

A Hybrid Factored Frontier Algorithm for Dynamic Bayesian Networks with a Biopathways Application

Succheendra Kumar Palaniappan, Sundararaman Akshay, Bing Liu,
Blaise Genest, and P.S. Thiagarajan

Abstract—Dynamic Bayesian Networks (DBNs) can serve as succinct probabilistic dynamic models of biochemical networks [1]. To analyze these models, one must compute the probability distribution over system states at a given time point. Doing this exactly is infeasible for large models; hence one must use approximate algorithms. The Factored Frontier algorithm (FF) is one such algorithm [2]. However FF as well as the earlier Boyen-Koller (BK) algorithm [3] can incur large errors. To address this, we present a new approximate algorithm called the Hybrid Factored Frontier (HFF) algorithm. At each time slice, in addition to maintaining probability distributions over local states—as FF does—HFF explicitly maintains the probabilities of a number of global states called spikes. When the number of spikes is 0, we get FF and with all global states as spikes, we get the exact inference algorithm. We show that by increasing the number of spikes one can reduce errors while the additional computational effort required is only quadratic in the number of spikes. We validated the performance of HFF on large DBN models of biopathways. Each pathway has more than 30 species and the corresponding DBN has more than 3,000 nodes. Comparisons with FF and BK show that HFF is a useful and powerful approximate inferencing algorithm for DBNs.

Index Terms—Probability and statistics, symbolic and algebraic manipulation—algorithms, life and medical sciences—biology and genetics.

1 INTRODUCTION

BIOLOGICAL pathways are usually described by a network of biochemical reactions. The dynamics of these reaction networks are often modeled and analyzed as a set of Ordinary Differential Equations (ODEs). The ODEs will be nonlinear due to the nature of the kinetic laws governing the reactions. Except for the toy examples, the ODEs system will also be high dimensional. Hence, closed-form solutions will not be obtainable. One must instead resort to large scale numerical simulations to perform analysis tasks such as parameter estimation and sensitivity analysis. Further, only a small amount of noisy data of limited precision will be available to support model calibration and validation. Guided by these considerations a method for approximating the ODE dynamics of biological pathways as a Dynamic Bayesian Network (DBN) was developed in [1]. This approximation—explained in more detail later—is derived

by discretizing both the time and value domains, sampling an assumed set of initial states and using numerical integration to generate a large number of representative trajectories. Then by exploiting the network structure and simple counting, the generated trajectories are compactly stored as a DBN. One then studies the biochemical network by applying Bayesian inferencing techniques to the much simpler DBN model. This approach significantly reduces the computational burden of tasks such as parameter estimation and sensitivity analysis as demonstrated in [1] and [4]. It also opens up the possibility of incrementally updating the DBN model using belief propagation techniques when additional experimental data becomes available (using the ideas developed in [5]). Finally, one can also apply probabilistic verification methods [6] to the DBN model to gain insights into the dynamics of the biopathway.

In this scheme, the random variables associated with each time point t represent the discretized concentrations of the associated molecular species (protein, gene, etc.) at t . To perform tasks such as parameter estimation and sensitivity analysis one must repeatedly compute the probability of a random variable assuming a specific value at a time point. Due to conditional dependencies between the variables, this will require an effort exponential in the number of species in the biochemical network. Hence for large networks, exact computation is infeasible and one must resort to approximate methods. There is an efficient approximate inferencing technique for DBNs which is already available and it is called the Factored Frontier algorithm (FF) [2], [7].

FF was used extensively in [1], [4] and this subsequently led us to consider its error behavior. Surprisingly, we could

- S.K. Palaniappan, B. Liu, and P.S. Thiagarajan are with the School of Computing, National University of Singapore, 13 Computing Drive, Singapore 117417. E-mail: {succhee, liubing, thiagu}@comp.nus.edu.sg.
- S. Akshay is with IRISA/ENS Cachan Bretagne, C-212, Distribcom, IRISA, Campus de Beaulieu, Rennes Cedex 35042, France. E-mail: akshay@irisa.fr.
- B. Genest is with the Department of Computer Science at IPAL, UMI CNRS (Unité Mixte Internationale), Institute for Infocomm Research (I2R), 1 Fusionopolis Way, #21-01 Connexis (South Tower), Singapore 138632. E-mail: bgenest@irisa.fr.

Manuscript received 16 Nov. 2011; revised 28 Mar. 2012; accepted 30 Mar. 2012; published online 16 Apr. 2012.

For information on obtaining reprints of this article, please send e-mail to: tcbb@computer.org, and reference IEEECS Log Number TCBBSI-2011-11-0306.

Digital Object Identifier no. 10.1109/TCBB.2012.60.

not find in the literature a rigorous error analysis. Further, we found that in the models we considered—as detailed later—though the Factored Frontier algorithm performed well for many variables, there were a few for which it incurred significant errors. This motivated us to construct an improved inferencing algorithm for DBNs called the *Hybrid Factored Frontier* algorithm (HFF) and it is the focus of this paper. HFF is a parameterized version of FF and to bring out its main features, it will be helpful to first see how FF works.

A DBN has a finite set of random variables with each variable having a finite domain of values. The value of a variable at time t only depends on the values of its parents at time $t - 1$. The probabilistic dynamics is captured by a Conditional Probability Table (CPT) associated with each variable at each time point (see Fig. 1c for an example). This table will specify how the value of a variable at t is conditioned by the values of its parent variables at time $t - 1$. The global state of the system at time t is a tuple of values with each component denoting the value assumed by the corresponding variable at t . One is interested in the marginal probability, i.e., the probability of a variable X taking value v at time t . To compute this exactly, we need P^t , the joint probability distribution over global states at time t . This can be computed by propagating P^{t-1} through the CPTs. FF maintains and propagates these joint probability distributions approximately. Such approximate distributions are usually called *belief states*.

FF is a simplified and more efficient version of the earlier Boyen-Koller algorithm (BK) [3]. In BK, a belief state is maintained compactly as a product of the probability distributions of independent clusters of variables. This belief state is then propagated *exactly* at each step through the CPTs. Then the new belief state is compacted again into a product of the probability distributions of the clusters. In contrast, the FF algorithm maintains a belief state as a product of the marginal distributions of the individual variables. Instead of computing first the new belief state as done by BK, the FF algorithm computes the new marginal distributions directly via the propagation of the current marginal distributions through the CPTs.

FF is attractive in terms of simplicity and efficiency but unlike BK [3], it lacks a rigorous error analysis. More importantly, as we observe in Section 4, FF can exhibit significant errors. As for BK, its accuracy crucially depends on how one clusters the variables. Further, computing the next belief state exactly is computationally infeasible when the clusters are large. Identifying the right set of clusters is a difficult problem and there seem to be no efficient techniques for doing this well. One could avoid the problem of identifying clusters by just using singleton clusters (the so called fully factored BK algorithm). However, as we report in Section 4, this can also result in significant errors.

The error analysis for BK [3] suggests that the key to minimizing the overall error is to reduce the error incurred in one step. HFF attempts to achieve this by maintaining the current belief state as a hybrid entity; for a small number of global states called *spikes*, their current probabilities are maintained. The probability distribution over the remaining states is represented, as in FF, as a product of the marginal probability distributions. The key insight underlying this

idea is that when the error produced by one step of the inferencing algorithm is large for a global state, then either the probability of this state or its estimate must itself be high. If such states are chosen to be the spikes then since the total probability is bounded by 1, the number of spikes at each time point must be small. The main technical component of HFF is to explicitly identify and approximately compute the probabilities of the spikes.

A pleasing feature of HFF is that it is a parametrized version of FF with σ , the number of spikes, being the parameter. When $\sigma = 0$, we get FF and when $\sigma = N$ where N is the total number of global states, we get the exact inferencing algorithm. Thus by tuning σ , one can gain control over the error behavior. We have derived the single step error bound for HFF, which then also leads to an error analysis for FF. We show that the worst case one step error of HFF is lower than that of FF. The time complexity of HFF is $O(n \cdot (\sigma^2 + K^{D+1}))$ where n is the number of nodes in the DBN, σ is the number of spikes, K is the maximum number of values that a random variable (associated with each node) can assume and D is the maximum number of parents that a node can have. This compares favorably with the time complexity of FF which is $O(n \cdot K^{D+1})$. Since the running time of HFF is *linear* in n , it scales well in terms of network size. The factor D is determined by the maximum number of reactions that a species takes part in as a product or reactant. For most of the networks we have encountered, D is much smaller than n .

Our experimental results confirm that HFF is a useful and efficient algorithm. We considered four large DBNs. Three of them arise from the EGF-NGF pathway [8] with one model for NGF stimulation, the second for EGF stimulation and the third for EGF-NGF costimulation. The fourth DBN captures the behavior of the Epo mediated ERK signaling pathway [9]. Starting from their ODE-based models (each of which had more than 30 species), we applied the technique developed in [1] to obtain the DBNs each of which had more than 3,000 nodes. In all four cases, we found that the errors suffered by FF and BK (with singleton clusters) were high for the marginal distributions of some biologically significant species. The errors incurred by HFF were always lower and they reduced monotonically when the number of spikes was increased.

1.1 Related Work

DBNs are used extensively in domains such as AI, computer vision, signal processing and computational biology [7], [10], [11], [12]. HFF is a generic inferencing algorithm and it can be applied to compute and maintain belief states in these settings as well. As done in HFF, capturing a large probability mass using only a few values has been considered in [13] and [14]. The main idea of [13] is to use stochastic sampling methods to look for cut sets in the graph structure that have high probability mass. The approach of [14] consists of predictive pruning to remove all but a few high probability nodes. Loosely speaking, these methods select the spikes with methods that differ from HFF's. Further, they are not guaranteed to improve the accuracy in theory as is the case for HFF.

Probabilistic dynamical models arise in many other settings in the study of biochemical networks. In particular,

there is a rich vein of work based on Continuous Time Markov Chains [15], [16], [17], [18], [19], [20], [21], [22]. Typically, in these studies one maintains the *number* of molecules of each species in the network and tracks the changes in these numbers as reactions one at a time. This approach is mandated when the number of molecules involved in the pathway is too low to support the “smoothness” conditions required by ODE-based models. As might be expected, the exact inferencing problem for large CTMCs is also computationally infeasible. Analysis methods based on Monte Carlo simulations [23], [24], [21], [25], probabilistic model checking [15], [26] as well as numerically solving the Chemical Master Equation describing a CTMC [20], [27] are being developed. In these studies, the CTMCs are presented implicitly but the analysis is carried out on the full state space of the CTMCs. Consequently, most of the current analysis methods based on CTMCs do not cope well with large biochemical networks (say with 30 or more species). Further, the associated parameter estimation problem is rarely addressed.

Our overall approach may be viewed as Monte Carlo based but with the crucial additional feature that we generate a pool of representative trajectories just *once* and compactly store these trajectories as a DBN. Then using this representation and *approximate* Bayesian inferencing, we carry out tasks such as parameter estimation and sensitivity analysis [4]. At present it is not clear whether approximate inferencing algorithms such as FF, BK, and HFF can be deployed to analyze the CTMC-based models of biochemical networks cited above.

1.2 Plan of the Paper

In the next section, we introduce DBNs and explain in more detail how they arise in our context as models of biochemical networks. In Section 3, we sketch the FF algorithm and then present the HFF algorithm. This is followed by an error analysis for the two algorithms. The experimental results are presented in Section 4 and we conclude with a discussion in Section 5. Additional technical material and details concerning the case studies can be found in [28].

This is an expanded and improved version of the conference paper [29]. The present paper contains the main proofs and many additional technical details including a complete error analysis of FF and HFF. In the experimental part of the work, we have used here a more sophisticated sampling method—as explained in Section 4—to generate the initial states of the trajectories. Further, we have validated the quality of the DBN approximations in all our case studies.

