

Infernal development

Eric Nawrocki and Claude
Benasque RNA meeting, 16 July 2026

National Library of Medicine
National Institutes of Health



Talk outline

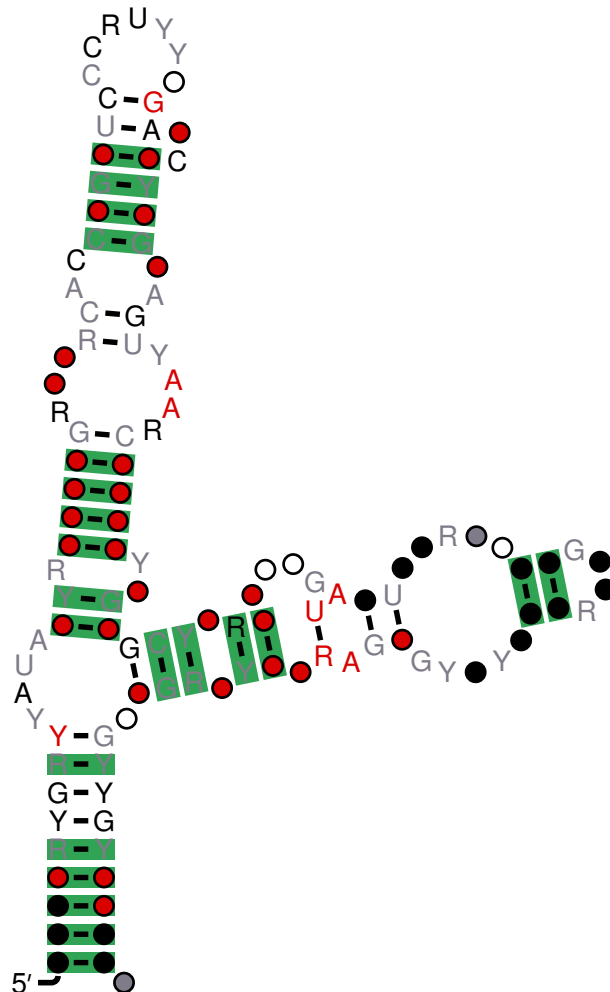
1. Introduction to CMs/Infernal
2. Eliminating cmcalibrate
3. Pseudoknot alignment markup

Infernal implements CMs: statistical models of an RNA family

Search sequence databases for (or align) homologs of an RNA family using *both* primary sequence conservation *and* secondary structure conservation

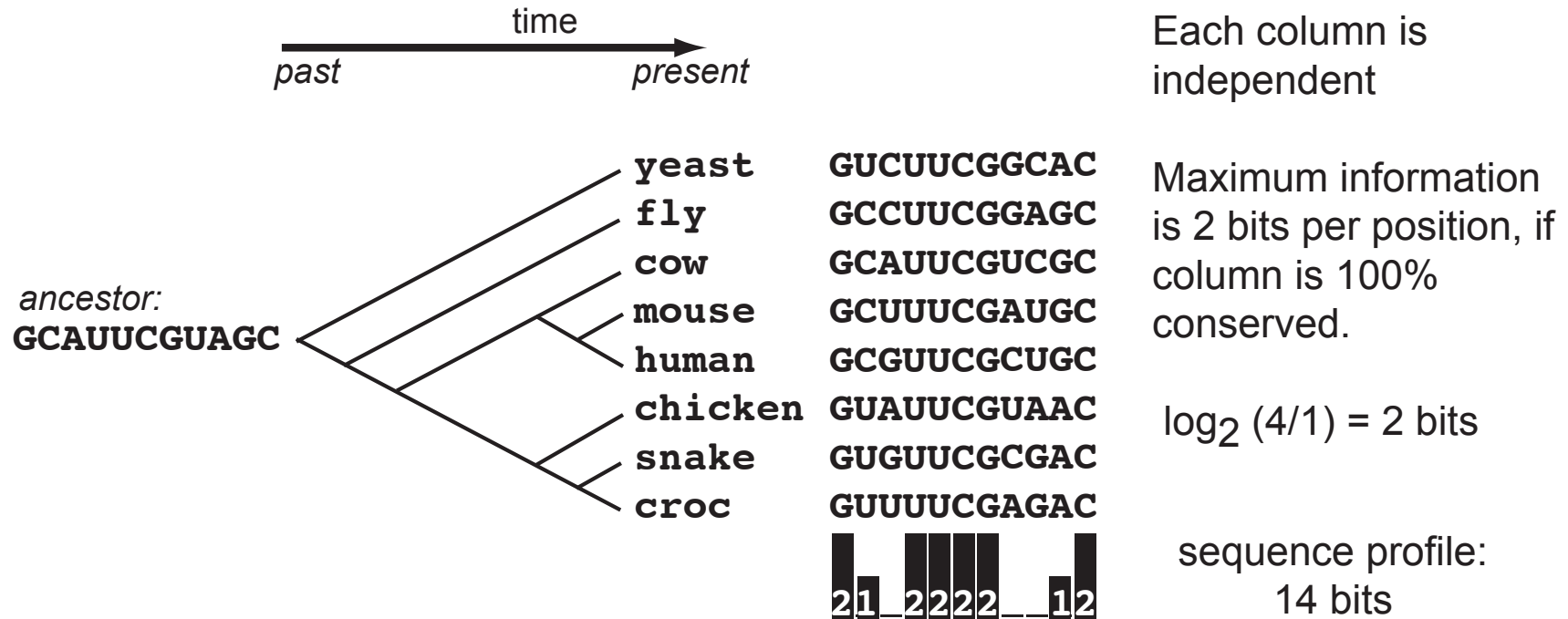
using covariance models (CMs) - profile stochastic context-free grammars

RF00001_5S_rRNA



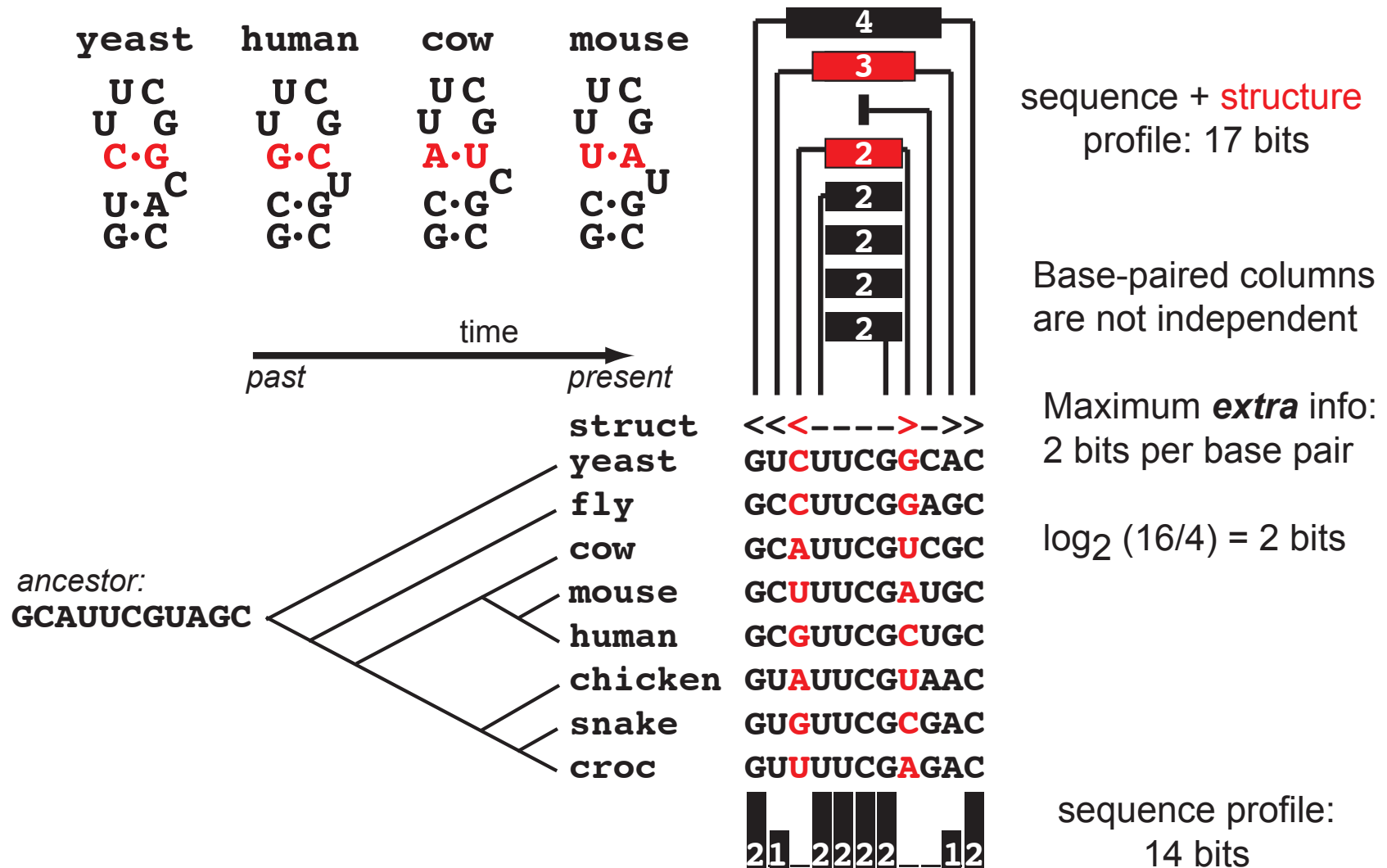
Nucleotide present		Nucleotide identity	
● 97%	● 75%	N	97%
● 90%	○ 50%	N	90%
		N	75%

Amount of information in a profile can be measured in bits



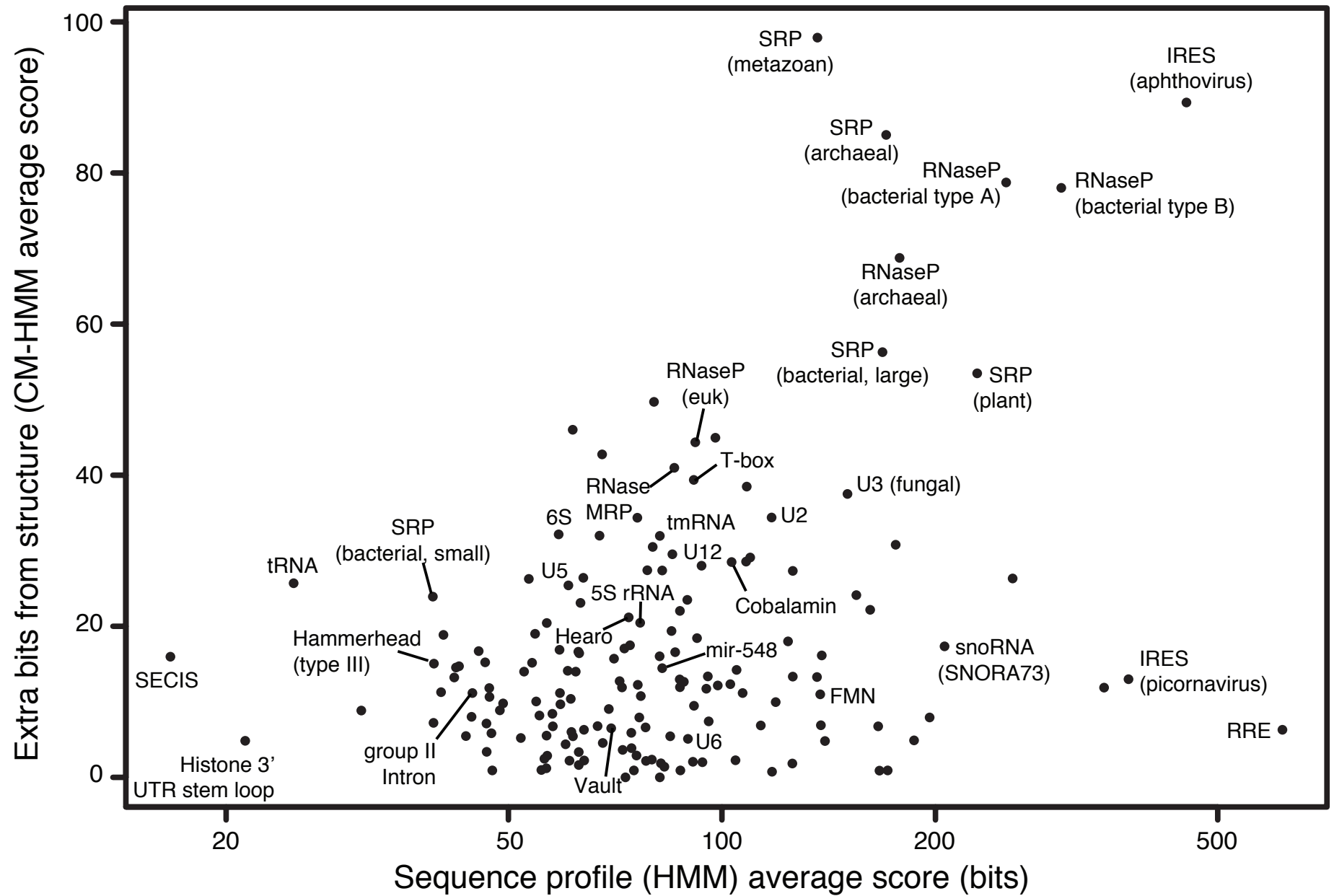
expect a match by chance: 1 in 2^{14} nt $\approx$ 16 Kb

Structure contributes additional information from covariation

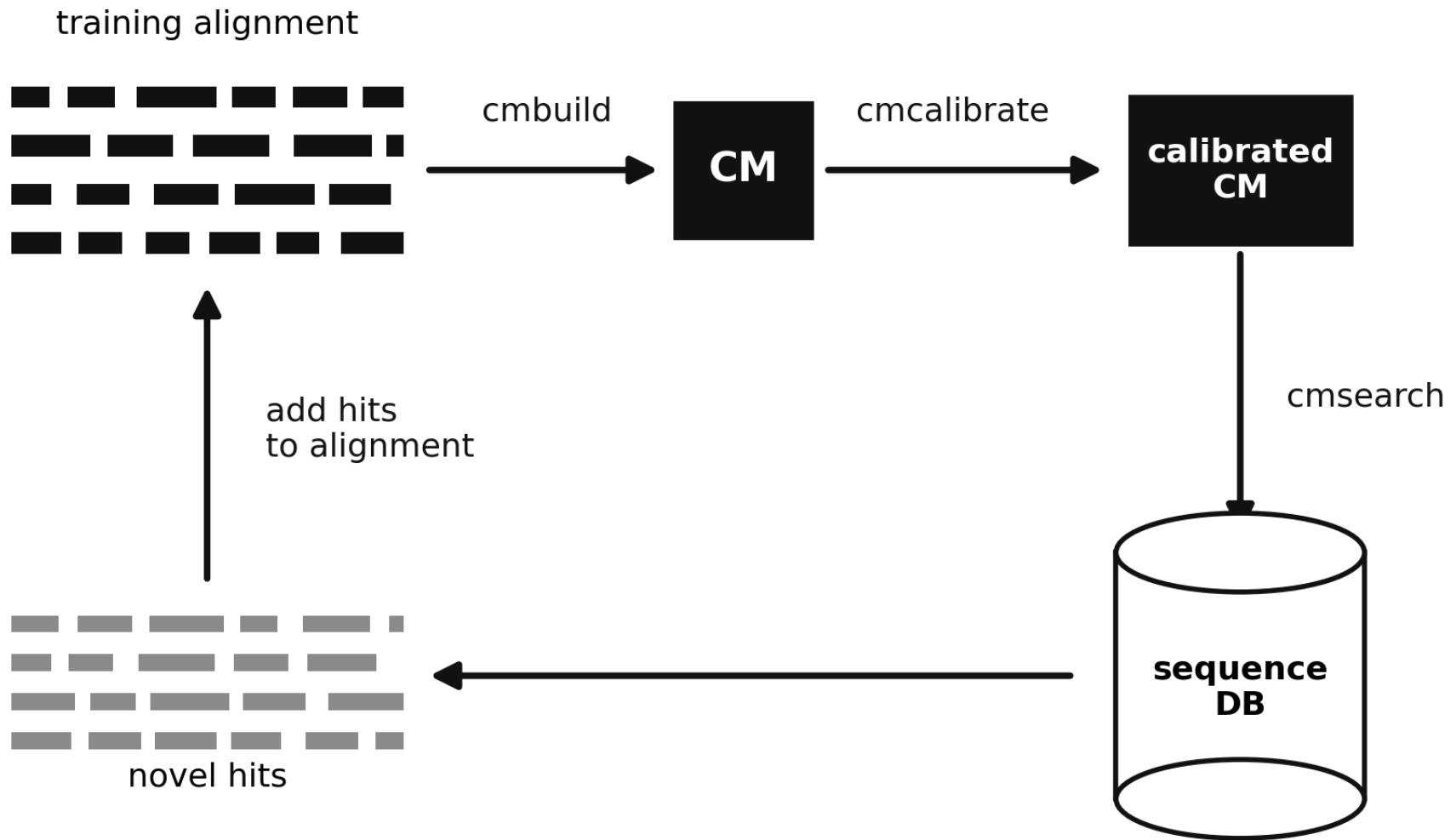


expect a match by chance: 1 in 2^{17} nt $\approx$ 130 Kb
 reducing expected false positives by $2^3 = 8$ -fold

