelli, F. Marchesoni, F. Marion, A. Marquina, S. Marsat, F. Martelli, V. Martinez, A. Masserot, S. Mastrogiovanni, E. Mejuto Villa, L. Mereni, M. Merzougui, R. Metzdorff, A. Miani, C. Michel, L. Milano, A. Miller, E. Milotti, O. Minazzoli, Y. Minenkov, M. Montani, F. Morawski, B. Mours, F. Muciaccia, A. Nagar, I. Nardecchia, L. Naticchioni, J. Neilson, G. Nelemans, C. Nguyen, D. Nichols, S. Nissanke, F. Nocera, G. Oganessian, C. Olivetto, G. Pagano, G. Pagliaroli, C. Palomba, P. T. H. Pang, F. Pannarale, F. Paoletti, A. Paoli, D. Pascucci, A. Pasqualetti, R. Passaquieti, D. Passuello, B. Patricelli, A. Perego, M. Pegoraro, C. Périgois, A. Perreca, S. Perriès, K. S. Phukon, O. J. Piccinni, M. Pichot, M. Piendibene, F. Piergiovanni, V. Pierro, G. Pillant, L. Pinard, I. M. Pinto, K. Piotrkowski, W. Plastino, R. Poggiani, P. Popolizio, E. K. Porter, M. Prevedelli, M. Principe, G. A. Prodi, M. Punturo, P. Puppo, G. Raaijmakers, N. Radulesco, P. Rapagnani, M. Razzano, T. Regimbau, L. Rei, P. Rettengo, F. Ricci, G. Riemenschneider, F. Robinet, A. Rocchi, L. Roland, M. Romanelli, R. Romano, D. Rosińska, P. Ruggi, O. S. Salafia, L. Salconi, A. Samajdar, N. Sanchis-Gual, E. Santos, B. Sassolas, O. Sauter, S. Sayah, D. Sentenac, V. Sequino, A. Sharma, M. Sieniawska, N. Singh, A. Singhal, V. Sipala, V. Sordini, F. Sorrentino, M. Spera, C. Stachie, D. A. Steer, G. Stratta, A. Sur, B. L. Swinkels, M. Tacca, A. J. Tanasijczuk, E. N. Tapia San Martin, S. Tiwari, M. Tonelli, A. Torres-Forné, I. Tosta e Melo, F. Travasso, M. C. Tringali, A. Trovato, K. W. Tsang, M. Turconi, M. Valentini, N. van Bakel, M. van Beuzekom, J. F. J. van den Brand, C. Van Den Broeck, L. van der Schaaf, M. Vardaro, M. Vasúth, G. Vedovato, D. Verkindt, F. Vetrano, A. Viceré, J.-Y. Vinet, H. Vocca, R. Walet, M. Was, A. Zadrożny, T. Zelenova, J.-P. Zendri, H. Vahlbruch, M. Mehmet, H. Lück, and K. Danzmann (Virgo Collaboration), *Phys. Rev. Lett.* **123**, 231108 (2019)
- [13] M. Tse, H. Yu, N. Kijbunchoo, A. Fernandez-Galiana, P. Dupej, L. Barsotti, C. D. Blair, D. D. Brown, S. E. Dwyer, A. Effler, M. Evans, P. Fritschel, V. V. Frolov, A. C. Green, G. L. Mansell, F. Matichard, N. Mavalvala, D. E. McClelland, L. McCuller, T. McRae, J. Miller, A. Mullavey, E. Oelker, I. Y. Phinney, D. Sigg, B. J. J. Slagmolen, T. Vo, R. L. Ward, C. Whittle, R. Abbott, C. Adams, R. X. Adhikari, A. Ananyeva, S. Appert, K. Arai, J. S. Areeda, Y. Asali, S. M. Aston, C. Austin, A. M. Baer, M. Ball, S. W. Ballmer, S. Banagiri, D. Barker, J. Bartlett, B. K. Berger, J. Betzwieser, D. Bhattacharjee, G. Billingsley, S. Biscans, R. M. Blair, N. Bode, P. Booker, R. Bork, A. Bramley, A. F. Brooks, A. Buikema, C. Cahillane, K. C. Cannon, X. Chen, A. A. Ciobanu, F. Clara, S. J. Cooper, K. R. Corley, S. T. Countryman, P. B. Covas, D. C. Coyne, L. E. H. Datrier, D. Davis, C. Di