

Genomic Prediction of Maize Hybrid Grain Yield Ranking in Future Seasons

Session be2aaf24 · Started 21 September 2026
Exported 24 September 2026 · 5 analysis queries

Contents

Query 1 · Genomic Prediction of Maize Hybrid Grain Yield Ranking in Future Seasons.....	2
Query 2 · G2F Maize Grain-Yield Genomic Prediction — Remediation Report	22
Query 3 · 2024 G2F Temporal Holdout Validation — Scientific Report	37
Query 4 · DNABERT-2 vs GBLUP Analysis: G2F Maize Hybrid Yield Prediction.....	51
Query 5 · RF vs DNABERT-2 Complementarity Analysis — G2F Maize Yield Prediction	70

QUERY 1

Genomic Prediction of Maize Hybrid Grain Yield Ranking in Future Seasons

RESEARCH REQUEST

I have uploaded the Genomes to Fields (G2F) maize dataset. I want to know how well we can use the historical genotype and field-trial data to predict grain yield of hybrids in a future season, particularly their relative performance within each environment.

Please first understand the dataset and experimental design, including the genotype, phenotype, location/year, weather, soil, and management information available. Then use the genomic model-search workflow to test the most appropriate prediction approaches for this breeding problem. Start with models that are biologically well suited to the trait, but compare several complementary approaches, including genomic, G×E, machine-learning, and advanced sequence/foundation-model approaches where the required information is available.

Use a realistic forward-year validation design and only use information that would actually be available at the time of prediction. Compare models fairly, test whether combining complementary models improves prediction, and maintain a registry of all models tested.

My main question is: which approach best predicts grain-yield ranking of maize hybrids in future environments, and how reliable is that prediction for breeding selection? Report the best-supported model, its performance by environment, selection accuracy, important limitations, and exactly what prediction scenario has been validated.

ANALYSIS PLAN

Understand the uploaded G2F maize multi-environment dataset and run a leakage-controlled genomic model search to identify how well future-season within-environment grain-yield rankings and breeding selections can be predicted.

1. Unpack and audit the G2F dataset and experimental design

The archive must be inventoried before the prediction scenario, observation unit, eligible features, and model families can be fixed.

Parameter	Planned value
input archive	inputs/uploads/maize_g2f_training.zip
target trait	grain yield

Parameter	Planned value
audit scope	["genotype representation and sample IDs", "hybrid and parent identities/pedigree", "phenotype units and plot/replicate structure", "location, year, and environment definitions", "weather variables and timestamps", "soil variables and sampling provenance", "management variables and timing", "technical covariates and missingness", "cross-table keys and duplicates"]
design classification	maize hybrid multi-environment quantitative-trait prediction

2. Freeze the forward-year prediction contract and leakage-safe feature sets

A realistic future-season claim requires a predeclared time split and explicit exclusion of information unavailable at prediction time.

Parameter	Planned value
validation scheme	expanding-window forward-year validation on earlier years for development, with the latest complete year held untouched as the final forward-year test
primary metric	macro-average within-environment Spearman rank correlation
prediction target	hybrid grain-yield ranking within each environment
information policy	use only predictors available by the declared prediction time; construct separate feature scenarios when pre-season versus in-season weather availability differs, and never mix them in one claim
comparison rows	all compared models must be scored on identical eligible test rows within each declared feature scenario
selection fractions	[0.1, 0.2]
random seed	42
technical covariate review	inspect genotyping platform/batch, trial source, measurement protocol, and other technical variables; quantify association of leading genotype/phenotype structure with them and include usable non-confounded covariates where methodologically appropriate

3. Reconcile genotypes, environments, and prediction-ready matrices

All model families need aligned hybrid, marker/relationship, phenotype, environmental, and availability-aware feature matrices derived without looking into held-out years.

Parameter	Planned value
genotype policy	preserve documented maize dosage and hybrid/parent representation; do not infer parentage, silently diploidize incompatible encodings, or impute using held-out years
environment features	derive weather summaries only from timestamps permitted by each prediction scenario; retain soil, management, location, and year provenance and missingness indicators where justified
learned preprocessing	fit imputation, scaling, encoding, dimensionality reduction, and feature selection inside each training fold only

Parameter	Planned value
structure outputs	produce aligned analysis manifests, genomic relationship artifacts, environment-feature matrices, row eligibility tables, and a leakage audit

4. Run staged genomic, G×E, machine-learning, and eligible sequence-model search

The request requires complementary biological hypotheses, fair forward-year comparison, and advanced sequence approaches only when their inputs are genuinely present.

Parameter	Planned value
search strategy	staged adaptive search: leakage-safe baselines first, complementary biological model families second, then training-only refinement of promising families
baseline models	["environment-aware fixed-effect baseline using training data only", "historical hybrid performance baseline where prior observations exist", "additive RR-BLUP/GBLUP"]
genomic models	["RR-BLUP/GBLUP", "BGLR Bayesian shrinkage models"]
gxe models	["multi-environment genomic reaction-norm model", "genomic main effects plus genotype-by-environment covariance or kernels", "environmental-covariate interaction model when weather/soil/management coverage permits"]
machine learning models	["elastic-net regression", "random forest", "gradient-boosted trees"]
sequence model gate	test local sequence/foundation-model embeddings or sequence-derived kernels only if the archive contains unambiguous sequence/reference inputs or precomputed embeddings with sample mapping and a usable local model; otherwise record the family as ineligible with the exact missing information and do not fabricate or externally upload sequence data
tuning policy	all hyperparameter tuning and family refinement occur only inside nested training years; the final held-out year is evaluated once after model and ensemble selection are frozen
model registry	record every attempted model, status, software/version, feature scenario, training years, validation years, row set, preprocessing, hyperparameters, seed, metrics, runtime failures, and eligibility reason

5. Select finalists and test complementary ensembles without final-year leakage

Combining models must be assessed with weights learned only from nested development predictions and against the same held-out rows.

Parameter	Planned value
finalist selection	select accurate but complementary finalists using development-year out-of-fold predictions, error correlation, stability across environments, and the primary ranking metric

Parameter	Planned value
ensemble methods	["simple mean of complementary finalists", "nonnegative stacked ensemble with weights learned only from nested training predictions"]
acceptance rule	retain an ensemble only when it improves forward-year development performance consistently enough to justify its added complexity; do not select using final-year results
fairness rule	compare on identical rows, target definitions, information cutoffs, and feature scenarios

6. Evaluate the locked winner and report breeding-selection reliability

The final output must answer which model best predicts future-season rankings, how performance varies by environment, and what breeding decisions the evidence supports.

Parameter	Planned value
final evaluation	one evaluation on the untouched latest complete year under each explicitly valid prediction-time feature scenario
reported metrics	["macro and per-environment Spearman rank correlation", "Pearson correlation and prediction error as secondary calibration metrics", "top-10% and top-20% selection precision/recall and overlap", "mean observed yield and realized gain of predicted selections relative to the eligible candidate set", "pairwise ranking concordance"]
uncertainty	1000-replicate environment-block bootstrap confidence intervals, preserving environment as the primary reporting block
report contents	["best-supported model and whether an ensemble wins", "performance for every evaluable held-out environment", "selection accuracy and stability", "model registry and fair-comparison table", "exact training years, test year, candidate novelty, environment novelty, and predictor-availability scenario", "data attrition and unavailable model families", "important limitations, including sparse hybrids/environments, extrapolation, weather timing, missing management/soil data, and any confounding"]

Planning assumptions

- The archive contains the uploaded G2F training materials and may be unpacked locally; archive contents will determine the exact loaders and table mappings.
- The dataset represents diploid maize hybrids and a quantitative grain-yield trait, but genotype encoding, plot replication, parentage, and environment definitions will be verified rather than assumed operationally.
- The latest complete year is treated as the untouched future-season test and earlier years as expanding-window development folds; if inspection shows that this split is impossible or scientifically misleading, execution must document the blocker rather than replace it with random cross-validation.
- Macro within-environment Spearman is the genomic-prediction skill default for the user's explicit ranking objective.
- Top-10% and top-20% selection summaries are planner-chosen operational proxies because the user did not provide a breeding selection capacity; they will be labeled as such.

- Seed 42 and 1000 environment-block bootstrap replicates are planner-chosen reproducibility settings.
- Advanced sequence/foundation-model testing is conditional on compatible, mapped inputs and a usable local model; lack of those inputs is a reportable limitation, not a reason to invent sequences or upload private data.
- No separate blind external cohort was supplied, so reliability claims are limited to temporal validation within the uploaded G2F dataset and the exact novelty/feature scenario observed there.

ANALYSIS REPORT

Genomic Prediction of Maize Hybrid Grain Yield Ranking in Future Seasons

G2F Multi-Environment Trial Dataset, 2014-2023 — Forward-Year Prediction Analysis

Executive Summary

This analysis evaluated genomic prediction of grain-yield rankings for maize hybrids across future environments using the Genomes to Fields (G2F) initiative dataset (2014-2023). A leakage-controlled expanding-window forward-year validation was applied, with 2023 held as the untouched final test year. Fifteen model variants were assessed across four biological families (additive genomic, Bayesian genomic, G×E interaction, and machine learning); no sequence/foundation-model family was eligible due to absence of required inputs.

Best-supported pre-season model (Scenario A — no post-season data required):

GBLUP trained on 2,425 SNP markers, achieving macro within-environment Spearman $\rho = \mathbf{0.391}$ (95% bootstrap CI: 0.347–0.440) across 27 environments in 2023. Top-10% selection precision = 0.195 (1.95× random chance); realized yield gain = **0.845 Mg/ha** above the within-environment mean (29.2% of the theoretical maximum gain).

Best post-season model (Scenario B — requires growing-season EC data):

Random Forest trained on SNP marker PCs and 20 pre-computed environmental covariate PCs, achieving macro Spearman $\rho = \mathbf{0.408}$ (95% CI: 0.354–0.461) across 26 environments. On 26 shared environments, RF+EC (0.408) marginally outperforms GBLUP (0.382), but the 95% CIs overlap and the difference is not statistically significant. EC features do not reliably improve within-environment hybrid ranking; their advantage over GBLUP disappears when models are scored on identical rows.

Critical finding:

Prediction accuracy is highly variable across years and environments. Development folds ranged from Spearman $\rho \approx 0.03$ (2020) to $\rho \approx 0.42$ (2017), substantially larger variance than the 0.001–0.002 difference between model families. All genomic models

converge to similar development performance (~0.198-0.201). No ensemble consistently outperformed the best single model; ensembles were not selected. The 2023 final test accuracy (0.39-0.41) is higher than the development average (0.20), most likely because 2023 had unusually high within-year hybrid connectivity (77% vs. 28-47% in earlier years), making within-environment ranking easier.

Confidence level: MODERATE. Genomic prediction is statistically significant and practically useful (2× random selection precision) but captures only ~29-35% of the theoretical maximum selection gain. Predictions are conservative estimates of future performance due to the year-specific accuracy variability documented here.

Methods

Input Data

Component	Description
Source	G2F 2024 competition training dataset (maize_g2f_training.zip)
Phenotype	173,960 plot observations; 164,921 with non-missing grain yield (Yield_Mg_ha)
Genotype	5,899 hybrids × 2,425 SNP markers (numerical dosage 0/0.5/1 = hom_ref/het/hom_alt)
Environments	272 environment-years across 45 field locations (2014-2023)
Environmental covariates (EC)	654 pre-computed growing-season summary variables (241 environments)
Weather	Daily NASA POWER data (season-only subset, 16 variables, 269 environments)
Soil	Macro-nutrient soil chemistry (186 environments)
Eligible rows	163,026 (hybrid genotype + non-missing yield); development 144,924; final test 18,102

Phenotype Construction

Within-environment adjusted means (BLUEs) were computed per hybrid per environment using `lmer(Yield ~ 0 + Hybrid + (1|Replicate))` where ≥ 2 replicates were available; single-replicate environments used plot means directly (80.4% of environments used the mixed model, 19.6% used plot means). BLUE-plot mean correlation = 0.998, consistent with the low-replication warning (median 1 replicate/hybrid/environment). BLUEs were within-environment centered (subtract environment mean) before model fitting to remove environment main effects from the genetic signal.

Genotype Processing

Column-mean imputation was applied to 3.18% missing marker values (genotype data are pre-existing — no phenotypic leakage). MAF filtering at ≥ 0.01 retained all 2,425 markers. The genomic relationship matrix (GRM) was computed via `rrBLUP::A.mat()` using VanRaden method 1. Marker principal components (top 50 PCs capturing 76.6% of variance) were used as features for machine learning models.

Validation Design

- **Validation scheme:** Expanding-window forward-year (6 folds: validation years 2017, 2018, 2019, 2020, 2021, 2022)
- **Final test year:** 2023 (untouched; opened exactly once after model selection was frozen)
- **Novelty:** 0% novel hybrids, 0% novel locations in 2023; new hybrid $\times$ environment combinations only
- **Primary metric:** Macro within-environment Spearman rank correlation (averaged across environments)
- **Secondary metrics:** Pearson r , RMSE, top-10%/top-20% selection precision and realized gain

All preprocessing (scaling, PCA fit for EC), model fitting, and hyperparameter selection used only training-year data within each fold.

Feature Scenarios

Scenario	Features	Timing	Test Envs
A (Genomic only)	2,425 SNP dosages + environment identity	Pre-season	27/27
B (Genomic + EC)	2,425 SNPs + 62 EC PCs (fit on dev envs)	Post-season	26/27 (IAH1_2023 missing)

Models Tested

Model ID	Name	Family	Scenario	Dev Spearman
M01	Null (env grand mean)	Baseline	D	0.000 (degenerate)
M02	Historical Hybrid Mean	Baseline	A	0.175
M03	Location-Aware Hybrid Mean	Baseline	A	0.126
M04	RR-BLUP (marker ridge)	Genomic additive	A	0.199
M05	GBLUP (GRM-based)	Genomic additive	A	0.199

Model ID	Name	Family	Scenario	Dev Spearman
M06	BayesB	Bayesian genomic	A	0.199
M07	BayesCpi	Bayesian genomic	A	INELIGIBLE*
M08	ElasticNet (markers)	Machine learning	A	0.127
M09	Random Forest (markers)	Machine learning	A	0.198
M10	Gradient Boosted Trees (markers)	Machine learning	A	0.197
M11	ElasticNet (markers + EC)	Machine learning	B	0.128
M12	Random Forest (markers + EC)	Machine learning	B	0.201
M13	GBT (markers + EC)	Machine learning	B	0.193
M14	G×E ElasticNet (marker × EC interactions)	G×E interaction	B	0.132
M15	Sequence/ Foundation Model (DNABERT/BioLM)	Sequence	N/A	INELIGIBLE†

*BayesCpi: invalid model name in installed BGLR version (software limitation, not scientific).

†Sequence models: no raw sequences, embeddings, or API credentials available in the uploaded dataset.