2 DBN MODELS OF BIOPATHWAYS

We start with an introduction to DBNs and related notions. We fix an ordered set of n random variables $\{X_1, \dots, X_n\}$ and let i, j range over $\{1, 2, \dots, n\}$. We denote by $\mathbf{X}$ the tuple $(X_1, \dots, X_n)$. The random variables are assumed to take values from the finite set V of cardinality K . We let x_i, u_i, v_i to denote a value taken by X_i . In the present setting, we need to consider only finite DBNs with no hidden variables.

A *Dynamic Bayesian Network* is a structure $\mathcal{D} = (\mathcal{X}, T, Pa, \{C_i^t\})$ where,

- T is a positive integer with t ranging over the set of time points $\{0, 1, \dots, T\}$.
- $\mathcal{X} = \{X_i^t \mid 1 \leq i \leq n, 0 \leq t \leq T\}$ is the set of random variables. As usual, these variables will be identified with the nodes of the DBN. X_i^t is the instance of X_i at time slice t .
- Pa assigns a set of parents to each node and satisfies: 1) $Pa(X_i^0) = \emptyset$, 2) If $X_j^{t'} \in Pa(X_i^t)$ then $t' = t - 1$. 3) If $X_j^{t-1} \in Pa(X_i^t)$ for some t then $X_j^{t-1} \in Pa(X_i^{t'})$ for every $t' \in \{1, 2, \dots, T\}$. Thus, the way nodes at the $(t-1)$ th time slice are connected to nodes at the t th time slice remains invariant as t ranges over $\{1, 2, \dots, T\}$.
- C_i^t is the *Conditional Probability Table* associated with node X_i^t specifying the probabilities $P(X_i^t \mid Pa(X_i^t))$. Suppose $Pa(X_i^t) = \{X_{j_1}^{t-1}, X_{j_2}^{t-1}, \dots, X_{j_m}^{t-1}\}$ and $(x_{j_1}, x_{j_2}, \dots, x_{j_m}) \in V^m$. Then as usual we require,

$$\sum_{x_i \in V} C_i^t(x_i \mid x_{j_1}, x_{j_2}, \dots, x_{j_m}) = 1.$$

Thus, the *structure* of our DBNs will be time invariant, in the sense that, the way nodes at the $(t-1)$ th time slice are connected to nodes at the t th time slice will remain the same as t ranges over $\{1, 2, \dots, T\}$. However, we do not require the local probabilistic dynamics as captured by the CPTs to be time invariant. More precisely, for the random variable X_i , we allow the CPT C_i^t to be different from the CPT $C_i^{t'}$ if $t \neq t'$. This is despite the fact that $Pa(X_i^t)$ can be obtained from $Pa(X_i^{t'})$ by replacing every appearance of t' by t in $Pa(X_i^{t'})$. This flexibility is required due to the discretization of the value space and the fact that the CPTs will be representing local probabilistic dynamics induced the ODE dynamics. This will become clear in the next section.

A state of the DBN at t will be a member of V^n , say $\mathbf{x} = (x_1, x_2, \dots, x_n)$ specifying that $X_i^t = x_i$ for $1 \leq i \leq n$. This in turn stands for $X_i = x_i$ for $1 \leq i \leq n$ at t . Suppose $Pa(X_i^t) = \{X_{j_1}^{t-1}, X_{j_2}^{t-1}, \dots, X_{j_m}^{t-1}\}$. Then a CPT entry of the form $C_i^t(x_i \mid x_{j_1}, x_{j_2}, x_{j_m}) = p$ says that if the system is in a state at $t-1$ in which $X_{j_l} = x_{j_l}$ for $1 \leq l \leq m$, then the probability of $X_i = x_i$ being the case at t is p . In this sense the CPTs specify the probabilistic dynamics locally.

The regular structure of our DBNs induces the function PA given by: $X_j \in PA(X_i)$ iff $X_j^{t-1} \in Pa(X_i^t)$. We define $\hat{i} = \{j \mid X_j \in PA(X_i)\}$ to capture Pa in terms of the corresponding indices.

In the following, $\mathbf{x}_{\mathcal{I}}$ will denote a vector of values over the index set $\mathcal{I} \subseteq \{1, 2, \dots, n\}$. It will be viewed as a map $\mathbf{x}_{\mathcal{I}} : \mathcal{I} \rightarrow V$. We will often denote $\mathbf{x}_{\mathcal{I}}(i)$ as $\mathbf{x}_{\mathcal{I},i}$ or just $\mathbf{x}_i$ if $\mathcal{I}$ is clear from the context. If $\mathcal{I} = \{i\}$ and $\mathbf{x}_{\mathcal{I}}(i) = x_i$, we will identify $\mathbf{x}_{\mathcal{I}}$ with x_i . If $\mathcal{I}$ is the full index set $\{1, 2, \dots, n\}$, we will simply write $\mathbf{x}$. Further, we denote by $\mathbf{X}^t$ the vector of random variables $(X_1^t, \dots, X_n^t)$.

Using these notations, we can write $C_i^t(x_i \mid \mathbf{u}_{\hat{i}}) = p$ to mean that p is the probability that $X_i = x_i$ at time t given that at time $t-1$, $X_{j_1} = \mathbf{u}_{j_1}, X_{j_2} = \mathbf{u}_{j_2}, \dots, X_{j_m} = \mathbf{u}_{j_m}$ with $\hat{i} = \{j_1, j_2, \dots, j_m\}$.

The probability distribution $P(X_1^t, X_2^t, \dots, X_n^t)$ describes the possible states of the system at time t . In other words, $P(\mathbf{X}^t = \mathbf{x})$ is the probability that the system will reach the state $\mathbf{x}$ at t . Starting from $P(\mathbf{X}^0)$ at time 0, given by

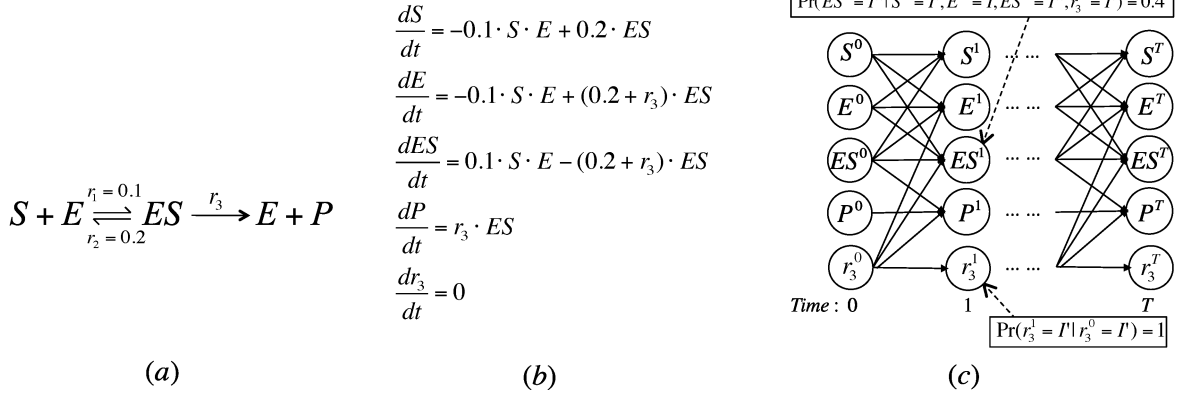


Fig. 1. (a) The enzyme catalytic reaction network. (b) The ODE model. (c) The DBN approximation.

$P(\mathbf{X}^0 = \mathbf{x}) = \prod_i C_i^0(\mathbf{x}_i)$, one would like to compute $P(X_1^t, \dots, X_n^t)$ for a given t .

We can use the CPTs to inductively compute this:

$$P(\mathbf{X}^t = \mathbf{x}) = \sum_{\mathbf{u}} P(\mathbf{X}^{t-1} = \mathbf{u}) \left(\prod_i C_i^t(\mathbf{x}_i | \mathbf{u}_i) \right), \quad (1)$$

with $\mathbf{u}$ ranging over V^n .

Since $|V| = K$, the number of possible states at t is K^n . Hence, explicitly computing and maintaining the probability distributions is feasible only if n is small or if the underlying graph of the DBN falls apart into many disjoint components. Neither restriction is realistic, and hence, one needs approximate ways to maintain $P(\mathbf{X}^t)$ compactly and compute it efficiently. Before we get into methods for doing this, we shall describe how we use DBNs to model the dynamics of biopathways.

2.1 DBN Models of Biopathways

Biological pathways are often described as a network of biochemical reactions. Fig. 1a shows a simple network consisting of 3 reactions; a pair of reversible reactions and one irreversible reaction. The dynamics of such a network can be modeled as a system of ODEs; one equation of the form $\frac{dy}{dt} = f(\mathbf{y}, \mathbf{r})$ for each molecular species y , with f describing the kinetics of the reactions that produce and consume y , while $\mathbf{y}$ is the set (vector) of molecular species taking part in these reactions and $\mathbf{r}$ is the rate constants associated with these reactions. The term corresponding to each reaction will be determined by the kinetic law governing this reaction. Fig. 1b illustrates this idea, where we have assumed that the kinetics of all three reactions is governed by the mass law [30], which states that the rate at which a reaction proceeds is directly proportional to the current concentration levels of the reactants taking part in the reaction. Thus the rate at which the forward reaction produces the enzyme substrate complex ES from the substrate S and the enzyme E is directly proportional to the current concentrations of E and S . Further the constant of proportionality, the *rate constant* for this reaction, is given to be 0.1 in this example. This produces the term $0.1 \times S \times E$ in the equation for S which will capture the rate at which S is being depleted due to the forward reaction. Similarly, the term $0.2 \times ES$ will capture the rate which S is being produced by the reverse reaction where we are given that

the rate constant for this reaction is 0.2. (the other aspects of Fig. 1 will be explained later in this section).

Due to the nature of kinetic laws governing biochemical reactions, the ODE models of biopathways will be non-linear. They will also be high dimensional and hence closed form solutions will not be available. Hence, one must resort to numerical simulations. For carrying tasks such as sensitivity analysis, the number of simulations one must carry out can be very large. Further, many of the rate constant values and initial concentrations will be unknown and will have to be estimated by searching through a high dimensional parameter value space. For each “guessed” value of the parameters one will have to compare the dynamics generated through numerical simulations with experimental data. For high dimensional systems this will again result in a very large number of simulations. Further the experimental data available for calibrating the parameters will be very limited in quantity and precision.

Motivated by these considerations, a method for approximating a system of ODEs as a dynamic Bayesian network was developed in [1]. The discretized nature of DBN helps to bridge the gap between the lack of precision in experimental data and the precise nature of ODE solutions. Further, the DBN can be analyzed using Bayesian inferencing techniques. As demonstrated in [1] the one-time cost of building the DBN is easily amortized by carrying out parameter estimation, sensitivity analysis, and other tasks on the DBN approximation.

The first step in this approximation procedure is to discretize the time domain. For biopathways, experimental data will be available only for a few time points with the value measured at the final time point typically signifying the steady state value. Hence, we assume the dynamics is of interest only for discrete time points and, furthermore, only up to a maximal time point. We denote these time points as $\{0, 1, \dots, T\}$. It is *not* necessary to uniformly discretize the time domain though we shall often do so for convenience.

Next, we assume that the values of the variables can be observed with only finite precision and accordingly partition the range of each variable y_i into a set of intervals I_i (I_i). Again, it is not necessary to partition the value range of each variable evenly. The initial values as well as the parameters of the ODE system are assumed to be distributions (usually uniform) over certain intervals. We then sample the initial

states of the system many times [1] and generate a trajectory by numerical integration for each sampled initial state. The resulting set of trajectories is then treated as an approximation of the dynamics of ODE system.