Fronzo, J. C. Driggers, T. Etzel, T. M. Evans, J. Feicht, P. Fulda, M. Fyffe, J. A. Giaime, K. D. Giardina, P. Godwin, E. Goetz, S. Gras, C. Gray, R. Gray, A. Gupta, E. K. Gustafson, R. Gustafson, J. Hanks, J. Hanson, T. Hardwick, R. K. Hasskew, M. C. Heintze, A. F. Helmling-Cornell, N. A. Holland, J. D. Jones, S. Kandhasamy, S. Karki, M. Kasprzack, K. Kawabe, P. J. King, J. S. Kissel, R. Kumar, M. Landry, B. B. Lane, B. Lantz, M. Laxen, Y. K. Lecoecuche, J. Leviton, J. Liu, M. Lormand, A. P. Lundgren, R. Macas, M. MacInnis, D. M. Macleod, S. Márka, Z. Márka, D. V. Martynov, K. Mason, T. J. Massinger, R. McCarthy, S. McCormick, J. McIver, G. Mendell, K. Merfeld, E. L. Merilh, F. Meylahn, T. Mistry, R. Mittleman, G. Moreno, C. M. Mow-Lowry, S. Mozzon, T. J. N. Nelson, P. Nguyen, L. K. Nuttall, J. Oberling, R. J. Oram, B. O'Reilly, C. Osthelder, D. J. Ottaway, H. Overmier, J. R. Palamos, W. Parker, E. Payne, A. Pele, C. J. Perez, M. Pirello, H. Radkins, K. E. Ramirez, J. W. Richardson, K. Riles, N. A. Robertson, J. G. Rollins, C. L. Romel, J. H. Romie, M. P. Ross, K. Ryan, T. Sadecki, E. J. Sanchez, L. E. Sanchez, T. R. Saravanan, R. L. Savage, D. Schaetzl, R. Schnabel, R. M. J. Schofield, E. Schwartz, D. Sellers, T. J. Shaffer, J. R. Smith, S. Soni, B. Sorazu, A. P. Spencer, K. A. Strain, L. Sun, M. J. Szczepańczyk, M. Thomas, P. Thomas, K. A. Thorne, K. Toland, C. I. Torrie, G. Traylor, A. L. Urban, G. Vajente, G. Valdes, D. C. Vander-Hyde, P. J. Veitch, K. Venkateswara, G. Venugopalan, A. D. Viets, C. Vorvick, M. Wade, J. Warner, B. Weaver, R. Weiss, B. Willke, C. C. Wipf, L. Xiao, H. Yamamoto, M. J. Yap, H. Yu, L. Zhang, M. E. Zucker, and J. Zweizig, *Phys. Rev. Lett.* **123**, 231107 (2019)
- [14] A. Einstein, Sitzungsberichte der Königlich Preußischen Akademie der Wissenschaften (Berlin), 688 (1916).
- [15] A. Einstein, Sitzungsberichte der Königlich Preußischen Akademie der Wissenschaften (Berlin), 154 (1918).
- [16] W. G. Anderson, P. R. Brady, J. D. E. Creighton, and É. É. Flanagan, *Phys. Rev. D* **63**, 042003 (2001) [arXiv:gr-qc/0008066 \[gr-qc\]](#)
- [17] S. Klimentenko and G. Mitselmakher, *Classical and Quantum Gravity* **21**, S1819 (2004)
- [18] W. G. Anderson, P. R. Brady, J. D. E. Creighton, and É. É. Flanagan, *International Journal of Modern Physics D* **9**, 303 (2000) [arXiv:gr-qc/0001044 \[gr-qc\]](#)
- [19] S. W. Hawking and W. Israel, *Three Hundred Years of Gravitation* (1989).
- [20] B. J. Owen and B. S. Sathyaprakash, *Phys. Rev. D* **60**, 022002 (1999)
- [21] C. Cutler, T. A. Apostolatos, L. Bildsten, L. S. Finn, E. E. Flanagan, D. Kennefick, D. M. Markovic, A. Ori, E. Poisson, G. J. Sussman, and K. S. Thorne, *Phys. Rev. Lett.* **70**, 2984 (1993) [arXiv:astro-](#)

