

RESEARCH ARTICLE SUMMARY

SYSTEMS BIOLOGY

Genetic interaction mapping informs integrative structure determination of protein complexes

Hannes Braberg*, Ignacia Echeverria*, Stefan Bohn*, Peter Cimermancic*, Anthony Shiver, Richard Alexander, Jiewei Xu, Michael Shales, Raghuvor Dronamraju, Shuangying Jiang, Gajendradhar Dwivedi, Derek Bogdanoff, Kaitlin K. Chaung, Ruth Hüttenhain, Shuyi Wang, David Mavor, Riccardo Pellarin, Dina Schneidman, Joel S. Bader, James S. Fraser, John Morris, James E. Haber, Brian D. Strahl, Carol A. Gross, Junbiao Dai, Jef D. Boeke†, Andrej Sali†, Nevan J. Krogan†

INTRODUCTION: Determining the structures of protein complexes is crucial for understanding cellular functions. Here, we describe an integrative structure determination approach that relies on *in vivo* quantitative measurements of genetic interactions. Genetic interactions report on how the effect of one mutation is altered by the presence of a second mutation and have proven effective for identifying groups of genes or residues that function in the same pathway. The point mutant epistatic miniarray profile (pE-MAP) platform allows for rapid measurement of genetic interactions between sets of point mutations and deletion libraries. A pE-MAP is made up of phenotypic profiles, each of which contains all genetic interactions between a single point mutant and the entire deletion library.

RATIONALE: We observe a statistical association between the distance spanned by two mutated residues in a protein complex and the similarity of their phenotypic profiles

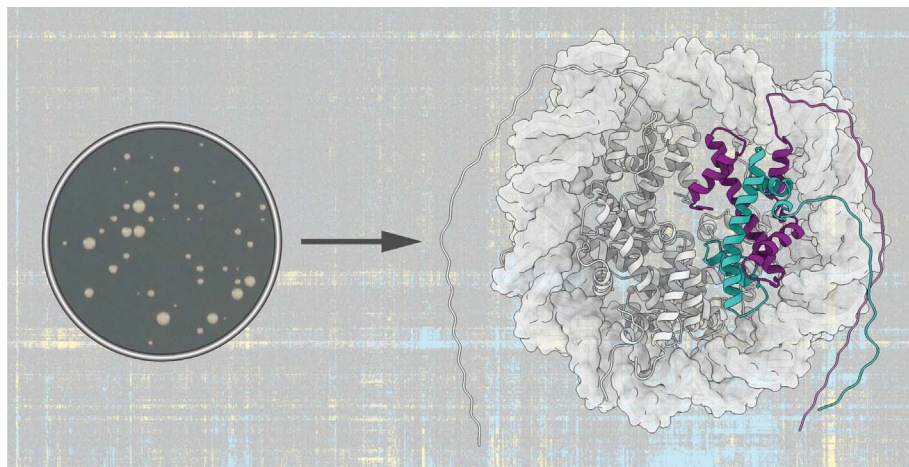
(phenotypic similarity) in a pE-MAP. This observation is in agreement with the expectation that mutations within the same functional region (e.g., active, allosteric, and binding sites) are likely to share more similar phenotypes than those that are distant in space. Here, we explore how to use these associations for determining *in vivo* structures of protein complexes using integrative modeling.

RESULTS: We generated a large pE-MAP by crossing 350 mutations in yeast histones H3 and H4 against 1370 gene deletions (or hypomorphic alleles of essential genes). The phenotypic similarities were then used to generate spatial restraints for integrative modeling of the H3-H4 complex structure. The resulting ensemble of H3-H4 configurations is accurate and precise, as evidenced by its close similarity to the crystal structure. This finding indicates the utility of the pE-MAP data for integrative structure determination. Furthermore, we show that the pE-MAP provides a wealth of biological

insight into the function of the nucleosome and can connect individual histone residues and regions to associated complexes and processes. For example, we observe very high phenotypic similarities between modifiable histone residues and their cognate enzymes, such as H3K4 and COMPASS, or H3K36 and members of the Set2 pathway. Furthermore, the pE-MAP reveals several residues involved in DNA repair and others that function in cryptic transcription.

We demonstrate that the approach is transferable to other complexes and other types of phenotypic profiles by determining the structures of two complexes of known structure: (i) subunits Rpb1 and Rpb2 of yeast RNA polymerase II, using a pE-MAP of 53 point mutants crossed against 1200 deletions and hypomorphic alleles; and (ii) subunits RpoB and RpoC of bacterial RNA polymerase, using a chemical genetics map of 44 point mutants subjected to 83 environmental stresses. The accuracy and precision of the models are comparable to those based on chemical cross-linking, which is commonly used to determine protein complex structures. Moreover, the accuracy and precision improve when using pE-MAP and cross-linking data together, indicating complementarity between these methods and demonstrating a premise of integrative structure determination.

CONCLUSION: We show that the architectures of protein complexes can be determined using quantitative genetic interaction maps. Because pE-MAPs contain purely phenotypic measurements, collected in living cells, they generate spatial restraints that are orthogonal to other commonly used data for integrative modeling. The pE-MAP data may also enable the characterization of complexes that are difficult to isolate and purify, or those that are only transiently stable. Recent advances in CRISPR-Cas9 genome editing provide a means for extending our platform to human cells, allowing for identification and characterization of functionally relevant structural changes that take place in disease alleles. Expanding this analysis to look at structural changes in host-pathogen complexes and how they affect infection will also be feasible by introducing specific mutations into the pathogenic genome and studying the phenotypic consequences using genetic interaction profiling of relevant host genes. ■



In vivo structure determination using genetic interactions. pE-MAPs are generated by measuring the growth of yeast colonies (left) and visualized as a heatmap (background). We present an application of pE-MAPs to determine protein complex structures, using integrative modeling, and apply it to histones H3 and H4 (right) and other complexes. H3 (purple) and H4 (teal) are highlighted in the context of the nucleosome [gray, modified Protein Data Bank (PDB) 1ID3].

The list of author affiliations is available in the full article online.
*These authors contributed equally to this work.

†Corresponding author. Email: nevan.krogan@ucsf.edu (N.J.K.); sali@salilab.org (A.Sa.); jef.boeke@nyulangone.org (J.D.B.)

Cite this article as H. Braberg *et al.*, *Science* **370**, eaaz4910 (2020). DOI: 10.1126/science.aaz4910

S READ THE FULL ARTICLE AT
<https://doi.org/10.1126/science.aaz4910>

RESEARCH ARTICLE

SYSTEMS BIOLOGY

Genetic interaction mapping informs integrative structure determination of protein complexes

Hannes Braberg^{1,2*}, Ignacia Echeverria^{1,2,3*}, Stefan Bohn^{1,2,4*}††, Peter Cimermancic^{3,5,§}, Anthony Shiver^{5,¶}, Richard Alexander¹, Jiewei Xu^{1,2,4}, Michael Shales^{1,2}, Raghuvardhan Dronamraju⁶, Shuangying Jiang⁷, Gajendradhar Dwivedi^{8,¶}, Derek Bogdanoff⁹, Kaitlin K. Chaung⁹, Ruth Hüttenhain^{1,2,4}, Shuyi Wang¹, David Mavor^{2,3}, Riccardo Pellarin^{3,¶}, Dina Schneidman³, Joel S. Bader¹⁰, James S. Fraser^{2,3}, John Morris¹¹, James E. Haber⁸, Brian D. Strahl⁶, Carol A. Gross¹², Junbiao Dai⁷, Jef D. Boeke^{13,14,15,16}††, Andrej Sali^{2,3,11}††, Nevan J. Krogan^{1,2,4,17}††

Determining structures of protein complexes is crucial for understanding cellular functions. Here, we describe an integrative structure determination approach that relies on in vivo measurements of genetic interactions. We construct phenotypic profiles for point mutations crossed against gene deletions or exposed to environmental perturbations, followed by converting similarities between two profiles into an upper bound on the distance between the mutated residues. We determine the structure of the yeast histone H3-H4 complex based on ~500,000 genetic interactions of 350 mutants. We then apply the method to subunits Rpb1-Rpb2 of yeast RNA polymerase II and subunits RpoB-RpoC of bacterial RNA polymerase. The accuracy is comparable to that based on chemical cross-links; using restraints from both genetic interactions and cross-links further improves model accuracy and precision. The approach provides an efficient means to augment integrative structure determination with in vivo observations.

A mechanistic understanding of cellular functions requires structural characterization of the corresponding macromolecular assemblies (1). Traditional structural biology methods—such as x-ray crystallography, nuclear magnetic resonance (NMR) spectroscopy, and electron microscopy (EM)—rely on purified samples and are generally not applicable to heterogeneous samples, such as those of large, membrane-bound, or transient assemblies (2). Moreover, these methods do not determine the structures in their native environments, therefore increasing the risk of producing structures in nonfunctional states or missing relevant functional states.

Integrative structure determination has emerged as a powerful approach for determining the structures of biological assemblies (3). The motivation is that any system can be described most accurately, precisely, completely, and efficiently by using all available information about it, including varied experimental data (e.g., chemical cross-links, protein interaction data, small-angle x-ray scattering profiles) and prior models (e.g., atomic structures of the subunits). Integrative methods can often tackle protein assemblies that are difficult to characterize using traditional structural biology methods alone (1, 4–10). Spatial data generated by in vivo methods are especially useful for integrative structure determination (11). Therefore, high-throughput in vivo methods are needed to supplement low-throughput in vivo methods, such as single-molecule Förster resonance energy transfer spectroscopy (12).

Here, we describe how integrative structure modeling can benefit from spatial restraints derived from in vivo quantitative measurements of genetic interactions. A genetic interaction between two mutations occurs when the effect of one mutation is altered by the presence of the second mutation (Fig. 1A) (13). Positive genetic interactions (epistasis or suppression) arise when the double mutant is healthier than expected, whereas negative interactions (synthetic sickness) arise in relationships where the double mutant is sicker than expected. Single genetic interactions can often be difficult to interpret in isolation. A phenotypic profile, defined as a set of genetic interactions between a given mutation (e.g., a point mutation) and a library of secondary mutations (e.g., gene deletions), can be more informative (Fig. 1B) (14). A point mutant epistatic miniarray profile (pE-MAP) is composed of such phenotypic profiles for all mutations in the analysis (Fig. 1C) (15). We have previously found a statistical association between the distance between two mutated residues in the wild-type (WT) structure and the similarity between their phenotypic profiles (i.e., phenotypic similarity) (15, 16) (Fig. 1D). This observation is in agreement with the expectation that mutations within the same functional region (e.g., active, allosteric, and binding sites) are likely to share more similar phenotypes than those that are distant in space (17–19). Here, we explore how to use these associations for determining in vivo structures of macromolecular assemblies using integrative modeling (Fig. 1E). To enable this analysis,

we generated a large pE-MAP, by designing a comprehensive set of 350 mutations in histones H3 and H4 and crossing these against 1370 gene deletions (or hypomorphic alleles for essential genes). We describe this pE-MAP and illustrate integrative structure determination by its application to three complexes of known structure: (i) the yeast histones H3 and H4; (ii) subunits Rpb1 and Rpb2 of yeast RNA polymerase II (RNAPII), using a pE-MAP dataset of 53 point mutants crossed against a library of 1200 deletions and hypomorphic alleles (15); and (iii) subunits RpoB and RpoC of bacterial RNA polymerase (RNAP), using a chemical genetics miniarray profile (CG-MAP), where 44 point mutants were subjected to 83 different environmental stresses (e.g., treatments with chemicals and temperature shocks) (20).

A comprehensive pE-MAP of histones H3 and H4

Histones are central to chromatin structure and dynamics because they make up the core of the nucleosome, the fundamental repeating unit of chromatin. The state of the nucleosome

¹Department of Cellular and Molecular Pharmacology, University of California, San Francisco, San Francisco, CA 94158, USA. ²Quantitative Biosciences Institute, University of California, San Francisco, San Francisco, CA 94158, USA. ³Department of Bioengineering and Therapeutic Sciences, University of California, San Francisco, San Francisco, CA 94158, USA. ⁴Gladstone Institutes, San Francisco, CA 94158, USA. ⁵Graduate Group in Biophysics, University of California San Francisco, San Francisco, CA 94158, USA. ⁶Department of Biochemistry and Biophysics, University of North Carolina School of Medicine, Chapel Hill, NC 27599, USA. ⁷CAS Key Laboratory of Quantitative Engineering Biology, Guangdong Provincial Key Laboratory of Synthetic Genomics and Shenzhen Key Laboratory of Synthetic Genomics, Shenzhen Institute of Synthetic Biology, Shenzhen Institutes of Advanced Technology, Chinese Academy of Sciences, Shenzhen 518055, China. ⁸Department of Biology and Rosenstiel Basic Medical Sciences Research Center, Brandeis University, Waltham, MA 02454, USA. ⁹Center for Advanced Technology, Department of Biophysics and Biochemistry, University of California, San Francisco, San Francisco, CA 94158, USA. ¹⁰Department of Biomedical Engineering, Johns Hopkins University School of Medicine, Baltimore, MD 21205, USA. ¹¹Department of Pharmaceutical Chemistry, University of California, San Francisco, San Francisco, CA 94158, USA. ¹²Department of Microbiology and Immunology and Department of Cell and Tissue Biology, University of California, San Francisco, San Francisco, CA 94158, USA. ¹³NYU Langone Health, New York, NY 10016, USA. ¹⁴High Throughput Biology Center and Department of Molecular Biology & Genetics, Johns Hopkins University School of Medicine, Baltimore, MD 21205, USA. ¹⁵Institute for Systems Genetics and Department of Biochemistry and Molecular Pharmacology, NYU Langone Health, New York, NY 10016, USA. ¹⁶Department of Biomedical Engineering, NYU Tandon School of Engineering, Brooklyn, NY 11201, USA. ¹⁷Department of Microbiology, Icahn School of Medicine at Mount Sinai, New York, NY 10029, USA.

*These authors contributed equally to this work. †Present address: Department of Molecular Machines and Signaling, Max Planck Institute of Biochemistry, 82152 Martinsried, Germany. ‡Present address: Department of Molecular Structural Biology, Max Planck Institute of Biochemistry, 82152 Martinsried, Germany. §Present address: Verily Life Sciences, 269 E. Grand Ave., South San Francisco, CA 94080, USA. ¶Present address: Department of Bioengineering, Stanford University, Stanford, CA 94305, USA. #Deceased. **Present address: Institut Pasteur, Structural Bioinformatics Unit, Department of Structural Biology and Chemistry, CNRS UMR 3528, C3BI USR 3756 CNRS & IP, Paris, France. ††Corresponding author. Email: nevan.krogan@ucsf.edu (N.J.K.); sali@salilab.org (A.S.); jef.boeke@nyulangone.org (J.D.B.)

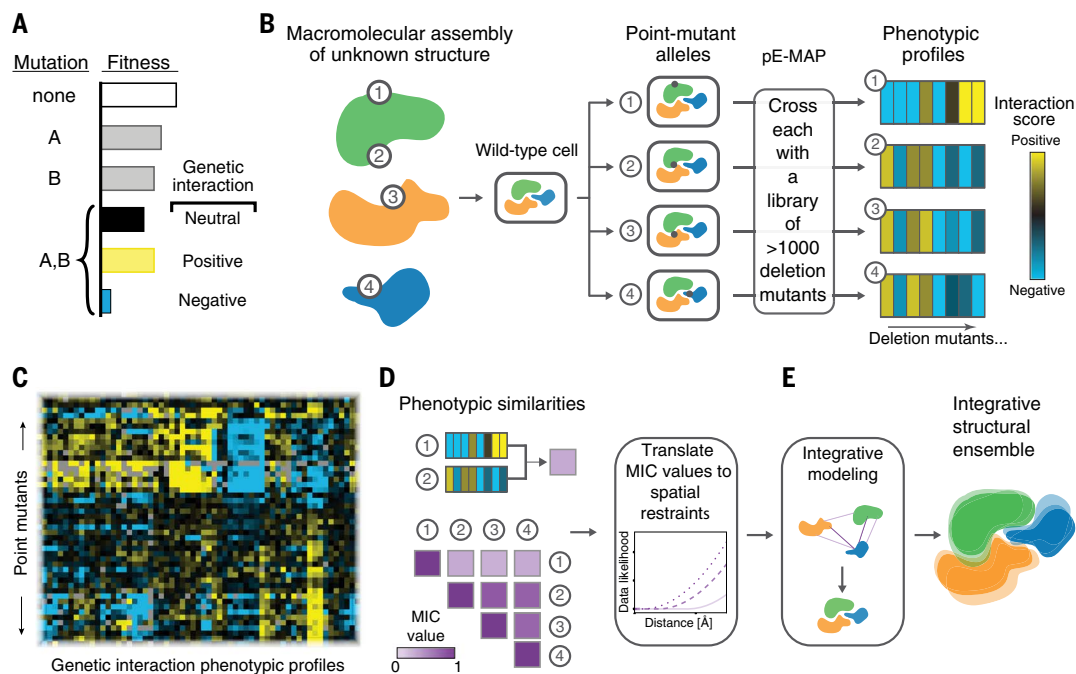


Fig. 1. Building spatial restraints from pairwise genetic perturbations.

(A) Genetic interactions arise when the combined fitness defect of a double mutant deviates from the expected multiplicative growth defect of the two single mutants. (B) The generation of a pE-MAP relies on a collection of point mutations, which is constructed by systematic mutagenesis of genes that encode the subunits of a macromolecular assembly (mutations labeled 1 to 4). The point-mutant strains are then crossed against a library of gene deletions, followed by fitness measurement and subsequent calculation of

genetic interaction scores to obtain the phenotypic profiles. (C) An example subset of a pE-MAP of point mutants crossed against a library of gene deletions. (D) Each pairwise combination of phenotypic profiles is transformed into a single MIC value that reflects the similarity between the two profiles. The MIC values are translated into spatial restraints for integrative modeling. (E) The MIC values and other input information are used for integrative structure modeling. An ensemble of structures that satisfy the input information is obtained.

is controlled by histone posttranslational modifications (PTMs) (21)—including acetylation, methylation, phosphorylation, and ubiquitination—that help maintain and regulate chromatin structure and transcription. Our library of point mutations in the core histones H3 and H4 was designed to comprise a comprehensive alanine scan, as well as context-specific mutations of modifiable residues (e.g., lysine and arginine), such as charge removal or reversal and substitutions mimicking PTMs (22, 23). Partial deletions of the N-terminal tails of H3 and H4 were also included, because these regions play important and sometimes redundant roles in chromatin biology (24, 25). In budding yeast, histones H3 and H4 are expressed from two loci each, *HHT1/HHT2* and *HHF1/HHF2*, respectively. To ensure preservation of the native expression levels, we engineered each strain to include identical point mutations in both relevant loci with separate selection markers (*HYG^R* and *URA3*) (Fig. 2A). In total, we designed 479 histone mutants, of which 350 were amenable to pE-MAP analysis (Fig. 2, B to D, and table S1); the remaining 129 mutants either were lethal or exhibited very poor growth, rendering them inaccessible to genetic analysis (fig. S1). The histone mutants were crossed against a library of 1370 gene deletions and hypomorphic alleles (table S1) using

our triple-mutant selection strategy (26, 27) involving three different selectable markers (*HYG^R* and *URA3* to select for both copies of the histone alleles and *KAN^R* for the knockout library strains) (Fig. 2A and Methods) (26). Genetic interactions were quantified using the S-score (28), which measures the deviation of the double mutant fitness from the expected combined effect of the individual mutations (Methods). The pE-MAP screen was carried out in three biological replicates (Methods), which exhibit a high reproducibility (Fig. 2E), and the final S-scores (as depicted in Fig. 1C) are the averages of these replicates.

