

Responder Profile

vemurafenib

Single-indication depth · basis: DepMap measured · confidence Elevated

Indication — Skin

01 What this report is, and is not

Status — verbatim

Hypothesis-generating, not a validated prediction. Per-drug node ranking is weak (LDO per-drug $r \sim 0.072$); node-mean accuracy is aggregate (~indication-level, pooled $r \sim 0.39$). Biomarkers are candidates to test, not validated markers. Cell-line (DepMap/PRISM) triage, not a clinical or efficacy prediction.

Profile (skin): vemurafenib is predicted at +41.7% viability change in skin (rank 1 of 17, +64.6% vs its own average) — clears the noise floor (2.6x, clear); indication state: true responder. Within skin, 40 of 67 skin-relevant node subgroups are true responders (better than this drug's own average and clearing the noise floor); 4 more are relatively better but below threshold. 30 candidate biomarker(s) are surfaced as leads to test, not validated markers.

Disclaimer — verbatim

Hypothesis-generating, not a validated prediction. Indication-scoped subgroups are computed on the node basis, whose per-drug ranking is weak (committed LDO per-drug node $r \sim 0.072$); node-mean accuracy is aggregate (pooled $r \sim 0.39$, indication-level). Relevant nodes merely contain the target indication's cells -- the predicted response is the node-mean over all of a node's (mixed-tissue) members, so a low indication purity means the score is not indication-pure. Candidate biomarkers and the full-panel relevance map are leads to test, not validated markers: committed imputed biomarker transfer is weak (top-k Jaccard ~ 0). Held-out AUC is deliberately not surfaced at this resolution. Cell-line (DepMap/PRISM) triage, not a clinical or efficacy prediction. Validate experimentally before relying on any item.

What counts as a responder

Definition — verbatim

A true responder is a unit that (relative) responds better than this drug's own average, and (absolute) clears the committed 0.30 log2 noise floor. The floor is a noise gate ('distinguishable from no effect'), not an efficacy bar. Units that meet the relative criterion but fail the floor are reported in full and labelled 'does not clear the response threshold — not a true responder'.

What the threshold is. The response threshold is a noise floor at 0.30 log2 — the point at which a predicted response becomes distinguishable from no effect at all. It is not an efficacy bar and not a claim about clinical usefulness.

Clearing it. Clearing the floor means the predicted response is real rather than noise. It does not mean the response is large, and it does not mean the compound will work.

Reading the magnitude. The absolute magnitude above the floor is the promise signal — that is the number to judge. A potent compound clears the floor easily and with high magnitude. A weak-but-real compound (a typical OTC drug or dietary

supplement) may sit just above the floor: real, potentially useful, and entirely your judgment to act on. Both are reported the same way, with the magnitude shown, so you can tell them apart.

Failing it. Failing the floor means the apparent response is within noise — not a true responder — even when it is relatively better than the drug's own average. A weak drug still has a 'best' subgroup; that is arithmetic, not activity. These are labelled 'does not clear the response threshold' and their full analysis is still reported, so you can see exactly what was and was not found.

Why "its own average". The relative criterion compares each unit against this drug's own average, not against other drugs. That is deliberately toxicity-adaptive: a mild supplement and a potent chemotherapeutic are each judged against their own baseline, so a modest compound is not dismissed for being modest and a broadly cytotoxic one earns no credit for killing everything uniformly.

02 Skin — indication level

+41.7%	+64.6%	2.60×	1/17	-0.0600
Absolute viability change	Relative effect	Multiples of noise floor	Rank among indications	Compound's own average (log ₂)

True responder. Responds better than this drug's own average and clears the noise floor — a true responder. Read the absolute magnitude below to judge how strong the effect is: clearing the floor means the response is real, not that it is large.

Clear

clears the noise floor by 2-3x — a clearly real predicted effect

Basis for this indication: Measured · predicted log₂ -0.7792 · margin over the compound's own average 0.7192 log₂ · pan-cancer cell lines in responder zones 24/38

Context note — verbatim

Indication-level context from the committed 17-indication model. The absolute viability change says whether the compound acts on this tissue at all; the relative effect only compares it with the compound's own average and cannot establish activity. The subgroup analysis below is computed on the node basis and can surface a responsive subgroup even where this indication average is unpromising — but it does not overturn a weak absolute reading.

