

Reviewed Preprint

Published from the original preprint after peer review and assessment by eLife.

About eLife's process

Reviewed preprint version 1

November 8, 2023 (this version)

Posted to preprint server

July 29, 2023

Sent for peer review

July 11, 2023

Interpreting roles of mutations associated with the emergence of *S. aureus* USA300 strains using transcriptional regulatory network reconstruction

Saugat Poudel, Jason Hyun, Ying Hefner, Jon Monk, Victor Nizet, Bernhard O. Palsson 

Department of Bioengineering, University of California San Diego • Palmona Pathogenomics • Collaborative to Halt Antibiotic-Resistant Microbes (CHARM), Department of Pediatrics, University of California San Diego • Department of Pediatrics, University of California San Diego

 https://en.wikipedia.org/wiki/Open_access

 Copyright information

Abstract

The *Staphylococcus aureus* clonal complex 8 (CC8) is made up of several subtypes with varying levels of clinical burden; from community-associated methicillin resistant *S. aureus* (CA-MRSA) USA300 strains to hospital-associated (HA-MRSA) USA500 strains and basal methicillin susceptible (MSSA) strains. This phenotypic distribution within a single clonal complex makes CC8 an ideal clade to study the emergence of mutations important for antibiotic resistance and community spread. Gene level analysis comparing USA300 against MSSA and HA-MRSA strains have revealed key horizontally acquired genes important for its rapid spread in the community. However, efforts to define the contributions of point mutations and indels have been confounded by strong linkage disequilibrium resulting from clonal propagation. To break down this confounding effect, we combined genetic association testing with a model of the transcriptional regulatory network (TRN) to find candidate mutations that may have led to changes in gene regulation. First, we used a De Bruijn graph genome-wide association study (DBGWAS) to enrich mutations unique to the USA300 lineages within CC8. Next, we reconstructed the TRN by using Independent Component Analysis on 670 RNA sequencing samples from USA300 and non-USA300 CC8 strains which predicted several genes with strain-specific altered expression patterns. Examination of the regulatory region of one of the genes enriched by both approaches, *isdH*, revealed a 38 base pair deletion containing a Fur binding site and a conserved SNP which likely led to the altered expression levels in USA300 strains. Taken together, our results demonstrate the utility of reconstructed TRNs to address the limits of genetic approaches when studying emerging pathogenic strains.

eLife assessment

This study presents **useful** insights into core genome mutations that could have contributed to the emergence of the *Staphylococcus aureus* lineage USA300, a frequent cause of community-acquired infections. The **solid** approach used is innovative in combining genome-wide association studies and RNA-expression analyses, both applied to extensive publicly available datasets. This strategy reduces the rate of false positives attributed to high genome-wide linkage disequilibrium. It is noted that this method cannot be used for most phenotype-genotype studies, especially those requiring essential population structure correction, and it can therefore not be readily replicated in different datasets.

Introduction

Comparative genomic methods represent an important approach to understand the emergence and evolution of new strains of pathogens. In *S. aureus* alone, whole genome comparisons have enabled rapid characterization of genetic basis for antibiotic resistance, increased virulence, host specificity and altered metabolic capabilities^{1,2,3,4,5}. However, genome-wide linkage disequilibrium and strong lineage structuring currently limits the differentiation of causative alleles from genetically linked ones. By calculating lineage level associations, methods like bugwas address these issues for single, recurring phenotypes like antibiotic resistance⁶.

Endemic strains, on the other hand, exhibit multiple complex phenotypes that may contribute to their emergence and proliferation. For example, USA300 strains carry antibiotic resistance cassettes, Panton Valentine Leukocidin (PVL) associated with pyomastitis, increased ability to colonize locations outside of the nasopharynx etc.^{1,7,8,9}. As these strains often emerge clonally from closely related 'basal strains,' efforts to discern causal mutations that lead to their increased clinical burden is hampered by strong population-stratification and genome-wide linkage disequilibrium^{9,10,11}. Though recombination at species level is common in *S. aureus*, within clade recombination rates tend to be lower, thus preserving the linkage between mutations^{10,12,13,14}. Due to this limitation, studies of emerging strains often focus on gene level analysis such as acquisition of mobile genetic elements or loss of gene function while determining the possible phenotypic effect (if any) of all enriched Single Nucleotide Polymorphisms (SNPs) remains challenging^{15,16}. Computational modeling methods can help tackle these challenges by predicting phenotype differences between strains, thus acting as a sieve to filter enriched mutations with potential phenotypic effects and therefore find candidate causal mutations^{17,18,19}. Even if experimentally intractable, the large possible phenotypic space of an organism can be explored quickly with computational models.

Here, we used De Bruijn graph GWAS (DBGWAS) to discover enriched mutations associated with the endemic USA300 strain within clonal complex 8 (CC8)²⁰. Due to clonal expansion of USA300 strains from their progenitors within CC8, the enriched USA300 specific mutations were in high linkage disequilibrium. Further complicating the matter, we found that almost all mutations enriched within open reading frames (ORFs) were unique to USA300 lineage and not found in any other clonal complexes, precluding identification of potential causative mutations by homoplasy. Instead, we turned to reconstruction of a transcriptional regulatory network (TRN) to identify genes that were both associated with an enriched mutation and had altered regulation in USA300 strains. We built an Independent Component Analysis (ICA)-based reconstruction of the TRN using 670 publicly available RNA sequencing samples from both USA300 and non-USA300 CC8 strains. By factoring the RNA sequencing data into a series of signals and their activities, the ICA-based

reconstruction of the TRN shows both the static gene-regulator interaction and the dynamic activity of these interactions in a sample specific manner²¹. However, ICA is a generalized signal extraction algorithm and therefore does not distinguish between biological sources of signals like regulatory elements and ‘artificial’ sources that can be created by sourcing data from multiple strains. Therefore, in addition to signals associated with gene regulators, ICA also outputs signals associated with strain-specific changes in gene regulation. Furthermore, by utilizing RNA sequencing data from hundreds of samples to identify genes with strain-specific expression patterns, this approach is more likely to find strain-specific differences than previous approaches that focus on specific conditions^{22,23}.

This analysis revealed several genes with distinct expression patterns in USA300 strains that were also associated with DBGWAS enriched mutations. One of these genes, *isdH*, which encodes a haptoglobin binding protein, showed several enriched mutations in the regulatory region of USA300 strains including the deletion of the Fur repressor binding site. Additionally, the *isdH* gene generally had higher expression levels in samples from USA300 strains, connecting the mutation enriched by DBGWAS with differences in TRN enriched by ICA. Overall, our analysis shows how the reconstruction of TRN can be used to extend the limits of current GWAS approaches when studying emerging and endemic populations of bacterial pathogens.

Results

Classifying USA300 and non-USA300 genomes based on genetic markers

We sought to compare the genetic differences between USA300 CA-MRSA strains and other subtypes within CC8 that have lower clinical and community burden. Given that both subtypes exist within the same clonal complex, this comparison allowed us to probe the genetic basis for the success of USA300 strains with limited confounding effects of different genetic backgrounds. We analyzed 2038 *S. aureus* CC8 genomes which formed a closed pangenome, suggesting that the sampled genomes mostly captured the gene level variations within the clonal complex (**Figure 1a**). The CC8 pangenome consisted of 19176 gene clusters with 2291 core genes that were present in at least ~95% of the genomes analyzed. Among the remainder of the genes, 931 were categorized as accessory genes and 15954 were found in less than 5% of the genomes (**Figure S1a**). To get a higher resolution view of these genomes, we surveyed for unique alleles within the ORFs and in the 300 base-pair 5’ upstream and 3’ downstream sequences. Interestingly, we found a larger number of mutations within the ORFs, indicating the presence of greater genetic variation in the ORFs than in the neighboring regulatory regions.