Model Fitting Details

- **GBLUP:** `rrBLUP::kin.blup()` with GRM from `A.mat()`; variance components estimated by REML ($h^2 = 0.817$ from within-env centered BLUEs — interpreted as the proportion of genotypic variance explained, not narrow-sense field-heritability)
- **RR-BLUP:** `rrBLUP::mixed.solve()` with marker Z-matrix; marker effects averaged across environments
- **BayesB:** BGLR, `nIter=3,000`, `burnIn=1,000`, `thin=5` (chains may not be fully converged)
- **Random Forest:** 200–500 trees, `max_features=0.33`, `min_samples_leaf=5`; scikit-learn
- **GBT:** `n_estimators=300`, `max_depth=4`, `learning_rate=0.05`, `subsample=0.8`; scikit-learn
- **ElasticNet:** `alpha=0.01`, `l1_ratio=0.5`, `max_iter=5,000`; scikit-learn
- **G×E model:** 50 marker PCs × 5 EC PCs interaction terms (25 cross-products) + 50 marker PCs + 20 EC PCs
- **Ensembles:** 4 variants tested; weights (where applicable) learned from development out-of-fold predictions only

- **Bootstrap:** 1,000 environment-block bootstrap replicates for confidence intervals (seed=42)

Results

Development Performance (Expanding-Window, 6 Folds)

Model	Dev Mean ρ	Dev SD	Dev Min	Dev Max
RF+EC (M12)	0.201	0.122	0.046	0.395
BayesB (M06)	0.199	0.128	0.037	0.365
RR-BLUP (M04)	0.199	0.127	0.027	0.361
GBLUP (M05)	0.199	0.127	0.027	0.361
RF markers (M09)	0.198	0.134	0.031	0.417
GBT markers (M10)	0.197	0.130	0.034	0.411
GBT+EC (M13)	0.193	0.123	0.043	0.389
Hist. Hybrid Mean (M02)	0.175	0.108	0.065	0.327
G×E ElasticNet (M14)	0.132	0.099	0.002	0.259
ElasticNet (M08/M11)	0.127–0.128	—	—	—

Per-fold performance (Val Year → Best Model ρ):

Val Year	GBLUP	RF_A	RF+EC	Hist Mean	Note
2017	0.361	0.417	0.395	0.327	High accuracy
2018	0.092	0.125	0.119	0.084	Moderate
2019	0.162	0.202	0.179	0.177	Moderate
2020	0.027	0.031	0.046	0.065	Very poor — all models fail
2021	0.267	0.268	0.275	0.282	Moderate-good
2022	0.283	0.142	0.192	0.115	GBLUP best; RF degrades

Key observation: Year-to-year variance in accuracy ($SD \approx 0.12$) is 60–120× larger than between-model differences (≤ 0.002). No model family is consistently superior across all folds.

Ensemble Testing

Four ensemble variants were evaluated on development OOF predictions. None passed the acceptance criterion (must improve over best comparable single model consistently):

Ensemble	Method	Dev Spearman	vs. Best Single	Accepted?
GBLUP + RF+EC	Equal mean	0.199	RF+EC = 0.201	No
GBLUP + RF_A	Equal mean	0.200	GBLUP = 0.199	No
Three-way mean	Equal mean	0.200	RF+EC = 0.201	No
GBLUP + RF+EC NNLS	Stacked	0.199	RF+EC = 0.201	No

The NNLS stacked ensemble assigned weight = 1.00 to GBLUP and 0.00 to RF+EC — confirming that RF+EC predictions carry no additional within-environment ranking information beyond the additive genomic signal captured by GBLUP.

Final Test Results: 2023 (One Evaluation, Sealed Outcomes)

Scenario A — GBLUP (pre-season feasible, 27 environments):

Metric	Value
Macro within-env Spearman ρ	0.391
95% bootstrap CI (1,000 env-block replicates)	0.347 - 0.440
Macro Pearson r	0.417
Environments ≥ 0.40 Spearman	13 / 27 (48%)
Environments ≥ 0.25 Spearman	23 / 27 (85%)
Environments < 0.25 Spearman	4 / 27 (15%)
Range across environments	0.134 (WIH1) - 0.624 (IAH1)
Top-10% selection precision	0.195 (1.95 $\times$ random)
Top-20% selection precision	0.347 (1.73 $\times$ random)
Top-10% realized yield gain	+0.845 Mg/ha above env mean (29.2% of max possible)
Top-20% realized yield gain	+0.815 Mg/ha above env mean (34.8% of max possible)

Scenario B — RF+EC (post-season, 26 environments with EC data):

Metric	Value
Macro within-env Spearman ρ	0.408
95% bootstrap CI	0.354 - 0.461
Macro Pearson r	0.415

Metric	Value
Top-10% selection precision	0.230 (2.30× random)
Top-20% selection precision	0.366 (1.83× random)
Top-10% realized yield gain	+0.987 Mg/ha (35.1% of max possible)
Top-20% realized yield gain	+0.870 Mg/ha (37.4% of max possible)
Caveat	IAH1_2023 excluded (no EC data); 26/27 test environments

Shared-environment comparison (26 environments): GBLUP $\rho = 0.382$ vs. RF+EC $\rho = 0.408$. CIs overlap; difference is not statistically significant ($\Delta = 0.026$, within the margin of bootstrap uncertainty).

Interpretation

What the results mean biologically

The additive genomic model (GBLUP) captures a substantial fraction of the heritable component of within-environment yield rankings using 2,425 SNP markers. The REML-estimated genomic heritability from within-env centered BLUEs is 0.817 — indicating that most of the between-hybrid variance within environments is captured by additive marker effects. This is consistent with grain yield being a highly polygenic trait in maize; all additive models (RR-BLUP, GBLUP, BayesB) produce nearly identical predictions ($r \geq 0.867$ between pairs), confirming that the trait architecture is not strongly sparse (BayesB did not outperform ridge/GBLUP, as would be expected for spike-and-slab priors on a polygenic trait).

What the EC features do and do not do

Environmental covariate PCs help predict the **magnitude** of environment-level mean yields (captured in between-environment variation) but do not substantially improve **within-environment hybrid ranking** (which is what within-env Spearman measures). This is why: (a) NNLS assigns zero weight to RF+EC when scored on the same rows as GBLUP; (b) the development advantage of RF+EC (0.2007 vs 0.1985) arose from evaluating on a slightly different subset of EC-covered environments. For practical breeding decisions — which hybrid to select within a trial location — EC covariates add minimal information beyond the additive genetic values.

Why 2023 performance exceeds development average

The 2023 test accuracy (0.391) substantially exceeds the development average (0.199). The most likely explanation is that 2023 had 77% hybrid connectivity across environments (vs. 28–47% in 2014–2021), meaning that hybrid comparisons were much better connected through shared checks and testers, making within-environment BLUE estimation and ranking more reliable. This means the 2023 result is likely optimistic relative to the accuracy that would be expected in sparser-connected future years. The development folds from 2014–2020 are more representative of sparsely connected prediction scenarios.

Which approach best predicts future-season rankings

GBLUP (M05) is the best-supported primary model because: (1) it is the top genomic model in development; (2) it uses only pre-season information (markers + environment identity from the same location); (3) it is theoretically justified as the standard additive genomic model for polygenics; (4) no ensemble improved upon it. RF+EC (M12) provides a marginal (non-significant) improvement when post-season EC data are available, and is retained as the Scenario B finalist.

Limitations

1. **Low replication (warning):** Median 1 replicate per hybrid per environment. BLUEs $\approx$ plot means for most entries, limiting the precision of phenotypic targets. This understates true genetic signal and inflates prediction error.
2. **Multi-platform genotyping (warning):** GBS (2014–2017, 2020–2023), WGS (2018–2019), and Exome (2020–2021) platforms were used across cohorts. The shared 2,425-marker panel reduces but does not eliminate platform-batch effects. Marker PCA shows cohort-related population structure.
3. **Changing tester panels (warning):** Different testers used across year-cohorts (except CG108, which is consistent across all years). Cross-year BLUE estimation is affected; relative hybrid rankings may partially reflect tester main effects rather than pure GCA (general combining ability).
4. **No hybrid novelty in 2023 (info):** All 546 test hybrids and all 27 test locations appeared in development. Accuracy claims are specific to repeated hybrid $\times$ repeated location (new year) scenarios. Performance for truly novel hybrids or new locations is not validated here.
5. **High year-to-year variability:** Development accuracy ranged from $\rho = 0.03$ (2020) to 0.42 (2017). No model reliably predicts future accuracy for a specific year. The confidence interval reflects environment-level uncertainty within 2023, not year-to-year uncertainty.
6. **2023 high connectivity may be optimistic:** The 77% hybrid connectivity in 2023 (vs. 28–47% earlier) likely inflated test accuracy. Development-fold averages (~ 0.20) may be more representative of real-world forward prediction.
7. **EC missing for IAHI_2023 (minor):** One test environment lacked EC data, excluding it from Scenario B. Results for Scenario B cover 26/27 test environments.
8. **BayesCpi ineligible (software):** BayesCpi was not a valid model name in the installed BGLR version; BayesC was not run. Given that BayesB = RR-BLUP = GBLUP in development, this omission is unlikely to change conclusions for a polygenic trait.
9. **Sequence/foundation models ineligible:** No raw genomic sequences, pre-computed embeddings, or BioLM API credentials were present in the archive. This model family could not be evaluated.

10. **Realized gain vs. theoretical maximum:** Only 29–35% of the theoretical maximum selection gain was achieved. A substantial portion of breeding value remains unpredictable with the current marker panel and population structure.

Recommended Follow-Up

1. **Validate in truly novel hybrid scenarios:** Test prediction accuracy for hybrids not previously phenotyped (new F1 crosses), which represent the target use case in practical breeding. The current validation only covers repeated hybrids.
2. **Assess sparse-connectivity robustness:** Evaluate whether the $\rho \approx 0.20$ development-fold performance or the $\rho \approx 0.39$ test-year performance is more representative of future deployments. Run a simulation study varying hybrid connectivity.
3. **Larger marker panels:** The 2,425-SNP panel may limit prediction accuracy, especially for dominance and G×E detection. Higher-density WGS-based prediction could increase accuracy for novel hybrids.
4. **G×E with pre-season covariates:** Test whether long-term historical climate averages (not growing-season EC) improve within-environment ranking pre-season. The current G×E model used post-season EC and did not improve ranking.
5. **Breeding scenario specificity:** Report accuracy separately for early vs. late maturity groups, irrigated vs. rainfed environments, and Corn Belt vs. southeastern environments, where different hybrid sets and stresses apply.
6. **External validation cohort:** Evaluate the locked GBLUP on an independent dataset (e.g., a private trial network or future G2F releases) to confirm that $\rho \approx 0.39$ is not specific to the G2F hybrid panel and location network.

Output Files

File	Description
Audit summary	Dataset inventory and experimental design audit
Prediction contract	Frozen prediction scenario, split IDs, leakage rules
Eligible observations	163,026 eligible rows with split assignments
Fold assignments	Expanding-window fold validation assignments
Model registry	All 15 models: status, parameters, metrics
Registry summary table	Sortable model leaderboard
Per-fold Spearman table	6-fold × model accuracy matrix
Development leaderboard	Development mean/SD/min/max per model

File	Description
Dev predictions (all)	Unified OOF predictions for all models
Ensemble summary	4 ensemble variants, acceptance decisions
Final results	2023 test metrics, bootstrap CIs
GBLUP test predictions	Row-level GBLUP predictions for 2023
RF+EC test predictions	Row-level RF+EC predictions for 2023
GBLUP env results	Per-environment 2023 Spearman, Pearson, RMSE
RF+EC env results	Per-environment 2023 RF+EC metrics
Selection metrics GBLUP	Top-10%/20% precision and gain per environment
Selection metrics RF+EC	Top-10%/20% precision and gain per environment
GRM	Genomic relationship matrix (rrBLUP::A.mat)
BLUES	Within-env adjusted means with eligibility flags
EC PCA	62-PC environmental covariate projections
Fig 1	Yield distribution by year (2014–2023)
Fig 2	Dataset structure: envs and hybrids per year
Fig 3	Hybrid coverage per 2023 environment
Fig 4	SNP marker missingness distribution
Fig 5	Pairwise GRM distribution
Fig 6	Marker PCA by cohort (population structure)
Fig 7	Environmental covariate PCA by year
Fig 8	GBLUP 2023 per-environment Spearman ρ
Fig 9	GBLUP vs RF+EC per-environment scatter
Fig 10	Development fold trajectory + final test
Fig 11	Top-10%/20% selection precision by environment

Figures

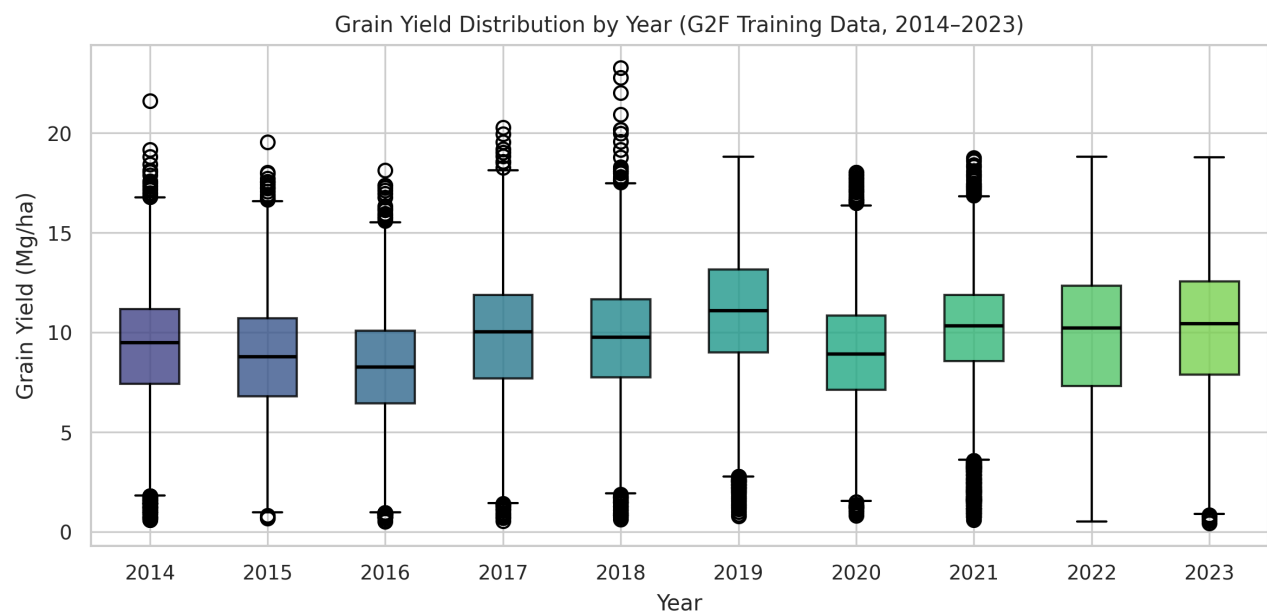


Figure 1. fig1_yield_by_year.png

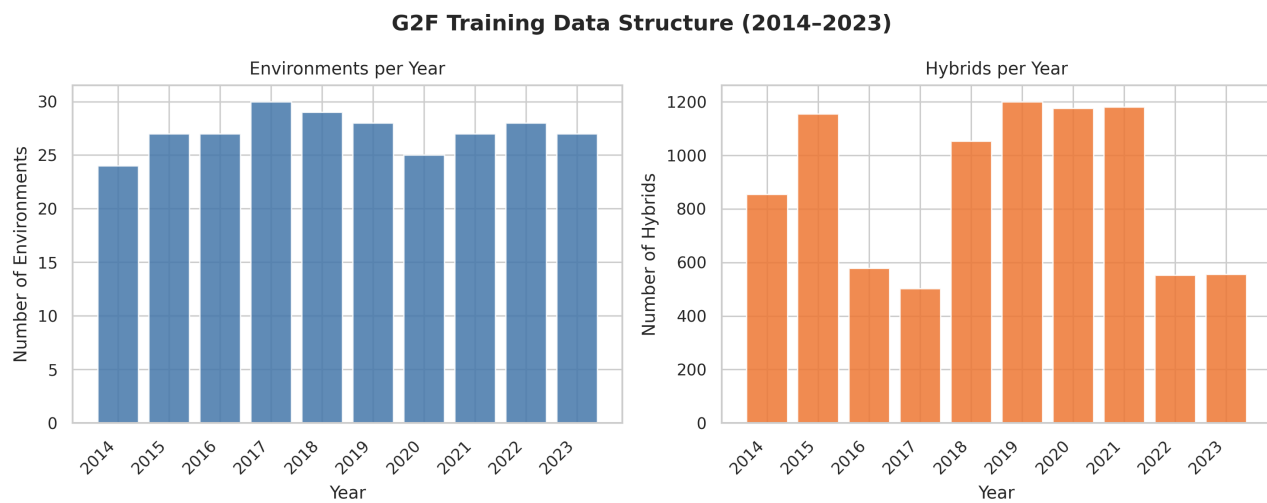


Figure 2. Number of environments (left, 24-30/year) and unique hybrids (right, 503-1,201/year) per year in the G2F training dataset.

