

RESEARCH

Open Access



Differentially mutated subnetworks discovery

Morteza Chalabi Hajkarim¹, Eli Upfal² and Fabio Vandin^{3*} 

Abstract

Problem: We study the problem of identifying differentially mutated subnetworks of a large gene–gene interaction network, that is, subnetworks that display a significant difference in mutation frequency in two sets of cancer samples. We formally define the associated computational problem and show that the problem is NP-hard.

Algorithm: We propose a novel and efficient algorithm, called DAMOKLE, to identify differentially mutated subnetworks given genome-wide mutation data for two sets of cancer samples. We prove that DAMOKLE identifies subnetworks with statistically significant difference in mutation frequency when the data comes from a reasonable generative model, provided enough samples are available.

Experimental results: We test DAMOKLE on simulated and real data, showing that DAMOKLE does indeed find subnetworks with significant differences in mutation frequency and that it provides novel insights into the molecular mechanisms of the disease not revealed by standard methods.

Keywords: Network analysis, Somatic mutations, Differential analysis

Introduction

The analysis of molecular measurements from large collections of cancer samples has revolutionized our understanding of the processes leading to a tumour through somatic mutations, changes of the DNA appearing during the lifetime of an individual [1]. One of the most important aspects of cancer revealed by recent large cancer studies is *inter-tumour genetic heterogeneity*: each tumour presents hundreds-thousands mutations and no two tumours harbour the same set of DNA mutations [2].

One of the fundamental problems in the analysis of somatic mutations is the identification of the handful of *driver mutations* (i.e., mutations related to the disease) of each tumour, detecting them among the thousands or tens of thousands that are present in each tumour genome [3]. Inter-tumour heterogeneity renders the identification of driver mutations, or of driver genes (genes containing driver mutations), extremely difficult, since only few genes are mutated in a relatively large

fraction of samples while most genes are mutated in a low fraction of samples in a cancer cohort [4].

Recently, several analyses (e.g. [5, 6]) have shown that interaction networks provide useful information to discover driver genes by identifying groups of interacting genes, called *pathways*, in which each gene is mutated at relatively low frequency while the entire group has one or more mutations in a significantly large fraction of all samples. Several network-based methods have been developed to identify groups of interacting genes mutated in a significant fraction of tumours of a given type and have been shown to improve the detection of driver genes compared to methods that analyze genes in isolation [5, 7–9].

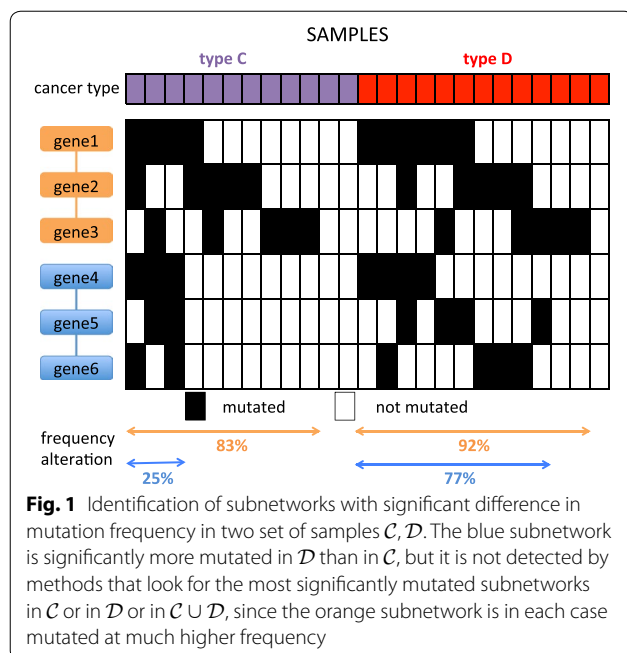
The availability of molecular measurements in a large number of samples for different cancer types have also allowed *comparative* analyses of mutations in cancer [5, 10, 11]. Such analyses usually analyze large cohorts of different cancer types as a whole employing methods to find genes or subnetworks mutated in a significant fraction of tumours in *one* cohort, and also analyze each cancer type individually, with the goal to identify:

*Correspondence: fabio.vandin@unipd.it

³ Department of Information Engineering, University of Padova, Padova, Italy

Full list of author information is available at the end of the article





1. pathways that are common to various cancer types;
2. pathways that are specific to a given cancer type.

For example, [5] analyzed 12 cancer types and identified subnetworks (e.g., a TP53 subnetwork) mutated in most cancer types as well as subnetworks (e.g., a MHC subnetwork) enriched for mutations in one cancer type. In addition, comparative analyses may also be used for the identification of mutations of clinical relevance [12]. For example: comparing mutations in a patients that responded to a given therapy with mutations in patients (of the same cancer type) that did not respond to the same therapy may identify genes and subnetworks associated with response to therapy; comparing mutations in patients whose tumours metastasized with mutations in patients whose tumours did not metastasize may identify mutations associated with the insurgence of metastases.

Pathways that are significantly mutated only in a specific cancer type may not be identified by analyzing one cancer type at the time or all samples together (Fig. 1), but, interestingly, to the best of our knowledge no method has been designed to *directly* identify sets of interacting genes that are significantly more mutated in a set of samples compared to another. The task of finding such sets is more complex than the identification of subnetworks significantly mutated in a set of samples, since subnetworks that have a significant difference in mutations in two sets may display relatively modest frequency of mutation in both set of samples, whose difference can

be assessed as significant only by the joint analysis of both sets of samples.

Related work

Several methods have been designed to analyze different aspects of somatic mutations in a large cohort of cancer samples in the context of networks. Some methods analyze mutations in the context of known pathways to identify the ones significantly enriched in mutations (e.g., [13]). Other methods combine mutations and large interaction networks to identify cancer subnetworks [5, 14, 15]. Networks and somatic mutations have also been used to prioritize mutated genes in cancer [7, 8, 16–18] and for patients stratification [6, 19]. Some of these methods have been used for the identification of common mutation patterns or subnetworks in several cancer types [5, 10], but to the best of our knowledge no method has been designed to identify mutated subnetworks with a significant difference in two cohorts of cancer samples.

Few methods studied the problem of identifying subnetworks with significant differences in two sets of cancer samples using data other than mutations. [20] studied the problem of identifying optimally discriminative subnetworks of a large interaction network using gene expression data. Mall et al. [21] developed a procedure to identify statistically significant changes in the topology of biological networks. Such methods cannot be readily applied to find subnetworks with significant difference in mutation frequency in two sets of samples. Other related work use gene expression to characterize different cancer types: [22] defined a pathway-based score that clusters samples by cancer type, while [23] defined pathway-based features used for classification in various settings, and several methods [24–28] have been designed for finding subnetworks with differential gene expression.

Our contribution

In this work we study the problem of finding subnetworks with frequency of mutation that is significantly different in two sets of samples. In particular, our contributions are fourfold. First, we propose a combinatorial formulation for the problem of finding subnetworks significantly more mutated in one set of samples than in another and prove that such problem is NP-hard. Second, we propose Differentially Mutated subnetwOrKs analysis in canCEr (DAMOKLE), a simple and efficient algorithm for the identification of subnetworks with a significant difference of mutation in two sets of samples, and analyze DAMOKLE proving that it identifies subnetworks significantly more mutated in one of two sets of samples under reasonable assumptions for the data. Third, we test DAMOKLE on simulated data, verifying

experimental that DAMOKLE correctly identifies subnetworks significantly more mutated in a set of samples when enough samples are provided in input. Fourth, we test DAMOKLE on large cancer datasets comprising two cancer types, and show that DAMOKLE identifies subnetworks significantly associated with one of the two types which cannot be identified by state-of-the-art methods designed for the analysis of one set of samples.