Next, we classified the CC8 genomes into USA300 and non-USA300 strains using Genetic Marker Inference (GMI). GMI was previously developed to rapidly and systematically identify different subclades within inner-CC8 strictly based on genetic markers²⁴. In this scheme, USA300 genomes can be differentiated from non-USA300 CC8 genomes by the presence of either SCCMecIVa or the presence of Pantone-Valentine Leukocidin (PVL) in case of methicillin sensitive *S. aureus* (MSSA).

To identify the root of the USA300 clades, we first traversed up nodes of the phylogenetic tree starting from known USA300 strain TCH1516 and determined the number of strains, fraction PVL positive and fraction SCCMecIVa positive for each node during traversal. The root was placed at the last node where >90% of the strains within the subclade represented by the daughter nodes were SCCMecIVa and PVL positive (**Figure 1b**). As phylogenetic trees are nested, root finding with this procedure is not dependent on the starting USA300 strain. Same root was identified when the procedure was initialized with another well known USA300 reference strain FPR3757 (**Figure S1b**). Combining the genetic markers with phylogenetic grouping led to the classification of 1449

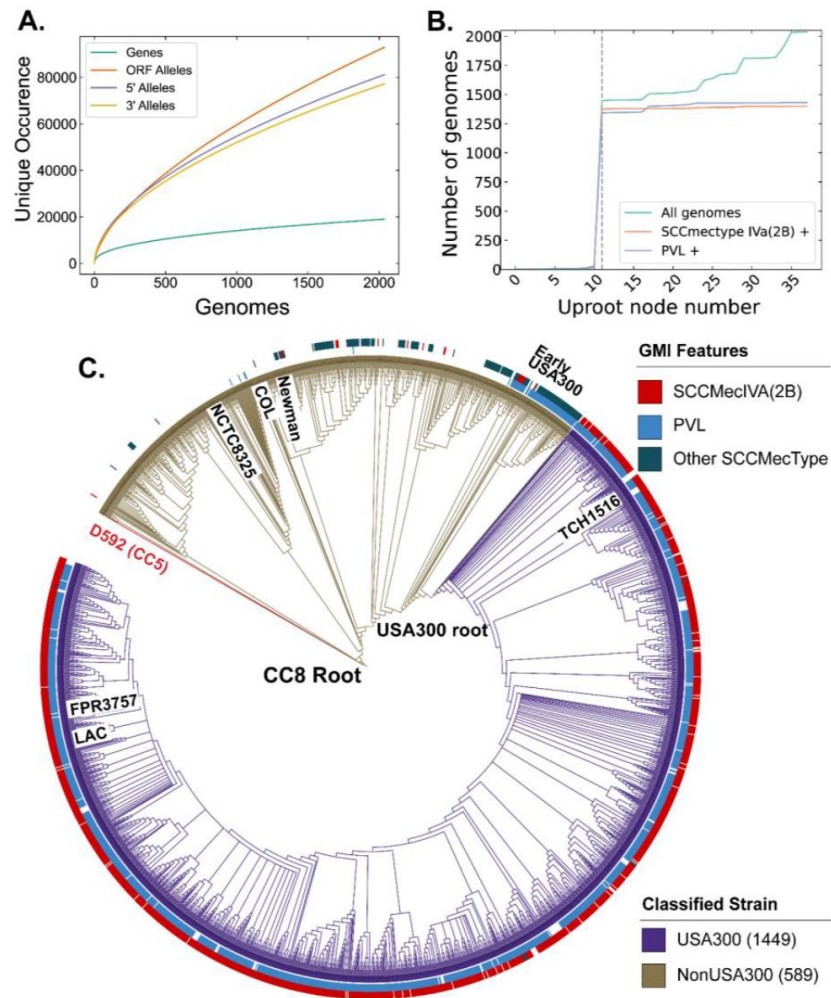


Figure 1.

CC8 pangenome and phylogeny.

(a) Pangenomic analysis of CC8 genomes shows the distribution of genes and mutations in ORFs and regulatory regions. (b) Prevalence of USA300 specific genetic markers, PVL and SCCMecIVA, as you traverse up the phylogenetic tree from TCH1516. The gray dashed line represents the node where the USA300 root is placed. (c) Phylogenetic tree of CC8 genomes classified into USA300 and non-USA300 strains.

genomes as USA300 and 589 genomes as non-USA300 (**Figure 1c** [↗](#), **Supplementary Table 1**). Strains previously identified as ‘early USA300’ were not part of our USA300 classification²⁴ [↗](#). While many of these strains are PVL positive, they have variable SCCMec types and therefore are likely to be clinically distinct from the endemic USA300 strains.

Enriching USA300 specific genes and mutations using DBGWAS

After classifying the genomes into USA300 and non-USA300 strains, we identified genes and mutations associated with each subtype by using the De Bruijn graph Genome Wide Association Study (DBGWAS)²⁰ [↗](#). DBGWAS provides a reference genome-free method for conducting GWAS analysis in prokaryotes by building a compacted De Bruijn Graph to represent the pan genome of input sequences. The nodes of the graph represent unique compacted k-mers that are joined by edges to other nodes with k-mers that appear adjacent to it in genomes. The procedure enriches unique k-mers that appear with different frequencies in each classification and outputs the enriched k-mer as well as its genetic neighborhood (called ‘components’) from the De Bruijn graph. Visualizing the components associated with the enriched k-mers makes it easier to interpret the k-mers and makes it easy to identify large structural variations (e.g. cassette acquisition) which are often represented by multiple enriched k-mers that fall within the same component. We took the output component graphs and automatically extracted the enriched genetic changes e.g. indels, SNPs, phage insertions etc (see **Supplementary Note 1**).

Many of the components were associated with genes and genetic elements expected to be enriched with USA300 strains-SCCMecIVA (the GMI marker), Arginine Catabolite Repressor Element (ACME), cap5E point mutation, multiple prophages etc (**Supplementary Table 2**). In total, we found k-mers in 149 components associated with 137 unique TCH1516 genes that were enriched in this analysis, pointing to a large array of mutations that are unique to the USA300 lineage (**Figure 2a** [↗](#)). Significant k-mers in some components did not uniquely match to TCH1516 genes or only matched to genes in non-USA300 reference genome, NCTC8325.

Genome-wide linkage and de novo mutations obfuscate identification of causal mutations

Though these mutations were enriched in USA300 strains with DBGWAS, we could not attribute the prevalence of any particular mutation to selection due to strong genome-wide linkage. We quantified the linkage disequilibrium by calculating the square of the correlation coefficient (r^2) for each of the enriched k-mer not associated with mobile genetic elements. High correlation coefficient indicates tight co-occurrence of k-mers in the genomes and therefore high linkage disequilibrium between the sequences. There was a strong linkage between the k-mers that were enriched in USA300 strains. Surprisingly, even k-mers that were 1.4 million base pairs away (the maximum distance between two sites in the circular 2.8 million base pairs long *S. aureus* genome) still had r^2 over 0.9 (**Figure 3a** [↗](#)).

