
1

INTRODUCTION TO PATHWAY ANALYSIS

ANTON YURYEV

Table of Contents

1.1	Introduction	1
1.2	Methods to Construct the Pathway Analysis Knowledge Base	2
1.3	Organizational Challenges for Constructing the Knowledge Base for Pathway Analysis	5
1.4	From Molecular Interaction Database to Pathway Collection	6
1.5	Pathway Analysis Software and the Scientific Community	10
1.6	Pathway Analysis and Systems Biology	12
1.7	Pathway Analysis and Network Analysis	13
1.8	Pathway Analysis of Disease	13
1.9	Pathway Analysis and Dynamic Modeling in Drug Development	16
1.10	Steady-State Analysis of Metabolic Networks	18
1.11	So What Is Pathway Analysis?	19
	References	20

1.1 INTRODUCTION

Pathway analysis is a rapidly developing discipline that combines software tools, database models, and computational algorithms—all of which help molecular biologists to convert molecular interaction data into a set of computational models. The models are developed for better prediction of cell behavior in response to a drug, nutrients, or other outside stimuli. The

development of pathway analysis was triggered by the expansion of high-throughput methods and the completion of human genome sequencing project. Because of these technological advances, the emphasis of molecular biology has shifted from reductionism to system integration. Suddenly, nearly all of the components of a living cell became known and the new goal of “putting them all together” into a working computational model of the living cell is awaiting the scientific community. This model must be built by a consensus effort of all molecular biologists and will be constantly refined for a significant period of time. The first and most important goal of pathway analysis is to provide tools and infrastructure that facilitate building a consensus cell model by the collective effort of the scientific community. These tools must enable adequate data exchange, automatic data integration, communication with central public depositories of pathways, and molecular interaction information supporting consensus knowledge base building. In this review, I discuss current approaches for constructing a molecular interaction database, explain the available pathway analysis methods from the drug discovery point of view, and place pathway analysis into the historical perspective of advances in molecular biology.

1.2 METHODS TO CONSTRUCT THE PATHWAY ANALYSIS KNOWLEDGE BASE

Several layers of consensus information are necessary for pathway analysis: (1) a generally agreed-upon list of molecules; (2) a consensus global molecular interaction network; and (3) a collection of consensus pathways for known biological processes. Even though significant portion of the work to create the molecular “inventory” of the human organism has been accomplished by sequencing the human genome, it is still far from being completed. The alternative splicing and protein modification isoforms are yet to be fully cataloged. This next major challenge for this knowledge level is being addressed now by development of exon and phosphorylation microarray technologies. The consensus interaction network is being created by a combination of efforts in high-throughput experiments, prediction of interactions, and classical molecular biological and genetic techniques aimed at elucidating the function of individual proteins. Because the results of numerous small-scale experiments usually are available only in the form of scientific publications and because no central depository for molecular interactions exists in the scientific community, special text-mining techniques have been developed to extract this information into machine-readable format [1].

Every available technique to record interactions for a global network database has some degree of a false-positive rate. High-throughput methods for detection of protein–protein interactions, such as a two-hybrid screen or an identification of protein complexes by co-immunoprecipitation followed by mass spectrometry, are currently being reassessed in a panic due to an appar-

ently worrisome 50–70% false-positive rate [2]. Even though the sources of errors are well understood for these methods, the only way to reduce them at present is to scale down these experiments, effectively reducing them back to one of the laboratory techniques and preventing their high-throughput application. Consequently, attempts to improve the reliability of these methods continue amidst criticism [3]. As a reconciliation note for all high-throughput methods, I emphasize that proteins were designed by nature to interact with each other so as to provide the structural backbone for a living organism. They literally stick to each other to be alive. Therefore, it is not surprising that every protein can interact with many different partners, and many interactions that appear as false positives are in fact true physical interactions. Yet, many of these interactions are not biologically meaningful. Some of them never occur inside the living cell because two proteins never meet each other in space or time. Even if an interaction does occur *in vivo*, it simply may not perform a biological function—it is not followed by the cascade of molecular events that is called a “cell process.” In my opinion, high-throughput methods probably produce mostly correct interaction data, but additional evaluation is necessary to sort out biologically functional and meaningful interactions. The identification of biologically meaningful interactions is a separate task from measuring them. If this is the case, our frustrated energy should be diverted away from these methods and refocused on remedying our general inability to understand what each interaction means for cell physiology. In Chapter 3, about automatic pathway inference, I show that measuring and recording regulatory interactions is one way to calculate biological meaning for physical interactions.

The prediction of physical interactions is typically accomplished using sequence homology [4,5], and regulatory interactions can be calculated from time-series gene expression data using Bayesian inference [6,7]. By holding the powerful promise to “reverse engineer” biological objects, Bayesian inference has been even proclaimed as the principal method of systems biology. However, it currently suffers from the noisy and dispersed experimental data available for analysis, a lack of understanding how to construct good training sets [8], and the emerging realization of the plasticity of biological regulatory networks [9].

Peer-reviewed scientific literature is still the most important source of reliable molecular interaction data. The sheer number of scientific publications and the fact that they are written in machine non-readable format necessitated the development of methods to extract this information into machine-readable format that can be used by computational algorithms. These attempts focused on the manual recording of interactions into a database by an army of curators and, at the same time, on the development of natural language processing algorithms to read scientific papers automatically. Manual curation turned out to be a slow and expensive process that is not error-free: humans do make mistakes when both writing and reading papers. The rate of manual recording appears unable to keep up with the rate of new articles being published.

Automatic extraction of the interaction from peer-reviewed scientific literature faces its own challenges. The main advantage of automatic text-mining algorithms is the speed that allows the processing of hundreds of thousands of articles in minutes. For example, MedScan technology from Ariadne Genomics can process the entire PubMed database that contains more than 14 million abstracts in 2 days on a regular personal computer. This speed allows comprehensive coverage of the entire body of scientific literature, as well as a fairly good assessment of interaction confidence. Interaction confidence can be estimated using the frequency with which the interaction was recovered from the literature. The high misinterpretation rate that occurs during fact extraction is the biggest problem of text-mining technologies. Natural language processing algorithms that use a full sentence parsing approach have the lowest error rate, but they also have the lowest recovery rate among all methods for automatic extraction of interactions [1]. Because errors are evenly distributed among all sentences, false-positive interactions appear as interactions with a low number of supporting references. This technique can be used as a filter to eliminate erroneous interactions at the end of an extraction. Unfortunately, it effectively increases the error rate among relations with a low reference count. Recently discovered true interactions by definition have a low number of references. Thus, automatic extraction methods make the effective confidence of novel interactions even lower due to contamination by false-positive facts. For example, some linguistic patterns used by the MedScan technology extract interactions with one supporting reference with only 70% accuracy. Hence, it is difficult to use the extracted data immediately for analyzing experimental data and building pathways. Additional efforts to curate automatically extracted data are required prior to an analysis such as this one [10,11].

Improvement of all methods for recording of molecular interactions will have to continue for some time until a clear winner can emerge. The most likely outcome, however, is that all interactions for human proteins will be found by the combined effort of all these methods before a winner is determined. After the best method is apparent, interactions for new organisms will be predicted mostly from the consensus interaction network for the human organism and other model organisms. Despite a seemingly overwhelming challenge to measure all physical interactions between proteins in the human genome, this goal will be achieved within the life span of most readers of this book. Indeed, the total number of unique interactions between N proteins is equal to $N(N - 1)/2$. There are 30,000–35,000 genes in the human genome, making the total number of all possible pairwise interactions around 500 million. This number is the upper estimate for human interactome size, which includes both true- and false-positive interactions regardless of the methods used to measure and record them to the global database. The interactions for alternatively spliced proteins can be found relatively easily by calculation using protein sequence information, known interactions of the longest iso-

forms, and interactions from the homologs to determine the protein interacting domains. Thus, even though alternatively splicing greatly increases the number of possible protein–protein interactions, measuring and recording the interactions between splicing isoforms should not require much time and investment. To measure all 500 million possible interactions, each of 500,000 molecular biologists will have to measure about 100 interactions to achieve this goal in 1 year. Assuming there are about 50,000 research projects worldwide actually measuring protein–protein interactions, all physical interactions will be measured in 10 years. The speed of measuring the new physical interaction should gradually increase, and the actual number of interactions is smaller than 500 million. Therefore, 10 years is a very safe upper estimate for time required for recording of all physical interactions.

