

Technical Guide

Getting started with ChIP-seq: experimental design to data analysis

Introduction to ChIP-Seq

Chromatin-immunoprecipitation (ChIP) followed by next generation sequencing (ChIP-seq) of the immunoprecipitated DNA is a powerful tool for the investigation of protein:DNA interactions. To perform ChIP-seq, chromatin is isolated from cells or tissues (with or without chemical crosslinking) and fragmented. Antibodies recognizing chromatin-associated proteins of interest are used to enrich the sample for specific chromatin fragments. The DNA is recovered, sequenced on various NGS platforms, and aligned to a reference genome to determine specific protein binding loci. ChIP-seq studies have increased our knowledge of transcription factor biology, DNA methylation and histone modifications.

In this guide, we will introduce ChIP-seq and outline key steps of the experimental process, including:

- Experimental design
- Controls for ChIP-seq experiments
- Reference genome alignment of ChIP-seq reads (mapping)
- Background estimation
- Peak finding
- Quality control of ChIP-seq experiments
- Differential binding analysis
- Motif analysis

ChIP-seq was first described in 2007¹. ChIP-seq was among the first methods to make use of the power of massively parallel (or next-generation) sequencing (NGS) to significantly advance genome wide coverage relative to existing real-time PCR and array-based methods. ChIP-seq is a counting assay that uses only short reads to align to the genome, but requires millions of them to provide meaningful data. Fortunately, the Solexa® 1G NGS instrument (now a part of Illumina) provided up to 30 million 21-35 bp reads per run. Current NGS systems, such as those being developed by Illumina and Life Technologies' Ion Torrent™ Platform, are producing longer and deeper reads. However, most users still produce single-end 35 bp reads, albeit generating up to 1.5 billion reads of individually barcoded samples per run, as is typically obtained using a Illumina HiSeq™ single flow cell system.

This advance in technology has allowed projects like the Encyclopedia of DNA Elements (ENCODE) to generate almost 1250 ChIP-seq datasets². The ENCODE consortium also put significant efforts into standardizing experimental procedures. For instance, they produced a set of working standards and reporting guidelines designed to test that an antibody is specific to its antigen and has minimal cross-reactivity to other proteins³. Other groups have published comprehensive method descriptions to which we refer readers if they would like a more detailed description of the experimental steps⁴.

ChIP-seq may have evolved from microarray analysis, but it has required the development of a completely new set of analysis tools to make the most of the platform. ChIP-seq analysis begins with mapping of trimmed unique sequence reads to a reference genome. Next, peaks are found using peak-calling algorithms. To further analyze the data, differential binding or motif analyses are included as common endpoints of ChIP-seq workflows. At every stage, the choice of method or algorithm and the parameters used affect the downstream results.

Further complicating analysis options is the fact that ChIP-seq experiments can be divided into different classes⁵. Some experiments produce clearly defined peaks of 100-200 bp, as typified by transcription factors like estrogen receptor α . Other experiments, such as those investigating H3K27me3 binding, produce wider "smears" or regions composed of sequences that range between a few to several kilobases. Lastly, some experiments produce a mix of clearly defined peaks and wider smears, such as ChIP-seq using antibodies to RNA polymerase II. Most algorithms have been developed for analysis of clearly defined peaks, because these present the greatest opportunity to determine nucleotide resolution of transcription factor binding and motif analysis.

Experimental Design

All experiments should be designed to meet the goals of the study and make best use of the resources available. Novices to ChIP-seq, or investigators that rely on outside sources for sequencing and data analysis, should consult with a bioinformaticist to ensure that proper model system setup, controls, experimental parameters and data formats are in place prior to beginning a ChIP-seq project.

When choosing which sequencing platform to use, although any NGS platform will work, most ChIP-seq users are concerned with generating as many reads as they can, as cheaply as possible. Most experiments require 5-10 million reads, widely considered the minimum, and many users regularly generate 20-40 million reads. Biological replication is another important consideration in planning the scale of the experiment. Biological replication reveals variation within sample groups and enables the analysis of differential binding⁶.