The compound's full 17-indication context

All 17 modelled indications, ranked, including those with no predicted activity. The table continues over the page if it does not fit.

#	Indication	Absolute	Relative	log ₂	Basis
1	Skin this report	+41.7%	+64.6%	-0.7792	Measured
2	Kidney	+18.0%	+17.0%	-0.2869	Measured
3	Ovary	+13.4%	+10.8%	-0.2078	Measured
4	Uterus	+10.0%	+6.6%	-0.1520	Measured
5	Thyroid	+5.7%	+1.7%	-0.0848	Measured
6	Liver	+4.3%	+0.2%	-0.0634	Measured
7	Lung	+2.0%	-2.1%	-0.0291	Measured
8	Central Nervous System	+0.2%	-3.8%	-0.0035	Measured
9	Rhabdoid	-0.3%	-4.3%	0.0038	Measured
10	Bone	-0.5%	-4.6%	0.0079	Measured

#	Indication	Absolute	Relative	log ₂	Basis
11	Pancreas	-1.1%	-5.1%	0.0157	Measured
12	Esophagus	-1.2%	-5.2%	0.0178	Measured
13	Urinary Tract	-4.8%	-8.4%	0.0671	Measured
14	Upper Aerodigestive	-6.5%	-9.9%	0.0906	Measured
15	Breast	-7.0%	-10.4%	0.0980	Measured
16	Gastric	-10.0%	-12.8%	0.1374	Measured
17	Colorectal	-10.8%	-13.4%	0.1483	Measured

† Imputed from chemical structure, not measured. Imputed and measured values are not on the same scale. · Basis reasons, abbreviated: not in DepMap = no measured data for this compound; too few lines = too few measured cell lines in this tissue; structure unmatched = structure could not be matched to DepMap. · Section confidence — elevated (chemistry similarity); committed per-indication $r \sim 0.39$ (indication-level).

03 Responder subgroups within Skin

Node basis

Method. indication-scoped node subgroups · **scope** indication-scoped (skin) · **mode** responder subgroups (node)

491-node vector (chemistry->node-mean regressor; varies within indication)

Relevant nodes: 67 (those containing at least 3 Skin cell line(s)). Selection rule: two-criterion responder definition; compound's own node average -0.2800 log₂.

A true responder is a unit that (relative) responds better than this drug's own average, and (absolute) clears the committed 0.30 log₂ noise floor. The floor is a noise gate ('distinguishable from no effect'), not an efficacy bar. Units that meet the relative criterion but fail the floor are reported in full and labelled 'does not clear the response threshold — not a true responder'.

Purity caveat — verbatim, read before the table below

Relevant nodes are those containing at least 3 skin cell lines; predicted response is the node-mean over all of a node's cells, which are mixed-tissue, so a low indication purity means the score is not indication-pure.

40	4	23	67	368
True-responder nodes	Below threshold	No differential	Relevant nodes	Cell lines characterized

Top 15 nodes by predicted response

Node	Cells	Skin cells	Indication purity	Absolute	log ₂	× floor	Tier	Responder state
Subgroup 9228015a	10	5	50.0%	+43.9%	-0.8343	2.78×	Clear	True responder
Subgroup f2fb262b	6	3	50.0%	+40.3%	-0.7437	2.48×	Clear	True responder
Subgroup 27ea3821	8	5	62.5%	+36.1%	-0.6462	2.15×	Clear	True responder
Subgroup dfd9f943	17	4	23.5%	+34.9%	-0.6191	2.06×	Clear	True responder
Subgroup cc15b8d1	25	8	32.0%	+33.0%	-0.5776	1.93×	Modest	True responder
Subgroup 3376c2b6	12	6	50.0%	+31.9%	-0.5543	1.85×	Modest	True responder
Subgroup 2d3d76ec	8	3	37.5%	+31.6%	-0.5471	1.82×	Modest	True responder
Subgroup 073f6127	12	4	33.3%	+30.1%	-0.5161	1.72×	Modest	True responder
Subgroup 10d5aa6e	19	4	21.1%	+29.3%	-0.4994	1.67×	Modest	True responder
Subgroup bb99f3d2	15	5	33.3%	+28.9%	-0.4926	1.64×	Modest	True responder
Subgroup 27cc7dd4	18	4	22.2%	+28.4%	-0.4814	1.60×	Modest	True responder
Subgroup afbcef57	55	14	25.5%	+27.6%	-0.4661	1.55×	Modest	True responder
Subgroup 6817e159	48	12	25.0%	+27.3%	-0.4590	1.53×	Modest	True responder
Subgroup 8e28944e	93	16	17.2%	+26.3%	-0.4409	1.47×	Modest	True responder
Subgroup fe8177ae	10	3	30.0%	+26.0%	-0.4347	1.45×	Modest	True responder

Indication purity is the purity signal: the fraction of the node's cell lines that are Skin. A low fraction means the node's predicted response is driven substantially by other tissues — read it together with the purity caveat above. This table shows the 15 strongest of 67 relevant nodes; the remainder are available on request.