To differentiate potential causal mutations from genetically linked alleles, we searched for mutation hotspots by comparing the positions of USA300 mutations in open reading frames (ORFs) to mutations in other clonal complexes. Barring recombination events, presence of mutation hotspots in the same position in multiple clades could point to selection acting on the sequence. Therefore, we searched for prevalence of enriched mutations in other non-CC8 clades. We identified 61 SNPs within open reading frames (ORFs) that were enriched in USA300 strains. To identify mutational hotspots in other clades, we downloaded all the amino acid sequences belonging to the PATRIC genus protein family of each of the gene products encoded by the selected ORFs²⁵ [↗](#). The PATRIC local protein family consists of sequences of homologous proteins within the same genus which were further filtered down to *S. aureus* species specific sequences. After filtering, each protein family comprised 2,000 to 16,000 unique sequences and the strains from

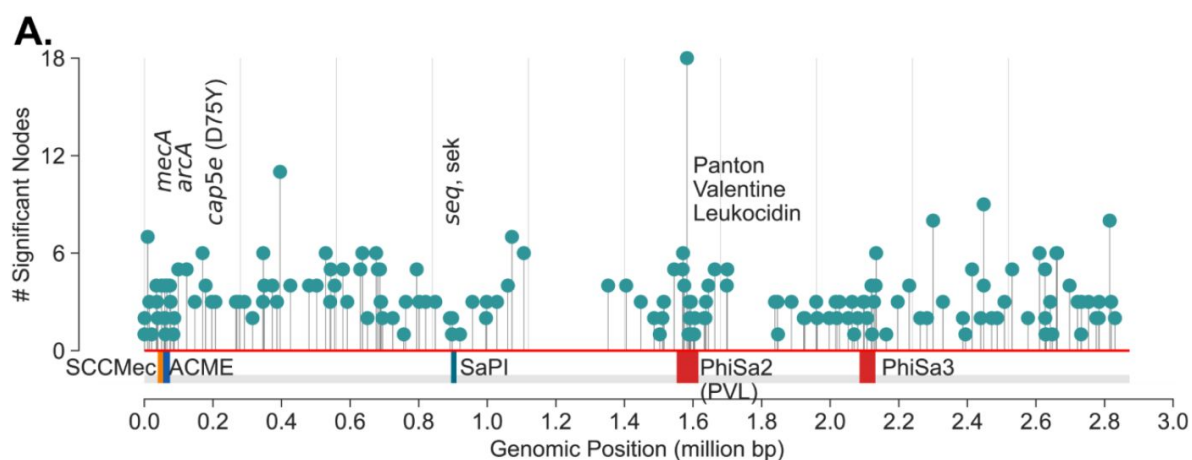


Figure 2.

USA300 strains associated mutations.

(a) DBGWAS recovers components associated with USA300 previously described markers of USA300 strains including *mecA* (SCCMec IVa), *arcA* (ACME), *cap5e* mutation, *seq, sek* and Phi-PVL. In addition, components with many other mutations scattered throughout the genome (NC_010079) are also enriched. Each 'significant node' represents a k-mer sequence (with minimum size of 31 nucleotides) that are associated with USA300 strains (adjusted p-value < 0.05).

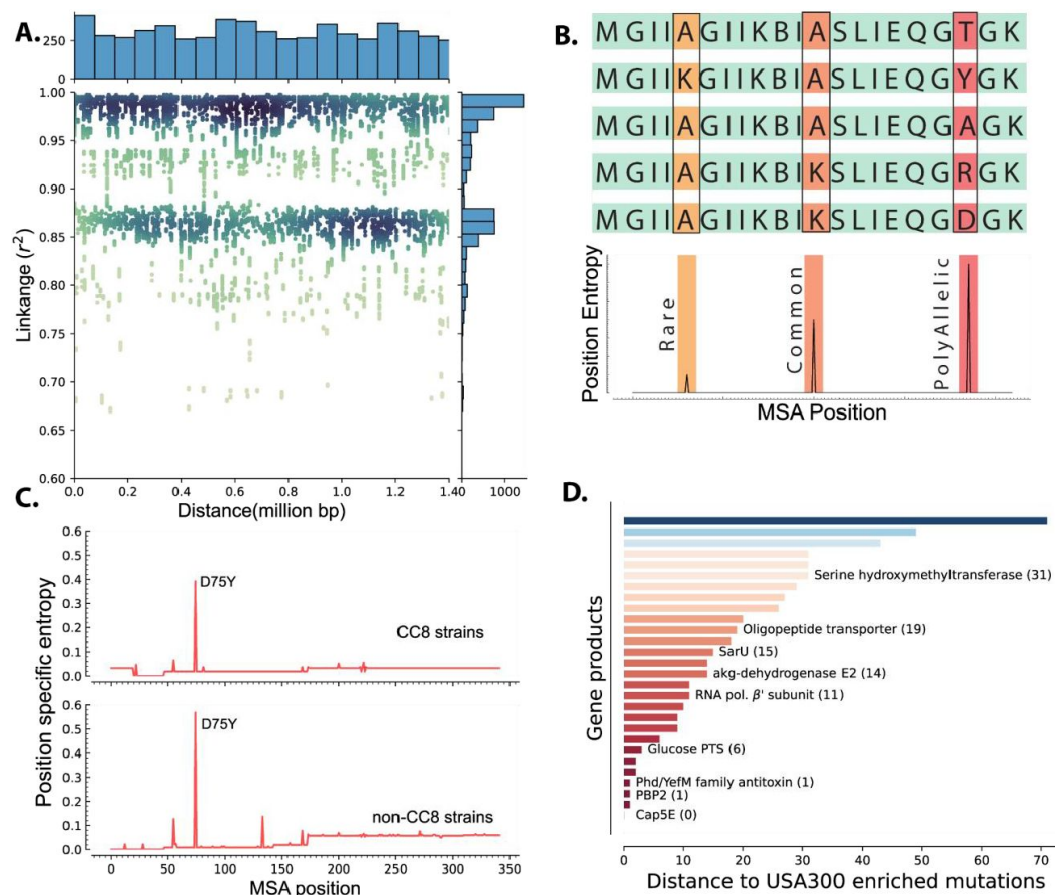


Figure 3.

Linkage Disequilibrium and de novo mutations in USA300 strains.

(a) Enriched k-mers showed high linkage disequilibrium, with some k-mers at 1.4 Mbp distance still having r^2 of greater than 0.98. (b) Schematic of position specific entropy analysis. Positions with heterogeneous sequences have higher calculated entropy than more conserved sequences with fewer mutations. (c) Using position specific entropy, we only found one example of shared enriched mutation in ORFs of USA300 and non-USA300 strains. (d) Distance (in base pairs) between the position of enriched mutation in USA300 strains and the position of the nearest entropy peak in other non-CC8 strains.

which the amino acid sequences were derived spanned dozens of clades allowing for broad comparisons (**Figure S2**). Lastly, we removed sequences associated with ST239 as it is thought to have emerged from large-scale recombination of ST8 and ST30 strains ²⁶.