Methods and algorithms

This section presents the problem we study, the algorithm we propose for its solution, and the analysis of our algorithm. In particular, "Computational problem" section formalizes the computational problem we consider; "Algorithm" section presents Differentially Mutated subnetworks analysis in cancer (DAMOKLE), our algorithm for the solution of the computational problem; "Analysis of DAMOKLE" section describes the analysis of our algorithm under a reasonable generative model for mutations; "Statistical significance of the results" section presents a formal analysis of the statistical significance of subnetworks obtained by DAMOKLE; and "Permutation testing" section describes two permutation tests to assess the significance of the results of DAMOKLE for limited sample sizes.

Computational problem

We are given measurements on mutations in m genes $\mathcal{G} = \{1, \dots, m\}$ on two sets $\mathcal{C} = \{c_1, \dots, c_{n_C}\}$, $\mathcal{D} = \{d_1, \dots, d_{n_D}\}$ of samples. Such measurements are represented by two matrices C and D , of dimension $m \times n_C$ and $m \times n_D$, respectively, where n_C (resp., n_D) is the number of samples in $\mathcal{C}$ (resp., $\mathcal{D}$). $C(i, j) = 1$ (resp., $D(i, j) = 1$) if gene i is mutated in the j -th sample of $\mathcal{C}$ (resp., $\mathcal{D}$) and $C(i, j) = 0$ (resp., $D(i, j) = 0$) otherwise. We are also given an (undirected) graph $G = (V, E)$, where vertices $V = \{1, \dots, m\}$ are genes and $(i, j) \in E$ if gene i interacts with gene j (e.g., the corresponding proteins interact).

Given a set of genes $S \subset \mathcal{G}$, we define the indicator function $c_S(c_i)$ with $c_S(c_i) = 1$ if at least one of the genes of S is mutated in sample c_i , and $c_S(c_i) = 0$ otherwise. We define $c_S(d_i)$ analogously. We define the *coverage* $c_S(\mathcal{C})$ of S in $\mathcal{C}$ as the fraction of samples in $\mathcal{C}$ for which at least one of the genes in S is mutated in the sample, that is

$$c_S(\mathcal{C}) = \frac{\sum_{i=1}^{n_C} c_S(c_i)}{n_C}$$

and, analogously, define the *coverage* $c_S(\mathcal{D})$ of S in $\mathcal{D}$ as

$$c_S(\mathcal{D}) = \frac{\sum_{i=1}^{n_D} c_S(d_i)}{n_D}.$$

We are interested in identifying sets of genes S , with $|S| \leq k$, corresponding to connected subgraphs in G and displaying a *significant* difference in coverage between $\mathcal{C}$ and $\mathcal{D}$, i.e., with a high value of $|c_S(\mathcal{C}) - c_S(\mathcal{D})|$.

We define the *differential coverage* $dc_S(\mathcal{C}, \mathcal{D})$ as $dc_S(\mathcal{C}, \mathcal{D}) = c_S(\mathcal{C}) - c_S(\mathcal{D})$.

In particular, we study the following computational problem.

The differentially mutated subnetworks discovery problem: given a value θ with $\theta \in [0, 1]$, find all connected subgraphs S of G of size $\leq k$ such that $dc_S(\mathcal{C}, \mathcal{D}) \geq \theta$.

Note that by finding sets that maximize $dc_S(\mathcal{C}, \mathcal{D})$ we identify sets with significantly more mutations in $\mathcal{C}$ than in $\mathcal{D}$, while to identify sets with significantly more mutations in $\mathcal{D}$ than in $\mathcal{C}$ we need to find sets maximizing $dc_S(\mathcal{D}, \mathcal{C})$. In addition, note that a subgraph S in the solution may contain genes that are not mutated in $\mathcal{C} \cup \mathcal{D}$ but that are needed for the connectivity of S .

We have the following.

Theorem 1 *The differentially mutated subnetworks discovery problem is NP-hard.*

Proof The proof is by reduction from the connected maximum coverage problem [14]. In the connected maximum coverage problem we are given a graph G defined on a set $V = \{v_1, \dots, v_n\}$ of n vertices, a family $\mathcal{P} = \{P_1, \dots, P_n\}$ of subsets of a universe I (i.e., $P_i \in 2^I$), with P_i being the subset of I covered by $v_i \in V$ and value k , and we want to find the subgraph $C^* = \{v_{i_1}, \dots, v_{i_k}\}$ with k nodes of G that maximizes $|\cup_{j=1}^k P_{i_j}|$.

Given an instance of the connected maximum coverage problem, we define an instance of the differentially mutated subnetworks discovery problem as follows: the set $\mathcal{G}$ of genes corresponds to the set V of vertices of G in the connected maximum coverage problem, and the graph G is the same as in the instance of the maximum coverage instance; the set $\mathcal{C}$ is given by the set I and the matrix C is defined as $C_{i,j} = 1$ if $i \in P_j$, while $\mathcal{D} = \emptyset$.

Note that for any subgraph S of G , the differential coverage $dc_D(\mathcal{C}, \mathcal{D}) = c_S(\mathcal{C}) - c_S(\mathcal{D}) = c_S(\mathcal{C})$ and $c_S(\mathcal{C}) = |\cup_{g \in S} P_g|/|I|$. Since $|I|$ is the same for all solutions, the optimal solution of the differentially mutated subnetworks discovery instance corresponds to the optimal solution to the connected maximum coverage instance, and viceversa. $\square$

Algorithm

We now describe Differentially Mutated subnetworks analysis in cancer (DAMOKLE), an algorithm to solve the differentially mutated subnetworks discovery problem. DAMOKLE takes in input mutation matrices C and D for two sets $\mathcal{C}, \mathcal{D}$ of samples, a (gene-gene) interaction graph G , an integer $k > 0$, and a real value $\theta \in [0, 1]$, and returns subnetworks S of G with $\leq k$ vertices and differential coverage $dc_S(\mathcal{C}, \mathcal{D}) \geq \theta$. Subnetworks reported by DAMOKLE are also *maximal* (no vertex can be added to

S while maintaining the connectivity of the subnetwork, $|S| \leq k$ and $dc_S(\mathcal{C}, \mathcal{D}) \geq \theta$). DAMOKLE is described in Algorithm 1. DAMOKLE starts by considering each edge $e = \{u, v\} \in E$ of G with differential coverage $dc_{\{u,v\}}(\mathcal{C}, \mathcal{D}) \geq \theta/(k-1)$, and for each such e identifies subnetworks including e to be reported in output using Algorithm 2.