Within-indication spread across the 67 relevant nodes: -0.8343 to 0.0649 log₂.

Responder counts at each committed threshold

Threshold (log ₂)	Responder nodes	Proportion
-0.30	40	59.7%
-0.50	8	11.9%
-1.00	0	0.0%

Section confidence — Node basis — aggregate accuracy matches the indication basis (pooled r ~0.386) but per-drug node ranking is weak (committed LDO r ~0.072). A triage substrate for within-indication structure, not a validated within-indication predictor. Held-out AUC is deliberately not surfaced.

04 Candidate biomarkers

Label — verbatim

Candidates to test — not validated biomarkers

Characterized on nodes better than the drug's own average (44 of 67); 40 of them also clear the noise floor. State **True responder** · 368 responder cell lines vs 100 background (indication-matched background). Content is reported in every responder state — a below-threshold result is labelled, not withheld.

Expression — candidates to test

#	Feature	Direction	RF importance	CV mean $\pm$ sd	Cohen's d	Nominal p	In CART
1	DHDDS	Higher expression → stronger response	0.0055	0.0054 $\pm$ 0.0028	0.477	1.72e-05	—
2	ARHGAP1	Higher expression → stronger response	0.0049	0.0023 $\pm$ 0.0014	0.399	5.05e-04	—
3	GRB7	Lower expression → stronger response	0.0046	0.0038 $\pm$ 0.0023	-0.131	1.26e-01	yes
4	TFAP2A	Lower expression → stronger response	0.0040	0.0030 $\pm$ 0.0009	-0.185	2.80e-01	—
5	APP	Higher expression → stronger response	0.0039	0.0025 $\pm$ 0.0011	0.310	0.0077	—
6	GNA15	Lower expression → stronger response	0.0037	0.0031 $\pm$ 0.0016	-0.416	4.74e-05	—
7	MST1R	Lower expression → stronger response	0.0036	0.0011 $\pm$ 0.0007	-0.446	1.14e-04	—
8	PLS1	Lower expression → stronger response	0.0035	0.0040 $\pm$ 0.0009	-0.404	9.07e-04	—
9	EPCAM	Lower expression → stronger response	0.0034	0.0027 $\pm$ 0.0011	-0.419	0.0011	yes
10	ATP6V0B	Higher expression → stronger response	0.0034	0.0024 $\pm$ 0.0014	0.413	5.85e-04	—
11	CDKN2C	Higher expression → stronger response	0.0032	0.0015 $\pm$ 0.0005	0.053	3.65e-01	—
12	TMEM50A	Higher expression → stronger response	0.0032	0.0036 $\pm$ 0.0014	0.247	0.0120	—
13	HERPUD1	Higher expression → stronger response	0.0031	0.0023 $\pm$ 0.0011	0.280	0.0061	—
14	CAST	Lower expression → stronger response	0.0031	0.0031 $\pm$ 0.0013	-0.365	0.0026	—
15	PIK3C3	Higher expression → stronger response	0.0030	0.0029 $\pm$ 0.0015	0.403	1.14e-04	—