It has been shown that a pE-MAP can be used to predict protein-protein interactions (PPIs) by comparing the genetic interaction patterns between pairs of deletion mutants across all the point mutants (15, 29). On a global level, this is only possible if the point mutant set affects a broad group of processes and exhibits genetic interactions with the many different deletion mutants that encode the PPI proteins. Because the histone mutant collection perturbs only two proteins (H3 and H4), we set out to investigate whether the resulting phenotypic profiles are sufficient to predict PPIs among the 1370 deletion mutants. Using a receiver operating characteristic (ROC) curve, we find that the histone pE-MAP predicts PPIs

similarly to previous E-MAPs that affect more genes (15, 29) (Fig. 2F and Methods). This finding indicates that the combined set of histone point mutants affects a broad set of cellular processes, reflecting the multifunctional nature of histones H3 and H4 and their central role in controlling the global genetic environment of cells.

To gain insight into the regulatory hierarchy that drives the widespread functional effects of histone perturbations, we set out to examine the relationship between genetic interactions and gene expression changes. To this end, we determined the genome-wide gene expression levels for 29 representative histone mutants using RNA sequencing (RNA-seq) and found no correlation between the expression change of a gene resulting from a given histone mutation and the corresponding S-score (Fig. 2G and table S2). This indicates that observed genetic interactions between histone mutations and deletion mutants are due to complex regulatory patterns, rather than the histone mutation directly modulating the expression of the interacting gene.

The pE-MAP was clustered hierarchically along both dimensions (Fig. 3A and data S1) and effectively recapitulates known protein complex and pathway memberships. For example, the pE-MAP identified COMPASS (30, 31),

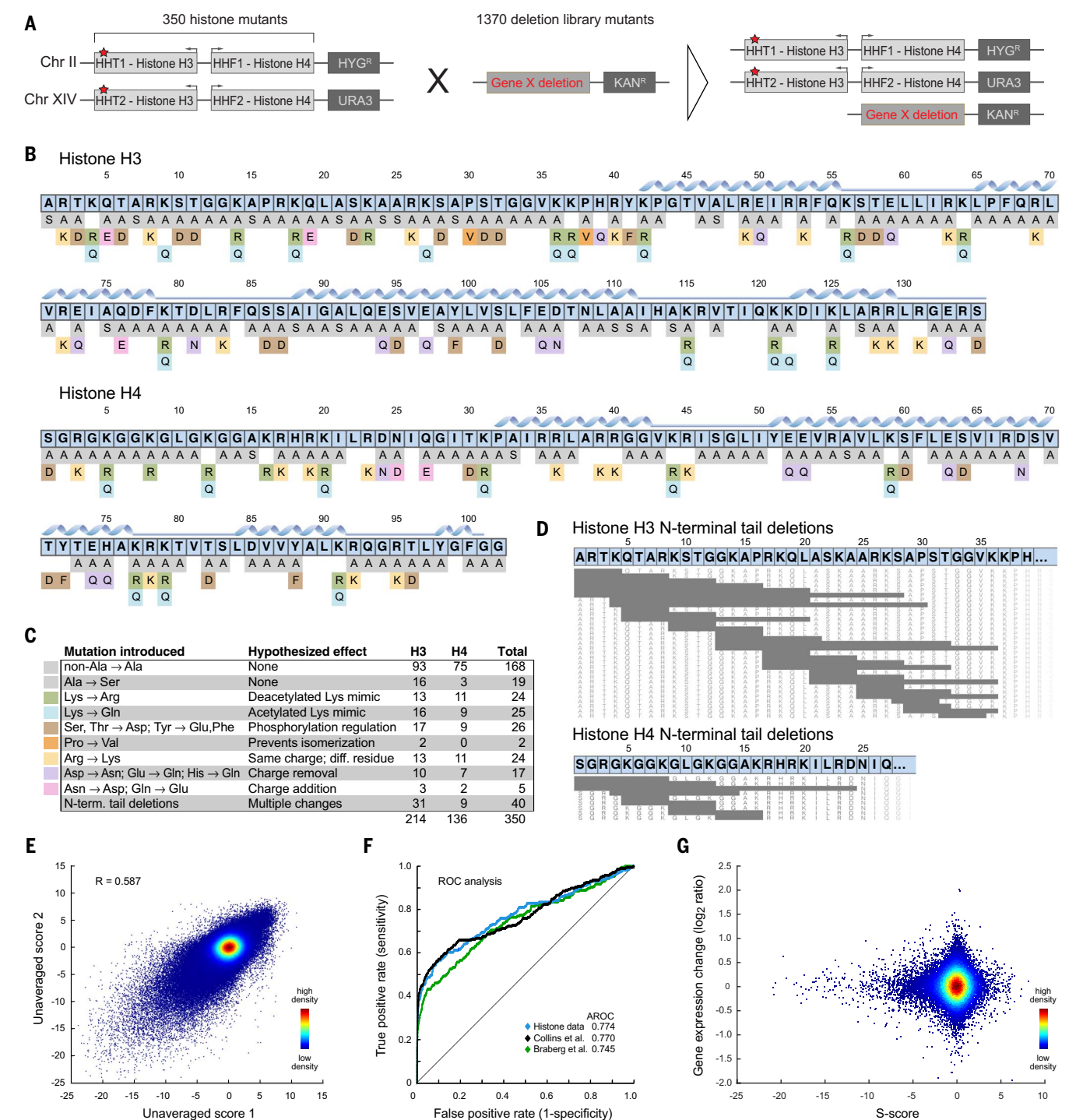


Fig. 2. Genetic interrogation of histones H3 and H4 at a residue-level resolution. (A) Each histone mutant strain was modified at both native loci (*HHT1* and *HHT2* for H3 or *HHF1* and *HHF2* for H4, red stars) and crossed against a library of 1370 different deletion mutants (or hypomorphic alleles for essential genes). Chr, chromosome. (B) Schematic of the histone point mutants analyzed in this study (table S1). Secondary structure elements are indicated as ribbons above the amino acid sequence. The mutations are color coded according to the mutation introduced (C). Mutations resulting in inviable strains or strains too sick for genetic analysis are shown in fig. S1. (C) Table of histone mutant categories and their hypothesized effects [color coding as in (B)]. (D) Overview of viable H3 and H4 tail deletion mutants amenable to pE-MAP analysis. The amino acid sequences of the WT alleles are shown on top (residues 1 to 39 of histone H3 and 1 to 27 of histone H4). Gray

bars represent the deleted residues in H3 and H4. (E) Reproducibility of histone pE-MAP S-scores between biological replicates. Plotted are all S-score pairs among the biological replicas, which include triplicate measurements for 346 histone alleles and duplicates of four alleles (H4E73Q, H4H18A, H4I21A, and H4K44Q). (F) ROC curves showing the power to predict physical interactions between pairs of proteins from this pE-MAP (blue) as well as a previously published pE-MAP [green, (15)] and E-MAP [black, (29)] data. (G) Relationship between gene expression (log₂ fold change over WT) and S-scores of 29 H3 and H4 alleles (table S2). Data from all 1256 deletion library mutants that were measured in both RNA-seq expression and pE-MAP analysis are plotted. Single-letter abbreviations for the amino acid residues are as follows: A, Ala; C, Cys; D, Asp; E, Glu; F, Phe; G, Gly; H, His; I, Ile; K, Lys; L, Leu; M, Met; N, Asn; P, Pro; Q, Gln; R, Arg; S, Ser; T, Thr; V, Val; W, Trp; and Y, Tyr.

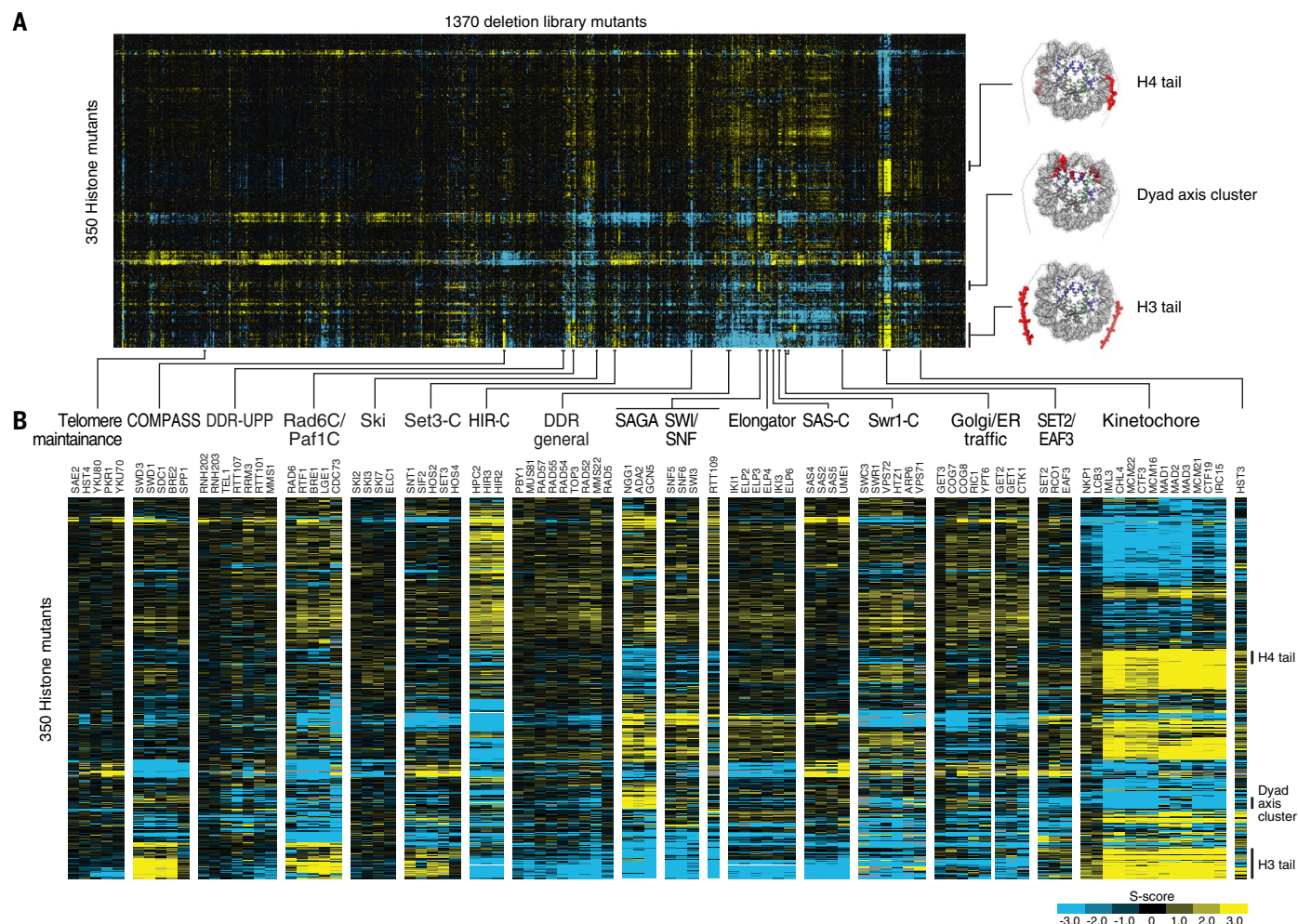


Fig. 3. The genetic interaction landscape of histones H3 and H4.

(A) Hierarchically clustered pE-MAP of 350 histone H3 and H4 alleles screened against a library of 1370 deletion mutants or hypomorphic alleles. The pE-MAP consists of more than 479,000 genetic interactions. Positive (suppressive or epistatic) and negative (synthetic sick) genetic interactions are colored in yellow or blue, respectively. Examples of histone alleles with similar genetic interaction profiles are highlighted on the right side in the context of the nucleosome structure. The nucleosome structure is modified

from PDB 1ID3 (data S2), with H3 in purple, H4 in green, and mutated or deleted residues highlighted in red. N-terminal tail residues of H3 and H4 not included in PDB 1ID3 are visualized as strings on the periphery. (B) Examples of genetic interaction profiles of gene clusters belonging to known protein complexes or biological pathways are highlighted and their genetic interaction profiles enlarged from (A). DDR, DNA damage or repair; UPP, ubiquitin proteasome pathway; SAGA, Spt-Ada-Gcn5 acetyltransferase; SWI/SNF, Switch/sucrose non-fermentable.

Swr1-C (32), and the Set2-Eaf3 pathway (33–35), as well as clusters of genes linked to telomere maintenance and Golgi-ER traffic (Fig. 3B and data S1). Furthermore, mutations of histone residues in close proximity to each other (e.g., mutants of the H3 or H4 N-terminal tails) tend to show similar phenotypic profiles (Fig. 3 and fig. S2A). Overall, we find that histone tail deletion mutants give rise to stronger phenotypic profiles than the point mutants (fig. S2B), reflecting the multiple residue perturbations and the importance of functional histone tails for cell homeostasis.

Phenotypic profile similarities are correlated with structural proximity

Similarities between pairs of phenotypic profiles in the histone H3-H4 pE-MAP were quan-

tified by the maximal information coefficient (MIC) (36, 37) (Fig. 1D, fig. S3, and Methods). The MIC values between pairs of phenotypic profiles do not linearly correlate with the distances between the mutated residues in the WT structure (Pearson correlation coefficient of -0.07 , C α -C α distances) but are informative about an upper distance bound between the residues (Fig. 4A and fig. S3C). The upper distance bound was obtained by binning the MIC values into 20 intervals and selecting the maximum distance spanned by any pair of residues in each bin, followed by fitting a logarithmic decay function to these maximum distances (Fig. 4A, fig. S3C, and Methods). The data show that a pair of proximal point mutations are more likely to have a high MIC value than a pair of distal point mutations. How-

ever, not all proximal mutations have a high MIC value: Most pairs of phenotypic profiles, even those for residues that are less than 16 Å apart, are highly dissimilar (94% of all pairs exhibit a MIC value <0.3). These observations justify converting the pE-MAP data into a Bayesian data likelihood that provides an upper bound on the distance spanned by the mutated residues (Fig. 4B and Methods). This Bayesian term objectively interprets the noise in the experimental data and allows us to quantify the uncertainty of the resulting structural models. The complete scoring function for evaluating any structural model also includes simple terms accounting for excluded volume and sequence connectivity, in addition to the Bayesian terms for all pairs of profiles in the pE-MAP with a MIC value above 0.3 (Fig. 5).

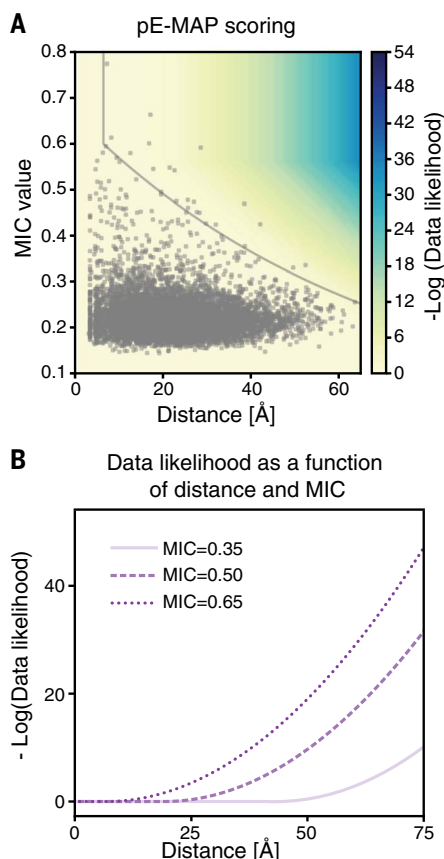


Fig. 4. Generation of the scoring function.

(A) Relationship between pairwise distances and MIC values. The solid gray line represents the logarithmic decay fit to the upper distance bounds (Methods, Eq.1). The background color gradient reflects how the data likelihood depends on MIC value and distance. (B) -Log of the data likelihood as a function of distance for different MIC values (Methods).

Spatial restraints derived from pE-MAP data can be used for integrative structure determination

An ensemble of the H3-H4 dimer configurations that satisfy the input information (i.e., the model) was found by exhaustive Monte Carlo sampling guided by the scoring function, starting with random initial configurations of the rigid comparative models of the H3 and H4 subunits (Fig. 5 and Methods). The resulting ensemble is accurate and precise, as demonstrated by the similarity between the x-ray structure [Protein Data Bank (PDB) 1ID3, (38)] and model contact maps (Fig. 6, A and B). Specifically, the mean accuracy is 3.8 Å (Fig. 6C); the accuracy is defined as the average Ca root-mean-square deviation (RMSD) between the x-ray structure and each of the structures in the ensemble. The precision is 1.0 Å (Fig. 6C), which is defined as the average RMSD between all solutions in the ensemble. As a control, we also computed a model from randomly shuffled MIC values. The resulting model (Fig. 6D) is incorrect (mean accuracy of

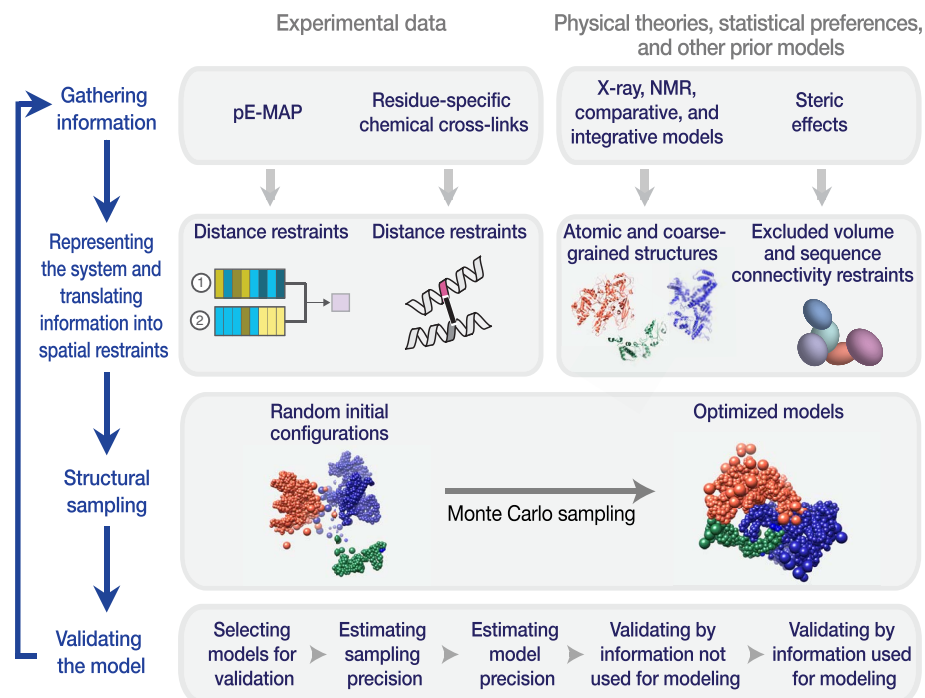


Fig. 5. Description of the integrative modeling workflow. The four stages include (i) gathering all available experimental data and prior information, (ii) translating all information into a representation of the assembly components and a scoring function for ranking alternative assembly structures, (iii) sampling structural models, and (iv) validating the model. In this example, the representation of the components of a complex is based on comparative models of its components. The scoring function consists of spatial restraints that are obtained from pE-MAP and/or cross-linking experiments (evolutionary coupling analysis is not indicated in this scheme) as well as excluded volume and sequence connectivity restraints. The sampling explores the configurations of rigid components, searching for those assembly structures that satisfy the spatial restraints as well as possible. The goal is to obtain an ensemble of structures that satisfy the input data within the uncertainty of the data used to compute them. The sampling precision is estimated and models are clustered and evaluated by the degree to which they satisfy the input information used to construct them as well as omitted information. The protocol can iterate through the four stages until the models are judged to be satisfactory, most often based on their precision and the degree to which they satisfy the data.