Unknown rate constants are handled as additional variables. We assume that the minimum and maximum values of these unknown rate constant variables are known. We then partition these ranges of values also into a finite numbers of intervals, and fix a uniform distribution over *all* the intervals. For each such variable r_j we add the equation $\frac{dr_j}{dt} = 0$ to the system of ODEs. This will reflect the fact that once the initial value of a rate constant has been sampled, this value will not change during the process of generating a trajectory. Naturally, different trajectories can have different initial values.

A key idea is to compactly store the generated set of trajectories as a dynamic Bayesian network. This is achieved by exploiting the network structure and by simple counting. First, we specify one random variable Y_i for each variable y_i of the ODE model. For each unknown rate constant r_j , we add one random variable R_j . The node Y_k^{t-1} will be in $Pa(Y_i^t)$ iff $k = i$ or y_k appears in the equation for y_i . Further, the node R_j^{t-1} will be in $Pa(Y_i^t)$ iff r_j appears in the equation for y_i . On the other hand, R_j^{t-1} will be the only parent of the node R_j^t . In Fig. 1, we show a simple enzymatic reaction network, its ODE model and the structure of its DBN approximation. In this example, we have assumed that r_3 is the only unknown rate constant.

Suppose $Pa(Y_i^t) = \{Z_1^{t-1}, Z_2^{t-1}, \dots, Z_k^{t-1}\}$. Then a CPT entry of the form $C_i^t(I | I_1, I_2, \dots, I_k) = p$ says that p is the probability of the value of y_i falling in the interval I at time t , given that the value of Z_j was in I_j for $1 \leq j \leq k$. The probability p is calculated through simple counting. Suppose N is the number of generated trajectories. We record, for how many of these trajectories, the value of Z_j falls in the interval I_j simultaneously for each $j \in \{1, 2, k\}$. Suppose this number is J . We then determine for how many of these J trajectories, the value of Y_i falls in the interval I at time t . If this number is J' then p is set to be $\frac{J'}{J}$ (It should now be clear why $C_i^t(I | I_1, I_2, \dots, I_k)$ will be in general different from $C_i^{t'}(I | I_1, I_2, \dots, I_k)$ if $t \neq t'$).

If r_j is an unknown rate constant, in the CPT of R_j^t we will have $P(R_j^t = I | R_j^{t-1} = I') = 1$ if $I = I'$ and $P(R_j^t = I | R_j^{t-1} = I') = 0$, otherwise. This is because the sampled initial value of r_j does not change during numerical integration. Suppose r_j appears on the right hand side of the equation for y_i and $Pa(Y_i^t) = \{Z_1^{t-1}, Z_2^{t-1}, \dots, Z_\ell^{t-1}\}$ with $Z_\ell^{t-1} = R_j^{t-1}$. Then for each choice of interval values for nodes other than R_j^{t-1} in $Pa(Y_i^t)$ and for each choice of interval value $\hat{I}$ for r_j there will be an entry in the CPT of Y_i^t of the form $P(y_i^t = I | Z_1^{t-1} = I_1, Z_2^{t-1} = I_2, \dots, R_j^{t-1} = \hat{I}) = p$. This is so, since we will sample for all possible initial interval values for r_j . In this sense the CPTs record the approximated dynamics for all possible combinations of interval values for the unknown rate constants. These features are illustrated in Fig. 1c for the unknown rate constant r_3 .

After building this DBN, we use a Bayesian inference-based technique to perform parameter estimation to complete the construction of the model (the details can be

found in [1]). This will result in a calibrated DBN in which each unknown rate constant will have a specific interval value assigned to it. In a similar way unknown initial concentrations can also be handled.

The one time cost of constructing the DBN can be easily recovered through the substantial gains obtained in doing parameter estimation and sensitivity analysis [1]. This method can cope with large biochemical networks with many unknown parameters. It has been used to uncover new biological facts about the complement system [4] where the starting point was a ODE-based model with 71 unknown parameters. In these studies FF was used as the core inferencing algorithm. It was then that we began to consider its shortcomings and started to seek an improved version.

3 THE HYBRID FACTORED FRONTIER ALGORITHM

In this section, we present the HFF algorithm. Since it is a parametrized version of FF, we first present FF. In doing so we will use the notations developed in Section 2.

Approximate probability distributions will be called belief states and denoted by B, B^t , etc. Exact probability distributions will be denoted by P, P^t , etc. Formally, a belief state B is a map from $V^n \rightarrow [0, 1]$ such that $\sum_{\mathbf{u} \in V^n} B(\mathbf{u}) = 1$. Thus a belief state is just a probability distribution but it will be convenient to linguistically separate them.

The FF algorithm uses marginal functions to represent belief states. A marginal function is a map $M : \{1, \dots, n\} \times V \rightarrow [0, 1]$ such that $\sum_{v \in V} M(i, v) = 1$ for each i . In what follows, u, v will range over V while $\mathbf{u}$ and $\mathbf{v}$ will range over V^n . A belief state B induces the marginal function M_B via $M_B(i, v) = \sum_{\mathbf{u} | \mathbf{u}_i = v} B(\mathbf{u})$. On the other hand, from a marginal function M , one can obtain a belief state B_M via $B_M(\mathbf{u}) = \prod_i M(i, \mathbf{u}_i)$. From the above definitions it follows that for a marginal function M , we have $M_{B_M} = M$. That is, for any

$$\begin{aligned} i, v, M_{B_M}(i, v) &= \sum_{\mathbf{u} | \mathbf{u}_i = v} B_M(\mathbf{u}) = \sum_{\mathbf{u} | \mathbf{u}_i = v} \prod_j M(j, \mathbf{u}_j) \\ &= \prod_j \sum_{\mathbf{u} | \mathbf{u}_i = v} M(j, \mathbf{u}_j) = \left(\prod_{j | j \neq i} \sum_{\mathbf{u}_j} M(j, \mathbf{u}_j) \right) \cdot M(i, v) \\ &= M(i, v). \end{aligned}$$

On the other hand, for a belief state B , unless $B = B_M$, we may have $B_{M_B} \neq B$.

For a DBN $\mathcal{D} = (\mathcal{X}, T, Pa, \{C_i^t\})$ recall that $\hat{i} = \{j | X_j \in Pa(X_i)\}$ captures the set of indices of the parents of i . In what follows, $V_{\hat{i}}$ will denote the tuple of values defined by $\hat{i}$. Thus, with a slight abuse of notation, $\mathbf{u}, \mathbf{v}$ will be used to denote $|\hat{i}|$ -dimensional vectors of values over V .

Given a DBN $\mathcal{D} = (\mathcal{X}, T, Pa, \{C_i^t\})$, FF computes inductively a sequence M^t of marginal functions as:

- $M^0(i, u) = C_i^0(u)$,
- $M^t(i, u) = \sum_{\mathbf{v} \in V_{\hat{i}}} [\prod_{j \in \hat{i}} M^{t-1}(j, \mathbf{v}_j)] C_i^t(u | \mathbf{v})$.

It is easy to check that these are indeed marginal functions, i.e., $\sum_{u \in V} M^t(i, u) = 1$ for all t and i . Thus FF maintains B^t , the belief state at t , compactly via the marginal function M^t . More precisely, $B^t(\mathbf{u}) = \prod_j M^t(j, \mathbf{u}_j) = B_{M^t}(\mathbf{u})$.

Let $t \geq 1$. Suppose that the DBN transforms the belief state B^{t-1} into the new belief state $\hat{B}^t$. In other words, $\hat{B}^t$ is the belief state obtained by performing $t - 1$ steps of FF and exact computation at the t th step. Then by (1), we have

$$\hat{B}^t(\mathbf{x}) = \sum_{\mathbf{u}} B^{t-1}(\mathbf{u}) \left(\prod_i C_i^t(\mathbf{x}_i | \mathbf{u}_i) \right). \quad (2)$$

However, the t th step of FF computes directly the marginal function M^t , which then represents the new belief state at time t as $B^t = B_{M^t}$. In general, $B^t \neq \hat{B}^t$, that is, the belief state B^t represented via M^t is an approximation of the belief state $\hat{B}^t$ as defined above. However, the computation of M^t is itself accurate in the following sense.

Proposition 1. For all $t \in \{1, \dots, T\}$, $M^t(i, v) = M_{\hat{B}^t}(i, v)$ for each i and v .

Proof. For $t > 0$, we have

$$\begin{aligned} M_{\hat{B}^t}(i, v) &= \sum_{\mathbf{v} | \mathbf{v}_i = v} \hat{B}^t(\mathbf{v}) \\ &= \sum_{\mathbf{v} | \mathbf{v}_i = v} \sum_{\mathbf{u}} \prod_j B^{t-1}(\mathbf{u}) (C_j^t(\mathbf{v}_j | \mathbf{u}_j)) \\ &\quad (\text{by (2)}) \\ &= \sum_{\mathbf{u}} B^{t-1}(\mathbf{u}) \sum_{\mathbf{v} | \mathbf{v}_i = v} \prod_j (C_j^t(\mathbf{v}_j | \mathbf{u}_j)) \\ &= \sum_{\mathbf{u}} B^{t-1}(\mathbf{u}) \left(\sum_{\mathbf{v}_n} C_n^t(\mathbf{v}_n | \mathbf{u}_n) \right) \dots \\ &\quad (C_i^t(v | \mathbf{u}_i)) \dots \left(\sum_{\mathbf{v}_1} C_1^t(\mathbf{v}_1 | \mathbf{u}_1) \right) \\ &= \sum_{\mathbf{u}} B^{t-1}(\mathbf{u}) (C_i^t(v | \mathbf{u}_i)). \end{aligned}$$

The last of the above equalities follows since each of the summands within the expression add up to one. Now, using $B^{t-1}(\mathbf{u}) = \prod_k M^{t-1}(k, \mathbf{u}_k)$ and splitting the above summation, we obtain

$$\begin{aligned} M_{\hat{B}^t}(i, v) &= \sum_{\mathbf{u} \in V_i} \sum_{\mathbf{u} \notin V_i} \prod_k M^{t-1}(k, \mathbf{u}_k) (C_i^t(v | \mathbf{u}_i)) \\ &= \sum_{\mathbf{u} \in V_i} \prod_{k \in \hat{i}} M^{t-1}(k, \mathbf{u}_k) (C_i^t(v | \mathbf{u}_i)) \\ &\quad \sum_{\mathbf{u} \notin V_i} \prod_{k \notin \hat{i}} M^{t-1}(k, \mathbf{u}_k) \\ &= \sum_{\mathbf{u} \in V_i} \prod_{k \in \hat{i}} M^{t-1}(k, \mathbf{u}_k) (C_i^t(v | \mathbf{u}_i)) \\ &\quad \prod_{k \notin \hat{i}} \sum_{\mathbf{u}_k} M^{t-1}(k, \mathbf{u}_k) \\ &= \sum_{\mathbf{u} \in V_i} \prod_{k \in \hat{i}} M^{t-1}(k, \mathbf{u}_k) (C_i^t(v | \mathbf{u}_i)) \\ &= M^t(i, v). \end{aligned}$$

The second factor above is just a product of 1s (by the definition of marginals) and the proposition follows. $\square$

As B^0 is accurate by definition, M^1 will also be accurate but not necessarily B^1 . A simple but crucial observation is that whenever the error $\max_{\mathbf{u} \in V^n} \{|\hat{B}^t(\mathbf{u}) - B^t(\mathbf{u})|\}$ incurred by FF

at step $t > 0$ (ignoring the error made in the previous steps) is large for some $\mathbf{u}$ then $M^t(i, \mathbf{u}_i)$ is large for every i . This is so since, $M^t(j, \mathbf{u}_j) = M_{\hat{B}^t}(j, \mathbf{u}_j) \geq \max(\hat{B}^t(\mathbf{u}), B^t(\mathbf{u}))$, which follows from Proposition 1 and the definition of marginals.