Mutation — candidates to test

#	Feature	Direction	RF importance	CV mean $\pm$ sd	Risk diff	Nominal p	In CART
1	KRAS	Mutation absent → stronger response	0.0247	0.0252 $\pm$ 0.0063	-0.149	0.0019	yes
2	BRAF	Mutation present → stronger response	0.0193	0.0178 $\pm$ 0.0033	0.108	0.0060	—
3	TP53	Mutation absent → stronger response	0.0127	0.0121 $\pm$ 0.0015	-0.062	2.63e-01	—
4	EP300	Mutation absent → stronger response	0.0108	0.0113 $\pm$ 0.0034	-0.089	0.0446	—
5	CREBBP	Mutation absent → stronger response	0.0107	0.0095 $\pm$ 0.0025	-0.047	2.60e-01	—
6	ZFHX3	Mutation absent → stronger response	0.0096	0.0087 $\pm$ 0.0020	-0.036	3.88e-01	yes
7	CDKN2A	Mutation absent → stronger response	0.0092	0.0105 $\pm$ 0.0017	-0.069	1.19e-01	—
8	SPEN	Mutation present → stronger response	0.0090	0.0081 $\pm$ 0.0015	0.031	5.09e-01	—
9	HLA-B	Mutation absent → stronger response	0.0086	0.0079 $\pm$ 0.0015	-0.064	0.0732	—
10	NF2	Mutation absent → stronger response	0.0086	0.0085 $\pm$ 0.0024	-0.028	3.31e-01	—
11	SMARCA4	Mutation present → stronger response	0.0082	0.0089 $\pm$ 0.0009	0.018	7.62e-01	—

#	Feature	Direction	RF importance	CV mean ± sd	Risk diff	Nominal p	In CART
12	BRCA1	Mutation present → stronger response	0.0080	0.0070 ± 0.0008	0.058	0.0711	—
13	ERCC5	Mutation present → stronger response	0.0077	0.0067 ± 0.0016	0.004	1.00e+00	—
14	PTPRT	Mutation absent → stronger response	0.0074	0.0072 ± 0.0020	-0.078	0.0446	—
15	KDM6A	Mutation absent → stronger response	0.0072	0.0074 ± 0.0021	-0.035	2.76e-01	—

p-values are **nominal and uncorrected**. These are leads to test experimentally, not validated markers, and no held-out discrimination statistic is reported — method 3a is AUC-free by design. · Section confidence — exploratory — committed imputed biomarker transfer is weak (top-k Jaccard ~0) and exact responder-cell membership does not transfer. Leads to test.

05 Molecular relevance map — appendix

Label — verbatim

Hypothesis prioritisation (rank-ordered univariate associations to prioritise experimental work) — not validated biomarker prediction

An **appendix-grade univariate scan** over the whole feature space, ranked for prioritisation only. Ranking is by univariate score; the Bonferroni column is reported alongside the nominal p and never reorders or removes a row. A tick marks a feature clearing Bonferroni at 0.05. **The two spaces differ in size and in origin.** Expression — 1,353 of 1,353 genes ranked (0 excluded), the L1000 landmark plus MSK-CHORD cancer-gene universe. Mutation — 458 of 470 genes ranked (12 excluded), the MSK-CHORD clinical panel, deliberately left narrower. Each space is corrected against its own count, so the same nominal p yields a different corrected value in each. The top 15 of each space are shown; the full ranked map is available on request.

Expression — top 15 of 1,353 ranked

Rank	Feature	Direction	Univariate score	Nominal p	Bonferroni p	Percentile	In headline top-k
1	DDR2	Higher expression → stronger response	0.5034	3.19e-04	4.31e-01	100.0	—
2	SFN	Lower expression → stronger response	0.4878	3.59e-05	0.0486 ✓	99.9	—
3	SLC1A4	Higher expression → stronger response	0.4862	6.85e-05	0.0926	99.9	—
4	LAMA3	Lower expression → stronger response	0.4819	2.91e-05	0.0394 ✓	99.8	—
5	DHDDS	Higher expression → stronger response	0.4771	1.72e-05	0.0233 ✓	99.7	yes
6	KCNK1	Lower expression → stronger response	0.4627	1.01e-04	1.37e-01	99.6	—
7	PDLIM1	Lower expression → stronger response	0.4509	0.0026	1.00e+00	99.6	—
8	EGFR	Lower expression → stronger response	0.4485	9.22e-04	1.00e+00	99.5	—
9	MST1R	Lower expression → stronger response	0.4457	1.14e-04	1.55e-01	99.4	yes
10	DSG2	Lower expression → stronger response	0.4411	0.0028	1.00e+00	99.3	—
11	IGF1	Higher expression → stronger response	0.4364	6.09e-05	0.0824	99.3	—
12	SMAD4	Higher expression → stronger response	0.4253	1.77e-04	2.39e-01	99.2	—
13	EPCAM	Lower expression → stronger response	0.4193	0.0011	1.00e+00	99.1	yes
14	GNA15	Lower expression → stronger response	0.4163	4.74e-05	0.0641	99.0	yes
15	CDH3	Lower expression → stronger response	0.4157	4.13e-04	5.59e-01	99.0	—