1.3 ORGANIZATIONAL CHALLENGES FOR CONSTRUCTING THE KNOWLEDGE BASE FOR PATHWAY ANALYSIS

The major barrier separating humankind from measuring all physical protein–protein interactions in its own species is the lack of organization and communication among individual scientists. I have reason to assume that this book will not change this situation, so research will continue as usual by measuring the same interactions multiple times in different laboratories that are trying to prove or disprove each other's theories. Typically, the authors of these studies publish only interactions that seem to support their respective favorite scientific theory or model in hopes of securing funding in the future. High-throughput methods will also continue to contribute to a significant amount of interaction data while attempting to improve their accuracy. Taking the opportunity given to me by this book, I want to join the call to release all interaction data into the public domain [2]. The relatively small organizational challenges that accompany this call include having a central authority to maintain an interaction depository, the tools for data submission and building a consensus global network database, and a method for calculating the confidence of a specific interaction from supporting evidence submitted by multiple sources.

Several public institutions have taken an early lead in the attempt to become a central authority for pathway and molecular interaction databases. Kyoto University provides the Kyoto Encyclopedia of Genes and Genomes (KEGG) database curated by its own staff. Its main disadvantage is that a system cannot be used as a depository by external users outside KEGG. The Signal Transduction Knowledge Environment (STKE) database is maintained by the American Association for the Advancement of Science (AAAS) and contains a collection of pathways curated by scientists considered to be the top experts in the field. This database contains a small collection of highly reliable and canonical pathways and accurately reflects the current state of the art of the

pathway analysis field: very few pathways are actually known and experimentally verified at present. The slow rate of curation and the absence of any formal method to create pathways, including the absence of universal identifiers for pathway components, are the main disadvantages of the STKE database. Unfortunately, the usual leaders in storage of biological information, the National Center for Biotechnology Information (NCBI) in the United States and the European Bioinformatics Institute (EBI) in the European Union, seem to be overwhelmed by the amount of sequencing data they need to maintain. Currently, they lag behind in creating a central resource for pathway information. NCBI, for example, has limited itself by integrating protein physical interaction information from public databases: Biomolecular Interaction Network Database (BIND), Human Protein Reference Database (HPRD), BioGRID, and EcoCyc. Moreover, the constant introduction of new protein identifiers by these organizations unnecessarily complicates the issue even further. Among other public sources of pathway information, I must mention the Reactome database maintained by Cold Spring Harbor Laboratory in collaboration with EBI and Gene Ontology, and the Database of Interacting Proteins (DIP) at the University of California at Los Angeles. Currently, the largest pathway and molecular interaction databases are only available commercially from privately held companies such as Ingenuity Systems, GeneGo, and Ariadne Genomics.

1.4 FROM MOLECULAR INTERACTION DATABASE TO PATHWAY COLLECTION

The collection of physical interactions is not, however, the pathway database. It merely provides the underlying network or pool of interactions necessary for pathway and network building. This pool is not likely to be 100% accurate because of all the reasons I previously mentioned. Software tools and methods developed for pathway building must take into account the reliability of each interaction when working with such a database. At this point, it is worth emphasizing the difference between a network and a pathway, because both can be built by the same software tools. A network represents a static image of all possible physical and/or regulatory interactions between biological entities, while a pathway represents how the information propagates through the network. Because information propagation is a directional process, a pathway must have entry nodes where the information flow starts and terminal points where the information flow ends. The pathway components represent the sequence of molecular events in space or time while the biological process is occurring. A network, in contrast to a pathway, can contain any relation or entity, including those that do not participate in the information flow. For example, a physical interaction network can include structural interactions, while regulatory networks can include indirect relations that actually are mediated by a set of physical interactions. Network analysis can provide

important insights into biological functions. It can, for example, identify major regulators and targets in a biological process that appear as hubs—nodes with many connections for this process in the network. Hubs also can provide an idea about the information flows within the network, starting from major hub regulators and propagating toward major hub targets. Network analysis can also identify protein complexes involved in a process. Yet, a network cannot be used for dynamic modeling because it lacks one essential ingredient of any pathway: an initial signal or input.

Because a pathway is a way to represent specific events that take place after exposing a cell to an extracellular signal or environmental condition, the main task of pathway analysis is to establish methods for converting network information into a pathway. Any biological pathway is essentially an abstraction or an approximation describing the major channels of information flow through the physical interaction network. The goal of pathway inference is generating a diagram simple enough to be used in kinetic simulations yet adequately describing what happens inside a cell after the stimulation. That known canonical pathways represent the preferred path through the network was proposed some time ago [12]. If this premise is true, then a pathway is simply a path through hubs in the global regulatory network. The evidence taken from scientific literature certainly supports this view: essentially every component in any currently known canonical pathway is a hub in the global and regulatory network. Hubs should be also the first experimentally detectable components of a pathway because they are the best targets for pathway inhibition, which is the favored method to study pathways *in vivo*. For this reason, the currently known connectivity of a hub must be artificially elevated relative to other “non-hub” proteins in both physical and regulatory networks derived from experimental literature: historically, researchers first identified hubs as principal pathway components and then began identifying other proteins interacting with them. Thus, in reality, the relative connectivity of hubs may not be as high as it appears to be in the currently known network. Nevertheless, currently known hubs will still probably remain hubs even after the entire network is known.

Figure 1.1 represents the principal task of pathway analysis: converting a network into a pathway suitable for dynamic modeling. The input for pathway analysis is a network and a stimulus used to invoke the information flow. The output of pathway analysis is a pathway suitable for dynamic modeling with adequate predictive power. Where does the input network come from? From another class of high-throughput experiments aimed at measuring the state of an entire biological system, such as gene expression microarray experiments. In spite of being noisy like all other biological high-throughput methods, they provide a snapshot of a system upon exposing cells to a stimulus. Ideally, the state of a cell must be measured on several different levels such as cell transcriptome, proteome, and metabolome [13]. The current state of the art, however, produces only gene expression data with acceptable quantity and accuracy. Yet recent developments in proteomic methods measuring protein