Levels of sequence and structure conservation in RNA families

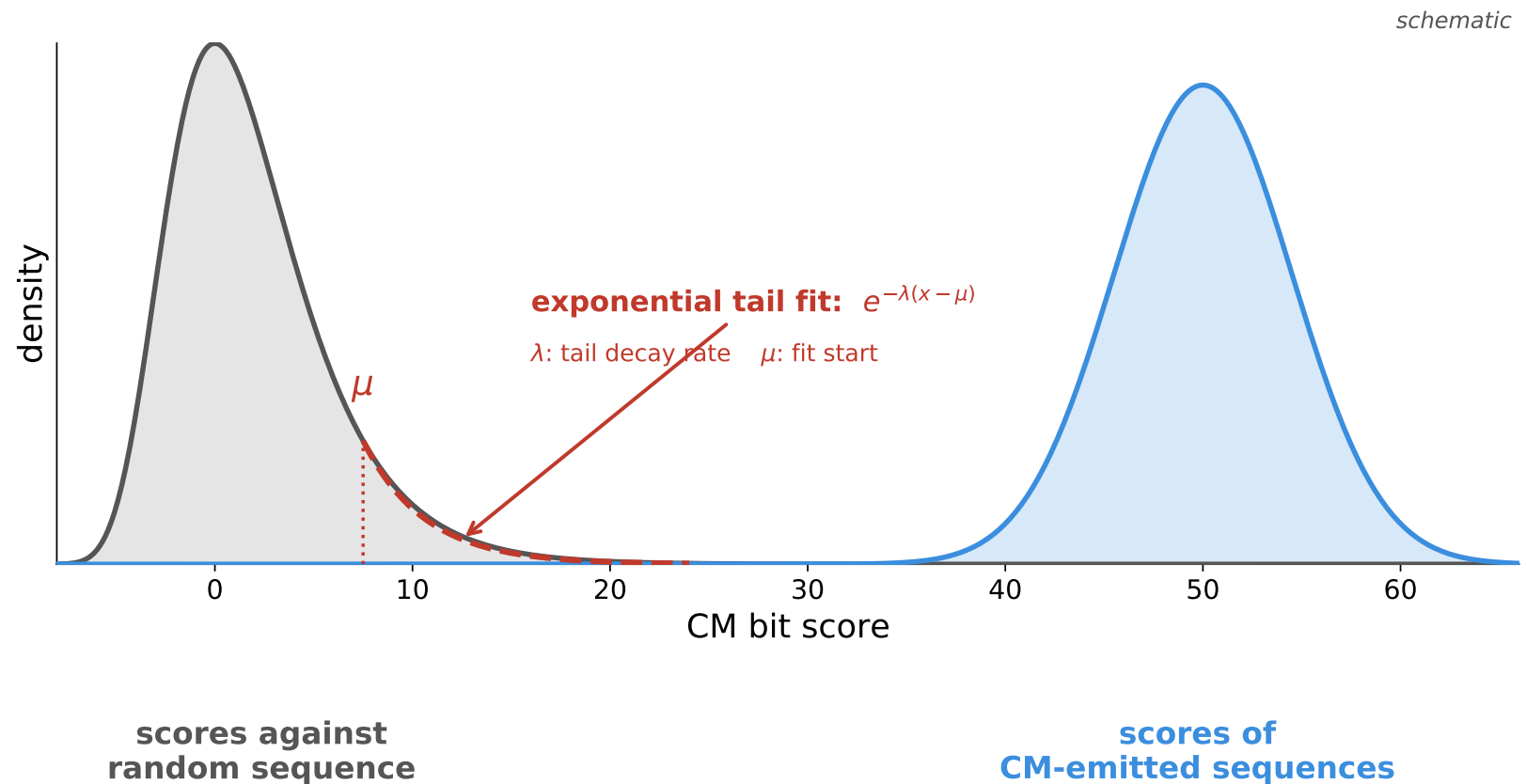


CM workflow



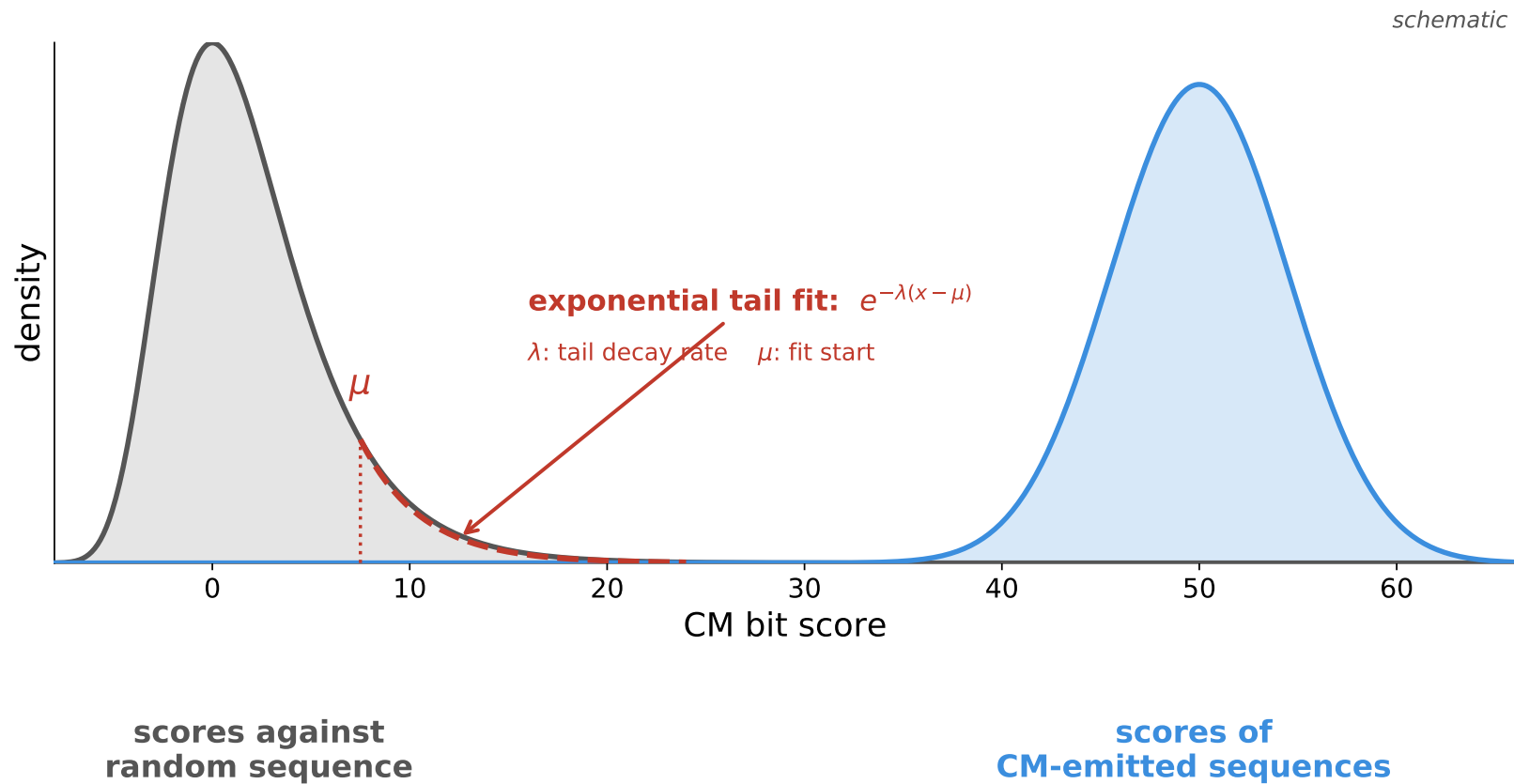
Calibration enables E-value reporting

- to assign *E-values* to search hits
- E-value = expected number of hits at a given score in random, nonhomologous sequence



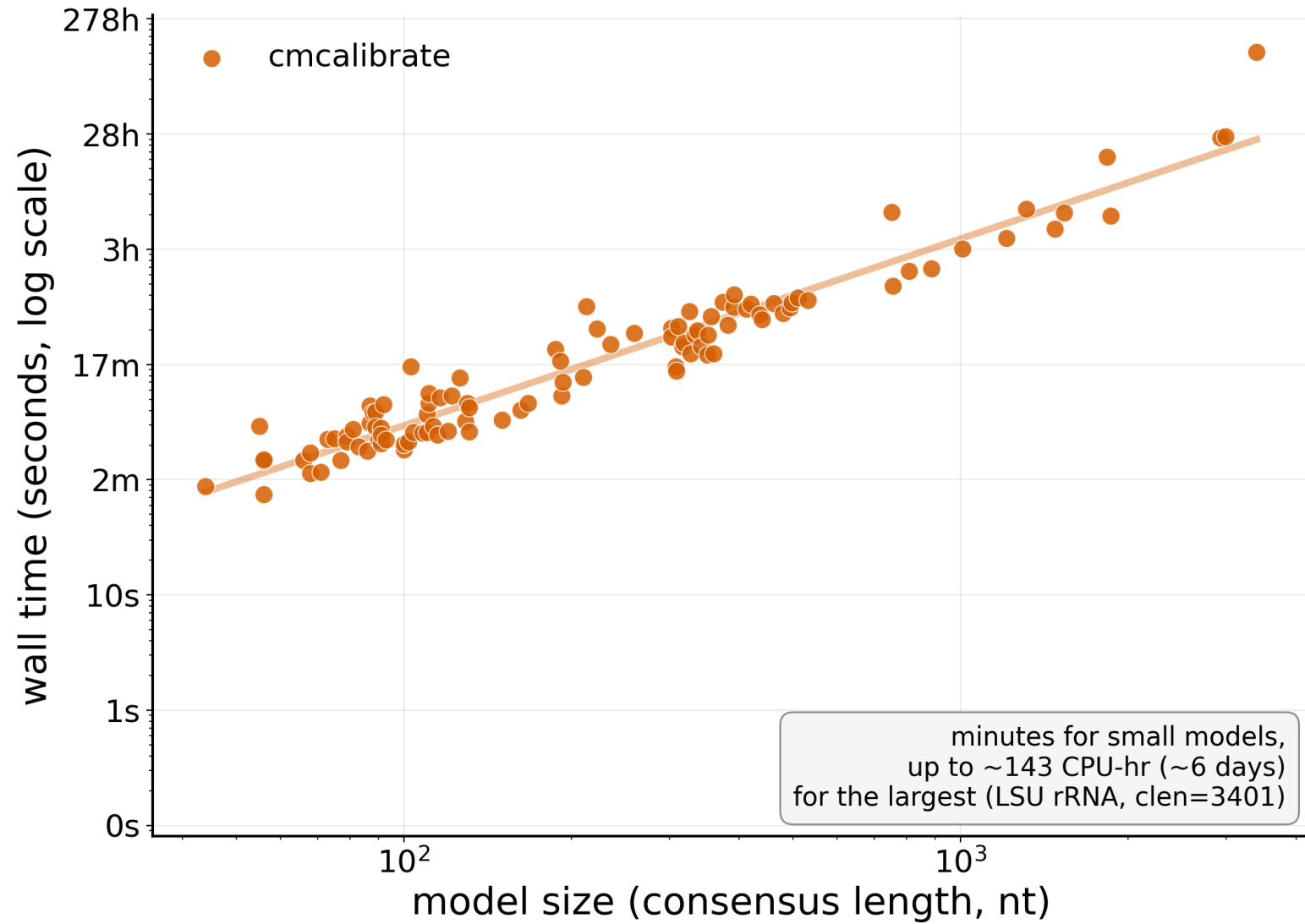
Calibration is slow

- to fit far tail's slope (λ) we need search a lot of random sequence
- CM search DP (CYK / Inside) scales as $O(LN^2 \log N)$, vs $O(LN)$ for a profile HMM (N = model length, L = target length)
- we speed up *cmsearch* with HMM filters and HMM bands* - but calibration can't use them: the filters remove nearly all random sequence

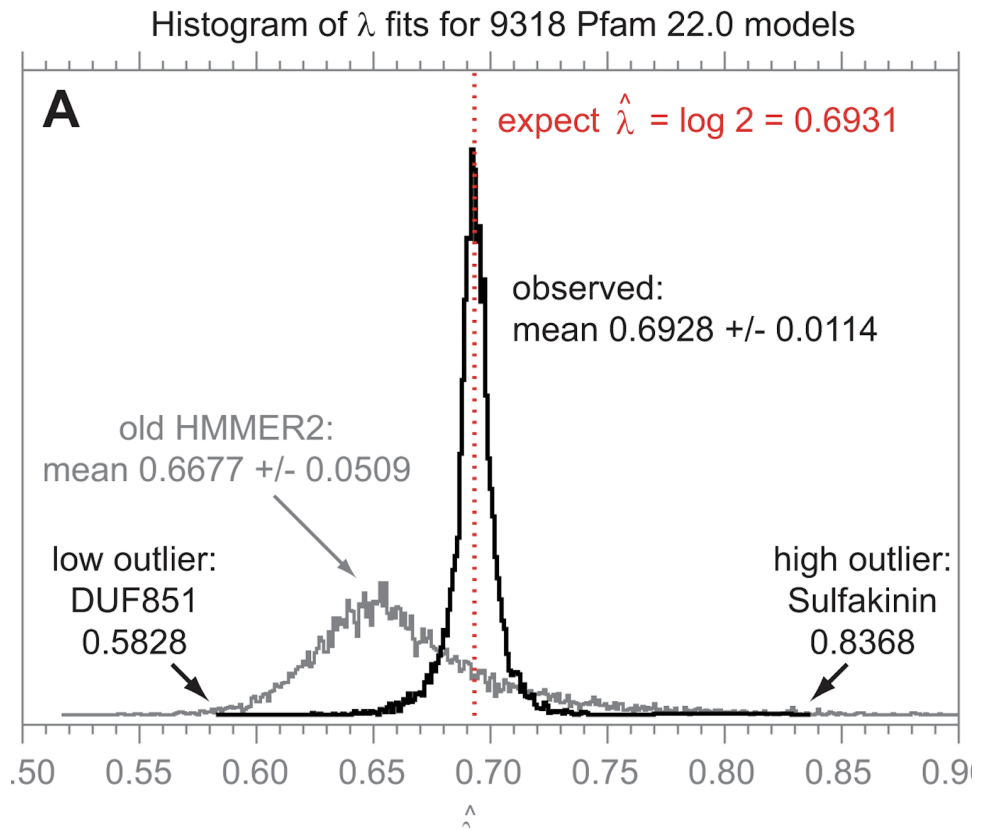
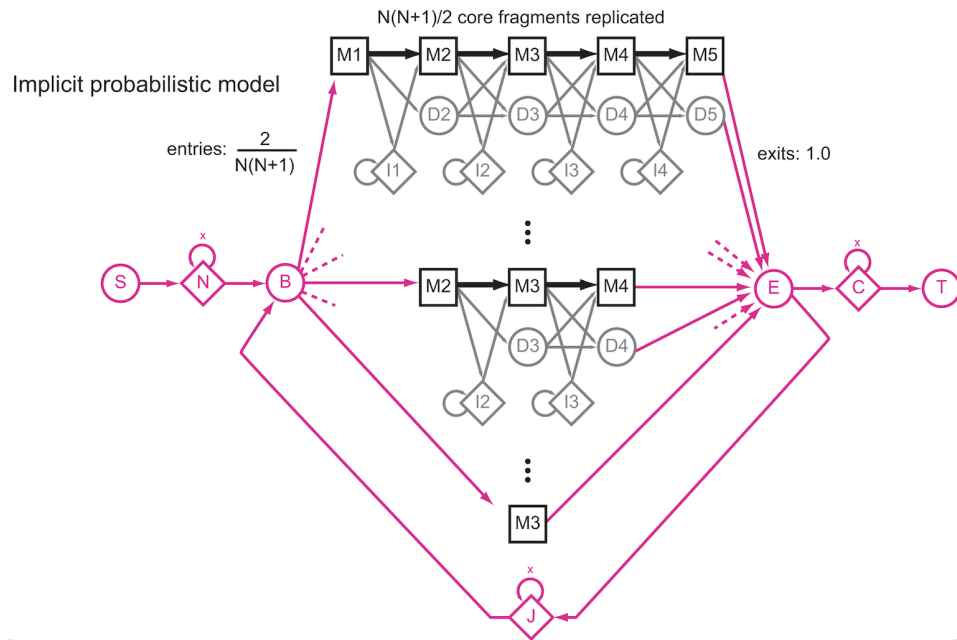


*E.P. Nawrocki & S.R. Eddy, "Infernal 1.1: 100-fold faster RNA homology searches," *Bioinformatics* **29**(22):2933–2935 (2013).