Mutation — top 15 of 458 ranked

Rank	Feature	Direction	Univariate score	Nominal p	Bonferroni p	Percentile	In headline top-k
1	KRAS	Mutation absent → stronger response	0.1488	0.0019	8.52e-01	100.0	yes

Rank	Feature	Direction	Univariate score	Nominal p	Bonferroni p	Percentile	In headline top-k
2	BRAF	Mutation present → stronger response	0.1085	0.0060	1.00e+00	99.8	yes
3	EP300	Mutation absent → stronger response	0.0887	0.0446	1.00e+00	99.6	yes
4	PTPRT	Mutation absent → stronger response	0.0776	0.0446	1.00e+00	99.3	yes
5	EPHA5	Mutation absent → stronger response	0.0721	0.0257	1.00e+00	99.1	—
6	CDKN2A	Mutation absent → stronger response	0.0687	1.19e-01	1.00e+00	98.9	yes
7	HLA-B	Mutation absent → stronger response	0.0639	0.0732	1.00e+00	98.7	yes
8	TP53	Mutation absent → stronger response	0.0616	2.63e-01	1.00e+00	98.5	yes
9	BRCA1	Mutation present → stronger response	0.0578	0.0711	1.00e+00	98.2	yes
10	FAT1	Mutation absent → stronger response	0.0523	1.93e-01	1.00e+00	98.0	—
11	PPM1D	Mutation present → stronger response	0.0516	0.0184	1.00e+00	97.8	—
12	PTCH1	Mutation present → stronger response	0.0488	1.16e-01	1.00e+00	97.6	—
13	ARID1B	Mutation absent → stronger response	0.0486	2.26e-01	1.00e+00	97.4	—
14	DR0SHA	Mutation present → stronger response	0.0479	0.0884	1.00e+00	97.2	—
15	PGR	Mutation present → stronger response	0.0479	0.0884	1.00e+00	96.9	—

Section confidence — appendix-grade univariate scan over the whole feature space; nominal (uncorrected) p-values are reported with a Bonferroni-corrected value alongside (Mann-Whitney for expression, Fisher for mutation). Each space is corrected against its own test count, which differs between them. Ranking uses the nominal value. A prioritisation aid, not a significance claim.

06 Gene dependency — an independent check

What this section is

CRISPR gene-effect data is an independent measurement. It is not used in response prediction, in zone or subgroup selection, in the two-criterion responder definition, or in any score in this report. It is laid alongside the response hypothesis as corroboration to weigh, and nothing in it feeds back into the analysis above.

Any dependency reported here is a mechanistic hypothesis consistent with the response hypothesis — not an established mechanism, and not a validated target. Chronos scores are loss-of-function fitness effects in cell lines, which is a different experiment from the drug-response data used above.

This compound's responder population, pooled

43	13	0.099	0.90	0
Cell lines with CRISPR data	Dropped — no CRISPR data	Smallest effect resolvable (Chronos)	Same, in standard deviations	Genes past correction

0 significant gene(s) against a matched null of 1000 random cell groups of the same size, same test and same correction, which produced a median of 0.0, a 95th percentile of 1.0 and at most 2. 1000 of 1000 null groups reached 0 or more, giving an empirical p of 1. That does not clear the 0.05 threshold, so this result is indistinguishable from noise and must not be read as a finding.

Matched null — 1000 random groups of 43 cell lines, same test and same correction: median 0.0, 95th percentile 1.0, most 2, 114 of 1000 groups produced at least one. 1000 of 1000 reached the observed count or more, giving an empirical p of 1.0000 against a threshold of 0.05. Seed 42.

Tissue make-up of this population. Skin 16; Ovary 5; Kidney 4; Colorectal 3; Upper Aerodigestive 3; Esophagus 3; Thyroid 2; Urinary Tract 2; Fibroblast 1; Lung 1; Liver 1; Uterus 1; Gastric 1 · Skin is the largest share at 37%.

A pooled responder population is drawn from mixed-tissue subgroups, so any dependency specific to one tissue is diluted by the others and can vanish entirely. The worked case in this data: vemurafenib's 43 pooled responder cell lines return no significant gene at all, while the 28 skin cell lines on their own return 51 — the same biology, visible in one framing and averaged away in the other. A null pooled result is therefore not evidence that no dependency exists.