Controls for ChIP-Seq Experiments

Two types of controls are often used in ChIP-seq studies, primarily because DNA fragmentation by sonication is not a truly random process. An "input" DNA sample is one that has been cross-linked and sonicated but not immunoprecipitated. An IgG "mock"-ChIP uses an antibody that will not bind to nuclear proteins. Any IgG immunoprecipitated DNA should represent a random, nonspecific population. Because "mock" ChIPs often produce relatively little amplifiable DNA, input controls are more widely used to normalize signals from ChIP enrichment.

Reference Genome Alignment of ChIP-Seq Reads (Mapping)

The millions of reads generated in each experiment need to be analyzed and that analysis begins with alignment to a reference genome. The SEQanswers SEQwiki (<http://seqanswers.com/wiki/Software>), which hosts a table of common tools for ChIP-seq analysis, lists 94 tools with sequence alignment capabilities.

The most widely used for ChIP-seq have been ELAND, MAQ⁷ and Bowtie⁸. Mapping is generally performed while allowing for a small number (1-3) of sequence mismatches.

Different alignment algorithms trade speed for quality of the final alignment. This tradeoff is partly determined by how each algorithm uses quality values in the sequence data or aligns sequences to more repetitive regions of the genome.

- MAQ makes use of the sequence quality values, so that a mismatch at low quality bases is treated differently from a mismatch at high quality bases, assuming that a low quality base-call is more likely a sequencing error.
- Bowtie is one of the fastest mapping algorithms. The algorithms also differ in their handling of reads that map to multiple locations, positioning them randomly or arbitrarily. If ChIP-seq experiments recover sequences from highly repetitive regions, then the use of paired-end sequencing presents the opportunity to anchor read-pairs in a non-repeat region of the genome, thereby increasing confidence in the final mapping.

Background Estimation

ChIP-seq generates sequences from regions specifically, or indirectly, bound to the antibody target (the signal) as well as from background binding of genomic DNA and regions non-specifically bound to the antibody (the noise). Consequently, ChIP-seq libraries need to be sufficiently complex, consisting of billions of unique molecules with distinct 3' and 5' ends.

Even high-quality ChIP-seq libraries can contain high levels of noise relative to signal. Thus, peak-calling becomes a signal-to-noise problem. The choice of analysis algorithm and parameters also affects the specificity and sensitivity of the experiment. Mapped reads used for downstream analysis can be restricted to reads that map to unique genome regions only (high specificity) or can include reads that are more "promiscuous," mapping to multiple sites in the genome (high sensitivity). Of note, the complexity of the library and noise can be influenced by the size of fragments of ChIP'd DNA. Smaller fragments are more readily clonable; therefore, complexity increases when chromatin is highly fragmented.

Peak Finding

Probably the most discussed issue in ChIP-seq experiments is the best method for finding true "peaks" in the data. A peak is a site where multiple reads have mapped, producing a pileup. ChIP sequencing is most often performed with single-end reads, and ChIP fragments are sequenced from their 5' ends only. This creates two distinct peaks; one on each strand with the binding site falling in the middle of these peaks. The distance from the middle of the peaks to the binding site is often referred to as the "shift".

A good understanding of ChIP fragment size helps in locating the specific nucleotide-resolved binding site. This can be done in the wet lab by gel-based methods; alternatively, paired-end sequence data allow the fragment size to be calculated directly from the data. This suggests that a mix of primarily single-end reads with a small percentage of paired-end reads could provide the best data set for analysis. The large number of different peak-finding software programs is a testament to the importance of finding true peaks in ChIP-seq datasets. Choosing the best is almost impossible; a comparison of eleven different peak detecting algorithms did not show that any one algorithm exhibited overall superiority⁹. The parameters chosen for peak calling can significantly affect the outcomes, so care must be taken that data sets are analyzed using the same methods.