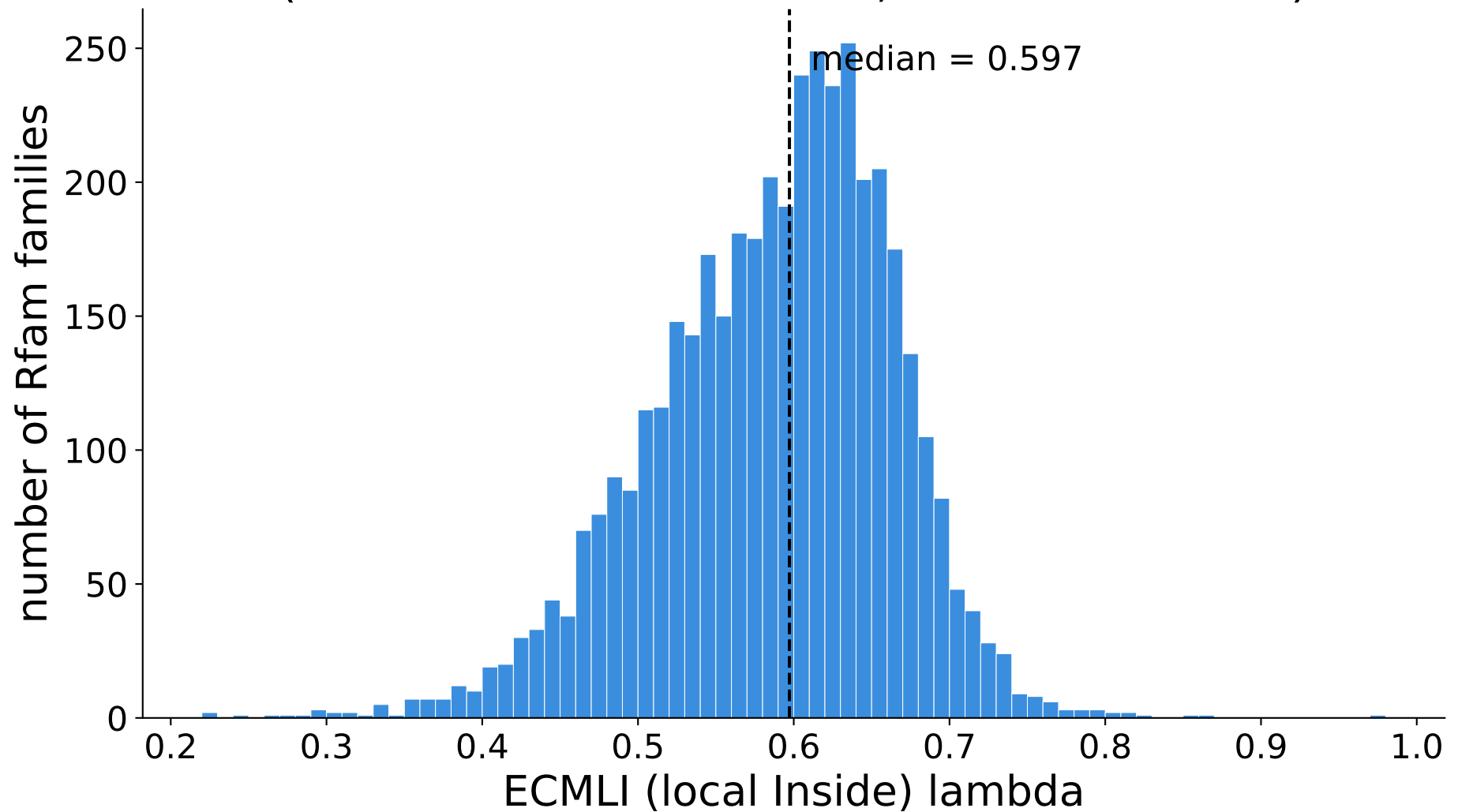
Calibration is slow



HMMER3 gets E-value statistics analytically

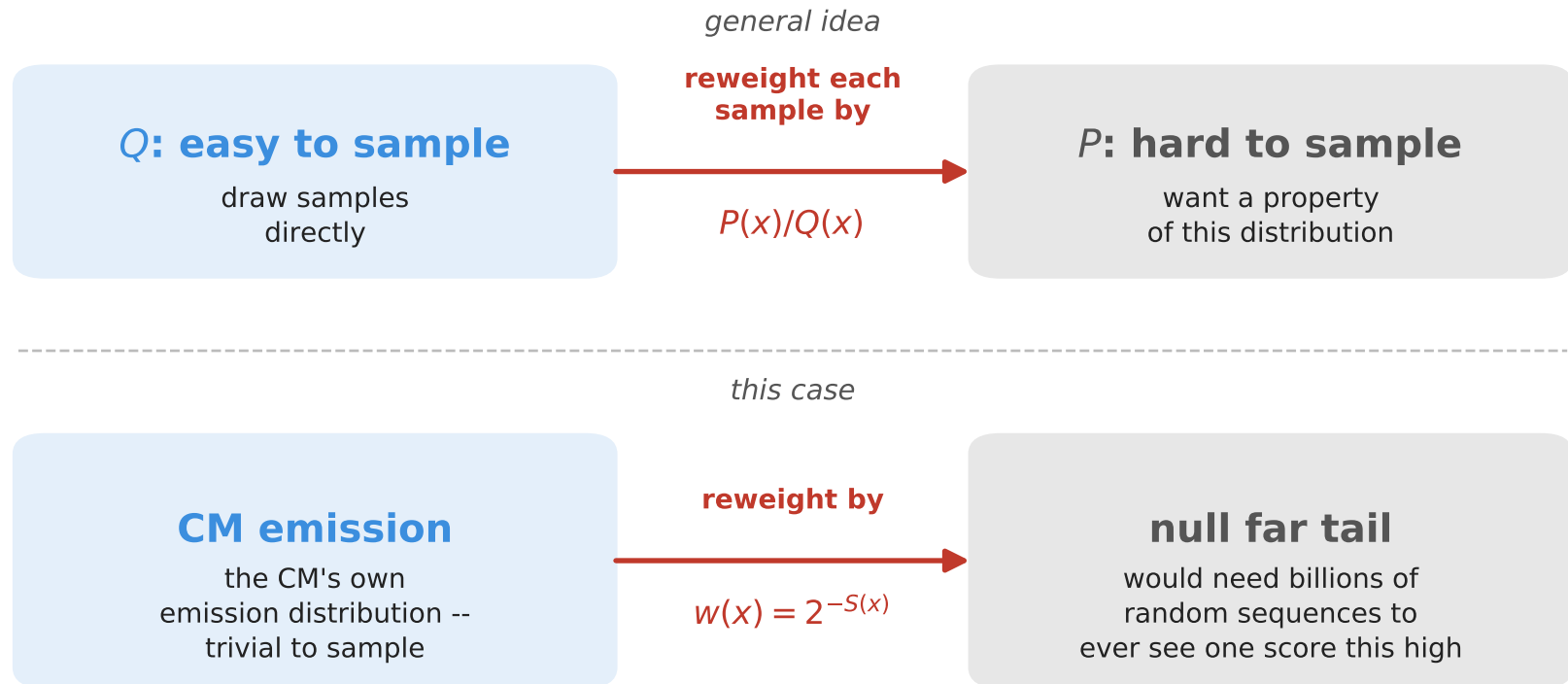


Rfam 15.1: local-Inside E-value lambda
(N = 4227 calibrated CMs, bin width = 0.01)



Importance sampling: sample the easy distribution, reweight to the hard one

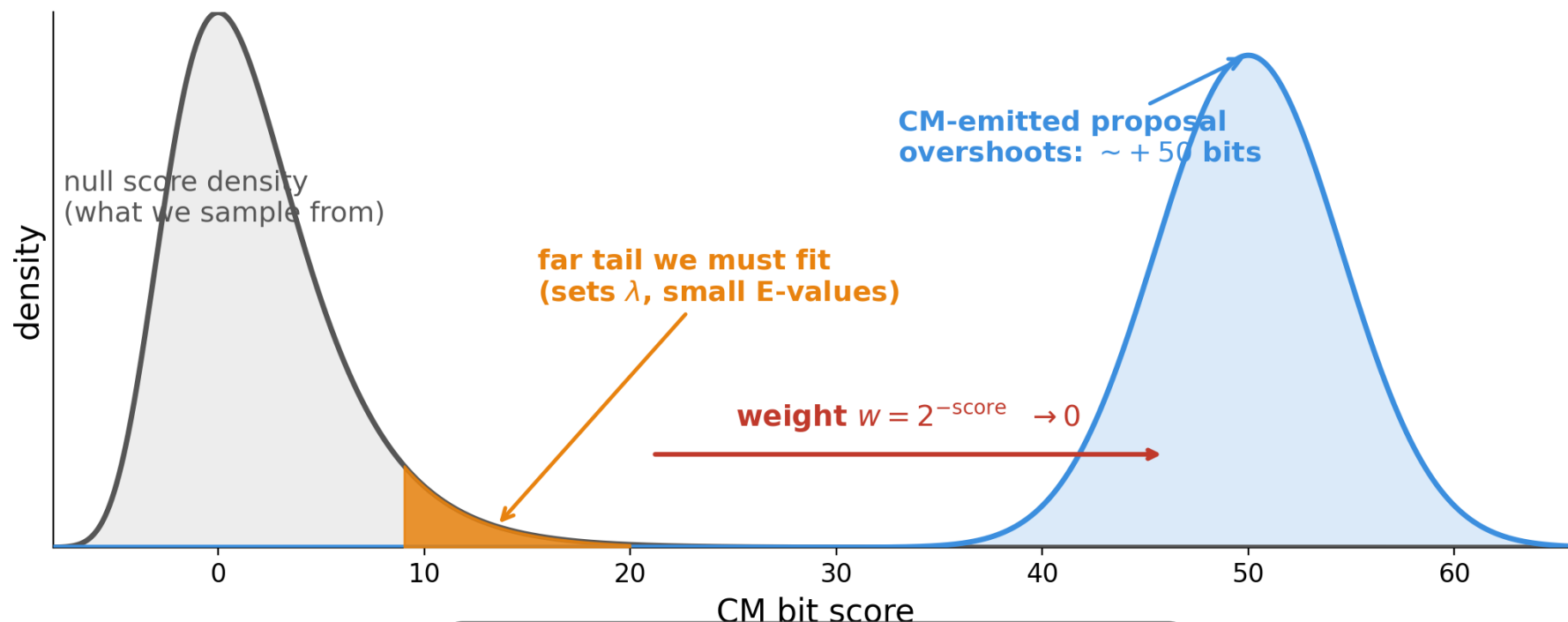
- want a property of a distribution P that's hard to sample directly
- instead sample an easy Q , and reweight by $P(x)/Q(x)$



$$S(x) = \log_2 \frac{P_{\text{CM}}(x)}{P_{\text{null}}(x)} \quad (\text{the CM bit score is itself the log-odds of emission vs. null})$$

Importance sampling fails: distributions too far apart

schematic

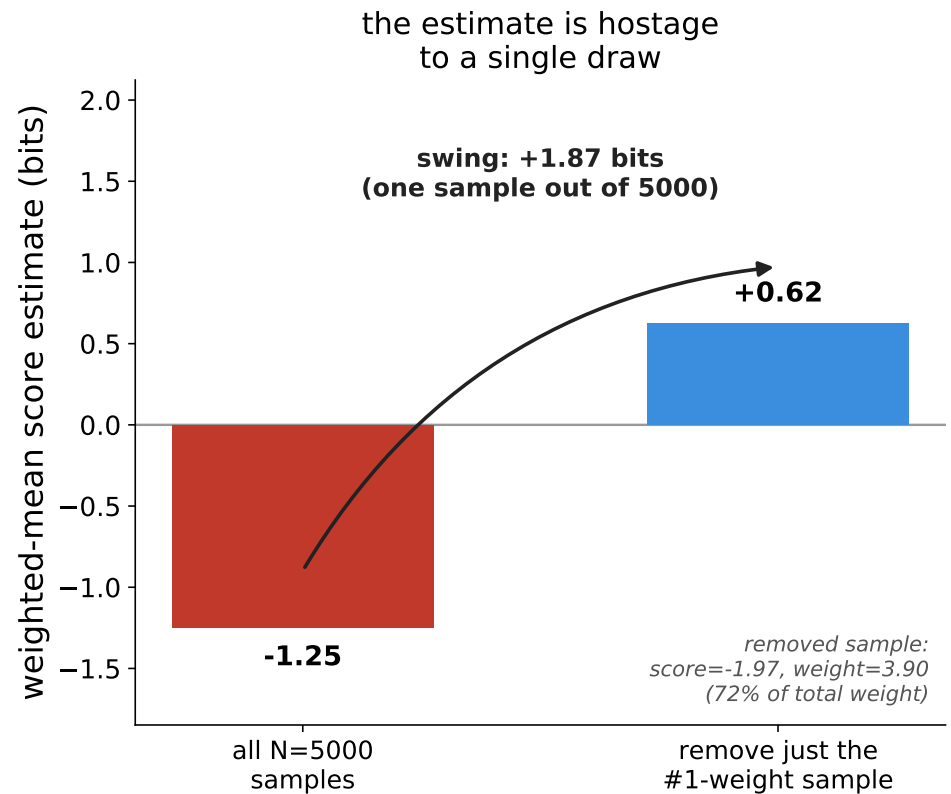
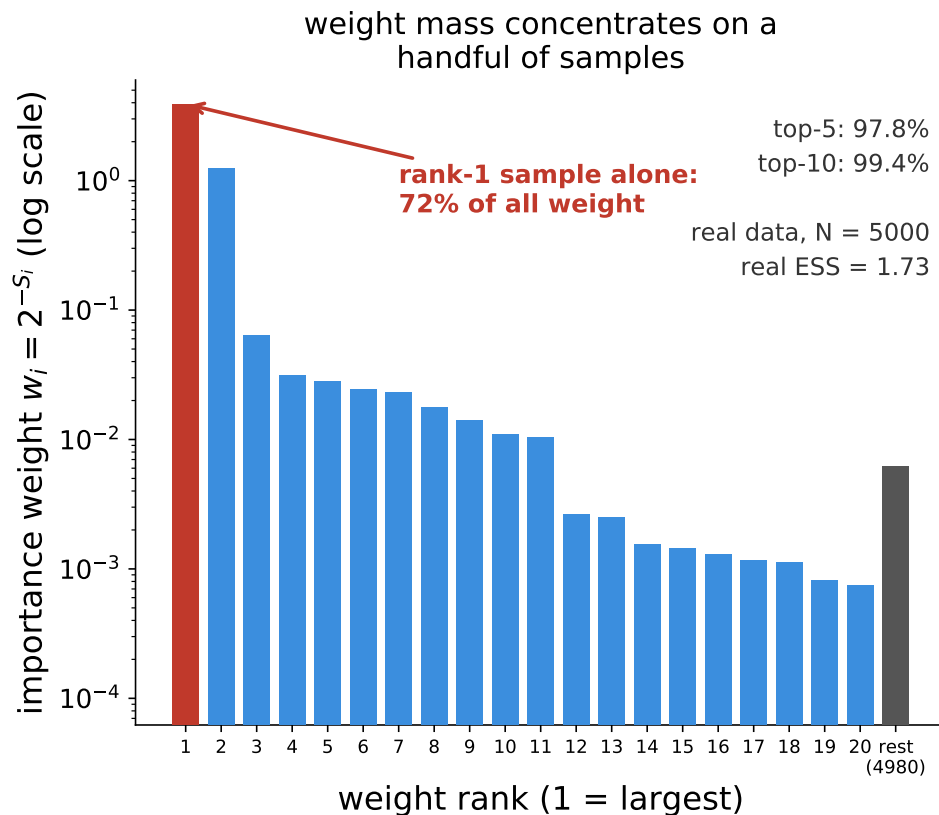


$$w(x) = \frac{P_{\text{null}}(x)}{q(x)} = 2^{-\text{score}} \quad \text{ESS} = \frac{(\sum_i w_i)^2}{\sum_i w_i^2} \approx 1$$

a few sequences carry all the weight $\Rightarrow$ effective sample size ≈ 1 — worst for large / global models

Few samples hold the whole estimate hostage

- real N=5000 sequences emitted from the 5S rRNA CM (RF00001); weight $w_i = 2^{-S_i}$ per sample; real ESS = 1.73 out of N=5000
- the single largest-weight sample carries 72% of all weight – removing just that one sample swings the weighted-mean estimate by +1.87 bits



I tried to fix this various ways — none worked

attempted fix	result
flatten the CM's probabilities, to sample closer to random sequence	still collapses: ESS stays ≤ 2
reject any emitted sequence outside a target score window	biased the fitted slope $3\times$ (0.14 vs. true 0.42)
design/mutate sequences to force them into the tail	fitted slope always converges to the same fixed value, regardless of the true one
sample directly from the tail with MCMC (no reweighting)	accurate, but cost grows quadratically with model size — slower than what it replaces for exactly the largest models

Skip sampling and predict the E-value parameters (λ , μ) directly

- λ , μ aren't random, they're a function of the CM (its composition, structure, and length). So they should be predictable from model properties.

Skip sampling and predict the E-value parameters (λ , μ) directly

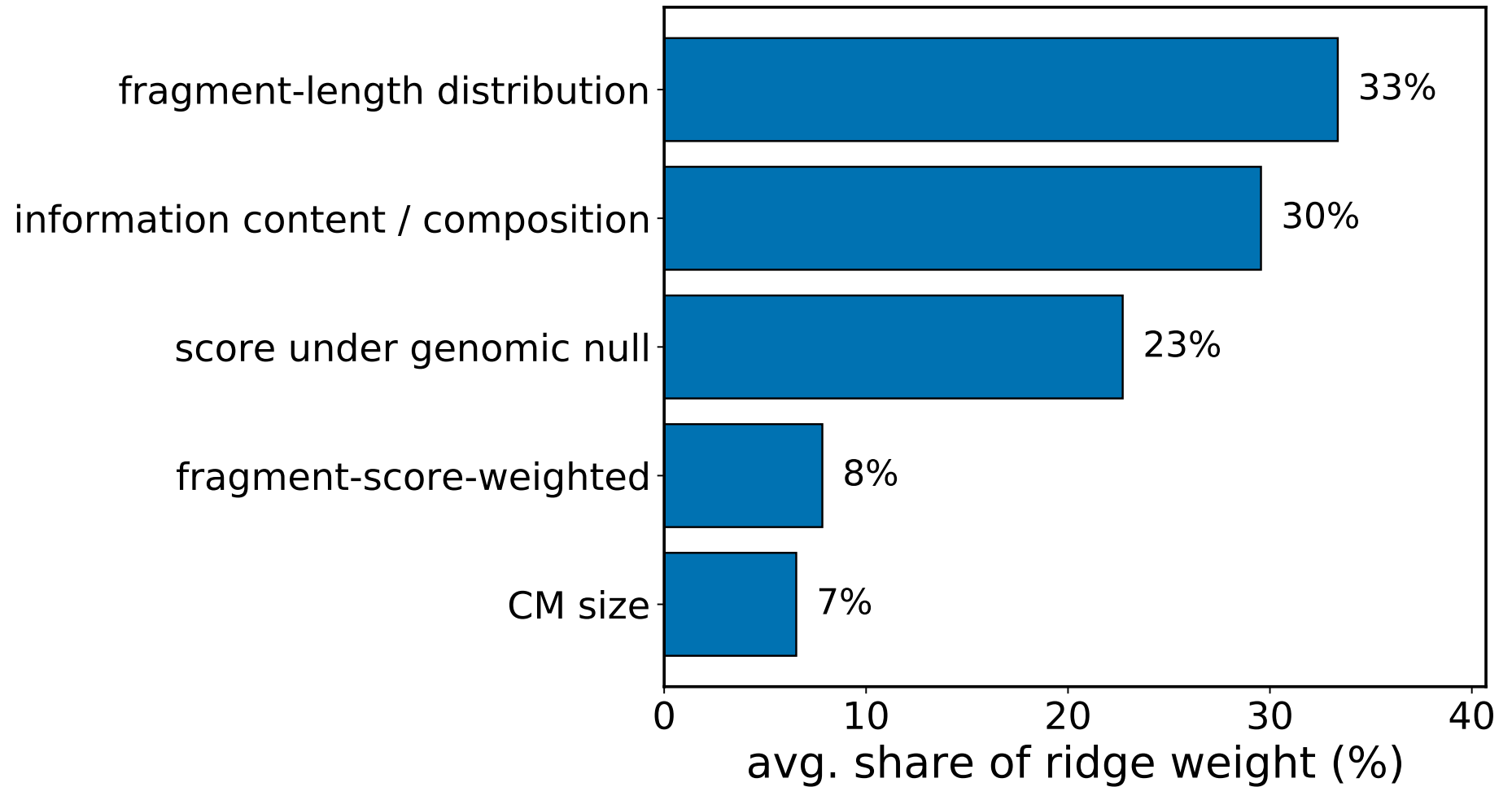
- λ , μ aren't random, they're a function of the CM (its composition, structure, and length). So they should be predictable from model properties.
- Thousands of already-calibrated Rfam CMs give us a free training/testing set: (CM $\rightarrow$ λ , μ) data points.