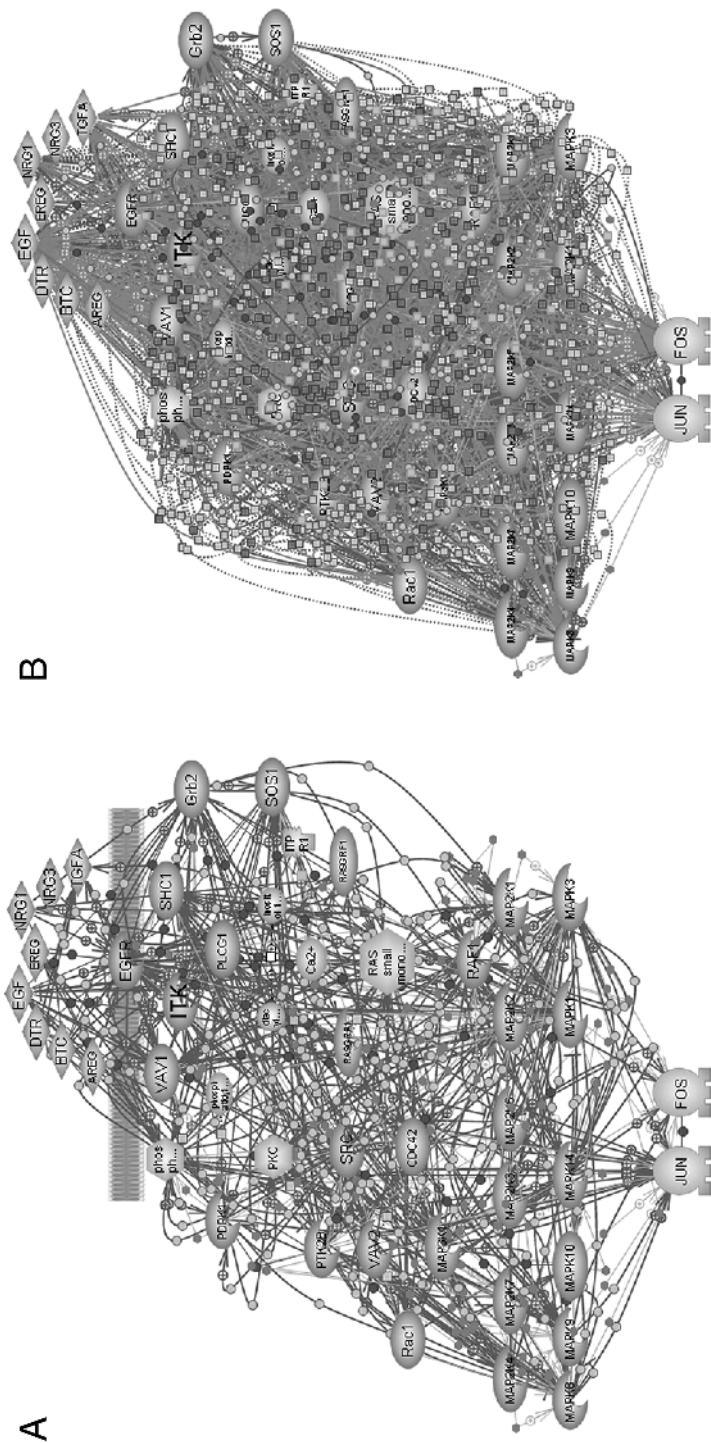
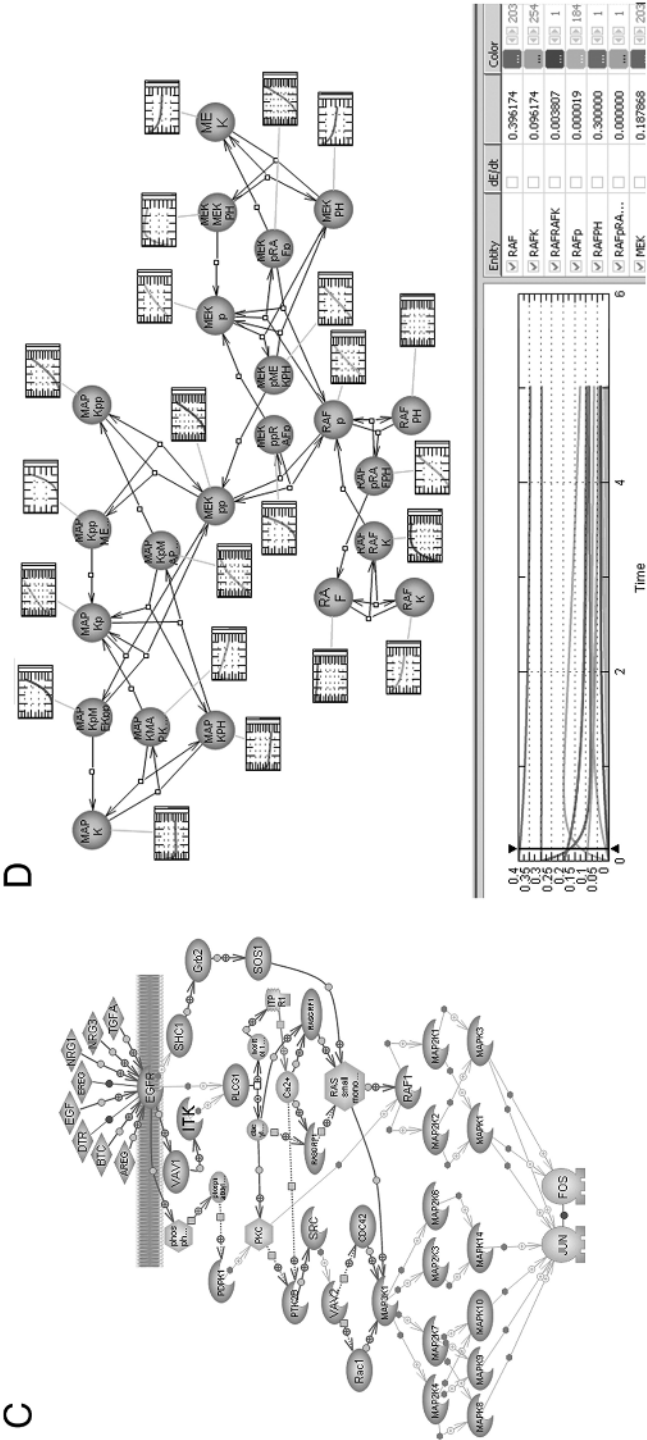


Figure 1.1 Canonical EGFR pathway shown in three ways: (A) physical interaction network; (B) physical interaction network combined with regulatory network; (C) as a pathway ready for kinetic modeling; and (D) kinetic model of MAP kinase cascade, which is a portion of EGFR pathway. One way to visualize the goal of pathway analysis is to imagine that it changes the contrast of the network image to make the main information flow more visible by hiding relations that are nonessential for depicting the principal information flow. See color insert.



concentration and modification already necessitate integration of information from different experimental types into software for pathway analysis. The important workflow described above must be supported by all pathway analysis software: it must provide access to the molecular interaction database, permit analysis of high-throughput data to identify molecular networks appearing in response to an experiment, and subsequently allow calculation of pathways suitable for molecular modeling.

An additional approach for building pathways for the human organism is by using orthologous pathway information from model organisms and paralogous information about known pathways in human tissues. One of the most important achievements of network biology in the last decade is providing further support to the duplication-divergence theory of molecular evolution (see reference [14] and references therein). The best way to evolve is to duplicate an existing mechanism and then modify one or both copies to develop new functions while keeping older functions in one of the copies, if necessary. Evolution constantly duplicates individual genes and occasionally makes a copy of entire genomes in order to mutate genes later and to develop new interactions and functions [15]. As further evidence in support of the duplication-divergence model of evolution, current efforts in studying model organisms provide crucial insights into general rules for the modular and pathway organization of a cell. They have revealed and will reveal more of the conserved, “must have” mechanisms in molecular signaling and cell physiology [16]. The following list enumerates currently known conserved principles of pathway organization:

- Pathway sub-compartmentalization using clustering in physical interaction networks [17] and scaffolding [18,19]
- Fast decay of crosstalk mediated by binding interactions [20]
- Feedback loops providing positive self-activation of a pathway [18,21]
- Feed-forward loops providing noise tolerance for a pathway [22,23]
- Cross-pathway inhibition [18,22,24]

In addition, model organisms reveal the conserved molecular interaction and regulatory blocks necessary for biologically meaningful propagation of information [25], such as the MAPK-kinase cascade, for example [18].