Model

For each gene $i \in \mathcal{G} = \{1, 2, \dots, m\}$ there is an a-priori probability p_i of observing a mutation in gene i . Let $H \subset \mathcal{G}$ be the connected subnetwork of up to k genes that is differentially mutated in samples of $\mathcal{C}$ w.r.t. samples of $\mathcal{D}$. Mutations in our samples are taken from two related distributions. In the “control” distribution F a mutation

Algorithm 1: DAMOKLE

Input: mutation matrices $\mathcal{C}, \mathcal{D}$; gene-gene interaction graph $G = (V, E)$; integer $k > 0$;
 $\theta \in [0, 1]$
Output: maximal connected subgraphs with $dc_S(\mathcal{C}, \mathcal{D}) \geq \theta$

```

1 solutions  $\leftarrow \emptyset$ ;
2 foreach  $\{u, v\} \in E$  do
3   if  $dc_{\{u,v\}}(\mathcal{C}, \mathcal{D}) \geq \theta/(k-1)$  then
4     solutions  $\leftarrow$  solutions  $\cup$  GETSOLUTIONS( $E, \{u, v\}$ );
5   end
6 end
7 return solutions;
```

GETSOLUTIONS, described in Algorithm 2, is a recursive algorithm that, given a current subgraph S , identifies all maximal connected subgraphs S' , $|S'| \leq k$, containing S and with $dc_{S'}(\mathcal{C}, \mathcal{D}) \geq \theta$. This is obtained by expanding S one edge at the time and stopping when the number of vertices in the current solution is k or when the addition of no vertex leads to an increase in differential coverage $dc_S(\mathcal{C}, \mathcal{D})$ for the current solution S . In Algorithm 2, $N(S)$ refers to the set of edges with exactly one vertex in the set S .

in gene i is observed with probability p_i independent of other genes' mutations. The second distribution F_H is analogous to the distribution F but we condition on the event $E(H)$ = “at least one gene in H is mutated in the sample”.

For genes not in H , all mutations come from distribution F . For genes in H , in a perfect experiment with no noise we would assume that samples in $\mathcal{C}$ are taken from F_H and samples from $\mathcal{D}$ are taken from F . However, to

Algorithm 2: GETSOLUTIONS

Input: set E of edges of the graph; current subgraph (solution) S
Output: maximal connected subgraphs containing S with $dc_S(\mathcal{C}, \mathcal{D}) \geq \theta$

```

1 nextEdges  $\leftarrow \emptyset$ ;
2 foreach  $e \in N(S)$  do
3   if  $dc_{S \cup \{e\}}(\mathcal{C}, \mathcal{D}) \geq dc_S(\mathcal{C}, \mathcal{D})$  then nextEdges  $\leftarrow$  nextEdges  $\cup \{e\}$ ;
4 end
5 if |nextEdges| = 0 OR  $|S| = k$  then
6   if  $dc_S(\mathcal{C}, \mathcal{D}) \geq \theta$  then return  $S$ ;
7 end
8 newSols  $\leftarrow \emptyset$ ;
9 foreach  $e \in$  nextEdges do newSols  $\leftarrow$  newSols  $\cup$  GETSOLUTIONS( $E, S \cup \{e\}$ );
10 return newSols;
```

The motivation for design choices of DAMOKLE are provided by the results in the next section.

Analysis of DAMOKLE

The design and analysis of DAMOKLE are based on the following generative model for the underlying biological process.

model realistic, noisy data we assume that with some probability q the “true” signal for a sample is lost, that is the sample from $\mathcal{C}$ is taken from F . In particular, samples in $\mathcal{C}$ are taken with probability $1 - q$ from F_H and with probability q from F .

Let p be the probability that H has at least one mutation in samples from the control model F , $p = 1 - \prod_{j \in H} (1 - p_j) \approx \sum_{j \in H} p_j$. Clearly, we are only interested in sets $H \subset \mathcal{G}$ with $p \ll 1$.

If we focus on individual genes, the probability gene i is mutated in a sample from $\mathcal{D}$ is p_i , while the probability that it is mutated in a sample from $\mathcal{C}$ is $\frac{(1-q)p_i}{1 - \prod_{j \in H} (1 - p_j)} + qp_i$.

Such a gap may be hard to detect with a small number of samples. On the other hand, the probability of $E(H)$ (i.e., of at least one mutation in the set H) in a sample from $\mathcal{C}$ is $(1 - q) + q(1 - \prod_{j \in H} (1 - p_j)) = 1 - q + qp$, while the probability of $E(H)$ in a sample from $\mathcal{D}$ is $1 - \prod_{j \in H} (1 - p_j) = p$ which is a more significant gap, when $p \ll 1$.

The efficiency of DAMOKLE is based on two fundamental results. First we show that it is sufficient to start the search only in edges with relatively high differential coverage.

Proposition 1 *If $dc_S(\mathcal{C}, \mathcal{D}) \geq \theta$, then, in the above generating model, with high probability (asymptotic in n_C and n_D) there exist an edge $e \in S$ such that $dc_{\{e\}}(\mathcal{C}, \mathcal{D}) \geq (\theta - \epsilon)/(k - 1)$, for any $\epsilon > 0$.*

Proof For a set of genes $S' \subset \mathcal{G}$ and a sample $z \in \mathcal{C} \cup \mathcal{D}$, let $Count(S', z)$ be the number of genes in S' mutated in sample z . Clearly, if for all $z \in \mathcal{C} \cup \mathcal{D}$, we have $Count(S, z) = 1$, i.e. each sample has no more than one mutation in S , then

$$\begin{aligned} dc_S(\mathcal{C}, \mathcal{D}) &= c_S(\mathcal{C}) - c_S(\mathcal{D}) = \frac{\sum_{i=1}^{n_C} c_S(c_i)}{n_C} - \frac{\sum_{i=1}^{n_D} c_S(d_i)}{n_D} \\ &= \frac{\sum_{i=1}^{n_C} \sum_{j \in S} Count(\{j\}, c_i)}{n_C} - \frac{\sum_{i=1}^{n_D} \sum_{j \in S} Count(\{j\}, d_i)}{n_D} \\ &= \sum_{j \in S} \left(\frac{\sum_{i=1}^{n_C} Count(\{j\}, c_i)}{n_C} - \frac{\sum_{i=1}^{n_D} Count(\{j\}, d_i)}{n_D} \right) \\ &\geq \theta. \end{aligned}$$

Thus, there is a vertex $j^* = \arg \max_{j \in S} \left(\frac{\sum_{i=1}^{n_C} Count(\{j\}, c_i)}{n_C} - \frac{\sum_{i=1}^{n_D} Count(\{j\}, d_i)}{n_D} \right)$ such that $dc_{\{j^*\}}(\mathcal{C}, \mathcal{D}) = \frac{\sum_{i=1}^{n_C} Count(\{j^*\}, c_i)}{n_C} - \frac{\sum_{i=1}^{n_D} Count(\{j^*\}, d_i)}{n_D} \geq \theta/k$.