We determined mutation hotspots by calculating position specific allelic entropy. Allelic entropy at a given amino acid position is a function of the number of unique amino acids found in that position and the frequency of the mutation ²⁷. Positions where all queried sequences have the same amino acid have low entropy, while positions that have frequent amino acid substitutions (hotspots) have high entropy (**Figure 3b**). This measure allows us to quickly determine the positions of mutation hotspots while accounting for multiple possible amino acid substitutions and rare mutations. Before calculating the position-specific entropy, all sequences within each of the PATRIC local protein families were aligned with Multiple Sequence Alignment (MSA). This alignment ensures proper comparison of amino acids even when there are deletions or insertions in some of the genes in the family.

Of the 36 enriched ORF mutations only the Asp75Tyr mutation in the *cap5E* gene, which was previously shown to ablate capsule production in USA300 strains, was found in other strains (**Figure 3c**) ¹⁶. Peaks in entropy corresponding to this mutation position were present in both the CC8 and non-CC8 strains while all other mutation positions were unique to CC8. Despite not having any perfect matches outside of the *cap5E* mutation, we found that for 28 of the mutations, a peak was present in sequences from other clades within 71 MSA positions. Together, our data suggests that mutations within ORFs in USA300 strains are likely de novo mutations and are not acquired through horizontal gene transfer though many of these mutations have occurred in hotspot regions (**Figure 3d**).

iModulon in the CC8 TRN points to mutations associated with differential regulation

The presence of genome-wide linkage and de novo mutations in ORFs severely limited the ability to distinguish causal SNPs contributing to increased pathogenesis in USA300 strains. The effect of some mutations, especially in ORFs, has been successfully linked to distinct phenotypes such as the absence of a capsule in USA300 and USA500 strains ¹⁶. However, the effect of mutations associated with changes in gene regulations can be much more difficult to assess ¹⁵. To look for mutations that may be associated with changes in transcriptional regulation, we used ICA to model gene-regulation in CC8 strains which can predict strain specific differences in expression patterns.

We collected CC8 strains associated RNA-sequencing data from Sequence Read Archive (SRA). After stringent QC/QA and curation, 291 non-USA300 strains (e.g. Newman, NCTC8325) and 379 USA300 (e.g. LAC, TCH1516) samples were used to create a single reconstruction of TRN across CC8 using ICA ^{21,28} (**Supplementary Table 3**) ICA calculates independently modulated sets of genes, iModulons, and the activities of those gene sets in each sample. iModulons calculated by ICA represent distinct sources of signals in the RNA-sequencing data. While most of the signals can be associated with different regulatory elements, iModulons associated with other biological features such as mobile genetic elements, genomic backgrounds are also enriched. In *Escherichia coli* and *Salmonella enterica Typhimurium*, multi-strain ICA has been used to calculate strain-specific iModulons that represent differences in gene expression ^{21,29}.

In our reconstruction, two iModulons captured a large number of genes with different expression levels in the non-USA300 and USA300 strains (**Figure 4a**; **Figure S3a**). Most of the genes in the strain-specific iModulons belonged to mobile elements associated with USA300 strains such as ACME, SCCMec, Phi-PVL etc. However, the iModulons also contained core genes that are present in both strains, pointing to possible differences in gene regulation (**Figure 4b**; **Figure S3b**).

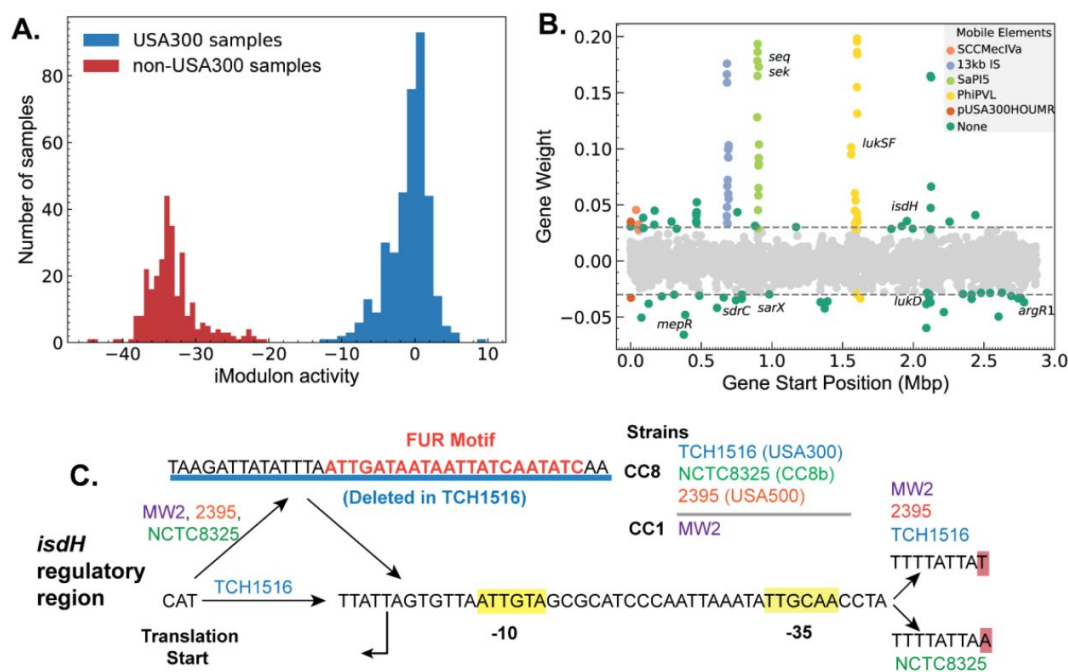


Figure 4.

Strain-specific regulatory changes in the CC8 clade.

(a) ICA analysis of USA300 and non-USA300 RNA-sequencing data identified an iModulon with strain specific activity.(b) The strain-specific iModulon contained various horizontally acquired elements (e.g. ACME, PhiPVL) that are prevalent in USA300 lineage as well as conserved genes with strain-specific expression patterns. (c) Comparing the 5' regulatory region of the gene *isdH* from various *S. aureus* strains revealed a unique deletion containing Fur binding site in USA300 reference strain TCH1516.

We mapped the enriched mutations from DBGWAS onto the core genes enriched in the strain-specific iModulon. 10 core genes-including *isdH*, *argR1* and *araC* family regulator-with mutations in the ORF or in the regulatory region were also enriched in the strain-specific iModulon (**Supplementary Table 4**). Of these genes, gene *isdH*, encoding a heme scavenger molecule showed distinct strain-specific expression levels and had enriched k-mers that are mapped to the upstream regulatory region. Therefore, we compared the upstream regulatory region of several reference strains including TCH1516 (USA300), NCTC8325 (CC8b), 2395 (USA500). Additionally, we included MW2 (CC1 CA-MRSA) as the transcription start site (TSS) in the region has been experimentally confirmed in this strain³⁰. Comparisons showed a 38 base-pair deletion in the 5' untranslated region containing a transcription factor Fur binding site (q-val=0.033e-4) (**Figure 4c**).

This deletion was detected in all of the 1385 USA300 genomes, but only present in 95 of the 589 non-USA300 genomes. As Fur is a repressor that blocks expression in presence of iron concentration, this deletion in the Fur binding site may be responsible for the general increase in *isdH* expression observed in USA300 samples (**Figure S4a**). We also found a second mutation upstream of the predicted -35 binding site that was also enriched in USA300. Interestingly, while the MW2 strain did not have the 38 bp deletion, it contained the exact upstream A->T mutation. All other base-pairs in the region were perfectly matched in between all the reference genomes. The combination of evidence from genetic and transcriptomic analysis suggests that regulation of *isdH* is altered in USA300 strains compared to its non-USA300 progenitors.