15.8 Å and incorrect contact map; Fig. 6, C and D) and imprecise (7.6 Å; Fig. 6C). As another control, we computed a model by a state-of-the-art protein-protein docking method (39), resulting in a model with an inferior accuracy of 6.9 Å (fig. S4). Finally, we also mapped the accuracy and precision of the model as a function of the fraction of the pE-MAP data used (Methods). As expected, the more pE-MAP data that are used, the more accurate and precise is the model (Fig. 6C).

To compute the structure of a protein complex for which the structures of the components are known, we estimate that 35 to 40 mutations per component are necessary to generate a complex model with precision sufficient to map the positions and relative orientations of the components (fig. S5 and Methods). What is a useful model precision depends on the questions asked (40). Fortunately, many questions can often be answered by models as precise as those obtained based on pE-MAP data (RMSD range of 1 to 15 Å). Some examples include describing the architecture and

evolution of protein assemblies (8, 41), designing interface mutations (42), characterizing structural heterogeneity of protein complexes (42, 43), and mapping binding-induced structural changes (44). Importantly, the estimate of 35 to 40 mutations is an upper bound, and, in many cases, the number might be reduced by specifically exploiting point mutations that target surface residues and/or residues known to be functionally important and by choosing substitutions likely to give rise to functional perturbations. The outcome of these calculations indicates the utility of the pE-MAP data for integrative structure determination.

The pE-MAP connects individual histone residues and regions to other associated complexes and processes

To examine whether the pE-MAP can identify interactions with complexes that are not stably associated with histones, we investigated the relationships between modifiable histone residues and their cognate enzymes (modifier pairs). Interestingly, we observed a dramatic

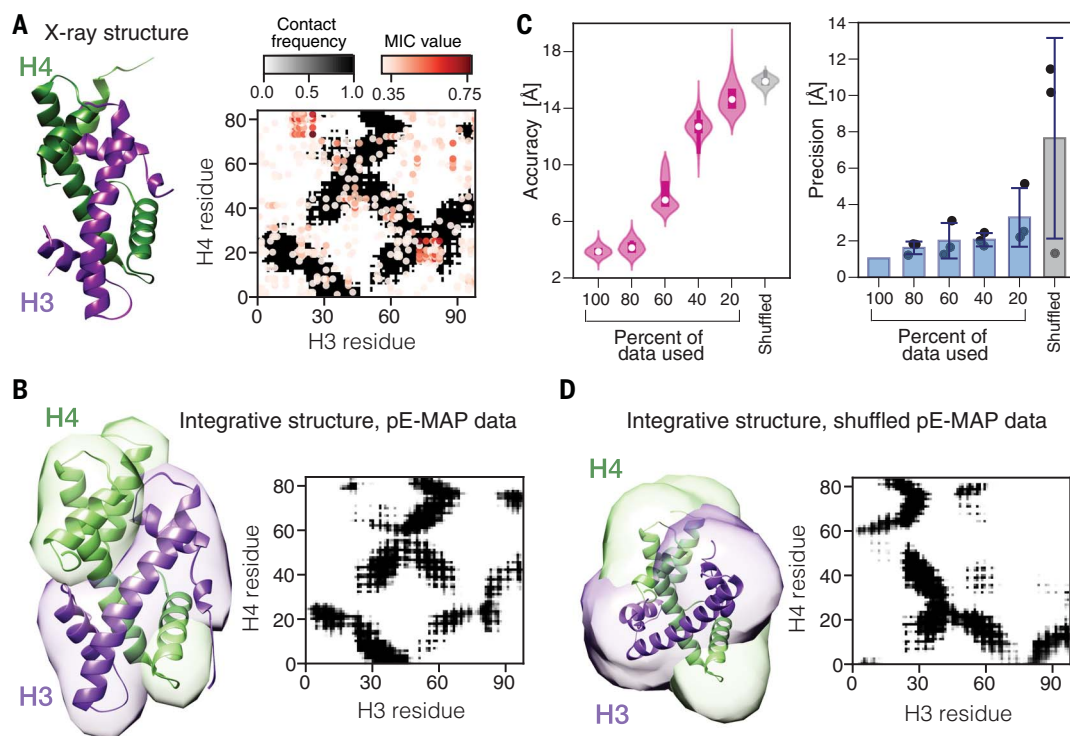


Fig. 6. Integrative structure determination of histones H3 and H4. (A) The native structure of the histone H3-H4 dimer (PDB 1ID3, left) and its contact map (right). In contact maps, the intensity of gray is proportional to the relative frequency of residue-residue contacts in the models (cutoff distance of 12 Å). For x-ray structures, the contact frequency is either 0 (white) or 1 (black). The circles correspond to the pairs of restrained residues, with the intensity of red proportional to the MIC value (MIC > 0.3), showing that the pairs of residues with high MIC values are distributed throughout the proteins. (B) The localization probability density of the ensemble of structures is shown with a representative (centroid) structure from the computed ensemble embedded within it (left)

and the corresponding contact map (right). The localization probability density map represents the probability of any volume element being occupied by a given protein. (C) Distributions of accuracy (left) of structures in the ensembles and model precisions (right) based on the full pE-MAP dataset, resampled datasets that consider fractions of the data, and using shuffled MIC values. On the left, the white dots represent median accuracies, and on the right, the error bars represent the standard deviations of model precision over three independent realizations (shown as dots). (D) Localization probability density and centroid structure (left) and contact map (right) computed with shuffled MIC values (Methods).

increase in S-scores within specific modifier pairs, as compared with the overall genetic interaction distribution (Fig. 7A and table S4). The positive S-scores reflect that a modifier and its target residue often are epistatic or suppressive because they function in the same pathway. To test if this pattern extends to phenotypic profile similarities, we integrated the histone pE-MAP into a merged map of previously collected genetic interaction data for gene deletions and hypomorphic alleles (45). We computed Pearson correlation coefficients for each histone mutant phenotypic profile across the merged map, generating a correlation map of 350 histone mutants against 4414 whole-gene perturbations (data S3 and Methods). In agreement with the individual S-scores, the specific modifier pairs exhibit significantly higher phenotypic profile correlations than the overall map (Fig. 7A and table S4). These findings show that the pE-MAP can be used to pair specific residues to their respective modifiers, even though these are not stably associated with the histones. For example, when components of COMPASS, which methylates histone H3K4 (30, 31, 46),

are deleted (*swd1Δ*, *swd3Δ*, *sdclΔ*, *bre2Δ*), we observe strong positive S-scores with both H3K4 mutants [K4R (Lys⁴→Arg) and K4Q (Lys⁴→Gln)] as well as high correlations of the phenotypic profiles between the H3K4 mutants and COMPASS deletions (Fig. 7B and data S1 and S2).

To explore these relationships in a structural context, we developed a Cytoscape (47) app named stE-MAP (structure E-MAP) that interactively maps the genetic interactions of pE-MAP gene clusters onto the point-mutated protein structure. stE-MAP connects Cytoscape to ChimeraX (48) and displays connections between a predefined set of genes and all mutated residues for which the underlying interactions pass user-defined criteria (fig. S6A). We mapped the genetic connections between COMPASS and all histone residues with which it exhibits >0.2 median correlation (Methods). Only five residues pass this threshold, and the strongest connection is displayed by H3K4. The other four residues (H3K1 to H3K3 and H3K5) are proximal and are thus likely to interfere with the interaction between COMPASS and H3K4 (Fig. 7C). This finding is particularly

notable because these residues reside in the most distal region of the unstructured H3 N-terminal tail. Given that we do not have COMPASS point mutations in our dataset, we did not attempt to model this interaction. However, analysis of the MIC values associated with the H3 and H4 tails, and their relationship with the core domains, indicates that distance restraints for the histone tails could be derived from the pE-MAP data. Specifically, the MIC value distributions for the tail-core and tail-tail pairs of mutations are similar to that of the core-core mutations (fig. S5C). This similarity indicates that we can derive distance restraints for the histone tails, thus, in principle, supporting the feasibility of integrative structure modeling of disordered regions. In such modeling, to avoid overinterpretation of the data and to account for a possibility that the pE-MAP data of the histone tails reflect interactions with neighboring nucleosomes, we would have to include multiple nucleosome copies in the model and allow for assignment ambiguity, just as we do for distances inferred from chemical cross-links and protein proximity inferred from affinity copurification (49).

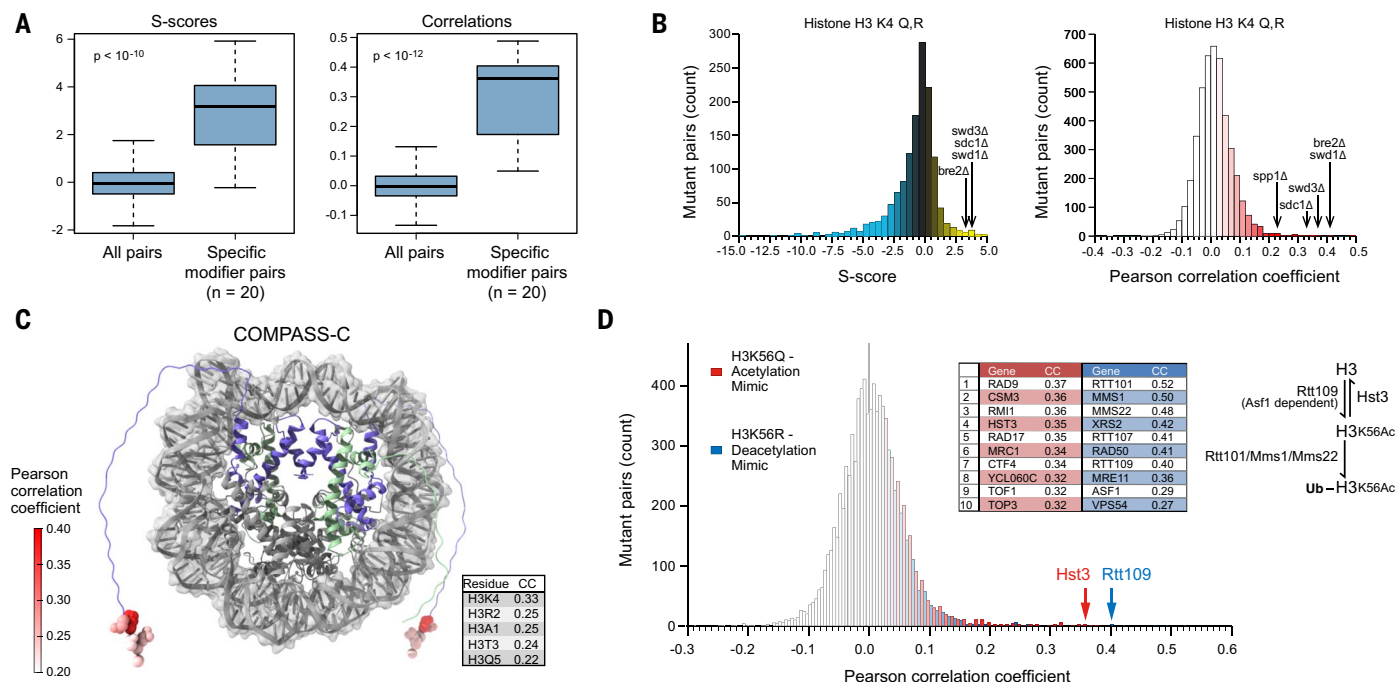


Fig. 7. Connecting individual histone residues and regions to other associated complexes. (A) Comparison of S-scores and Pearson correlation coefficients of phenotypic profiles of modifier-residue pairs to the overall data. Only residues with a single known modifier and modifiers with a single known target residue were included (table S4). p values were calculated using two-sided Wilcoxon rank sum tests. The whiskers of the boxplots extend to a maximum of $1.5 \times \text{IQR}$ (interquartile range), and outliers are not plotted. (B) Average distributions of S-scores (left) and phenotypic profile correlations (right) of H3K4 mutants (mean of H3K4Q and H3K4R). Members of the COMPASS complex that exhibit a mean S-score >2.5 or a mean genetic interaction profile correlation >0.2 with H3K4 mutants are highlighted. The COMPASS complex is responsible for H3K4 methylation. (C) Mapping of genetic interaction profile correlations to COMPASS complex members on the structure of the nucleosome (modified PDB 1ID3, data S2). N-terminal tail residues of H3 and H4 not included in 1ID3 are visualized as strings on the periphery. Only residues that exhibit a median genetic

profile correlation >0.2 with the COMPASS subunits are highlighted (Methods). H3 is depicted in purple, H4 in light green, and H2A/H2B and DNA in gray. The red color gradient reflects the strength of the correlation between each residue and the COMPASS members, calculated as the median correlation between the residue's tested mutations and the COMPASS members. (D) Distributions of genetic interaction profile correlations of H3K56Q (acetylation mimic) and H3K56R (deacetylation mimic). Correlations of key H3K56ac-level regulators, Rtt109 (acetylating) and Hst3 (deacetylating), are highlighted. The cartoon outlines the H3K56 acetylation pathway and its role in H3 ubiquitylation. Rtt109 acetylates H3K56 through an Asf1-dependent mechanism, which promotes ubiquitylation of H3 by Rtt101-Mms1 and Mms22. These five gene deletions are all found among the top 10 most similar to the deacetylation mimic H3K56R, whereas deletion of the H3K56 deacetylase Hst3 instead gives rise to a profile similar to the acetylation mimic H3K56Q (table inset). CC, Pearson correlation coefficient.

We observe similar trends for other histone modifiers. For example, members of the Set2 pathway (Set2, Eaf3, Rco1, Ctk1) (33–35) rank highly in the distributions of S-scores or correlations for mutations of their target residue, H3K36 (fig. S6, B and C). Interestingly, we also found instances where different mutations of a single residue identify connections to different modifiers. For example, the phenotypic profile of the deacetylation mimic H3K56R is similar to that of deletion of *RTT109*, which encodes the H3K56 acetylase (Pearson correlation coefficient of 0.4) (29), whereas the acetylation mimic H3K56Q instead correlates with the profile generated from the deletion of the corresponding deacetylase, *HST3* (Pearson correlation coefficient of 0.35) (50) (Fig. 7D). H3K56R further correlates with *asf1Δ*, *rtt101Δ*, *mms1Δ*, and *mms22Δ*, whose corresponding proteins play key roles in the H3K56 acetylation pathway and downstream H3 ubiquitylation (51, 52) (Fig. 7D). Accordingly, the stE-MAP app identified strong links between Hst3-Hst4

and H3K56Q, as well as Rtt109-Asf1 and H3K56R (fig. S6D). Although we find that it is often informative to group different mutations of the same residue together, these examples highlight the potential of these maps for deeper mechanistic insights where required.

Expanding on these findings, we built a gene set enrichment map connecting the modifiable histone residues to nuclear processes (fig. S7A, table S5, and Methods). We observe both known and previously unknown connections. For example, “DNA recombination and repair” is connected to four residues, and two of these, H3K56 and H3K79, have been shown to play key roles in yeast’s DNA repair (53–57). Interestingly, we find that mutations of the other two residues (H3R63K and H4R36K) result in increased spontaneous mutation frequencies at the *URA3* locus, indicating that these residues also function in DNA repair (fig. S7, B and C; table S6; and Methods).

The gene set enrichment analysis also identified 13 residues connected to cryptic

transcription (fig. S7A). The pE-MAP includes 24 different mutations of these residues, and we tested their involvement in cryptic transcription by quantifying the abundance of transcripts at the 5′ and 3′ end of the *STE11* gene, using quantitative polymerase chain reaction (qPCR) (fig. S7D and Methods). In total, 16 mutations, distributed among 10 residues, increase 3′ transcript abundance by $>50\%$ compared with WT (table S7), and nine mutants among five residues increase 3′ transcription more than twofold, without major changes in 5′ transcription (fig. S7E). As expected, H3K36A, H3K36R, H3K36Q, and *set2Δ* increase 3′ transcript abundance strongly, as do mutations of H4K44, which is a residue known to affect cryptic transcription (58). Interestingly, H3K122A increases 3′-transcript abundance >15 -fold and, using ATAC-seq, we find that the mutation gives rise to nucleosome-free regions in *STE11* and other genes known to produce cryptic transcripts (fig. S7, F to H). H3K122A exhibits positive genetic interactions with deletion of the

histone chaperone *SPT2* (S-score = 2.4) and the nucleosome remodeling factor *CHD1* (S-score = 4.9), which are both involved in cryptic transcription (59, 60). Accordingly, we find that deletion of either *SPT2* or *CHD1* suppresses the cryptic transcription phenotype observed in H3K122A to WT levels, even though *spt2Δ* or *chd1Δ* alone has no effect (fig. S7, I to K).

The integrative structure determination approach is transferable to other complexes

To test whether genetic interaction mapping can be used to determine the structure of other complexes, we examined a pE-MAP of RNAPII in budding yeast (15). This pE-MAP consists of 53 point mutants crossed against a library of 1200 gene deletions and hypomorphic alleles. Interestingly, the association between MIC values and the upper distance bound is also apparent in this dataset (fig. S8, A and B), even though the protein sizes and mutational coverage of the polymerase system (up to ~1700 residues and 1 to 2%, respectively) are vastly different from those of the histones (<140 residues and 85 to 90%). These observations suggest that our parameterization of the pE-MAP spatial restraint based on the histone data may be generally applicable. To evaluate this expectation directly, we next modeled subunits Rpb1 and Rpb2 of RNAPII using the Bayesian likelihood parametrization based on the histone pE-MAP. To illustrate the modeling of higher-order complexes, we divided Rpb1 into two domains, thereby representing the system with three rigid bodies (Methods and table S8). We obtained a model with a mean accuracy of 16.8 Å and precision of 9.8 Å (Fig. 8, A to D, and table S8). This positive result illustrates the generality of the pE-MAP-based spatial restraints.

To further assess the utility of pE-MAP data for structure determination, we compared the RNAPII model obtained using the pE-MAP to a model using 22 previously published chemical cross-links (67). Cross-linking is widely used for integrative structure determination of macromolecular assemblies (2, 8). Interestingly, a model of yeast RNAPII based on the pE-MAP data is as accurate as that based on the cross-links (16.8 and 16.7 Å, respectively; Fig. 8D). Moreover, the accuracy and precision of the model improves if both datasets are used simultaneously (10.2 and 3.7 Å, respectively; Fig. 8D), indicating complementarity between the two types of data and demonstrating a premise of integrative structure determination (Fig. 5). Although a cross-link between two residues may provide more direct structural information than the corresponding pE-MAP pair, the number of possible cross-links is limited by the number of proximal reactive residue pairs, whereas the number of pE-MAP pairs grows quadratically with every additional point mutation introduced. There-

fore, the larger number of less precise pE-MAP restraints can lead to a more accurate model than a smaller number of more precise cross-links.

The integrative structure determination approach is transferable to other types of phenotypic profiles

To examine the applicability of our approach to other types of phenotypic profiles, we turned to a CG-MAP of 44 bacterial RNAP point mutations exposed to 83 different environmental stresses (e.g., chemical perturbations, temperature stress, and pH change) (20). We observe an association between MIC values and the upper distance bound, similar to that of the pE-MAP datasets (fig. S8 and Methods). We modeled the structure of subunits RpoB and RpoC of the bacterial RNAP with a mean accuracy of 15.0 Å and precision of 6.6 Å (Fig. 8, E to H, and table S9). This result suggests that maps with relatively small numbers of orthogonal phenotypes per point mutation can be used to accurately predict the architecture of macromolecular assemblies. Considering that constructing large gene deletion libraries and crossing them against point mutations can be laborious, environmental phenotypic profiles may be a more efficient alternative for generating spatial restraints for integrative structure determination than genetic interaction phenotypic profiles.