- ph/9208005 [astro-ph]
- [22] C. Cutler and É. E. Flanagan, *Phys. Rev. D* **49**, 2658 (1994) [arXiv:gr-qc/9402014 [gr-qc]].
- [23] É. É. Flanagan and S. A. Hughes, *Phys. Rev. D* **57**, 4535 (1998) [arXiv:gr-qc/9701039 [gr-qc]].
- [24] É. É. Flanagan and S. A. Hughes, *Phys. Rev. D* **57**, 4566 (1998) [arXiv:gr-qc/9710129 [gr-qc]].
- [25] B. e. a. Abbott, *Phys. Rev. D* **69**, 122001 (2004) [arXiv:gr-qc/0308069 [gr-qc]].
- [26] “Pycbc software releases,” <https://github.com/gwastro/pycbc/releases>
- [27] B. Allen, W. G. Anderson, P. R. Brady, D. A. Brown, and J. D. E. Creighton, *Phys. Rev. D* **85**, 122006 (2012) [arXiv:gr-qc/0509116 [gr-qc]].
- [28] C. M. Biwer, C. D. Capano, S. De, M. Cabero, D. A. Brown, A. H. Nitz, and V. Raymond, *PASP* **131**, 024503 (2019) [arXiv:1807.10312 [astro-ph.IM]].
- [29] B. Allen, *Phys. Rev. D* **71**, 062001 (2005) [arXiv:gr-qc/0405045 [gr-qc]].
- [30] T. Dal Canton, A. H. Nitz, A. P. Lundgren, A. B. Nielsen, D. A. Brown, T. Dent, I. W. Harry, B. Krishnan, A. J. Miller, K. Wette, K. Wiesner, and J. L. Willis, *Phys. Rev. D* **90**, 082004 (2014) [arXiv:1405.6731 [gr-qc]].
- [31] S. A. Usman, A. H. Nitz, I. W. Harry, C. M. Biwer, D. A. Brown, M. Cabero, C. D. Capano, T. Dal Canton, T. Dent, S. Fairhurst, M. S. Kehl, D. Keppel, B. Krishnan, A. Lenon, A. Lundgren, A. B. Nielsen, L. P. Pekowsky, H. P. Pfeiffer, P. R. Saulson, M. West, and J. L. Willis, *Classical and Quantum Gravity* **33**, 215004 (2016) [arXiv:1508.02357 [gr-qc]].
- [32] A. H. Nitz, T. Dal Canton, D. Davis, and S. Reyes, *Phys. Rev. D* **98**, 024050 (2018) [arXiv:1805.11174 [gr-qc]].
- [33] C. Messick, K. Blackburn, P. Brady, P. Brockill, K. Cannon, R. Cariou, S. Caudill, S. J. Chamberlin, J. D. E. Creighton, R. Everett, C. Hanna, D. Keppel, R. N. Lang, T. G. F. Li, D. Meacher, A. Nielsen, C. Pankow, S. Privitera, H. Qi, S. Sachdev, L. Sadeghian, L. Singer, E. G. Thomas, L. Wade, M. Wade, A. Weinstein, and K. Wiesner, *Phys. Rev. D* **95**, 042001 (2017) [arXiv:1604.04324 [astro-ph.IM]].
- [34] S. Sachdev, S. Caudill, H. Fong, R. K. L. Lo, C. Messick, D. Mukherjee, R. Magee, L. Tsukada, K. Blackburn, P. Brady, P. Brockill, K. Cannon, S. J. Chamberlin, D. Chatterjee, J. D. E. Creighton, P. Godwin, A. Gupta, C. Hanna, S. Kapadia, R. N. Lang, T. G. F. Li, D. Meacher, A. Pace, S. Privitera, L. Sadeghian, L. Wade, M. Wade, A. Weinstein, and S. Liting Xiao, arXiv e-prints, arXiv:1901.08580 (2019), [arXiv:1901.08580 [gr-qc]].
- [35] C. Hanna, S. Caudill, C. Messick, A. Reza, S. Sachdev, L. Tsukada, K. Cannon, K. Blackburn, J. D. E. Creighton, H. Fong, P. Godwin, S. Kapadia, T. G. F. Li, R. Magee, D. Meacher, D. Mukherjee, A. Pace, S. Privitera, R. K. L. Lo, and L. Wade, *Phys. Rev. D* **101**, 022003 (2020) [arXiv:1901.02227 [gr-qc]].
- [36] Q. Chu, *Low-latency detection and localization of gravitational waves from compact binary coalescences*, [Ph.D. thesis], The University of Western Australia (2017).