As with many point mutations detected in our analysis, the absence of Fur binding site upstream of *isdH* gene is prevalent only in the USA300 lineage. We searched for Fur binding motif in the 100 base-pairs upstream regions from 3515 non-CC8 strains spanning multiple clonal complexes (**Figure S2**). We detected the binding motif in all but 21 strains. Of the 21 strains with no detectable Fur binding sites, 6 belonged to ST72 (out of 28 total from this type) and 6 had uncharacterized MLST type. The rest were distributed among types 121, 1750, 375, 1, 3317, 15, 7, 398, and 4803, with one positive strain per type.

Discussion

Emergence of CA-MRSA USA300 strains from HA-MRSA USA500 progenitors presents a natural experiment to probe the genetic basis for the establishment of the USA300 lineage. However, in studying these groups, genetic methods like GWAS were limited in finding causal mutations due to genome-wide linkage disequilibrium and presence of an unexpectedly large number of de novo mutations unique to the USA300 lineage. Here, we demonstrated how a model of transcriptional regulation with iModulons can be used to break through the impasse created by the high linkage disequilibrium and predict candidate causal mutations. From the combined RNA sequencing dataset of USA300 and non-USA300 strains, ICA calculated a single iModulon that captured strain specific variation in gene expression. As expected, most genes in the iModulons were part of mobile genetic elements such as ACME and SCCMec because they have zero expression level in non-USA300 samples. However, the iModulon also contained several core genes that are present in both groups but are differentially regulated. A deeper analysis of the regulatory region of one of these genes with enriched mutation, *isdH*, revealed a deletion of a DNA segment containing the binding site of the Fur repressor. In congruence with this observation we also found that USA300 strains with the deleted Fur binding site showed general increase in *isdH* expression level. Combining GWAS with large-scale transcriptomic modeling was therefore able to predict potential causal mutations that led to the increased clinical burden of the USA300 lineage.

The current analysis utilized the available DNA and RNA sequencing data and the methods used here are scalable to the rapidly growing number of data in the public repositories. Indeed, with the greater scale, we can get more granular insight into subclade specific differences. The

transcriptomic analysis consisted of samples primarily from the USA300 (CC8e and CC8f) clades, the CC8a clade represented by Newman and the CC8b clade represented by NCTC8325 and its derivatives. However, the CC8b and CC8a clades are currently undersampled due to its minimal clinical burden compared to USA300. We therefore combined strains from all non-USA300 clades into a single group for GWAS. The misalignment of RNA sequencing samples from GWAS samples may explain the low number of hits that were enriched by both methods when many other unique gene expression patterns have been observed in USA300 strains. The scalability of the methods used herein will enable granular and more in-depth analysis as these sequencing databases continue to expand rapidly.

With time, the scaling of databases may be able to resolve the issue of imbalanced sampling. On the other hand, resolving the confounding effect of linkage disequilibrium inherent in emerging and endemic strains will require a new generation of modeling methods¹¹. Our current approach focuses on modeling the changes in gene regulation at the transcriptional level, but causal mutations can have any number of effects on the phenotype of the organism. New modeling methods that can systematically predict these other phenotypes are now rapidly emerging. Our recent work with *Mycobacterium tuberculosis* utilized a metabolic allele classifier (MACs) which combines genome scale metabolic models with machine learning to estimate biochemical effects of alleles thus mapping mutations to changes in metabolic fluxes¹⁹. Similarly, advances in protein structure prediction with AlphaFold2 and RosettaFold puts us at the cusp of predicting the effects of mutations on protein folding^{31,32}. Combination of these modeling techniques may therefore prove to be the breakthrough required to advance solutions to the current challenges in population genetics of emerging pathogens.

Materials and Methods

Pangenomic analysis

The pangenome analysis was run as described in detail before²⁷. Briefly, “complete” or “WGS” samples from CC8/ST8 were downloaded from the PATRIC database²⁵. Sequences with lengths that were not within 3 standard deviations of the mean length or those with more than 100 contigs were filtered out. A non-redundant list of CDSs from all genomes was created and clustered by protein sequence using CD-HIT with minimum identity and minimum alignment length of 80%³³. To get the 5' and 3' sequences, non-redundant 300 nucleotide upstream and downstream sequences from the CDS were extracted for each gene.

The CDSs were divided into core, accessory and unique genes based on the frequency of genes as previously described²⁷. To calculate the frequency thresholds for each category, $P(x)$, the number of genes with frequency x and its integral $F(x)$, the cumulative frequency less than or equal to x were calculated. The multimodal gene distribution can be estimated by sum of two power laws as:

$$P(x) = c_1 x^{-\alpha_1} + c_2 (N + 1 - x)^{-\alpha_2} \quad x = 1, 2, \dots, N$$

where N is the total number of genomes, x is the gene frequency and $(c_1, c_2, -\alpha_1, -\alpha_2)$ are parameters fit based on the data. The cumulative distribution is then the integral of $P(x)$ with additional parameter k :

$$F(x) = k + \frac{c_1}{1 - \alpha_1} x^{1-\alpha_1} - \frac{c_2}{1 - \alpha_2} (N + 1 - x)^{1-\alpha_2}$$

The parameters $(c_1, c_2, -\alpha_1, -\alpha_2, \text{ and } k)$ were fitted based on the data using nonlinear least squares regression from `scipy`³⁴. The frequency threshold of core genomes was defined as greater than $0.9 N + 0.1 x^*$ and the threshold for unique genome was defined as $0.1 x^*$, where x^* represents the

inflection point of the fitted cumulative distribution.

Reconstructing CC8 phylogenetic tree

The phylogenetic tree was reconstructed using the standardized PHaME pipeline on the PATRIC sequences that passed the QC/QA³⁵. Using the pipeline, the contigs and sequences were aligned to the reference TCH1516 genome NC_010079 and plasmids NC_012417, NC_010063³⁶ and core SNPs were calculated. The core SNPs were then used to estimate the phylogenetic tree using IQ-TREE run with 1000 bootstraps and utilizing the ultrafast bootstrap^{37,38}. The tree was built using the “TVMe+ASC+G4” model as suggested by the IQ-TREE ModelFinder³⁹. Finally, iTOL was used to visualize, annotate and root the tree with the USA100 D592 (NZ_CP035791) from CC5 as the outgroup⁴⁰.

Classification of USA300 and non-USA300 strains

The USA300 and non-USA300 strains were classified based on a previously proposed and validated CC8 subtyping scheme²⁴. In this scheme, USA300 strains can be identified from the whole genome if they are PVL positive MSSA or MRSA with SCCMec IVa cassette. We detected SCCMec types using SCCMecFinder, and only those genomes where the cassette could be identified by both BLAST and k-mer based methods were marked as positive⁴¹. PVL was detected using protein BLAST. To find the root of the USA300 strains in the phylogenetic tree, the genomes in the tree were annotated by their PVL and SCCMec status. We added additional criteria that all genomes identified as USA300 by GMI form a distinct subclade before they are labeled as USA300 i.e. PVL or SCCMECIVa positive genomes that grouped separately from other USA300 strains in the phylogenetic tree were not labeled as USA300.