Spatial restraints derived from pE-MAP data are comparable to other commonly used data types

Coevolution information can also be used to predict the structure of protein assemblies (18, 62, 63). However, the success of such modeling is heavily dependent on the number of sequences in the input sequence alignments and the ability to discriminate interacting from noninteracting homologs in genomes with multiple paralogs (64). Using the RaptorX protein complex contact prediction server (65, 66), we predicted the interfacial contacts between RpoB and RpoC of the bacterial RNAP; the numbers of homologous sequences were insufficient for the yeast histones and RNAPII (Methods). Importantly, RaptorX is based on a combination of coevolution analysis and a deep-learning algorithm that reduces the requirement for sequence homologs and improves accuracy (67). Other commonly used coevolution methods (18, 62) did not identify any interfacial contacts. Similar to the pE-MAP and CG-MAP datasets, we observe a negative statistical association between the residue pair coupling strengths and the upper distance bounds (fig. S9). To mimic the pE-MAP restraint, we converted the top coupling strengths into upper distance bound restraints (Methods). The model ensemble computed from coevolution-derived restraints includes two different sets of configurations (mean accuracy of 22.6 Å; model precision of 9.0 Å; Fig. 8H). Only a frac-

tion of the bacterial RNAP structures computed using coevolution-derived restraints are as accurate as those computed using the CG-MAP restraints. The model precision and accuracy of the model improve slightly if both types of restraints are combined (mean accuracy of 14.5 Å; model precision of 6.5 Å; Fig. 8H).

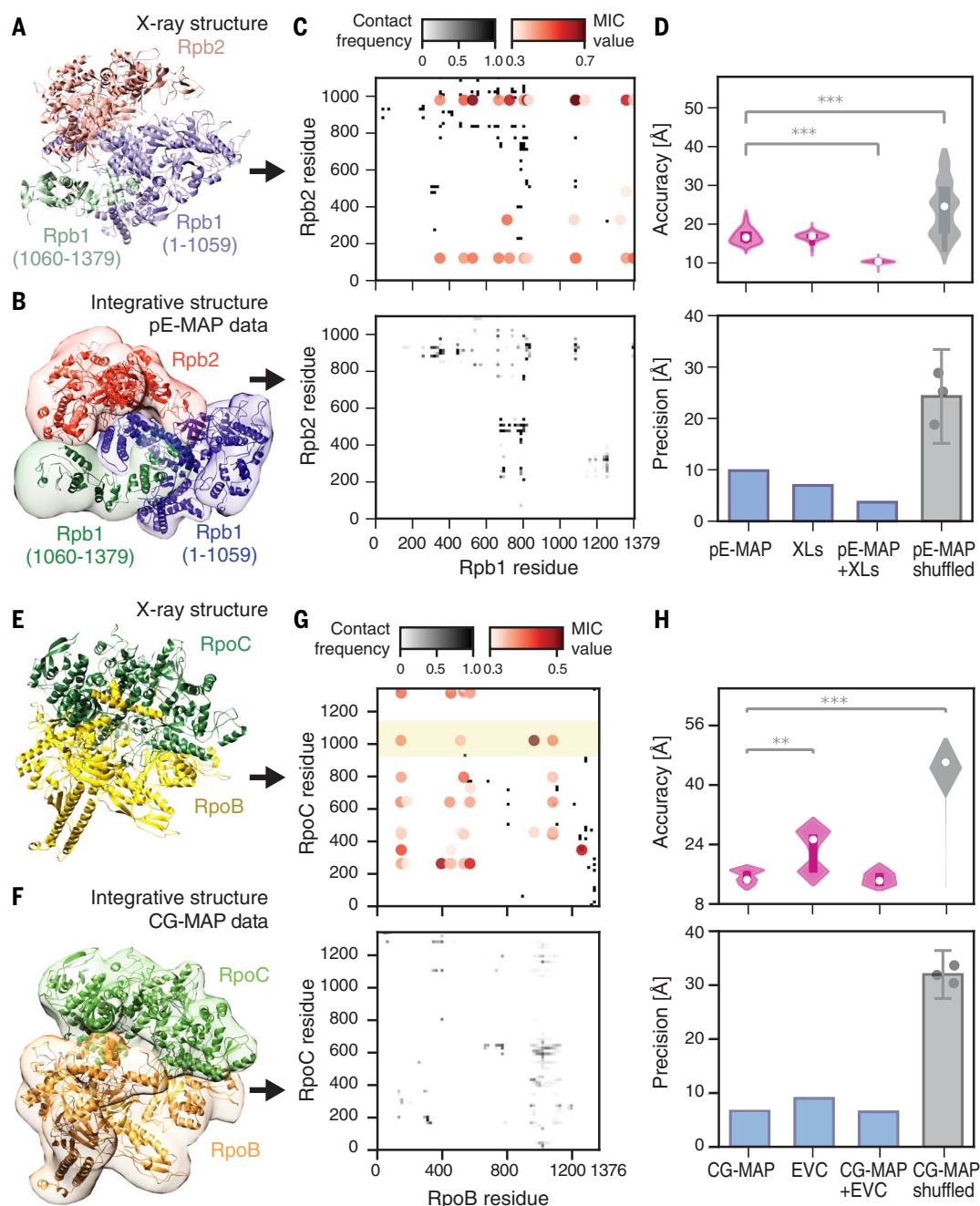
Discussion

We show that the architectures of macromolecular assemblies can be determined using quantitative genetic interaction data collected in vivo. The accuracy and precision of such models are comparable to those of models based on chemical cross-linking or coevolution analysis. A key premise of integrative modeling is that using several different types of data improves the accuracy and precision of the model. Because the pE-MAPs and CG-MAPs contain purely phenotypic measurements, collected in living cells, these datasets generate spatial restraints that are orthogonal to other commonly used data for integrative modeling. Because these data reflect in vivo structures, and are thus unlikely to share artifacts of biophysical methods, they could be of particularly high value in the integrative modeling process. The genetic interaction data may also allow for the characterization of complexes that are difficult to isolate and purify or those that are only transiently stable. Importantly, the equipment required for generating these data is basic, and, in particular, the CG-MAPs can be generated efficiently. Recent developments in CRISPR-Cas9 based approaches have paved the way for multiplexed precision genome editing in yeast (68), allowing for rapid generation of CG-MAPs. Together, these methods make feasible the proteome-wide modeling of protein complex structures, guided by global protein-protein interaction maps (69). In addition to proteins, the approach is also applicable to assemblies containing nucleic acids, thus further expanding the scope of integrative structural biology. pE-MAPs and CG-MAPs are complementary to other high-throughput functional assays. For example, two recent studies have shown that deep mutational scans can be used for determining the fold of a small protein domain (70, 71). If these methods prove useful for multidomain proteins or disordered regions (72), then the measured phenotype changes could be used to derive additional restraints for the integrative structure determination.

The relationship between phenotypic pE-MAP measurements and structure can be uncertain. The reasons for this include mutations in distant positions that are part of an allosteric network and could give rise to similar profiles, mutations that are functionally irrelevant, and mutations that perturb gene expression, mRNA stability, or translation. Additionally, the approach relies on the introduction of point mutations

Fig. 8. Integrative structure determination of yeast RNAPII and bacterial RNAP.

(A) The native structure of Rpb1-Rpb2 (PDB 2E2H) showing its three rigid-body components. Rpb1 was split into two domains, as shown. **(B)** The localization probability density of the ensemble of the three rigid-body structures is shown with a representative (centroid) structure from the computed ensemble embedded within it. **(C)** Contact maps computed for the x-ray structure (top) and model using the pE-MAP dataset (bottom). The circles correspond to the pairs of restrained residues, with the intensity of red proportional to the MIC value ($\text{MIC} > 0.3$). **(D)** Distributions of accuracy (top) for all structures in the ensemble and model precisions (bottom) for the computed ensembles based on pE-MAP and cross-link data. The white dots represent median accuracies. Error bars represent the standard deviations of model precisions over three independent realizations (shown as dots); $***p < 10^{-12}$. **(E)** Structure of subunits RpoB and RpoC from bacterial RNAP (PDB 4YG2). **(F)** The localization probability density of the ensemble of the RpoB-RpoC structures with a representative (centroid) structure from the computed ensemble embedded within it. **(G)** Contact maps computed for the x-ray structure (top) and model using the CG-MAP dataset (bottom). The shaded yellow band represents a region missing in the x-ray structure. **(H)** Distributions of accuracy (top) for all structures in the ensemble and model precisions (bottom) for the ensembles based on CG-MAP and evolutionary coupling (EVC) data. The white dots represent median accuracies. The error bars represent the standard deviations of model precision over three independent realizations (shown as dots). $**p < 10^{-6}$; $***p < 10^{-12}$.



into the proteins of interest, which may result in structural changes. However, proteins often adapt to mutations by small local changes in their structure, maintaining their overall fold and function (73). Mutations that cause major misfolding of essential proteins and/or assemblies are uncommon in pE-MAPs because the resulting fitness defects typically prevent successful screening. The method could be improved by specifically designing point mutants that do not alter the structure and/or lead to aggregation, by selecting commonly allowed

mutations as determined by divergent protein sequence alignments (74).

The aim of integrative structure determination is to model the structures of macromolecular assemblies. This often requires the structures of the individual components [from x-ray crystallography, NMR, cryo-EM, comparative modeling, or, increasingly, ab initio structure prediction (65, 67, 75, 76)]. The quality of the structures of the individual components and input data are crucial for integrative (or indeed any other) structural approaches,

and one cannot achieve a precise structure from low-quality starting structures or data. Even so, there are numerous examples of utility of structural models at lower resolution (3). For example, these models can be used to explain the architectural principles of large assemblies (8, 41, 77, 78), describe the structural dynamics of protein complexes (42, 43), or rationalize the impact of many mutations (41). A lower-resolution structure is also often a useful starting point for higher-resolution structure characterization.

CRISPR-Cas9 genome editing (79) has proven highly effective for high-throughput genetic interaction mapping in mammalian cells (80, 81). To date, these efforts have relied on whole-gene perturbations, but methods for systematic generation of point mutants using CRISPR-Cas9 have recently been developed (82, 83), paving the way for mammalian pE-MAP screening. This advance provides a means for integrative structure determination of assemblies in human cells and also allows for identification and characterization of functionally relevant structural changes that take place in disease alleles. Expanding this analysis to host-pathogen complexes (84–86) will be feasible by introducing specific mutations into the pathogenic genome and studying the phenotypic consequences using genetic interaction profiling of relevant host genes (87). Furthermore, several efforts are under way to generate multi-scale models of entire cells (88–92). In such instances, high-throughput genetic interaction mapping could provide global insights into cellular organization and dynamics of different components while also informing on the structures of individual assemblies.

Methods

Histone mutant strain construction

The histone H3 and H4 mutant strain library was constructed essentially as described (22, 23). Briefly, the mutants (tail deletions, complete alanine-scan, and context specific point mutations) were generated in the YMS196 background (MAT α *his3 Δ leu2 Δ ura3 Δ can1::STE2pr-spHIS5 lyp1::STE3pr-LEU2*) (table S1). First, the base strains were created by replacing the *HHT2-HHF2* locus with a *URA3*-containing cassette carrying a mutated *HHT2-HHF2* locus with their endogenous promoters. We randomly picked a few base strains, and the mutated *HHT2-HHF2* loci were PCR amplified and validated by sequencing. Then, the *HHT1-HHF1* locus was replaced with a *HYG^R*-containing cassette carrying a mutated version of the *HHT1-HHF1* locus, resulting in pE-MAP-amenable strains (Mata *his3 Δ leu2 Δ ura3 Δ hht1-hhf1::HYG^R hht2-hhf2::URA3 can1::STE2pr-spHIS5 lyp1::STE3pr-LEU2*).

Histone mutant library validation

Libraries were validated in three steps: (i) Each mutant was constructed, transformed into bacteria, and sequenced; 100% sequence identity was required to pass quality control. (ii) After the correct integration of histone mutants in the *HHT2-HHF2* locus, 5 to 10 yeast base strains from each 96-well plate were randomly selected, and corresponding histone fragments were amplified and sequenced to ensure the identity of each mutant in the well and no cross-contamination during plasmid preparation and yeast transformation; 100% of these were correct. (iii) After obtaining the yeast library with

the second (*HYG^R*) cassette integrated, 5 to 10 yeast strains from each 96-well plate were also randomly selected. Both copies of histone mutants in each strain were amplified and sequenced to confirm the identity of mutations. All of them were correct.

pE-MAP analysis

Each of the histone H3 and H4 mutant strains was crossed with 1370 MAT α KAN^R marked deletion (nonessential genes) or DAmP (decreased abundance by mRNA perturbation; essential genes) strains by pinning on solid media as described (15). Sporulation was induced, and MAT α haploid spores were selected by replica plating onto media containing canavanine (selecting *can1 Δ* haploids) and S-AEC (selecting *lyp1 Δ* haploids) and lacking histidine (selecting MAT α spores). Triple-mutant haploids were isolated on media containing hygromycin (selecting *hht1-hhf1 mutant cassette*) and G418 (selecting KAN^R marked deletion or DAmP), and lacking uracil (selecting *hht2-hhf2 mutant cassette*). Finally, triple-mutant colony sizes were extracted using imaging software. The screen was carried out in three biological replicates with three technical replicates in each. Four mutants (H4E73Q, H4H18A, H4I21A, and H4K44Q) failed screening in one biological replicate, and the results for these are based on the two successful replicates. Detailed E-MAP experimental procedures are described in (26, 29, 93). Genetic interactions were quantified using S-scores (28), which are closely related to *t* values. The S-score quantifies the deviation of the double (or triple) mutants from the expected combined fitness effects of the individual mutants and incorporates the reproducibility between technical replicates. The published S-scores represent the average S-scores across biological replicates.

Design of the pE-MAP spatial restraints

The distance restraint based on pE-MAP data was designed using the 308 single-point mutants from the histone pE-MAP and the structure from the PDB 1ID3, as follows: (i) postprocessing of the genetic interaction phenotypic profiles, (ii) devising a phenotypic similarity metric between the phenotypic profiles, and (iii) designing spatial restraints for integrative structure modeling using the phenotypic similarity values and the known nucleosome x-ray structure. Next, we describe each of these three steps in turn.

(i) Postprocessing of the genetic interaction phenotypic profiles: All missing values in the pE-MAP were imputed as the mean of the S-scores between the corresponding deletion mutant and all histone point mutants. To increase the signal-to-noise ratio of the pE-MAP, gene deletion mutants that mostly exhibited weak genetic interactions with the histone mutants were filtered out. To this end, the gene deletion profiles were ranked in descend-

ing order based on the counts of their S-scores that fell in either the top 2.5% of positive S-scores or the bottom 5% of negative S-scores, from the complete point mutant pE-MAP (cutoffs calculated after imputation). The more stringent cutoff for positive S-scores was chosen to reflect the smaller dynamic range for positive genetic interactions compared to negative genetic interactions. Gene deletions with the same count were then ranked in descending order by the mean of the absolute values of their highest and lowest score (fig. S3A). The top fraction of the deletions, determined in step (iii) below, were retained for computing the histone point mutant phenotypic profile similarities (below, fig. S3).

(ii) Devising a phenotypic similarity metric between the phenotypic profiles: We computed the similarity between all pairs of histone phenotypic profiles using the maximal information coefficient (MIC; fig. S3B), with the MIC parameters *alpha* and *c* set to 0.6 and 15, respectively, as suggested (36, 37). Many positions in the histones were mutated to several different residue types, giving rise to several phenotypic profiles for each of these positions. As a result, more than one MIC value would often be computed for a single residue pair. In such cases, only the highest MIC value was retained.

(iii) Designing spatial restraints for integrative structure modeling: Using the histone x-ray structure [PDB 1ID3; (38)], we measured the C α -C α distance between all pairs of residues for which we computed a MIC phenotypic similarity score. The percentage of the top scoring phenotypic profiles [ranked by the genetic interaction scores; step (i)] retained for further analysis was determined as follows. We compared the statistical association of the distances between two mutated residues with their phenotypic similarity by selecting the top 10, 25, 50, and 100% of the ranked deletions (fig. S3C). Although MIC values between phenotypic profiles do not linearly correlate with the distances spanned by the mutated residues in the WT structure (Pearson correlation coefficient of -0.07 when using the top 25% or top 50% of deletions), the MIC values provide an upper distance bound between the residues. The upper distance bound was obtained by binning the MIC values into 20 intervals and selecting the maximum distance spanned by any pair of residues in each bin, followed by fitting a logarithmic decay function (d_U) to the upper distance bounds:

$$d_U(\text{MIC}) = \begin{cases} \frac{\log(\text{MIC}) - n}{k} & \text{if MIC} \leq 0.6 \\ 6.84 & \text{if MIC} > 0.6 \end{cases} \quad (1)$$

where *k* and *n* are -0.0147 and -0.41 , respectively (fig. S3C). We find that selecting the top 25 or 50% of the deletions had a comparable

association between the upper distance bounds and the computed MIC values. The association was determined by computing the R and p values of the Pearson correlation coefficient and association significance, respectively, for the log-transformed MIC values. In this work, we retained the top 25% of the ranked phenotypic profiles for computing the phenotypic profile similarities.

To effectively handle the uncertain relationship between the data and model, we use Bayesian inference for scoring alternative models by formulating spatial restraints as Bayesian data likelihoods (94). Formally, the posterior probability of model M given data D and prior information I is $p(M|D, I) \propto p(D|M, I) \cdot p(M|I)$. The model, M , consists of a structure X and unknown parameters Y , such as noise in the data. The prior $p(M|I)$ is the probability density of model M given I . The prior can in general reflect information such as statistical potentials or a molecular mechanics force field; here, we only used excluded volume and sequence connectivity. The likelihood function $p(D|M, I)$ is the probability density of observing data D given M and I . The pE-MAP data was used to compute phenotypic similarities (i.e., MIC values) that inform distances between mutated residues pairs i, j . The likelihood of the entire pE-MAP dataset is the product over the individual observations between residue pairs i, j : $p(D|M, I) = \prod_{i,j} N[d_{i,j} | f_{i,j}(X), \sigma_{i,j}]$, where $f_{i,j}(X)$ is a forward model that predicts the data point $d_{i,j}$ in D that would have been observed for structure X in an experiment without noise; $N[d_{i,j} | f_{i,j}(X), \sigma_{i,j}]$ is a noise model that quantifies the deviation between the predicted and observed data points.

We defined the forward model by inverting the relation between the upper distance bound and observed MIC values [$d_U(\text{MIC})$, Eq. 1]:

$$f_{i,j}(X) = \text{MIC}(d_{i,j}) = \begin{cases} \exp(k \cdot d_{i,j} + n) & \text{if } d_{i,j} \leq d_0 \\ 0.6 & \text{if } d_{i,j} > d_0 \end{cases} \quad (2)$$

where $d_0 = d_U(0.6)$. Our choice of a noise model is a lognormal distribution with a flat plateau for MIC values below the upper bound on the experimentally observed MIC values (MIC^{obs}):

$$P(\text{MIC}_{i,j}^{\text{obs}} | \text{MIC}_{i,j}, X, \sigma_{i,j}) = \begin{cases} \frac{1}{N} & \text{if } \text{MIC}_{i,j}^{\text{obs}} \geq \text{MIC}_{i,j} \\ \frac{1}{M \sqrt{2\pi\sigma_{i,j}^2}} \exp\left[-\frac{1}{2\sigma_{i,j}^2} \log^2\left(\frac{\text{MIC}_{i,j}^{\text{obs}}}{\text{MIC}_{i,j}}\right)\right] & \text{if } \text{MIC}_{i,j}^{\text{obs}} < \text{MIC}_{i,j} \end{cases} \quad (3)$$

Here, $\sigma_{i,j}$ are the noise parameters that can optionally be determined as part of the model, and N and M are normalization factors necessary to make the likelihood continuous. Log-normal noise models have previously been used to describe errors of inherently positive

quantities (95). For computational efficiency, we used a single σ value for all residue pairs. An uninformative Jeffrey's prior is applied to σ to represent a lack of information on the bounds and distribution of this parameter (96).