All Skin cell lines

28	10	0.120	1.09	51
Cell lines with CRISPR data	Dropped — no CRISPR data	Smallest effect resolvable (Chronos)	Same, in standard deviations	Genes past correction

51 significant gene(s) against a matched null of 1000 random cell groups of the same size, same test and same correction, which produced a median of 0.0, a 95th percentile of 1.0 and at most 3. 0 of 1000 null groups reached 51 or more, giving an empirical p of 0.000999. That clears the 0.05 threshold, so the list is not simply what groups of this size produce by chance.

Matched null — 1000 random groups of 28 cell lines, same test and same correction: median 0.0, 95th percentile 1.0, most 3, 163 of 1000 groups produced at least one. 0 of 1000 reached the observed count or more, giving an empirical p of 0.0010 against a threshold of 0.05. Seed 42.

Selective — not common-essential

Gene	Entrez	Difference (Chronos)	Difference (SD)	p
SOX10	6663	-1.192	-2.92	9.38e-11
BRAF	673	-0.928	-2.58	1.83e-08
MAPK1	5594	-0.707	-2.17	4.46e-07
GRB2	2885	0.679	1.14	4.84e-08
ZEB2	9839	-0.553	-2.39	1.74e-08
MITF	4286	-0.456	-2.44	8.79e-07
PPP2R2A	5520	-0.450	-1.42	1.13e-07
ELOA	6924	-0.405	-1.77	1.33e-06

Gene	Entrez	Difference (Chronos)	Difference (SD)	p
LCMT1	51451	-0.395	-1.34	9.08e-07
CFLAR	8837	0.385	0.71	7.23e-07
RAB6A	5870	0.334	0.85	2.85e-07
SOX9	6662	-0.314	-1.22	5.67e-09
EFR3A	23167	0.292	0.73	4.80e-07
HEATR3	55027	-0.289	-1.00	1.18e-07
GGNBP2	79893	0.289	1.15	1.93e-07

Common-essential — read separately

These genes are required in most cell lines. They are listed apart from the selective genes above because general fragility and selective dependency are different claims, and ranking them together would confuse the two.

Gene	Entrez	Difference (Chronos)	Difference (SD)	p
YRDC	79693	-0.690	-0.78	6.74e-10
CHMP4B	128866	-0.526	-0.83	1.32e-06
TXN	7295	-0.510	-1.34	5.55e-07
PTPN11	5781	0.427	1.27	3.44e-10
KRAS	3845	0.388	0.59	9.73e-13
PPP1R15B	84919	0.303	0.64	2.44e-08
POLA2	23649	0.195	0.76	1.30e-06

This list is mostly non-essential, so it is not simply restating general fragility.

What this comparison can and cannot say

This comparison comes from the tissue, not from the compound. It contrasts every cell line of this indication against the rest of the cohort, so it describes what this tissue depends on whether or not the compound does anything to it — melanoma lines depending on SOX10 is a fact about melanoma, not about the drug. Read it as background biology the responder hypothesis sits inside, never as evidence that these genes explain the predicted response.

Pre-specified check — BRAF

In the responder population BRAF has a mean gene effect of -0.470 against -0.099 in the rest of the cohort — a difference of -0.370 Chronos (-1.03 standard deviations), stronger in the responder population, at an uncorrected p of 4.71e-04.

One test, not a search

One pre-specified test of a supplied target gene — not a discovery scan. The p-value is uncorrected because exactly one hypothesis was tested, and it therefore cannot be compared with, or ranked against, the genome-wide results above.

What was deliberately not computed

- **Dependency differences per subgroup.** Not computed: at the median zone size of 5 cell lines the minimum detectable effect is 2.49 SD, and random groups that size produce a median of 17 false findings. The comparison would manufacture results.
- **A genome-wide hunt for subgroup-selective dependencies.** Not computed: same power limit, with no pre-specified hypothesis to constrain the search.

DepMap CRISPR (Chronos) gene effect · release 23Q2, verified by exact MD5 match against the published figshare record for that release; DOI, licence and attribution are in *Method and limitations*. · 377 of 468 cell lines in this analysis have CRISPR data; the other 91 are dropped, never imputed. · Sign convention — more negative = stronger dependency; 0 = no effect; about -1 = the median common essential. · No non-essentials control list is present in the directory; only the inferred common-essentials list is available. · Comparisons run only at 25 cell lines or more: measured, not assumed: 1000 random cell groups per size, same test and correction, give a median of 2 false findings at n=7, 1 at n=11, and 0 from about n=25 upward. Below the floor the comparison is not run.