What Do Peak Callers Do?

Peak calling programs help to define sites of Protein:DNA binding by identifying regions where sequence reads are enriched in the genome after mapping. The common assumption is that the ChIP-seq process is relatively unbiased so reads should accumulate at sites of protein binding more frequently than in background regions of the genome. The millions of sequencing reads generated in a ChIP-seq experiment are first aligned to a reference genome using tools such as BWA⁷ and Bowtie⁸. The choice of alignment algorithm and the parameters used can impact peak calling. The number of mismatches allowed can affect the percentage of sequences that can be successfully aligned and the use and placement of reads that map to multiple locations e.g. in repeat regions, can mask true binding events. It is important to understand if and how the alignment algorithm and peak calling algorithm will work together.

Peak calling requires that several distinct analyses be carried out to generate the final peak list: read shifting, background estimation, identification of enriched peaks, significance analysis and removal of artifacts. A 2009 review by Pepke et al. details each of these steps and discusses how peak finding tools approach the separate steps very differently⁵. A follow-up review by Wilbanks et al. evaluated the performance of 11 ChIP-seq peak callers, nearly all of which are still widely used today⁹. Each step can have parameters that can be adjusted by the user, but changing these can significantly affect the final peak lists. Care must be taken that data sets are analyzed using the same methods. The ENCODE consortium produced guidelines for analysis of the dispersed data sets to avoid issues created by analysis parameter differences³. This project used MACS, PeakSeq and SPP.

Read shifting

The aligned reads are from fragments of 150-300 bp in length and, as most ChIP-seq data is from single-end sequencing, only one end of a fragment is read. Reads therefore align to either the sense or antisense strands and the 3' or 5' extremes of the DNA fragments pulled down in the immunoprecipitation. The reads are shifted and the data from both strands combined to determine the bases most likely to be involved in protein binding. How much to "shift" is determined by the fragment size generated in the ChIP-seq library preparation; this can be determined empirically or estimated from the sequence data. Comparison of these two measurements can be an effective quality control, as can the ratio of reads from different strands, where one would expect the ratio to be close to 1.

Background estimation

Control ChIPs are processed in the same way to allow either a genomic background to be determined (input controls) or for regions enriched through the ChIP process with no antibody specificity to be identified (IgG controls). Some peak callers work without control data and assume an even background signal, others make use of blacklist tools (such as RepeatMasker and the "Duke excluded regions" list developed for ENCODE), which mask regions of the genome.

Peak identification

A peak is called where either the number of reads exceeds a pre-determined threshold value or where there is a minimum enrichment compared to background signal, often in a sliding window across the genome. Some tools apply both methods. The parameters for identifying peaks can be adjusted, sometimes leading to very different numbers of peaks being called. The user must determine if fewer high-quality peaks are preferred over lower-quality peaks.

Significance analysis

Many peak callers compute a P value for called peaks, while others use the height of the peaks and/or enrichment over background to rank peaks. However, these calculations do not provide statistical significance values. The false discovery rate (FDR) is often used to provide a truer peak list, and this can be computed from the P values provided. Some packages make use of the control data to determine an empirical FDR and generate a ratio of peaks in the sample vs. control.

Artifact removal

Two major classes of ChIP-seq artifacts are generally removed before the final peak list is used in downstream applications. First, peaks containing either a single read, or just a few reads are assumed to be PCR amplification artifacts and discarded. Second, peaks in which there is a significant imbalance between the numbers of reads on each strand are removed. This second filtering is more difficult in complex regions where binding may be occurring at multiple co-located sites.