Skip sampling and predict the E-value parameters (λ , μ) directly

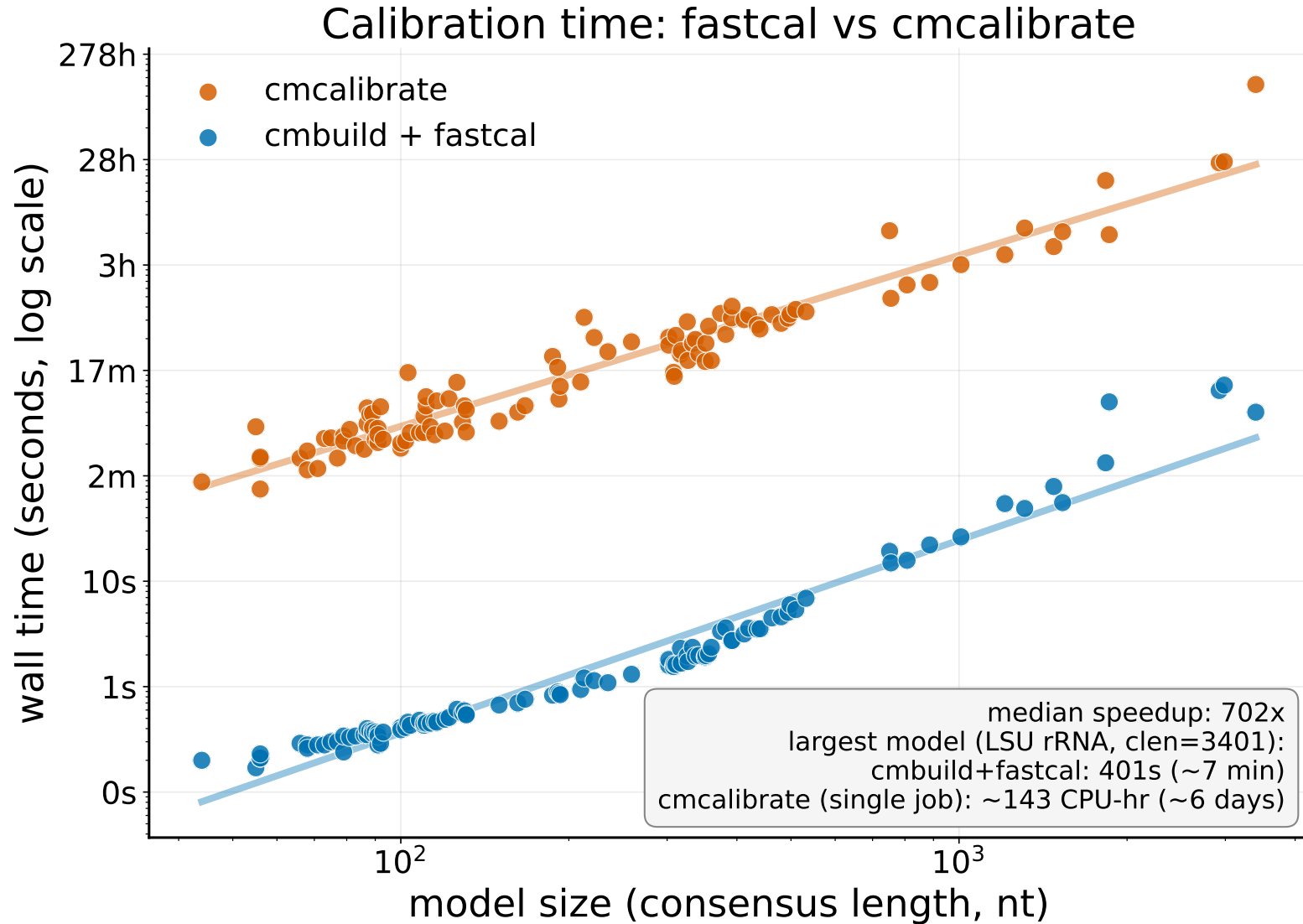
- λ , μ aren't random, they're a function of the CM (its composition, structure, and length). So they should be predictable from model properties.
- Thousands of already-calibrated Rfam CMs give us a free training/testing set: (CM $\rightarrow$ λ , μ) data points.
- Ridge regression: linear, L2-regularized (robust with many correlated features, hard to overfit), and interpretable (you can read off what predicts what)

The most important CM features

What predicts the E-value parameters?

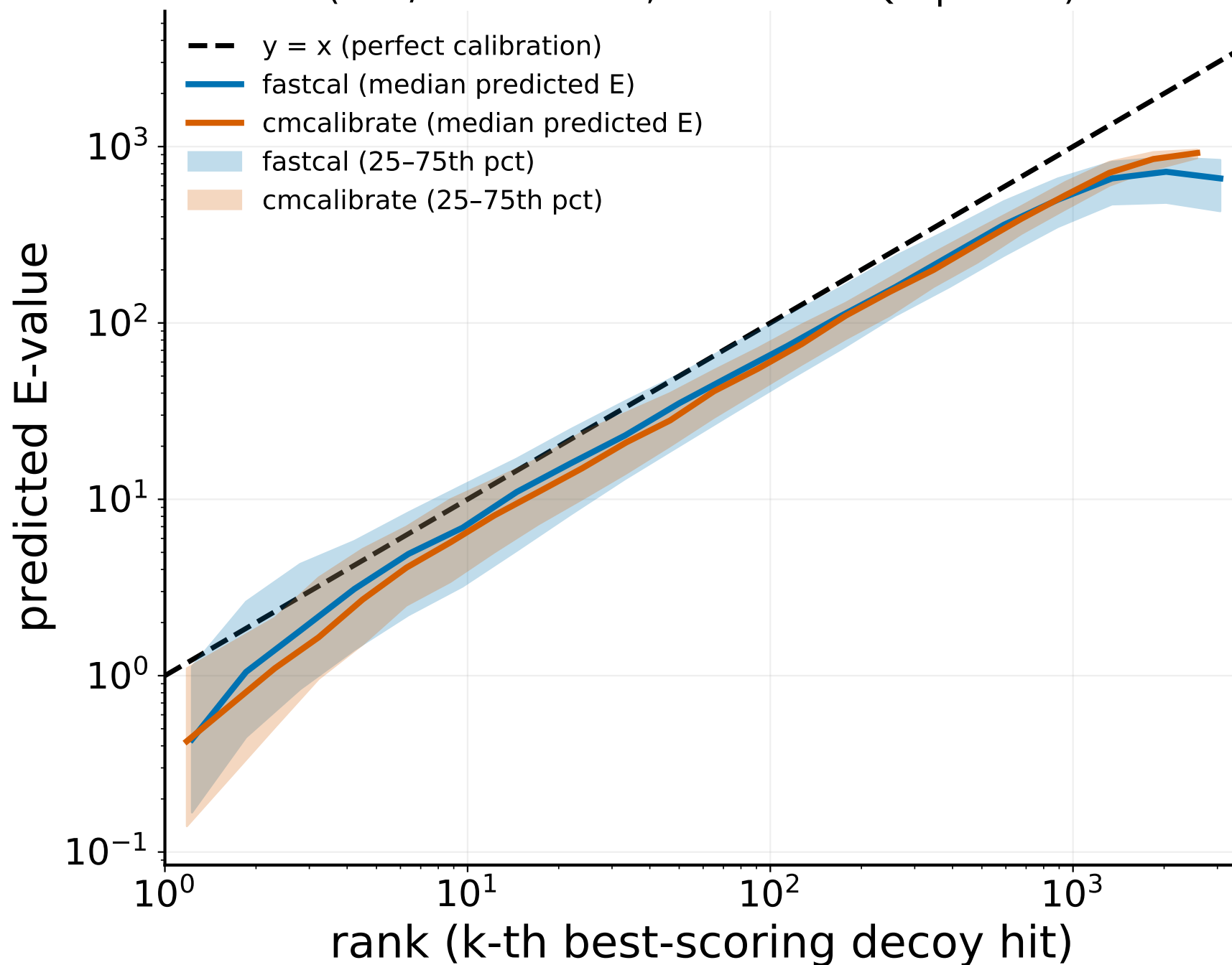


Fast



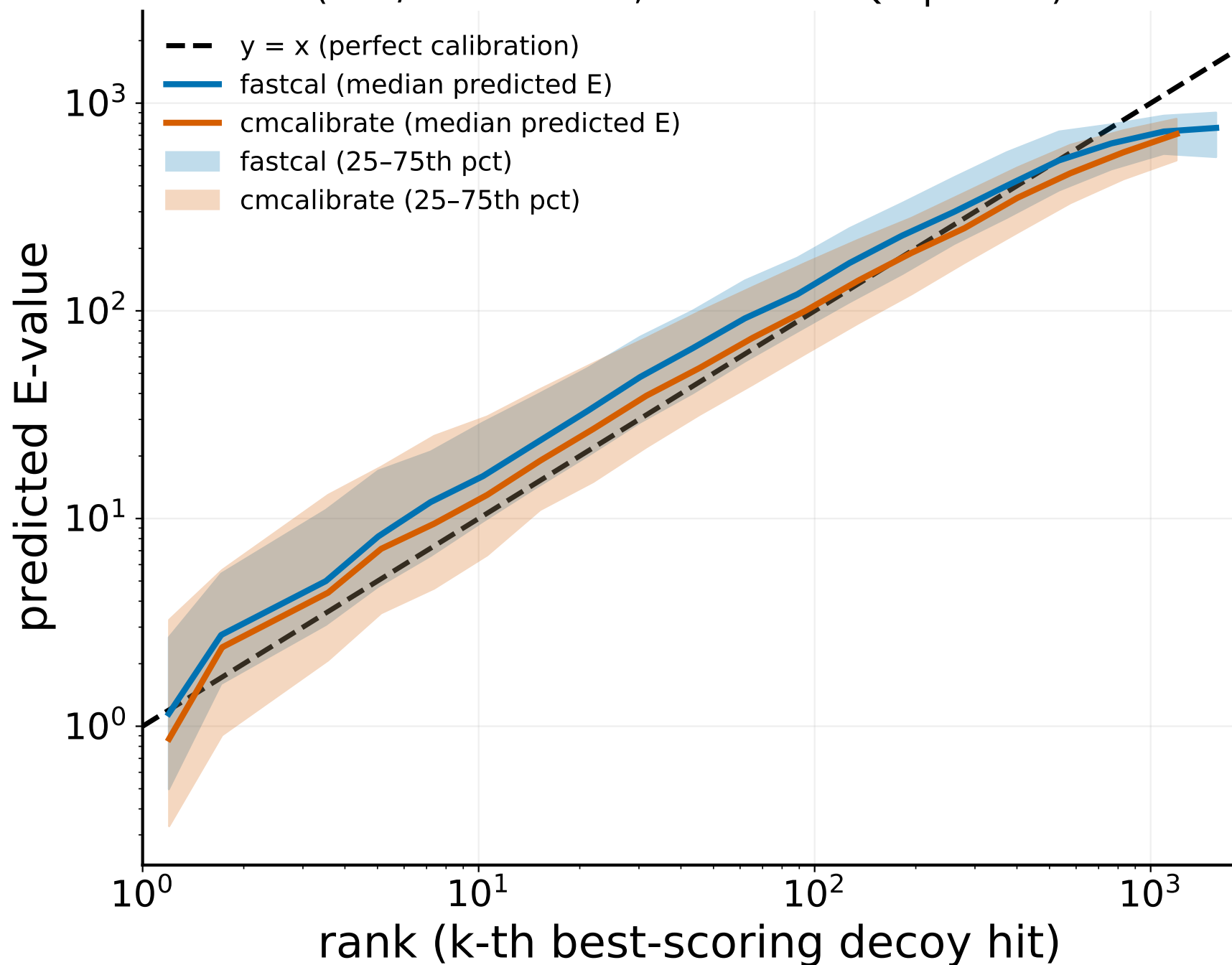
104 models, clen 44–3401 — **median speedup** $\sim 702\times$ (range 39–2634 $\times$),
conservative in fastcal's favor (cmcalibrate ran `--cpu 4`, fastcal ran
single-threaded)

Pooled calibration curve — synthetic genomic decoys (104/104 families, median $\pm$ IQR per bin)



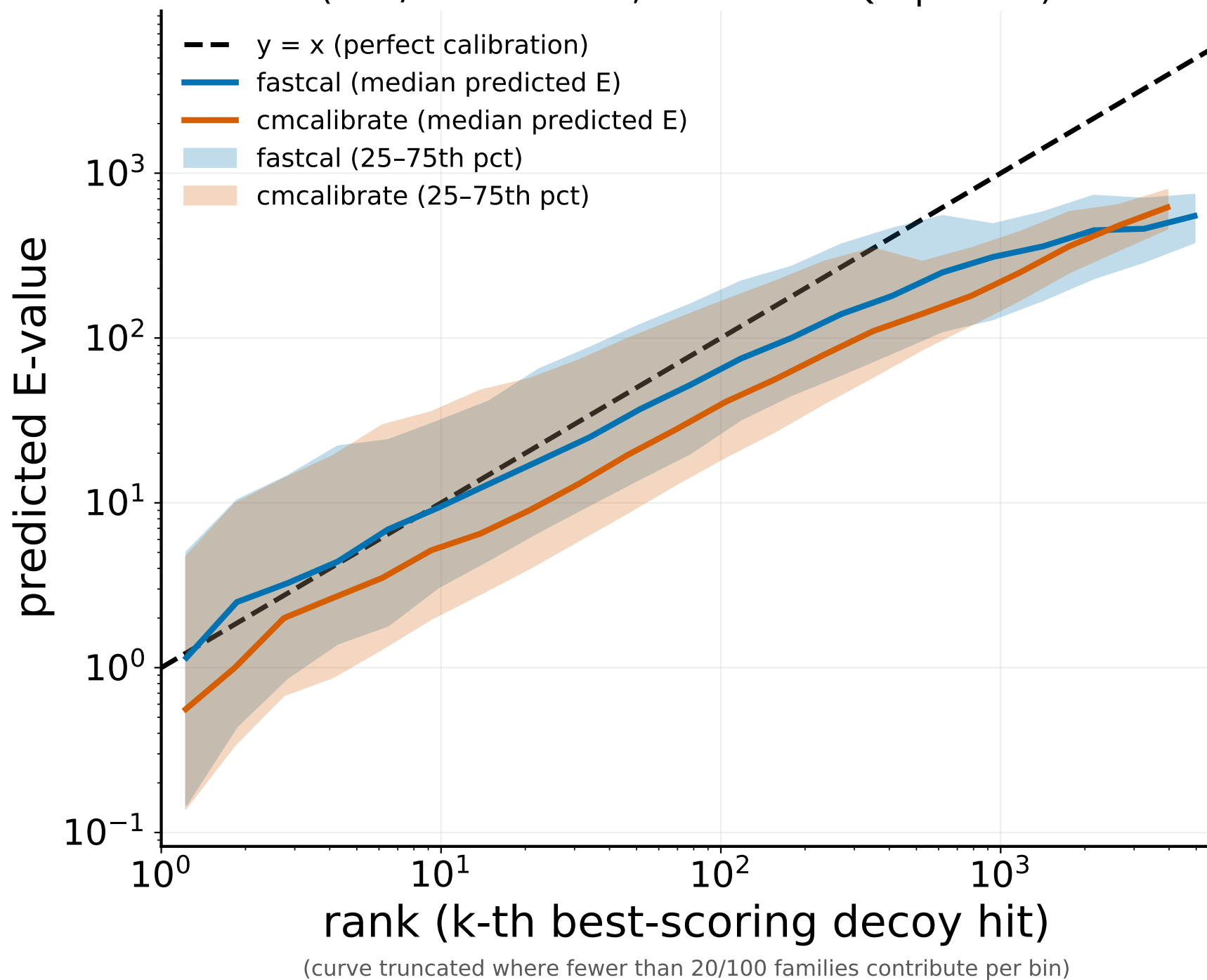
(curve truncated where fewer than 20/104 families contribute per bin)

Pooled calibration curve — chondrus genome decoys (~53% GC)
(104/104 families, median $\pm$ IQR per bin)