Finally, a Bayesian term in the scoring function is defined as the negative logarithm of the posterior probability density: $S(M) = -\log p(M|D, I)$. In the Bayesian view, the output model is in fact best equated to the posterior model density that specifies a distribution of alternative single models M with varying probability density, not a single model, although single representative or average models can always be proposed based on the posterior model density.

Calculation of similarity metrics for yeast RNAPII and bacterial RNAP datasets

Steps (i) and (ii) from "Design of the pE-MAP spatial restraints" were repeated for the yeast RNAPII and bacterial RNAP datasets to generate the similarity metrics (MIC values) for these two systems, with the following modifications:

For yeast RNAPII, before imputing missing values in the pE-MAP, any deletion mutants that exhibit missing values with more than 15% of the point mutants were filtered out. This step is part of our pipeline but had no effect on the histone pE-MAP (because this pE-MAP does not contain any deletion mutant with more than 15% values missing). The number of ranked deletion mutants retained at the end of pE-MAP postprocessing was chosen to be 25% of the number of deletions in the original unfiltered pE-MAP (in accordance with the histone pE-MAP processing).

For bacterial RNAP, owing to the very small number of perturbations in this dataset, all the perturbations (instead of the top 25%) were retained for computing point mutant phenotypic profile similarities. In addition, owing to differences in the experimental design for generating the yeast pE-MAPs and the bacterial RNAP CG-MAP, the bacterial RNAP MIC distribution had a ~2-fold higher median and greater spread than the other datasets. Correspondingly, the bacterial RNAP MIC distribution was normalized using linear scaling, decreasing its median to match that of the histone MIC distribution. Importantly, this step was based solely on the MIC distributions, without reliance on any structural information.

Integrative structure determination

Integrative structure determination for each system proceeded through the standard four stages (3, 4, 5, 8, 41, 97) (Fig. 5 and tables S3, S8, and S9): (i) gathering data, (ii) representing subunits and translating data into spatial restraints, (iii) configurational sampling to produce an ensemble of structures that satisfies the restraints, and (iv) analyzing and validating the ensemble structures and data. The integra-

tive structure modeling protocol [i.e., stages (ii), (iii), and (iv)] was scripted using the *Python Modeling Interface* (PMI) package, a library for modeling macromolecular complexes based on our open-source *Integrative Modeling Platform* (IMP) package (5), version 2.8 (<https://integrativemodeling.org>). Files containing the input data, scripts, and output results are available at <https://integrativemodeling.org/systems/pemap> and the nascent integrative methods benchmarking section of the worldwide Protein Data Bank (wwPDB) PDB-Dev repository for integrative structures and corresponding data (<http://pdb-dev.wwpdb.org>) (98).

(i) Gathering data

To mimic realistic integrative structure determination, we did not rely on the known atomic structures of the subunits in the actual modeled complex [correct docking of exact bound structures based on geometric complementarity is easy, (99)]. Instead, we computed comparative models of histones H3 and H4 based on their alignments with structures of the 1TZY (100) (89 and 92% sequence identity, respectively), using MODELLER, version 9.21 (107). The Ca -atom RMSDs between the crystal structures and comparative models is 2.8 and 5.5 Å for H3 and H4, respectively, corresponding to medium- and low-accuracy comparative models. The second major input information source was a pE-MAP dataset of 308 point mutations in histones H3 and H4 crossed against an array of ~1370 gene deletion alleles, resulting in 946 MIC values above 0.3. Of these, 170 MIC values were converted into distance restraints between H3 and H4 residues (fig. S8 and table S3).

Comparative models of subunits Rpb1 and Rpb2 of yeast RNAPII were computed based on template structures 6GMH (102) (54% sequence identity) and 4AYB (103) (43% sequence identity), respectively. The Ca RMSD between the crystal structures of subunit Rpb1 and Rpb2 [2E2H (104)] and their comparative models are 7.3 and 5.2 Å, respectively. A pE-MAP dataset of 53 single point mutants in yeast RNAPII (44 of which reside in subunits Rpb1 and Rpb2) and a library of ~1200 gene-deletions resulted in 195 MIC values above 0.3. Of these, 123 MIC values were converted into distance restraints (fig. S8 and table S8). In addition, we compared the RNAPII model based on the pE-MAP to a model based on 22 previously published chemical cross-links (61).

The structures of subunits RpoB and RpoC of bacterial RNAP were obtained from the x-ray structure of the entire complex (4YG2) (105). A CG-MAP of 44 single point mutants of the two subunits and a library of 83 conditions (e.g., treatments with chemicals and temperature shocks) resulted in 109 MIC values above 0.3. Of these, 63 MIC values were converted into distance restraints between the subunits

(fig. S8 and table S9). In addition, we compared the bacterial RNAP model based on the CG-MAP to a model computed based on distance restraints derived from the interfacial contacts predicted using the RaptorX protein complex contact prediction server (65, 66)

(ii) Representing subunits and translating data into spatial restraints

To maximize computational efficiency while avoiding using too coarse a representation, we represented each complex in a multiscale fashion. In particular, the subunits and domains of each complex were coarse-grained using beads of varying sizes representing either a rigid body or a flexible string, based on the available comparative models, as follows (tables S3, S8, and S9). The comparative models were coarse-grained into two representations at different resolutions. First, we identified loop regions of at least eight residues using DSSP (106, 107) and represented them by flexible strings of beads of up to 10 residues each. Second, for the remaining residues each bead corresponded to an individual residue, centered at the position of its C α atom. With this representation in hand, we next translated the input information into spatial restraints as follows.

The defining and most important restraint for our method is extracted from the pE-MAP and CG-MAP data. The collected pE-MAP and CG-MAP MIC values were used to construct the Bayesian term in the scoring function that restrained the distances spanned by the mutated residues as described above. The pE-MAP restraint was applied to the one residue-per-bead representation for the comparative models as well as to the flexible beads. To improve computational efficiency, we only considered point mutation pairs with MIC values greater than 0.3. This restraint was applied to all three complexes (tables S3, S8, and S9). In addition to the pE-MAP data, integrative modeling can benefit from many other types of input information. Here, we have supplemented the pE-MAP and CG-MAP data by additional simple terms accounting for excluded volume and sequence connectivity. First, the excluded volume restraints were applied to each bead in the one-residue (or the closest) bead representations, using the statistical relationship between the volume and the number of residues that it covered (4, 108). Second, we applied the sequence connectivity restraint, using a harmonic upper bound on the distance between consecutive beads in a subunit, with a threshold distance equal to four times the sum of the radii of the two connected beads. The bead radius was calculated from the excluded volume of the corresponding bead, assuming standard protein density (4, 108). Moreover, we evaluated the utility of the pE-MAP and CG-MAP data by considering two additional types of restraints. First, the 22 previ-

ously determined BS3 RNAPII cross-links (67) were used to construct a Bayesian term that restrained the distances spanned by the cross-linked residues (30 Å) (109, 110). The cross-link restraints were applied to the one residue-per-bead representation for the comparative models as well as flexible beads, only for RNAPII (table S8). Second, we applied the evolutionary coupling restraints to determine the structures of the RpoB and RpoC subunits of bacterial RNA polymerase. Coupling strengths between residue pairs were obtained using the RaptorX ComplexContact server (<http://raptorx.uchicago.edu/ComplexContact/>) (65, 66) with default parameters. The top L/50 coupling strengths (fig. S9) with sequence separation of three or greater were converted into distance restraints using a harmonic upper bound on the distances between the residues. The threshold distance was set to 12 Å. This restraint was applied only to a subset of bacterial RNAP modeling instances (table S9).

Configurational sampling to produce an ensemble of structures that satisfy the restraints

The initial positions and orientations of rigid bodies and flexible beads were randomized. The generation of structural models was performed using Replica Exchange Gibbs sampling, based on the Metropolis Monte Carlo (MC) algorithm (110, 111). Each MC step consisted of a series of random transformations (i.e., rotation and translation) of the positions of the flexible beads and rigid bodies. Details about the MC runs for each system are in tables S3, S8, and S9.

Analyzing and validating the ensemble structures and data

Model validation follows four major steps (3, 112): (i) selection of the models for validation, (ii) estimation of sampling precision, (iii) estimation of model precision, and (iv) quantification of the degree to which a model satisfies the information used to compute it. These validations are based on the nascent wwPDB effort on archival, validation, and dissemination of integrative structures (98, 113). We now discuss each one of these validations in turn.

(i) Selection of models for validation: The first step is to objectively define the ensemble of models that will be further analyzed. For each trajectory, we automatically determined the MC step at which all data likelihoods and priors have equilibrated (run equilibration step), and all prior frames are discarded (114). Discarding the initial, nonequilibrated steps of each run is helpful because nontypical early configurations (e.g., a random configuration of beads, an extended configuration of beads, and beads far apart from each other) are removed from the statistical sample used for posterior model estimates.

With this ensemble of sampled structures and their corresponding scores in hand, we analyze the data likelihoods and priors. We used HDBSCAN clustering, a hierarchical density-based clustering algorithm, to identify all high-density regions in the likelihoods and priors (115). If a single cluster was identified, we consider all the models after discarding the initial steps; otherwise, we consider all models in the clusters that satisfy the input information, within the uncertainty of the data, for further analysis (below).

(ii) Estimation of sampling precision: Next, we estimate the precision at which sampling sampled the selected structures (sampling precision) (112); the sampling precision must be at least as high as the precision of the structure ensemble consistent with the input data (model precision). As a proxy for testing the thoroughness of sampling, we performed four sampling convergence tests: (i) verify that the scores of refined structures do not continue to improve as more structures are computed, (ii) confirm that the selected structures in independent sets of sampling runs (sample A and sample B) satisfy the data equally well, (iii) cluster the structural models and determine the sampling precision at which the structural features can be interpreted (fig. S10), and (iv) compare the localization probability density maps for each protein obtained from independent sets of runs. Details about all the tests are described in (112). For each modeling instance, the results from the convergence tests are summarized in tables S3, S8, and S9.

(iii) Estimation of model precision: In the third step, the model uncertainty (precision) is estimated. The most explicit description of model uncertainty is provided by the set of all models that are sufficiently consistent with the input information (i.e., the ensemble). Model precision can be quantified by the variability among the models in the ensemble; in the end, the ensemble can be described by one or more representative models and their uncertainties. For example, if the structures in the ensemble are clustered into a single cluster, the model precision is defined as the RMSD between models in the cluster. Importantly, the uncertainty may not be distributed evenly across the ensemble, such that some regions are determined at a higher precision than others.

(iv) Quantification of the degree to which a model satisfies the data used to compute it: An accurate structure needs to satisfy the input information used to compute it; all structures at computed precision that are consistent with the data are provided in the ensemble. A pE-MAP derived restraint is satisfied by a cluster of structures if the corresponding C α -C α distance in any of the structures in the cluster is lower than the distance predicted by the MIC value (Eq. 1). A BS3 cross-link restraint is satisfied

by a cluster of structures if the corresponding $\text{C}\alpha$ - $\text{C}\alpha$ distance in any of the structures in the cluster is less than 30 Å (116). The remainder of the restraints are harmonic, with a specified standard deviation. Therefore, a restraint is satisfied by a cluster of structures if the restrained distance in any structure in the cluster is violated by less than three standard deviations, specified for the restraint. Tables S3, S8, and S9 show that all models satisfy the input information within its uncertainty.

Benchmark

To benchmark the four-stage protocol described above, we computed the distribution of the accuracy for each structure in the ensemble of solutions obtained by integrative modeling. The accuracy is defined as the mean of $\text{C}\alpha$ RMSD between the x-ray structure and each of the structures in the ensemble. The PDB accession code and accuracies for each modeling instance are summarized in tables S3, S8, and S9.

To assess the information content of the histone pE-MAP, we computed the models of the H3-H4 complex based on random subsets of the data. To this end, from the dataset of computed MIC values for pairs of mutated residues, we performed three independent random selections of 80, 60, 40, and 20% of the data each. As expected, the more pE-MAP data that are used, the more accurate and precise are the models (Fig. 6C).

Finally, as another test, we computed the model based on datasets with randomly shuffled MIC values for the same pE-MAP or CG-MAP residue pairs, for each of the complexes.

Estimation of the number of mutations per protein

To estimate the suggested number of mutated positions per protein for integrative structure determination, we computed the number of mutations that would result in four or more MIC values above a 0.4, 0.45, 0.5, or 0.55 threshold. Based on our scoring function, MIC values above these thresholds will result in distance restraints with an upper distance bound in the 12- to 34-Å range. These distances are comparable to the upper distance bounds used for chemical cross-links (e.g., DSS, DSSO, EDC). A previous systematic study established that at least four chemical cross-links are needed to determine the binding mode of protein dimers if the subunit structures are known (e.g., from x-ray, NMR, or comparative models) (110). In general, adding more chemical cross-links does not further improve the accuracy, although it increases the precision of the resulting ensemble. By analogy, we estimate that, for systems in which the structures of the components are known, a good number of mutations per protein is 35 to 40 (fig. S5A). This data can be used as a guideline to decide on the number of mutations to use for

generating a pE-MAP or CG-MAP. Importantly, this estimate is an upper bound on the number of mutations, and in many cases, the number might be smaller for the following two reasons. First, this estimation was done assuming protein-wide mutations of residues, often to alanine. In practice, the number of necessary mutations can be reduced by specifically designing point mutations that target surface residues and/or residues known to be functionally important, and by choosing substitutions likely to give rise to functional perturbations. In general, we did not find a correlation between the secondary structure of the residue pairs and their associated MIC value (fig. S5B). Second, this estimation only relies on the residue pairs with high MIC values. In contrast to chemical cross-links, the upper distance bound of pE-MAP derived restraints are obtained from the statistical association between the MIC values and distance between residues. Consequently, residue pairs with low MIC values still carry structural information, even if at low resolution. Consistent with these considerations, the RNAPII dataset contains only 31 and 9 mutated residues for Rpb1 and Rpb2, respectively. Similarly, the bacterial RNAP dataset contains 23 and 15 mutated residues for RpoB and RpoC, respectively.

Docking

To assess the relative value of pE-MAP restraints for structure determination, we computed the structures of the H3-H4 and RNAPII complexes by molecular docking. Specifically, we followed an integrative docking protocol (117) using the rigid-body docking program PatchDock (39). In each case, we used the same comparative models and rigid-body definitions used for integrative modeling (figs. S4 and S11) and default parameter values (figs. S4 and S11).

Visualizations

The pE-MAP was hierarchically clustered in both histone mutant and gene deletion dimensions using Cluster 3.0 (118) and displayed using Java Treeview (119) (Fig. 3). Images highlighting histone residues in context of the nucleosome structure (PDB IID3 or its modified version in data S2) were created using ChimeraX (Figs. 3A and 7C; and figs. S1D and S6) (48).

Distance of clustered pE-MAP profiles

First, all histone alleles affecting residues not included in the structural reference (PDB IID3, H3A1-H3K37 and H4S1-H4R17) were removed and the remaining data ($n = 222$) clustered hierarchically using Cluster 3.0 (118). For each node of the clustergram, the mean distance among member residues was calculated and plotted versus the normalized branch length (where the first node is set to branch length = 0 and the last node to branch length = 1) of the

respective node (fig. S2A, red dots, and random distribution plotted in black).

Generation of the correlation map

Pearson correlation coefficients were computed for each of the 350 H3 and H4 mutants against all genes and alleles (rows) in a merged map of previously published genetic interaction data [dataset S4 from (45)]. If the overlap between a histone mutant and a S-score vector from the merged map was <150 scores, the resulting correlation was not considered (i.e., replaced by “NaN”) (data S3). Pearson correlation coefficients were chosen over MIC for this analysis because we found Pearson correlation more robust than MIC when many missing values were present.

Structural mapping of genetic interactions—stEMAP app

The hierarchically clustered pE-MAP data was imported into Cytoscape (47), creating an initial network, and then linked to a modified version of the nucleosome structure IID3 (data S2) using the stEMAP app, developed to facilitate interactive exploration of the pE-MAP. The original nucleosome structure was modified by adding the N-terminal disordered regions of histone H3 and H4 and manually positioning them for clarity. The linking proceeds as follows: First, the structure is opened in ChimeraX (48) by structureVizX (120) and positioned in response to commands from the stEMAP app. Then, a residue interaction network (RIN) is created by the structureVizX app where nodes are positioned to reflect the nucleosome structure through the help of the RINalyzer app (121). Finally, the RIN network and the network created by the original cluster files are merged, and edges are drawn between genes and residues with interactions passing a user-defined threshold (fig. S6A). All of the preceding steps happen automatically through the stEMAP app interface, which takes as input any given PDB file and a short user-defined JSON configuration file defining interaction thresholds (here: correlation > 0.2), colors of edges (here: color-gradient from white to red for positive correlations), and display style of the structure in ChimeraX.

Selection of individual genes triggers the interacting residues to be selected, and, in the ChimeraX window, those residues are shown as space-filling atoms, which are colored according to the edge colors. When multiple genes are selected (e.g., genes belonging to the same complex), there might be multiple edges connecting an individual residue. In this case, the color reflects the significance and consistency of the interactions (see below). To assist in interpretation and interactive exploration of complex data sets (i) colors are quantized into 10 bins, five positive and five negative; (ii) a heatmap is presented that shows only the

values for the selected genes and their interacting residues; (iii) sets of genes belonging to a complex can be selected using the setsApp (122); and (iv) a slider provides a filter to restrict the selection to only those mutations with a minimum number of interactions.

To determine if a gene set is connected to a given residue, the stEMAP app calculates the median Pearson correlation coefficient across all genes of the gene set against all different mutations at that residue. If this median correlation is above the threshold of 0.2 (defined in the JSON file), the respective residue is colored according to the median. To instead determine if a gene set is connected to an individual mutation, the same method is used, except the median correlation coefficient is now calculated across all genes of the gene set against the single given mutation (instead of all mutations of the residue).

ROC curves

Only library deletion mutants that exist in both this study and the two previously published E-MAP datasets [Braberg *et al.* (15) and Collins *et al.* (29)] were included ($n = 389$) in this analysis. Based on their pE-MAP profiles, Pearson correlation coefficients were calculated for all pairwise combinations of these 389 deletion mutants. To determine the power of these correlations to predict physical interactions between encoded proteins, an ROC curve was computed, where a physical interaction between proteins was defined if their PE score is greater than 2 (69). From the Collins *et al.* E-MAP, query strain profiles with more missing data than the sparsest histone mutant were removed, as were query mutants that also existed in the library mutant set. Because the Braberg *et al.* pE-MAP only includes 53 query mutants (rows), we used subsets of 53 query mutants each for the histone and Collins *et al.* E-MAPs when generating their ROC curves, to make all three systems comparable. To this end, for the Collins *et al.* E-MAP and histone pE-MAP, 53 query mutant profiles were randomly selected 1000 times, and a ROC curve was generated for each run. The median areas under the ROC curves (AROCs) and corresponding ROC curves are reported together with the ROC curve of the pE-MAP from Braberg *et al.* in Fig. 2F.

RNA-seq expression analysis

Ten ml of overnight cultures of 29 histone mutant strains (table S2) were harvested in mid-log phase [optical density at 600 nm (OD_{600}) ≈ 1.0] and washed with DEPC-ddH₂O. RNA was extracted with hot acidic phenol as described previously (123). RNA-seq libraries were generated using the QuantSeq 3' mRNA-Seq Library Prep Kit FWD for Illumina (Lexogen). Single-end, 50-base reads were sequenced using an Illumina HiSeq 4000 se-

quencer. Reads were filtered for quality and aligned to the yeast genome using tophat (124). Nonunique reads and reads mapping to ribosomal RNA were removed before analysis. Transcript counts were extracted using htseq-count (125), and differential expression was measured using the Dseq2 package in R (126).

