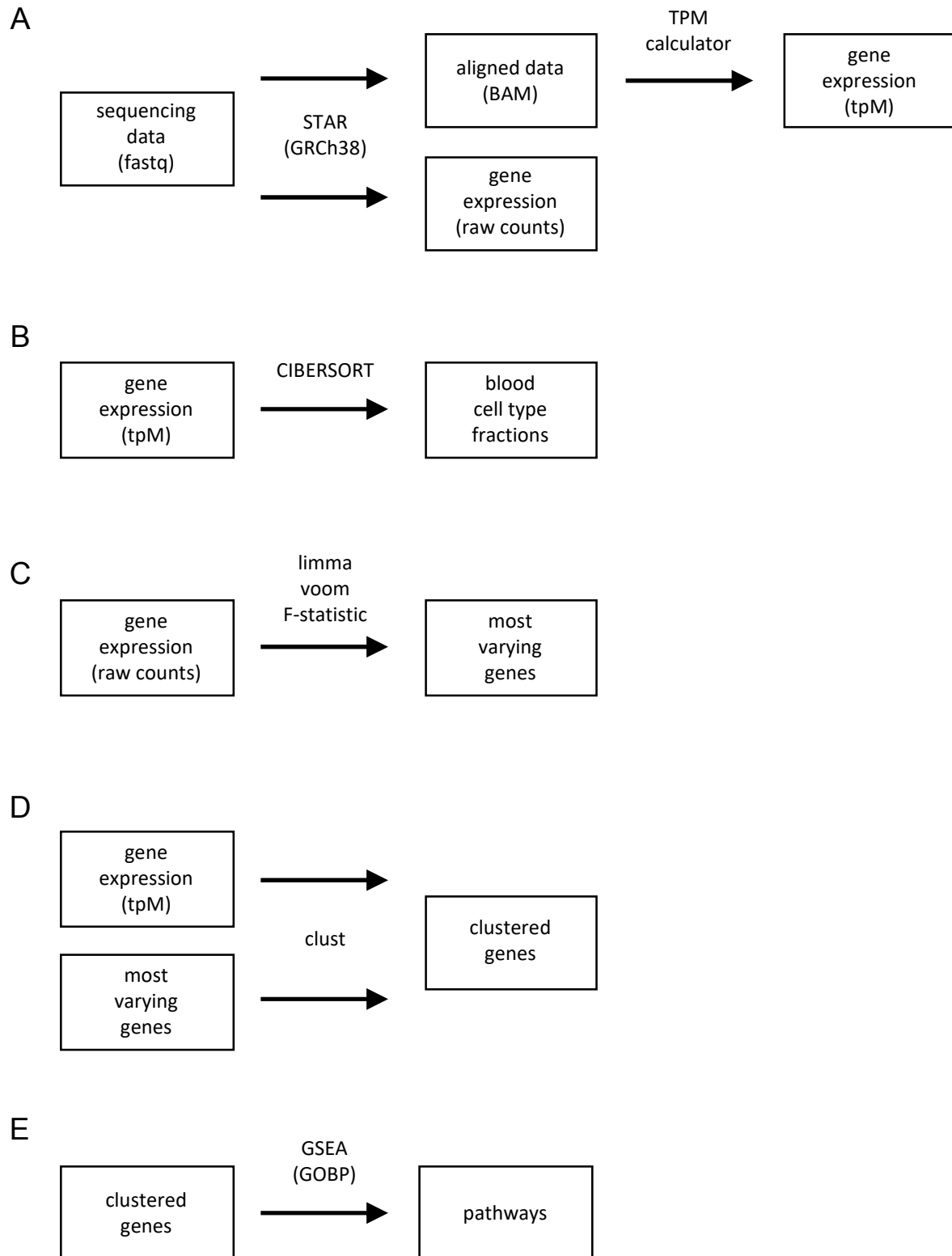




Supplementary Figure 1. Blood transcriptomics bioinformatics pipeline. (A) Deep-sequencing data stored in fastq files were mapped on the human genome version GRCh38 using STAR. STAR produced two output files per sample: the sequencing data

aligned on the human genome stored in a BAM file, and the raw number of aligned reads on each gene. The BAM files were further analyzed by TPMcalculator to quantify expression in transcripts-per-million (tpM), an expression unit normalized by RNA library size and gene lengths, for each gene. **(B)** The CIBERSORT algorithm inferred immune cell fractions in each sample from a deconvolution of the bulk gene expression data (in tpM units) on a database of 22 different and single-cell immune cell transcriptomes. **(C)** The raw counts were used for input to the limma-voom pipeline tandem for the detection of gene expression change with time, based on a moderated analysis-of-variance (ANOVA) test. **(D)** The top 2,500 genes most varying with time (largest F-statistics) were analysed with clust, the method designed to identify genes that are co-regulated in the same way, defining clusters of genes, from gene expression data in tpM units. **(E)** Each gene cluster, usually comprised of several hundreds of genes, was streamed into gene sets enrichment analyses (GSEA) to detect regulated pathways with an overrepresentation of genes in the cluster.