1.5 PATHWAY ANALYSIS SOFTWARE AND THE SCIENTIFIC COMMUNITY

While providing the tools and methods for pathway building and data analysis, pathway analysis software provides additional important functions for scientific enterprise: enabling fast communication, data exchange, and education

among members of the scientific community. It has been recognized that graphical and other visual information is more effective than text for learning molecular biological concepts [26]. For example, the image of the double-helix DNA structure is the most common form used by people to learn, think, and teach about DNA. This image migrates from one textbook to another. Any text describing the double-helix is merely a caption for the image. Similarly, diagram visualization in pathway analysis software allows scientists to exchange information about biological networks and work with them more efficiently. I want to demonstrate how important visualization is for pathway analysis by suggesting the following virtual experiment. First, take any pathway diagram that you find in this book and describe it in writing. Be warned that this exercise may be very boring. Second, find two of your colleagues who have never seen this pathway before. Show the diagram to the first colleague and show the text to the second one. Give both of them the same amount of time to inspect and memorize the pathway information. Then, ask them both to reproduce the pathway. You will discover that the colleague who saw the pathway will reproduce it more accurately than the person who read about the same pathway. Now, try to draw the same pathway yourself and you will realize that it is much faster to describe a pathway as text than to make a good drawing of it. Describing a pathway is easier because you most likely have a text processor program on your computer but do not have an application for pathway drawing. Imagine, however, that this pathway-drawing program exists and also enables you to send a pathway diagram to your colleagues so that they can reproduce an exact copy of your pathway, add their information to it, or compare it to their own experimental results or to other pathways. It should now be readily apparent that a pathway analysis tool can increase your ability to communicate with the scientific community and speed up your collaboration projects many times over.

Biologists usually visualize three major classes of processes as diagrams: biochemical pathways, molecular signaling cascades, and various cellular mechanisms such as a cell cycle and apoptosis. As all scientific papers are written these days using computers, pathway diagrams are drawn with computer programs as well. The degree of sophistication of these pathway-drawing programs ranges from simple vector graphics drawing tools like Microsoft PowerPoint to the database programs like Pathway Studio from Ariadne Genomics that link every pathway to the underlying global molecular interaction network and to the functional annotation of biological molecules. At the same time, these programs allow the comparison of thousands of data points from high-throughput experiments with the pathway collection in the database. The increased sophistication of pathway drawing tools has put the term “pathway analysis” on the same level with other scientific methods and disciplines that study the propagation of information inside the living cell, such as: systems biology, molecular network analysis, and dynamic modeling or kinetic simulation. In the remaining part of this

introduction, I will position pathway analysis relative to these approaches in an attempt to show how pathway analysis both differs from and complements them.

1.6 PATHWAY ANALYSIS AND SYSTEMS BIOLOGY

In short, systems biology is a discipline and pathway analysis is one of its methods. Historically, the term “systems biology” was used as an umbrella to describe various attempts to understand and to model the behavior of an entire cell or organism. Since our current biological knowledge is still incomplete, systems biology focuses on the development of computational methods for analysis of high-throughput data and on designing databases and data models to store and refine the information necessary for achieving the ultimate goal of modeling cell physiology. Pathway analysis is not different in this respect from any other methods of systems biology. It allows compiling, maintaining, classification, and utilization of pathway information. The reason why pathway analysis must be isolated from other methods is evident from the following estimates. There are more than 520 signaling ligands in the human genome and 232 tissues in the human organism. Even though many tissues do not have receptors for every ligand, one hormone often can bind different types of receptors and sometimes activate different pathways [27]. Therefore, we can estimate about 100,000 signaling diagrams that are necessary in order to have a comprehensive collection of signaling pathways for the human organism. Variations of about 50 canonical biochemical pathways in 232 human tissues add another 10,000 diagrams. The number of various intracellular and physiological intercellular processes can be estimated to be about 1,000. This estimate increases the number of necessary diagrams to roughly about 200,000. Finally, there are about 2,500 complex diseases that are usually depicted as diagrams of defective pathways and cellular processes. We also must remember that pharmaceutical research typically uses animal models that require a separate pathway collection for each model organism: mouse, rat, dog, etc. All of the previous numbers estimate a daunting collection of nearly 500,000 pathways needed to build a comprehensive database for drug discovery. To create and maintain this vast pathway collection, we need a rather sophisticated software infrastructure. This software must provide interactive access to pathway diagrams, enable fast and seamless data exchange between users and databases, and allow the comparison of pathway collection with high-throughput and other experimental data. One of the goals for pathway analysis is to devise strategies and algorithms to compose such a collection and to develop an appropriate software infrastructure for its maintenance, for its utilization in analysis, and for drug discovery. The effort to create this pathway collection for drug discovery is comparable in scale to the Human Genome Project (HGP). Continuing this analogy, I would say that the current technological level of pathway analysis is about the same as the level of

sequence analysis soon after Frederick Sanger proposed his method for DNA sequencing.

1.7 PATHWAY ANALYSIS AND NETWORK ANALYSIS

Briefly, network analysis studies the global properties of biological networks, while pathway analysis studies the propagation of information through a network. I have already described the pathway diagram as an intermediate step between isolating the network modules and building the dynamic model of a biological process, as shown in Figure 1.1. The method of network analysis gained momentum in 1999 after a publication by Hartwell et al. [28] that presented an intuitively clear paradigm about the modular organization inside a living cell. Since then, biological networks have undergone intense investigation. The recent efforts to analyze biological networks have yielded several important findings relevant to pathway analysis. I will list them here, since this book does not cover network analysis in much detail. First, it has been demonstrated that both physical and regulatory biological networks indeed have a modular structure that correlates well with known functional modules, as defined by humans in protein annotation [17]. Second, it was found that biological protein networks tend to have power law degree distribution, meaning that they have hubs—highly connected proteins with many interaction partners. This discovery appears to be true for both physical and regulatory networks (Figure 1.2). The main reason for having a hub or scale-free network topology is to provide robustness to the network so it does not break into independent components when a link or node is occasionally removed [29]. A node or link removal can occur because of the following: genetic mutations in evolution; somatic mutations occurring in a disease and throughout the life of the organism; and environmental conditions such as diet, trauma, and stress. The scale-free topology of a network allows biological systems to randomly try new ways of evolutionary adaptation without the danger of disintegration and to survive rather significant damage during a disease.

1.8 PATHWAY ANALYSIS OF DISEASE

Network biology logic suggests that, in order to become stable and robust, a disease network must acquire a scale-free topology with hubs. A disease is developed if its sub-network or module that carries out the malignant function becomes robust and perpetual. These perpetual networks may take years to develop in an organism. Genetic predisposition, diet, lifestyle, infection, accidental trauma, and stress all may contribute to the development of a disease network. Because of these multiple contributions, the networks causing the same disease most likely differ among individual patients, thus necessitating personalized drug intervention. To cause the same disease in different

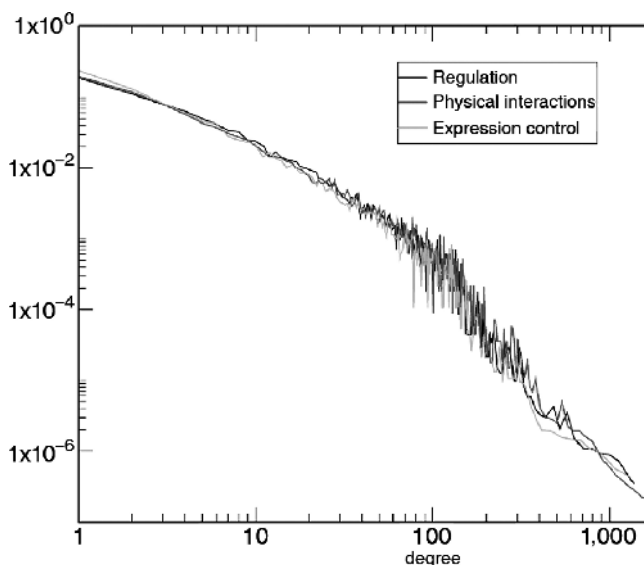


Figure 1.2 Degree distribution among known physical and regulatory interactions in ResNet database from Ariadne Genomics. The networks were built by automatically extracting relations from scientific literature using MedScan, a natural language processing technology [16].

individuals, individualized disease networks must, however, have in common the misdirected information flow that should occur through the sharing of the pathway components. Since hubs are the best targets to disrupt a scale-free network, it is yet to be seen whether individualized disease networks have similar hub composition or they overlap only in the “peripheral” nodes that are responsible for malignancy. Once formed, the disease network should be resilient to drugs precisely because of the robustness that has enabled its existence in the first place [30]. Therefore, from the network biology perspective, a drug design strategy must be a strategy to disrupt the robustness of a disease network. The network disruption itself cannot be adequate and must be supplemented by other changes that will make the disruption irreversible. Irreversibility can be achieved by either using other drugs or changes in lifestyle and diet. It is highly unlikely, however, that drugs will restore the original “normal” *network*, due to the general complexity and evolutionary nature of the development of the malignant network. Nonetheless, drugs must restore the “healthy” information flow, which is a desirable clinical outcome. Therefore, from the pathway analysis perspective, drugs must redirect a malignant information flow to restore the normal, healthy *pathway*. In many cases, this restoration can mean returning to the original pathway disrupted by a disease. In some cases, for example, in diabetes or obesity, it may mean “improving” the original pathway and making it even “healthier.”