Identification of functional links between H3 and H4 mutants and biological processes

The correlation map (data S3) was used as the basis for this analysis. First, a curated annotation of all genes in the correlation map relevant to nuclear function was devised. Biological process definitions for genes in nuclear processes were assigned manually based on literature and annotations from previous genetic interaction maps (29, 127, 128) (table S5). To identify links between H3 and H4 residues and nuclear processes that were highly correlated, we used a one-sided Mann-Whitney *U* test to compare the correlation distribution between the mutants of each H3 and H4 residue and the members of each process to (i) the correlations between the same H3 and H4 mutants and all genes not in that process and to (ii) the correlations between the same process and all other H3 and H4 mutants. The highest *p* value of the comparison to (i) or (ii) was recorded. False discovery rates (FDRs) for the links were then computed using the method of Benjamini and Hochberg (129) and are reported in table S5. The most significant links with $FDR < 10^{-6}$ were used for follow-ups.

Spontaneous mutation frequency

Cells were grown to saturation and then plated on yeast extract peptone dextrose (YEPD)- and 5-fluoroorotic acid (5-FOA)-supplemented media. Mutants growing on 5-FOA were counted only after confirming that colonies growing on YEPD for all the strains were of equal size. The assay was repeated three times independently (table S6).

MS quantification of H3K56ac levels

Sample preparation

Histone mutant cultures (WT, H3R63K, and H4R36K) were harvested in mid-log phase ($OD_{600} \approx 1.0$) using a 250-mm ceramic filter funnel and 30- μ m nitrocellulose membranes connected to high continuous wall suction. Yeast were removed from the nitrocellulose membrane and flash frozen for storage or used immediately for protein extraction. Per gram of yeast pellet, 3 ml of Yeast-Protein Extract Reagent (Y-PER; ThermoFisher Scientific) with added protease inhibitors (cOmplete Sigma-Aldrich, one tablet per 50 ml), phosphatase inhibitors (PhosSTOP Sigma-Aldrich; one tablet per 50 ml), histone deacetylase inhibitors (sodium butyrate 100 mM and nicotinamide 100 mM), and β -mercaptoethanol (15 mM) were added. The suspension was mixed

on a gyrator at 4°C for 30 min and centrifuged. Pellets were resuspended in fresh Y-PER medium, and extraction was repeated two additional times for a total of three extractions. Pellets were sequentially washed twice with 3 ml ddH₂O per gram of yeast. Histone extraction was performed in the presence of 2.5 ml of 8M urea/0.4N sulfuric acid per gram of yeast protein pellets, incubated for 1 hour, centrifuged, and supernatants collected. Proteins were precipitated using a methanol-chloroform precipitation as previously described (130). Extracted proteins were trypsin digested; desalted and acetylated peptides were enriched as previously described (131).

Generation of selected reaction monitoring (SRM) assays for acetylation sites

Peptide mixtures (obtained from ThermoFisher) were analyzed by liquid chromatography-tandem mass spectrometry (LC-MS/MS) on a Thermo Scientific Orbitrap Fusion mass spectrometry system equipped with a Proxeon Easy nLC 1200 ultra-high-pressure LC and autosampler system. Samples were injected onto a C18 column (25 cm $\times$ 75 μ m I.D. packed with ReproSil Pur C18 AQ 1.9- μ m particles) in 0.1% formic acid and then separated with a 60-min gradient from 5 to 40% buffer B (90% ACN/10% water/0.1% formic acid) at a flow rate of 300 nl/min. The mass spectrometer collected data in a data-dependent fashion, collecting one full scan in the Orbitrap followed by collision-induced dissociation MS/MS scans in the dual linear ion trap for the 20 most intense peaks from the full scan. Dynamic exclusion was enabled for 30 s with a repeat count of 1. Charge state screening was used to reject analysis of singly charged species or species for which a charge could not be assigned. The raw data was matched to protein sequences using the MaxQuant algorithm (version 1.5.2.8) (132). Data were searched against a database containing SwissProt Human protein sequences concatenated to a decoy database where each protein sequence was randomized in order to estimate the FDR. Variable modifications were allowed for methionine oxidation and protein N-terminal acetylation and lysine acetylation. A fixed modification was indicated for cysteine carbamidomethylation. Full trypsin specificity was required. The first search was performed with a mass accuracy of ± 20 parts per million and the main search was performed with a mass accuracy of ± 4.5 parts per million. A maximum of five modifications were allowed per peptide. A maximum of two missed cleavages were allowed. The maximum charge allowed was 7+. Individual peptide mass tolerances were allowed. For MS/MS matching, a mass tolerance of 0.8 Da was allowed, and the top eight peaks per 100 Da were analyzed. MS/MS matching was allowed for higher charge states and water and ammonia

loss events. The data were filtered to obtain a peptide, protein, and site-level FDR of 0.01. The minimum peptide length was seven amino acids. Selected reaction monitoring (SRM) assays were generated for selected acetylation sites. SRM assay generation was performed using Skyline (133). For all targeted proteins, proteotypic peptides and optimal transitions for identification and quantification were selected based on a spectral library generated from the shotgun MS experiments. The Skyline spectral library was used to extract optimal coordinates for the SRM assays, for example, peptide fragments and peptide retention times. For each peptide, the five best SRM transitions were selected based on intensity and peak shape.

Acquisition and quantification of acetylation sites by SRM

Digested peptide mixtures were analyzed by LC-SRM on a Thermo Scientific TSQ Quantiva MS system equipped with a Proxeon Easy nLC 1200 ultra-high-pressure LC and autosampler system. Samples were injected onto a C18 column (25 cm × 75 µm I.D. packed with ReproSil Pur C18 AQ 1.9-µm particles) in 0.1% formic acid and then separated with a 60-min gradient from 5 to 40% buffer B (90% ACN/10% water/0.1% formic acid) at a flow rate of 300 nL/min. SRM acquisition was performed operating Q1 and Q3 at 0.7-unit-mass resolution. For each peptide, the best five transitions were monitored in a scheduled fashion with a retention time window of 5 min and a cycle time fixed to 2 s. Argon was used as the collision gas at a nominal pressure of 1.5 mTorr. Collision energies were calculated by, $CE = 0.0348 \times (m/z) + 0.4551$ and $CE = 0.0271 \times (m/z) + 1.5910$ (CE , collision energy and m/z , mass to charge ratio) for doubly and triply charged precursor ions, respectively. SRM data was processed using Skyline (133). Protein significance analysis was performed using MSstats (134). Normalization of the intensities across samples was performed using the acetylated peptides H3K9 H3K14 (peptide containing both acetylation sites), H3K23 and H3K14 as global standards, which did not show any change across the mutants. Log₂ fold changes were calculated from three independent runs and plotted (fig. S7C and table S6).

Cryptic transcription—qPCR

Total RNA was extracted from 10 OD₆₀₀ units of mid-log phase cells (WT and respective mutant strains) using hot acid phenol-chloroform extraction method as described. Ten µg of total RNA was DNase I treated (Promega) followed by purification using an RNeasy minikit (Qiagen). One µg of DNase I-treated total RNA was used to synthesize cDNA using SuperScript III first strand synthesis system (Life Technologies) and random hexamer primers. cDNA was diluted 1:25 before amplification by PCR using primers

designed for the 5' and the 3' ends of the *STE11* gene. qPCR was performed using SYBR green (Biorad) as described previously (135). Relative change in the transcript levels were estimated using the $\Delta\Delta C_t$ method described in (136) and were normalized to *ACT1* transcript (table S7). Primer sequences are available upon request.

Western blotting

Whole-yeast-cell lysates were prepared using trichloroacetic acid (TCA) lysis as described previously (137). Lysates were subjected to immunoblotting according to standard procedures, and proteins were detected using ECL Prime (Amersham Biosciences). Membranes were probed with α H3K36me3 antibody purchased from (Abcam, catalog no. 9050). Glyceraldehyde-phosphate dehydrogenase (GAPDH) was used as a loading control and detected using an antibody purchased from Sigma (catalog no. A9521).

ATAC-seq

Yeast cells (2.5×10^6) were grown to mid-log phase, pelleted, washed with SB-buffer (1.4 M Sorbitol, 40 mM HEPES-KOH pH 7.5, 0.5 mM MgCl₂), resuspended in 200 µl SB buffer + 10 mM DTT with 10 µl of 10 mg/ml 100T zymolyase (MP Biomedicals) solution and incubated for 5 min at 30°C. Spheroblasted cells were washed with SB-buffer and incubated for 15 min at 37°C in 25 µl of transposase solution (12.5 µl 2x TD buffer, 1.25 µl Nextera enzyme, 11.25 µl water). DNA was purified (Qiagen MinElute DNA Purification Kit), amplified, and barcoded by PCR. Purified PCR products were sequenced using an Illumina HiSeq 4000 sequencer. Sequence reads were trimmed and aligned to the genome of *Saccharomyces cerevisiae* (version SacCer3 from hgdownload.cse.ucsc.edu/downloads.html), and reads with a length <100 basepairs (bp) were removed. Replicates belonging to an allele (WT, H3K36A, H3K122A, *set2Δ*) were merged and normalized to the smallest read number. For visualization of *STE11* read coverage using the IGV genome browser (138) (fig. S7G), each track was scaled linearly so that the largest peak in the displayed window is the same height for all tracks. Count files were generated with “featureCounts v1.5.3” (139).

Gene body plots (fig. S7H) were generated as follows: First, counts from genes reported to be targets of cryptic transcription ($n = 11$; *FLO8*, *AVO1*, *LCB5*, *SMC3*, *SPB4*, *APM2*, *DDC1*, *SYF1*, *OMS1*, *PUS4*, *STE11*), as well as counts 400 bp up- and downstream of the respective gene bodies, were extracted. Second, up- and downstream regions were split into 50 bins of equal size (8 bp), whereas the gene body was split into 300 equal bins, resulting in 400 bins for each gene in each tested strain (WT, *set2Δ*, H3K36A and H3K122A). Third, for each of the 400 bins, the average for the 11 target genes

was calculated. Fourth, each mutant allele was then scaled linearly so that the first bin (i.e., 400 bp upstream of the gene body start) was equal to that of WT. Finally, the WT counts were subtracted from the mutant counts for each bin.

REFERENCES AND NOTES

- F. Alber, F. Förster, D. Korkin, M. Topf, A. Sali, Integrating diverse data for structure determination of macromolecular assemblies. *Annu. Rev. Biochem.* **77**, 443–477 (2008). doi: [10.1146/annurev.biochem.77.060407.135530](https://doi.org/10.1146/annurev.biochem.77.060407.135530); pmid: [18318657](https://pubmed.ncbi.nlm.nih.gov/18318657/)
- F. Herzog et al., Structural probing of a protein phosphatase 2A network by chemical cross-linking and mass spectrometry. *Science* **337**, 1348–1352 (2012). doi: [10.1126/science.1221483](https://doi.org/10.1126/science.1221483); pmid: [22984071](https://pubmed.ncbi.nlm.nih.gov/22984071/)
- M. P. Rout, A. Sali, Principles for integrative structural biology studies. *Cell* **177**, 1384–1403 (2019). doi: [10.1016/j.cell.2019.05.016](https://doi.org/10.1016/j.cell.2019.05.016); pmid: [31150619](https://pubmed.ncbi.nlm.nih.gov/31150619/)
- F. Alber et al., Determining the architectures of macromolecular assemblies. *Nature* **450**, 683–694 (2007). doi: [10.1038/nature06404](https://doi.org/10.1038/nature06404); pmid: [18046405](https://pubmed.ncbi.nlm.nih.gov/18046405/)
- D. Russel et al., Putting the pieces together: Integrative modeling platform software for structure determination of macromolecular assemblies. *PLOS Biol.* **10**, e1001244 (2012). doi: [10.1371/journal.pbio.1001244](https://doi.org/10.1371/journal.pbio.1001244); pmid: [2272186](https://pubmed.ncbi.nlm.nih.gov/2272186/)
- A. B. Ward, A. Sali, I. A. Wilson, Integrative structural biology. *Science* **339**, 913–915 (2013). doi: [10.1126/science.1228565](https://doi.org/10.1126/science.1228565); pmid: [23430643](https://pubmed.ncbi.nlm.nih.gov/23430643/)
- F. Alber et al., The molecular architecture of the nuclear pore complex. *Nature* **450**, 695–701 (2007). doi: [10.1038/nature06405](https://doi.org/10.1038/nature06405); pmid: [18046406](https://pubmed.ncbi.nlm.nih.gov/18046406/)
- K. Lasker et al., Molecular architecture of the 26S proteasome holocomplex determined by an integrative approach. *Proc. Natl. Acad. Sci. U.S.A.* **109**, 1380–1387 (2012). doi: [10.1073/pnas.1120559109](https://doi.org/10.1073/pnas.1120559109); pmid: [22307589](https://pubmed.ncbi.nlm.nih.gov/22307589/)
- A. Loquet et al., Atomic model of the type III secretion system needle. *Nature* **486**, 276–279 (2012). doi: [10.1038/nature11079](https://doi.org/10.1038/nature11079); pmid: [22699623](https://pubmed.ncbi.nlm.nih.gov/22699623/)
- Z. Duan et al., A three-dimensional model of the yeast genome. *Nature* **465**, 363–367 (2010). doi: [10.1038/nature08973](https://doi.org/10.1038/nature08973); pmid: [20436457](https://pubmed.ncbi.nlm.nih.gov/20436457/)
- J. M. Plitzko, B. Schuler, P. Selenko, Structural biology outside the box—inside the cell. *Curr. Opin. Struct. Biol.* **46**, 110–121 (2017). doi: [10.1016/j.sbi.2017.06.007](https://doi.org/10.1016/j.sbi.2017.06.007); pmid: [28735108](https://pubmed.ncbi.nlm.nih.gov/28735108/)
- M. Dimura et al., Quantitative FRET studies and integrative modeling unravel the structure and dynamics of biomolecular systems. *Curr. Opin. Struct. Biol.* **40**, 163–185 (2016). doi: [10.1016/j.sbi.2016.11.012](https://doi.org/10.1016/j.sbi.2016.11.012); pmid: [27939973](https://pubmed.ncbi.nlm.nih.gov/27939973/)
- S. R. Collins, A. Roguev, N. J. Krogan, Quantitative genetic interaction mapping using the E-MAP approach. *Methods Enzymol.* **470**, 205–231 (2010). doi: [10.1016/S0076-6879\(10\)70009-4](https://doi.org/10.1016/S0076-6879(10)70009-4); pmid: [20946812](https://pubmed.ncbi.nlm.nih.gov/20946812/)
- P. Beltrao, G. Cagny, N. J. Krogan, Quantitative genetic interactions reveal biological modularity. *Cell* **141**, 739–745 (2010). doi: [10.1016/j.cell.2010.05.019](https://doi.org/10.1016/j.cell.2010.05.019); pmid: [20510918](https://pubmed.ncbi.nlm.nih.gov/20510918/)
- H. Braberg et al., From structure to systems: High-resolution, quantitative genetic analysis of RNA polymerase II. *Cell* **154**, 775–788 (2013). doi: [10.1016/j.cell.2013.07.033](https://doi.org/10.1016/j.cell.2013.07.033); pmid: [23932120](https://pubmed.ncbi.nlm.nih.gov/23932120/)
- H. Braberg, E. A. Moehle, M. Shales, C. Guthrie, N. J. Krogan, Genetic interaction analysis of point mutations enables interrogation of gene function at a residue-level resolution: Exploring the applications of high-resolution genetic interaction mapping of point mutations. *BioEssays* **36**, 706–713 (2014). doi: [10.1002/bies.201400044](https://doi.org/10.1002/bies.201400044); pmid: [24842270](https://pubmed.ncbi.nlm.nih.gov/24842270/)
- N. Halabi, O. Rivoire, S. Leibler, R. Ranganathan, Protein sectors: Evolutionary units of three-dimensional structure. *Cell* **138**, 774–786 (2009). doi: [10.1016/j.cell.2009.07.038](https://doi.org/10.1016/j.cell.2009.07.038); pmid: [19703402](https://pubmed.ncbi.nlm.nih.gov/19703402/)
- D. S. Marks et al., Protein 3D structure computed from evolutionary sequence variation. *PLOS ONE* **6**, e28766 (2011). doi: [10.1371/journal.pone.0028766](https://doi.org/10.1371/journal.pone.0028766); pmid: [21613331](https://pubmed.ncbi.nlm.nih.gov/21613331/)
- G. Diss, B. Lehner, The genetic landscape of a physical interaction. *eLife* **7**, e32472 (2018). doi: [10.7554/eLife.32472](https://doi.org/10.7554/eLife.32472); pmid: [29638215](https://pubmed.ncbi.nlm.nih.gov/29638215/)
- A. L. Shiver, H. Osadnik, J. M. Peters, R. A. Mooney, P. I. Wu, J. C. Hu, R. Landick, K. C. Huang, C. A. Gross,