- [37] T. Adams, D. Buskulic, V. Germain, G. M. Guidi, F. Marion, M. Montani, B. Mours, F. Piergiovanni, and G. Wang, *Classical and Quantum Gravity* **33**, 175012 (2016) [arXiv:1512.02864 [gr-qc]].
- [38] A. H. Nitz, T. Dent, T. Dal Canton, S. Fairhurst, and D. A. Brown, *The Astrophysical Journal* **849**, 118 (2017).
- [39] A. C. Searle, arXiv e-prints, arXiv:0804.1161 (2008), [arXiv:0804.1161 [gr-qc]].
- [40] R. Biswas, P. R. Brady, J. Burguet-Castell, K. Cannon, J. Clayton, A. Dietz, N. Fotopoulos, L. M. Goggin, D. Keppel, C. Pankow, L. R. Price, and R. Vaulin, *Phys. Rev. D* **85**, 122008 (2012) [arXiv:1201.2959 [gr-qc]].
- [41] T. Dent and J. Veitch, *Phys. Rev. D* **89**, 062002 (2014) [arXiv:1311.7174 [gr-qc]].
- [42] B. J. Owen and B. S. Sathyaprakash, *Phys. Rev. D* **60**, 022002 (1999) [arXiv:gr-qc/9808076 [gr-qc]].
- [43] B. J. Owen, *Phys. Rev. D* **53**, 6749 (1996) [arXiv:gr-qc/9511032 [gr-qc]].
- [44] B. S. Sathyaprakash and S. V. Dhurandhar, *Phys. Rev. D* **44**, 3819 (1991).
- [45] S. V. Dhurandhar and B. S. Sathyaprakash, *Phys. Rev. D* **49**, 1707 (1994).
- [46] B. P. Abbott *et al.*, *Classical and Quantum Gravity* **33**, 134001 (2016) [arXiv:1602.03844 [gr-qc]].
- [47] B. P. Abbott *et al.*, *Classical and Quantum Gravity* **37**, 055002 (2020) [arXiv:1908.11170 [gr-qc]].
- [48] V. Gayathri, J. Healy, J. Lange, B. O’Brien, M. Szczepanczyk, I. Bartos, M. Campanelli, S. Klimentko, C. Lousto, and R. O’Shaughnessy, arXiv e-prints, arXiv:2009.05461 (2020), [arXiv:2009.05461 [astro-ph.HE]].
- [49] T. Gebhard, N. Kilbertus, G. Parascandolo, I. Harry, and B. Schölkopf, in *Workshop on Deep Learning for Physical Sciences (DLPS) at the 31st Conference on Neural Information Processing Systems (NIPS)* (2017) pp. 1–6.
- [50] D. George and E. Huerta, *Physics Letters B* **778**, 64 (2018).
- [51] H. Gabbard, M. Williams, F. Hayes, and C. Messenger, *Physical review letters* **120**, 141103 (2018).
- [52] D. George and E. Huerta, *Physical Review D* **97**, 044039 (2018).
- [53] X. Fan, J. Li, X. Li, Y. Zhong, and J. Cao, *SCIENCE CHINA Physics, Mechanics & Astronomy* **62**, 969512 (2019).
- [54] F. Morawski, M. Beijer, and P. Cieciela, *Machine Learning: Science and Technology* **1**, 025016 (2020).
- [55] C. Dreissigacker, R. Sharma, C. Messenger, R. Zhao, and R. Prix, *Physical Review D* **100**, 044009 (2019).
- [56] P. G. Krastev, *Physics Letters B*, 135330 (2020).
- [57] B.-J. Lin, X.-R. Li, and W.-L. Yu, *Frontiers of Physics* **15**, 1 (2020).
- [58] Y.-C. Lin and J.-H. P. Wu, arXiv preprint arXiv:2007.04176 (2020).
- [59] C. Bresten and J.-H. Jung, arXiv preprint arXiv:1910.08245 (2019).
- [60] P. Astone, P. Cerdá-Durán, I. Di Palma, M. Drago, F. Muciaccia, C. Palomba, and F. Ricci, *Physical Review D* **98**, 122002 (2018).
- [61] T. S. Yamamoto and T. Tanaka, arXiv preprint arXiv:2011.12522 (2020).
- [62] C. Dreissigacker and R. Prix, *Physical Review D* **102**, 022005 (2020).
- [63] R. Corizzo, M. Ceci, E. Zdravetski, and N. Japkowicz, *Expert Systems with Applications* **151**, 113378 (2020).