A second important observation is that there can only be a few instances of $\mathbf{u}$ such that $M^t(i, \mathbf{u}_i)$ is large for every i . For instance, there can be only one such $\mathbf{u}$ if we want $M^t(i, \mathbf{u}_i) > \frac{1}{2}$ for every i . Hence, by computing $\hat{B}^t(\mathbf{u})$ for a small subset of V^n for which M^t is high for all dimensions and maintaining it explicitly, one can hope to reduce the one step error incurred FF and hence the overall error too. This is the intuition underlying the HFF algorithm.

3.1 The Hybrid Factored Frontier Algorithm

The overall structure of HFF is as follows: starting with $t = 0$, we inductively compute and maintain the tuple $(M^t, S^t, B_H^t, \alpha^t)$, where

- M^t is a marginal function.
- $S^t \subseteq V^n$ is a set of tuples called *spikes*.
- $B_H^t : V^n \rightarrow [0, 1]$ is a function s.t. $B_H^t(\mathbf{u}) = 0$ if $\mathbf{u} \notin S^t$ and $\sum_{\mathbf{u} \in S^t} B_H^t(\mathbf{u}) < 1$.
- $\alpha^t = \sum_{\mathbf{u} \in S^t} B_H^t(\mathbf{u})$.

This hybrid state $(M^t, S^t, B_H^t, \alpha^t)$ represents the following belief state B^t :

$$\begin{aligned} B^t(\mathbf{u}) &= B_H^t(\mathbf{u}) + (1 - \alpha^t) \prod_i M_H^t(i, \mathbf{u}_i), \text{ where} \\ M_H^t(i, v) &= [M^t(i, v) - \sum_{\{\mathbf{u} \in S^t | \mathbf{u}_i = v\}} B_H^t(\mathbf{u})] / (1 - \alpha^t). \end{aligned}$$

The first component of $B^t(\mathbf{u})$ is the probability mass $B_H^t(\mathbf{u})$ of the spike (if $\mathbf{u}$ is not a spike, $B_H^t(\mathbf{u}) = 0$). The second component is the product of (uniformized) marginals $M_H^t(i, v)$, as in FF. Notice that we need to use M_H^t rather than M^t since the cumulative weight of the contribution made by the spikes needs to be discounted from M^t . The coefficient $(1 - \alpha^t)$ must be used first to ensure that M_H^t is a marginal function, and second to ensure that B^t is a belief state, as will be demonstrated subsequently.

The HFF Algorithm. We initialize with $M^0 = C^0$, $S^0 = \emptyset$, $B_H^0 = \mathbf{0}$, and $\alpha^0 = 0$ and fix a parameter σ . This σ will be the number of spikes we choose to maintain. It is a crucial parameter as our results will show. We inductively compute $(M^{t+1}, S^{t+1}, B_H^{t+1}, \alpha^{t+1})$ from $(M^t, S^t, B_H^t, \alpha^t)$ as follows.

Step 1: We first compute M^{t+1} as

$$\begin{aligned} M^{t+1}(i, x) &= \sum_{\mathbf{u} \in S^t} [B_H^t(\mathbf{u}) \times C_i^{t+1}(x | \mathbf{u}_i)] \\ &\quad + (1 - \alpha^t) \left(\sum_{\mathbf{u}_i} \left[\prod_{j \in \hat{i}} M_H^t(j, \mathbf{u}_j) \times C_i^{t+1}(x | \mathbf{u}_i) \right] \right). \end{aligned}$$

Step 2: We next compute a set S^{t+1} of at most σ spikes using M^{t+1} . We want to consider as spikes $\mathbf{u} \in V^n$ where $M^{t+1}(i, \mathbf{u}_i)$ is large for every i . To do so, we find a constant η^{t+1} such that $M^{t+1}(i, \mathbf{u}_i) \geq \eta^{t+1}$ for every i for a subset of V^n containing σ elements and for all other $\mathbf{u}'$, there exists i with $M^{t+1}(i, \mathbf{u}'_i) < \eta^{t+1}$. We compute η^{t+1} via binary search. First we fix the precision with which we want to compute

η^{t+1} to be ξ . We have found $\xi = 10^{-6}$ to be a good choice. For this choice there will be at most 20 iterations of the loop described below. The search for η^{t+1} proceeds as follows:

- $\eta_1 = 0$ and $\eta_2 = 1$.
- While $\eta_2 - \eta_1 > \xi$ do
 - $\eta = \frac{\eta_1 + \eta_2}{2}$.
 - Determine the set of values U_i such that $v \in U_i$ iff $M^{t+1}(i, v) > \eta$.
 - Set a_i to be the cardinality of U_i .
 - If $\prod_i (a_i) > \sigma$ then $\eta_1 = \eta$; otherwise $\eta_2 = \eta$
- endwhile
- Return $\eta^{t+1} = \eta_2$ and $S^{t+1} = \prod_i U_i$.

Step 3: Finally, we compute $B_H^{t+1}(\mathbf{u})$ for each $\mathbf{u}$ in S^{t+1} as follows, by only taking into account the contribution of the current spikes.

$$B_H^{t+1}(\mathbf{u}) = \sum_{\mathbf{v} \in S^t} \left(B^t(\mathbf{v}) \times \prod_i C_i^{t+1}(\mathbf{u}_i | \mathbf{v}_i) \right).$$

End of Algorithm. As in the case of FF, we denote by $\hat{B}^{t+1}$ the belief state obtained from B^t through an exact step of the DBN

$$\hat{B}^{t+1}(\mathbf{u}) = \sum_{\mathbf{v} \in V^n} \left(B^t(\mathbf{v}) \times \prod_i C_i^{t+1}(\mathbf{u}_i | \mathbf{v}_i) \right).$$

Notice that $B_H^{t+1}(\mathbf{u}) \leq \hat{B}^{t+1}(\mathbf{u})$ for all $\mathbf{u}$. We recall that T is the number of time points, σ the number of spikes, n the number of variables, V is the set of values with $K = |V|$, and D be the maximum in degree of the DBN graph.

Theorem 2. HFF has the following properties.

1. if $\sigma = 0$, the HFF algorithm is the same as FF and if $\sigma = K^n$, it is the exact algorithm.
2. $M^t(i, v) = M_{B^t}^t(i, v)$ for every v . Further, B^t is a belief state while M_H^t and M^t are marginal functions, for every t .
3. The time complexity of HFF is $O(T \cdot n \cdot (\sigma^2 + K^{D+1}))$.

Proof. $\sigma = 0$ implies that the set of spikes $S^t = \emptyset$ for all t . This implies that $\alpha^t = 0$ and the computation done by HFF is the same as FF. If $\sigma = K^n$, then $S^t = V$ for all t and $\alpha^t = 1$ (of course, M_H^t is then not computed). Thus, this boils down to perform exact inferencing. We have now established part (1).

We prove that for all $t \geq 1$, if B^{t-1} is a belief state and M^{t-1}, M_H^{t-1} are marginals, then $M^t = M_{B^t}^t$ and B^t is a belief state and M^t, M_H^t are marginals. We thus obtain part (2) by induction on t , using the fact that B^0 is a belief state and M^0, M_H^0 are marginals by definitions. For $t \geq 0$, let $M^t(i, v), M_H^t(i, v)$ be marginals and B^t be a belief state. Then at $t+1$, let us start by proving $M_{B^{t+1}}^{t+1}(i, v) = M^{t+1}(i, v)$. The first step is the same as in Proposition 1

$$\begin{aligned} M_{B^{t+1}}^{t+1}(i, v) &= \sum_{\mathbf{v} | \mathbf{v}_i = v} \hat{B}^{t+1}(\mathbf{v}) \\ &= \sum_{\mathbf{u}} B^t(\mathbf{u}) (C_i^{t+1}(v | \mathbf{u}_i)) \text{ (by 2)}. \end{aligned}$$

Now, however, the definition of B^t is different for HFF and so we have from part

$$\begin{aligned} M_{B^{t+1}}^{t+1}(i, v) &= \sum_{\mathbf{u}} B_H^t(\mathbf{u}) (C_i^{t+1}(v | \mathbf{u}_i)) \\ &\quad + (1 - \alpha^t) \sum_{\mathbf{u}} \left(\prod_{k=1}^n M_H^t(k, \mathbf{u}_k) \right) \\ &\quad (C_i^{t+1}(v | \mathbf{u}_i)). \end{aligned} \quad (3)$$

But if $\mathbf{u} \notin S^t$, then $B_H^t(\mathbf{u}) = 0$. Further splitting the second term as in Proposition 1, we obtain

$$\begin{aligned} M_{B^{t+1}}^{t+1}(i, v) &= \sum_{\mathbf{u} \in S^t} B_H^t(\mathbf{u}) (C_i^{t+1}(v | \mathbf{u}_i)) \\ &\quad + (1 - \alpha^t) \sum_{\mathbf{u} \in V_i} \sum_{\mathbf{u} \notin V_i} \prod_k M_H^t(k, \mathbf{u}_k) \\ &\quad (C_i^{t+1}(v | \mathbf{u}_i)) \\ &= \sum_{\mathbf{u} \in S^t} B_H^t(\mathbf{u}) (C_i^{t+1}(v | \mathbf{u}_i)) + (1 - \alpha^t) \\ &\quad \sum_{\mathbf{u} \in V_i} \prod_{k \in \hat{i}} M_H^t(k, \mathbf{u}_k) (C_i^{t+1}(v | \mathbf{u}_i)) \\ &\quad \sum_{\mathbf{u} \notin V_i} \prod_{k \notin \hat{i}} M_H^t(k, \mathbf{u}_k) \\ &= \sum_{\mathbf{u} \in S^t} B_H^t(\mathbf{u}) (C_i^{t+1}(v | \mathbf{u}_i)) + (1 - \alpha^t) \\ &\quad \sum_{\mathbf{u} \in V_i} \prod_{k \in \hat{i}} M_H^t(k, \mathbf{u}_k) (C_i^{t+1}(v | \mathbf{u}_i)) \\ &\quad \prod_{k \notin \hat{i}} \sum_{\mathbf{u}_k} M_H^t(k, \mathbf{u}_k) \\ &= \sum_{\mathbf{u} \in S^t} B_H^t(\mathbf{u}) (C_i^{t+1}(v | \mathbf{u}_i)) + (1 - \alpha^t) \\ &\quad \sum_{\mathbf{u} \in V_i} \prod_{k \in \hat{i}} M_H^t(k, \mathbf{u}_k) (C_i^{t+1}(v | \mathbf{u}_i)) \times 1 \\ &= M^{t+1}(i, v). \end{aligned}$$