We detected SCCMec cassettes in 1588 genomes of which 1358 were SCCMecIVa positive. We also found 1431 PVL positive genomes using BLASTn search with PVL encoding genes from USA300 TCH1516 (USA300HOU_RS07645, USA300HOU_RS07650) as reference (**Supplementary Table 5**). Lastly, we reconstructed the CC8 phylogenetic tree based on core Single Nucleotide Polymorphisms (SNPs) and rooted the tree using strain D592 (CC5) as an outgroup. The tree was then traversed from reference strain TCH1516 to the CC8 root using ete3, while tracking the total number of genomes, the total number of SCCMec IVa positive genomes and the number of PVL positive genomes in each root⁴². The root of USA300 was placed manually where the number of total genomes kept increasing while the number of PVL and SCCMec positive genomes plateaued. All strains in the clade represented by the USA300 root were classified as USA300 regardless of their SCCMec or PVL status.

DBGWAS and k-mer linkage calculations

DBGWAS was used to enrich mutations unique to USA300 strains. Alleles with frequency less than 0.05 were filtered (-maf 0.05) and all components enriched with q-values less than 0.05 were documented (-SFF q0.05). Genome-wide linkage was estimated by Pearson correlation of the presence/ absence of enriched k-mers and distance was measured based on the k-mer alignment to the reference TCH1516 genome.

To determine the enriched ‘genetic event’ (e.g. SNP, indel, mobile genetic element etc), the graph output from DBGWAS was first loaded onto a networkX model⁴³. All nodes in the graph with frequency lower than 0.05 were discarded. MGEs were identified if all significant nodes from DBGWAS had higher frequencies in one strain, e.g. all nodes associated with SCCMec had higher frequencies in USA300 strains. To find SNPs and smaller indel events, the networkx was used to find cycles in the graph, which results from bifurcation and eventual re-collapse of Debruijn graphs around mutations. For each cycle, the ‘end nodes’ representing the start and end of the bifurcation were identified by finding the nodes in the cycle with highest frequency across all samples. As ‘end nodes’ are present in both case and control samples, they will have higher

frequency than other nodes in the cycle which are specific to either case or control. Once the end nodes are identified, the two paths around the bifurcations representing the case and control specific sequences were identified using the shortest path algorithm in networkx. The sequences from nodes of each path were concatenated, changing the sequences to reverse complements and removing overlaps in sequences when required. The concatenated sequences from each path were then compared using BioPython pairwise global alignments to find the SNPs or indels that differentiate the sequences from case and control⁴⁴. If reference sequences are passed, the concatenated sequences are aligned to the reference sequences using BLAST and mutation positions were converted from k-mer positions to positions in the reference genomes. The code used for this analysis can be found in https://github.com/SBRG/dbgwas_network_analysis.

Mapping mutation hotspots with position specific Shannon entropy

For each of the CDS with enriched mutations, the PATRIC local protein family (PLfam) was identified based on the reference TCH1516 genome. All available protein sequences for each CDS PLfam were downloaded and filtered for *S. aureus* sequences. The multilocus sequence type (MLST) of the source genome of each downloaded sequence was mapped using the PATRIC database. The sequences were divided into ST8 and non ST8 and ST239 sequences were filtered. MAFFT was used for multiple alignment and position-specific Shannon entropy was calculated on the aligned file⁴⁵. The entropy is calculated as:

$$H(X) = - \sum_{i=1}^n P(x_i) \log_2 P(x_i)$$

where n is the total number of unique amino acids in the position and $P(x_i)$ is the probability of finding the given amino acid.

Calculating strain-specific iModulons with independent component analysis

ICA of RNA sequencing data was performed using the pymodulon package²⁸. Using the package, all available RNA sequencing data for non-USA300 and USA300 strains were downloaded, run through the QC/QA pipeline, manually curated for metadata and aligned to the TCH1516 genome (NC_010079, NC_012417, NC_010063). The combined data was then transformed into log-TPM and normalized to a single reference condition (SRX3760886, SRX3760891). This is in contrast to other ICA models that normalize the data to project specific reference conditions to reduce batch effects. However, normalizing to project specific control conditions also erases the strain specific information as almost all BioProjects contain data from only one isolate (e.g. NCTC8325, TCH1516, LAC etc). ICA was then run as previously described in chapter 3. The activities of the output iModulons were manually parsed to look for iModulons with the largest strain specific differences.

Fur box motif search

isdH genes in all the genomes were first clustered using CD-HIT with identity and coverage minimum of 0.8³³. All annotated *isdH* genes fell within a single cluster. For each genome, the 100 bp upstream region was then extracted and used for motif search. Motif search for the Fur box was conducted using the FIMO package from the MEME suite with default settings⁴⁶. The *S. aureus* strain NCTC8325 Fur motif from collectTF was used as a reference⁴⁷.

Code and Data Availability

The genome sequences and the RNA-sequencing data used in this study are publicly available (See Supplementary Table 1 and 3 respectively). The code used for analysis, the intermediate files and models are available on github (<https://github.com/sapoudel/USA300GWASPUB>).

References

1. Young B. C., et al. (2019) **Panton-Valentine leucocidin is the key determinant of *Staphylococcus aureus* pyomyositis in a bacterial GWAS** *Elife* **8**
2. Bosi E., et al. (2016) **Comparative genome-scale modelling of *Staphylococcus aureus* strains identifies strain-specific metabolic capabilities linked to pathogenicity** *Proc. Natl. Acad. Sci. U. S. A* **113**:E3801–9
3. Choudhary K. S., et al. (2018) **The *Staphylococcus aureus* Two-Component System AgrAC Displays Four Distinct Genomic Arrangements That Delineate Genomic Virulence Factor Signatures** *Front. Microbiol* **9**
4. Copin Correction for, et al. (2019) **Sequential evolution of virulence and resistance during clonal spread of community-acquired methicillin-resistant *Staphylococcus aureus*** *Proc. Natl. Acad. Sci. U. S. A* **116**
5. Krishna A., Holden M. T. G., Peacock S. J., Edwards A. M., Wigneshweraraj S. (2018) **Naturally occurring polymorphisms in the virulence regulator Rsp modulate *Staphylococcus aureus* survival in blood and antibiotic susceptibility** *Microbiology* **164**:1189–1195
6. Earle S. G., et al. (2016) **Identifying lineage effects when controlling for population structure improves power in bacterial association studies** *Nat Microbiol* **1**
7. Diep B. A., et al. (2008) **The arginine catabolic mobile element and staphylococcal chromosomal cassette mec linkage: convergence of virulence and resistance in the USA300 clone of methicillin-resistant *Staphylococcus aureus*** *J. Infect. Dis* **197**:1523–1530
8. Faden H., et al. (2010) **Importance of colonization site in the current epidemic of staphylococcal skin abscesses** *Pediatrics* **125**:e618–24
9. Steinig E. J., et al. (2019) **Evolution and Global Transmission of a Multidrug-Resistant, Community-Associated Methicillin-Resistant *Staphylococcus aureus* Lineage from the Indian Subcontinent** *MBio* **10**
10. Challagundla L., et al. (2018) **Phylogenomic Classification and the Evolution of Clonal Complex 5 Methicillin-Resistant *Staphylococcus aureus* in the Western Hemisphere** *Front. Microbiol* **9**
11. Bal A. M., et al. (2016) **Genomic insights into the emergence and spread of international clones of healthcare-, community- and livestock-associated methicillin-resistant *Staphylococcus aureus*: blurring of the traditional definitions** *Journal of Global Antimicrobial Resistance* **6**:95–101
12. Uhlemann A.-C., et al. (2014) **Molecular tracing of the emergence, diversification, and transmission of *S. aureus* sequence type 8 in a New York community** *Proc. Natl. Acad. Sci. U. S. A* **111**:6738–6743
13. Challagundla L., et al. (2018) **Range Expansion and the Origin of USA300 North American Epidemic Methicillin-Resistant *Staphylococcus aureus*** *MBio* **9**