Unfortunately ChIP-seq does have biases, but these are gradually being understood. In experiments where de-proteinized, sheared, non-cross-linked DNA was used as the template for ChIP-seq studies, it has been possible to identify some of the factors affecting background noise¹⁰. The authors of this study also developed a model-based approach called MOSAiCS (MOdel-based one and two Sample Analysis and inference for ChIP-Seq Data) to find peaks more reliably, although this has not yet been widely adopted.

Choosing a Peak Calling Algorithm

There does not appear to be a clear winner among the thirty or more peak calling algorithms available today. Ask ten bioinformaticians which is best, and you will likely get ten different answers. The answer very much depends on the type of experiment being analyzed. Some peak callers, such as MACS, are better for studying transcription factor binding, while others, such as SICER, produce more reliable data for long-range interactions like polymerase binding. However a large factor influencing the success of peak-calling software is user experience. Many peak callers have multiple parameters that can affect the number of peaks called, and understanding these parameters takes time. Once users become comfortable with a particular setup, many are unlikely to change parameters. This is perhaps one of the reasons that MACS is still dominant. Although it is one of the oldest peak callers, it compares well to newer tools and many people have experience with it.

How Do Peak Callers Compare

Papers that compare the various peak calling algorithms are typically out of date the moment they are written, let alone published, but they point out important areas for consideration. The comparison methods used in different papers could be usefully updated and presented in a non-static electronic format. The winners, as far as the number of citations the primary publication received to date are E-Range (from the Wold group at Caltech), ChIP-seq peak finder (from the Genome Institute of Singapore) and MACS (from the Liu lab at Dana Farber), with over 4000 citations between them. However, these are also three of the oldest packages, released in the early days of ChIP-seq analysis.

At least one group has tried to produce benchmark datasets that can be used for comparison of peak callers¹¹. One of their aims was to provide datasets that were independent of those used to develop analysis tools, making an unbiased comparison easier. Their analysis of five programs showed that control data were essential for reduction of false-positive peaks, but that even without this, a manual visual inspection allowed 80% of false-positives to be removed, suggesting that the shape of the peaks could be used to improve analysis methods. They suggested a meta-approach that used features from four of the programs tested, which gave improved results for the benchmark dataset. Other groups have also suggested a multi-tool approach using several peak callers to generate consensus peak lists.

ChIP-Seq Peak Callers

Rather than giving a detailed description of all peak-finding packages, here we have picked four: MACS, which is one of the most popular tools, and three others that offer something different compared the majority of programs. A more comprehensive list of current packages can be found below.

MACS

MACS, used often for transcription factor binding site peaks, is one of the most popular peak callers, it is also one of the oldest, and its age probably contributes to its success. It is a good method, good enough for many experimental conditions and requires very little justification if cited as the tool used in a publication. MACS removes redundant reads and performs read-shifting to account for the offset in forward or reverse strand reads. It uses control samples and local statistics to minimize bias and calculates an empirical FDR.

SICER

Not all ChIP-seq users are interested in the "peaky" data as seen with transcription factors. However, nearly all peak callers were developed for exactly this kind of data. SICER was developed for more diffuse chromatin modifications that can span kilobases or megabases of the genome. The SICER method scans the genome in windows and identifies clusters of spatial signals that are unlikely to appear by chance. These clusters or "islands" are used instead of fixed length windows. Gaps in the islands are allowed in order to overcome technical issues (under-saturated experiments, repeat regions, etc). This gap size can be adjusted for different types of chromatin modifications. The program makes use of control data or a random background model¹².

T-PIC

This package uses the shape of putative peaks to identify true peaks from the background noise. The authors compared their approach to MACS and PeakSeq and demonstrated improved results. The package first extends short reads to the estimated fragment length; it then divides the genome into regions for which it constructs "trees" for shape analysis and uses the tree shape statistic to identify true peaks¹³.