In the step above, $\sum_{\mathbf{u}_k} M_H^t(k, \mathbf{u}_k) = 1$ follows from our inductive hypothesis that M_H^t is a marginal. Now, we will prove the remainder of this part, i.e., M^{t+1}, M_H^{t+1} are marginals and B^{t+1} is a belief state. For all i , again from $\alpha^t = \sum_{\mathbf{u} \in S^t} B_H^t(\mathbf{u})$ and $\sum_{v \in V} C_i^{t+1}(v | \mathbf{u}_i) = 1$, we have

$$\begin{aligned} \sum_{v \in V} M^{t+1}(i, v) &= \sum_{\mathbf{u} \in S^t} \left[\sum_{v \in V} C_i^{t+1}(v | \mathbf{u}_i) \times B_H^t(\mathbf{u}) \right] \\ &\quad + (1 - \alpha^t) \left(\sum_{\mathbf{u}_i} \left[\sum_{v \in V} C_i^{t+1}(v | \mathbf{u}_i) \times \prod_{j \in \hat{i}} M_H^t(j, \mathbf{u}_j) \right] \right) \\ &= \sum_{\mathbf{u} \in S^t} B_H^t(\mathbf{u}) + (1 - \alpha^t) \sum_{\mathbf{u}_i} \prod_{j \in \hat{i}} M_H^t(j, \mathbf{u}_j) \\ &= \alpha^t + (1 - \alpha^t) \prod_{j \in \hat{i}} \sum_{\mathbf{u}_j} M_H^t(j, \mathbf{u}_j) = 1. \end{aligned}$$

Now, using the above and $\alpha^{t+1} = \sum_{\mathbf{u} \in S^{t+1}} B_H^{t+1}(\mathbf{u})$ (assuming $\alpha^{t+1} \neq 1$), we have

$$\begin{aligned}
& \sum_{v \in V} M_H^{t+1}(i, v) \\
&= \left(\sum_{v \in V} M^{t+1}(i, v) - \sum_{v \in V} \sum_{\mathbf{u} \in S^{t+1} | \mathbf{u}_i = v} B_H^{t+1}(\mathbf{u}) \right) \times \frac{1}{1 - \alpha^{t+1}} \\
&= \left(1 - \sum_{\mathbf{u} \in S^{t+1}} B_H^{t+1}(\mathbf{u}) \right) \frac{1}{1 - \alpha^{t+1}} = 1 \\
&\quad \sum_{\mathbf{u} \in V^n} B^{t+1}(\mathbf{u}) = \sum_{\mathbf{u} \in V^n} B_H^{t+1}(\mathbf{u}) \\
&\quad + (1 - \alpha^{t+1}) \sum_{\mathbf{u} \in V^n} \prod_i M_H^{t+1}(i, \mathbf{u}_i) \\
&= \sum_{\mathbf{u} \in S^{t+1}} B_H^{t+1}(\mathbf{u}) + (1 - \alpha^{t+1}) \times 1 = 1.
\end{aligned}$$

It now follows easily that for any i, v ,

$$1 \geq M^{t+1}(i, v) \geq 0,$$

and $1 \geq M_H^{t+1}(i, v)$. It remains to prove that

$$M_H^{t+1}(i, v) \geq 0,$$

that is $M^{t+1}(i, v) \geq \sum_{\mathbf{u} \in S^{t+1} | \mathbf{u}_i = v} B_H^{t+1}(\mathbf{u})$. As $B_H^{t+1}(\mathbf{u}) \leq \hat{B}^{t+1}(\mathbf{u})$ for all $\mathbf{u}$, we have,

$$\begin{aligned}
\sum_{\mathbf{u} \in S^{t+1} | \mathbf{u}_i = v} B_H^{t+1}(\mathbf{u}) &\leq \sum_{\mathbf{u} \in S^{t+1} | \mathbf{u}_i = v} \hat{B}^{t+1}(\mathbf{u}) \\
&= \hat{B}_{B^{t+1}}^{t+1}(i, v) = M^{t+1}(i, v),
\end{aligned}$$

which completes the proof of part (2).

Turning to part (3), we note that at each time point, the step 1 of HFF has the same complexity as FF together with the spikes contributing: $O(K \cdot n \cdot (K^D + \sigma))$. Step 2 makes at most $K \times n$ comparisons for each iteration of the loop and there are only a constant number of iterations of the loop. Thus the complexity of this step per time point is $O(K \times n)$. Step 3 computes for each spike, $B^t(u)$ from the values of $B_H^t(u)$ and $M^t(i, u)$ which takes $O(Kn + \sigma n)$. Then, we sum over all the spikes the value computed by multiplying n values of the CPT which takes $O(\sigma \times n)$. Thus, this step overall takes $O(\sigma Kn + n\sigma^2)$. Hence the overall time complexity of HFF is the sum of all these quantities which is $O(T \cdot n \cdot (K^{D+1} + K\sigma + \sigma^2))$ which is bounded by $O(T \cdot n \cdot (K^{D+1} + \sigma^2))$. $\square$

HFF gathers in one sweep—just as FF does—the required information about the belief states. However, it can take more time than FF depending on the number of spikes but the added complexity is only quadratic in the number of spikes.

3.2 Error Analysis

It is easy to see that with each time slice t of the DBN one can associate a stochastic matrix $\mathcal{T}_t$. This stochastic matrix will capture the transformation of probability distributions effected by the n CPTs associated with the time slice t as dictated by (1) in Section 2. In particular, we will have $P(X^t) = \mathcal{T}_t(P(X^{t-1}))$.

We now denote the cumulative error at t as Δ^t and define it to be: $\Delta^t = \max_{\mathbf{u} \in V^n} (|P(X^t = \mathbf{u}) - B^t(\mathbf{u})|)$. Toward deriving an upper bound for Δ^t , we first note that Markov

chain theory (for instance, using the Dobrushin's coefficient, see chapter 6.7 in [31]) guarantees the following:

Theorem 3. Let $\mathcal{T}$ be an n -dimensional stochastic matrix. Then for two probability distributions A, B , we have $\|\mathcal{T}(A) - \mathcal{T}(B)\|_\infty \leq \beta_{\mathcal{T}} \|A - B\|_\infty$ where $0 \leq \beta_{\mathcal{T}} \leq 1$ is a constant that depends only on $\mathcal{T}$.

$\beta_{\mathcal{T}}$ is called the contraction factor. In what follows we shall write β^t for the contraction factor associated with $\mathcal{T}_t$ and set $\beta = \max_t \beta^t$.

An implicit assumption in what follows is that $\beta < 1$. As pointed out in [3] this is a very reasonable assumption since it fails for the extreme case where the variables are completely decoupled and are independent. The case studies we report in Section 4 also easily satisfy this assumption. When $\beta < 1$, due to the theorem above the maximum error will reduce by a factor of β at each step as we step through t starting from $t = 0$. Hence the cumulative error will stabilize rapidly.

Now following a reasoning similar to [3] we shall show that Δ^t can be bounded by $\epsilon_0 (\sum_{j=0}^t \beta^j)$ where ϵ_0 is the maximum one step error given by $\epsilon_0 = \max_t \|\mathcal{T}_t(B^{t-1}) - \mathcal{T}_t(B^{t-1})\|_\infty$. Notice that $\mathcal{T}_t(B^{t-1})$ was denoted as $\hat{B}^t$ in previous sections.

Lemma 4. $\Delta^t \leq \epsilon_0 (\sum_{j=0}^t \beta^j)$. Further if $\beta < 1$, we have $\Delta^t \leq \frac{\epsilon_0}{1-\beta}$.

Proof. By definition of overall error and the above stated property of Markov chains,

$$\begin{aligned}
\Delta^t &= |B^t - P(X^t)| \\
&\leq |B^t - \mathcal{T}_t(B^{t-1})| + |\mathcal{T}_t(B^{t-1}) - \mathcal{T}_t(P(X^{t-1}))| \\
&\leq \epsilon_0 + \beta_t \Delta^{t-1}.
\end{aligned}$$

Then by recursively computing the second factor, we obtain,

$$\begin{aligned}
\Delta^t &\leq \epsilon_0 + \beta_t \epsilon_0 + \beta_t \beta_{t-1} \epsilon_0 + \dots + (\beta_t \beta_{t-1} \dots \beta_1) \epsilon_0 \\
&\leq \epsilon_0 (\sum_{j=0}^t \beta^j).
\end{aligned}$$

Further if $\beta < 1$, we have

$$\Delta^t \leq \epsilon_0 \left(\sum_{j=0}^t \beta^j \right) \leq \epsilon_0 \left(\sum_{j=0}^{\infty} \beta^j \right) = \frac{\epsilon_0}{1-\beta}.$$

$\square$

We note that $\sum_{j=0}^t \beta^j$ depends only on the DBN. Hence, theoretically comparing the error behaviors of FF and HFF amounts to comparing their single step errors. To do so, we shall next analyze single step error of FF followed by that of HFF.

Recall that for FF, $B^t = B_{M_{\wedge}^t}$. Thus the one-step error incurred by FF at step t is $\max_{\mathbf{u} \in V^n} \{|\hat{B}^t(\mathbf{u}) - B_{M_{\wedge}^t}(\mathbf{u})|\}$. We can bound this from above by ϵ_0 where $\epsilon_0 = \max_{\mathbf{u} \in V^n} \{|\mathcal{B}(\mathbf{u}) - B_{M_B}(\mathbf{u})|\}$ with $\mathcal{B}$ ranging over the set of all possible belief states and $\mathbf{u}$ ranging over V^n . It turns out that ϵ_0 can be made arbitrarily close to 1 as n , the number of variables, tends to ∞ . To see this, fix $0 < \delta < 1$ and consider the belief state B defined by $B(\mathbf{u}) = 1 - \delta$, $B(\mathbf{u}') = \delta$ for some $\mathbf{u}, \mathbf{u}' \in V^n$ such

that for all i , $\mathbf{u}_i \neq \mathbf{u}_j$ and $B(\mathbf{v}) = 0$ for all $\mathbf{v} \in V^n \setminus \{\mathbf{u}, \mathbf{u}^t\}$. Then, $M_B(i, \mathbf{u}_i) = 1 - \delta$ for all $i \in \{1, \dots, n\}$ and so $B_{M_B}(\mathbf{u}) = (1 - \delta)^n$. As a result, we have $\epsilon_0 = \max |B - B_{M_B}| \geq (1 - \delta) - (1 - \delta)^n$ which tends to $1 - \delta$ as n tends to ∞ . Now if we choose δ to be close to 0, $1 - \delta$ is close to 1. Thus ϵ_0 can be made as close to 1 as we want, with n tending to ∞ . We found that the cumulative errors made by FF can be large in practice too as shown in the next section.

Notice that for HFF too we have $B^t = B_{M_B}$, and its one step error can be bound by ϵ_0 . However, the spikes can be used to bound the single step error of HFF more precisely as follows:

Claim 1. *The one step error made by HFF is bounded by $\hat{\epsilon}_0$ with $\hat{\epsilon}_0 \leq \min\{(1 - \alpha), \eta\}$, where $\alpha = \min_t(\alpha^t)$ and $\eta = \max_t(\eta^t)$.*

Proof sketch. If α is large, then the value of $\hat{B}^t(\mathbf{u}) \leq 1 - \alpha$ for $\mathbf{u} \notin S^t$. Also, as $B^t(\mathbf{u}) \leq \hat{B}^t(\mathbf{u})$, we have $\hat{B}^t(\mathbf{u}) - B^t(\mathbf{u}) \leq 1 - \alpha$. Finally, if η is small, then by construction for all $\mathbf{u} \notin S^t$, $M^{t+1}(i, v) \leq \eta$ for some i with $\mathbf{u}_i = v$, and hence $\hat{B}^t(\mathbf{u}) \leq M^{t+1}(i, v) \leq \eta$. Also, $\hat{B}^t(\mathbf{u}) - B^t(\mathbf{u}) \leq \eta \times \sum_{v \notin S^t} \prod_i C_i^{t+1}(\mathbf{u}_i | \mathbf{v}_i) \leq \eta$ for $\mathbf{u} \in S^t$. $\square$

Thus, the worst case error for HFF with at least two spikes (implying $\eta < 1/2$) is smaller than for FF. Taking more spikes will increase α and decrease η , reducing the worst case error. Experiments in the next section show that the practical accuracy is also improved as we increase the number of spikes.