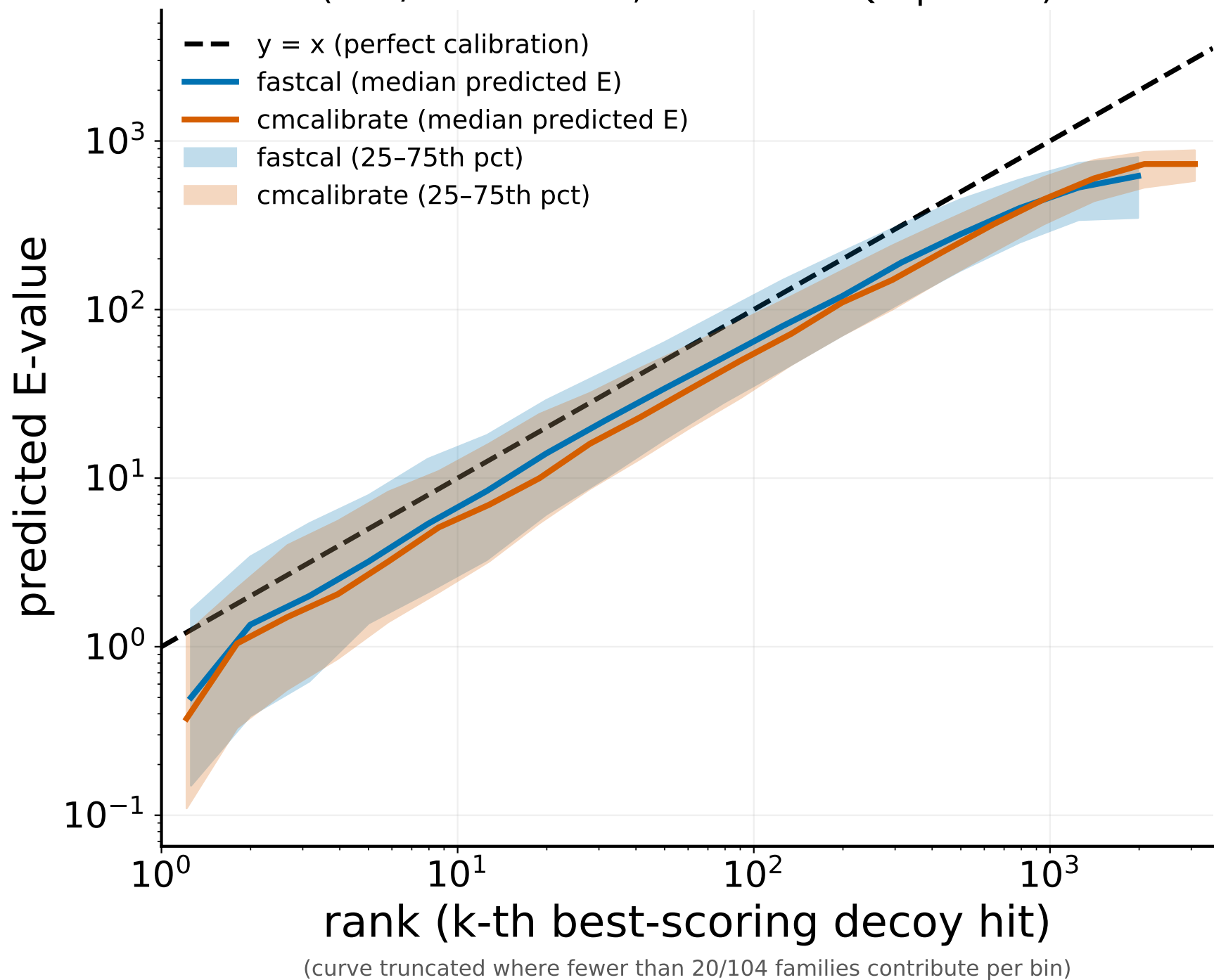


(curve truncated where fewer than 20/104 families contribute per bin)

Pooled calibration curve — chlamy genome decoys (~64% GC)
(100/104 families, median $\pm$ IQR per bin)

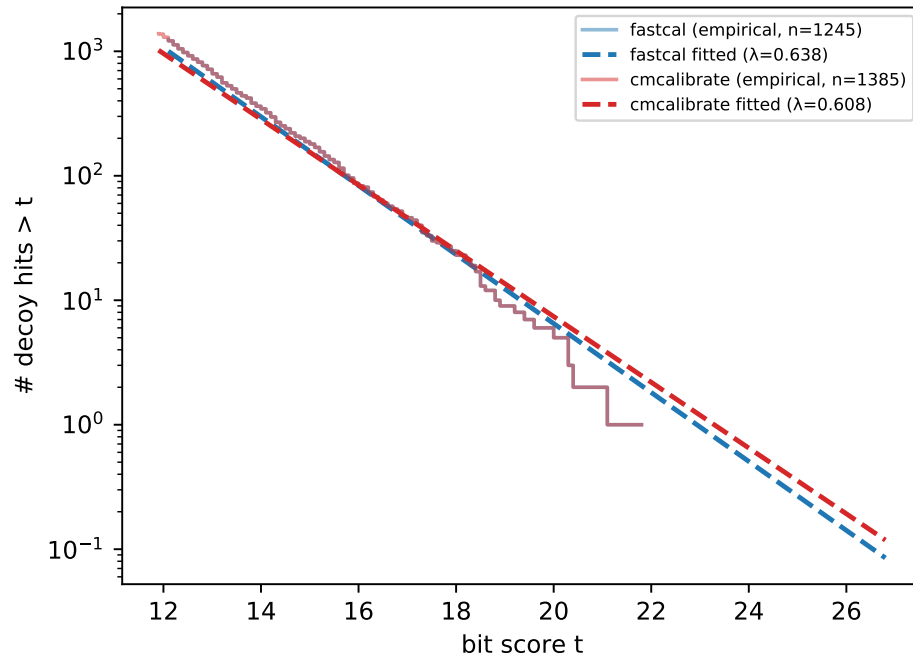


Pooled calibration curve — duck genome decoys
(104/104 families, median $\pm$ IQR per bin)

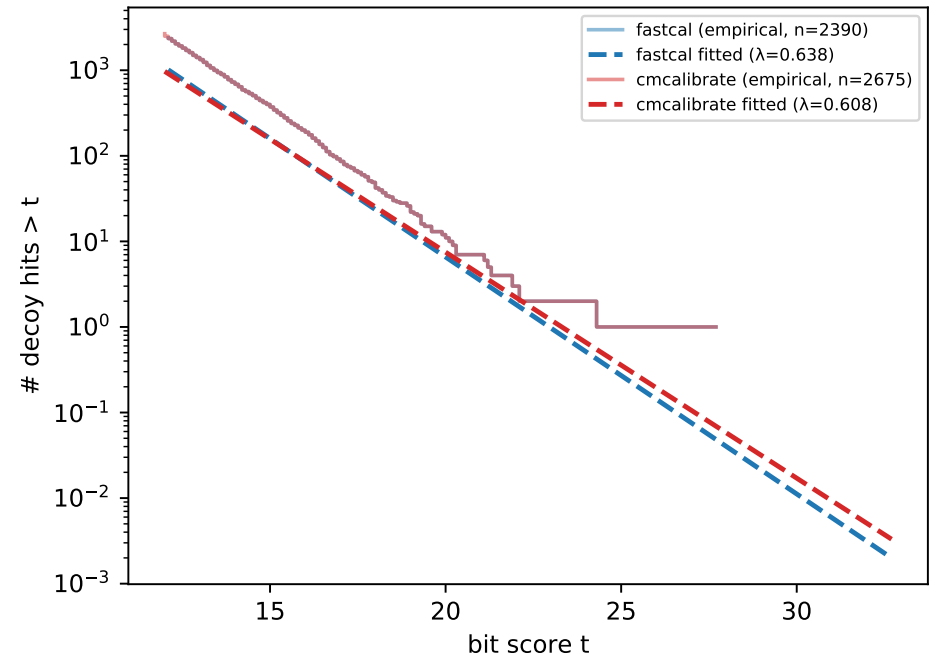


mir-5429 — calibration curve across 4 decoys (2 new real-genome + 2 original)

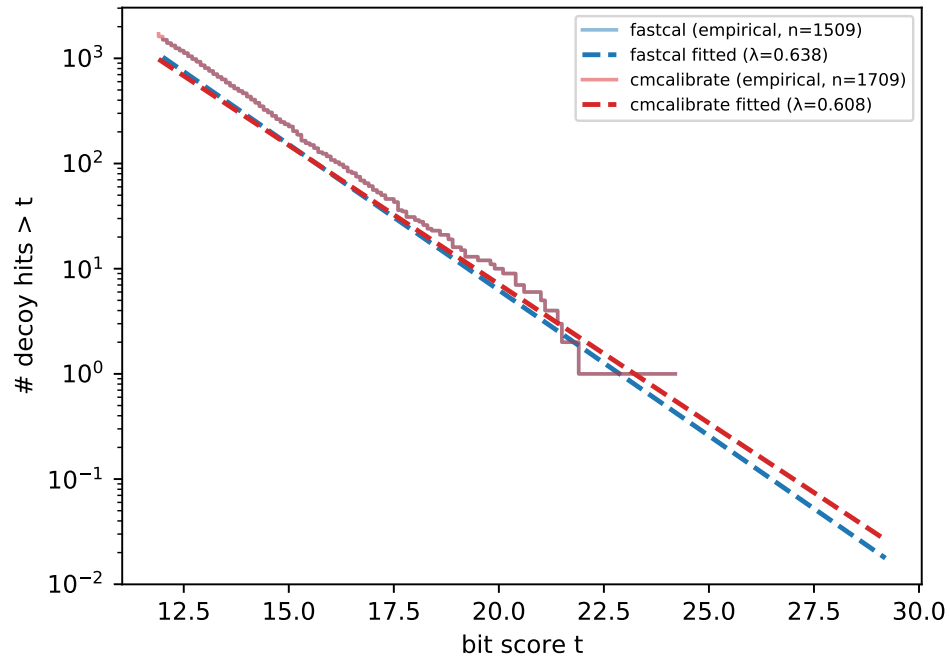
chondrus (~53% GC, real genome)



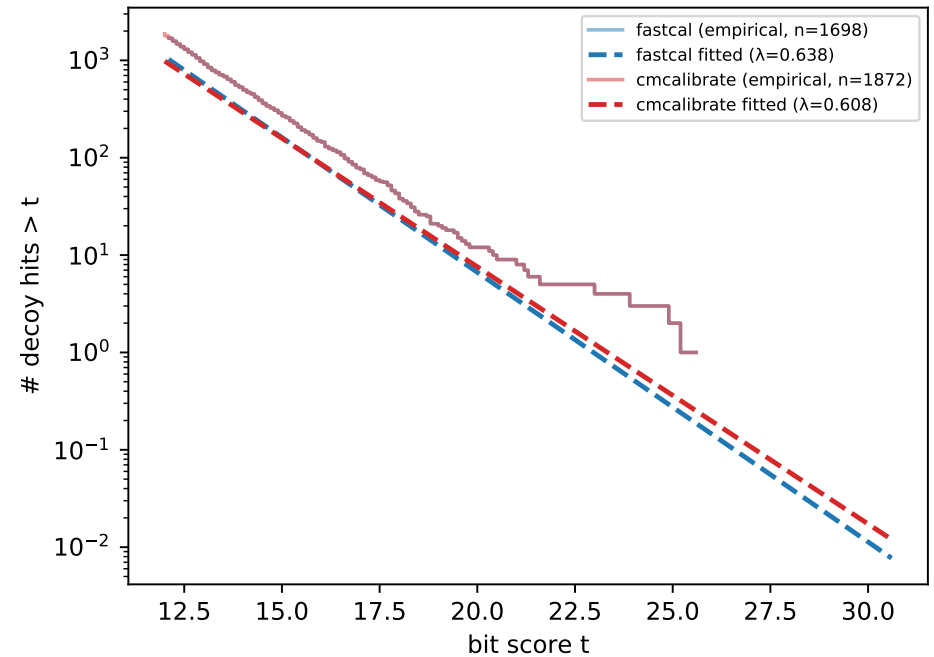
chlamy (~64% GC, real genome)



synth (~35-40% GC, HMM-generated)

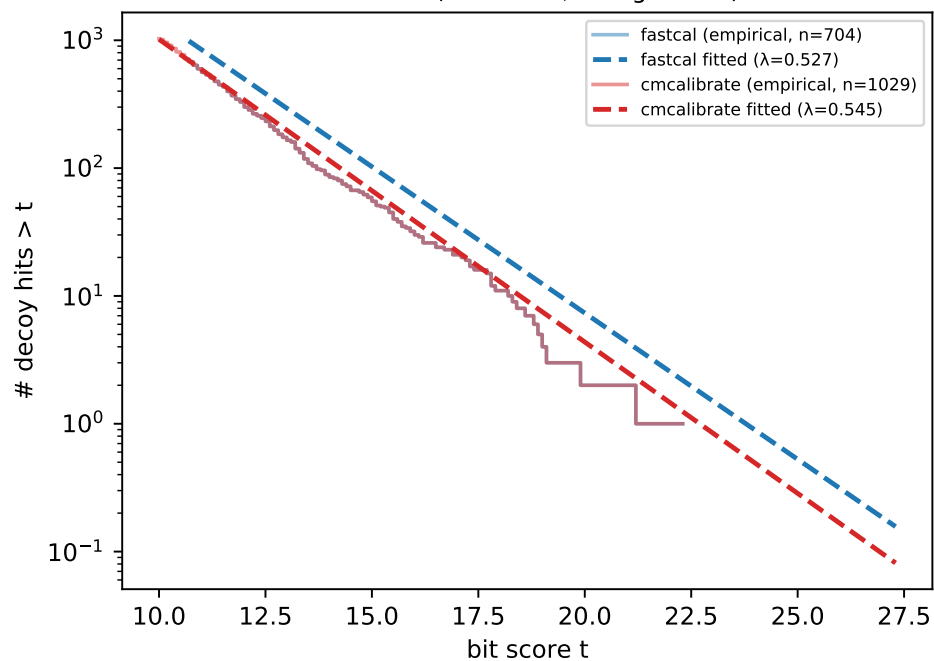


duck (~40% GC, real genome)

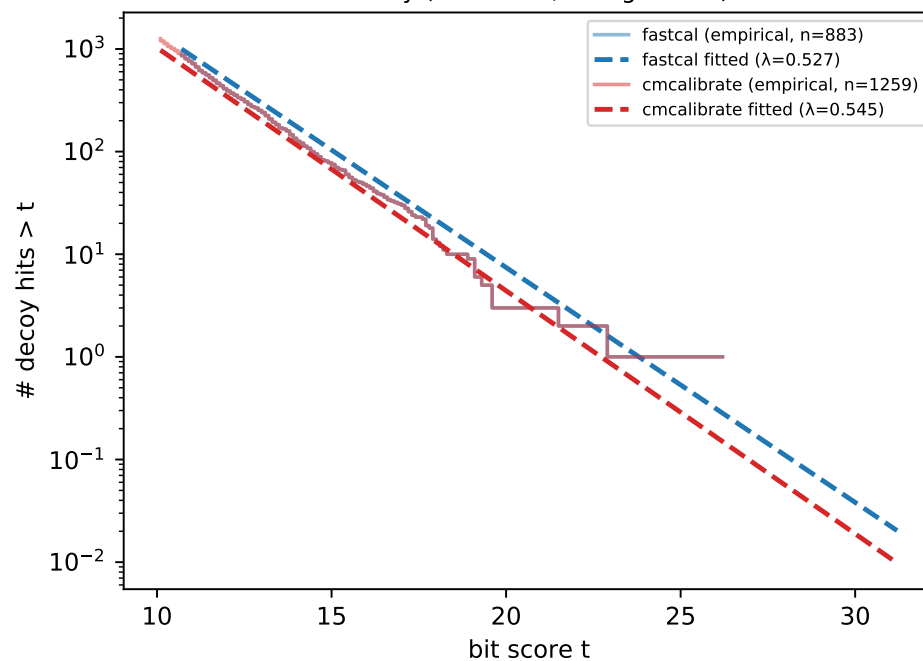


SNORA30 — calibration curve across 4 decoys (2 new real-genome + 2 original)

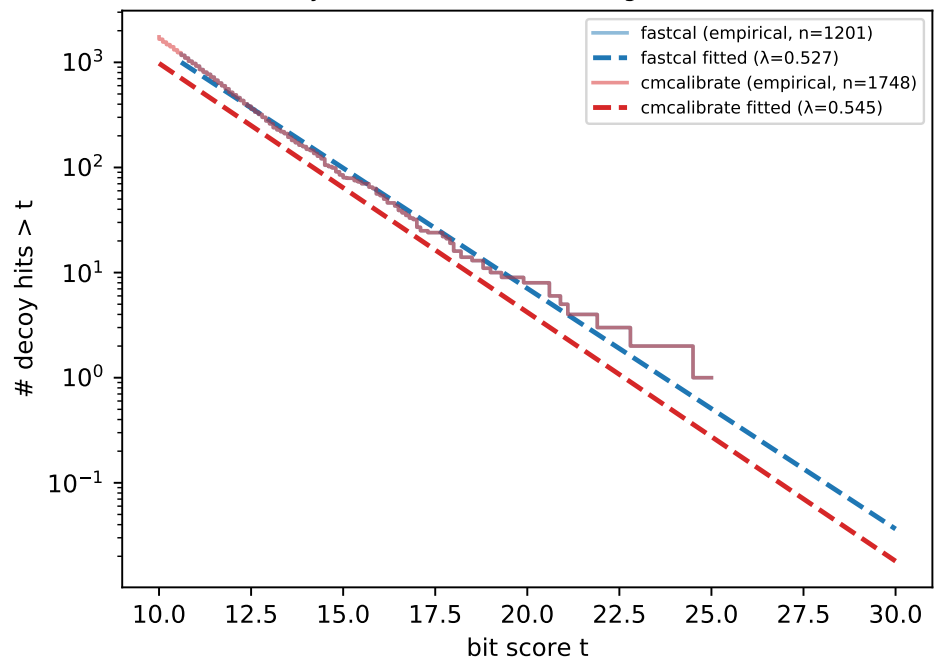
chondrus (~53% GC, real genome)