- [64] A. L. Miller, P. Astone, S. D’Antonio, S. Frasca, G. Intini, I. La Rosa, P. Leaci, S. Mastrogiovanni, F. Muciaccia, A. Mitidis, *et al.*, *Physical Review D* **100**, 062005 (2019).
- [65] J. Bayley, C. Messenger, and G. Woan, *Physical Review D* **102**, 083024 (2020).
- [66] P. G. Krastev, K. Gill, V. A. Villar, and E. Berger, *arXiv preprint arXiv:2012.13101* (2020).
- [67] H.-M. Luo, W. Lin, Z.-C. Chen, and Q.-G. Huang, *Frontiers of Physics* **15**, 1 (2020).
- [68] G. R. Santos, M. P. Figueiredo, A. d. P. Santos, P. Protopapas, and T. A. Ferreira, *arXiv preprint arXiv:2003.09995* (2020).
- [69] M. L. Chan, I. S. Heng, and C. Messenger, *Physical Review D* **102**, 043022 (2020).
- [70] H. Xia, L. Shao, J. Zhao, and Z. Cao, *Physical Review D* **103**, 024040 (2021).
- [71] R. Biswas, L. Blackburn, J. Cao, R. Essick, K. A. Hodge, E. Katsavounidis, K. Kim, Y.-M. Kim, E.-O. Le Bigot, C.-H. Lee, *et al.*, *Physical Review D* **88**, 062003 (2013).
- [72] D. George, H. Shen, and E. Huerta, *arXiv preprint arXiv:1706.07446* (2017).
- [73] N. Mukund, S. Abraham, S. Kandhasamy, S. Mitra, and N. S. Philip, *Physical Review D* **95**, 104059 (2017).
- [74] M. Razzano and E. Cuoco, *Classical and Quantum Gravity* **35**, 095016 (2018).
- [75] S. Coughlin, S. Bahaadini, N. Rohani, M. Zevin, O. Patane, M. Harandi, C. Jackson, V. Noroozi, S. Allen, J. Areeda, *et al.*, *Physical Review D* **99**, 082002 (2019).
- [76] R. E. Colgan, K. R. Corley, Y. Lau, I. Bartos, J. N. Wright, Z. Márka, and S. Márka, *Physical Review D* **101**, 102003 (2020).
- [77] H. Nakano, T. Narikawa, K.-i. Oohara, K. Sakai, H.-a. Shinkai, H. Takahashi, T. Tanaka, N. Uchikata, S. Yamamoto, and T. S. Yamamoto, *Physical Review D* **99**, 124032 (2019).
- [78] S. R. Green, C. Simpson, and J. Gair, *Physical Review D* **102**, 104057 (2020).
- [79] J. P. Marulanda, C. Santa, and A. E. Romano, *Physics Letters B* **810**, 135790 (2020).
- [80] A. Caramete, A. Constantinescu, L. Caramete, T. Popescu, R. Balasov, D. Felea, M. Rusu, P. Stefanescu, and O. Tintareanu, *arXiv preprint arXiv:2009.06109* (2020).
- [81] A. Delaunoy, A. Wehenkel, T. Hinderer, S. Nissanke, C. Weniger, A. R. Williamson, and G. Louppe, *arXiv preprint arXiv:2010.12931* (2020).
- [82] H. Shen, D. George, E. A. Huerta, and Z. Zhao, in *ICASSP 2019-2019 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)* (IEEE, 2019) pp. 3237–3241.
- [83] W. Wei and E. Huerta, *Physics Letters B* **800**, 135081 (2020).
- [84] T. D. Gebhard, N. Kilbertus, I. Harry, and B. Schölkopf, *Physical Review D* **100**, 063015 (2019).
- [85] R. Balestriero and R. Baraniuk, *arXiv preprint arXiv:1805.06576* (2018).
- [86] J. Cheng, Y. Wu, W. AbdAlmageed, and P. Natarajan, in *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition* (2019) pp. 11553–11562.
- [87] J. M. Bower, Y.-F. Wong, and J. Banik, in *Neural Information Processing Systems* (1988) pp. 103–113.
- [88] D. Tank and J. Hopfield, *IEEE Transactions on Circuits and Systems* **33**, 533 (1986).
- [89] D. Buniatyan, T. Macrina, D. Ih, J. Zung, and H. S. Seung, *arXiv preprint arXiv:1705.08593* (2017).