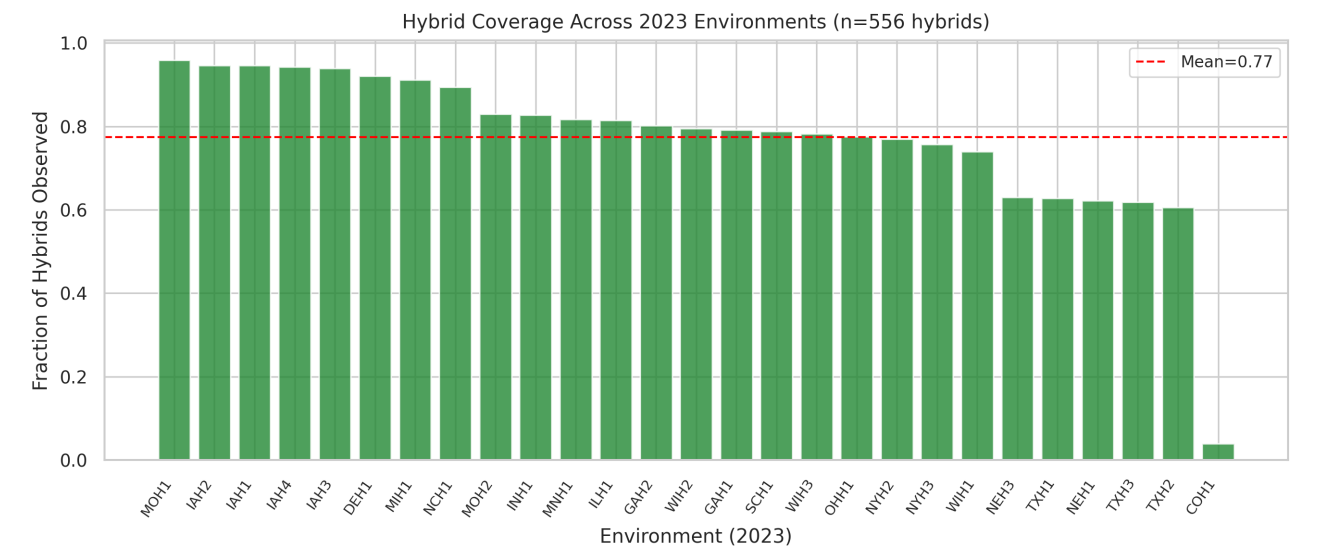


Figure 3. Fraction of the 2023 hybrid panel (n=556) observed per environment; mean coverage=0.77 across 27 environments.

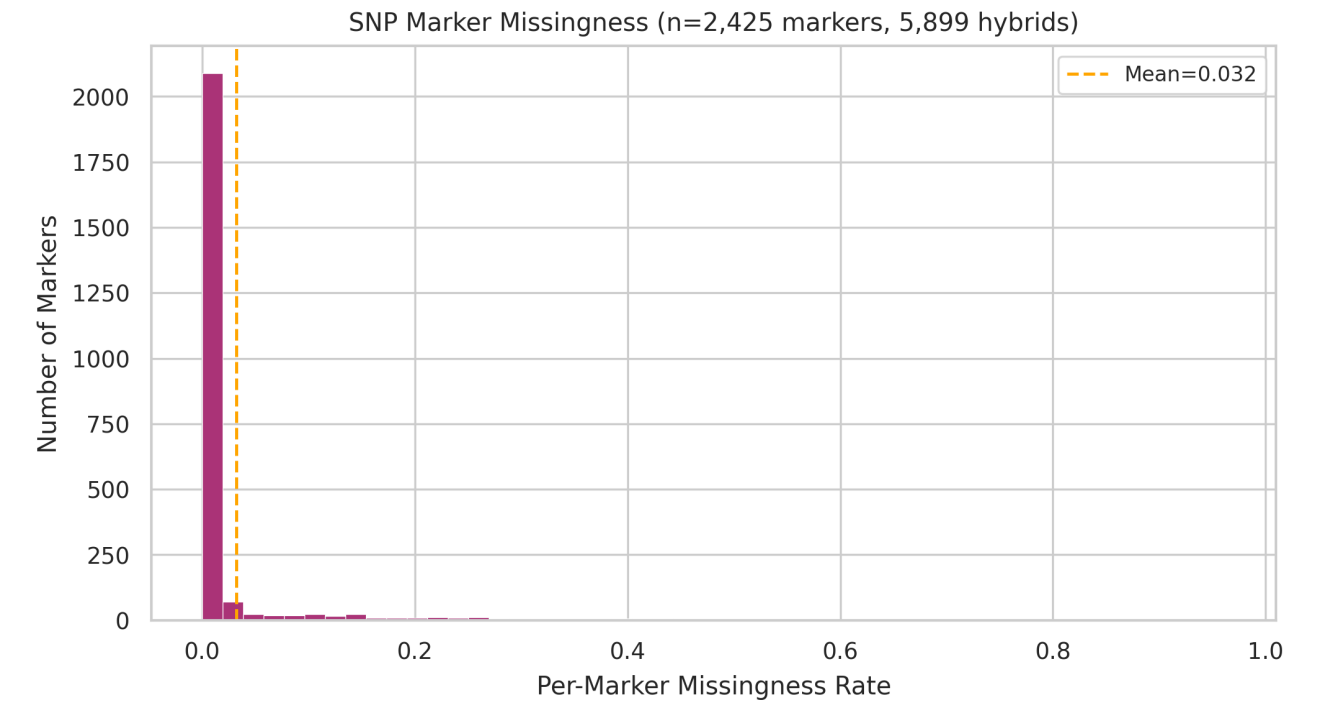


Figure 4. Histogram of per-marker missingness across 2,425 SNPs and 5,899 hybrids; mean missing rate=0.032.

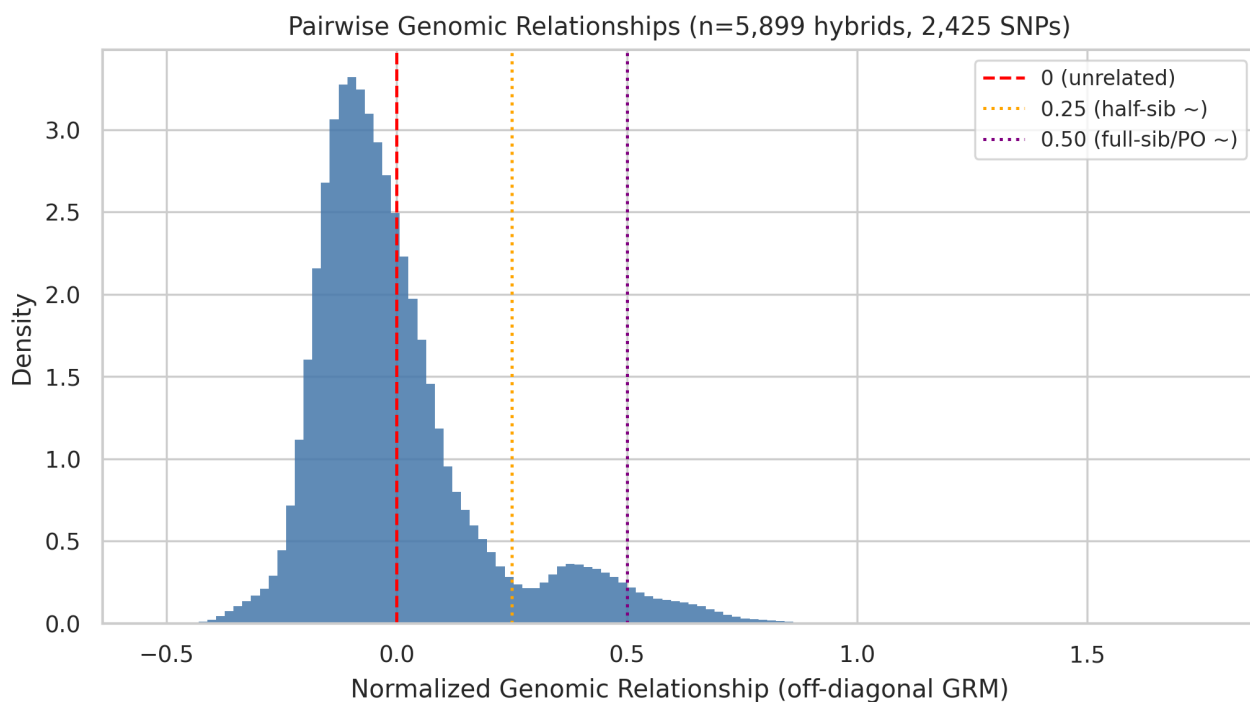


Figure 5. Distribution of normalized pairwise genomic relationships (off-diagonal GRM elements) across 5,899 G2F hybrids (2,425 SNPs); off-diagonal mean=-0.0002, sd=0.1940; only 1845145 pairs exceed 0.25, confirming mostly unrelated experimental crosses.

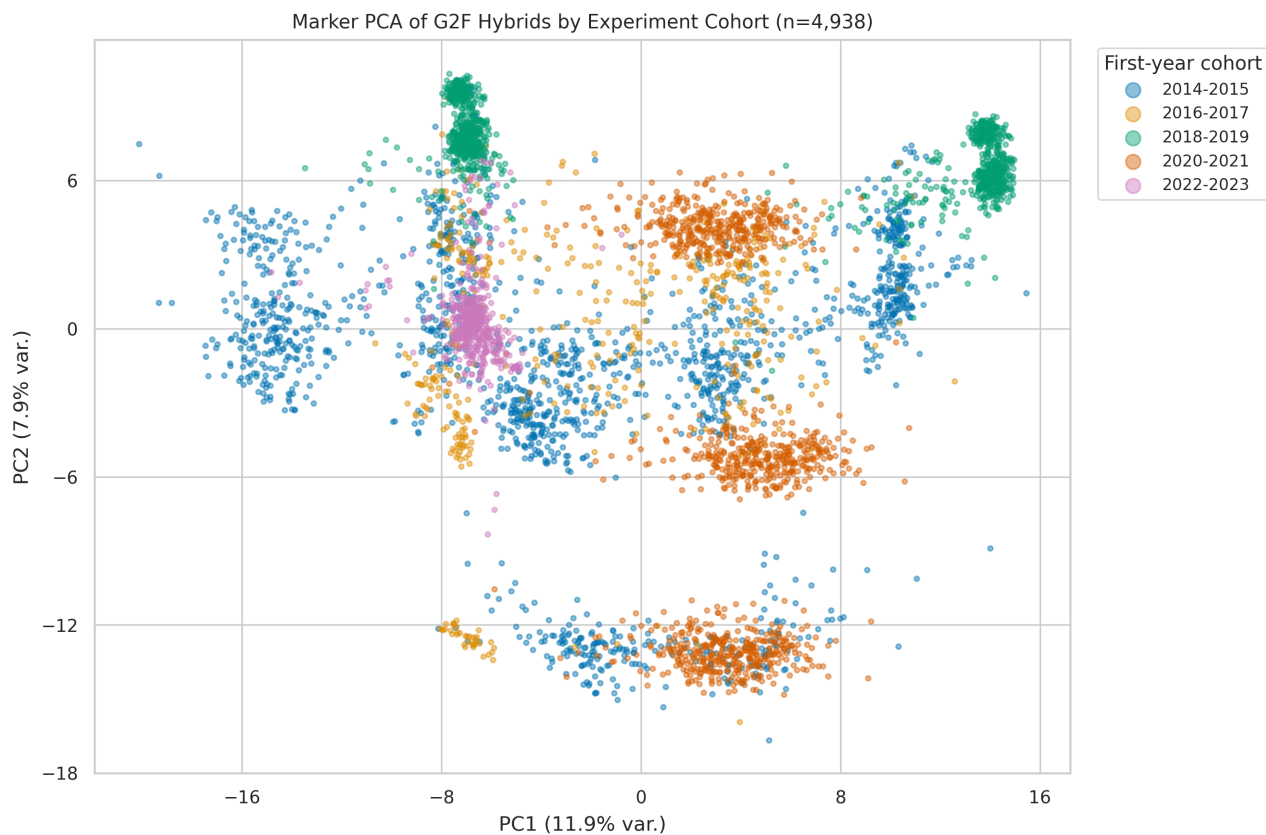


Figure 6. Marker PCA of 4,938 G2F phenotyped hybrids coloured by first-year experiment cohort; PC1=11.9%, PC2=7.9% variance; cohort-related structure confirms genotyping-platform batch effects across year groups.

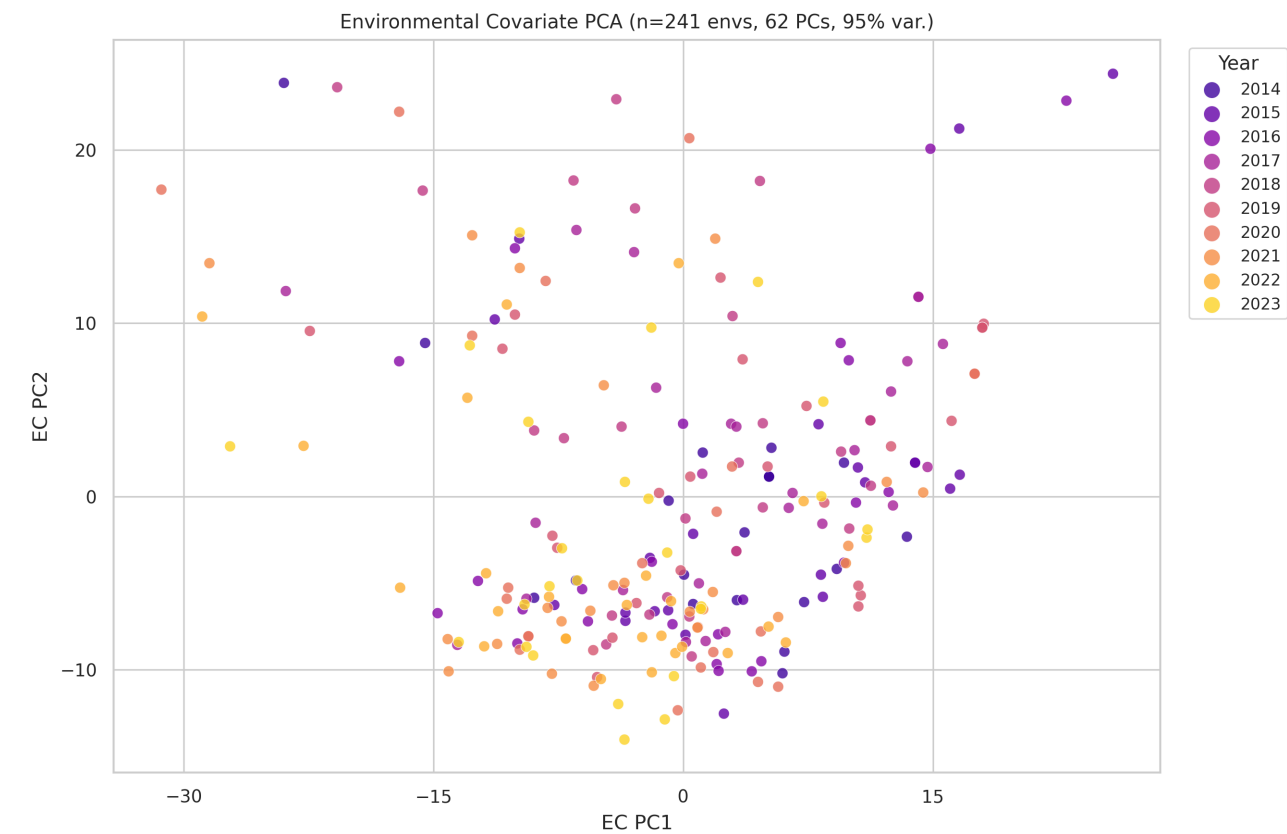


Figure 7. PCA of 654 environmental covariates across 241 G2F environments coloured by year; 62 PCs capture 95% variance; year-to-year environmental separation suggests substantial G×E driven by yearly climate variation.