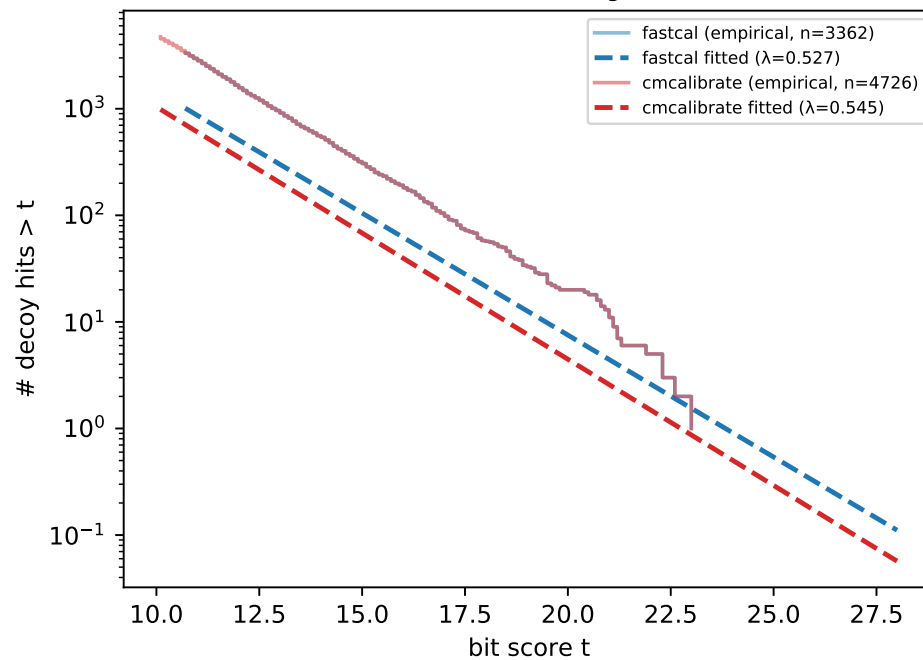
chlamy (~64% GC, real genome)



synth (~35-40% GC, HMM-generated)

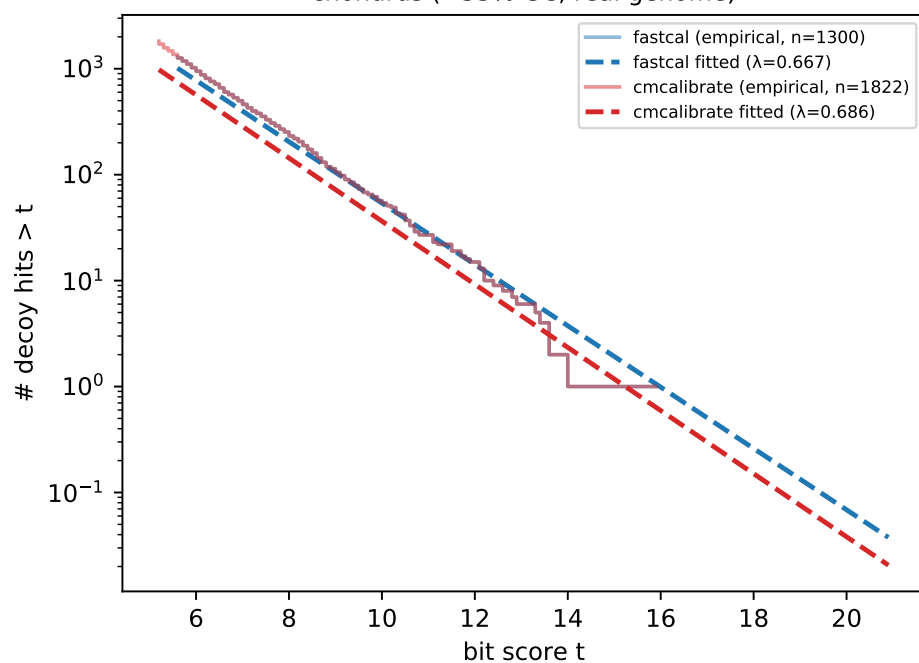


duck (~40% GC, real genome)

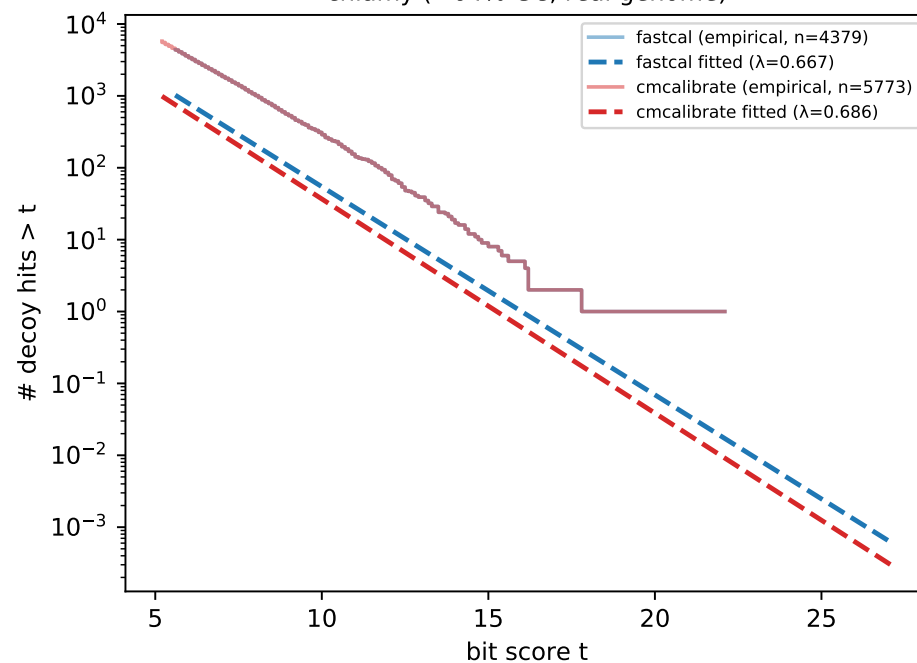


IRES_c-myc — calibration curve across 4 decoys (2 new real-genome + 2 original)

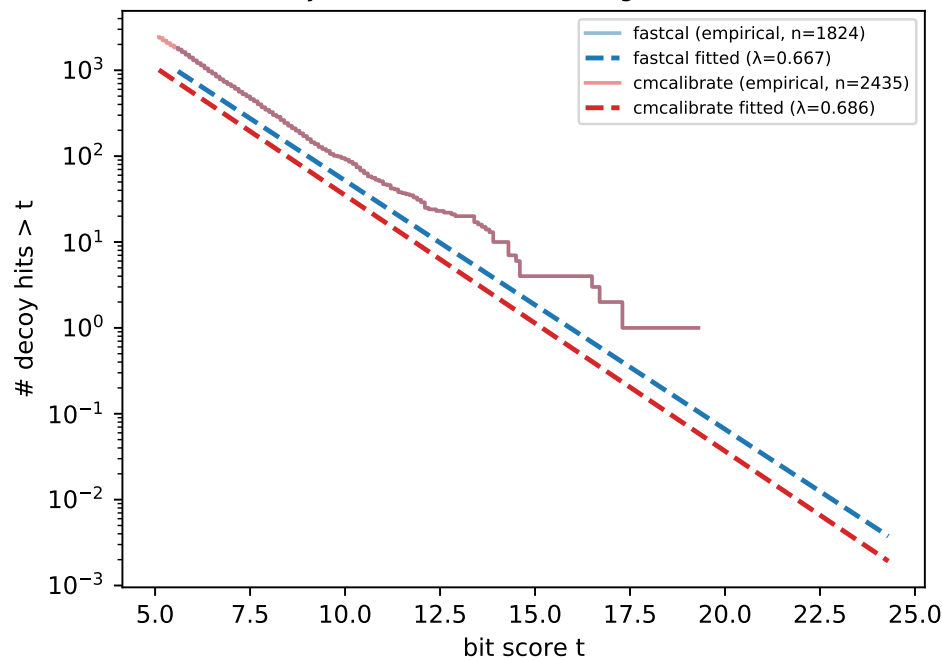
chondrus (~53% GC, real genome)



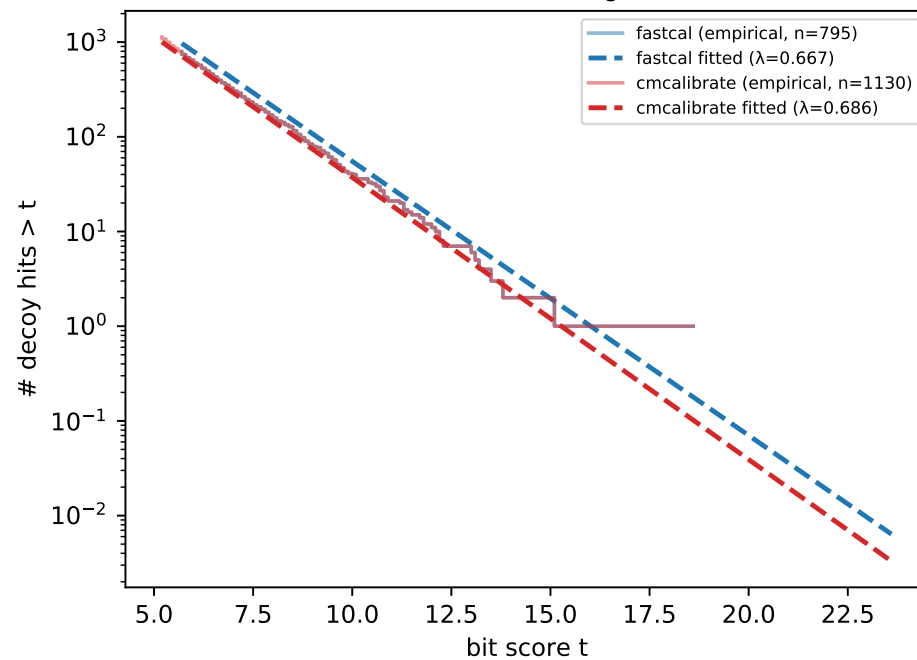
chlamy (~64% GC, real genome)



synth (~35-40% GC, HMM-generated)

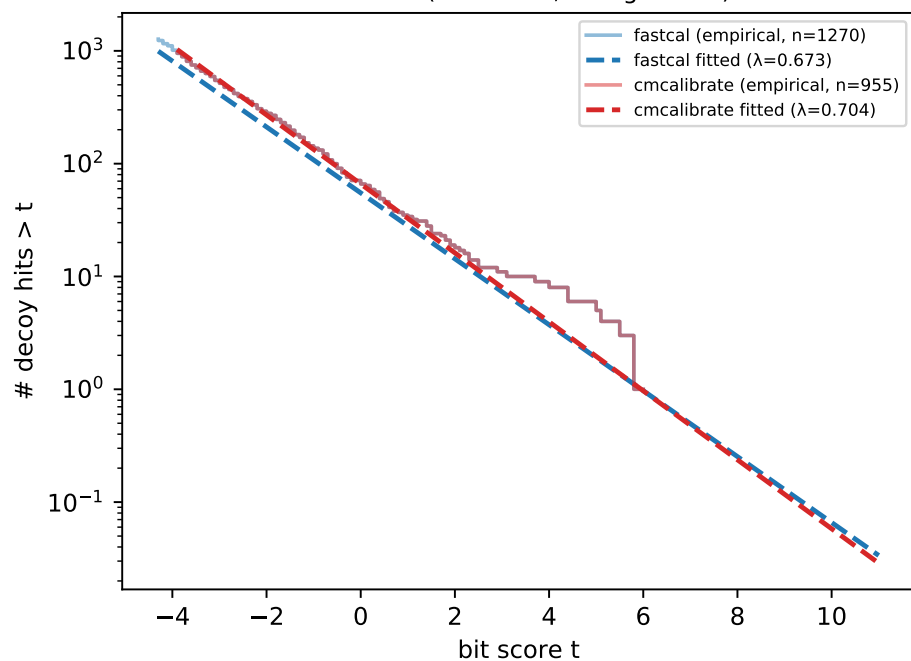


duck (~40% GC, real genome)

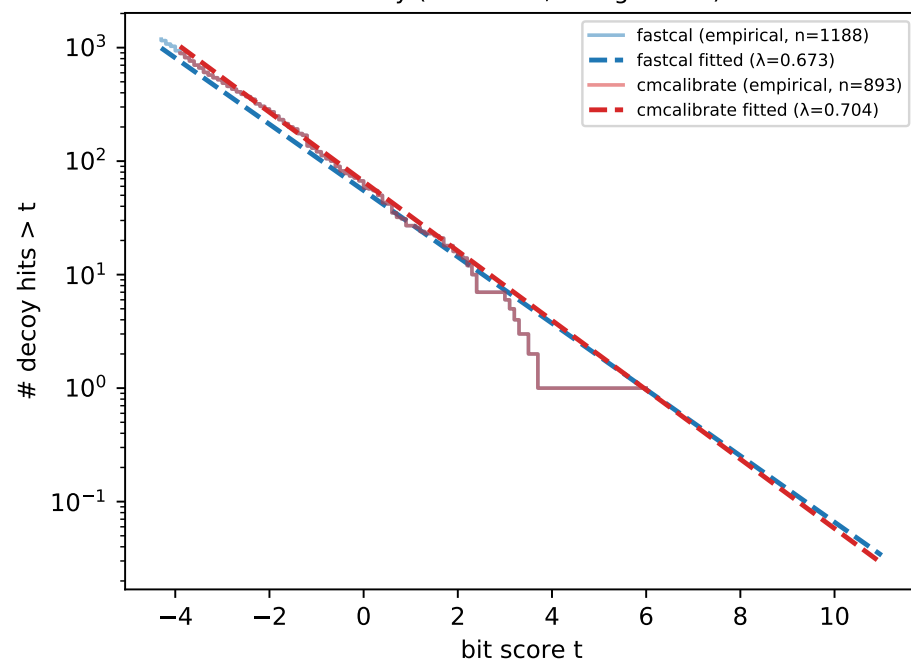


SSU_rRNA_eukarya — calibration curve across 4 decoys (2 new real-genome + 2 original)

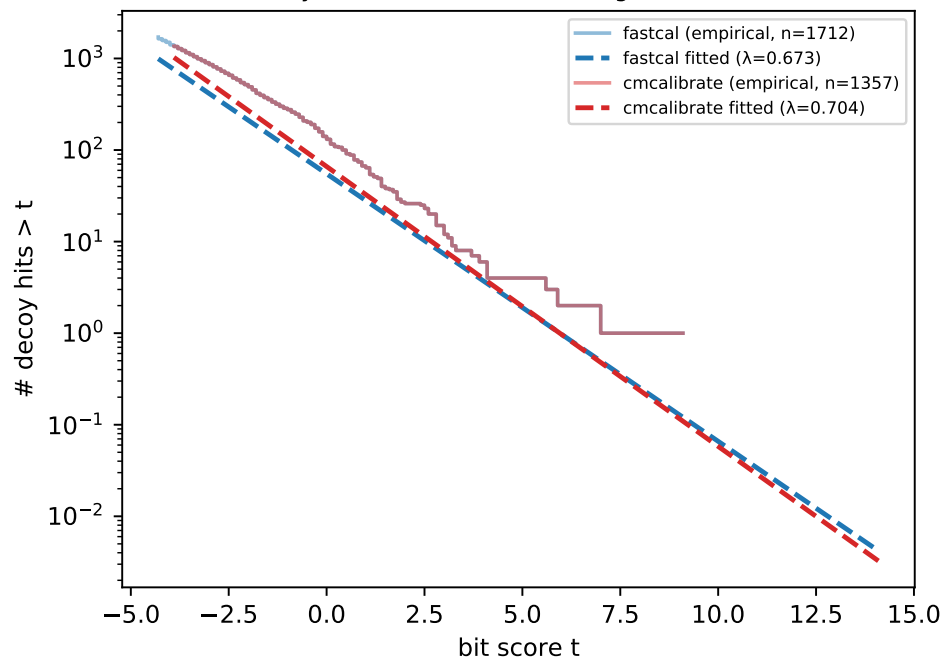
chondrus (~53% GC, real genome)



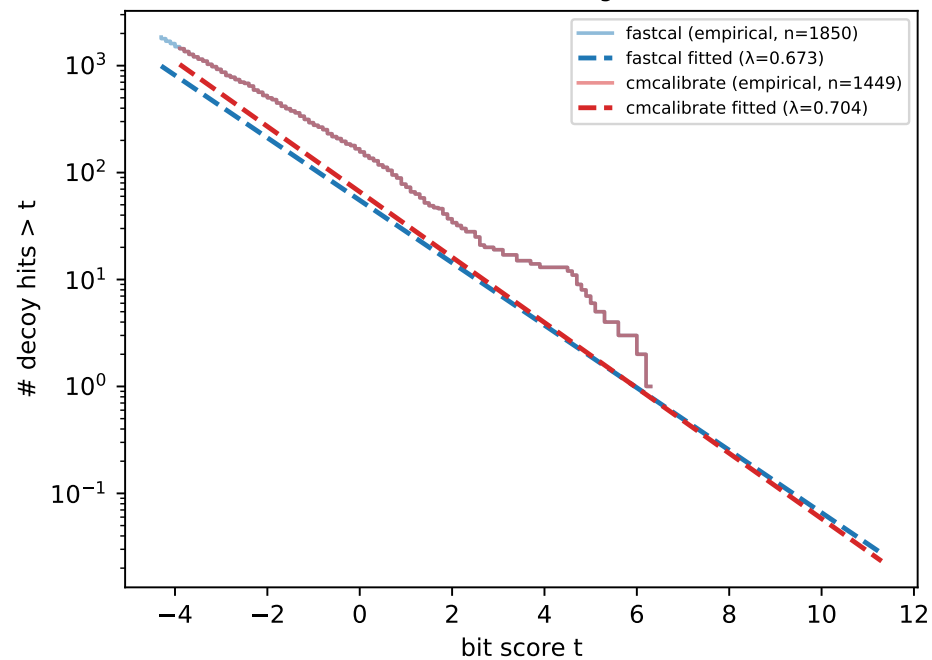
chlamy (~64% GC, real genome)