14. Everitt R. G., et al. (2014) **Mobile elements drive recombination hotspots in the core genome of *Staphylococcus aureus*** *Nat. Commun* **5**
15. Thurlow L. R., Joshi G. S., Richardson A. R. (2012) **Virulence strategies of the dominant USA300 lineage of community-associated methicillin-resistant *Staphylococcus aureus* (CA-MRSA)** *FEMS Immunol. Med. Microbiol* **65**:5–22
16. Boyle-Vavra S., et al. (2015) **USA300 and USA500 clonal lineages of *Staphylococcus aureus* do not produce a capsular polysaccharide due to conserved mutations in the cap5 locus** *MBio* **6**
17. Nishizaki S. S., et al. (2020) **Predicting the effects of SNPs on transcription factor binding affinity** *Bioinformatics* **36**:364–372
18. Choi Y., Sims G. E., Murphy S., Miller J. R., Chan A. P. (2012) **Predicting the functional effect of amino acid substitutions and indels** *PLoS One* **7**
19. Kavas E. S., Yang L., Monk J. M., Heckmann D., Palsson B. O. (2020) **A biochemically-interpretable machine learning classifier for microbial GWAS** *Nat. Commun* **11**
20. Jaillard M., et al. (2018) **A fast and agnostic method for bacterial genome-wide association studies: Bridging the gap between k-mers and genetic events** *PLoS Genet* **14**
21. Sastry A. V., et al. (2019) **The *Escherichia coli* transcriptome mostly consists of independently regulated modules** *Nat. Commun* **10**
22. Jones M. B., et al. (2014) **Genomic and transcriptomic differences in community acquired methicillin resistant *Staphylococcus aureus* USA300 and USA400 strains** *BMC Genomics* **15**
23. Iqbal Z., et al. (2016) **Comparative virulence studies and transcriptome analysis of *Staphylococcus aureus* strains isolated from animals** *Sci. Rep* **6**
24. Bowers J. R., et al. (2018) **Improved Subtyping of *Staphylococcus aureus* Clonal Complex 8 Strains Based on Whole-Genome Phylogenetic Analysis** *mSphere* **3**
25. Wattam A. R., et al. (2014) **PATRIC, the bacterial bioinformatics database and analysis resource** *Nucleic Acids Res* **42**:D581–91
26. Robinson D. A., Enright M. C. (2004) **Evolution of *Staphylococcus aureus* by large chromosomal replacements** *J. Bacteriol* **186**:1060–1064
27. Hyun J. C., Monk J. M., Palsson B. O. (2022) **Comparative pangenomics: analysis of 12 microbial pathogen pangenomes reveals conserved global structures of genetic and functional diversity** *BMC Genomics* **23**
28. Sastry A. V., et al. (2021) **Mining all publicly available expression data to compute dynamic microbial transcriptional regulatory networks** *bioRxiv* 2021.07.01.450581 <https://doi.org/10.1101/2021.07.01.450581>
29. Yuan Y., et al. (2022) **Pan-genomic analysis of transcriptional modules across *Salmonella Typhimurium* reveals the regulatory landscape of different strains** *bioRxiv* 2022.01.11.475931 <https://doi.org/10.1101/2022.01.11.475931>

30. Prados J., Linder P., Redder P. (2016) **TSS-EMOTE, a refined protocol for a more complete and less biased global mapping of transcription start sites in bacterial pathogens** *BMC Genomics* **17**
31. Baek M., et al. (2021) **Accurate prediction of protein structures and interactions using a three-track neural network** *Science* **373**:871–876
32. Jumper J., et al. (2021) **Highly accurate protein structure prediction with AlphaFold** *Nature* **596**:583–589
33. Fu L., Niu B., Zhu Z., Wu S., Li W. (2012) **CD-HIT: accelerated for clustering the next generation sequencing data** *Bioinformatics* **28**:3150–3152
34. Jones Eric **Travis Oliphant, Pearu Peterson and others**
35. Shakya M., et al. (2020) **Standardized phylogenetic and molecular evolutionary analysis applied to species across the microbial tree of life** *Sci. Rep* **10**
36. Highlander S. K., et al. (2007) **Subtle genetic changes enhance virulence of methicillin resistant and sensitive *Staphylococcus aureus*** *BMC Microbiol* **7**
37. Minh B. Q., et al. (2020) **IQ-TREE 2: New Models and Efficient Methods for Phylogenetic Inference in the Genomic Era** *Mol. Biol. Evol* **37**:1530–1534
38. Hoang D. T., Chernomor O., von Haeseler A., Minh B. Q., Vinh L. S. (2018) **UFBoot2: Improving the Ultrafast Bootstrap Approximation** *Mol. Biol. Evol* **35**:518–522
39. Kalyaanamoorthy S., Minh B. Q., Wong T. K. F., von Haeseler A., Jermini L. S. (2017) **ModelFinder: fast model selection for accurate phylogenetic estimates** *Nat. Methods* **14**:587–589
40. Letunic I., Bork P. (2021) **Interactive Tree Of Life (iTOL) v5: an online tool for phylogenetic tree display and annotation** *Nucleic Acids Res* **49**:W293–W296
41. Kaya H., et al. (2018) **SCCmecFinder, a Web-Based Tool for Typing of Staphylococcal Cassette Chromosome mec in *Staphylococcus aureus* Using Whole-Genome Sequence Data** *mSphere* **3**
42. Huerta-Cepas J., Serra F., Bork P. (2016) **ETE 3: Reconstruction, Analysis, and Visualization of Phylogenomic Data** *Mol. Biol. Evol* **33**:1635–1638
43. Hagberg A., Swart P., Chult D. (2008) **Hagberg, A., Swart, P. & S Chult, D. Exploring network structure, dynamics, and function using networkx.** <https://www.osti.gov/biblio/960616> (2008).
44. Cock P. J. A., et al. (2009) **Biopython: freely available Python tools for computational molecular biology and bioinformatics** *Bioinformatics* **25**:1422–1423
45. Katoh K., Misawa K., Kuma K.-I., Miyata T. (2002) **MAFFT: a novel method for rapid multiple sequence alignment based on fast Fourier transform** *Nucleic Acids Res* **30**:3059–3066
46. Bailey T. L., et al. (2009) **MEME SUITE: tools for motif discovery and searching** *Nucleic Acids Res* **37**:W202–8

47. Kiliç S., White E. R., Sagitova D. M., Cornish J. P., Erill I. (2013) **CollectTF: a database of experimentally validated transcription factor-binding sites in Bacteria** *Nucleic Acids Res* **42**:D156–D160

Article and author information

Saugat Poudel

Department of Bioengineering, University of California San Diego

Jason Hyun

Department of Bioengineering, University of California San Diego

Ying Hefner

Department of Bioengineering, University of California San Diego

Jon Monk

Palmona Pathogenomics

Victor Nizet

Collaborative to Halt Antibiotic-Resistant Microbes (CHARM), Department of Pediatrics, University of California San Diego, Department of Pediatrics, University of California San Diego
ORCID ID: [0000-0003-3847-0422](https://orcid.org/0000-0003-3847-0422)

Bernhard O. Palsson

Department of Bioengineering, University of California San Diego, Collaborative to Halt Antibiotic-Resistant Microbes (CHARM), Department of Pediatrics, University of California San Diego, Department of Pediatrics, University of California San Diego
For correspondence: palsson@eng.ucsd.edu

Copyright

© 2023, Poudel et al.