- [90] Q. Xue, Y. H. Hu, and W. J. Tompkins, *IEEE Transactions on Biomedical Engineering* **39**, 317 (1992).
- [91] R. P. Lippmann and P. Beckman, in *Advances in neural information processing systems* (1989) pp. 124–132.
- [92] R. A. et al, *SoftwareX* **13**, 100658 (2021).
- [93] L. S. Finn, *Phys. Rev. D* **46**, 5236 (1992).
- [94] In general, a global Earth and Space based gravitational-wave detector network can be treated as a composite data stream [42] [119]. However, that added complexity is not necessary when discussing the principles of the paper. Therefore, we constrain ourselves to a single datastream in this proof of principle analysis.
- [95] G. Casella and R. L. Berger, *Statistical inference* (Cengage Learning, 2021).
- [96] Q. Yu, *The Canadian Journal of Statistics/La Revue Canadienne de Statistique*, 281 (1992).
- [97] Y. LeCun, Y. Bengio, and G. Hinton, *nature* **521**, 436 (2015).
- [98] D. Scherer, A. Müller, and S. Behnke, in *International conference on artificial neural networks* (Springer, 2010) pp. 92–101.
- [99] Y. Bengio, I. Goodfellow, and A. Courville, *Deep learning*, Vol. 1 (MIT press Massachusetts, USA:, 2017).
- [100] G. Cybenko, *Mathematics of control, signals and systems* **2**, 303 (1989).
- [101] M. Leshno, V. Y. Lin, A. Pinkus, and S. Schocken, *Neural networks* **6**, 861 (1993).
- [102] Z. Lu, H. Pu, F. Wang, Z. Hu, and L. Wang, *Advances in neural information processing systems* **30**, 6231 (2017).
- [103] S. Buchanan, D. Gilboa, and J. Wright, “Deep networks and the multiple manifold problem,” (2020), [arXiv:2008.11245 \[stat.ML\]](https://arxiv.org/abs/2008.11245)
- [104] K. Gregor and Y. LeCun, in *Proceedings of the 27th international conference on machine learning* (2010) pp. 399–406.
- [105] J. Bruna and S. Mallat, *IEEE transactions on pattern analysis and machine intelligence* **35**, 1872 (2013).
- [106] Z. Wang, D. Liu, J. Yang, W. Han, and T. Huang, in *Proceedings of the IEEE international conference on computer vision* (2015) pp. 370–378.
- [107] J. Sun, H. Li, Z. Xu, *et al.*, in *Advances in neural information processing systems* (2016) pp. 10–18.
- [108] We note that the representation is not unique, and can be subject to shift and scale to produce essentially the same decision rule. Specifically, $\mathbf{b}^1$ can be identity vector times a constant (including zero) and $\mathbf{W}^1$ can be scaled by an arbitrary positive constant. However we choose the form given here for simplicity.
- [109] P. Kidger and T. Lyons, in *Conference on Learning Theory* (PMLR, 2020) pp. 2306–2327.
- [110] P. Orponen *et al.*, *Nordic Journal of Computing* (1994).
- [111] M. Bianchini and F. Scarselli, *IEEE transactions on neural networks and learning systems* **25**, 1553 (2014).
- [112] In fact it is straightforward to show that the conclusion of Proposition 2 holds for more general classes of loss functions, including the logistic loss.
- [113] J. Creighton and W. Anderson, *Gravitational-Wave Physics and Astronomy: An Introduction to Theory, Experiment and Data Analysis*. (2011).