- Chemical-genetic interrogation of RNA polymerase mutants reveals structure-function relationships and physiological tradeoffs. *bioRxiv* 2020.2006.2016.155770 [Preprint]. 17 June 2020; doi: [10.1101/2020.06.16.155770](https://doi.org/10.1101/2020.06.16.155770)
21. H. Huang, S. Lin, B. A. Garcia, Y. Zhao, Quantitative proteomic analysis of histone modifications. *Chem. Rev.* **115**, 2376–2418 (2015). doi: [10.1021/cr500491u](https://doi.org/10.1021/cr500491u); pmid: [25688442](https://pubmed.ncbi.nlm.nih.gov/25688442/)
 22. J. Dai *et al.*, Probing nucleosome function: A highly versatile library of synthetic histone H3 and H4 mutants. *Cell* **134**, 1066–1078 (2008). doi: [10.1016/j.cell.2008.07.019](https://doi.org/10.1016/j.cell.2008.07.019); pmid: [18805098](https://pubmed.ncbi.nlm.nih.gov/18805098/)
 23. S. Jiang *et al.*, Construction of comprehensive dosage-matching core histone mutant libraries for *Saccharomyces cerevisiae*. *Genetics* **207**, 1263–1273 (2017). pmid: [29084817](https://pubmed.ncbi.nlm.nih.gov/29084817/)
 24. J. E. Brownell *et al.*, Tetrahymena histone acetyltransferase A: A homolog to yeast Gcn5p linking histone acetylation to gene activation. *Cell* **84**, 843–851 (1996). doi: [10.1016/S0092-8674\(00\)81063-6](https://doi.org/10.1016/S0092-8674(00)81063-6); pmid: [8601308](https://pubmed.ncbi.nlm.nih.gov/8601308/)
 25. L. K. Durrin, R. K. Mann, P. S. Kayne, M. Grunstein, Yeast histone H4 N-terminal sequence is required for promoter activation in vivo. *Cell* **65**, 1023–1031 (1991). doi: [10.1016/0092-8674\(91\)90554-C](https://doi.org/10.1016/0092-8674(91)90554-C); pmid: [2044150](https://pubmed.ncbi.nlm.nih.gov/2044150/)
 26. H. Braberg *et al.*, Quantitative analysis of triple-mutant genetic interactions. *Nat. Protoc.* **9**, 1867–1881 (2014). doi: [10.1038/nprot.2014.127](https://doi.org/10.1038/nprot.2014.127); pmid: [25010907](https://pubmed.ncbi.nlm.nih.gov/25010907/)
 27. J. E. Haber *et al.*, Systematic triple-mutant analysis uncovers functional connectivity between pathways involved in chromosome regulation. *Cell Rep.* **3**, 2168–2178 (2013). doi: [10.1016/j.celrep.2013.05.007](https://doi.org/10.1016/j.celrep.2013.05.007); pmid: [23746449](https://pubmed.ncbi.nlm.nih.gov/23746449/)
 28. S. R. Collins, M. Schuldiner, N. J. Krogan, J. S. Weissman, A strategy for extracting and analyzing large-scale quantitative epistatic interaction data. *Genome Biol.* **7**, R63 (2006). doi: [10.1186/gb-2006-7-7-r63](https://doi.org/10.1186/gb-2006-7-7-r63); pmid: [16859555](https://pubmed.ncbi.nlm.nih.gov/16859555/)
 29. S. R. Collins *et al.*, Functional dissection of protein complexes involved in yeast chromosome biology using a genetic interaction map. *Nature* **446**, 806–810 (2007). doi: [10.1038/nature05649](https://doi.org/10.1038/nature05649); pmid: [17314980](https://pubmed.ncbi.nlm.nih.gov/17314980/)
 30. T. Miller *et al.*, COMPASS: A complex of proteins associated with a trithorax-related SET domain protein. *Proc. Natl. Acad. Sci. U.S.A.* **98**, 12902–12907 (2001). doi: [10.1073/pnas.231473398](https://doi.org/10.1073/pnas.231473398); pmid: [11687631](https://pubmed.ncbi.nlm.nih.gov/11687631/)
 31. A. Roguev *et al.*, The *Saccharomyces cerevisiae* Set1 complex includes an Ash2 homologue and methylates histone 3 lysine 4. *EMBO J.* **20**, 7137–7148 (2001). doi: [10.1093/emboj/20.24.7137](https://doi.org/10.1093/emboj/20.24.7137); pmid: [11742990](https://pubmed.ncbi.nlm.nih.gov/11742990/)
 32. G. Mizuguchi *et al.*, ATP-driven exchange of histone H2AZ variant catalyzed by SWR1 chromatin remodeling complex. *Science* **303**, 343–348 (2004). doi: [10.1126/science.1090701](https://doi.org/10.1126/science.1090701); pmid: [14645854](https://pubmed.ncbi.nlm.nih.gov/14645854/)
 33. M. J. Carrozza *et al.*, Histone H3 methylation by Set2 directs deacetylation of coding regions by Rpd3S to suppress spurious intragenic transcription. *Cell* **123**, 581–592 (2005). doi: [10.1016/j.cell.2005.10.023](https://doi.org/10.1016/j.cell.2005.10.023); pmid: [16286007](https://pubmed.ncbi.nlm.nih.gov/16286007/)
 34. S. Venkatesh *et al.*, Set2 methylation of histone H3 lysine 36 suppresses histone exchange on transcribed genes. *Nature* **489**, 452–455 (2012). doi: [10.1038/nature11326](https://doi.org/10.1038/nature11326); pmid: [22914091](https://pubmed.ncbi.nlm.nih.gov/22914091/)
 35. M. C. Keogh *et al.*, Cotranscriptional set2 methylation of histone H3 lysine 36 recruits a repressive Rpd3 complex. *Cell* **123**, 593–605 (2005). doi: [10.1016/j.cell.2005.10.025](https://doi.org/10.1016/j.cell.2005.10.025); pmid: [16286008](https://pubmed.ncbi.nlm.nih.gov/16286008/)
 36. D. N. Reshef *et al.*, Detecting novel associations in large data sets. *Science* **334**, 1518–1524 (2011). doi: [10.1126/science.1205438](https://doi.org/10.1126/science.1205438); pmid: [21724245](https://pubmed.ncbi.nlm.nih.gov/21724245/)
 37. D. Albanese *et al.*, Minerva and minepy: A C engine for the MINE suite and its R, Python and MATLAB wrappers. *Bioinformatics* **29**, 407–408 (2013). doi: [10.1093/bioinformatics/bts707](https://doi.org/10.1093/bioinformatics/bts707); pmid: [23242262](https://pubmed.ncbi.nlm.nih.gov/23242262/)
 38. C. L. White, R. K. Suto, K. Luger, Structure of the yeast nucleosome core particle reveals fundamental changes in internucleosome interactions. *EMBO J.* **20**, 5207–5218 (2001). doi: [10.1093/emboj/20.18.5207](https://doi.org/10.1093/emboj/20.18.5207); pmid: [11566884](https://pubmed.ncbi.nlm.nih.gov/11566884/)
 39. D. Schneidman-Duhovny, Y. Inbar, R. Nussinov, H. J. Wolfson, PatchDock and SymmDock: Servers for rigid and symmetric docking. *Nucleic Acids Res.* **33**, W363–W367 (2005). doi: [10.1093/nar/gki481](https://doi.org/10.1093/nar/gki481); pmid: [15980490](https://pubmed.ncbi.nlm.nih.gov/15980490/)
 40. D. Baker, A. Sali, Protein structure prediction and structural genomics. *Science* **294**, 93–96 (2001). doi: [10.1126/science.1065659](https://doi.org/10.1126/science.1065659); pmid: [11588250](https://pubmed.ncbi.nlm.nih.gov/11588250/)
 41. S. J. Kim *et al.*, Integrative structure and functional anatomy of a nuclear pore complex. *Nature* **555**, 475–482 (2018). doi: [10.1038/nature26003](https://doi.org/10.1038/nature26003); pmid: [29539637](https://pubmed.ncbi.nlm.nih.gov/29539637/)
 42. C. Gutierrez *et al.*, Structural dynamics of the human COP9 signalosome revealed by cross-linking mass spectrometry and integrative modeling. *Proc. Natl. Acad. Sci. U.S.A.* **117**, 4088–4098 (2020). doi: [10.1073/pnas.1915542117](https://doi.org/10.1073/pnas.1915542117); pmid: [32034103](https://pubmed.ncbi.nlm.nih.gov/32034103/)
 43. K. S. Molnar *et al.*, Cys-scanning disulfide crosslinking and bayesian modeling probe the transmembrane signaling mechanism of the histidine kinase, PhoQ. *Structure* **22**, 1239–1251 (2014). doi: [10.1016/j.str.2014.04.019](https://doi.org/10.1016/j.str.2014.04.019); pmid: [25087511](https://pubmed.ncbi.nlm.nih.gov/25087511/)
 44. Y. Kwon *et al.*, Structural basis of CD4 downregulation by HIV-1 Nef. *Nat. Struct. Mol. Biol.* **27**, 822–828 (2020). doi: [10.1038/s41594-020-0463-z](https://doi.org/10.1038/s41594-020-0463-z); pmid: [32719457](https://pubmed.ncbi.nlm.nih.gov/32719457/)
 45. C. J. Ryan *et al.*, Hierarchical modularity and the evolution of genetic interactomes across species. *Mol. Cell* **46**, 691–704 (2012). doi: [10.1016/j.molcel.2012.05.028](https://doi.org/10.1016/j.molcel.2012.05.028); pmid: [22681890](https://pubmed.ncbi.nlm.nih.gov/22681890/)
 46. J. Wysocka *et al.*, A PHD finger of NURF couples histone H3 lysine 4 trimethylation with chromatin remodelling. *Nature* **442**, 86–90 (2006). doi: [10.1038/nature04815](https://doi.org/10.1038/nature04815); pmid: [16728976](https://pubmed.ncbi.nlm.nih.gov/16728976/)
 47. P. Shannon *et al.*, Cytoscape: A software environment for integrated models of biomolecular interaction networks. *Genome Res.* **13**, 2498–2504 (2003). doi: [10.1101/gr.1239303](https://doi.org/10.1101/gr.1239303); pmid: [14597658](https://pubmed.ncbi.nlm.nih.gov/14597658/)
 48. T. D. Goddard *et al.*, UCSF ChimeraX: Meeting modern challenges in visualization and analysis. *Protein Sci.* **27**, 14–25 (2018). doi: [10.1002/pro.3235](https://doi.org/10.1002/pro.3235); pmid: [28710774](https://pubmed.ncbi.nlm.nih.gov/28710774/)
 49. D. Schneidman-Duhovny, R. Pellarin, A. Sali, Uncertainty in integrative structural modeling. *Curr. Opin. Struct. Biol.* **28**, 96–104 (2014). doi: [10.1016/j.sbi.2014.08.001](https://doi.org/10.1016/j.sbi.2014.08.001); pmid: [25173450](https://pubmed.ncbi.nlm.nih.gov/25173450/)
 50. N. L. Maas, K. M. Miller, L. G. DeFazio, D. P. Toczyski, Cell cycle and checkpoint regulation of histone H3 K56 acetylation by Hst3 and Hst4. *Mol. Cell* **23**, 109–119 (2006). doi: [10.1016/j.molcel.2006.06.006](https://doi.org/10.1016/j.molcel.2006.06.006); pmid: [16818235](https://pubmed.ncbi.nlm.nih.gov/16818235/)
 51. J. Han *et al.*, A Cul4 E3 ubiquitin ligase regulates histone hand-off during nucleosome assembly. *Cell* **155**, 817–829 (2013). doi: [10.1016/j.cell.2013.10.014](https://doi.org/10.1016/j.cell.2013.10.014); pmid: [24209620](https://pubmed.ncbi.nlm.nih.gov/24209620/)
 52. T. Tsubota *et al.*, Histone H3-K56 acetylation is catalyzed by histone chaperone-dependent complexes. *Mol. Cell* **25**, 703–712 (2007). doi: [10.1016/j.molcel.2007.02.006](https://doi.org/10.1016/j.molcel.2007.02.006); pmid: [17320445](https://pubmed.ncbi.nlm.nih.gov/17320445/)
 53. E. M. Hyland *et al.*, Insights into the role of histone H3 and histone H4 core modifiable residues in *Saccharomyces cerevisiae*. *Mol. Cell. Biol.* **25**, 10060–10070 (2005). doi: [10.1128/MCB.25.22.10060-10070.2005](https://doi.org/10.1128/MCB.25.22.10060-10070.2005); pmid: [16260619](https://pubmed.ncbi.nlm.nih.gov/16260619/)
 54. H. Masumoto, D. Hawke, R. Kobayashi, A. Verreault, A role for cell-cycle-regulated histone H3 lysine 56 acetylation in the DNA damage response. *Nature* **436**, 294–298 (2005). doi: [10.1038/nature03714](https://doi.org/10.1038/nature03714); pmid: [16015338](https://pubmed.ncbi.nlm.nih.gov/16015338/)
 55. M. Giannattasio, F. Lazzaro, P. Plevani, M. Muzi-Falconi, The DNA damage checkpoint response requires histone H2B ubiquitination by Rad6-Brel and H3 methylation by Dot1. *J. Biol. Chem.* **280**, 9879–9886 (2005). doi: [10.1074/jbc.M414453200](https://doi.org/10.1074/jbc.M414453200); pmid: [15632126](https://pubmed.ncbi.nlm.nih.gov/15632126/)
 56. I. Celic *et al.*, The SirTins Hst3 and Hst4p preserve genome integrity by controlling histone H3 lysine 56 deacetylation. *Curr. Biol.* **16**, 1280–1289 (2006). doi: [10.1016/j.cub.2006.06.023](https://doi.org/10.1016/j.cub.2006.06.023); pmid: [16815704](https://pubmed.ncbi.nlm.nih.gov/16815704/)
 57. I. Celic, A. Verreault, J. D. Boeke, Histone H3 K56 hyperacetylation perturbs replisomes and causes DNA damage. *Genetics* **179**, 1769–1784 (2008). doi: [10.1534/genetics.108.088914](https://doi.org/10.1534/genetics.108.088914); pmid: [18579506](https://pubmed.ncbi.nlm.nih.gov/18579506/)
 58. H. N. Du, S. D. Briggs, A nucleosome surface formed by histone H4, H2A, and H3 residues is needed for proper histone H3 Lys36 methylation, histone acetylation, and repression of cryptic transcription. *J. Biol. Chem.* **285**, 11704–11713 (2010). doi: [10.1074/jbc.M109.085043](https://doi.org/10.1074/jbc.M109.085043); pmid: [20139424](https://pubmed.ncbi.nlm.nih.gov/20139424/)
 59. S. Chen *et al.*, Structure-function studies of histone H3/H4 tetramer maintenance during transcription by chaperone Spt2. *Genes Dev.* **29**, 1326–1340 (2015). doi: [10.1101/gad.261115.115](https://doi.org/10.1101/gad.261115.115); pmid: [26109053](https://pubmed.ncbi.nlm.nih.gov/26109053/)
 60. H. G. Tran, D. J. Steger, V. R. Iyer, A. D. Johnson, The chromo domain protein chd1p from budding yeast is an ATP-dependent chromatin-modifying factor. *EMBO J.* **19**, 2323–2331 (2000). doi: [10.1093/emboj/19.10.2323](https://doi.org/10.1093/emboj/19.10.2323); pmid: [1081623](https://pubmed.ncbi.nlm.nih.gov/1081623/)
 61. Z. A. Chen *et al.*, Architecture of the RNA polymerase II-TFIIF complex revealed by cross-linking and mass spectrometry. *EMBO J.* **29**, 717–726 (2010). doi: [10.1038/emboj.2009.401](https://doi.org/10.1038/emboj.2009.401); pmid: [20094031](https://pubmed.ncbi.nlm.nih.gov/20094031/)
 62. S. Ovchinnikov, H. Kamisetty, D. Baker, Robust and accurate prediction of residue-residue interactions across protein interfaces using evolutionary information. *eLife* **3**, e02030 (2014). doi: [10.7554/eLife.02030](https://doi.org/10.7554/eLife.02030); pmid: [24842992](https://pubmed.ncbi.nlm.nih.gov/24842992/)
 63. Q. Cong, I. Anishchenko, S. Ovchinnikov, D. Baker, Protein interaction networks revealed by proteome coevolution. *Science* **365**, 185–189 (2019). pmid: [31296772](https://pubmed.ncbi.nlm.nih.gov/31296772/)
 64. T. Gueudré, C. Baldassi, M. Zamparo, M. Weigt, A. Pagnani, Simultaneous identification of specifically interacting paralogs and interprotein contacts by direct coupling analysis. *Proc. Natl. Acad. Sci. U.S.A.* **113**, 12186–12191 (2016). doi: [10.1073/pnas.1607570113](https://doi.org/10.1073/pnas.1607570113); pmid: [27729520](https://pubmed.ncbi.nlm.nih.gov/27729520/)
 65. S. Wang, S. Sun, Z. Li, R. Zhang, J. Xu, Accurate de novo prediction of protein contact map by ultra-deep learning model. *PLOS Comput. Biol.* **13**, e1005324 (2017). doi: [10.1371/journal.pcbi.1005324](https://doi.org/10.1371/journal.pcbi.1005324); pmid: [28056090](https://pubmed.ncbi.nlm.nih.gov/28056090/)
 66. H. Zeng *et al.*, ComplexContact: A web server for inter-protein contact prediction using deep learning. *Nucleic Acids Res.* **46**, W432–W437 (2018). doi: [10.1093/nar/gky420](https://doi.org/10.1093/nar/gky420); pmid: [29790960](https://pubmed.ncbi.nlm.nih.gov/29790960/)
 67. S. Wang, S. Sun, J. Xu, Analysis of deep learning methods for blind protein contact prediction in CASP12. *Proteins* **86** (suppl. 1), 67–77 (2018). doi: [10.1002/prot.25377](https://doi.org/10.1002/prot.25377); pmid: [28845538](https://pubmed.ncbi.nlm.nih.gov/28845538/)
 68. K. R. Roy *et al.*, Multiplexed precision genome editing with trackable genomic barcodes in yeast. *Nat. Biotechnol.* **36**, 512–520 (2018). doi: [10.1038/nbt.4137](https://doi.org/10.1038/nbt.4137); pmid: [29734294](https://pubmed.ncbi.nlm.nih.gov/29734294/)
 69. S. R. Collins *et al.*, Toward a comprehensive atlas of the physical interactome of *Saccharomyces cerevisiae*. *Mol. Cell. Proteomics* **6**, 439–450 (2007). doi: [10.1074/mcp.M600381-MCP200](https://doi.org/10.1074/mcp.M600381-MCP200); pmid: [17200106](https://pubmed.ncbi.nlm.nih.gov/17200106/)
 70. J. M. Schmiedel, B. Lehner, Determining protein structures using deep mutagenesis. *Nat. Genet.* **51**, 1177–1186 (2019). doi: [10.1038/s41588-019-0431-x](https://doi.org/10.1038/s41588-019-0431-x); pmid: [31209395](https://pubmed.ncbi.nlm.nih.gov/31209395/)
 71. N. J. Rollins *et al.*, Inferring protein 3D structure from deep mutation scans. *Nat. Genet.* **51**, 1170–1176 (2019). doi: [10.1038/s41588-019-0432-9](https://doi.org/10.1038/s41588-019-0432-9); pmid: [31209393](https://pubmed.ncbi.nlm.nih.gov/31209393/)
 72. R. W. Newberry, J. T. Leong, E. D. Chow, M. Kampmann, W. F. DeGrado, Deep mutational scanning reveals the structural basis for α -synuclein activity. *Nat. Chem. Biol.* **16**, 653–659 (2020). doi: [10.1038/s41589-020-0480-6](https://doi.org/10.1038/s41589-020-0480-6); pmid: [32125244](https://pubmed.ncbi.nlm.nih.gov/32125244/)
 73. E. Eyal, R. Najmanovich, M. Edelman, V. Sobolev, Protein side-chain rearrangement in regions of point mutations. *Proteins* **50**, 272–282 (2003). doi: [10.1002/prot.10276](https://doi.org/10.1002/prot.10276); pmid: [12486721](https://pubmed.ncbi.nlm.nih.gov/12486721/)
 74. R. Sasidharan, C. Chothia, The selection of acceptable protein mutations. *Proc. Natl. Acad. Sci. U.S.A.* **104**, 10080–10085 (2007). doi: [10.1073/pnas.0703737104](https://doi.org/10.1073/pnas.0703737104); pmid: [17540730](https://pubmed.ncbi.nlm.nih.gov/17540730/)
 75. J. Schaarschmidt, B. Monastyrsky, A. Kryshchak, A. M. J. J. Bonvin, Assessment of contact predictions in CASP12: Co-evolution and deep learning coming of age. *Proteins* **86** (suppl. 1), 51–66 (2018). doi: [10.1002/prot.25407](https://doi.org/10.1002/prot.25407); pmid: [29071738](https://pubmed.ncbi.nlm.nih.gov/29071738/)
 76. S. Ovchinnikov, H. Park, D. E. Kim, F. DiMaio, D. Baker, Protein structure prediction using Rosetta in CASP12. *Proteins* **86** (suppl. 1), 113–121 (2018). doi: [10.1002/prot.25390](https://doi.org/10.1002/prot.25390); pmid: [28940798](https://pubmed.ncbi.nlm.nih.gov/28940798/)
 77. P. J. Robinson *et al.*, Molecular architecture of the yeast Mediator complex. *eLife* **4**, e08719 (2015). doi: [10.7554/eLife.08719](https://doi.org/10.7554/eLife.08719); pmid: [26402457](https://pubmed.ncbi.nlm.nih.gov/26402457/)
 78. J. Luo *et al.*, Architecture of the human and yeast general transcription and DNA repair factor TFIIH. *Mol. Cell* **59**, 794–806 (2015). doi: [10.1016/j.molcel.2015.07.016](https://doi.org/10.1016/j.molcel.2015.07.016); pmid: [26340423](https://pubmed.ncbi.nlm.nih.gov/26340423/)
 79. M. Jinek *et al.*, A programmable dual-RNA-guided DNA endonuclease in adaptive bacterial immunity. *Science* **337**, 816–821 (2012). doi: [10.1126/science.1225829](https://doi.org/10.1126/science.1225829); pmid: [22745249](https://pubmed.ncbi.nlm.nih.gov/22745249/)
 80. J. P. Shen *et al.*, Combinatorial CRISPR-Cas9 screens for de novo mapping of genetic interactions. *Nat. Methods* **14**, 573–576 (2017). doi: [10.1038/nmeth.4225](https://doi.org/10.1038/nmeth.4225); pmid: [28319113](https://pubmed.ncbi.nlm.nih.gov/28319113/)
 81. D. Du *et al.*, Genetic interaction mapping in mammalian cells using CRISPR interference. *Nat. Methods* **14**, 577–580 (2017). doi: [10.1038/nmeth.4286](https://doi.org/10.1038/nmeth.4286); pmid: [28481362](https://pubmed.ncbi.nlm.nih.gov/28481362/)
 82. L. Ma *et al.*, CRISPR-Cas9-mediated saturated mutagenesis screen predicts clinical drug resistance with improved accuracy. *Proc. Natl. Acad. Sci. U.S.A.* **114**, 11751–11756 (2017). doi: [10.1073/pnas.1708268114](https://doi.org/10.1073/pnas.1708268114); pmid: [29078326](https://pubmed.ncbi.nlm.nih.gov/29078326/)
 83. A. V. Anzalone *et al.*, Search-and-replace genome editing without double-strand breaks or donor DNA. *Nature* **576**, 149–157 (2019). doi: [10.1038/s41586-019-1711-4](https://doi.org/10.1038/s41586-019-1711-4); pmid: [31634902](https://pubmed.ncbi.nlm.nih.gov/31634902/)