This article is distributed under the terms of the [Creative Commons Attribution License](https://creativecommons.org/licenses/by/4.0/), which permits unrestricted use and redistribution provided that the original author and source are credited.

Editors

Reviewing Editor

Marisa Nicolás

Laboratório Nacional de Computação Científica, Rio de Janeiro, Brazil

Senior Editor

Aleksandra Walczak

École Normale Supérieure - PSL, Paris, France

Reviewer #1 (Public Review):

Summary:

This is large-scale genomics and transcriptomics study of the epidemic community-acquired

methicillin-resistant *S. aureus* clone USA300, designed to identify core genome mutations that drove the emergence of the clone. It used publicly available datasets and a combination of genome-wide association studies (GWAS) and independent principal-component analysis (ICA) of RNA-seq profiles to compare USA300 versus non-USA300 within clonal complex 8. By overlapping the analyses the authors identified a 38bp deletion upstream of the iron-scavenging surface-protein gene *isdH* that was both significantly associated with the USA300 lineage and with a decreased transcription of the gene.

Strengths:

Several genomic studies have investigated genomic factors driving the emergence of successful *S. aureus* clones, in particular USA300. These studies have often focussed on acquisition of key accessory genes or have focussed on a small number of strains. This study makes a smart use of publicly available repositories to leverage the sample size of the analysis and identify new genomics markers of USA300 success.

The approach of combining large-scale genomics and transcriptomics analysis is powerful, as it allows to make some inferences on the impact of the mutations. This is particularly important for mutations in intergenic regions, whose functional impact is often uncertain. The statistical genomics approaches are elegant and state-of-the-art and can be easily applied to other contexts or pathogens.

Weaknesses:

The main weakness of this work is that these data don't allow a casual inference on the role of *isdH* in driving the emergence of USA300. It is of course impossible to prove which mutation or gene drove the success of the clone, however, experimental data would have strengthened the conclusions of the authors in my opinion.

Another limitation of this approach is that the approach taken here doesn't allow to make any conclusions on the adaptive role of the *isdH* mutation. In other words, it is still possible that the mutation is just a marker of USA300 success, due to other factors such as PVL, ACMI or the SCCmecIVa. This is because by its nature this analysis is heavily influenced by population structure. Usually, GWAS is applied to find genetic loci that are associated with a phenotype and are independent of the underlying population structure. Here, authors are using GWAS to find loci that are associated with a lineage. In other words, they are simply running a univariate analysis (likely a logistic regression) between genetic loci and the lineage without any correction for population structure, since population structure is the outcome. Therefore, this approach can't be applied to most phenotype-genotype studies where correction for population structure is critical.

Finally, the approach used is complex and not easily reproduced in another dataset. Although I like DBGWAS and find the network analysis elegant, I would be interested in seeing how a simpler GWAS tool like Pyseer would perform.

- <https://doi.org/10.7554/eLife.90668.1.sa1>

Reviewer #2 (Public Review):

Summary:

The work of Poudel et al. identified potential causal mutations related to the successful emergence of the virulent USA300 community-associated MRSA clone within clonal complex 8. To achieve this, the authors employed a methodology that combines the genome-wide association studies (GWAS) with the inference of a transcriptional regulatory network (TRN) through the independent component analysis (ICA) method from publicly available transcriptomic data. Thus, they identified genes with altered expression in the iModulons calculated by ICA and enriched mutations obtained from the De Bruijn graph genome-wide association study (DBGWAS) in the USA300 strains versus non-USA300 strains. The results revealed a deletion of 38 base pairs, containing a binding site for the Fur repressor, and an

A → T mutation, both occurring in the upstream region of the *isdH* gene, whose expression level in USA300 strains exhibited a general increase compared to the other group. *IsdH* encodes the iron-regulated surface determinant protein H, which plays a crucial role in iron acquisition from heme and immune system evasion - two essential processes for the pathogenicity of *S. aureus*.

Strengths:

The clonal complex 8 (CC8), one of the most prevalent among *S. aureus*, encompasses strains responsible for both community-associated MRSA infections (CA-MRSA) and healthcare-associated (HA) infections (HA-MRSA and HA-MSSA). Within the CC8, one of the most prominent lineages is USA300, which emerged in the early 2000s and has since become a leading cause of CA-MRSA infections in the United States. The key genetic traits that characterize USA300 strains include the presence of the Pantan-Valentine leukocidin (PVL) encoded by the genes *lukF-PV* and *lukS-PV*, the staphylococcal chromosomal cassette *mec IVa* (SCC*mecIVa*), and the arginine catabolic mobile element (ACME). Investigating the phenotypic impact of individual mutations on the success of epidemic strains through GWAS poses a challenge due to two main confounding factors: genome-wide linkage disequilibrium (LD) and population structure. The genome-wide LD is associated with false positives, where linked non-causal mutations are mistakenly identified as causal due to the same genomic backgrounds. Therefore, the strength of this work lies in the use of publicly available transcriptomic data to construct a TRN based on ICA. This approach validates the mutations enriched by GWAS and reduces the occurrence of false positives attributed to high genome-wide LD. By integrating various 'omics' data sources, this method enhances the reliability of the results and has successfully identified new potential genetic markers specific to USA300 strains. Furthermore, it revealed mutations within core genes and intergenic regulatory regions, findings that can be validated through experimental data.

Weaknesses:

GWAS aims to identify statistically significant associations that suggest a causal link between genotype and the specific phenotype of interest while simultaneously filtering out spurious associations caused by confounding factors. While the method described in this study minimizes the impact of genome-wide linkage disequilibrium (LD), it does not extend to addressing population structure. This is because the objective was precisely to identify mutations associated with the emergence of the USA300 clone. In this context, the confounding element arising from shared ancestry becomes the subject of analysis rather than an issue to be corrected. Therefore, it is essential to highlight that the method proposed in this work can not be applied to genome-wide association studies, where correction for population structure is critical for distinguishing genuine causal associations from spurious ones. This correction is crucial and necessary to most of the studied phenotypes of interest.

Another limitation is that, although the authors emphasize the mutation in the *isdH* gene, the analyses conducted in this study do not provide insight into any potential adaptive function associated with it. Similarly, like the other genes exhibiting distinct expression patterns associated with enriched mutations from DBGWAS in USA300 strains, *isdH* is among the potential markers related to the success of the clone. This group includes well-established markers, such as ACME, which carries relevant genes like the *arc* operon and the *speG* gene that contribute to virulence and survival at infection sites.

Finally, despite the availability of the codes on GitHub, the analysis itself is not easily reproducible or adaptable to other datasets.

- <https://doi.org/10.7554/eLife.90668.1.sa0>