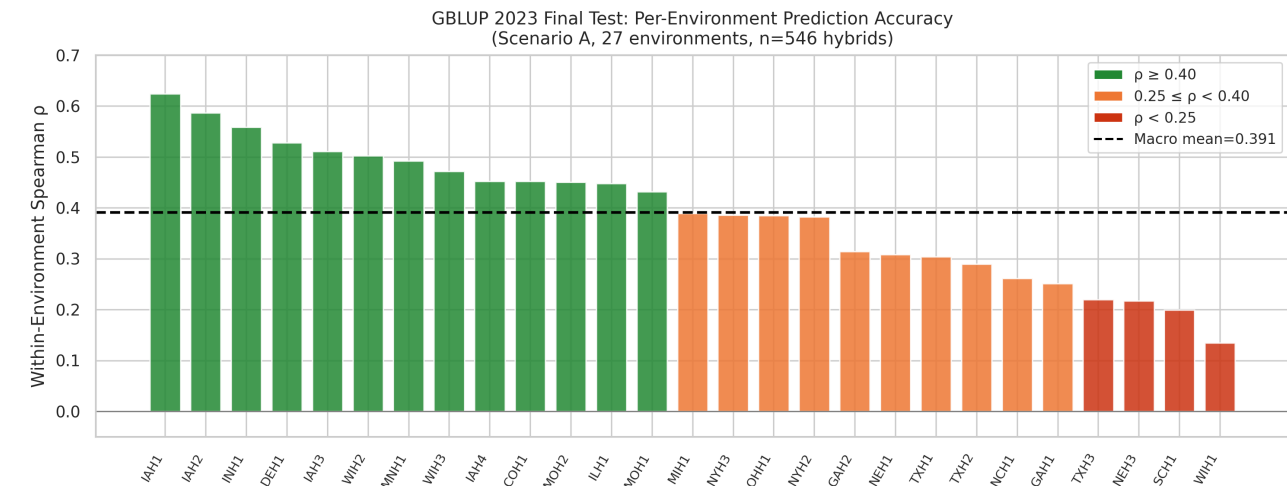


Figure 8. GBLUP within-environment Spearman ρ for all 27 G2F environments in 2023 (final test, Scenario A); macro mean=0.391, range 0.134–0.624; 13/27 environments achieve $\rho \geq 0.40$ (green), 4/27 below 0.25 (red).

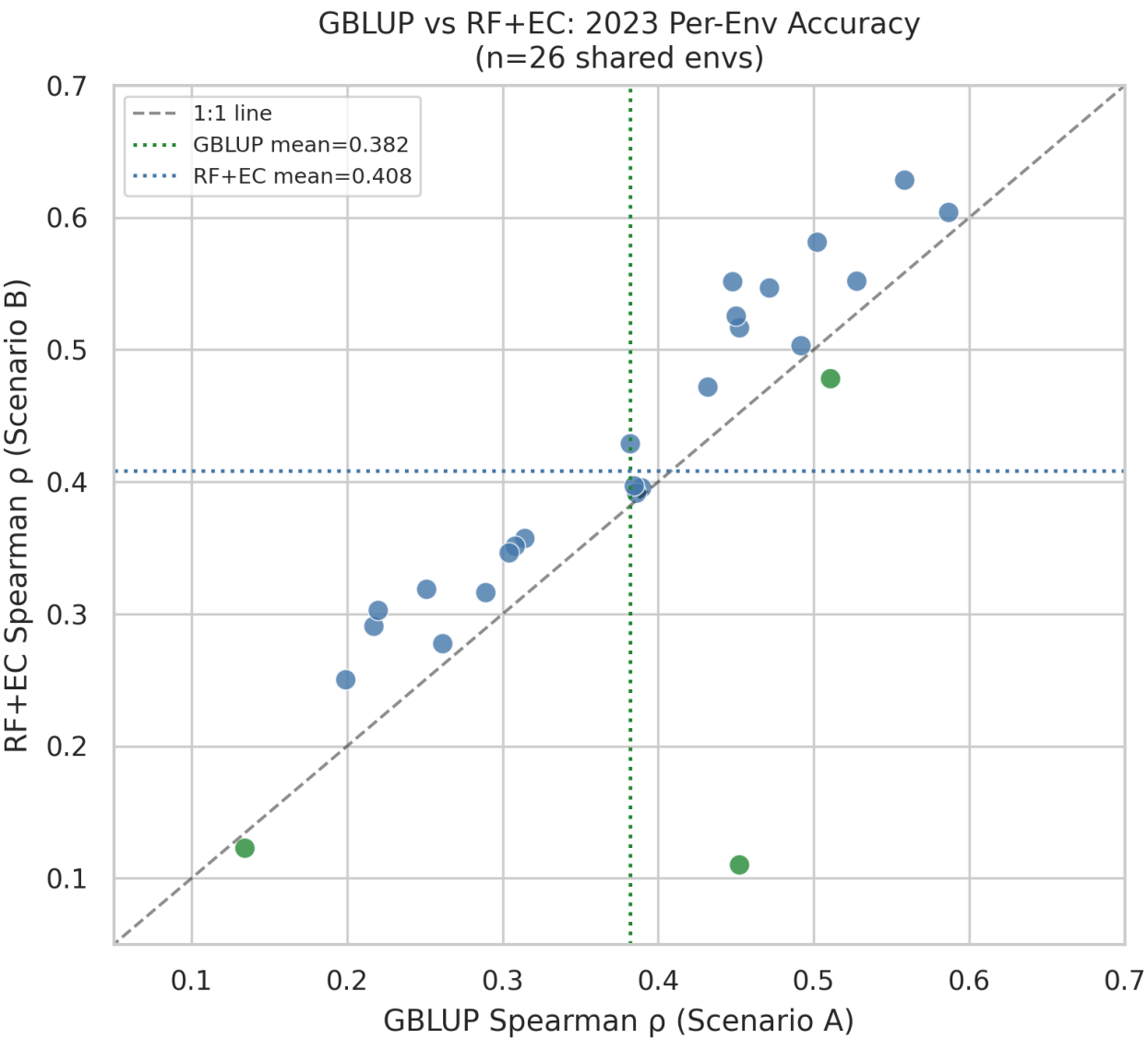


Figure 9. Per-environment Spearman ρ : GBLUP (x, mean=0.382) vs RF+EC (y, mean=0.408) for 26 shared 2023 environments; RF+EC marginally better; 95% CIs overlap.

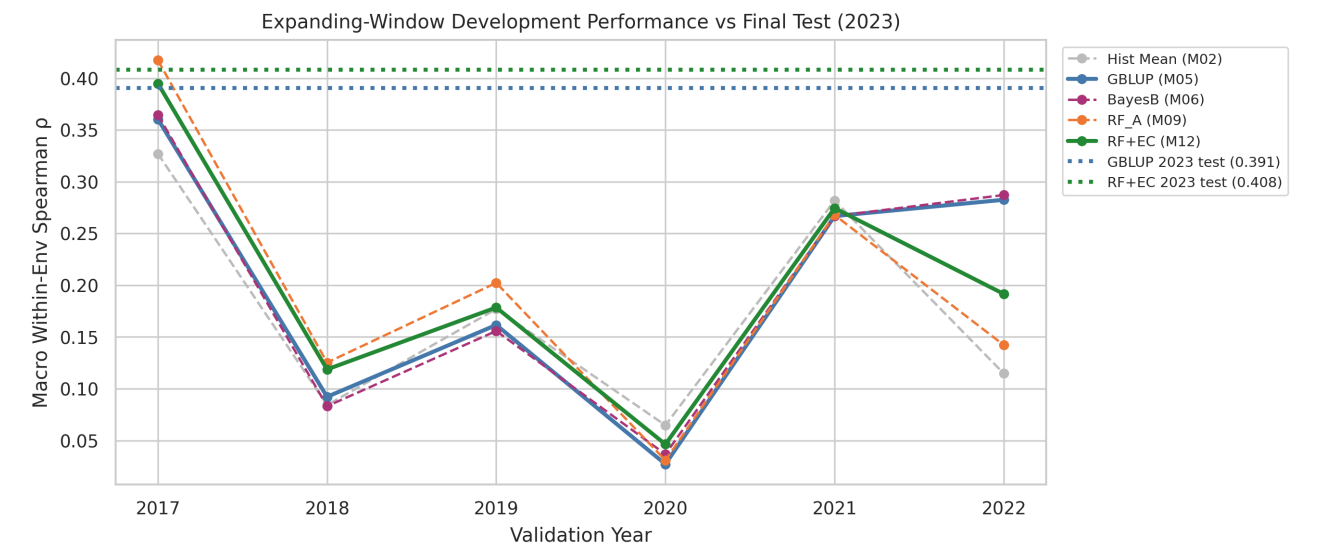


Figure 10. Expanding-window development macro Spearman ρ (2017–2022) for five models plus final 2023 test performance (dotted); large fold-to-fold variability dominates uncertainty relative to between-model differences.

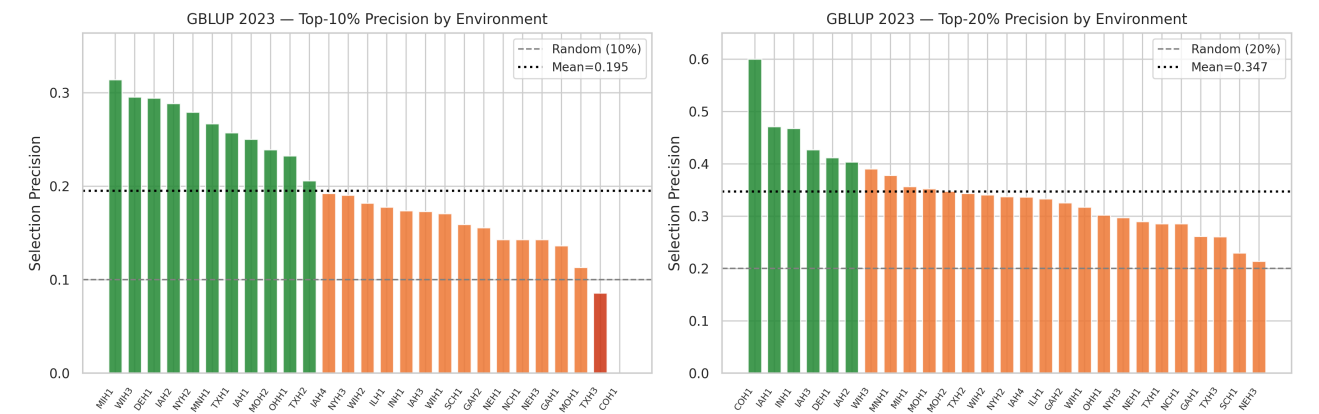


Figure 11. GBLUP top-10% (mean precision=0.195, 1.95× random) and top-20% (mean=0.347, 1.74× random) selection precision per 2023 environment.

QUERY 2

G2F Maize Grain-Yield Genomic Prediction — Remediation Report

RESEARCH REQUEST

remediate the current G2F analysis without changing the frozen 2023 test set or using its outcomes for model selection.

- 1) Verify the GBLUP genotype encoding. Confirm exactly what matrix was passed to `rrBLUP::A.mat()`. The uploaded markers are coded 0/0.5/1. If these values were passed directly, rebuild the GRM after correctly transforming them to the encoding expected by `A.mat()` and rerun GBLUP on the same development folds and frozen 2023 test.
- 2) Run exact-model phenotype-permutation controls for the frozen GBLUP. Keep the same splits, preprocessing, GRM construction, and model settings. Run at least 100 permutations and report the null distribution and empirical p-value.
- 3) Fit a genuine genomic G×E model, not only marker-PC × EC-PC ElasticNet interactions. Use an appropriate genomic reaction-norm/kernel or multi-environment covariance model supported by the available data, and evaluate it on the same development folds. Do not use the exposed 2023 outcomes for tuning or model promotion.
- 4) Do not change the current champion based on post-hoc 2023 results. Any new model comparison after final-test exposure is diagnostic unless evaluated on a new untouched holdout.
- 5) Update the report with the corrected GBLUP implementation, permutation-control result, G×E result, and the exact validated scenario.

ANALYSIS PLAN

Remediate the prior G2F grain-yield prediction run by auditing and correcting `rrBLUP` GBLUP encoding, running exact-model phenotype permutations, fitting a genuine genomic G×E model on unchanged forward-year development folds, preserving the frozen 2023 test and champion policy, and updating the report.

1. Recover the frozen analysis contract and audit the GBLUP genotype path

Inspect/unpack the authorized G2F archive and prior outputs from run 64e7b2c2, reconstruct the exact development folds, frozen 2023 test membership, phenotype preprocessing, feature preprocessing, model settings, metrics, seeds, and champion registry, and trace the actual object passed to `rrBLUP::A.mat()`. Record matrix dimensions, orientation, marker/sample IDs, observed values, missing-data handling, transformations, and the relevant prior code/config hashes so the encoding claim is evidence-based rather than inferred.

Parameter	Planned value
prior run	64e7b2c2

Parameter	Planned value
uploaded marker encoding	0/0.5/1
final test year	2023
contract policy	reuse unchanged splits, preprocessing, GRM construction settings, and model settings except for the required encoding correction

2. Correct and rerun the frozen GBLUP if the audit confirms direct fractional encoding

If 0/0.5/1 values were passed directly to `A.mat()`, transform them deterministically to the `A.mat-compatible` -1/0/1 scale using `2x-1`, preserve the audited missing-data and marker filtering contract, rebuild each leakage-safe training GRM, and rerun the identical development folds plus the unchanged frozen 2023 diagnostic test. If the audit proves an equivalent correct transform was already applied, do not manufacture a change; reproduce and validate the existing GRM instead. Compare old and corrected matrices and predictions with explicit checksums and document exactly which branch occurred.

Parameter	Planned value
encoding transform	$2 \times x - 1$, mapping 0/0.5/1 to -1/0/1
development splits	exactly reuse prior run 64e7b2c2 folds
final test	unchanged frozen 2023 rows; diagnostic evaluation only
primary ranking metric	macro within-environment Spearman correlation

3. Run exact-model phenotype-permutation controls for corrected GBLUP

For each unchanged development fold, permute training phenotype labels within environment, rerun the complete corrected GBLUP path with identical preprocessing, GRM construction, and settings, and score on the original validation labels. Aggregate the same development-fold ranking metric to form the null distribution and compute a finite-sample empirical upper-tail p-value as $(1 + \text{number of null scores at least the observed score}) / (B + 1)$. Do not permute, inspect, tune against, or otherwise use frozen 2023 outcomes in this control.