I must mention that no disease network has been completely established, and the previous paragraph is only a speculation. This book will offer examples of the current efforts toward building such networks. Here I want to mention the networks recently identified for inherited ataxia [31] and the angiogenic switch in human pancreatic cancer [32]. Most complex diseases, such as cancer or diabetes, are multistep processes of malignant transformation. Each step is probably characterized by a different robust malignant network. Therefore, it is not entirely correct to speak about a cancer disease network, for example. It is quite appropriate to speak about and study a cancer pathway, however. This pathway represents the path of cancer development that has a start, a direction, and a final stage. Every step on this pathway represents a biological process whose defect causes the progression of the disease (Figure 1.3). Each step has its own malignant network. This is another example of the difference between a pathway diagram and a network diagram: the depiction of disease history as a chain of events in time.

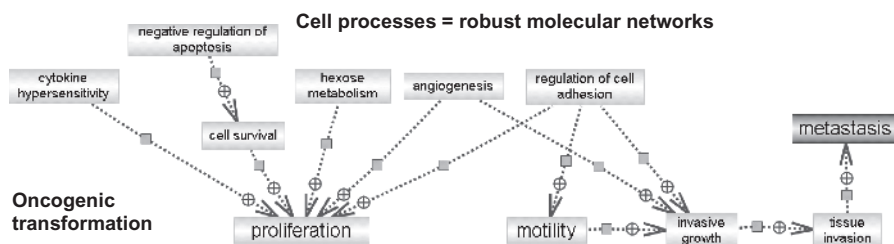


Figure 1.3 Typical cancer development pathway. This diagram depicts the sequence of cell transformation events occurring from the onset of the disease until the death of the patient. Each transformation is triggered by defects in one or several cellular processes shown at the top level of the diagram. Depending of the type of cancer, some steps may be irrelevant or skipped. For example, angiogenesis is not necessary for oncogenic transformation of blood cells. The exact order of events contributing to malignant cell proliferation is not known. In principle, this sequence can vary in different patients or cancer types, but tumor-induced angiogenesis occurs at the later stages of cancer after micro-tumors have reached a certain size. Angiogenesis in the context of cancer development means the induction of blood vessels growth by tumors. Therefore, the angiogenesis network in cancer cells mediates the production of cytokines responsible for angiogenesis in endothelial cells and the angiogenic switch network in endothelial cells mediates proliferation and differentiation of endothelial cells.

A defect in every process contributing to cancer development occurs through the establishment of the robust molecular network performing the malignant function in the cell [73]. These networks are beginning to be identified [31] for each step of cancer progression. In principle, these networks should be different in every human tissue, and each tissue may have several networks for each cancer step depending on individual genetic and environmental factors that lead to the development of cancer in the first place.

1.9 PATHWAY ANALYSIS AND DYNAMIC MODELING IN DRUG DEVELOPMENT

In essence, dynamic modeling or pathway kinetic simulation requires building a pathway diagram prior to developing a kinetic model. The mathematical model developed from experimental data is considered to be the final triumph of the effort to understand biological processes. The main criterion for successful mathematical modeling is the correct simulation of experimentally observed behavior. For this reason, good models can be built only when a system or a process is studied well enough to have known reproducible behavior in response to a stimulus. Very few examples of such processes in molecular biology exist, imposing a major limitation on the development of kinetic models. Experimentally observed cycling, oscillations, and threshold behavior appear as the most attractive targets for kinetic modeling. Hence, the first models developed in molecular biology were for the cell cycle [33], the circadian cycle [34], and the oscillation in NF-kappaB pathway [35]. Recently, the models simulating development of apoptosis [36–41], signaling in the EGFR receptor [42], and the JAK-STAT pathway [43] were developed. The criteria for a successful apoptosis model were selected based on known drug effects [44,45], threshold behavior or bistability between cell apoptotic responses and survival responses to cytotoxic stress [39], and tissue-specific differences in bistable (irreversible) and monostable (reversible) apoptotic responses [40]. For the EGFR signaling model, the criterion for success was the known activation profile of EGFR downstream targets. For the JAK-STAT pathway, the criterion for success was the cycling of STAT protein between cytoplasm and nuclei.

A mathematical theory for modeling of signaling pathways has been put forward by Heinrich et al. [46]. The theory was designed to model rate, duration, and amplitude of a signal in linear kinase-phosphatase cascades, coupled to feedback interactions and crosstalk with other signaling pathways. Undoubtedly, these modeling attempts contribute to our general understanding of biological processes dynamics. For example, they showed that phosphatases affected the rate and duration of signaling, whereas kinases controlled signal amplitude in the EGFR pathway [46]; that RAF-1 signaling was the most important regulator of EGFR phosphorylation [47]; that EGF-induced responses were stable over a wide range of ligand concentration; and that the initial velocity of receptor activation determines signaling efficiency through the EGFR pathway [42].

The behavior of a biological system, by definition, must be probabilistic in order to cope with novel environmental factors previously unmet in evolution. The only way a cell can find the optimal response to a novel stimulus is by presenting all possible responses while looking for the first optimal one that may become selected. For example, that cell transcriptional response to previously unknown challenges is fundamentally random was recently shown for yeast [9]. Drug treatment, by all means, represents the unmet challenge for a

cell and therefore the cell response to a drug must be a stochastic process, which varies among the cell population of one or several human tissues. Transcriptional and epigenetic reprogramming of an entire system in response to a drug may provide strong limitations on using dynamic models in drug development. Drug-induced cell reprogramming can be caused by both the inhibition of an intended drug target and the side effects of a drug. If a drug causes global transcriptional reprogramming in a cell, its efficacy cannot be predicted by solely using the kinetic models. Instead, the new cell state first must be determined through using experimental measurements in every individual patient. The reprogrammed cell will have changed relative concentrations of proteins involved in the process that is to be modeled. Therefore, mechanistic dynamic modeling can be useful in evaluating how close the new cell state and the desired clinical outcome are by calculating the dynamics of affected processes in a new state. In the case when a drug does not cause the global reprogramming, or reprogramming is mild, it should be possible to predict the cell response by a dynamic model. Due to the variability of the initial conditions among cells in the body, the biologically meaningful outcome of any modeling should be a solution space of all possible cell responses to a drug treatment with a specific probability score assigned to every response curve in the solution space. Even though computational approaches to address this problem are being developed [48–50], they currently suffer from the general lack of knowledge about intracellular events, which is necessary for creating the model. These approaches also have knowledge gaps about cell behavior that is necessary for model validation.

A global cell model is the ultimate goal of pathway analysis. Its general complexity with millions of freely adjustable parameters probably will require more experimental constraints than modern molecular biology can ever provide. Thus, the uncertainty due to the lack of knowledge about the system must be added to the fundamental variability of the cell response, making the predictions by available models even less reliable. Nevertheless, current efforts will ultimately yield a computational model for the entire cell. The useful application of the global cell models, however, will be made in the even more distant future than the creation of complete database of molecular interactions for the human cell. Therefore, the best success stories of dynamic modeling are likely to happen beyond the life span of most readers of this book.