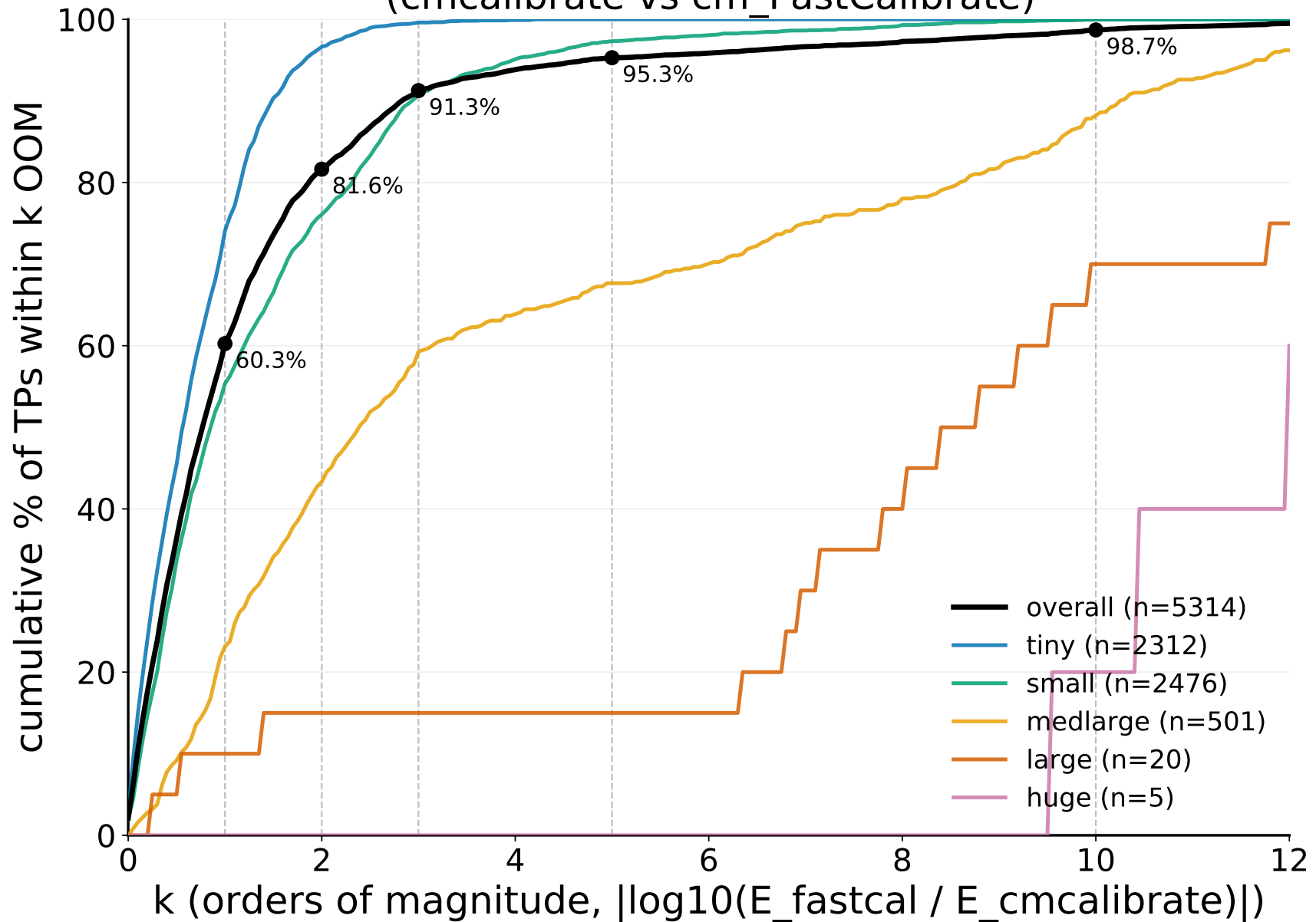
synth (~35-40% GC, HMM-generated)



duck (~40% GC, real genome)

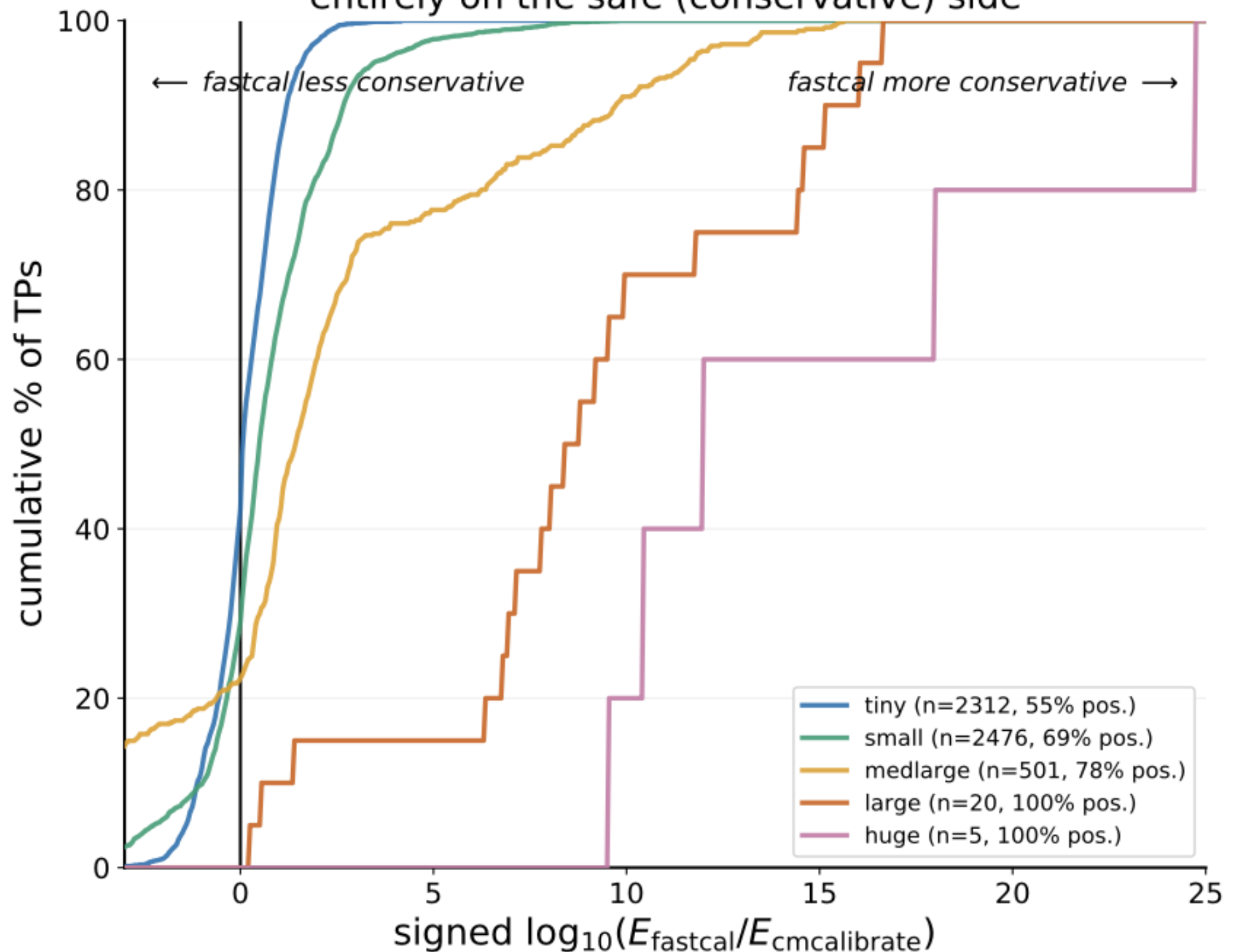


rmark4h TP E-value agreement: survival by order-of-magnitude
(cmcalibrate vs cm_FastCalibrate)

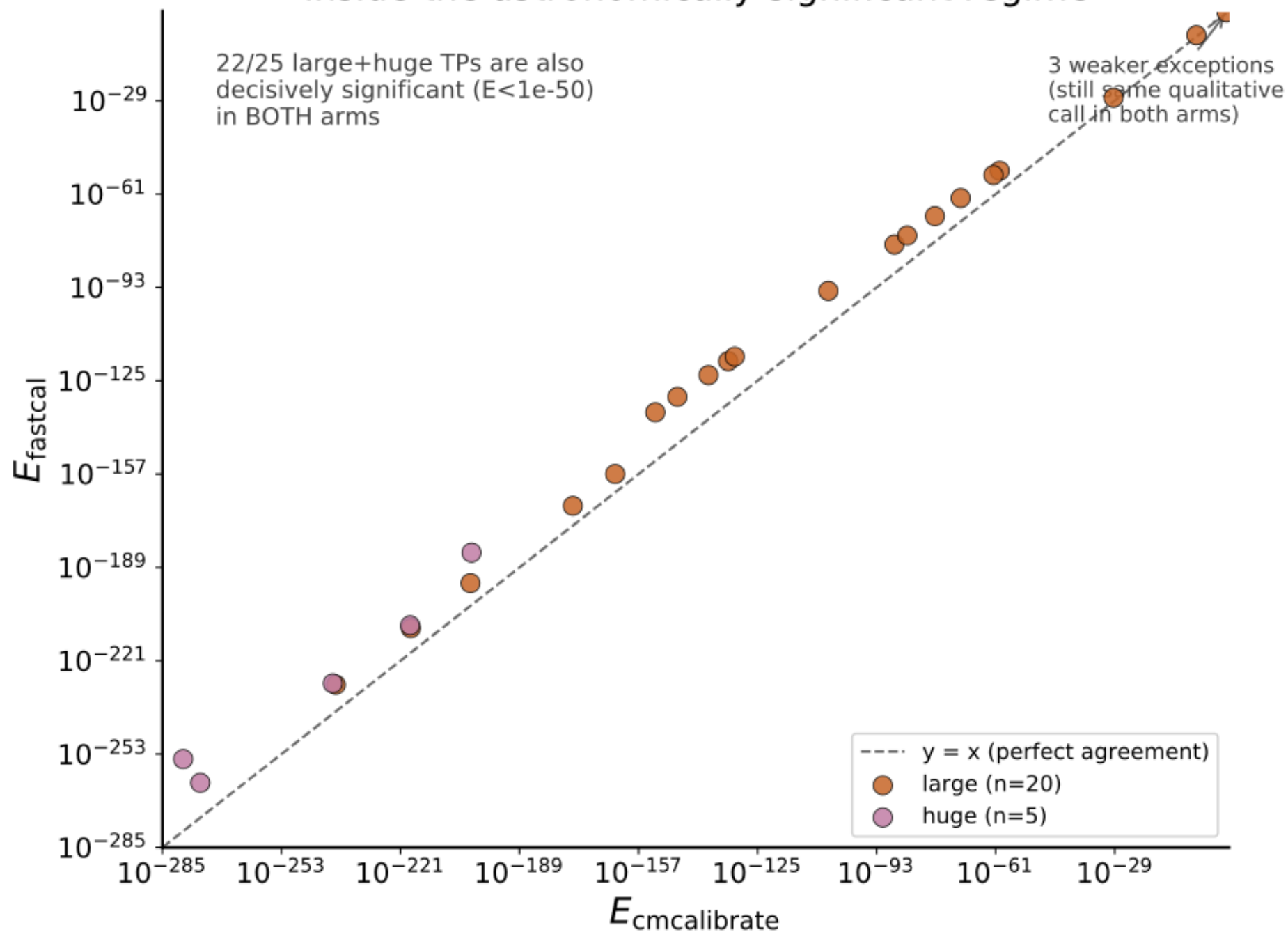


n=5314 matched true positives, rmark4h — Pearson $r=0.991$ on $\log_{10}(E)$;
91.3% agree within 3 orders of magnitude

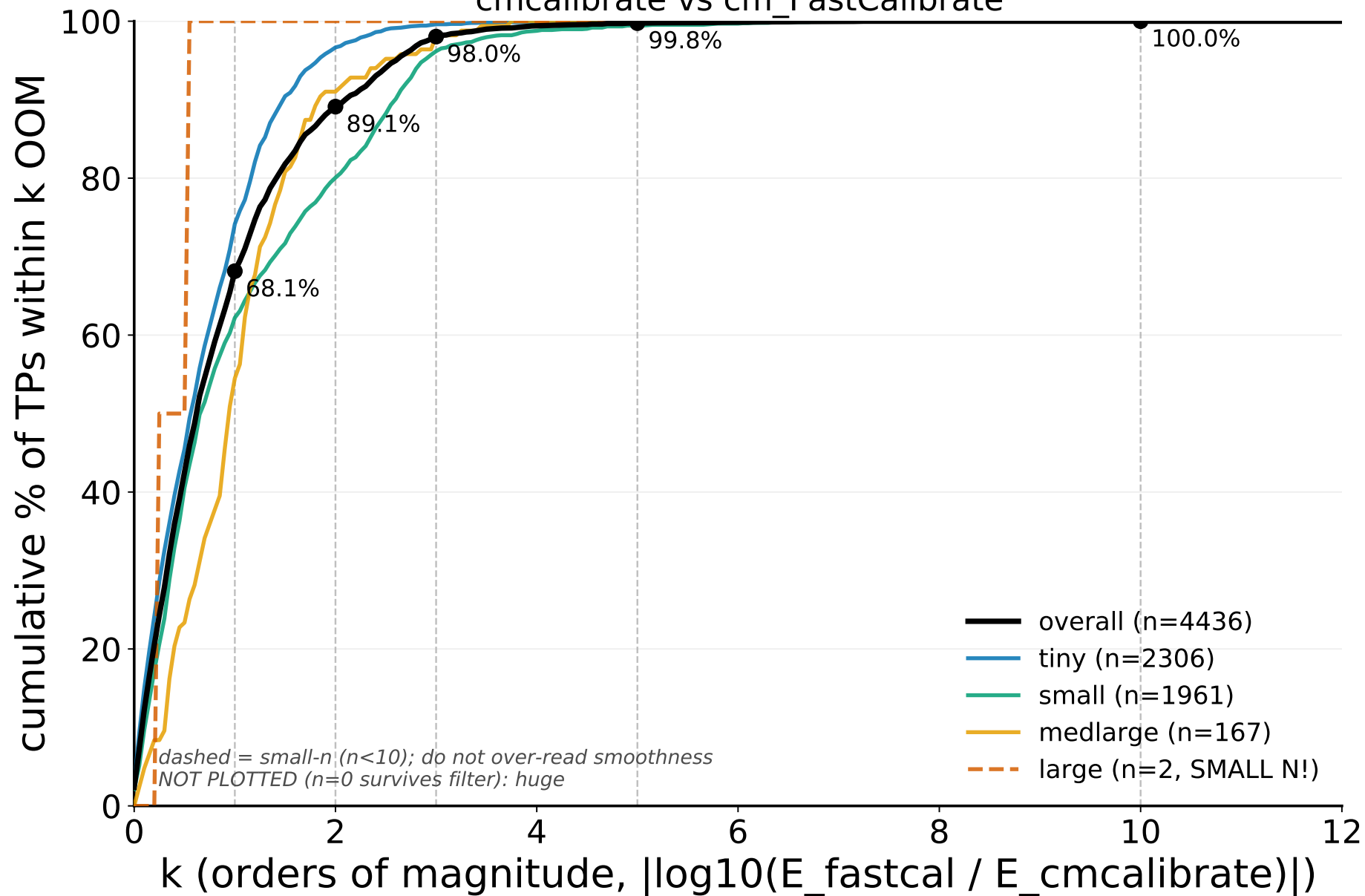
(A) Direction is visible: large/huge sit entirely on the safe (conservative) side

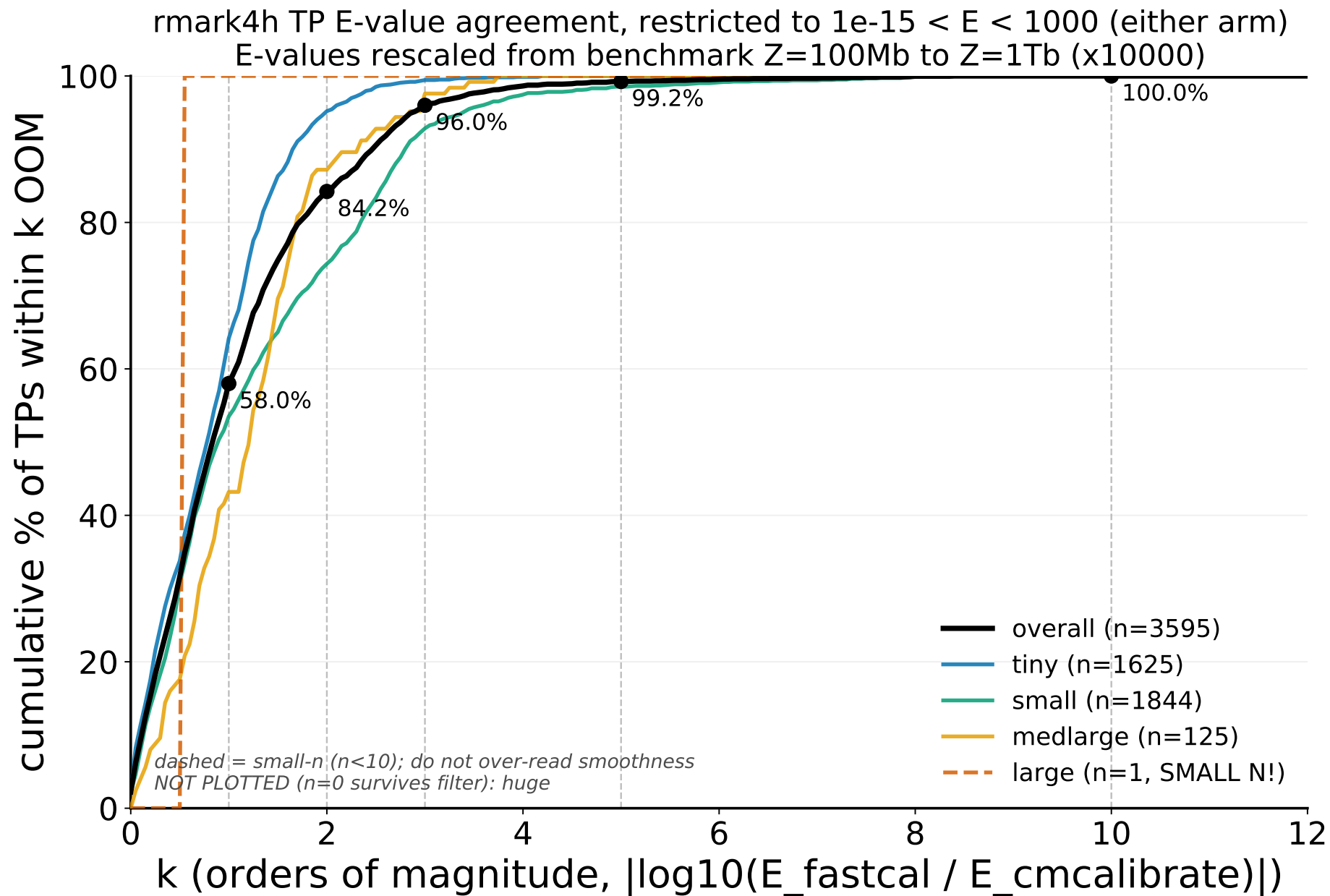


(B) The large/huge OOM gap lives entirely inside the astronomically-significant regime