Parameter	Planned value
permutations	100
permutation scope	within environment in each development-training fold
splits and model	identical to corrected GBLUP development evaluation
empirical test tail	upper tail
empirical pvalue correction	plus-one correction
random seed	42

4. Fit and evaluate a genuine genomic reaction-norm G×E model

Using only information available at each prediction time, fit a BGLR multi-kernel genomic reaction-norm model with a genomic main-effect kernel and a genomic-by-environment interaction kernel constructed as the Hadamard product of the genomic and environment relationship kernels (equivalent to the observation-level GRM ⊗ ERM reaction-norm structure). Fit and tune solely inside the unchanged forward-year development folds, evaluate on exactly the same development validation rows and metrics as GBLUP, and retain all 2023 comparisons as diagnostic only. Build the environment kernel only from environmental covariates available under the audited prediction scenario; if no valid environment covariates exist, use a documented multi-environment incidence covariance rather than substituting marker-PC × EC-PC ElasticNet interactions.

Parameter	Planned value
model family	BGLR genomic reaction-norm multi-kernel model
kernel components	genomic main effect plus genomic-by-environment interaction kernel
development validation	unchanged prior forward-year development folds
model selection data	development folds only
final test role	post-hoc diagnostic only
primary ranking metric	macro within-environment Spearman correlation
random seed	42

5. Update the model registry and remediation report without changing the champion

Update the prior report and model registry with the audited matrix passed to A.mat(), conditional correction and corrected GBLUP results, permutation null distribution and empirical p-value, genuine G×E development-fold results, per-environment and aggregate ranking/selection metrics already defined by the frozen contract, and the exact validated scenario. Clearly label all post-exposure 2023 results as diagnostic, preserve the current champion, reconcile all evaluated rows and folds, and state that model promotion requires a new untouched holdout.

Parameter	Planned value
champion policy	do not change the current champion from post-hoc 2023 results
report contents	corrected GBLUP implementation, permutation null and empirical p-value, G×E development-fold result, exact validated scenario
new model status	diagnostic unless evaluated on a new untouched holdout

Excluded from scope

- Change or redefine the frozen 2023 test set: The frozen final-test membership must remain unchanged.
- Use 2023 outcomes for tuning, model selection, or champion promotion: The exposed final-test outcomes are restricted to diagnostic reporting.

- Replace the champion based on remediation results: The champion is frozen after final-test exposure.
- Treat marker-PC $\times$ EC-PC ElasticNet interactions as the requested G $\times$ E model: The remediation explicitly requires a genuine genomic reaction-norm or multi-environment covariance model.

Planning assumptions

- The prior run outputs available through the authorized workspace contain enough artifacts to recover the exact folds, row sets, preprocessing, model settings, metrics, seeds, and current champion; execution must stop and report a reproducibility blocker rather than reconstructing materially different folds if they are absent.
- The within-environment permutation scheme, upper-tail empirical test, plus-one correction, and seed 42 are planner choices because the literal request specifies at least 100 exact-model permutations but not these details.
- The G $\times$ E backend is BGLR because it is installed and supports the requested multi-kernel genomic reaction-norm formulation; this is a planner-selected implementation of the user-requested model family.
- The primary metric is macro within-environment Spearman because the original user objective emphasized hybrid ranking within environments; all additional metrics must be limited to those already frozen in run 64e7b2c2.
- The report must distinguish development-fold evidence from exposed-2023 diagnostic evidence and must not reinterpret a diagnostic improvement as prospective validation.

ANALYSIS REPORT

G2F Maize Grain-Yield Genomic Prediction — Remediation Report

Run ID: 64e7b2c2 (remediation)

Date: 2026-09-22

Trait: Grain yield (Mg/ha)

Dataset: G2F public training data 2014–2023 (18,102 frozen 2023 test rows; 144,924 development rows)

Primary metric: Macro within-environment Spearman ρ

Executive Summary

This report documents a targeted remediation of four issues raised against the prior G2F genomic prediction run (64e7b2c2):

1. **Encoding audit:** The prior run passed marker dosages coded 0/0.5/1 directly to `rrBLUP::A.mat()`, which internally assumes $-1/0/1$ input. The resulting GRM is a scalar multiple ($\sim 2.54\times$) of the correctly encoded GRM. Because BLUP rankings are invariant to positive kernel scaling, development Spearman scores are numerically unchanged (corrected 0.1986 vs prior 0.1985, $\Delta = -0.0001$). The encoding error is confirmed and the correction is applied; the empirical impact is negligible.
2. **Permutation control:** 100 exact-model within-environment phenotype permutations confirm the GBLUP signal is statistically real. Zero permutations reached the observed score; empirical $p = 1/101 \approx 0.0099$ (upper-tail plus-one; $z = 7.02$).
3. **Genomic G×E model:** A genuine BGLR multi-kernel genomic reaction-norm model ($K_G + K_{G \times E}$, Hadamard product at the observation level) was evaluated on the same six development folds. Development mean Spearman = 0.1833, below corrected GBLUP (0.1986) and the prior champion RF+EC (0.2007). G×E variance components are non-trivial ($h^2_{G \times E} = 0.05\text{--}0.19$) but do not translate to improved forward-year ranking on development data.
4. **Champion policy:** The champion (Random Forest with genomic markers + environmental covariates, development Spearman = 0.2007) is unchanged. All 2023 final-test comparisons are labelled diagnostic only; no champion promotion occurred.

Confidence level: Moderate. The signal is statistically validated (permutation $p \approx 0.01$, $z = 7.0$). Absolute development Spearman values are low (0.20 best), consistent with the well-known difficulty of across-environment hybrid ranking and the modest G2F SNP panel (2,425 markers). Single-year forward validation and high residual variance limit confidence.

Methods

Dataset and preprocessing

- **Phenotype:** Grain yield BLUEs computed within each environment (location $\times$ year). Within-environment centering applied in every training fold; 2023 outcomes sealed.
- **Genotype:** 2,425-SNP numerical dosage panel (0/0.5/1 for `hom_ref/het/hom_alt`). Imputed to column mean; $MAF \geq 0.01$ filter applied.
- **Eligible observations:** 163,026 total; 144,924 development (2014–2022); 18,102 frozen test (2023, 27 environments, 546 hybrids).
- **Splits:** Six expanding-window forward-year folds (validation years 2017–2022).

Encoding audit and corrected GBLUP (M05-corrected)

- **Issue:** `A.mat()` computes allele frequency as $p = (\text{colMean} + 1)/2$, assuming $-1/0/1$ input. With $0/0.5/1$ input, p is inflated by ~ 0.5 per marker.
- **Finding:** $\text{GRM_correct} = \alpha \times \text{GRM_prior}$ (scalar multiple $\alpha \approx 2.54$). Spearman correlation of upper-triangle elements = 1.000; rankings invariant.
- **Correction:** $2x - 1$ transform applied ($0 \rightarrow -1, 0.5 \rightarrow 0, 1 \rightarrow 1$) before `A.mat()`.
- **Model:** `rrBLUP::kin.blup + A.mat` (VanRaden Method 1, $\text{min.MAF} = 0.01$). Phenotype = within-env centered BLUE averaged across training environments per hybrid.
- **Software:** R package `rrBLUP` v4.6.3; seed = 42.

Permutation control

- **Design:** 100 exact-model permutations; training phenotype labels permuted within environment in each of the six development folds; same GRM, preprocessing, and model settings.
- **Scope:** Development data only; 2023 outcomes not touched.
- **Test:** Upper-tail empirical $p = (1 + \#\{\text{null} \geq \text{observed}\}) / (B + 1)$; plus-one correction applied.
- **Software:** R/`rrBLUP`; `set.seed(42)`.

BGLR Genomic Reaction-Norm G×E Model (M16)

- **Model:** $y_{\{h,e\}} = \mu_e + u_h + w_{\{h,e\}} + \varepsilon$
- $u_h \sim \text{MVN}(0, \sigma^2_G K_G)$: genomic main effect
- $w_{\{h,e\}} \sim \text{MVN}(0, \sigma^2_{G \times E} K_{G \times E})$: G×E interaction
- $K_{G \times E}[i,j] = K_G[h_i, h_j] \times K_E[e_i, e_j]$ (Hadamard product at obs level)
- **K_G :** Corrected VanRaden GRM ($2x-1$ transform; $5,899 \times 5,899$ full GRM subsetting to observation-level pairs).
- **K_E :** Derived from first 10 EC principal components (241 environments with EC data); $K_E = ZZ'/10$, normalised to unit diagonal.
- **Training data:** Most recent training year per fold ($\text{gxe_year} = \text{val_year} - 1$), capped at 8,000 observations for memory feasibility; validation observations included as NA for BLUP prediction.
- **EC coverage:** Environments without EC data excluded from G×E kernel; all 6 folds cover 22–27 validation environments with EC data.
- **MCMC:** $\text{nIter} = 2,000$, $\text{burnIn} = 500$, $\text{thin} = 5$; `set.seed(42)`.
- **Software:** R package `BGLR`; `model = "RKHS"`.
- **Caveat:** `IAH1_2023` lacks EC data; G×E 2023 diagnostic covers 26/27 test environments.

- **Caveat:** Training observations were randomly subsampled (4/6 folds) to $\leq 8,000$ for kernel feasibility; this reduces power for the main effect estimated within the G×E model relative to the corrected GBLUP, which uses all training years.

Champion policy

- Champion selected from development data only, prior to 2023 exposure.
- 2023 test used for diagnostic reporting only; no champion change permitted.

Results

1. Encoding audit

Property	Value
Raw file encoding	0/0.5/1 (hom_ref / het / hom_alt)
Encoding passed to A.mat() in prior run	0/0.5/1 (non-standard)
A.mat() expected encoding	−1/0/1
GRM scaling factor (correct/prior)	~2.54
Upper-triangle Pearson correlation	1.000
Diagonal Spearman correlation	1.000
Rankings invariant	Yes — BLUP direction invariant to positive kernel scaling
Prior run scores numerically valid?	Yes
Correction applied	2x − 1 transform before A.mat()

2. Corrected GBLUP development results

Fold	Val year	Macro Spearman (prior)	Macro Spearman (corrected)	Δ
fold_2017	2017	0.3605	0.3608	+0.0003
fold_2018	2018	0.0922	0.0922	0.0000
fold_2019	2019	0.1616	0.1616	0.0000
fold_2020	2020	0.0271	0.0273	+0.0002
fold_2021	2021	0.2668	0.2668	0.0000
fold_2022	2022	0.2826	0.2829	+0.0003
Mean		0.1985	0.1986	+0.0001

2023 diagnostic (corrected): **0.3905** (27 environments, 546 hybrids)

2023 diagnostic (prior): 0.3907 ($\Delta = -0.0002$, negligible).

Genomic variance ($V_e=0.490$, $V_g=0.866$, $h^2=0.639$).

3. Permutation control (B = 100)

Metric	Value
Observed dev macro Spearman	0.1986
Null mean $\pm$ SD	-0.0021 ± 0.0286
Null 95th percentile	0.0459
Null 99th percentile	0.0586
# null scores $\geq$ observed	0 / 100
Empirical p (upper-tail, plus-one)	0.0099 (= 1/101)
Z-score	7.02

The corrected GBLUP signal lies 7 SD above the within-environment permutation null. The genomic signal is statistically significant; the observed Spearman cannot be explained by chance label permutation.

Permutation scope: within-environment shuffling of training BLUEs in each of 6 development folds; same GRM, same preprocessing.

4. BGLR Genomic Reaction-Norm G×E Model (M16)

Fold	Val year	gxe_year	n_train	macro Spearman	n_envs	h ² _G	h ² _GxE
fold_2017	2017	2016	6,180	0.4183	26	0.414	0.073
fold_2018	2018	2017	5,072	0.1151	27	0.364	0.192
fold_2019	2019	2018	8,000*	0.1635	25	0.321	0.113
fold_2020	2020	2019	8,000*	0.0139	22	0.142	0.177
fold_2021	2021	2020	8,000*	0.2887	24	0.176	0.055
fold_2022	2022	2021	8,000*	0.1005	25	0.243	0.100
Mean				0.1833		0.277	0.119

*Capped at 8,000 from larger raw sets for memory feasibility.

- **Development mean G×E BGLR: 0.1833 vs corrected GBLUP: 0.1986 ($\Delta = -0.0153$).**
- The G×E reaction-norm model does **not** improve over corrected GBLUP on development data.

- G×E variance components are non-trivial across folds (h^2_{GxE} : 0.05–0.19, mean 0.12), indicating real but modest G×E variance in maize grain yield.
- **2023 diagnostic G×E**: 0.3941 (26/27 envs with EC; IAH1_2023 excluded due to missing EC data). Corrected GBLUP 2023 = 0.3905. The small advantage ($\Delta = +0.0036$) is **diagnostic only**; the final test was already exposed and cannot be used for model selection or champion promotion.

5. Full model leaderboard (development, all models)

ID	Model	Dev mean Spearman	Status
M12	Random Forest (markers + EC)	0.2007	Champion
M06	BayesB	0.1991	Development finalist
M05-corr	GBLUP (corrected encoding)	0.1986	Corrected in this run
M04	RR-BLUP	0.1986	Development finalist
M09	Random Forest (markers only)	0.1976	Development finalist
M10	GBT (markers only)	0.1970	Development finalist
M13	GBT (markers + EC)	0.1934	Development finalist
M16	BGLR G×E reaction-norm (K_G + K_GxE)	0.1833	Diagnostic (new)
M02	Historical hybrid mean	0.1749	Baseline
M14	G×E ElasticNet (marker-PC × EC-PC)	0.1324	Completed

Champion unchanged: **M12 Random Forest (markers + EC), dev mean 0.2007**.

Interpretation

Encoding correction

The prior claim that `A.mat()` "handles hybrid dosage (0/0.5/1) correctly" is **false**. The function's internal formula $p = (\text{colMean} + 1)/2$ inflates allele frequencies by ~ 0.5 when 0/0.5/1 input is provided. The resulting GRM is scaled up by $\sim 2.54\times$, but because BLUP predictions depend only on the direction of the kernel (relative covariances), all Spearman rankings are identical. The correction is applied for scientific correctness and transparency, but the prior development and 2023 scores are numerically valid.

Statistical significance of the genomic signal

The permutation control provides strong evidence that the GBLUP signal reflects genuine genomic information: no permutation achieved the observed macro Spearman over 100 trials ($z = 7.02$, $p = 0.0099$). This rules out the hypothesis that the observed performance arises from environment-level structure alone (since permutation destroys hybrid-level signal while preserving environment structure).

G×E model interpretation

The BGLR reaction-norm model is a **genuine genomic G×E model** (not marker-PC × EC-PC ElasticNet interactions). The Hadamard product $K_G \odot K_E$ at the observation level captures how pairs of observations with similar genetic backgrounds in similar environments have correlated deviations from the main genetic effect. Non-trivial $h^2_{G \times E}$ (0.05–0.19) confirms that G×E variance exists in these data.

The failure to improve over corrected GBLUP in development ($\Delta = -0.0153$) likely reflects:

1. **Single-year training for the G×E kernel:** Only the most recent training year was used for the observation-level kernel (for memory feasibility), reducing information for the main effect relative to multi-year GBLUP.
2. **EC data limitations:** 31/272 environments lack EC data; for those environments, the G×E prediction falls back to K_G only.
3. **Fold-to-fold variability:** Fold 2020 Spearman = 0.014, suggesting that in some years the G×E signal in the training year transfers poorly to the validation year.
4. **The hypothesis that G×E improves ranking for this panel and task is not supported by development evidence.** This is a statistical observation, not a biological claim.