07 Named combinations — coverage and depth

What this section is

Combined response is the Bliss independence null: the two log₂ responses are added, which multiplies the linear viability fractions. That is the NON-INTERACTION expectation — what the pair covers together if neither drug changes what the other does. It is not a synergy claim, and drug-drug interaction, antagonism and pharmacokinetics are not modelled.

Coverage statistics describe reach and, by construction, treat overlap as a cost. On these statistics alone an irrelevant partner scores highly: a penicillin reaches the 90th percentile of its matched null against vemurafenib, above the canonical MEK partner. They are reported for completeness and must not be read on their own.

Partners appear in the order they were requested. They are described independently and are not ranked, scored against one another, or compared — the three patterns below are different kinds of result, not better and worse ones.

Possible outcomes:

- **Complementary** — the partner's responder subgroups are largely ones this compound misses; the value is broader coverage.
- **Amplifying** — the partner's responder subgroups overlap this compound's; the value is identifying which subgroup benefits most from the pair.
- **Mixed** — both patterns are present.
- **Dose-doubling** — the two compounds behave near-identically across subgroups, so the pair may be closer to more of the same drug than to a combination.
- **No pattern** — this compound has no responder subgroup in that scope, so there is nothing for a partner to overlap or extend.
- **Ambiguous** — the quantities do not sit clearly in one of the patterns above, or there are too few subgroups to tell them apart; reported as ambiguous rather than forced into one.
- **Null** — the partner adds no subgroup and overlaps none.

vemurafenib with trametinib

Mixed

Pattern within Skin

193

Subgroups the partner reaches alone

-0.254

Partner's own average (log₂)

0.06

Similarity to this compound

60

Similarity percentile of 805 partners

Not flagged as near-identical. This partner sits at the 59.9th percentile of similarity to vemurafenib, below the 99th this analysis treats as near-identical. Validated over 648,830 compound pairs from the approved pool, labelled same-class by shared mechanism or target. At this threshold the flag is 3.9 times more likely than chance to mark a same-class pair, and the pairs it marks are textbook duplicates. It is a ONE-WAY flag: it catches about one same-class pair in twenty, so a pair that is not flagged has NOT been shown to be mechanistically distinct. Read a raised flag as informative and an absent one as no information either way.

Scope	Subgroups in scope	This compound	Partner	Together	Shared	Overlap share	Depth inside (log ₂)	Partner only	Newly clearing
Pan-cancer	482	36	193	227	23	0.64	-0.498	170	195
Skin	56	11	28	33	10	0.91	-0.694	18	23

Pan-cancer — mixed. The partner responds in 23 of the query's 36 responder subgroup(s) (overlap fraction 0.64), adds 170 subgroup(s) of its own, and reaches -0.498 log₂ of additional depth inside the query's subgroups — both an overlapping and an extending contribution.

Matched comparison — **all 38** approved compounds of similar reach (within the 154–232 subgroup window), not a sample of them. Depth inside this compound's subgroups, median -0.266 log₂ against this partner's -0.498. Newly-clearing subgroups, median 222 against 195. · Reported for completeness. It does NOT cleanly separate a mechanistically relevant partner from an irrelevant one: in the reference set an irrelevant control reached the 65th percentile against a rational partner's 70th. · Secondary to the raw log₂ magnitude above, and NOT separating on its own: the irrelevant control reaches the 90th-95th percentile here. Read the magnitude.

Skin — mixed. The partner responds in 10 of the query's 11 responder subgroup(s) (overlap fraction 0.91), adds 18 subgroup(s) of its own, and reaches -0.694 log₂ of additional depth inside the query's subgroups — both an overlapping and an extending contribution.

Matched comparison — all 38 approved compounds of similar reach (within the 154–232 subgroup window), not a sample of them. Depth inside this compound's subgroups, median -0.281 log₂ against this partner's -0.694. Newly-clearing subgroups, median 28 against 23. · Reported for completeness. It does NOT cleanly separate a mechanistically relevant partner from an irrelevant one: in the reference set an irrelevant control reached the 65th percentile against a rational partner's 70th. · Secondary to the raw log₂ magnitude above, and NOT separating on its own: the irrelevant control reaches the 90th-95th percentile here. Read the magnitude.