- [114] R. Balasubramanian, B. S. Sathyaprakash, and S. V. Dhurandhar, *Phys. Rev. D* **53**, 3033 (1996) [arXiv:gr-qc/9508011 [gr-qc]]
- [115] P. R. Brady, T. Creighton, C. Cutler, and B. F. Schutz, *Phys. Rev. D* **57**, 2101 (1998) [arXiv:gr-qc/9702050 [gr-qc]]
- [116] C. Messenger, R. Prix, and M. A. Papa, *Phys. Rev. D* **79**, 104017 (2009) [arXiv:0809.5223 [gr-qc]]
- [117] D. P. Kingma and J. Ba, arXiv preprint arXiv:1412.6980 (2014).
- [118] GWOSC, “Gravitational wave open science center,”
- [119] L. S. Finn, *Phys. Rev. D* **63**, 102001 (2001) [arXiv:gr-qc/0010033 [gr-qc]]

Appendix A: Proofs of Key Technical Claims

1. Proof of Proposition 1

Combining the definitions of the likelihood ratio $\lambda(\mathbf{x})$ and the probability densities $\rho_0(\mathbf{x})$ and $\rho_1(\mathbf{x})$, we have

$$\lambda(\mathbf{x}) = \frac{\int \rho_{\text{noise}}(\mathbf{x} - \mathbf{s}_\gamma) d\nu(\gamma)}{\rho_{\text{noise}}(\mathbf{x})} \quad (\text{A1})$$

$$= \int \frac{\rho_{\text{noise}}(\mathbf{x} - \mathbf{s}_\gamma)}{\rho_{\text{noise}}(\mathbf{x})} d\nu(\gamma). \quad (\text{A2})$$

When the noise is Gaussian $\mathcal{N}(\mathbf{0}, \sigma^2 \mathbf{I})$, the integrand equals

$$\frac{\rho_{\text{noise}}(\mathbf{x} - \mathbf{s}_\gamma)}{\rho_{\text{noise}}(\mathbf{x})} = \exp\left(\frac{\langle \mathbf{x}, \mathbf{s}_\gamma \rangle - \|\mathbf{s}_\gamma\|^2/2}{\sigma^2}\right), \quad (\text{A3})$$

which is a convex function of $\mathbf{x}$. Hence after integrating over γ , the resulting function $\lambda(\mathbf{x})$ is still a convex function of $\mathbf{x}$. The optimal decision region is a sublevel set of $\lambda(\mathbf{x})$, and is hence a convex set.