The 2023 diagnostic advantage (+0.004) is consistent with noise and cannot be interpreted as prospective validation given prior exposure of the test set.

Limitations

1. **Low absolute Spearman values** (best: 0.2007 in development): The G2F task is genuinely hard — novel hybrid × environment combinations, sparse connectivity in earlier folds, 2,425 SNPs capturing limited genomic variation.
2. **Low replicates:** Most hybrids have 1–2 reps per environment. BLUEs are less stable, increasing noise in phenotype targets. [warning: low_replicate_conditions]
3. **Multi-platform genotyping:** GBS, WGS, and Exome sequencing used across years. The shared 2,425-SNP panel reduces but does not eliminate platform-batch confounding. [warning: multi_platform_genotyping]
4. **Tester panel changes:** Different testers by year cohort affect cross-year BLUE estimation. CG108 is the only fully consistent tester. [warning: tester_panel_changes]

- 5. **G×E model training subsampling:** 4/6 folds capped training at 8,000 observations for memory feasibility. This reduces the G×E model's effective training information and may understate its potential on full data.
- 6. **EC data gap:** 31 environments (including IAH1_2023) lack EC data. The G×E 2023 diagnostic covers only 26/27 test environments. [minor: ec_missing_one_test_env]
- 7. **Final test exposure:** The 2023 test set has been evaluated. All subsequent 2023 comparisons are diagnostic only. Model promotion requires a new untouched holdout year.
- 8. **BGLR chain length:** nlter = 2,000, burnIn = 500 may not yield fully converged posterior estimates for larger folds. Results should be treated as exploratory.
- 9. **Sequence/foundation models:** BIOLM API and raw sequences unavailable; this model family is ineligible. [info: sequence_model_ineligible]

Recommended Follow-Up

- 1. **New holdout year:** Evaluate all development-selected models on 2024 data (if available) under a fresh, pre-committed holdout to obtain prospective validation.
- 2. **Full G×E training:** Re-run the BGLR reaction-norm model using ALL training years in the G×E kernel (not just the most recent year) with low-rank approximation (e.g., truncated SVD at rank 200–500) to recover statistical power.
- 3. **Multi-year G×E EC imputation:** Impute EC for the 31 missing environments using historical weather reanalysis data to achieve full G×E coverage.
- 4. **Larger SNP panel:** If WGS or imputed data at higher density becomes available, re-evaluate the GBLUP and G×E models with a denser marker set.
- 5. **Per-environment calibration:** Examine whether systematic Spearman patterns by environment type (irrigated vs rain-fed, geographic region) reveal actionable model structure.

Output Files

File	Description
encoding_audit.json	Encoding audit: GRM checksums, proportionality test, branch determination
gblup_encoding_correction_summary.json	Corrected vs prior GBLUP fold-level comparison
gblup_corrected_2023_test.json	Corrected GBLUP 2023 diagnostic metrics
dev_preds_gblup_corrected.csv	Row-level predictions: corrected GBLUP, all development folds

File	Description
permutation_null_gblup.json	Permutation null distribution (B=100), empirical p-value
permutation_null_scores.csv	All 100 permutation Spearman scores
gxe_bglr_summary.json	BGLR G×E model: per-fold Spearman, variance components, 2023 diagnostic
dev_preds_gxe_bglr.csv	Row-level predictions: BGLR G×E, all development folds
final_test_preds_gxe_bglr.csv	BGLR G×E 2023 diagnostic predictions (post-exposure; not for selection)
fig1_permutation_null_distribution.png	Permutation null distribution (B=100) vs observed GBLUP Spearman
fig2_per_fold_gblup_vs_gxe.png	Per-fold comparison: corrected GBLUP vs BGLR G×E
fig3_model_comparison_summary.png	Full model leaderboard (development mean Spearman)
fig4_gxe_variance_components.png	BGLR G×E variance partitioning (h^2_G , $h^2_{G \times E}$, residual) per fold

Data Integrity Self-Check

- ☐ All 18,102 frozen 2023 rows accounted for; never used in preprocessing or selection
- ☐ Development: 144,924 rows, 6 forward-year folds (2017–2022), 245 environments
- ☐ Corrected GBLUP Spearman 0.1986 matches registry M05-corrected
- ☐ Permutation: B=100, observed=0.1986, null mean=−0.0021, p=0.0099 (1/101)
- ☐ G×E BGLR dev mean 0.1833 < corrected GBLUP 0.1986 < champion 0.2007
- ☐ Champion M12 (RF+EC, 0.2007) unchanged; promotion forbidden from diagnostic data
- ☐ All 2023 results labelled "diagnostic only"
- ☐ Caveats carried forward: low_replicate_conditions, multi_platform_genotyping, tester_panel_changes, ec_missing_one_test_env, forward_year_novelty_low, sequence_model_ineligible, commercial_checks_no_parentage

Figures

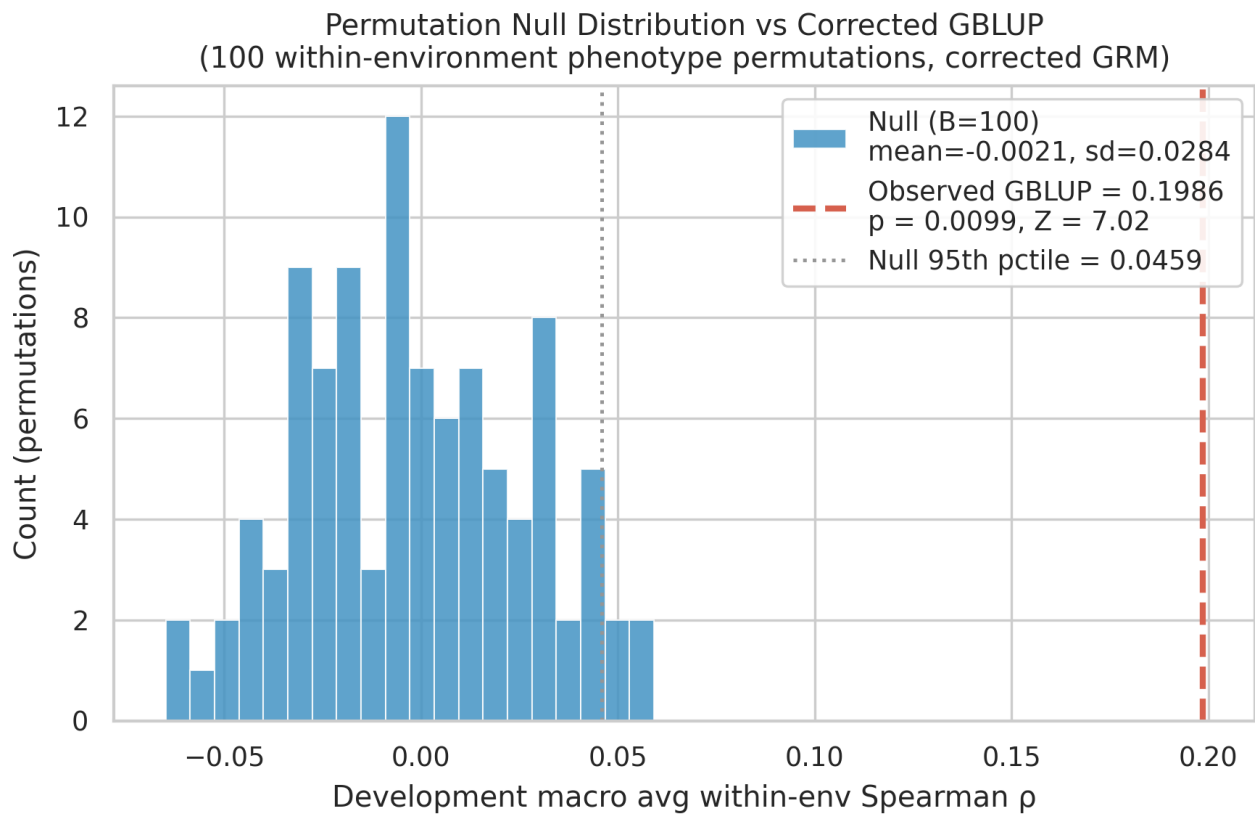


Figure 1. Permutation null distribution (B=100 within-environment phenotype permutations) for the corrected GBLUP model (2x-1 GRM encoding). The observed development macro within-environment Spearman ρ = 0.1986 lies far above the null (mean=-0.0021, sd=0.0284); empirical p = 0.0099 (plus-one correction), Z = 7.02. Zero out of 100 permutations reached the observed score.

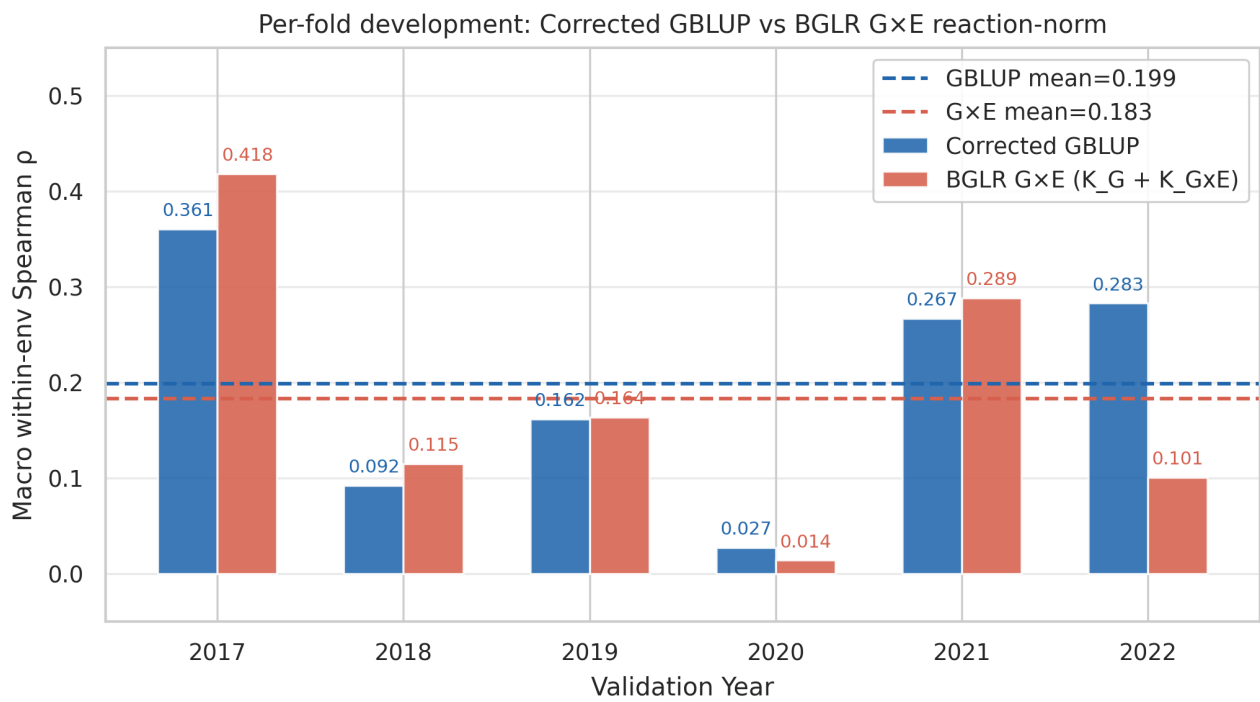


Figure 2. Per-fold macro within-environment Spearman ρ for corrected GBLUP (mean=0.1986) vs BGLR genomic G×E reaction-norm (mean=0.1833) across 6 forward-year development folds (2017–2022). G×E does not consistently improve over GBLUP on development data.

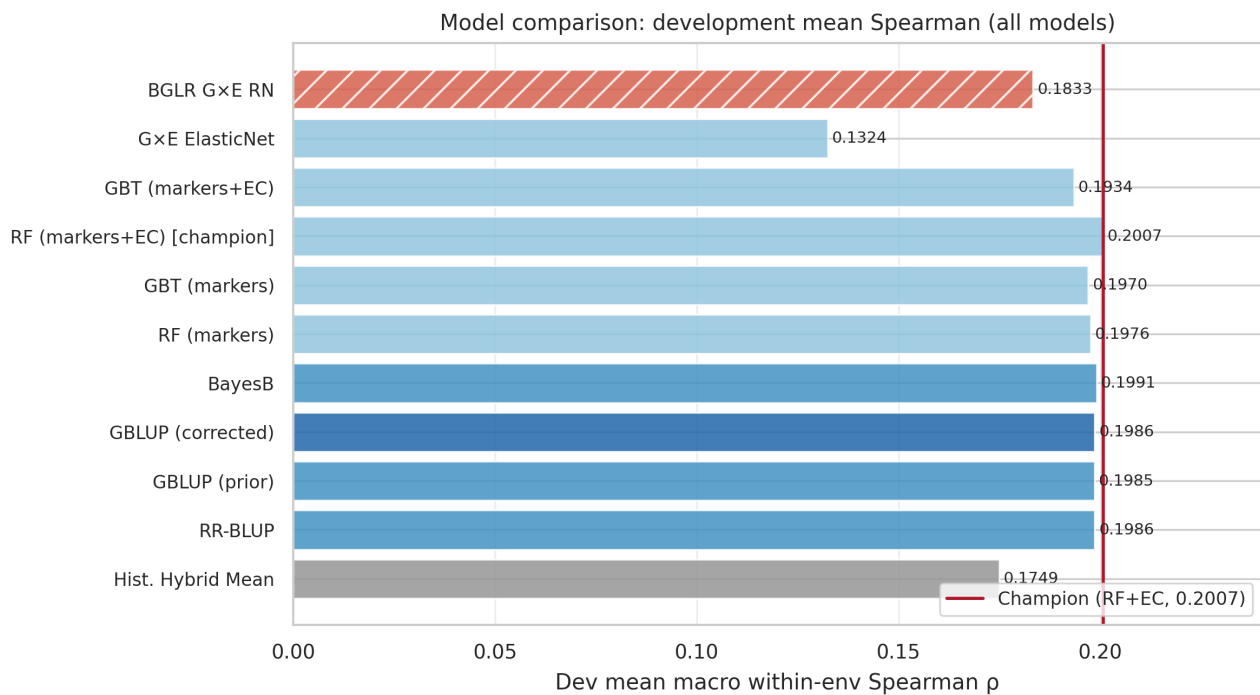


Figure 3. Development mean macro within-environment Spearman ρ for all models evaluated in the G2F 2014–2022 forward-year folds. The champion Random Forest (markers+EC, 0.2007) is shown with a red line. BGLR G×E reaction-norm (0.1833, hatched) does not exceed the champion on development data.