How the pattern is decided

Rule applied, in order. (1) If node-vector similarity reaches the provisional percentile threshold of the query's own distribution, the pair is dose-doubling and no further pattern is assigned. (2) Otherwise the overlap fraction — the share of the query's responder subgroups the partner also responds in — and the count of partner-only subgroups are read together: a high overlap fraction with enrichment above its matched null is amplifying; a low overlap fraction with partner-only and crossing subgroups dominating is complementary; both present is mixed. A pair that does not sit clearly in one of these is reported as ambiguous rather than forced.

Depth inside the query's responder subgroups is arithmetically the partner's own response restricted to those subgroups: under Bliss the combination minus the query is exactly the partner. It therefore measures how strongly the partner acts in the subgroup the query already reaches — not an interaction between the two drugs.

Matched comparisons come from curated approved set (measured, >=1 true responder zone alone) — 804 compounds, of which every one inside a partner's breadth window is used; the comparison is exhaustive over that window, so these percentiles are exact for it rather than estimated from a sample. Combined response is the sum of the two log₂ responses, which is the product of their linear viability fractions.

08 Method and limitations

Basis

- Cell-line data (DepMap / PRISM single-dose viability), resolved measured-first: 17 of 17 indications measured (vemurafenib), 0 imputed from chemical structure.
- Subgroups: indication-scoped node subgroups — 491-node vector (chemistry->node-mean regressor; varies within indication)
- Model measured, schema profile-1.7. Compound confidence **Elevated** — chemistry-similarity tier (max ECFP4 Tanimoto to v0.3 training set): elevated >=0.5 (committed node-rank rho ~0.62, 62.5% reach rho>0.5); moderate 0.4-0.5 (rho ~0.50, 50%); exploratory <0.4 (rho ~0.30, ~18%). Estimated, not guaranteed (committed outlier: doxorubicin tanimoto 1.0 yet rho 0.46).

Data attribution. Cell-line viability and dependency data from the Cancer Dependency Map (DepMap), Broad Institute. Compound viability data: PRISM Repurposing 19Q4, doi:10.6084/m9.figshare.9393293. Expression and CRISPR (Chronos) gene effect data: DepMap 23Q2 Public, doi:10.6084/m9.figshare.22765112. Both licensed CC BY 4.0 (creativecommons.org/licenses/by/4.0/). Data were modified: subset, reshaped and aggregated. This work is not endorsed by the Broad Institute.

Limitations

- **Node purity.** Relevant nodes are those containing at least 3 skin cell lines; predicted response is the node-mean over all of a node's cells, which are mixed-tissue, so a low indication purity means the score is not indication-pure.
- No held-out discrimination statistic is reported. Method 3a is AUC-free by design; any accuracy figure quoted elsewhere is an aggregate, not a per-node claim.
- Biomarker p-values are reported nominal (uncorrected) with a Bonferroni-corrected value alongside, across 1,353 expression genes (L1000 landmark plus MSK-CHORD cancer genes) and 470 mutation genes (MSK-CHORD clinical panel). Ranking uses the nominal value; correction is informational. Candidates clearing correction are marked in the relevance map.
- Single-dose viability conflates cytotoxic and cytostatic effects.
- Measured and imputed values are not on one scale.

Disclaimer — verbatim

Hypothesis-generating, not a validated prediction. Indication-scoped subgroups are computed on the node basis, whose per-drug ranking is weak (committed LDO per-drug node r ~0.072); node-mean accuracy is aggregate (pooled r ~0.39,

indication-level). Relevant nodes merely contain the target indication's cells -- the predicted response is the node-mean over all of a node's (mixed-tissue) members, so a low indication purity means the score is not indication-pure. Candidate biomarkers and the full-panel relevance map are leads to test, not validated markers: committed imputed biomarker transfer is weak (top-k Jaccard ~ 0). Held-out AUC is deliberately not surfaced at this resolution. Cell-line (DepMap/PRISM) triage, not a clinical or efficacy prediction. Validate experimentally before relying on any item.

Report PRO-DEMURAFENIB-SKIN-20260831 · 2026-08-31 · The Responder Lab · Public sample report. Hypothesis-generating, not a validated prediction, and not medical advice.