Since the set of genes S is connected, there is an edge $e = (j^*, \ell)$ for some $\ell \in S$. For that edge,

$$dc_{\{e\}}(\mathcal{C}, \mathcal{D}) \geq \frac{\theta - dc_{\{\ell\}}(\mathcal{C}, \mathcal{D})}{k - 1} + dc_{\{\ell\}}(\mathcal{C}, \mathcal{D}) \geq \frac{\theta}{k - 1}.$$

For the case when the assumption $Count(S, z) = 1$ for all $z \in \mathcal{C} \cup \mathcal{D}$ does not hold, let

$$\begin{aligned} Mul(S, \mathcal{C}, \mathcal{D}) &= \frac{\sum_{i=1}^{n_C} \sum_{j \in S} Count(\{j\}, c_i)}{n_C} - \frac{\sum_{i=1}^{n_C} c_S(c_i)}{n_C} \\ &\quad + \frac{\sum_{i=1}^{n_D} \sum_{j \in S} Count(\{j\}, d_i)}{n_D} - \frac{\sum_{i=1}^{n_D} c_S(d_i)}{n_D}. \end{aligned}$$

Then

$$\sum_{j \in S} \left(\frac{\sum_{i=1}^{n_C} Count(\{j\}, c_i)}{n_C} - \frac{\sum_{i=1}^{n_D} Count(\{j\}, d_i)}{n_D} \right) - Mul(S, \mathcal{C}, \mathcal{D}) \geq \theta$$

and

$$dc_{\{e\}}(\mathcal{C}, \mathcal{D}) \geq \frac{\theta + Mul(S, \mathcal{C}, \mathcal{D})}{k - 1}.$$

Since the probability of having more than one mutation in S in a sample from $\mathcal{C}$ is at least as high as from a sample from $\mathcal{D}$, we can normalize (similar to the proof of Theorem 2 below) and apply Hoeffding bound (Theorem 4.14 in [29]) to prove that

$$Prob(Mul(S, \mathcal{C}, \mathcal{D}) < -\epsilon) \leq 2e^{-2\epsilon^2 n_C n_D / (n_C + n_D)}.$$

□

The second result motivates the choice, in Algorithm 2, of adding only edges that increase the score of the current solution (and to stop if there is no such edge).

Proposition 2 *If subgraph S can be partitioned as $S = S' \cup \{j\} \cup S''$, and $dc_{S' \cup \{j\}}(\mathcal{C}, \mathcal{D}) < dc_{S'}(\mathcal{C}, \mathcal{D}) - pp_j$, then with high probability (asymptotic in n_D) $dc_{S \setminus \{j\}}(\mathcal{C}, \mathcal{D}) > dc_S(\mathcal{C}, \mathcal{D})$.*

Proof We first observe that if each sample in $\mathcal{D}$ has no more than 1 mutation in S then $dc_{S' \cup \{j\}}(\mathcal{C}, \mathcal{D}) < dc_{S'}(\mathcal{C}, \mathcal{D})$ implies that $dc_{\{j\}}(\mathcal{C}, \mathcal{D}) < 0$, and therefore, under this assumption, $dc_{S \setminus \{j\}}(\mathcal{C}, \mathcal{D}) > dc_S(\mathcal{C}, \mathcal{D})$.

To remove the assumption that a sample has no more than one mutation in S , we need to correct for the fraction of samples in $\mathcal{D}$ with mutations both in j and S'' . With high probability (asymptotic in n_D) this fraction is bounded by $pp_j + \epsilon$ for any $\epsilon > 0$. □

Statistical significance of the results

To compute a threshold that guarantees statistical confidence of our finding, we first compute a bound on the gap in a non significant set.

Theorem 2 Assume that S is not a significant set, i.e., $\mathcal{C}$ and $\mathcal{D}$ have the same distribution on S , then

$$\text{Prob}(dc_S(\mathcal{C}, \mathcal{D}) > \epsilon) \leq 2e^{-2\epsilon^2 n_C n_D / (n_C + n_D)}.$$

Proof Let $X_1, \dots, X_{n_C}$ be independent random variables such that $X_i = 1/n_C$ if sample c_i in $\mathcal{C}$ has a mutation in S , otherwise $X_i = 0$. Similarly, let $Y_1, \dots, Y_{n_D}$ be independent random variables such that $Y_i = -1/n_D$ if sample d_i in $\mathcal{D}$ has a mutation in S , otherwise $Y_i = 0$.

Clearly $dc_S(\mathcal{C}, \mathcal{D}) = \sum_{i=1}^{n_C} X_i + \sum_{i=1}^{n_D} Y_i$, and since S is not significant $E[\sum_{i=1}^{n_C} X_i + \sum_{i=1}^{n_D} Y_i] = 0$.

To apply Hoeffding bound (Theorem 4.14 in [29]), we note that the sum $\sum_{i=1}^{n_C} X_i + \sum_{i=1}^{n_D} Y_i$ has n_C variables in the range $[0, 1/n_C]$, and n_D variables in the range $[-1/n_D, 0]$. Thus,

$$\begin{aligned} \text{Prob}(dc_S(\mathcal{C}, \mathcal{D}) > \epsilon) &\leq 2e^{(-2\epsilon^2)/(n_C/n_C^2 + n_D/n_D^2)} \\ &= 2e^{-2\epsilon^2 n_C n_D / (n_C + n_D)}. \end{aligned}$$

□

Let N_k be the set of subnetworks under consideration, or the set of all connected components of size $\leq k$. We use Theorem 2 to obtain guarantees on the statistical significance of the results of DAMOKLE in terms of the Family-Wise Error Rate (FWER) or of the False Discovery Rate (FDR) as follows:

- FWER: if we want to find just the subnetwork with significant maximum differential coverage, to bound the FWER of our method by α we use the maximum ϵ such that $N_k 2e^{-2\epsilon^2 n_C n_D / (n_C + n_D)} \leq \alpha$.
- FDR: if we want to find several significant subnetworks with high differential coverage, to bound the FDR by α we use the maximum ϵ such that $N_k 2e^{-2\epsilon^2 n_C n_D / (n_C + n_D)} / n(\alpha) \leq \alpha$, where $n(\alpha)$ is the number of sets with differential coverage $\geq \epsilon$.

Permutation testing

While Theorem 2 shows how to obtain guarantees on the statistical significance of the results of DAMOKLE by appropriately setting θ , in practice, due to relatively small sample sizes and to inevitable looseness in the theoretical guarantees, a permutation testing approach may be more effective in estimating the statistical significance of the results of DAMOKLE and provide more power for the identification of differentially mutated subnetworks.

We consider two permutation tests to assess the association of mutations in the subnetwork with the highest differential coverage found by DAMOKLE. The first

test assesses whether the observed differential coverage can be obtained under the independence of mutations in genes by considering the null distribution in which each gene is mutated in a random subset (of the same cardinality as observed in the data) of all samples, independently of all other events. The second test assesses whether, under the observed marginal distributions for mutations in sets of genes, the observed differential coverage of a subnetwork can be obtained under the independence between mutations and samples' memberships (i.e., being a sample of $\mathcal{C}$ or a sample of $\mathcal{D}$), by randomly permuting the samples memberships.

Let $dc_S(\mathcal{C}, \mathcal{D})$ be the differential coverage observed on real data for the solution S with highest differential coverage found by DAMOKLE (for some input parameters). For both tests we estimate the p -value as follow:

1. generate N (permuted) datasets from the null distribution;
2. run DAMOKLE (with the same input parameters used on real data) on each of the N permuted datasets;
3. let x be the number of permuted datasets in which DAMOKLE reports a solution with differential coverage $\geq dc_S(\mathcal{C}, \mathcal{D})$: then the p -value of S is $(x + 1)/(N + 1)$.