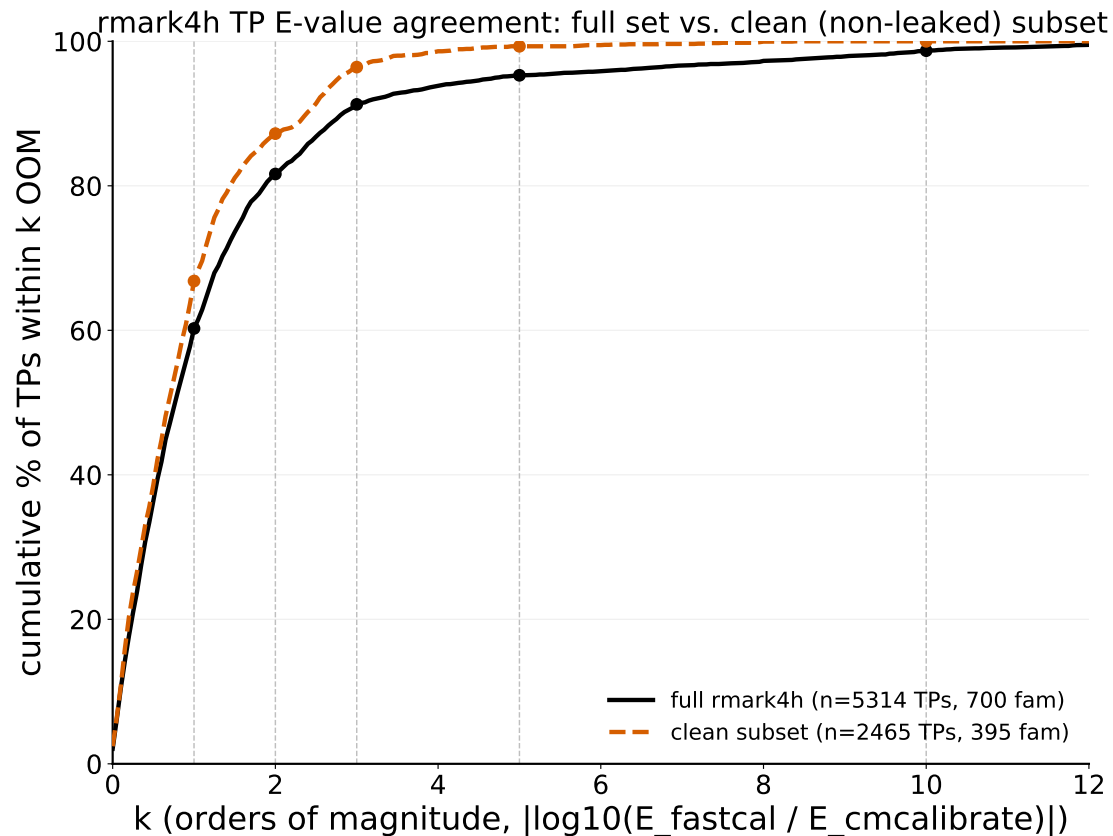


rmark4h TP E-value agreement, restricted to $1e-15 < E < 1000$ (either arm)
cmcalibrate vs cm_FastCalibrate





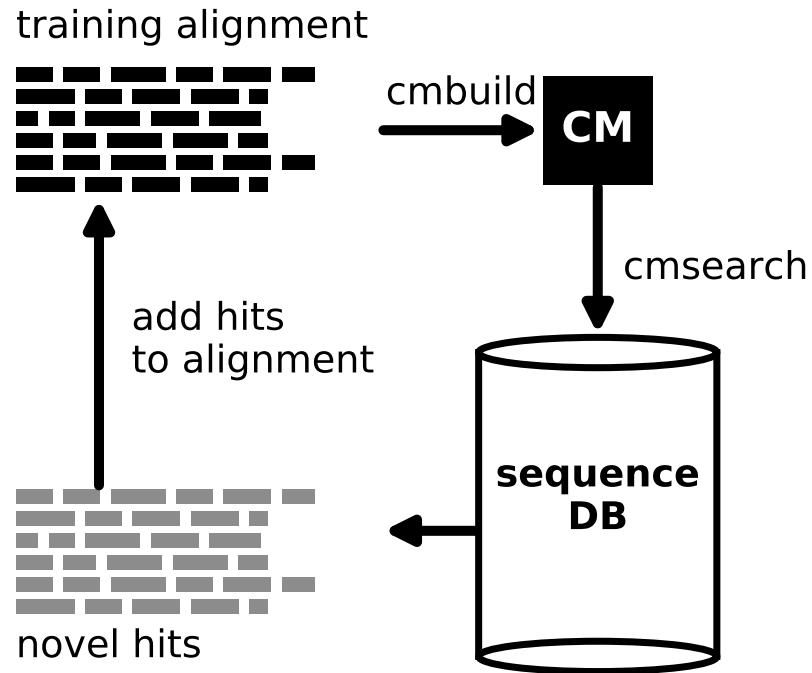
Does this hold up on non-training families?



395/700 rmark4h families with zero training overlap (2465/5314 TPs) - agreement equal-or-better than the full set: ≤ 1 OOM 60.3% $\rightarrow$ 66.8%, ≤ 3 OOM 91.3% $\rightarrow$ 96.4%, ≤ 5 OOM 95.3% $\rightarrow$ 99.3% (full $\rightarrow$ clean)

*caveat: this clean subset has **zero** large/huge-bucket coverage (every real large/huge Rfam family was already in training) - it validates tiny/small, not the biggest models; the LOOCV-over-synthetic-CMs result covers large/huge separately*

cmcalibrate eliminated!?



I'm currently trying to decide if `cmcalibrate` should be left in Infernal as a fallback for very highly specialized use cases like AT-rich models (and for Claude-skeptics).

Pseudoknot alignment markup

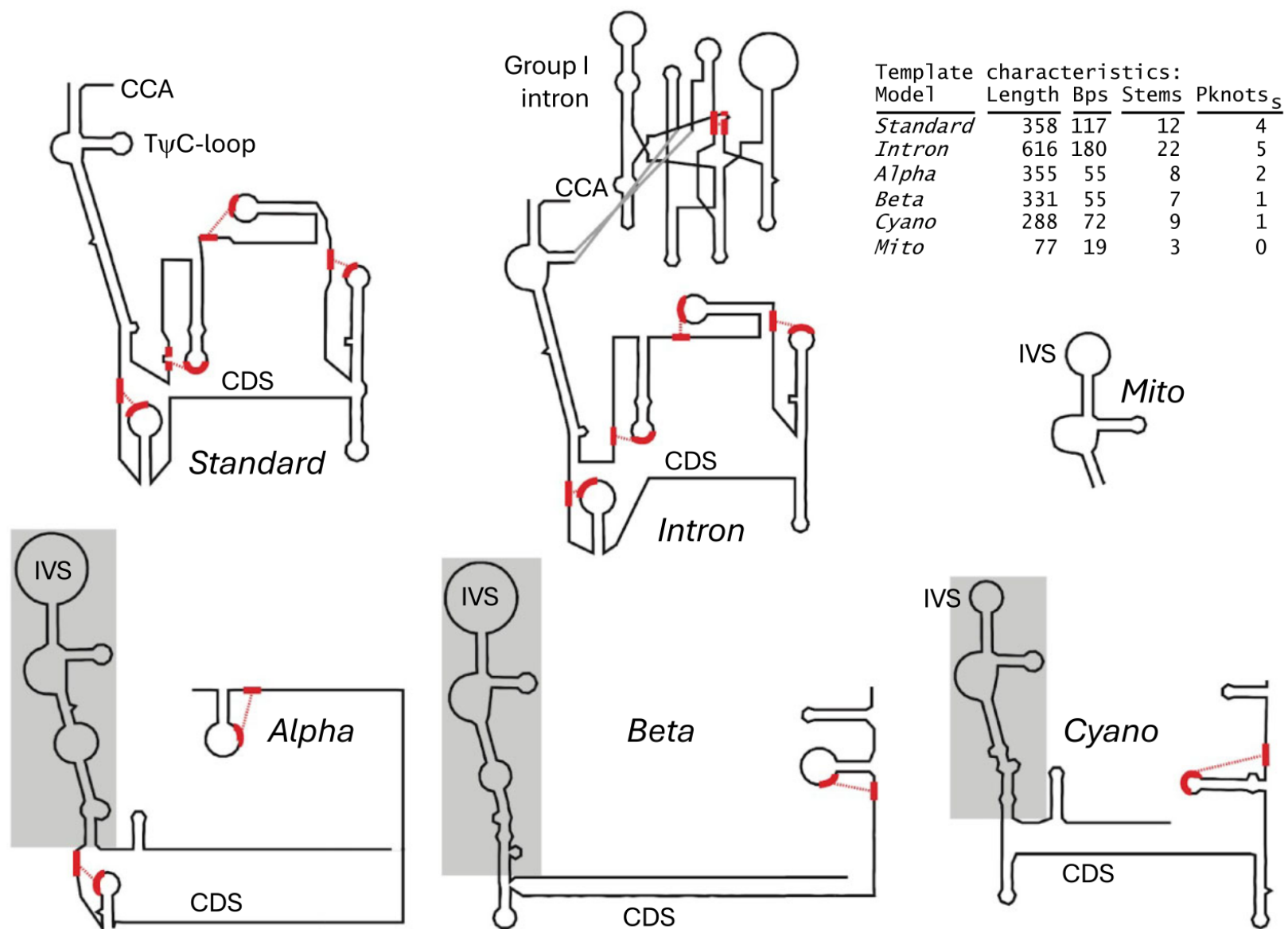


Figure 3. The six tmRNA secondary structure templates, generated by R2DT in “thumbnail” format using the consensus sequence for each template’s CM. Pseudoknot stems not modeled by the CM are highlighted, with each half connected by dashed lines. The boxed regions in the alpha, beta, and cyano templates are the tRNA domains, which include all base pairs modeled by the new, more general tmRNA_permuted Rfam model. Counts of base pairs (Bps) and stems in the table include pseudoknots.

AP009178.1 Nitratiruptor sp. SB155-2 genomic DNA, complete genome														
rank	E-value	score	bias	mdl	mdl from	mdl to	seq from	seq to	acc	trunc	gc			
(1)	!	1.1e-65	225.6	0.0	cm	1	358	[]	1	375	+	[]	0.94	no 0.54
tmRNA.i1ra	1	1	1	1	1	1	1	1	1	1	1	1	1	1
AP009178.1	1	1	1	1	1	1	1	1	1	1	1	1	1	1
tmRNA.i1ra	80	75	141	154	213	289	358	375	305	375	305	375	305	375
AP009178.1	75	154	229	299	368	437	506	575	644	713	782	851	920	989

[illegible]

[illegible]

cmalign Stockholm output: pseudoknot A/a (well-conserved)

AP009178.1:366342-366716/1-375		.-UGUCACAGGCUGCGUGCCGGCU.-UGAGGA-.-ACGCCGUA AAAAU-UCCUCA...AA
#=GR AP009178.1:366342-366716/1-375	PP	..*****999..9****9...9*****65.9*****..**
#=GR AP009178.1:366342-366716/1-375	PS	.v:::__:_:__:__:__==\$\$x.-__:_:_-._x\$\$==____+_-:__:_.
CP001085.1:311523-311892/370-1		.AGGUUGAAGUAGCAUGUCGAGC.-AGA----.-UUCUCGUUAAA--A---UCU...-A
#=GR CP001085.1:311523-311892/370-1	PP	.88888889*****.555.....8*****.4...555....*
#=GR CP001085.1:311523-311892/370-1	PS	.v:::vv_:v____:===x.-_:----.-_x===+_____-_-_:_.
CP002116.1:2157433-2157788/356-1		aCACUCACCGGUUGCAAGCGGAGG.-UGCCCUU.CUCCUCCUUAUCC-AGGGCA..aAC
#=GR CP002116.1:2157433-2157788/356-1	PP	989*****.*****.*****77.*****.***
#=GR CP002116.1:2157433-2157788/356-1	PS	.v:::__:v:__:__:__==\$==.-____:_.==\$+_+__-:__:_.
CP001744.1:2441287-2441647/1-361		aUACGGAAUGGUGGCGUGUCGUGGuUGAUCAGAcGGCCACGUAAAAAU-CUGAUC...AA
#=GR CP001744.1:2441287-2441647/1-361	PP	6*****877777779*****66.777777...8*
#=GR CP001744.1:2441287-2441647/1-361	PS	.v:v:__:_:v__:__:==\$==.:::_.+_.+==\$==____+_-:__:_.
FR891194.1:2245-2609/1-365		.-AAGCAGGGUAAGCGGCAGAGC.-CGGGACA.-AUCUCUUUAAACU-GUUCCAaccUA
#=GR FR891194.1:2245-2609/1-365	PP	..*****.8888887..*****76.8888889****
#=GR FR891194.1:2245-2609/1-365	PS	.v:::_v:~::~:__:__\$==x.-:~:::~+.-_x==\$+_____-:~:::~.
#=GC SS_cons		.(((((((((((((-((((AAAAA.,<<<<<_.aaaa____>>>>>...,,
#=GC RF		.auuccaaaagcuGCAuGcCGAGg.uugccggu.gacCUCGUaAAAcaccggca...AA
#=GC bp_cons		.066**8*68*82.*****4..8*****.4*****.*****8.....
#=GC bp_cov		.735695**5534.0690690956..95*7*~.....65909.....*7*59.....

cmalign Stockholm output: pseudoknot A/a (well-conserved)

```

AP009178.1:366342-366716/1-375      .-UGUCACAGGCUGCGUGCCGGCU.-UGAGGA-.-ACGCCGUAAAAAU-UCCUCA...AA
#=GR AP009178.1:366342-366716/1-375  PP ..*****999..9****9...9*****65.9*****...
#=GR AP009178.1:366342-366716/1-375  PS .v:::__:__:__:$x.-_:_-._x$==+_-:__:___.
CP001085.1:311523-311892/370-1      .AGGUUUGAAGUAGCAUGUCGAGC.-AGA---.-UUCUCGUUAAA--A--UCU...-A
#=GR CP001085.1:311523-311892/370-1  PP .888888889*****.555.....8*****.4...555....*
#=GR CP001085.1:311523-311892/370-1  PS .v:::::vv_:v____:===x.-_:_-._x===+_---_:___.
CP002116.1:2157433-2157788/356-1      aCACUCACCGGUUGCAAGCGGAGG.-UGCCCUU.CUCCUCCUUAUCC-AGGGCA..aAC
#=GR CP002116.1:2157433-2157788/356-1 PP 989*****.*****77.*****.***
#=GR CP002116.1:2157433-2157788/356-1 PS .v:::__:v:__:__:$===.-_:_-._$===+_+_-:__:___.
CP001744.1:2441287-2441647/1-361      aUACGGAAUGGUGGCGUGUCGUGGuUGAUCAGAcGGCCACGUAAAAAU-CUGAUC...AA
#=GR CP001744.1:2441287-2441647/1-361 PP 6*****877777779*****66.777777...8*
#=GR CP001744.1:2441287-2441647/1-361 PS .v:v:__:__:v_:__:==$=.-_:_-._+=$=+_-:__:___.
FR891194.1:2245-2609/1-365          .-AAGCAGGGUAAGCGGGCAGAGC.-CGGGACA.-AUCUCUUUAAACU-GUUCCAaccUA
#=GR FR891194.1:2245-2609/1-365      PP ..*****.8888887..*****76.8888889****
#=GR FR891194.1:2245-2609/1-365      PS .v:::_v:::::__:__$==x.-_:_-._+=$+_---_:__:___.
#=GC SS_cons                          .(((((((((((((-(((((AAAAA.,<<<<<_..aaaaa_____>>>>>...,,
#=GC RF                               .auuccaaaagcuGCAuGcCGAGg.uugccggu.gacCUCGUaAAACaaccggca...AA
#=GC bp_cons                          .066**8*68*82.*****4..8*****.4*****8.....
#=GC bp_cov                           .735695**5534.0690690956..95*7*...65909.....*7*59.....

```

-
- o 'bp_cons' (base-pair conservation): per consensus base-pair column, how often that pair is actually formed as a canonical Watson-Crick/G-U pair across the alignment's sequences (0-9,*);
 - o 'bp_cov' (base-pair covariation): a separate, complementary measure, whether substitutions on the two sides of a pair are statistically coupled (compensatory), i.e. mutual information between the paired columns (0-9,*);

In this alignment 'bp_cons' is * across nearly the whole A/a pair (this pair is structurally almost always formed as WC/GU), and 'bp_cov' ranges 0-9, showing some of the pairs covary across the 5 sequences.

Infernal 1.2 development is ongoing

- We always take feature requests, but now is an especially good time
- Longstanding goal: continue to make Infernal more efficient
 - target of using Infernal for 10-35Kb viral genomes to drive development

Acknowledgements

- David Landsman
- Sean Eddy
- Tom Jones

Thanks, Eric and Elena!



This research was supported by the Intramural Research Program of the NIH, National Library of Medicine