4 EXPERIMENTAL EVALUATION

We have implemented our algorithm in C++. The experiments reported here were carried out on an Opteron 2.2 Ghz Processor, with 3 GB memory. The algorithms were evaluated on five DBN models of biochemical networks: the small enzyme catalytic reaction network shown in Fig. 1 for initial experimentation, the EGF-NGF pathway [8] under 1) EGF-stimulation, 2) NGF-stimulation, and 3) costimulation of EGF and NGF, and the Epo mediated ERK signaling pathway. The ODE model for the EGF-NGF pathway was obtained from the BioModels database [32] and the Epo mediated ERK signaling pathway from [9]. For all these models, there were no unknown parameters and this enabled us to focus on the main issue of evaluating the performance of HFF. The DBNs were constructed using the method presented in Section 2 [1]. To improve the quality of the approximations for the large pathway models, we constructed the DBNs using the *equation-based subinterval sampling* method explained in more detail below. In what follows, we highlight the main findings of our experiments. More details can be found in the supplementary materials [28].

4.1 Enzyme Catalytic Kinetics

For initial validation, we started with the enzyme catalytic reaction network shown in Fig. 1 which has only four species/equations and three rate constants. The value space of each variable was divided into five equally wide intervals ($\{[0, 1], [1, 2], \dots, [4, 5]\}$). We assumed the initial distributions of variables to be uniform over certain intervals. We then fixed the time horizon of interest to be 10 minutes and divided

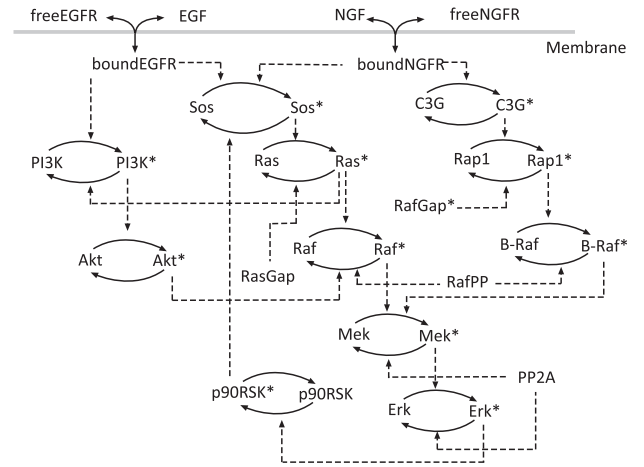


Fig. 2. EGF-NGF pathway.

this interval evenly into $[0, 1, \dots, 100]$ time points. The conditional probability tables associated with each node of the DBN were filled by generating 10^6 trajectories by direct random sampling over the initial states [1].

This being a small example, we could compute the marginal distributions for each species exactly. We ran FF and HFF (σ) with various choices of σ , the number of spikes. The resulting estimates were then compared against the exact marginals. We also ran the fully factored version of BK (which we call BK in this section), using the implementation provided in the Bayes Net Toolbox of MATLAB [33].

In what follows we report the errors in terms of the absolute difference between the marginal probabilities computed by the exact and approximate methods. Thus, if we say the error is 0.15 then this means that the actual marginal probability was p and the marginal probability computed by the approximate algorithm was p' with $|p - p'| = 0.15$.

Even for this small network, FF and BK deviated from some of the exact marginals by as much as 0.169. Fig. 2 shows the profile of the marginal distribution of E (the enzyme) assuming a value in the first interval as computed by FF, BK, HFF(64), and the exact method. The profiles of exact and HFF(64) were almost the same while FF and BK (whose curve practically coincides with that of FF and is hence not shown) make noticeable errors. The computation times for all the algorithms were negligible. The maximum error incurred for the four species taken over all the interval values and all time points was 0.169 for FF and 0.024 for HFF(16) and 0.003 for HFF(64). Further, the number of errors greater than 0.1 taken over all the species, intervals, and time points reduced from 72 for FF to 0 for HFF(16).

4.2 The Large Pathway Models

As explained in Section 2, during the construction of the DBN we assume that the initial values are distributed along certain predefined intervals of a variable's value space. The vector of initial states for large systems will hence be high dimensional. To ensure that the ODE dynamics is well explored, one needs to draw a large number of representative trajectories. Naive direct sampling where we randomly pick values from the initial intervals vector cannot ensure that all parts of the initial states region are sufficiently

probed. Hence, we used a more sophisticated sampling method called *equation-based subinterval sampling* which is a variant of the method proposed in [1]. Suppose the ODE equation for the variable x_i involves variables x_j and x_k . We then subdivide the initial intervals of the variables x_i , x_j , and x_k into J finer subintervals. Then for *every combination of subintervals* say, (I_i, I_j, I_k) , we pick H samples each of which will have its x_i -value falling in I_i , x_j -value falling in I_j and its x_k -value falling in I_k while the values for the other variables are picked randomly from within their initial intervals. This ensures a coverage of at least H samples for every combination of the subintervals of the variables governing each equation which in turn ensures that ODE dynamics is being explored systematically along each dimension at least. In general, if an equation has R variables on its right hand side, and there are n equations and H is the required degree of coverage per equation, we pick $n \cdot H \cdot J^{R+1}$ samples.

To assess the quality of the constructed DBNs in terms of the original ODE dynamics, we used Monte Carlo integration to generate random trajectories from the prior (initial states distribution) using the ODE. We then computed the average values of each variable at the time points $0 \leq t \leq T$. We term the resulting time series for each variable as a *nominal profile*. We then used marginal probability values derived from the DBN approximation to compute expected values as follows $E(M^t(i, u)) = \sum_{u=u_j} (M^t(i, u_j) \cdot L)$, where L is the mid-point of the interval u_j . For each variable, the resulting time series of expected values was compared with its nominal profile. For all the models studied below the quality of the DBN approximation measured this way was high. Due to space limitations, the comparison plots will be shown in what follows here only for a few chosen species in the case of the NGF stimulated EGF-NGF pathway and the Epo mediated ERK pathway. The detailed comparison plots for other DBNs can be found in [28].

Finally, for the DBNs arising from EGF-NGF pathway and Epo mediated ERK pathway, exact inference is infeasible due to the large sizes of the corresponding DBNs. To get around this, we used simulation-based inferencing of the DBN to obtain an estimate of the exact marginal distribution. We generated around 200 million trajectories from the underlying DBN to obtain the various marginal probabilities. This took around 2 days for each DBN. These marginals were used—in place of exact marginals—as benchmarks to compare the performance of the various algorithms. Here again we compared HFF(σ) for various choices of σ with FF and BK. We discuss toward the end of this section the performance of the clustered version of BK. In what follows, we write HFF(cK) to mean the HFF(σ) with $\sigma = c \cdot 1,000$.

4.2.1 The EGF-NGF Pathway

The EGF-NGF pathway describes the behavior of PC12 cells under multiple stimulations. In response to EGF stimulation they proliferate but differentiate into sympathetic neurons in response to NGF stimulation. This phenomenon has been intensively studied [34] and the network structure of this pathway is as shown in Fig. 3. The ODE model of this pathway [32] consists of 32 differential equations and

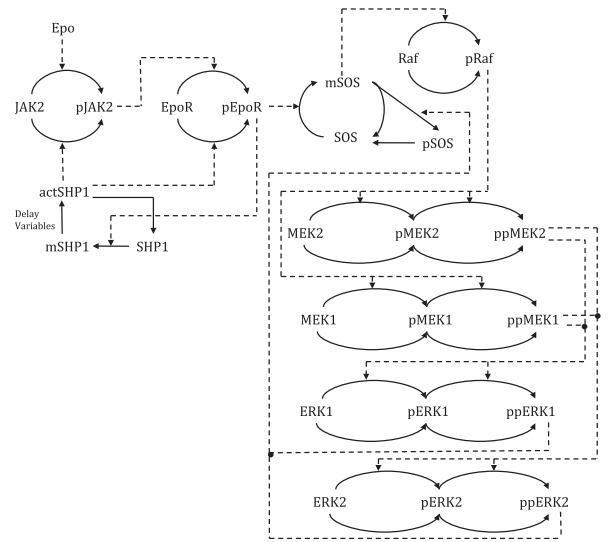


Fig. 3. Epo mediated ERK Signaling pathway

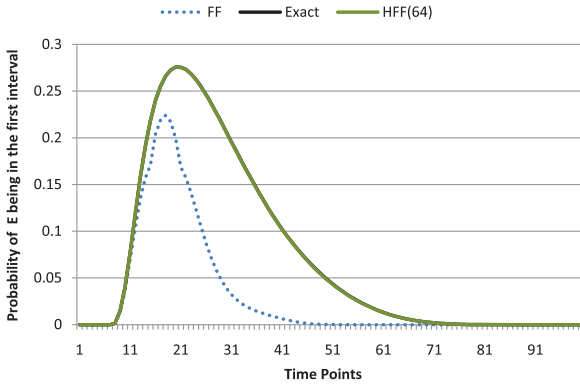
48 associated rate constants (estimated from multiple sets of experimental data as reported in [32]).

To construct the three DBNs arising out of EGF, NGF, and costimulation, we divided as before the value domains of the variables into five equally wide intervals and assumed the initial distributions to be uniformly distributed over some of these intervals. The time horizon of each model was set at 10 minutes which was evenly divided into 100 time points. To fill up the conditional probability tables, we used the equation-based subinterval sampling. We subdivided each of the initial states into four subintervals. 2.1 million trajectories were generated to get a coverage of 500 per combination of the subinterval. As shown in Fig. 5a, the quality of the approximations relative to the original ODE dynamics was high. Once we had the DBNs, we ran FF, BK, and HFF(σ) for various choices of σ .

For the DBN obtained for the pathway under NGF-stimulation, for 6 of the 32 species there were significant differences between FF and BK on one hand and HFF on the other, including some biologically important proteins such as *Sos* and *Erk*. In Fig. 6, we show for *Erk*, the marginal probability of the concentration falling in the interval $[1, 2)$ at various time points as computed by FF, BK, HFF(3K), and HFF(32K) as well as the pseudoexact marginals obtained via massive Monte Carlo simulations. We observe that HFF tends to the exact values as the number of spikes increases.

To measure the overall error behavior, noting that HFF always did better than FF, we fixed the error incurred by FF as the base (100 percent) and normalized all other errors relative to this base. Under this regime, the relationship between computation time and normalized mean error for *Erk*'s value to fall in $[1, 2)$ is shown in Fig. 7. We observe that the mean error reduces to 22 percent for HFF(32K) at the cost of approximately 10^4 seconds increase in running time. For HFF(σ) the errors did not decrease linearly as the number of spikes were increased. This is to be expected since the probability mass captured by the additional spikes will be less than what is captured by the initial spikes.