One goal of current dynamic modeling efforts is to isolate the minimal set of components and relations that can be used to correctly predict the behavior of a system and then to predict a system response, such as drug action, to the perturbations. While creating a model that uses pathway analysis methods, one can identify essential interactions from the pool of all cellular interactions mediating the modeled process. Once principal interactions are identified from the pool of all known interactions, it may become evident that the pathway is incomplete and actually misses some interactions. Pathway analysis can point an investigator to missing pathway components and help in the design of an appropriate experiment for identifying missing components. An

assumption that the correct evaluation and prediction of drug efficacy requires the kinetic simulation of pathways containing drug targets automatically necessitates simulating hundreds of thousands of pathways for the same reasons described in previous sections. Such a large number of models will require automatic assembly using pathway analysis tools. As you will see in upcoming chapters of this book, pathway analysis may provide some alternatives to brute force kinetic modeling in evaluating drug efficacy and toxicity, as well as in selecting drug targets.

1.10 STEADY-STATE ANALYSIS OF METABOLIC NETWORKS

Flux balance analysis (FBA) [51] and extreme pathway analysis [52] are two main methods of global analysis of metabolic networks developed by Bernard Palsson and his colleagues in the last 10 years. They take advantage of a basic physical principle of mass balance with the reasonable assumption that intracellular metabolic reactions are in a steady state and therefore have constant fluxes. Certainly, metabolism changes multiple times throughout the day in the human organism [53]. These changes are rapid and the periods between changes can be viewed as a steady state. Steady-state analysis predicts a set of allowed metabolic phenotypes or phenotypic planes within the space of all possible fluxes while avoiding any kinetic modeling. The allowed flux space is calculated using standard linear algebraic methods and has the form of a convex polyhedral cone. The cone edges are called “extreme pathways” [54] and the space between them is a phenotypic plane [55]. The initial FBA solution space turned out to be very large, and the problem of calculating extreme pathways was found to be NP-hard. For genome-scale networks, this solution space has proven to be infeasible [56]. For those reasons, additional constraints on thermodynamics [57], expression regulation [58], compartmentalization [59], maximum capacity, and reaction irreversibility [60] had to be taken into account to narrow down the solution space.

In my view, metabolic reconstruction represents the most advanced method available for pathway analysis. It uses the power of the complete genome sequence, modern mathematical and computational approaches to model metabolism using basic physical principles. Yet, it possesses several limitations for drug discovery that do not allow it to become the main focus of this book. First, metabolic control in multicellular organisms is under tight hormonal control. Second, the scale of changes detectable by this method may be too detailed and irrelevant for drug discovery. It is more important to predict global metabolism homeostasis of an entire organism than the metabolism of the local cell population. The difficulties of finding a physiological interpretation of the results from extreme pathway analysis were acknowledged by Bernard Palsson himself [61].

FBA has proven to be a very good method to model the metabolic state and to predict growth conditions and metabolite yields of the microorganisms

[62,63]. But the applications of FBA to drug discovery are yet to be found. FBA proponents suggested using this technique to model the metabolic state of mammalian mitochondria [64], red blood cells [56], and even the classical JAK-STAT signaling pathway [65]. One obvious application of FBA in drug discovery is predicting the response of microbial flora in the human body to drugs. The human body contains about 1.5 kilograms of symbiotically living microorganisms. Bacteria such as *E. coli* [58], *H. influenza* [66], and *H. pylori* [67] are among them and were modeled by Palsson and his colleagues. A better understanding of bacterial metabolism will likely help to develop better drugs and antibiotics against diseases linked to bacteria, such as diarrhea and gastritis.

Metabolic control analysis (MCA) is complementary to the FBA method for modeling biochemical reactions in a steady state. Its goal is to calculate a control coefficient for each enzymatic step in a pathway, reflecting the extent to which the component is rate-limiting [68]. MCA has been used to identify transketolase as a rate-limiting enzyme responsible for thiamine deficiency in the Ehrlich's ascites tumor model [69]. Interestingly, the authors of this paper failed to explain why a very high concentration of thiamine exhibits an inhibitory effect on the tumor, which is the opposite of its effect predicted by MCA and observed at modest concentrations. The metabolic control analysis methodology has been recently modified for transient signaling through the classical MAP kinase cascade [47,70].

1.11 SO WHAT IS PATHWAY ANALYSIS?

To answer this question, I want to remind you that pathway analysis is a method or a tool. As with any tool, it has a finite life span: it has appeared out of necessity to solve several biological problems currently facing the scientific community. Once these problems are solved, most pathway analysis methods will become obsolete, as it happened with many sequence analysis methods after human genome sequencing was completed. I have outlined problems that have to be solved by pathway analysis in previous sections and will summarize them in the next paragraph, while putting them into historical and social perspective.

In order to build a comprehensive computational model of the living cell, we need to know all of the principal components of the system plus the interactions that enable information flow through the system. The sequencing of the human genome has moved our understanding of the molecular mechanisms that govern life processes from the phase when many components and their interactions were unknown to the phase when most components are known but the interactions between them largely are not. More and more molecular interactions are being discovered as a result of the development of new high-throughput methods to measure individual interactions and because of the individual efforts of biologists worldwide. In 2004, biologists estimated

that the discovering of all possible biological interactions would take about 20 years [71]. This estimate is consistent with my previous estimates of having these goals completed within the life span of most readers of this book. While working toward this goal for another 15–17 years, we also must learn how to use this information to better diagnose diseases and to improve our drugs. Let me direct you again to Figure 1.1 for an illustration of the ultimate goal of pathway analysis: converting a huge, unstructured molecular interaction network into well-defined modules describing the major information flows that can be used for predictive dynamic modeling. These modules will describe individual biological processes and signaling cascades with enough accuracy so that they can be combined into the computational model of an entire cell. The global cell model will ultimately consist of these modules linked by inter-model interactions and will predict cell response to a drug. Because generally speaking, drug development and the state of public health are the most important practical validations for any achievements in biological science, the goal of pathway analysis must be to yield improvements in drug development and disease prevention.

The community of billions of interacting proteins yields the visible living structure that we call the cell. Community of cells in turn forms human organism. Both organism and an individual cell can also be described as an information system that constantly follows and responds to environmental signals. The information inside a cell is propagated through physical interaction network. Therefore, a disease can be described as broken pathways or as a disruption of a physical structure maintained by the network of interacting proteins. Effort in pathway analysis shall eventually yield the representation of human organism as an information system via the collection of pathways. This collection will allow quick assessment of the personal health state and efficiency of a drug action.

The collection of hundreds of thousand of pathways can exist and be used only in electronic form, which makes pathway analysis a discipline of computer science. However, the roots of pathway analysis are in experimental molecular biology and genetics. The classical deductive path in molecular biology starts with the definition of a phenotype or biological process. In principle, it is the most important step that defines the overall goal of research. The rest of the effort is “simply” elucidation of the molecular mechanisms driving the process or explaining the phenotype. The word “simply” is in quotation marks because, in reality, the understanding of the molecular mechanism for every biological process once was a daunting task that involved thousands of researchers, millions of dollars spent on salaries and reagents, and usually produced several Nobel laureates. The magnitude of the effort was multiplied by the general lack of understanding of basic principles governing the function and structure of biological molecules.

The biological knowledge base has dramatically changed after Human Genome Project (HGP) was finished. First, all missing pathway components

have suddenly become known. Second, the structure and function of proteins have suddenly become predictable using sequence homology and new computational algorithms. Finally, protein–protein interactions have also become abundant and predictable with a certain amount of accuracy, thanks to the development of high-throughput methods. Currently, more than 70% of all known physical protein–protein interactions for human proteins were measured by high-throughput experiments. By my estimate, there are about 200,000 experimentally determined protein–protein interactions and more than 5,000,000 predicted interactions for human proteins. The next challenge is to create a multitude of pathways and dynamic models. The lack of tools and talent to develop pathway infrastructure was indicated in 2000 by the same Bernard Palsson [72]. The situation has changed 7 years later in part due to Palsson’s efforts, but also due to the efforts of the scientific community, several commercial bioinformatics companies, and scientists in big pharmaceutical companies. With these changes in mind, this book provides an overview of current tools, methods, and software for pathway analysis.