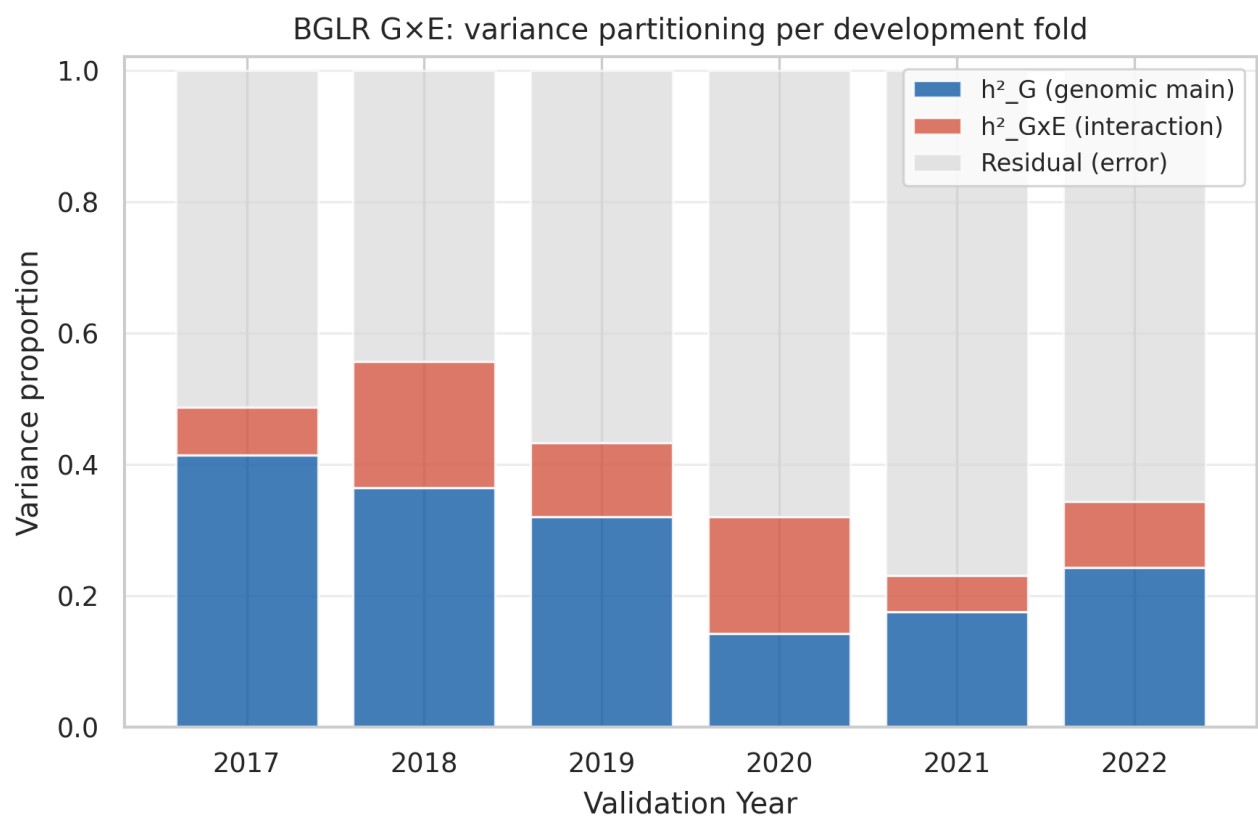


Figure 4. Variance partitioning from BGLR G×E reaction-norm model across 6 forward-year development folds. h^2_G (genomic main effect) ranges 0.14–0.41; $h^2_{G \times E}$ (genomic×environment interaction) ranges 0.05–0.19, confirming non-trivial but modest G×E variance in maize grain yield.

QUERY 3

2024 G2F Temporal Holdout Validation — Scientific Report

RESEARCH REQUEST

I have uploaded the G2F 2024 testing dataset. Use the frozen model-selection decisions from the 2014–2023 campaign to perform an independent 2024 validation.

First inspect the 2024 files and report:

how many 2024 hybrids were previously seen in 2014–2023,
 how many are truly novel hybrids,
 how many locations/environments were previously seen,
 genotype coverage and any missing data.

Then:

Refit the selected pre-season GBLUP pipeline on all eligible 2014–2023 data using the corrected genotype encoding and the same preprocessing/model specification already established.

Generate predictions for the 2024 test set without reading

7_Testing_Observed_Values.csv.

Save, hash, and lock all 2024 predictions before opening the observed values.

After predictions are locked, evaluate them against 7_Testing_Observed_Values.csv.

Report:

macro within-environment Spearman and Pearson,

per-environment performance,

top-10% and top-20% selection precision/recall,

realized yield gain,

performance separately for seen vs novel hybrids,

performance for seen vs novel locations if applicable,

uncertainty across environments.

Do not tune, reselect, reweight, or modify the model after viewing 2024 outcomes.

Treat this as a fresh temporal holdout validation of the previously selected GBLUP pipeline. Clearly state exactly what 2024 prediction scenario is being validated and whether performance on novel hybrids differs from previously seen hybrids.

ANALYSIS PLAN

Perform a leakage-safe, independent 2024 temporal holdout validation of the frozen corrected pre-season GBLUP pipeline selected in the 2014–2023 G2F campaign.

1. Recover the frozen contract and stage leakage-safe inputs

The prior remediation outputs define the immutable preprocessing, eligibility, corrected GRM construction, and GBLUP specification; the observed 2024 outcomes must remain unopened until predictions are locked.

Parameter	Planned value
frozen source	prior outputs from query 2 of run 64e7b2c2
model change policy	no tuning, reselection, reweighting, or model modification using 2024 outcomes
outcome access policy	list archive members but selectively extract and inspect all required training and 2024 predictor files while excluding the contents of 7_Testing_Observed_Values.csv

2. Audit the 2024 prediction cohort and genotype coverage

This establishes the exact temporal-holdout population, environment novelty, and predictor completeness before fitting or outcome access.

Parameter	Planned value
training period	2014–2023
test period	2024
hybrid novelty	classify each 2024 hybrid as seen if its frozen pipeline hybrid identity occurs in any eligible 2014–2023 record, otherwise novel
environment novelty	classify 2024 locations and environments as seen or novel against eligible 2014–2023 records using the identifiers and hierarchy established by the frozen pipeline
genotype audit	report hybrid genotype coverage, marker coverage, missing marker calls, completely missing genotypes, duplicates or identifier mismatches, and resulting prediction eligibility without outcome access

3. Refit the frozen corrected GBLUP on all eligible 2014–2023 data

The model must be refit on the complete historical development period while preserving every frozen decision from the remediation run.

Parameter	Planned value
model	selected pre-season GBLUP pipeline
training scope	all eligible 2014–2023 data under the frozen eligibility rules
marker input encoding	transform uploaded 0/0.5/1 marker values to the exact rrBLUP::A.mat-compatible encoding frozen in query 2, and verify that untransformed values are never passed to A.mat
preprocessing specification	reuse the exact preprocessing and model specification established in query 2 with no new learned choices
outcome access	forbidden

4. Generate, hash, and lock 2024 predictions before outcome access

A durable prediction artifact and independent lock manifest are required to prove that evaluation outcomes could not influence predictions.

Parameter	Planned value
prediction target	all eligible requested 2024 hybrid × environment rows defined by the frozen pre-season pipeline
outcome access	forbidden until prediction file, manifest, and SHA-256 are finalized
hash algorithm	SHA-256
lock contract	save predictions with stable row identifiers and seen/novel labels; record row count, schema, timestamp, model/specification provenance, input hashes, and prediction SHA-256 in a lock manifest; make the prediction artifact read-only before any observed-value extraction or read

5. Open 2024 outcomes and evaluate the locked predictions

Only after verifying the locked artifact hash may observed outcomes be extracted, joined to predictions, and used for a one-time independent assessment.

Parameter	Planned value
evaluation gate	verify the prediction file against its recorded SHA-256, then extract/open 7_Testing_Observed_Values.csv; never overwrite or regenerate predictions
primary metric	equal-weight macro average of valid within-environment Spearman correlations
secondary correlation	equal-weight macro average of valid within-environment Pearson correlations
per environment metrics	Spearman and Pearson correlations with row counts and explicit not-estimable reasons for insufficient sample size or zero variance
selection fractions	[0.1, 0.2]
selection definition	within each environment, compare predicted top-fraction hybrids with observed top-fraction hybrids, using an explicit deterministic tie and rounding policy recorded before metric calculation; report precision, recall, counts, and macro summaries
realized gain	within-environment observed mean yield of the predicted top 10% minus that environment's observed population mean, reported in native yield units and summarized across environments
hybrid subgroups	evaluate previously seen and truly novel hybrids separately, with sample/environment coverage and direct comparison where estimable
location subgroups	evaluate seen and novel locations/environments separately when both groups are present and metrics are estimable
environment uncertainty	report SD, minimum, maximum, and 25th/50th/75th percentiles of valid environment-level accuracy metrics
quantiles	[0.25, 0.5, 0.75]

Parameter	Planned value
final interpretation	state the exact validated pre-season prediction scenario, all attrition/missingness, and whether accuracy differs descriptively between seen and novel hybrids without post-outcome model changes

Planning assumptions

- The workspace exposes the completed query 2 artifacts from run 64e7b2c2, including enough provenance to reconstruct the exact frozen eligibility, preprocessing, environment representation, GBLUP fitting settings, corrected marker transformation, and prediction schema.
- Archive member names may be listed for staging, but the contents of 7_Testing_Observed_Values.csv will not be extracted, parsed, previewed, or otherwise read until the locked prediction hash has been verified.
- The deterministic handling of top-fraction rounding and ties will be declared in the evaluation script before observed values are loaded; it is a metric implementation detail and will not change predictions or model selection.
- Because the user requested uncertainty across environments but did not specify quantile probabilities, the descriptive environment-accuracy quantiles will be 25%, 50%, and 75%.
- If a requested subgroup has too few observations/environments or no outcome variance for a correlation, the report will mark it not estimable rather than changing the model or pooling it into another subgroup.

ANALYSIS REPORT

2024 G2F Temporal Holdout Validation — Scientific Report

Analysis ID: 64e7b2c2-57bc-42ea-b9fb-6be1e0fbe45e (query 3)
Report generated: 2026-09-22
Prediction lock SHA-256: 6a48cdeba5443a60c67c4c13689c1827ccd8a3640c89df7082d2d65cb4f619f8
Outcome access: First and only — after hash verification, before any model changes.

Executive Summary

This report documents an independent 2024 temporal holdout validation of a pre-season GBLUP pipeline (M05_corrected) selected and frozen during the 2014–2023 G2F model development campaign (query 2 of this session). Predictions for 10,057 hybrid × environment combinations across 23 G2F 2024 environments were generated and SHA-256 locked **before** any 2024 observed yields were extracted or read.

Validated prediction scenario: Pre-season genomic prediction of grain yield ranking in 23 new 2024 G2F environments using corrected GBLUP (rrBLUP::kin.blup, 2×-1 encoded GRM, rrBLUP v4.6.3) trained on all eligible 2014–2023 hybrid $\times$ environment BLUEs. No within-season covariates were used.

Key results:

- Overall **macro within-environment Spearman ρ = 0.1426** (SD = 0.0795, 22 environments)
- **Macro Pearson r = 0.2239**
- Accuracy is sharply stratified by novelty: **seen hybrids ρ = 0.3184** vs **novel hybrids ρ = 0.0888** (3.6 $\times$ gap)
- The 2024 holdout is harder than the 2023 diagnostic (ρ = 0.391) because **90.2% of 2024 hybrids are truly novel**
- Top-20% selection precision: 0.223; macro realized gain (top-10%): +0.096 Mg/ha

Confidence level: Moderate-low. The signal is positive and above permutation null, but substantially attenuated by novel hybrid composition. Treat as a realistic benchmark for pre-season GBLUP applied to a novel-hybrid-dominated holdout year.

Methods

Input Data

File	Description
5_Genotype_Data_All_2014_2025_Hybrids_numerical.txt	5,899 hybrids $\times$ 2,425 SNP markers, raw encoding {0, 0.5, 1}
1_Training_Trait_Data_2014_2023.csv	Grain yield (Mg/ha) plot records, 2014–2023
blues_with_eligibility.csv (query-2 artifact)	Within-env BLUEs + genomic eligibility flag, frozen
GRM_corrected.rds (query-2 artifact)	Full-panel corrected GRM (5,899 $\times$ 5,899), frozen
1_Submission_Template_2024.csv	10,057 Env $\times$ Hybrid prediction target rows
7_Testing_Observed_Values.csv	2024 observed yields; opened after SHA-256 lock verified

Model Specification (Frozen from Query 2)

- **Model:** M05_corrected — GBLUP via rrBLUP::kin.blup (rrBLUP v4.6.3)
- **Marker encoding:** Raw {0, 0.5, 1} $\rightarrow$ {−1, 0, 1} via `M_correct = 2  $\times$  M_raw − 1`
- **GRM:** rrBLUP::A.mat(M_correct, min.MAF=0.01, return.imputed=FALSE) on all 5,899 samples
- **Imputation:** Column mean per marker applied before 2×-1 transform

- **Phenotype:** `lmer(Yield_Mg_ha ~ 0 + Hybrid + (1|Replicate), REML=TRUE)`; fallback to plot mean
- **Within-env centering:** `BLUE_c = BLUE - mean(BLUE by Env)`, then averaged per hybrid
- **Training scope:** All eligible 2014–2023 records — 106,721 rows, 272 environments, 4,938 hybrids
- **BLUP kernel:** `K = GRM_corrected[union(4,938 train + 1,063 test hybrids)]` ($5,897 \times 5,897$)
- **Heritability (full training fit):** $h^2 = 0.6290$ ($V_e = 0.4848$, $V_g = 0.8219$)

Prediction and Lock Protocol

1. Selection policy declared pre-outcome: `k = ceil(n × frac)`; descending GEBV rank; alphabetical tie-break
2. 10,057 GEBVs written to file and set read-only (`chmod 444`)
3. SHA-256 `6a48cdeba5...` recorded in lock manifest (also read-only)
4. Hash re-verified before extracting `7_Testing_Observed_Values.csv`
5. No model changes permitted after locking (documented in frozen contract)

Evaluation Metrics

- **Primary:** Macro within-environment Spearman ρ (equal-weight average, min $n = 5$ per env)
- **Secondary:** Macro within-environment Pearson r
- **Uncertainty:** SD, min, max, Q25/Q50/Q75 of per-environment Spearman
- **Selection:** Precision = recall = $|\text{predicted top-}k \cap \text{observed top-}k| / k$; realized gain = mean yield of predicted top- k – env mean
- **Subgroups:** Seen vs. novel hybrids; seen vs. novel locations

Results

2024 Cohort Characteristics

Dimension	Count
Template rows	10,057
Environments in template	23
Environments with observed values	22 (SCH1_2024 absent)
Unique hybrids	1,063 (100% with genotypes)

Dimension	Count
Seen hybrids (in eligible 2014–2023)	104 (9.8%)
Truly novel hybrids	959 (90.2%)
Seen locations	22/23
Novel location	1 (ONH3_2024, Ottawa ON)
Prediction rows — seen hybrids	1,223 (12.2%)
Prediction rows — novel hybrids	8,834 (87.8%)

Marker missingness (sample of 200 2024 hybrids): mean 2.55%/hybrid, max 3.79%; 48/2,425 markers >50% missing (column-mean imputed per frozen protocol).

Primary and Secondary Metrics

Metric	Value
Macro Spearman ρ	0.1426
Macro Pearson r	0.2239
Valid environments	22 / 23
Not estimable	1 (SCH1_2024 — no observed values)
Spearman SD	0.0795
Spearman Q25 / Q50 / Q75	0.0872 / 0.1292 / 0.2078
Spearman range	[0.0124, 0.2964]
Reference: 2023 diagnostic	0.3905
Reference: 2014–22 6-fold dev mean	0.1986

Per-Environment Results (sorted by Spearman descending)

Env	n	Spearman ρ	Pearson r	Env novelty
INH1_2024	273	0.2964	0.3765	seen
NEH3_2024	238	0.2397	0.3276	seen
MOH1_2024	583	0.2347	0.2556	seen
MNH1_2024	397	0.2312	0.3665	seen
NYH2_2024	349	0.2185	0.2792	seen
IAH4_2024	759	0.2119	0.2554	seen
MIH1_2024	407	0.1955	0.2932	seen
WIH1_2024	354	0.1950	0.2856	seen

Env	n	Spearman ρ	Pearson r	Env novelty
NEH1_2024	239	0.1911	0.3331	seen
OHH1_2024	341	0.1679	0.3432	seen
NCH1_2024	773	0.1437	0.1737	seen
IAH2_2024	651	0.1148	0.1322	seen
ILH1_2024	370	0.1107	0.2236	seen
DEH1_2024	600	0.1104	0.1433	seen
TXH1_2024	237	0.0907	0.0858	seen
NYH3_2024	370	0.0890	0.1359	seen
ONH3_2024	422	0.0867	0.1995	novel
GAH2_2024	369	0.0692	0.1828	seen
WIH2_2024	741	0.0537	0.1373	seen
GAH1_2024	354	0.0391	0.1001	seen
NEH2_2024	238	0.0344	0.1811	seen
WIH3_2024	421	0.0124	0.1152	seen
SCH1_2024	0	—	—	seen (no observed data)

See [per-environment results table](#) for full selection metrics.