Results

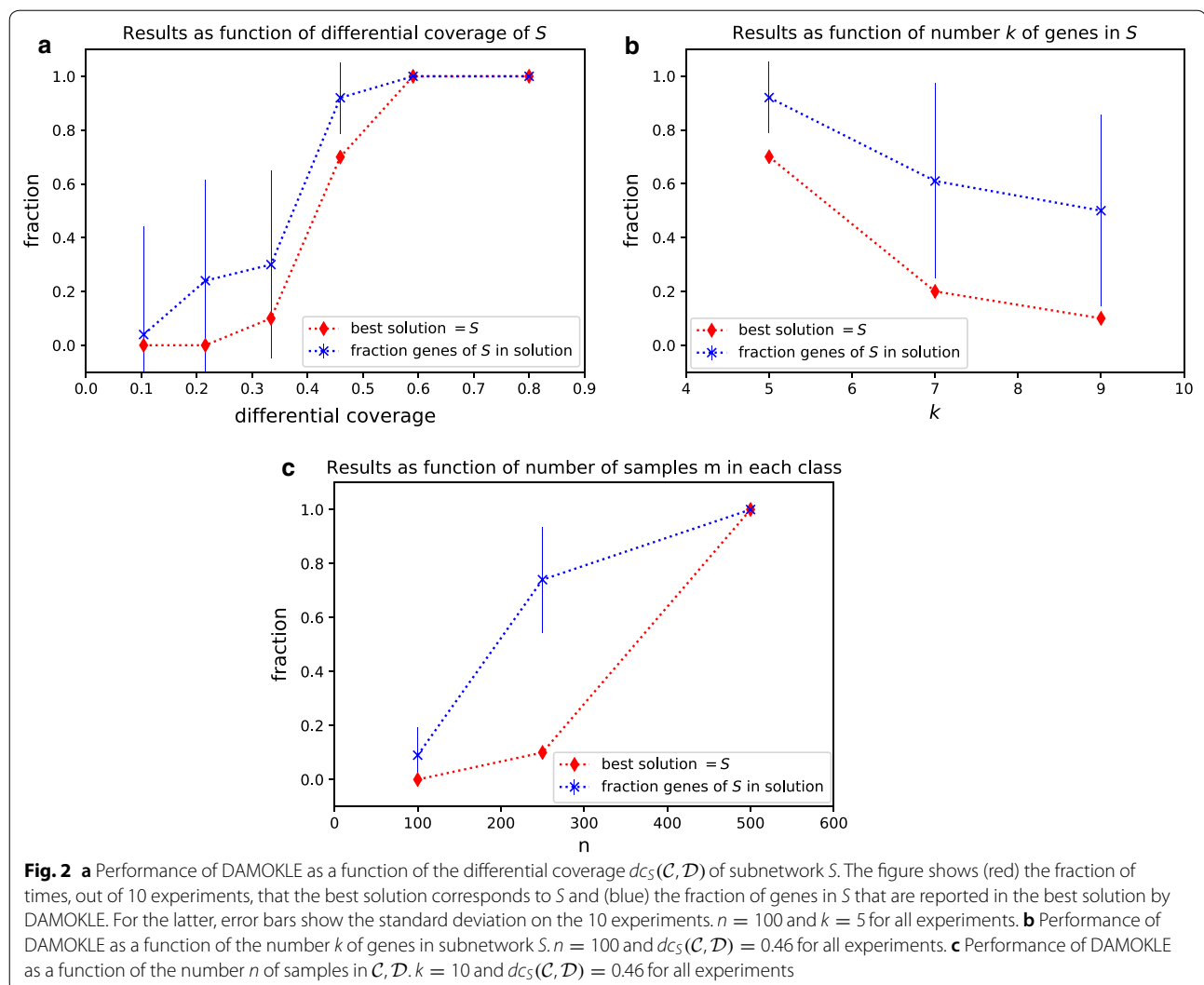
We implemented DAMOKLE in Python¹ and tested it on simulated and on cancer data. Our experiments have been conducted on a Linux machine with 16 cores and 256 GB of RAM. For all our experiments we used as interaction graph G the HINT+HI2012 network², a combination of the HINT network [30] and the HI-2012 [31] set of protein-protein interactions, previously used in [5]. In all cases we considered only the subnetwork with the highest differential coverage among the ones returned by DAMOKLE. We first present the results on simulated data ("Simulated data" section) and then present the results on cancer data ("Cancer data" section).

Simulated data

We tested DAMOKLE on simulated data generated as follows. We assume there is a subnetwork S of k genes with differential coverage $dc_S(\mathcal{C}, \mathcal{D}) = c$. In our simulations we set $|\mathcal{C}| = |\mathcal{D}| = n$. For each sample in $\mathcal{D}$, each gene g in G (including genes in S) is mutated with

¹ The implementation is available at <https://github.com/VandinLab/DAMOKLE>.

² <http://compbio-research.cs.brown.edu/pancancer/hotnet2/>.



probability p_g , independently of all other events. For samples in $\mathcal{C}$, we first mutated each gene g with probability p_g independently of all other events. We then considered the samples of $\mathcal{C}$ without mutations in S , and for each such sample we mutated, with probability c , one gene of S , chosen uniformly at random. In this way c is the *expectation* of the differential coverage $dc_S(\mathcal{C}, \mathcal{D})$. For genes in $G \setminus S$ we used mutation probabilities p_g estimated from oesophageal cancer data [32]. We considered only value of $n \geq 100$, consistent with sample sizes in most recent cancer sequencing studies. (The latest ICGC data release³ from April 30th, 2018 has data for ≥ 500 samples for 81% of the primary sites).

The goal of our investigation using simulated data is to evaluate the impact of various parameters on ability

of DAMOKLE to recover S or part of it. In particular, we studied the impact of three parameters: the differential coverage $dc_S(\mathcal{C}, \mathcal{D})$ of the planted subnetwork S ; the number k of genes in S ; and the number n of samples in each class. To evaluate the impact of such parameters, for each combination of parameters in our experiments we generated 10 simulated datasets and run DAMOKLE on each dataset with $\theta = 0.01$, recording

1. the fraction of times that DAMOKLE reported S as the solution with the highest differential coverage, and
2. the fraction of genes of S that are in the solution with highest differential coverage found by DAMOKLE.

We first investigated the impact of the differential coverage $c = dc_S(\mathcal{C}, \mathcal{D})$. We analyzed simulated

³ <https://dcc.icgc.org/>.

datasets with $n = 100$ samples in each class, where $k = 5$ genes are part of the subnetwork S , for values of $c = 0.1, 0.22, 0.33, 0.46, 0.6, 0.8$. We run DAMOKLE on each dataset with $k = 5$. The results are shown in Fig. 2a. For low values of the differential coverage c , with $n = 100$ samples DAMOKLE never reports S as the best solution found and only a small fraction of the genes in S are part of the solution reported by DAMOKLE. However, as soon as the differential coverage is ≥ 0.45 , even with $n = 100$ samples in each class DAMOKLE identifies the entire planted solution S most of the times, and even when the best solution does not entirely corresponds to S , more than 80% of the genes of S are reported in the best solution. For values of $c \geq 0.6$, DAMOKLE *always* reports the whole subnetwork S as the best solution. Given that many recent large cancer sequencing studies consider at least 200 samples, DAMOKLE will be useful to identify differentially mutated subnetworks in such studies.

We then tested the performance of DAMOKLE as a function of the number of genes k in S . We tested the ability of DAMOKLE to identify a subnetwork S with differential coverage $dc_S(\mathcal{C}, \mathcal{D}) = 0.46$ in a dataset with $n = 100$ samples in both $\mathcal{C}$ and $\mathcal{D}$, when the number k of genes in S varies as $k = 5, 7, 9$. The results are shown in Fig. 2b. As expected, when the number of genes in S increases, the fraction of times S is the best solution as well as the fraction of genes reported in the best solution by S decreases, and for $k = 9$ the best solution found by DAMOKLE corresponds to S only 10% of the times. However, even for $k = 9$, on average most of the genes of S are reported in the best solution by DAMOKLE. Therefore DAMOKLE can be used to identify relatively large subnetworks mutated in a significantly different number of samples even when the number of samples is relatively low.