Genome wide event finding and motif discovery (GEM): This is one of the newest tools, published in mid-2012. Its unique feature is the combination of peak finding and motif analysis to improve the resolution of the final peaks called. The published report presents an analysis of 63 transcription factors in 214 ENCODE experiments and demonstrates improved spatial resolution and motif

discovery when compared to previous tools. The tool also allows discovery of spatially-constrained binding events, which was demonstrated using the well-understood Sox2-Oct4 transcription factor complex. The GEM publication presents almost 400 spatially-constrained transcription factor binding events. This tool appears to be an exciting development for ChIP-seq studies.

Other ChIP-Seq Peak Callers

- AREM
- BayesPeak
- CEAS
- ChIP-Peak
- CisGenome
- CSDeconv
- E-RANGE: Dual-use package for RNA-seq and ChIP-seq, it is based on the ChIP-Seq mini peak finder published by the Wold group in 2007.
- EpiChip
- F-Seq
- FindPeaks: Is part of the Vancouver Short Read Analysis Package.
- HPeak
- MOSAiCS
- PeakSeq: Corrects for mappability and GC content biases to generate more accurate peak calls
- QuEST
- SIPeS: Uses paired-end data.
- SISRIS: Directional tool that identifies where reads "strand-shift" and can generate precise calls for sharp peaks. It is not very useful if you are interested in broader ChIP signals.
- Sole-Search
- SPP: Accounts for the read offset and read-shifts to improve results. The package makes use of background or control data, and estimate read saturation allowing the user to determine if more reads are required or not.
- SWEMBL
- Useq

Quality Control Of Chip-Seq Experiments

After sequencing, mapping and peak finding, several quality controls can be used to determine if further investigation and, ultimately, validation of the data are worthwhile. Packages such as FastQC allow raw sequence quality to be assessed. Read count enrichment can be calculated between ChIP and input samples and can help control for biases in the experimental methods. Finally, visual inspection of the data is a simple but effective tool for assessment of data quality.

Differential Binding Analysis

A relatively new technique is the analysis of differential binding, which draws much from the analysis of differential gene expression and has similar power to detect biologically meaningful binding changes between sample¹⁴. The DiffBind software package allows identification of genomic loci that are differentially bound between two conditions. It was developed based on algorithms used for differential gene expression analysis by RNA-seq. These differential methods allow researchers to assess ChIP peaks quantitatively using peak heights. Key to these methods is the normalization of read counts in ChIP-seq datasets and quantile normalization methods similar to those used in microarray analysis.

Motif Analysis

One of the most common aims of ChIP-seq experiments is to discover the sequence motifs for protein binding in the genome. The Multiple EM for Motif Elicitation (MEME) algorithm is the most widely adopted tool for motif discovery¹⁵. Often, multiple motifs can be found in a single data set, and motif analysis can be performed even on low quality ChIP-seq data, although the statistical significance of these motifs is likely to be lower.

Chromatin State

Another useful analysis of ChIP-seq data comes from a systematic approach used by the ENCODE consortium to characterize genomic regions based on histone modification content³. Various histone modifications are assayed using modification-specific histone antibodies in ChIP-seq experiments to obtain a profile of that histone mark within a sample. For its own experiments, the consortium has implemented rigorous specificity tests that use arrays of differentially modified histone tail peptides to ensure antibody specificity. They also share common cell sources which are collectively profiled and compared, ensuring consistency between individual experiments. Their current guidelines cover antibody validation, experimental replication, sequencing depth, data and metadata reporting, and data quality assessment¹⁷. You can access this information through the Human Epigenome Browser at Washington University¹⁸.

Summary

ChIP-seq is a powerful method and is yielding new biological insights¹⁶. Because of increased access to next generation sequencing platforms, ChIP-seq has almost entirely displaced earlier methods to investigate protein:DNA interactions. Being able to analyze these interactions genome-wide has increased our understanding of transcription factor biology, chromatin modification and transcription. In this article, we have attempted to present a broad overview of the major issues that need to be considered when designing and executing ChIP-seq experiments. Laboratory methods are now standardized, and kits such as Merck Millipore's Magna ChIP-seq™ chromatin IP and next generation library construction kit make it possible for virtually any lab to perform ChIP and construct an NGS library.