REFERENCES

1. Daraselia N, Yuryev A, Egorov S, Novichkova S, Nikitin A, Mazo I. Extracting human protein interactions from MEDLINE using a full-sentence parser. *Bioinformatics* 2004;20(5):604–611. Epub 2004 Jan 22.
2. Hart G, Ramani A, Marcotte E. How complete are current yeast and human protein-interaction networks? *Genome Biol* 2006;7:120.
3. Ewing RM, et al. Large-scale mapping of human protein–protein interactions by mass spectrometry. *Mol Syst Biol* 2007;3:89.
4. Kotelnikova E, Kalinin A, Yuryev A, Maslov S. Prediction of protein–protein interactions on the basis of evolutionary conservation of protein functions. *Evol Bioinform* 2007;3:323.
5. Yu H, Luscombe NM, Lu HX, Zhu X, Xia Y, Han JD, Bertin N, Chung S, Vidal M, Gerstein M. Annotation transfer between genomes: protein-protein interologs and protein-DNA regulogs. *Genome Res* 2004;14(6):1107–1118.
6. Heron EA, Finkenstadt B, Rand DA. Bayesian inference for dynamic transcriptional regulation; the Hes1 system as a case study. *Bioinformatics* 2007; 23:2596–2603.
7. Csete ME, Doyle JC. Reverse engineering of biological complexity. *Science* 2002;295(5560):1664–1669.
8. Michoel T, Maere S, Bonnet E, Joshi A, Saeys Y, van den Bulcke T, Van Leemput K, Van Remortel P, Kuiper M, Marchal K, Van De Peer Y. Validating module network learning algorithms using simulated data. *BMC Bioinformatics* 2007;May 3;8(Suppl 2):S5.
9. Stern S, Dror T, Stolovicki E, Brenner N, Braun E. Genome-wide transcriptional plasticity underlies cellular adaptation to novel challenge. *Mol Syst Biol* 2007;3:106. Epub 2007 Apr 24.

10. Yuryev A, Mulyukov Z, Kotelnikova E, Maslov S, Egorov S, Nikitin A, Daraselia N, Mazo I. Automatic pathway building in biological association networks. *BMC Bioinformatics* 2006;Mar 24;7:171.
11. Rodriguez-Esteban R, Iossifov I, Rzhetsky A. Imitating manual curation of text-mined facts in biomedicine. *PLoS Comput Biol* 2006;2(9):e118.
12. Marcotte EM. The path not taken. *Nat Biotechnol* 2001;19:826.
13. Ideker T, Thorsson V, Ranish JA, Christmas R, Buhler J, Eng JK, Bumgarner R, Goodlett DR, Aebersold R, Hood L. Integrated genomic and proteomic analyses of a systematically perturbed metabolic network. *Science* 2001;292(5518):929–934.
14. Ispolatov I, Krapivsky PL, Yuryev A. Duplication-divergence model of protein interaction network. *Phys Rev E Stat Nonlin Soft Matter Phys* 2005;71(6 Pt 1):061911.
15. Sankoff D. Gene and genome duplication. *Curr Opin Genet Dev* 2001;11(6):681–684.
16. Schwartz MA, Madhani HD. Principles of MAP kinase signaling specificity in *Saccharomyces cerevisiae*. *Annu Rev Genet* 2004;38(1):725–748.
17. Daraselia N, Yuryev A, Egorov S, Mazo I, Ispolatov I. Automatic extraction of gene ontology annotation and its correlation with clusters in protein networks. *BMC Bioinformatics* 2007;8:243.
18. Morrison D, Davis R. Regulation of MAP kinase signaling modules by scaffold proteins in mammals. *Annu Rev Cell Dev Biol* 2003;19(1):91–118.
19. Levchenko A, Bruck J, Sternberg P. Scaffold proteins may biphasically affect the levels of mitogen-activated protein kinase signaling and reduce its threshold properties. *Proc Natl Acad Sci U S A* 2000;97(11):5818–5823.
20. Maslov S, Ispolatov I. Propagation of large concentration changes in reversible protein-binding networks. *Proc Natl Acad Sci U S A* 2007;104(34):13655–13660.
21. Yeung K, Janosch P, McFerran B, Rose D, Mischak H, Sedivy J, Kolch W. Mechanism of suppression of the Raf/MEK/Extracellular signal-regulated kinase pathway by the RAF kinase inhibitor protein. *Mol Cell Biol* 2000;20:3079–3085.
22. Acar M, Becskei A, van Oudenaarden A. Enhancement of cellular memory by reducing stochastic transitions. *Nature* 2005;435(7039):228–232.
23. Brandman O, Ferrell J, Li R, Meyer T. Interlinked fast and slow positive feedback loops drive reliable cell decisions. *Science* 2005;310(5747):496–498.
24. Dang V, Bohn C, Bolotin-Fukuhara M, Daignan-Fornier B. The CCAAT box-binding factor stimulates ammonium assimilation in *Saccharomyces cerevisiae*, defining a new cross-pathway regulation between nitrogen and carbon metabolisms. *J Bacteriol* 1996;178(7):1842–1849.
25. Sharan R, Suthram S, Kelley RM, Kuhn T, McCuine S, Uetz P, Sittler T, Karp RM, Ideker T. Conserved patterns of protein interaction in multiple species. *Proc Natl Acad Sci U S A* 2005;102(6):1974–1979.
26. Mayer RE, Hegarty M, Mayer S, Campbell J. When static media promote active learning: annotated illustrations versus narrated animations in multimedia instruction. *J Exp Psychol Appl* 2005;11(4):256–265.
27. Deupi X, Kobilka B. Activation of G protein-coupled receptors. *Adv Protein Chem* 2007;74:137–166.

28. Hartwell LH, Hopfield JJ, Leibler S, Murray AW. From molecular to modular cell biology. *Nature* 1999;402:C47–52.
29. Barabási A-L, Oltvai Z. Network biology: understanding the cell's functional organization. *Nat Rev* 2004;5:101.
30. Kitano H. A robustness-based approach to systems-oriented drug design. *Nat Rev Drug Discov* 2007;6(3):202–210. Epub 2007 Feb 23.
31. Lim J, Hao T, Shaw C, Patel A, Szabó G, Rual J-F, Fisk C, Li N, Smolyar A, Hill DE, Barabási A-L, Vidal M, Zoghbi HY. A protein-protein interaction network for human inherited ataxias and disorders of Purkinje cell degeneration. *Cell* 2006;125:801–814.
32. Abdollahi A, Schwager C, Kleeff J, Esposito I, Domhan S, Peschke P, Hauser K, Hahnfeldt P, Hlatky L, Debus J, Peters JM, Friess H, Folkman J, Huber PE. Transcriptional network governing the angiogenic switch in human pancreatic cancer. *Proc Natl Acad Sci U S A* 2007;104(31):12890–12895.
33. Novak B, Tyson JJ. Modeling the control of DNA replication in fission yeast. *Proc Natl Acad Sci U S A* 1997;94(17):9147–9152.
34. Leloup JC, Gonze D, Goldbeter A. Limit cycle models for circadian rhythms based on transcriptional regulation in *Drosophila* and *Neurospora*. *J Biol Rhythms* 1999;14(6):433–448.
35. Covert MW, Leung TH, Gaston JE, Baltimore D. Achieving stability of lipopolysaccharide-induced NF-kappaB activation. *Science* 2005;309(5742):1854–1857.
36. Krammer PH, Kaminski M, Kiessling M, Gulow K. No life without death. *Adv Cancer Res* 2007;97C:111–138.
37. Fussenegger M, Bailey JE, Varner J. A mathematical model of caspase function in apoptosis. *Nat Biotechnol* 2000;18(7):768–774.
38. Bentele M, Lavrik I, Ulrich M, Stosser S, Heermann DW, Kalthoff H, et al. Mathematical modeling reveals threshold mechanism in CD95-induced apoptosis. *J Cell Biol* 2004;166(6):839–851.
39. Bagci EZ, Vodovotz Y, Billiar TR, Ermentrout GB, Bahar I. Bistability in apoptosis: roles of bax, bcl-2, and mitochondrial permeability transition pores. *Biophys J* 2006;90(5):1546–1559.
40. Legewie S, Bluthgen N, Herzel H. Mathematical modeling identifies inhibitors of apoptosis as mediators of positive feedback and bistability. *PLoS Comput Biol* 2006;2(9):e120.
41. Hua F, Cornejo MG, Cardone MH, Stokes CL, Lauffenburger DA. Effects of Bcl-2 levels on Fas signaling-induced caspase-3 activation: molecular genetic tests of computational model predictions. *J Immunol* 2005;175(2):985–995.
42. Schoeberl B, Eichler-Jonsson C, Gilles ED, Muller G. Computational modeling of the dynamics of the MAP kinase cascade activated by surface and internalized EGF receptors. *Nat Biotechnol* 2002;20(4):370–375.
43. Swameye I, Muller TG, Timmer J, Sandra O, Klingmüller U. Identification of nucleocytoplasmic cycling as a remote sensor in cellular signaling by databased modeling. *Proc Natl Acad Sci U S A* 2003;100(3):1028–1033.
44. Li B, Dou QP. Bax degradation by the ubiquitin/proteasome-dependent pathway: involvement in tumor survival and progression. *Proc Natl Acad Sci U S A* 2000;97(8):3850–3855.