Finally, we tested the performance of DAMOKLE as the number of samples n in each set $\mathcal{C}, \mathcal{D}$ increases. In particular, we tested the ability of DAMOKLE to identify a relatively large subnetwork S of $k = 10$ genes with differential coverage $dc_S(\mathcal{C}, \mathcal{D}) = 0.46$ as the number of samples n increases. We analyzed simulated datasets for $n = 100, 250, 500$. The results are shown in Fig. 2. For $n = 100$, when $k = 10$, DAMOKLE never reports S as the best solution and only a small fraction of all genes in S are reported in the solution. However, for $n = 250$, while DAMOKLE still reports S as the best solution only 10% of the times, on average 70% of the genes of S are reported in the best solution. More interestingly, already for $n = 500$, DAMOKLE *always* reports S as the best solution. These results show that DAMOKLE can reliably identify relatively large differentially mutated subnetworks from currently available datasets of large cancer sequencing studies.

Cancer data

We use DAMOKLE to analyze somatic mutations from The Cancer Genome Atlas. We first compared two similar cancer types and two very different cancer types to test whether DAMOKLE behaves as expected on these types. We then analyzed two pairs of cancer types where differences in alterations are unclear. In all cases we run DAMOKLE with $\theta = 0.1$ and obtained p -values with the permutation tests described in "Permutation testing" section.

Lung cancer

We used DAMOKLE to analyze 188 samples of lung squamous cell carcinoma (LUSC) and 183 samples of lung adenocarcinoma (LUAD). We only considered single nucleotide variants (SNVs)⁴ and use $k = 5$. DAMOKLE did not report any significant subnetwork, in agreement with previous work showing that these two cancer types have known differences in gene expression [33] but are much more similar with respect to SNVs [34].

Colorectal vs ovarian cancer

We used DAMOKLE to analyze 456 samples of colorectal adenocarcinoma (COADREAD) and 496 samples of ovarian serous cystadenocarcinoma (OV) using only SNVs.⁵ For $k = 5$, DAMOKLE identifies the significant ($p < 0.01$ according to both tests in "Permutation testing" section) subnetwork APC, CTNNB1, FBXO30, SMAD4, SYNE1 with differential coverage 0.81 in COADREAD w.r.t. OV. APC, CTNNB1, and SMAD4 are members of the WNT signaling and TFG- β signaling pathways. The WNT signaling pathway is one of the cascades that regulates stemness and development, with a role in carcinogenesis that has been described mostly for colorectal cancer [35], but altered Wnt signaling is observed in many other cancer types [36]. The TFG- β signaling pathway is involved in several processes including cell growth and apoptosis, that is deregulated in many diseases, including COADREAD [35]. The high differential coverage of the subnetwork is in accordance with COADREAD being altered mostly by SNVs and OV being altered mostly by copy number aberrations (CNAs) [37].

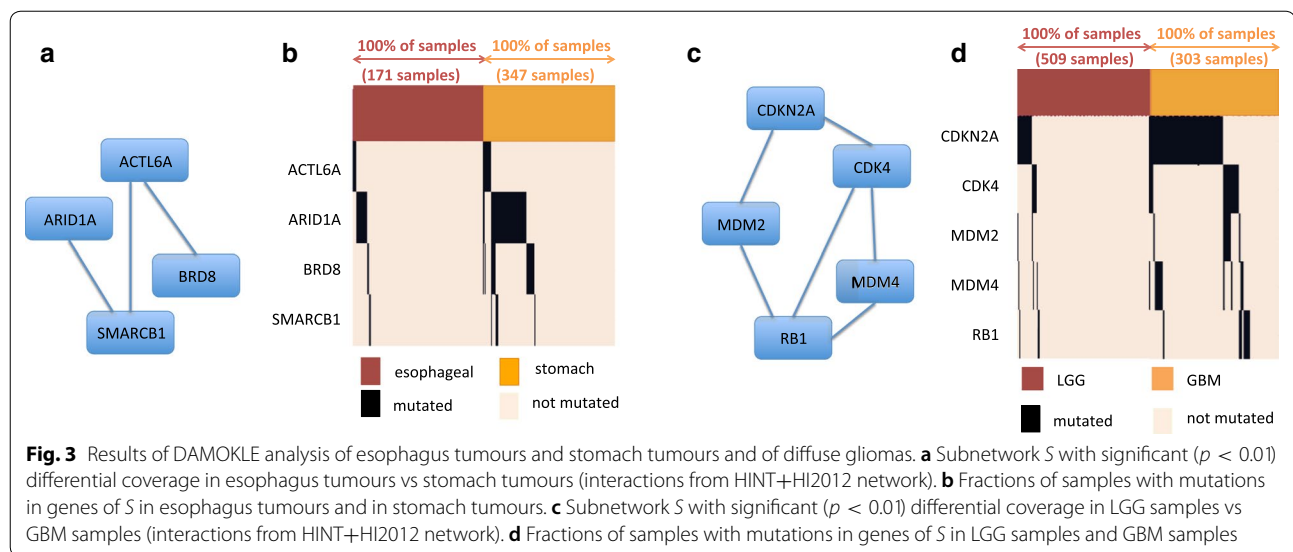
Esophagus-stomach cancer

We analyzed SNVs and CNAs in 171 samples of esophagus cancer and in 347 samples of stomach cancer [32].⁶ The number of mutations in the two sets is not significantly different (t-test $p = 0.16$). We first considered

⁴ http://cbio.mskcc.org/cancergenomics/pancan_tcga/.

⁵ http://cbio.mskcc.org/cancergenomics/pancan_tcga/.

⁶ http://www.cbioportal.org/study?id=stes_tcga_pub#summary.



single genes, identifying TP53 with high (> 0.5) differential coverage between the two cancer types. Alterations in TP53 have then be removed for the subsequent DAMOKLE analysis. We run DAMOKLE with $k = 4$ with $\mathcal{C}$ being the set of stomach tumours and $\mathcal{D}$ being the set of esophagus tumours. DAMOKLE identifies the significant ($p < 0.01$ for both tests in "Permutation testing" section) subnetwork $S = \{\text{ACTL6A}, \text{ARID1A}, \text{BRD8}, \text{SMARCB1}\}$ with differential coverage 0.26 (Fig. 3a, b). Interestingly, all four genes in the subnetwork identified by DAMOKLE are members of the chromatin organization machinery recently associated with cancer [38, 39]. Such subnetwork is not reported as differentially mutated in the TCGA publication comparing the two cancer types [32]. BRD8 is only the top-16 gene by differential coverage, while ACTL6 and SMARCB1 are not among the top-2000 genes by differential coverage. We compared the results obtained by DAMOKLE with the results obtained by HotNet2 [5], a method to identify significantly mutated subnetworks, using the same mutation data and the same interaction network as input: none of the genes in S appeared in significant subnetworks reported by HotNet2.