Although the focus of most current ChIP-seq experiments is on detecting the more dispersed class of protein:DNA interactions and on discovery of statistically significant differential binding, significant efforts are still underway to develop new analysis methods to enable improved analysis. Projects like ENCODE are showing that it is possible to produce very large data sets, as long as experiments are carefully controlled, while at the same time developing useful quality control metrics, analysis methods and parameters for the community.

References

- Robertson, G. et al. Genome-wide profiles of STAT1 DNA association using chromatin immunoprecipitation and massively parallel sequencing. *Nat Methods*. 2007 Aug; 4(8):651-7)
- Maier B. ENCODE: The human encyclopaedia. *Nature* 2012 Sep 6; 489(7414):46-8)
- Landt SG, et al. ChIP-seq guidelines and practices of the ENCODE and modENCODE consortia. *Genome Res*. 2012 Sep;22(9):1813-31
- Schmidt D et al. ChIP-seq: using high-throughput sequencing to discover protein-DNA interactions. *Methods* 2009 Jul;48(3):240-8).
- Pepke S, Wold B and Mortazavi A. Computation for ChIP-seq and RNA-seq studies. *Nature Methods*: 2009 Nov;6(11 Suppl):S22-32).
- Ross-Innes CS et al. Differential oestrogen receptor binding is associated with clinical outcome in breast cancer. *Nature*: 2012 481(7381): 389-393).
- Li H, Ruan J, Durbin R. Mapping short DNA sequencing reads and calling variants using mapping quality scores. *Genome Res*. 2008 Nov;18(11):1851-8.
- Langmead B, Trapnell C, Pop M, Salzberg SL. Ultrafast and memory-efficient alignment of short DNA sequences to the human genome. *Genome Biol*. 2009;10(3):R25.
- Wilbanks EG, Facciotti MT. Evaluation of algorithm performance in ChIP-seq peak detection. *PLoS One*. 2010 Jul 8;5(7):e11471.
- Kuan PF et al. A statistical framework for the analysis of ChIP-Seq data. *Journal of the American Statistical Association*. 2011; 106(495):891-903.
- Rye MB, Sætrum P, Drabløs F. A manually curated ChIP-seq benchmark demonstrates room for improvement in current peak-finder programs. *Nucleic Acids Res*. 2011 Mar;39(4):e25.
- Zang C, Schones DE, Zeng C, Cui K, Zhao K, Peng W. A clustering approach for identification of enriched domains from histone modification ChIP-Seq data. *Bioinformatics*. 2009 Aug 1;25(15):1952-1958.
- Hower V, Evans SN, Pachter L. Shape-based peak identification for ChIP-Seq. *BMC Bioinformatics*. 2011 Jan 12;12:15.
- Bardet AF, He Q, Zeitlinger J, Stark A. A computational pipeline for comparative ChIP-seq analyses. *Nat Protoc*. 2011 Dec 15;7(1):45-61.
- Bailey TL and Elkan C. Unsupervised learning of multiple motifs in biopolymers using expectation maximization. *Machine Learning*. 1995; 21(1-2): 51- 83.
- Rhee HS, Pugh BF. Comprehensive genome-wide protein-DNA interactions detected at single-nucleotide resolution. *Cell*. 2011 Dec 9;147(6):1408-19.
- Bernstein BE, Stamatoyannopoulos JA, Costello JF, Ren B, Milosavljevic A, Meissner A, Kellis M, Marra MA, Beaudet AL, Ecker JR, Farnham PJ, Hirst M, Lander ES, Mikkelsen TS, Thomson JA. The NIH Roadmap Epigenomics Mapping Consortium. *Nat Biotechnol*. 2010 Oct;28(10):1045-8.
- Zhou X et al. The Human Epigenome Browser at Washington University. *Nat Methods*. 2011 Nov 29;8(12):989-90.