84. D. E. Gordon *et al.*, A SARS-CoV-2 protein interaction map reveals targets for drug repurposing. *Nature* **583**, 459–468 (2020). doi: [10.1038/s41586-020-2286-9](#); pmid: [32353859](#)
85. M. Eckhardt, J. F. Hultquist, R. M. Kaake, R. Hüttenhain, N. J. Krogan, A systems approach to infectious disease. *Nat. Rev. Genet.* **21**, 339–354 (2020). doi: [10.1038/s41576-020-0212-5](#); pmid: [32606427](#)
86. D. E. Gordon *et al.*, Comparative host-coronavirus protein interaction networks reveal pan-viral disease mechanisms. *Science* **370**, eabe9403 (2020). doi: [10.1126/science.abe9403](#); pmid: [33060197](#)
87. D. E. Gordon *et al.*, A quantitative genetic interaction map of HIV infection. *Mol. Cell* **78**, 197–209.e7 (2020). doi: [10.1016/j.molcel.2020.02.004](#); pmid: [32084337](#)
88. S. R. McGuffee, A. H. Elcock, Diffusion, crowding & protein stability in a dynamic molecular model of the bacterial cytoplasm. *PLoS Comput. Biol.* **6**, e1000694 (2010). doi: [10.1371/journal.pcbi.1000694](#); pmid: [20221255](#)
89. S. Takamori *et al.*, Molecular anatomy of a trafficking organelle. *Cell* **127**, 831–846 (2006). doi: [10.1016/j.cell.2006.10.030](#); pmid: [17110340](#)
90. B. G. Wilhelm *et al.*, Composition of isolated synaptic boutons reveals the amounts of vesicle trafficking proteins. *Science* **344**, 1023–1028 (2014). doi: [10.1126/science.1252884](#); pmid: [24876496](#)
91. J. Singla *et al.*, Opportunities and challenges in building a spatiotemporal multi-scale model of the human pancreatic β cell. *Cell* **173**, 11–19 (2018). doi: [10.1016/j.cell.2018.03.014](#); pmid: [29570991](#)
92. P. J. Thul *et al.*, A subcellular map of the human proteome. *Science* **356**, eaal3321 (2017). doi: [10.1126/science.aal3321](#); pmid: [28495876](#)
93. M. Schuldiner, S. R. Collins, J. S. Weissman, N. J. Krogan, Quantitative genetic analysis in *Saccharomyces cerevisiae* using epistatic miniarray profiles (E-MAPs) and its application to chromatin functions. *Methods* **40**, 344–352 (2006). doi: [10.1016/j.ymeth.2006.07.034](#); pmid: [17101447](#)
94. W. Rieping, M. Habeck, M. Nilges, Inferential structure determination. *Science* **309**, 303–306 (2005). doi: [10.1126/science.1110428](#); pmid: [16002620](#)
95. W. Rieping, M. Habeck, M. Nilges, Modeling errors in NOE data with a log-normal distribution improves the quality of NMR structures. *J. Am. Chem. Soc.* **127**, 16026–16027 (2005). doi: [10.1021/ja055092c](#); pmid: [16287280](#)
96. H. Jeffreys, An invariant form for the prior probability in estimation problems. *Proc. R. Soc. London Ser. A* **186**, 453–461 (1946). doi: [10.1098/rspa.1946.0056](#); pmid: [20998741](#)
97. A. Sali *et al.*, Outcome of the first wwPDB Hybrid/Integrative Methods Task Force Workshop. *Structure* **23**, 1156–1167 (2015). doi: [10.1016/j.str.2015.05.013](#); pmid: [26095030](#)
98. S. K. Burley *et al.*, PDB-Dev: A prototype system for depositing integrative/hybrid structural models. *Structure* **25**, 1317–1318 (2017). doi: [10.1016/j.str.2017.08.001](#); pmid: [28877501](#)
99. R. Chen, J. Mintseris, J. Janin, Z. Weng, A protein-protein docking benchmark. *Proteins* **52**, 88–91 (2003). doi: [10.1002/prot.10390](#); pmid: [12784372](#)
100. C. M. Wood *et al.*, High-resolution structure of the native histone octamer. *Acta Crystallogr. Sect. F Struct. Biol. Cryst. Commun.* **61**, 541–545 (2005). doi: [10.1107/S1744309105013813](#); pmid: [16511091](#)
101. A. Šali, T. L. Blundell, Comparative protein modelling by satisfaction of spatial restraints. *J. Mol. Biol.* **234**, 779–815 (1993). doi: [10.1006/jmbi.1993.1626](#); pmid: [8254673](#)
102. S. M. Vos *et al.*, Structure of activated transcription complex Pol II-DNA-PAF-SPT6. *Nature* **560**, 607–612 (2018). doi: [10.1038/s41586-018-0440-4](#); pmid: [30135578](#)
103. M. N. Wojtas, M. Mogri, O. Millet, S. D. Bell, N. G. A. Abrescia, Structural and functional analyses of the interaction of archaeal RNA polymerase with DNA. *Nucleic Acids Res.* **40**, 9941–9952 (2012). doi: [10.1093/nar/gks692](#); pmid: [22848102](#)
104. D. Wang, D. A. Bushnell, K. D. Westover, C. D. Kaplan, R. D. Kornberg, Structural basis of transcription: Role of the trigger loop in substrate specificity and catalysis. *Cell* **127**, 941–954 (2006). doi: [10.1016/j.cell.2006.11.023](#); pmid: [17129781](#)
105. K. S. Murakami, X-ray crystal structure of *Escherichia coli* RNA polymerase $\sigma 70$ holoenzyme. *J. Biol. Chem.* **288**, 9126–9134 (2013). doi: [10.1074/jbc.M112.430900](#); pmid: [23389035](#)
106. R. P. Joosten *et al.*, A series of PDB related databases for everyday needs. *Nucleic Acids Res.* **39**, D411–D419 (2011). doi: [10.1093/nar/gkq1105](#); pmid: [21071423](#)
107. W. Kabsch, C. Sander, Dictionary of protein secondary structure: Pattern recognition of hydrogen-bonded and geometrical features. *Biopolymers* **22**, 2577–2637 (1983). doi: [10.1002/bip.360221211](#); pmid: [6667333](#)
108. M.-Y. Shen, A. Sali, Statistical potential for assessment and prediction of protein structures. *Protein Sci.* **15**, 2507–2524 (2006). doi: [10.1110/ps.062416606](#); pmid: [17075131](#)
109. J. P. Erzberger *et al.*, Molecular architecture of the 40S-eIF1-eIF3 translation initiation complex. *Cell* **158**, 1123–1135 (2014). doi: [10.1016/j.cell.2014.07.044](#); pmid: [25171412](#)
110. Y. Shi *et al.*, Structural characterization by cross-linking reveals the detailed architecture of a coatomer-related heptameric module from the nuclear pore complex. *Mol. Cell. Proteomics* **13**, 2927–2943 (2014). doi: [10.1074/mcp.M114.041673](#); pmid: [25161197](#)
111. R. H. Swendsen, J. S. Wang, Replica Monte Carlo simulation of spin glasses. *Phys. Rev. Lett.* **57**, 2607–2609 (1986). doi: [10.1103/PhysRevLett.57.2607](#); pmid: [10033814](#)
112. S. Viswanath, I. E. Chemmama, P. Cimermancic, A. Sali, Assessing exhaustiveness of stochastic sampling for integrative modeling of macromolecular structures. *Biophys. J.* **113**, 2344–2353 (2017). doi: [10.1016/j.bpj.2017.10.005](#); pmid: [2921988](#)
113. B. Vallat, B. Webb, J. D. Westbrook, A. Sali, H. M. Berman, Development of a prototype system for archiving integrative/hybrid structure models of biological macromolecules. *Structure* **26**, 894–904.e2 (2018). doi: [10.1016/j.str.2018.03.011](#); pmid: [29657133](#)
114. J. D. Chodera, A simple method for automated equilibration detection in molecular simulations. *J. Chem. Theory Comput.* **12**, 1799–1805 (2016). doi: [10.1021/acs.jctc.5b00784](#); pmid: [26771390](#)
115. L. McInnes, J. Healy, S. Astels, hdbscan: Hierarchical density based clustering. *J. Open Source Softw.* **2**, 205 (2017). doi: [10.21105/joss.00205](#)
116. E. D. Merkle *et al.*, Distance restraints from crosslinking mass spectrometry: Mining a molecular dynamics simulation database to evaluate lysine-lysine distances. *Protein Sci.* **23**, 747–759 (2014). doi: [10.1002/pro.2458](#); pmid: [24639379](#)
117. D. Schneidman-Duhovny *et al.*, A method for integrative structure determination of protein-protein complexes. *Bioinformatics* **28**, 3282–3289 (2012). doi: [10.1093/bioinformatics/bts628](#); pmid: [23093611](#)
118. M. J. de Hoon, S. Imoto, J. Nolan, S. Miyano, Open source clustering software. *Bioinformatics* **20**, 1453–1454 (2004). doi: [10.1093/bioinformatics/bth078](#); pmid: [14871861](#)
119. A. J. Saldanha, Java Treeview—Extensible visualization of microarray data. *Bioinformatics* **20**, 3246–3248 (2004). doi: [10.1093/bioinformatics/bth349](#); pmid: [15180930](#)
120. J. H. Morris, C. C. Huang, P. C. Babbitt, T. E. Ferrin, structureViz: Linking Cytoscape and UCSF Chimera. *Bioinformatics* **23**, 2345–2347 (2007). doi: [10.1093/bioinformatics/btm329](#); pmid: [17623706](#)
121. N. T. Doncheva, K. Klein, F. S. Domingues, M. Albrecht, Analyzing and visualizing residue networks of protein structures. *Trends Biochem. Sci.* **36**, 179–182 (2011). doi: [10.1016/j.tibs.2011.01.002](#); pmid: [21345680](#)
122. J. H. Morris *et al.*, setsApp: Set operations for Cytoscape Nodes and Edges. *F1000 Res.* **3**, 149 (2014). doi: [10.12688/f1000research.4392.1](#); pmid: [25352980](#)
123. M. A. Collart, S. Oliviero, Preparation of yeast RNA. *Curr. Protoc. Mol. Biol.* **23**, 13.12.1–13.12.5 (2001). doi: [10.12688/f1000research.4392.1](#); pmid: [25352980](#)
124. D. Kim *et al.*, TopHat2: Accurate alignment of transcriptomes in the presence of insertions, deletions and gene fusions. *Genome Biol.* **14**, R36 (2013). doi: [10.1186/gb-2013-14-4-r36](#); pmid: [23618408](#)
125. S. Anders, P. T. Pyl, W. Huber, HTSeq—A Python framework to work with high-throughput sequencing data. *Bioinformatics* **31**, 166–169 (2015). doi: [10.1093/bioinformatics/btu638](#); pmid: [25260700](#)
126. M. I. Love, W. Huber, S. Anders, Moderated estimation of fold change and dispersion for RNA-seq data with DESeq2. *Genome Biol.* **15**, 550 (2014). doi: [10.1186/s13059-014-0550-8](#); pmid: [25516281](#)
127. M. Costanzo *et al.*, The genetic landscape of a cell. *Science* **327**, 425–431 (2010). doi: [10.1126/science.1180823](#); pmid: [20093466](#)
128. G. M. Wilmes *et al.*, A genetic interaction map of RNA-processing factors reveals links between Sem1/Dss1-containing complexes and mRNA export and splicing. *Mol. Cell* **32**, 735–746 (2008). doi: [10.1016/j.molcel.2008.11.012](#); pmid: [19061648](#)
129. Y. H. Benjamini, Y. Hochberg, Controlling the false discovery rate: A practical and powerful approach to multiple testing. *J. R. Stat. Soc. B* **57**, 289–300 (1995). doi: [10.1111/j.2517-6161.1995.tb02031.x](#)
130. D. Wessel, U. I. Flügge, A method for the quantitative recovery of protein in dilute solution in the presence of detergents and lipids. *Anal. Biochem.* **138**, 141–143 (1984). doi: [10.1016/0003-2697\(84\)90782-6](#); pmid: [6731838](#)
131. M. Downey *et al.*, Acetylome profiling reveals overlap in the regulation of diverse processes by sirtuins, gcn5, and esa1. *Mol. Cell. Proteomics* **14**, 162–176 (2015). doi: [10.1074/mcp.M114.043411](#); pmid: [25381059](#)
132. J. Cox, M. Mann, MaxQuant enables high peptide identification rates, individualized p.p.b.-range mass accuracies and proteome-wide protein quantification. *Nat. Biotechnol.* **26**, 1367–1372 (2008). doi: [10.1038/nbt1511](#); pmid: [19029910](#)
133. B. MacLean *et al.*, Skyline: An open source document editor for creating and analyzing targeted proteomics experiments. *Bioinformatics* **26**, 966–968 (2010). doi: [10.1093/bioinformatics/btq054](#); pmid: [20147306](#)
134. M. Choi *et al.*, MSstats: An R package for statistical analysis of quantitative mass spectrometry-based proteomic experiments. *Bioinformatics* **30**, 2524–2526 (2014). doi: [10.1093/bioinformatics/btu305](#); pmid: [24794931](#)
135. R. Dronamraju *et al.*, Spt6 association with RNA polymerase II directs mRNA turnover during transcription. *Mol. Cell* **70**, 1054–1066.e4 (2018). doi: [10.1016/j.molcel.2018.05.020](#); pmid: [29932900](#)
136. K. J. Livak, T. D. Schmittgen, Analysis of relative gene expression data using real-time quantitative PCR and the $2^{-\Delta\Delta Ct}$ Method. *Methods* **25**, 402–408 (2001). doi: [10.1006/meth.2001.1262](#); pmid: [11846609](#)
137. R. Dronamraju, B. D. Strahl, A feed forward circuit comprising Spt6, Ctk1 and PAF regulates Pol II CTD phosphorylation and transcription elongation. *Nucleic Acids Res.* **42**, 870–881 (2014). doi: [10.1093/nar/gkt1003](#); pmid: [24163256](#)
138. J. T. Robinson *et al.*, Integrative genomics viewer. *Nat. Biotechnol.* **29**, 24–26 (2011). doi: [10.1038/nbt.1754](#); pmid: [21221095](#)
139. Y. Liao, G. K. Smyth, W. Shi, featureCounts: An efficient general purpose program for assigning sequence reads to genomic features. *Bioinformatics* **30**, 923–930 (2014). doi: [10.1093/bioinformatics/btt656](#); pmid: [24227677](#)
140. R. Edgar, M. Domrachev, A. E. Lash, Gene Expression Omnibus: NCBI gene expression and hybridization array data repository. *Nucleic Acids Res.* **30**, 207–210 (2002). doi: [10.1093/nar/30.1.207](#); pmid: [11752295](#)

ACKNOWLEDGMENTS

We thank M. Eckhardt, I. Chemmama, A. Verreault, D. Truong, P. Beltrao, and A. Roguev for helpful discussions. **Funding:** We acknowledge funding through NIH grants R01 GM084448, R01 GM084279, P50 GM081879, P50 AI150476, P01 HL089707, U54 NS100717, and R01 GM098101 to N.J.K.; R01 GM083960, P50 GM082250, U19 AI135990, S10 OD021596, and P41 GM109824 to A.Sa.; U54 RR020839 and U54 GM103520 to J.D.B.; R35 GM118061 to C.G.; GM126900 to B.D.S.; P01GM105473 to J.E.H.; U54 RR020839 and R01 GM084448 to J.S.B.; OD009180 and GM123159 to J.S.F.; and GM008284 to A.Sh. We also acknowledge funding through the National Key Research and Development Program of China (2017YFA0505103), Shenzhen Science and Technology Program (KQTD20180413181837372), Guangdong Provincial Key Laboratory of Synthetic Genomics (2019B030301006), and Bureau of International Cooperation, Chinese Academy of Sciences (I72644KYSB20180022) to J.D. and the National Natural Science Foundation of China (31800069) to S.J. **Author contributions:** H.B., I.E., P.C., J.D.B., A.Sa., and N.J.K. designed the research; H.B., S.B., A.Sh., R.A., J.X., R.D., S.J., G.D., D.B., K.K.C., R.H., and S.W. performed experimental assays; I.E., P.C., R.P., and D.S. developed and performed integrative modeling; H.B., S.B., M.S., and D.M. analyzed experimental data; J.M. developed software; J.S.B., J.S.F., J.E.H., B.D.S., C.A.G., J.D., J.D.B., A.Sa., and

N.J.K. supervised the research; H.B., I.E., S.B., A.Sa., and N.J.K. wrote the paper with input from all authors. **Competing interests:** J.D.B. is a founder and director of CDI Labs, Inc.; a founder of Neochromosome, Inc.; and a founder and SAB member of ReOpen Diagnostics. J.D.B. also serves or served on the scientific advisory board of Sangamo, Inc.; Modern Meadow, Inc.; Sample6, Inc.; and the Wyss Institute. J.S.B. is a founder and director of Neochromosome, Inc. **Data and materials availability:** IMP modeling scripts, input data, and output results are available at <https://integrativemodeling.org/systems/pemap>. IMP is an open-source program, available at <http://integrativemodeling.org>. The clustered pE-MAP and correlation maps are available in the supplementary data and can be visualized using the Java Treeview app (119), available at <http://jtreeview.sourceforge.net/>. The stEMAP app source is publicly available at <https://github.com/RBVI/StEMAPApp>, and stEMAP will be available through the

Cytoscape App store. stEMAP requires installation of Cytoscape and ChimeraX. Cytoscape is available at <https://cytoscape.org/download.html> and ChimeraX at www.rbvi.ucsf.edu/chimera/download.html. stEMAP also requires installation of the Sets, RINalyzer, StructureVizX, clusterMaker2, and ctdReader apps, all available through the Cytoscape apps menu. The modified version of PDB 1ID3 (including the unstructured H3 and H4 tails for illustrative purposes) is available in the supplementary data. The yeast RNAPII pE-MAP is available in the supplementary materials of (15). The bacterial RNAP point mutation data are available in the supplementary materials of (20). The library of histone mutant strains is available from J. Dai under a material transfer agreement with the Shenzhen Institutes of Advanced Technology, Chinese Academy of Sciences. RNA-seq and ATAC-seq data have been deposited in NCBI's Gene Expression Omnibus (140) and are accessible through GEO

Series accession number GSE160350 (www.ncbi.nlm.nih.gov/geo/query/acc.cgi?acc=GSE160350).

SUPPLEMENTARY MATERIALS

science.sciencemag.org/content/370/6522/eaaz4910/suppl/DC1

Figs. S1 to S11

Tables S1 to S9

References (141–144)

MDAR Reproducibility Checklist

Data S1 to S3

[View/request a protocol for this paper from Bio-protocol.](#)

14 October 2019; resubmitted 23 July 2020

Accepted 23 October 2020

10.1126/science.aaz4910