Diffuse gliomas

We analyzed single nucleotide variants (SNVs) and copy number aberrations (CNAs) in 509 samples of lower grade glioma (LGG) and in 303 samples of glioblastoma multiforme (GBM).⁷ We considered nonsilent SNVs,

short indels, and CNAs. We removed from the analysis genes with < 6 mutations in both classes. By single gene analysis we identified IDH1 with high (> 0.5) differential coverage, and removed alterations in such gene for the DAMOKLE analysis. We run DAMOKLE with $k = 5$ with $\mathcal{C}$ being the set of GBM samples and $\mathcal{D}$ being the set of LGG samples. The number of mutations in $\mathcal{C}$ and in $\mathcal{D}$ is not significantly different (t-test $p = 0.1$). DAMOKLE identifies the significant ($p < 0.01$ for both tests in "Permutation testing" section) subnetwork $S = \{\text{CDKN2A}, \text{CDK4}, \text{MDM2}, \text{MDM4}, \text{RB1}\}$ (Fig. 3c, d). All genes in S are members of the p53 pathway or of the RB pathway. The p53 pathway has a key role in cell death as well as in cell division, and the RB pathway plays a crucial role in cell cycle control. Both pathways are well known glioma cancer pathways [40]. Interestingly, [41] did not report any subnetwork with significant difference in mutations between LGG and GBM samples. CDK4, MDM2, MDM4, and RB1 do not appear among the top-45 genes by differential coverage. We compared the results obtained by DAMOKLE with the results obtained by HotNet2. Of the genes in our subnetwork, only CDK4 and CDKN2A are reported in a significantly mutated subnetwork ($p < 0.05$) obtained by HotNet2 analyzing $\mathcal{D}$ but not analyzing $\mathcal{C}$, while MDM2, MDM4, and RB1 are not reported in any significant subnetwork obtained by HotNet2.

Conclusion

In this work we study the problem of finding subnetworks of a large interaction network with significant difference in mutation frequency in two sets of cancer samples. This problem is extremely important to identify mutated

⁷ https://media.githubusercontent.com/media/cBioPortal/datahub/master/public/lgggbm_tcg_a_pub.tar.gz.

mechanisms that are specific to a cancer (sub)type as well as for the identification of mechanisms related to clinical features (e.g., response to therapy). We provide a formal definition of the problem and show that the associated computational problem is NP-hard. We design, analyze, implement, and test a simple and efficient algorithm, DAMOKLE, which we prove identifies significant subnetworks when enough data from a reasonable generative model for cancer mutations is provided. Our results also show that the subnetworks identified by DAMOKLE cannot be identified by methods not designed for the *comparative* analysis of mutations in two sets of samples. We tested DAMOKLE on simulated and real data. The results on simulated data show that DAMOKLE identifies significant subnetworks with currently available sample sizes. The results on two large cancer datasets, each comprising genome-wide measurements of DNA mutations in two cancer subtypes, shows that DAMOKLE identifies subnetworks that are not found by methods not designed for the *comparative* analysis of mutations in two sets of samples.

While we provide a first method for the differential analysis of cohorts of cancer samples, several research directions remain. First, differences in the frequency of mutation of a subnetwork in two sets of cancer cohorts may be due to external (or hidden) variables, as for example the mutation rate of each cohort. While at the moment we ensure before running the analysis that no significant difference in mutation rate is present between the two sets, performing the analysis while correcting for possible differences in such confounding variable or in others would greatly expand the applicability of our method. Second, for some interaction networks (e.g., functional ones) that are relatively more dense than the protein–protein interaction network we consider, requiring a minimum connectivity (e.g., in the form of fraction of all possible edges) in the subnetwork may be beneficial, and the design of efficient algorithms considering such requirement is an interesting direction of research. Third, different types of mutation patterns (e.g., mutual exclusivity) among two set of samples could be explored (e.g., extending the method proposed in [42]). Fourth, the inclusion of additional types of measurements, as for example gene expression, may improve the power of our method. Fifth, the inclusion of noncoding variants in the analysis may provide additional information to be leveraged to assess the significance of subnetworks.

Authors' contributions

FV designed the study. FV and EU designed the algorithms. FV and MCH implemented the software and performed the computational analysis. All authors interpreted the results. All authors wrote the manuscript. All authors read and approved the final manuscript.

Author details

¹ Biotech Research and Innovation Centre, University of Copenhagen, Copenhagen, Denmark. ² Department of Computer Science, Brown University, Providence, RI, USA. ³ Department of Information Engineering, University of Padova, Padova, Italy.

Acknowledgements

This work is supported, in part, by University of Padova projects SID2017 and STARS: Algorithms for Inferential Data Mining and by NSF grant IIS-1247581. The results presented in this manuscript are in whole or part based upon data generated by the TCGA Research Network: <http://cancergenome.nih.gov/>.

Competing interests

The authors declare that they have no competing interests.

Publisher's Note

Springer Nature remains neutral with regard to jurisdictional claims in published maps and institutional affiliations.

Received: 7 November 2018 Accepted: 19 March 2019

Published online: 30 March 2019

References

- Garraway LA, Lander ES. Lessons from the cancer genome. *Cell*. 2013;153(1):17–37. <https://doi.org/10.1016/j.cell.2013.03.002>.
- Cancer Genome Atlas Research Network. Integrated genomic characterization of pancreatic ductal adenocarcinoma. *Cancer Cell*. 2017;32(2):185.
- Vogelstein B, Papadopoulos N, Velculescu VE, Zhou S, Diaz LA Jr, Kinzler KW. Cancer genome landscapes. *Science*. 2013;339(6127):1546–58. <https://doi.org/10.1126/science.1235122>.
- Vandin F. Computational methods for characterizing cancer mutational heterogeneity. *Front Genet*. 2017;8:83.
- Leiserson MDM, Vandin F, Wu H-T, Dobson JR, Eldridge JV, Thomas JL, Papoutsaki A, Kim Y, Niu B, McLellan M, Lawrence MS, Gonzalez-Perez A, Tamborero D, Cheng Y, Ryslik GA, Lopez-Bigas N, Getz G, Ding L, Raphael BJ. Pan-cancer network analysis identifies combinations of rare somatic mutations across pathways and protein complexes. *Nat Genet*. 2015;47(2):106–14. <https://doi.org/10.1038/ng.3168>.
- Hofree M, Shen JP, Carter H, Gross A, Ideker T. Network-based stratification of tumor mutations. *Nat Methods*. 2013;10(11):1108–15. <https://doi.org/10.1038/nmeth.2651>.
- Shrestha R, Hodzic E, Sauerwald T, Dao P, Wang K, Yeung J, Anderson S, Vandin F, Haffari G, Collins CC, et al. HITnDRIVE: patient-specific multidriver gene prioritization for precision oncology. *Genome Res*. 2017;27(9):1573–88.
- Hristov BH, Singh M. Network-based coverage of mutational profiles reveals cancer genes. *arXiv preprint arXiv:1704.08544*. 2017.
- Cowen L, Ideker T, Raphael BJ, Sharan R. Network propagation: a universal amplifier of genetic associations. *Nat Rev Genet*. 2017;18(9):551.
- Hoadley KA, Yau C, Wolf DM, Cherniack AD, Tamborero D, Ng S, Leiserson MD, Niu B, McLellan MD, Uzunangelov V, et al. Multiplatform analysis of 12 cancer types reveals molecular classification within and across tissues of origin. *Cell*. 2014;158(4):929–44.
- Kandoth C, McLellan MD, Vandin F, Ye K, Niu B, Lu C, Xie M, Zhang Q, McMichael JF, Wyczalkowski MA, Leiserson MDM, Miller CA, Welch JS, Walter MJ, Wendl MC, Ley TJ, Wilson RK, Raphael BJ, Ding L. Mutational landscape and significance across 12 major cancer types. *Nature*. 2013;502(7471):333–9. <https://doi.org/10.1038/nature12634>.
- Zehir A, Benayed R, Shah RH, Syed A, Middha S, Kim HR, Srinivasan P, Gao J, Chakravarty D, Devlin SM, et al. Mutational landscape of metastatic cancer revealed from prospective clinical sequencing of 10,000 patients. *Nat Med*. 2017;23(6):703.
- Vaske CJ, Benz SC, Sanborn JZ, Earl D, Szeto C, Zhu J, Haussler D, Stuart JM. Inference of patient-specific pathway activities from multi-dimensional cancer genomics data using paradigm. *Bioinformatics*. 2010;26(12):237–45.
- Vandin F, Upfal E, Raphael BJ. Algorithms for detecting significantly mutated pathways in cancer. *J Comput Biol*. 2011;18(3):507–22.