Overall, the maximum error over *all* the marginals ($32 \times 5 \times 100 = 16,000$ data points) reduced from 0.42 for FF

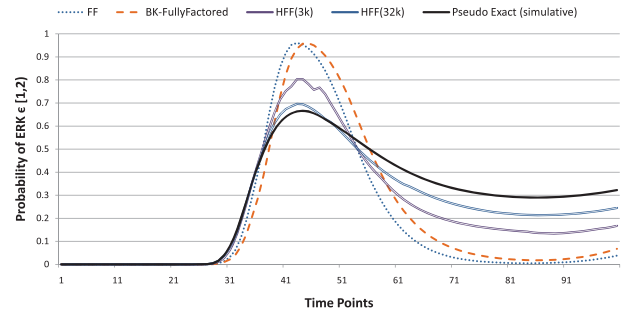
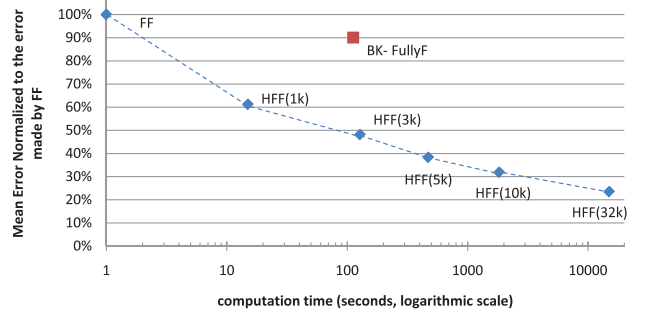
Fig. 4. $M^t(E \in [0, 1])$

to 0.3 for HFF(3K) and to 0.12 for HFF(32K). The normalized mean error over all marginals went down to 60 percent for HFF(3K) and 30 percent for HFF(32K) as shown in Fig. 8a which also displays the corresponding computation times. Further, when we computed the *number of marginals* with errors greater than 0.1, we found that this number reduced to about half for HFF(3K) and by more than a factor of 10 for HFF(32K) compared to FF as shown in Fig. 8b.

For the DBN obtained for the pathway under EGF-stimulation, we found similar results. Overall, the maximum error over *all* the marginals reduced from 0.35 for FF to 0.14 for HFF(3K) and to 0.07 for HFF(32K). The normalized mean error over all marginals went down to 40 percent for HFF(3K) spikes and 20 percent for HFF(32K) spikes as shown in Fig. 9a which also displays the corresponding computation times. Further, when we computed the number of marginals with errors greater than 0.1, we found that this number reduced by more than a factor of 4 for HFF(3K) and to 0 for HFF(32K) as shown in Fig. 9b. Similar results were obtained for the DBN describing the dynamics of the EGF-NGF pathway under co-stimulation of both NGF and EGF.

4.2.2 The Epo Mediated ERK Pathway

Next we considered the DBN model of Epo mediated ERK Signaling pathway as shown in Fig. 4. *Erk* and its related kinase isoforms play a crucial role in cell proliferation, differentiation, and survival. This pathway describes the

Fig. 6. $M^t(Erk \in [1, 2))$ under NGF-stimulation.Fig. 7. Normalized mean error for $M^t(Erk \in [1, 2))$ under NGF-stimulation.

effect of these isoforms on the Epo (cytokine) induced ERK cascade. The ODE model of this pathway [9] consists of 32 differential equations and 24 associated rate constants. To construct the DBN, we divided the value domain of variables into five intervals. Here the interval sizes for variables were not all kept equal. For 23 species that have very low basal concentration level, we set the first interval of the corresponding variables to be smaller ($\sim 20\%$) compared to the other four intervals (equal sized) (see online supplementary materials [28]). The rest nine variables all have equal sized intervals as before. Time horizon was fixed at 60 minutes which was then divided into 100 time points. We constructed the DBN using equation-based subinterval sampling. As Fig. 5b indicates, the quality of the approximation relative to the original ODE dynamics was again high (more comparison plots can be found in [28]). We then ran FF, BK, and HFF(σ) for various choices of σ .

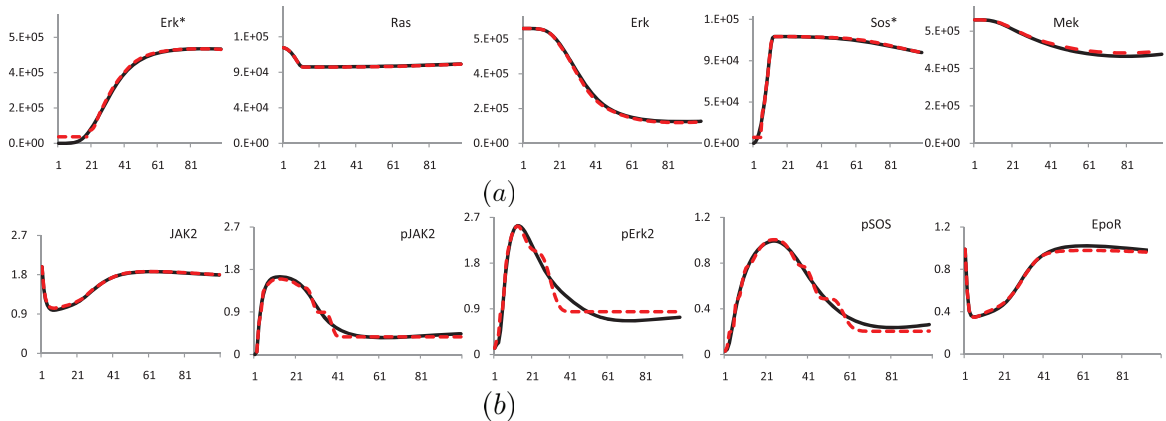


Fig. 5. Comparison of ODE dynamics with DBN approximation. Solid black line represents nominal ODE profiles and dashed red lines represent the DBN simulation profiles for (a) NGF stimulated EGF-NGF Pathway. (b) Epo mediated ERK pathway.

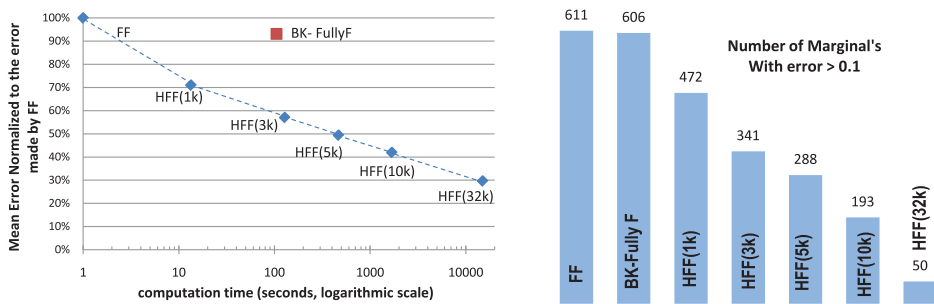


Fig. 8. (a) Normalized mean errors over all marginals. (b) Number of marginals with error greater than 0.1: NGF-stimulation.

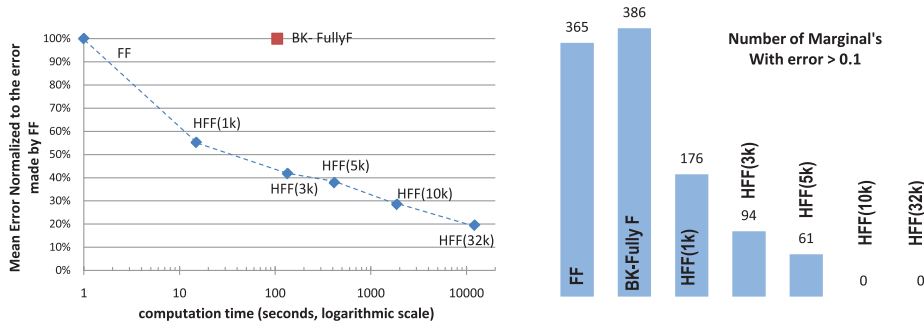


Fig. 9. (a) Normalized mean error over all marginals. (b) Number of marginals with error greater than 0.1: EGF-stimulation.

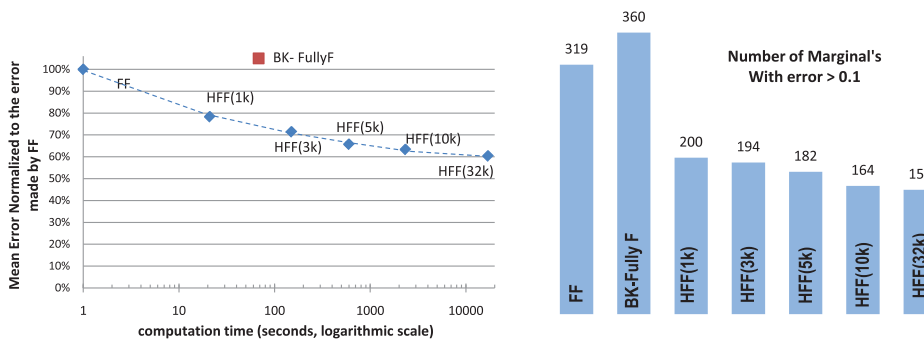


Fig. 10. (a) Normalized mean errors over all marginals. (b) Number of marginals with error greater than 0.1: Epo stimulated ERK pathway.

FF and BK were quite accurate for many of the species. However, for some species such as *JAK2*, phosphorylated *EpoR*, *SHP1*, *mSHP1*, etc. which are biologically relevant, FF and BK incurred a max error of 0.49. On the other hand, HFF(3K) incurred a max error of 0.45 while HFF(32K) incurred a max error of 0.31. The normalized mean error over all marginals went down to $\sim 70\%$ for HFF(3K) and $\sim 60\%$ for HFF(32K) as shown in Fig. 10a. Further, when we computed the number of marginals with errors greater than 0.1, we found that this number reduced by around half for HFF(32K) compared to FF as shown in Fig. 10b.

It is worth noting that our current implementation is quite crude and sequential. We believe significant performance gains can be expected from an optimized version.

4.3 Comparison with Clustered BK

An important component of the BK algorithm is the grouping of the variables into clusters. The idea is to choose the clusters in such a way that there is not much interaction between variables belonging to two different clusters. When this is done well, BK can also perform well. However, choosing the right clusters seems to be a difficult task. The easy option, namely, the fully factored BK in which each

cluster is a singleton performs in our case studies as badly (or well) as FF. The idea is to choose the clusters in such a way that there is not much interaction between variables belonging to two different clusters. When this is done well, BK can also perform well. However, choosing the right clusters seems to be a difficult task. The easy option, namely, the fully factored BK in which each cluster is a singleton performs in our case studies as badly (or well) as FF.

We tried to gain a better understanding of BK augmented with nontrivial clusters by using the structure of the pathway to come up with good clusters. A natural way to form 2-clusters seemed to be to pair together the activated (phosphorylated) and inactivated (dephosphorylated) counterparts of a species in the pathway. For the EGF-NGF pathway, this clustering indeed reduced overall errors compared to FF and HFF(3K). However, we found that HFF(σ) with $\sigma > 5,000$ outperformed this version of BK. We did not consider bigger clusters for two reasons: first, when we tried to increase the sizes and the number of clusters in different ways, BK ran out of the 3 GB memory. Second, there seemed to be no biological criterion using which one could improve the error performance of BK.

For the Epo mediated ERK pathway too we tried similar clustering. Here the natural clusters were of size 3. Unfortunately, the results were as bad as for fully factored BK. HFF, even with $1K$ spikes ($\sigma = 1,000$) was able to perform better than this clustered version of BK. This suggests that the clusters we chose were not the right ones. Hence in our setting, a clustered version of BK that performs well in terms of the computational resources required and the errors incurred appears to be difficult to realize.

5 DISCUSSION

DBNs are a widely used class of probabilistic graphical models and can often be a succinct and more natural representation of the underlying Markov chains. Computing the marginal probabilities of the random variables assuming specific values in a DBN is a basic analysis task. This can be done only approximately for high dimensional systems. FF and BK are two attractive approximate algorithms that have been proposed in this context. However they can incur significant errors. To cope with this, we have developed here the Hybrid Factored Frontier (HFF) algorithm. In HFF, in addition to maintaining and propagating beliefstates in a factored form, we also maintain a small number of full dimensional state vectors called spikes and their probabilities at each time slice. By tuning the number of spikes, one can gain accuracy at the cost of increased but polynomial (quadratic) computational cost. We have provided an error analysis for HFF as well as FF which shows that HFF is more accurate. Our experimental results confirm that HFF outperforms both FF and fully factored BK in realistic biological settings.