2. Proof of Proposition 2

Assume the training data is drawn iid from some distribution on $(\mathbf{x}, y) \in \mathbb{R}^n \times \{0, 1\}$. In this setting, the previous defined densities $p_0(\mathbf{x})$ and $p_1(\mathbf{x})$ can be expressed as $p_0(\mathbf{x}) = p(\mathbf{x}|y=0)$ and $p_1(\mathbf{x}) = p(\mathbf{x}|y=1)$. If the predictor function is $f: \mathbb{R}^n \rightarrow \mathbb{R}$, then the risk is

$$\mathcal{R}(f) = \mathbb{E}_{(\mathbf{x}, y)}[\ell(f(\mathbf{x}), y)] \quad (\text{A4})$$

$$= \mathbb{P}[y=0] \cdot \mathbb{E}_{\mathbf{x}|y=0}[\ell(f(\mathbf{x}), 0)] + \mathbb{P}[y=1] \cdot \mathbb{E}_{\mathbf{x}|y=1}[\ell(f(\mathbf{x}), 1)] \quad (\text{A5})$$

$$= \mathbb{P}[y=0] \int_{\mathbb{R}^n} \ell(f(\mathbf{x}), 0) p_0(\mathbf{x}) d\mathbf{x} + \mathbb{P}[y=1] \int_{\mathbb{R}^n} \ell(f(\mathbf{x}), 1) p_1(\mathbf{x}) d\mathbf{x} \quad (\text{A6})$$

$$= \int_{\mathbb{R}^n} \left((1-c)\ell(f(\mathbf{x}), 0)p_0(\mathbf{x}) + c\ell(f(\mathbf{x}), 1)p_1(\mathbf{x}) \right) d\mathbf{x}, \quad (\text{A7})$$

where $c := \mathbb{P}[y=1] \in (0, 1)$ is an exogenous constant that only depends on the data distribution. The function that minimizes the above risk is

$$f_\star(\mathbf{x}) = \arg \min_{\hat{y}} (1-c)\ell(\hat{y}, 0)p_0(\mathbf{x}) + c\ell(\hat{y}, 1)p_1(\mathbf{x}) \quad (\text{A8})$$

for all $\mathbf{x} \in \mathbb{R}^n$, or equivalently

$$f_\star(\mathbf{x}) = \arg \min_{\hat{y}} \ell(\hat{y}, 0) + \frac{c\lambda(\mathbf{x})}{1-c} \ell(\hat{y}, 1). \quad (\text{A9})$$

Therefore, the optimal predicted value at a point is the solution to an optimization problem that only depends on the likelihood ratio $\lambda(\mathbf{x})$.

Take an arbitrary fixed $\mathbf{x}$. From the assumption that $\ell(\hat{y}, y)$ is strictly convex and minimized at $\hat{y} = y$, it follows that $\ell(\hat{y}, 0) + \frac{c\lambda(\mathbf{x})}{1-c} \ell(\hat{y}, 1)$ is strictly convex in $\hat{y}$, strictly decreasing on $(-\infty, 0]$ and strictly increasing on $[1, \infty)$. Hence for any $\mathbf{x}$ the risk minimization problem of equation (A9) has a unique solution in $[0, 1]$. The optimal solution can be found from the first-order-condition (FOC). Noticing that $\hat{y}$ cannot be 0 or 1 under the FOC, we can rewrite the FOC as

$$\frac{\ell'(\hat{y}, 0)}{-\ell'(\hat{y}, 1)} = \frac{c\lambda(\mathbf{x})}{1-c}. \quad (\text{A10})$$

From the assumption of strong convexity, we know that on the interval $(0, 1)$ we have $\ell'(\hat{y}, 0) > 0$ and $\ell'(\hat{y}, 1) < 0$, where in ℓ' the derivative is taken with respect to the first argument. Hence the left-hand-side of (A10) is strictly increasing in $\hat{y}$.

This concludes that the optimal decision function $f_\star(\mathbf{x})$ is strictly increasing in $\lambda(\mathbf{x})$.