15. Ciriello G, Cerami E, Sander C, Schultz N. Mutual exclusivity analysis identifies oncogenic network modules. *Genome Res.* 2012;22(2):398–406.
16. Kim Y-A, Cho D-Y, Dao P, Przytycka TM. Memcover: integrated analysis of mutual exclusivity and functional network reveals dysregulated pathways across multiple cancer types. *Bioinformatics.* 2015;31(12):284–92.
17. Pulido-Tamayo S, Weytjens B, De Maeyer D, Marchal K. SSA-ME detection of cancer driver genes using mutual exclusivity by small subnetwork analysis. *Sci Rep.* 2016;6:36257.
18. Cho A, Shim JE, Kim E, Supek F, Lehner B, Lee I. MUFFINN: cancer gene discovery via network analysis of somatic mutation data. *Genome Biol.* 2016;17(1):129.
19. Le Morvan M, Zinovyev A, Vert J-P. Netnorm: capturing cancer-relevant information in somatic exome mutation data with gene networks for cancer stratification and prognosis. *PLoS Computat Biol.* 2017;13(6):1005573.
20. Dao P, Wang K, Collins C, Ester M, Lapuk A, Sahinalp SC. Optimally discriminative subnetwork markers predict response to chemotherapy. *Bioinformatics.* 2011;27(13):205–13.
21. Mall R, Cerulo L, Bensmail H, Iavarone A, Ceccarelli M. Detection of statistically significant network changes in complex biological networks. *BMC Syst Biol.* 2017;11(1):32.
22. Young MR, Craft DL. Pathway-informed classification system (PICS) for cancer analysis using gene expression data. *Cancer Inform.* 2016;15:40088.
23. Kim S, Kon M, DeLisi C. Pathway-based classification of cancer subtypes. *Biol Direct.* 2012;7(1):21.
24. Ideker T, Ozier O, Schwikowski B, Siegel AF. Discovering regulatory and signalling circuits in molecular interaction networks. *Bioinformatics.* 2002;18(suppl-1):233–40.
25. Dittrich MT, Klau GW, Rosenwald A, Dandekar T, Müller T. Identifying functional modules in protein–protein interaction networks: an integrated exact approach. *Bioinformatics.* 2008;24(13):223–31.
26. Gu J, Chen Y, Li S, Li Y. Identification of responsive gene modules by network-based gene clustering and extending: application to inflammation and angiogenesis. *BMC Syst Biol.* 2010;4(1):47.
27. Jiao Y, Widschwendter M, Teschendorff AE. A systems-level integrative framework for genome-wide DNA methylation and gene expression data identifies differential gene expression modules under epigenetic control. *Bioinformatics.* 2014;30(16):2360–6.
28. He H, Lin D, Zhang J, Wang Y-p, Deng H-w. Comparison of statistical methods for subnetwork detection in the integration of gene expression and protein interaction network. *BMC Bioinform.* 2017;18(1):149.
29. Mitzenmacher M, Upfal E. Probability and computing: randomization and probabilistic techniques in algorithms and data analysis. Cambridge: Cambridge University Press; 2017.
30. Das J, Yu H. HINT: high-quality protein interactomes and their applications in understanding human disease. *BMC Syst Biol.* 2012;6:92. <https://doi.org/10.1186/1752-0509-6-92>.
31. Yu H, Tardivo L, Tam S, Weiner E, Gebreab F, Fan C, Svrikapa N, Hirozane-Kishikawa T, Rietman E, Yang X, Sahalie J, Salehi-Ashtiani K, Hao T, Cusick ME, Hill DE, Roth FP, Braun P, Vidal M. Next-generation sequencing to generate interactome datasets. *Nat Methods.* 2011;8(6):478–80. <https://doi.org/10.1038/nmeth.1597>.
32. Cancer Genome Atlas Research Network. Integrated genomic characterization of oesophageal carcinoma. *Nature.* 2017;541(7636):169–75.
33. Sun F, Yang X, Jin Y, Chen L, Wang L, Shi M, Zhan C, Shi Y, Wang Q. Bioinformatics analyses of the differences between lung adenocarcinoma and squamous cell carcinoma using the cancer genome atlas expression data. *Mol Med Rep.* 2017;16(1):609–16.
34. Chen F, Zhang Y, Parra E, Rodriguez J, Behrens C, Akbani R, Lu Y, Kurie J, Gibbons DL, Mills GB, et al. Multiplatform-based molecular subtypes of non-small-cell lung cancer. *Oncogene.* 2017;36(10):1384.
35. Network Cancer Genome Atlas. Comprehensive molecular characterization of human colon and rectal cancer. *Nature.* 2012;487(7407):330.
36. Zhan T, Rindtorff N, Boutros M. Wnt signaling in cancer. *Oncogene.* 2017;36(11):1461.
37. Ciriello G, Miller ML, Aksoy BA, Senbabaoglu Y, Schultz N, Sander C. Emerging landscape of oncogenic signatures across human cancers. *Nat Genet.* 2013;45(10):1127.
38. Saladi SV, Ross K, Karaayvaz M, Tata PR, Mou H, Rajagopal J, Ramaswamy S, Ellisen LW. ACTL6A is co-amplified with p63 in squamous cell carcinoma to drive YAP activation, regenerative proliferation, and poor prognosis. *Cancer Cell.* 2017;31(1):35–49.
39. Lu C, Allis CD. SWI/SNF complex in cancer. *Nat Genet.* 2017;49(2):178–9.
40. Vogelstein B, Kinzler KW. Cancer genes and the pathways they control. *Nat Med.* 2004;10(8):789–99.
41. Ceccarelli M, Barthel FP, Malta TM, Sabetot TS, Salama SR, Murray BA, Morozova O, Newton Y, Radenbaugh A, Pagnotta SM, et al. Molecular profiling reveals biologically discrete subsets and pathways of progression in diffuse glioma. *Cell.* 2016;164(3):550–63.
42. Basso RS, Hochbaum DS, Vandin F. Efficient algorithms to discover alterations with complementary functional association in cancer. *arXiv preprint arXiv:1803.09721.* 2018.

Ready to submit your research? Choose BMC and benefit from:

- fast, convenient online submission
- thorough peer review by experienced researchers in your field
- rapid publication on acceptance
- support for research data, including large and complex data types
- gold Open Access which fosters wider collaboration and increased citations
- maximum visibility for your research: over 100M website views per year

At BMC, research is always in progress.

Learn more biomedcentral.com/submissions