45. Katiyar SK, Roy AM, Baliga MS. Silymarin induces apoptosis primarily through a p53-dependent pathway involving Bcl-2/Bax, cytochrome c release, and caspase activation. *Mol Cancer Ther* 2005;4(2):207–216.
46. Heinrich R, Neel BG, Rapoport TA. Mathematical models of protein kinase signal transduction. *Mol Cell* 2002;9(5):957–970.
47. Hornberg JJ, Binder B, Bruggeman FJ, Schoeberl B, Heinrich R, Westerhoff HV. Control of MAPK signaling: from complexity to what really matters. *Oncogene* 2005;24(36):5533–5542.
48. Christopher R, Dhiman A, Fox J, Gendelman R, Haberichter T, Kagle D, Spizz G, Khalil IG, Hill C. Data-driven computer simulation of human cancer cell. *Ann N Y Acad Sci* 2004;1020:132–153.
49. Aksenov SV, Church B, Dhiman A, Georgieva A, Sarangapani R, Helmlinger G, Khalil IG. An integrated approach for inference and mechanistic modeling for advancing drug development. *FEBS Lett* 2005;579(8):1878–1883.
50. Meng TC, Somani S, Dhar P. Modeling and simulation of biological systems with stochasticity. *In Silico Biol* 2004;4(3):293–309.
51. Schilling CH, Edwards JS, Palsson BO. Toward metabolic phenomics: analysis of genomic data using flux balances. *Biotechnol Prog* 1999;15(3):288–295.
52. Schilling CH, Letscher D, Palsson BO. Theory for the systemic definition of metabolic pathways and their use in interpreting metabolic function from a pathway-oriented perspective. *J Theor Biol* 2000;203(3):229–248.
53. Albrecht U, Eichele G. The mammalian circadian clock. *Curr Opin Genet Dev* 2003;13(3):271–277.
54. Schilling CH, Edwards JS, Letscher D, Palsson BO. Combining pathway analysis with flux balance analysis for the comprehensive study of metabolic systems. *Biotechnol Bioeng* 2000–2001;71(4):286–306.
55. Edwards JS, Ramakrishna R, Palsson BO. Characterizing phenotypic plasticity: a phase plane analysis. *Engineering in Medicine and Biology. BMES/EMBS Conference, 1999 Proceedings of the First Joint* 2:1217.
56. Barrett CL, Price ND, Palsson BO. Network-level analysis of metabolic regulation in the human red blood cell using random sampling and singular value decomposition. *BMC Bioinformatics* 2006;7:132.
57. Henry CS, Broadbelt LJ, Hatzimanikatis V. Thermodynamics-based metabolic flux analysis. *Biophys J* 2007;92(5):1792–1805. Epub 2006 Dec 15.
58. Covert MW, Palsson BO. Constraints-based models: regulation of gene expression reduces the steady-state solution space. *J Theor Biol* 2003;221(3):309–325.
59. Feist AM, Henry CS, Reed JL, Krummenacker M, Joyce AR, Karp PD, Broadbelt LJ, Hatzimanikatis V, Palsson BO. A genome-scale metabolic reconstruction for *Escherichia coli* K-12 MG1655 that accounts for 1260 ORFs and thermodynamic information. *Mol Syst Biol* 2007;3:121. Epub 2007 Jun 26.
60. Price ND, Papin JA, Schilling CH, Palsson BO. Genome-scale microbial in silico models: the constraints-based approach. *Trends Biotechnol* 2003;21:162–169.
61. Wiback SJ, Palsson BO. Extreme pathway analysis of human red blood cell metabolism. *Biophys J* 2002;83:808–818.

62. Edwards JS, Palsson BO. The *Escherichia coli* MG1655 in silico metabolic genotype: its definition, characteristics, and capabilities. *Proc Natl Acad Sci U S A* 2000;97:5528–5533.
63. Alvarez-Vasquez F, Sims K, Cowart L, Okamoto Y, Voit E, et al. Simulation and validation of modelled sphingolipid metabolism in *Saccharomyces cerevisiae*. *Nature* 2005;433:425–430.
64. Ramakrishna R, Edwards JS, McCulloch A, Palsson BO. Flux-balance analysis of mitochondrial energy metabolism: consequences of systemic stoichiometric constraints. *Am J Physiol Regul Integr Comp Physiol* 2001;280(3):R695–704.
65. Papin JA, Palsson BO. The JAK-STAT signaling network in the human B-cell: an extreme signaling pathway analysis. *Biophys J* 2004;87(1):37–46.
66. Papin JA, Price ND, Edwards JS, Palsson BO. The genome-scale metabolic extreme pathway structure in *Haemophilus influenzae* shows significant network redundancy. *J Theor Biol* 2002;215(1):67–82.
67. Price ND, Papin JA, Palsson BO. Determination of redundancy and systems properties of the metabolic network of *Helicobacter pylori* using genome-scale extreme pathway analysis. *Genome Res* 2002;12(5):760–769.
68. Hornberg JJ, Bruggeman FJ, Bakker BM, Westerhoff HV. Metabolic control analysis to identify optimal drug targets. *Prog Drug Res* 2007;64:171,173–189.
69. Comin-Anduix B, Boren J, Martinez S, Moro C, Centelles JJ, Trebukhina R, Petushok N, Lee WN, Boros LG, Cascante M. The effect of thiamine supplementation on tumour proliferation. A metabolic control analysis study. *Eur J Biochem* 2001;268(15):4177–4182.
70. Hornberg JJ, Bruggeman FJ, Binder B, Geest CR, De Vaate AJ, Lankelma J, et al. Principles behind the multifarious control of signal transduction. ERK phosphorylation and kinase/phosphatase control. *FEBS J* 2005;272(1):244–258.
71. Aloy P, Russell RB. Ten thousand interactions for the molecular biologist. *Nat Biotechnol* 2004;22(10):1317–1321.
72. Palsson B. The challenges of in silico biology. *Nat Biotechnol* 2000;18:1147–1150.
73. Kitano H. The theory of biological robustness and its implication in cancer. *Ernst Schering Res Found Workshop* 2007;(61):69–88.