As pointed out earlier, HFF is parametrized algorithm. In particular, as the number of spikes approaches V^n , the error will tend to 0. This does not mean however that the rate of reduction in error is uniform in the number of spikes or that it is independent of the nature of the spikes. In fact our error analysis suggests that the errors will be high when there are *high* spikes of which there cannot be too many. Hence, if we reduce the errors for such spikes—which HFF will do—we will substantially improve accuracy.

One may consider BK also to be a parametrized algorithm with the number of clusters and their sizes constituting the parameters. However identifying the clusters is a difficult problem and our experimental results suggest that as the sizes of the clusters increase the errors may reduce but the memory consumption could rise rapidly. In contrast, HFF's parameter can be computed in an efficient, approximate, and *automated* fashion. In our case studies, we have found that by using HFF, the accuracy of results can be improved with an increase in computational times. Further, for tasks such as parameter estimation that require repeated executions, one can first deploy FF to get an initial estimate and then run HFF with a suitable number of spikes just once to achieve a sufficiently accurate estimate.

In our future work, an important goal will be to deploy HFF to perform tasks such as parameter estimation and sensitivity analysis. An equally important goal will be to develop approximate probabilistic verification methods for DBN models of biochemical networks and evaluate them with respect to approaches developed in related settings [23], [21], [24], [25].

REFERENCES

- [1] B. Liu, D. Hsu, and P.S. Thiagarajan, "Probabilistic Approximations of ODEs Based Bio-Pathway Dynamics," *Theoretical Computer Science*, vol. 412, pp. 2188-2206, 2011.
- [2] K.P. Murphy and Y. Weiss, "The Factored Frontier Algorithm for Approximate Inference in DBNs," *Proc. 17th Int'l Conf. Uncertainty in Artificial Intelligence (UAI '01)*, pp. 378-385, 2001.
- [3] X. Boyen and D. Koller, "Tractable Inference for Complex Stochastic Processes," *Proc. 14th Int'l Conf. Uncertainty in Artificial Intelligence (UAI '98)*, pp. 33-42, 1998.
- [4] B. Liu, J. Zhang, P.Y. Tan, D. Hsu, A.M. Blom, B. Leong, S. Sethi, B. Ho, J.L. Ding, and P.S. Thiagarajan, "A Computational and Experimental Study of the Regulatory Mechanisms of the Complement System," *PLoS Computational Biology*, vol. 7, no. 1, p. e1001059, 2011.
- [5] G. Koh, D. Hsu, and P.S. Thiagarajan, "Incremental Signaling Pathway Modeling by Data Integration," *Proc. Int'l Conf. Research in Computational Molecular Biology (RECOMB '10)*, pp. 281-296, 2010.
- [6] C. Langmead, S. Jha, and E. Clarke, "Temporal Logics as Query Languages for Dynamic Bayesian Networks: Application to D. Melanogaster Embryo Development," technical report, Carnegie Mellon Univ., 2006.
- [7] D. Koller and N. Friedman, *Probabilistic Graphical Models: Principles and Techniques*. MIT Press, 2009.
- [8] K.S. Brown, C.C. Hill, G.A. Calero, C.R. Myers, K.H. Lee, and R.A. Cerione, "The Statistical Mechanics of Complex Signaling Networks: Nerve Growth Factor Signaling," *Physical Biology*, vol. 1, pp. 184-195, 2004.
- [9] M. Schilling, T. Maiwald, S. Hengl, D. Winter, C. Kreutz, W. Kolch, W.D. Lehmann, J. Timmer, and U. Klingmüller, "Theoretical and Experimental Analysis Links Isoform-Specific ERK Signalling to Cell Fate Decisions," *Molecular Systems Biology*, vol. 5, p. 334, 2009.
- [10] P.F. Felzenszwalb and D.P. Huttenlocher, "Efficient Belief Propagation for Early Vision," *Int'l J. Computer Vision*, vol. 70, pp. 41-54, 2006.
- [11] R.J. McEliece, D.J.C. Mackay, and J.-F. Cheng, "Turbo Decoding as an Instance of Pearl's "Belief Propagation" Algorithm," *IEEE J. Selected Areas in Comm.*, vol. 16, no. 2, pp. 140-152, Feb. 1998.
- [12] N. Friedman, "Inferring Cellular Networks Using Probabilistic Graphical Models," *Science*, vol. 303, pp. 799-805, 2004.
- [13] B. Bidyuk and R. Dechter, "An Anytime Scheme for Bounding Posterior Beliefs," *Proc. 21st Nat'l Conf. Artificial Intelligence (AAAI '06)*, pp. 1-6, 2006.
- [14] J. Biles and H. Lin, "Online Adaptive Learning for Speech Recognition Decoding," *Proc. Ann. Conf. Int'l Speech Comm. Assoc. (Interspeech '10)*, pp. 1958-1961, 2010.
- [15] M.Z. Kwiatkowska, G. Norman, and D. Parker, "PRISM: Probabilistic Symbolic Model Checker," *Proc. 12th Int'l Conf. Modelling Techniques and Tools for Computer Performance Evaluation (TOOLS '02)*, pp. 200-204, 2002.
- [16] M. Calder, V. Vyshemirsky, D. Gilbert, and R.J. Orton, "Analysis of Signalling Pathways Using Continuous Time Markov Chains," *Trans. Computational Systems Biology*, vol. 4220, pp. 44-67, 2006.
- [17] W.S. Hlavacek, J.R. Faeder, M.L. Blinov, R.G. Posner, M. Hucka, and W. Fontana, "Rules for Modeling Signal-Transduction Systems," *Science STKE*, vol. 2006, p. re6, 2006.
- [18] V. Danos, J. Feret, W. Fontana, R. Harmer, and J. Krivine, "Rule-Based Modelling of Cellular Signalling," *Proc. 18th Int'l Conf. Concurrency Theory (CONCUR '07)*, pp. 17-41, 2007.
- [19] B. Liu, P.S. Thiagarajan, and D. Hsu, "Probabilistic Approximations of Signaling Pathway Dynamics," *Proc. Seventh Int'l Conf. Computational Methods in Systems Biology (CMSB '09)*, pp. 251-265, 2009.
- [20] T.A. Henzinger, M. Mateescu, and V. Wolf, "Sliding Window Abstraction for Infinite Markov Chains," *Proc. 21th Int'l Conf. Computer Aided Verification (CAV '09)*, pp. 337-352, 2009.
- [21] S.K. Jha, E.M. Clarke, C.J. Langmead, A. Legay, A. Platzer, and P. Zuliani, "A Bayesian Approach to Model Checking Biological Systems," *Proc. Seventh Int'l Conf. Computational Methods in Systems Biology (CMSB '09)*, pp. 218-234, 2009.
- [22] M.Z. Kwiatkowska, G. Norman, and D. Parker, "Probabilistic Model Checking for Systems Biology," *Symbolic Systems Biology*, Jones and Bartlett, 2010.

- [23] R. Grosu and S.A. Smolka, "Monte Carlo Model Checking," *Proc. 11th Int'l Conf. Tools and Algorithms for Construction and Analysis of Systems (TACAS '05)*, pp. 271-286, 2005.
- [24] F. Fages and A. Rizk, "On the Analysis of Numerical Data Time Series in Temporal Logic," *Proc. Int'l Conf. Computational Methods in Systems Biology (CMSB '07)*, pp. 48-63, 2007.
- [25] R. Donaldson and D. Gilbert, "A Model Checking Approach to the Parameter Estimation of Biochemical Pathways," *Proc. Sixth Int'l Conf. Computational Methods in Systems Biology (CMSB '08)*, pp. 269-287, 2008.
- [26] F. Ciocchetta, A. Degasperis, J. Hillston, and M. Calder, "Some Investigations Concerning the CTMC and the ODE Model Derived from Bio-PEPA," *Electronic Notes in Theoretical Computer Science*, vol. 229, no. 1, pp. 145-163, 2009.
- [27] F. Didier, T.A. Henzinger, M. Mateescu, and V. Wolf, "Approximation of Event Probabilities in Noisy Cellular Processes," *Proc. Seventh Int'l Conf. Computational Methods in Systems Biology (CMSB '09)*, pp. 173-188, 2009.
- [28] "Supplementary Materials," <http://www.comp.nus.edu.sg/suchee/HFF/>, 2011.
- [29] S.K. Palaniappan, S. Akshay, B. Genest, and P.S. Thiagarajan, "A Hybrid Factored Frontier Algorithm for Dynamic Bayesian Network Models of Biopathways," *Proc. Ninth Int'l Conf. Computational Methods in Systems Biology (CMSB '11)*, pp. 35-44, 2011.
- [30] B.B. Aldridge, J.M. Burke, D.A. Lauffenburger, and P.K. Sorger, "Physicochemical Modelling of Cell Signalling Pathways," *Nature Cell Biology*, vol. 8, pp. 1195-1203, 2006.
- [31] P. Bremaud, *Markov Chains: Gibbs Fields, Monte Carlo Simulation and Queues*. Springer, 2010.
- [32] N. Le Novère, B. Bornstein, A. Broicher, M. Courtot, M. Donizelli, H. Dharuri, L. Li, H. Sauro, M. Schilstra, B. Shapiro, J. Snoep, and M. Hucka, "BioModels Database: A Free, Centralized Database of Curated, Published, Quantitative Kinetic Models of Biochemical and Cellular Systems," *Nucleic Acids Research*, vol. 34, pp. D689-D691, 2006.
- [33] K.P. Murphy, "Bayes Net Toolbox for Matlab," <http://bnt.googlecode.com>, 2012.
- [34] B.N. Kholodenko, "Untangling the Signalling Wires," *Nature Cell Biology*, vol. 9, pp. 247-249, 2007.



Suchendra Kumar Palaniappan received the BEng degree in biotechnology from Vishveshwaraya Technological University, India. Currently, he is working toward the PhD degree at the School of Computing, National University of Singapore. His main research interests include computational systems biology and probabilistic model checking.



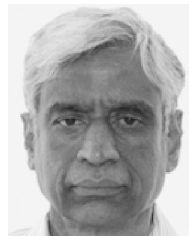
interests include verification of timed and communicating systems and computational systems biology.



Bing Liu received the BComp degree from the School of Computing, National University of Singapore (NUS), Singapore, in 2006 and the PhD degree in computational biology from NUS Graduate School for Integrative Sciences and Engineering, NUS, in 2011. Currently, he is a research fellow in the Computer Science Department, NUS. His research interests include computational systems biology and high-performance computing for computational biology.



Blaise Genest is a former student of ENS Cachan in mathematics and computer science in 1999-2003, and received the PhD degree in computer science from University Paris 7, France in 2004. He is a full time permanent researcher at French CNRS in computer science, and currently in the Joint French-Singapore Laboratory IPAL. His current research interests include verification of distributed systems and computational systems biology.



P.S. Thiagarajan received the BTech degree in electrical engineering from the Indian Institute of Technology (IIT), Madras, in 1970 and the PhD degree in computer science from Rice University, Houston, Texas, in 1973. He is working as a professor in the Computer Science Department of National University of Singapore. He has served as a member of the editorial boards of the journals *Theoretical Computer Science* and the *International Journal on Foundations of Computing*. Currently, he serves on the editorial boards of *The Real-Time Systems journal* and *Transactions on Petri nets and Other Models of Concurrency*. He is a fellow of the Indian Academy of Sciences and the Indian National Academy of Sciences. His current research interests include system-level design and analysis methods for real-time embedded systems and computational systems biology.

► For more information on this or any other computing topic, please visit our Digital Library at www.computer.org/publications/dlib.