Ordering Information

ChIP-Seq Kits, Antibodies and Reagents

Description	Cat. No.
Magna ChIP-Seq™ Chromatin Immunoprecipitation and Next Generation Sequencing Library Preparation Kit	17-1010
Magna ChIP® HiSens Chromatin Immunoprecipitation Kit	17-10460
EZ-Magna ChIP™ HiSens Chromatin Immunoprecipitation Kit	17-10461
Magna ChIP® A/G Chromatin Immunoprecipitation Kit	17-10085
EZ Magna ChIP™ A/G Chromatin Immunoprecipitation Kit	17-10086
Magna ChIP® Protein A+G Magnetic Beads	16-663
Magna ChIP® Protein A Magnetic Beads	16-661
Magna ChIP® Protein G Magnetic Beads	16-662
PureEpi™ Chromatin Preparation and Optimization Kit	17-10082

ChIP-Seq Qualified Antibodies

For a complete listing, visit: www.merckmillipore.com/antibodies

Description	Cat. No.
Anti-acetyl-Histone H3 (Lys 4)	07-539
Anti-acetyl-Histone H3 (Lys14)	07-353
Anti-acetyl-Histone H3 (Lys14), clone 13HH3-1A5	MABE351
Anti-acetyl-Histone H3	06-599
Anti-acetyl-Histone H4 (Lys12)	07-595
Anti-acetyl-Histone H4 (Lys16)	07-329
Anti-acetyl-Histone H4 (Lys5)	07-327
Anti-acetyl-Histone H4 (Lys8)	07-328
Anti-Androgen Receptor, PG-21	06-680
Anti-CTCF	07-729
Anti-dimethyl-Histone H3 (Lys27)	07-452
Anti-dimethyl-Histone H3 (Lys4)	07-030
Anti-E2F-4, clone GG22-2A6	05-312
Anti-EZH2	07-689
Anti-Histone H4, pan, clone 62-141-13	05-858
Anti-Methylcytosine dioxygenase TET1	09-872
Anti-monomethyl-Histone H3 (Lys27)	07-448
Anti-monomethyl-Histone H3 (Lys4)	07-436
Anti-Myc Tag, clone 4A6	05-724
Anti-phospho (Ser10)-acetyl (Lys14)-Histone H3	07-081
Anti-phospho-H2A.X (Ser139)	07-164
Anti-phospho-Histone H3 (Ser10), clone MC463	04-817
Anti-RNA polymerase II, clone CTD4H8	05-623
Anti-trimethyl-Histone H3 (Lys27)	07-449
Anti-trimethyl-Histone H3 (Lys4)	07-473
Anti-trimethyl-Histone H3 (Lys4), clone MC315	04-745
Anti-trimethyl-Histone H3 (Lys9), clone 6F12-H4	05-1242
ChIPAb+™ EZH2, clone AC22	17-662
ChIPAb+™ Trimethyl-Histone H3 (Lys4)	17-614

For our newest kits, assays, antibodies and proteins for ChIP, ChIP-seq and other epigenetics applications, visit: www.merckmillipore.com/epigenetics



www.merckmillipore.com/offices

To Place an Order or Receive Technical Assistance

In Europe, please call Customer Service:

France: 0825 045 645

Germany: 01805 045 645

Italy: 848 845 645

Spain: 901 516 645 Option 1

Switzerland: 0848 645 645

United Kingdom: 0870 900 4645

For other countries across Europe,
please call: +44 (0) 115 943 0840

Or visit: www.merckmillipore.com/offices

For Technical Service visit:

www.merckmillipore.com/techservice

Get Connected!

Join Merck Millipore Bioscience on your favorite social media outlet for the latest updates, news, products, innovations, and contests!



facebook.com/MerckMilliporeBioscience



twitter.com/Merck4Bio