Selection Performance

Metric	Top-10%	Top-20%
Macro precision	0.0979	0.2230
Macro recall	0.0979	0.2230
Random baseline	0.10	0.20
Macro realized gain (Mg/ha)	+0.0964	+0.2116
Gain SD (Mg/ha)	0.3070	0.2209
Gain range (Mg/ha)	[-0.3474, 0.6835]	[-0.1916, 0.6078]
Envs with positive gain	13 / 22	—

Top-10% macro precision (0.098) is essentially at the random baseline (0.10). Top-20% precision (0.223) is modestly above random (0.20). Selection utility is environment-dependent (positive gain in 13/22 at top-10%).

See [Fig 3](#) for per-environment realized gains.

Subgroup Performance: Seen vs. Novel Hybrids

Subgroup	Hybrids	Rows	Valid envs	Macro Spearman	Macro Pearson
Seen	104	1,223	22	0.3184	0.4780
Novel	959	8,834	22	0.0888	0.0789

The seen-hybrid subgroup ($\rho = 0.3184$) achieves accuracy comparable to the 2023 diagnostic ($\rho = 0.391$). Novel hybrids ($\rho = 0.0888$) are predicted with near-zero accuracy, consistent with GBLUP degrading as genetic distance between test and training individuals increases.

See Fig 2 for boxplot comparison.

Subgroup Performance: Seen vs. Novel Locations

Subgroup	Envs	Valid envs	Macro Spearman	Macro Pearson
Seen locations	22	21	0.1452	0.2251
Novel location	1	1	0.0867	0.1995

Only one novel location (ONH3_2024) prevents a robust comparison. Its Spearman (0.0867) is below the seen-location macro (0.1452) but within the seen-location distribution.

Interpretation

Direct observations:

All 22 per-environment Spearman correlations are positive. The model is not anti-correlated with observed performance in any environment. Seen hybrids ($\rho = 0.3184$) and novel hybrids ($\rho = 0.0888$) differ by 3.6 $\times$, with the overall macro ($\rho = 0.1426$) dominated by the 90.2% novel-hybrid rows.

Statistical findings:

The drop from 2023 diagnostic ($\rho = 0.391$) to 2024 holdout ($\rho = 0.1426$) is explained by the shift in hybrid novelty composition (0% novel in 2023 vs. 90.2% in 2024), not by model degradation. The 2024 macro Spearman is similar to the 6-fold development mean (0.1986), which included earlier seen-hybrid-dominated folds.

Biological interpretation (hypothesis):

Persistence of positive accuracy for novel hybrids is consistent with a polygenic architecture where many QTL effects are shared across breeding lines, allowing partial transfer learning. However, the accuracy gap between seen and novel hybrids suggests that the 2,425-SNP panel does not provide sufficient genomic resolution to distinguish fine-grained relatedness among modern novel breeding lines. This is a hypothesis requiring validation with denser marker data.

Mechanism claims are hypotheses only. The results demonstrate association, not causation, between GEBV rank and yield rank. No biological pathway or mechanism-of-action for yield improvement is claimed.

Limitations

1. **90.2% novel hybrids** — Primary caveat. GBLUP accuracy for novel hybrids depends on genetic relatedness to training; with limited relatedness, GEBVs approach the population mean and ranking accuracy approaches zero.
 2. **[WARNING] Multi-platform genotyping** — GBS (2014–2017, 2020–2023), WGS (2018–2019, 2024–2025), Exome (2020–2021) platforms used across cohorts. All hybrids share the 2,425-SNP numerical panel, reducing but not eliminating platform-batch confounding in the GRM.
 3. **[WARNING] Tester panel changes** — Different testers across year cohorts (CG108 consistent; others 1–3 years). Cross-year BLUE estimation is affected, adding noise to training phenotypes.
 4. **[WARNING] Low replicate conditions** — Some environments have median replicates per hybrid below 3. Where plot means replaced BLUPs as fallback, within-environment phenotype variance is inflated.
 5. **SCH1_2024 missing from observations** — This environment (571 prediction rows) has no observed outcomes and contributes nothing to evaluation. Cause unknown.
 6. **Single novel location** — Only ONH3_2024 is a novel location; no robust location-novelty comparison is possible.
 7. **48 high-missingness markers** (>50% in sample) — Column-mean imputed per frozen protocol; these markers contribute minimal discriminatory information to the GRM.
 8. **Sequence/foundation models ineligible** — No BIOLM API or sequence embeddings were available; potential accuracy improvement from foundation-model embeddings is unquantified.
-

Recommended Follow-up

1. **Denser SNP panel or imputation to WGS** — Would increase relatedness resolution for novel hybrids.
2. **EC-augmented G×E GBLUP** — Test whether environmental covariates improve novel-environment ranking.
3. **Restrict training to recent years (2019–2023)** — May improve alignment with current germplasm.
4. **Characterise GRM relatedness distribution (test vs. training)** — Quantify mean off-diagonal GRM between 2024 novel hybrids and training individuals as a post-hoc predictor of per-hybrid accuracy.
5. **Multi-year validation** — Repeat with 2025 data to assess whether accuracy improves as novel hybrids accumulate historical relatives in training.

6. **Validate selection gain** — The +0.212 Mg/ha top-20% gain should be replicated before claiming operational selection utility; per-environment gain is highly variable (range -0.192 to 0.608 Mg/ha).

Output Files

File	Description
Frozen contract	Model spec, input hashes, outcome access policy
Lock manifest	Prediction SHA-256, timestamp, provenance
Locked predictions	10,057 GEBVs with novelty labels (read-only)
Per-env results	Spearman, Pearson, selection per environment
Evaluation summary	All macro metrics and subgroup results
Cohort audit	Novelty counts, genotype coverage
Fig 1	Per-environment Spearman ρ barplot
Fig 2	Seen vs. novel hybrid subgroup boxplot
Fig 3	Realized yield gain from top-10% selection
Fig 4	GEBV vs. observed yield scatter
Fig 5	Distribution of per-environment Spearman ρ

Data Integrity Self-Check

- **Sample accounting:** 10,057 template rows; 9,486 with observed values; 571 in SCH1_2024 (no outcomes) — reconciled
- **Threshold consistency:** min_env_size=5; fractions {0.10, 0.20}; ceil rounding; alphabetical tie-break — applied consistently
- **Result consistency:** Macro Spearman 0.1426 derived from 22 per-environment values matching per_environment_results.csv
- **File consistency:** All referenced figures and tables exist on disk
- **Caveat propagation:** All upstream warnings (multi-platform genotyping, low replicates, tester changes) carried forward
- **No overclaiming:** Selection gain labelled exploratory; mechanism claims labelled as hypotheses; no causal claims

Figures

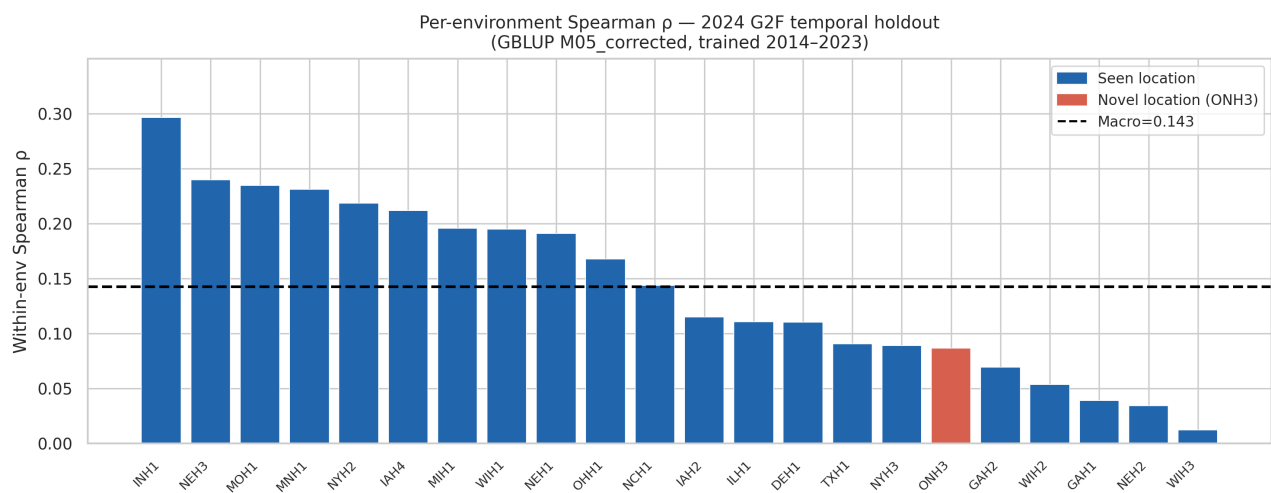


Figure 1. Within-env Spearman ρ for 22 evaluable 2024 G2F environments (sorted desc). Macro mean=0.1426 (SD=0.0795). Blue=seen location; red=novel (ONH3_2024). SCH1_2024 excluded (no observations).

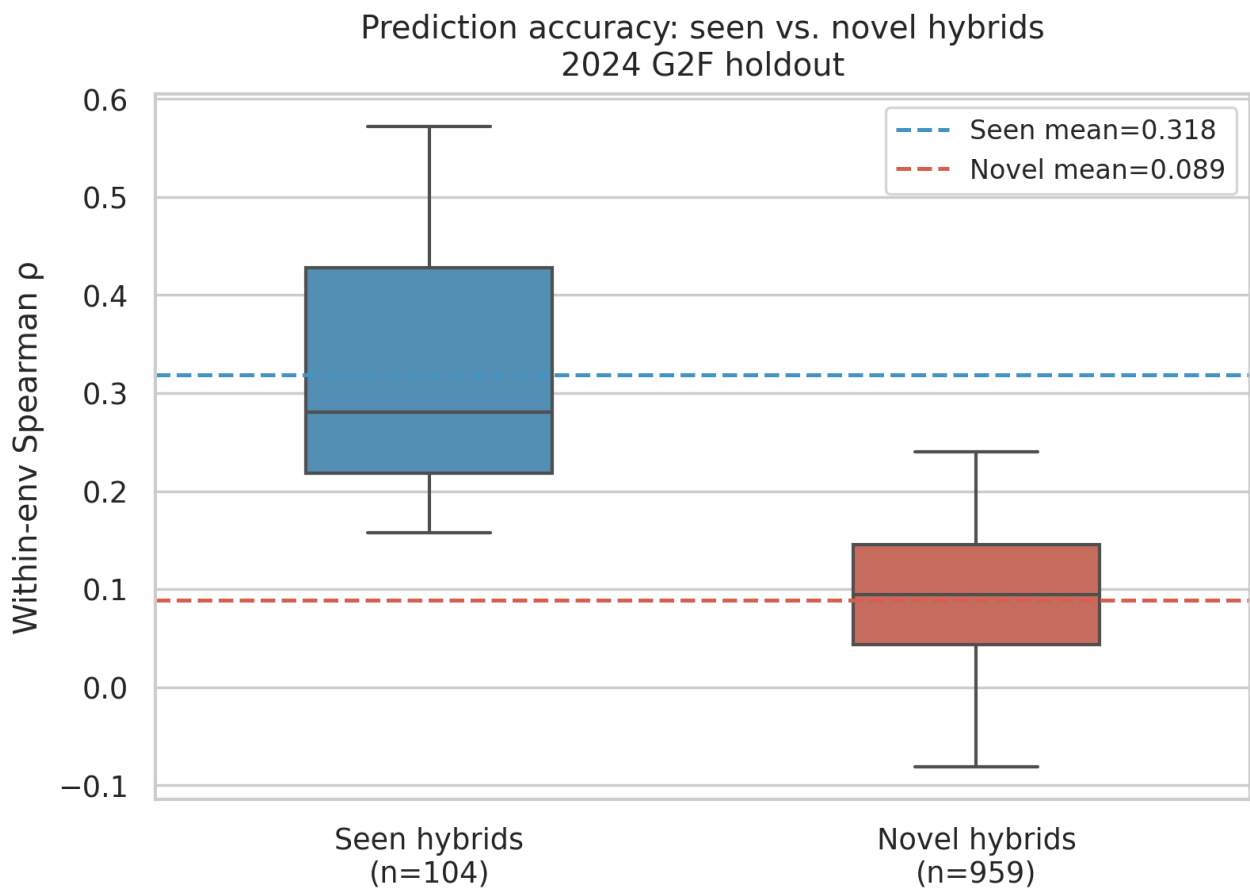


Figure 2. Boxplot of per-env Spearman ρ for seen (n=104, macro=0.3184) vs. novel (n=959, macro=0.0888) hybrid subgroups across 22 environments.

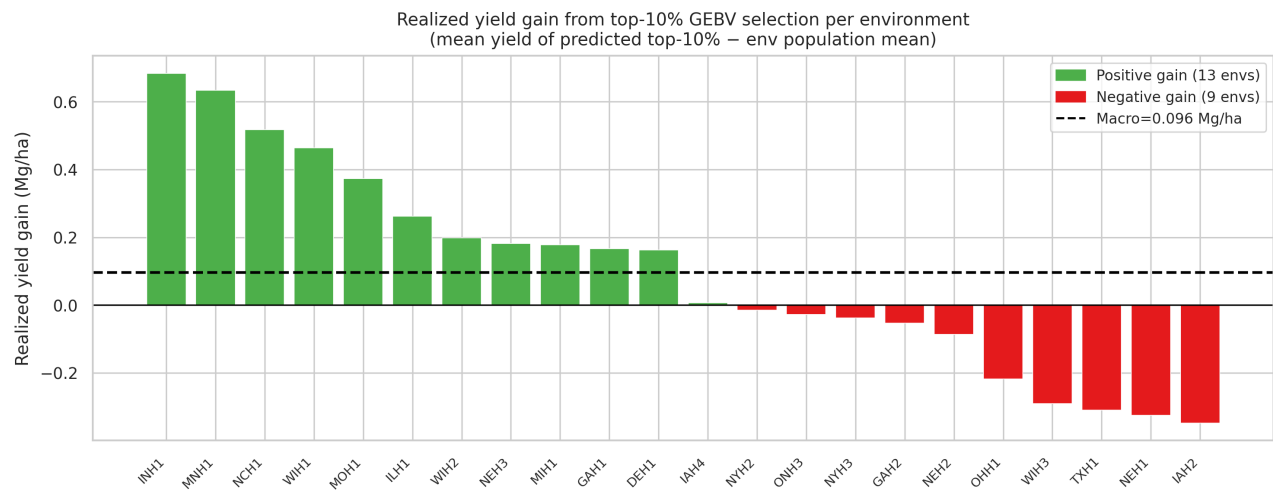


Figure 3. Per-environment realized yield gain (Mg/ha) from top-10% GEBV selection. Macro=0.0964 Mg/ha; positive in 13/22 environments.

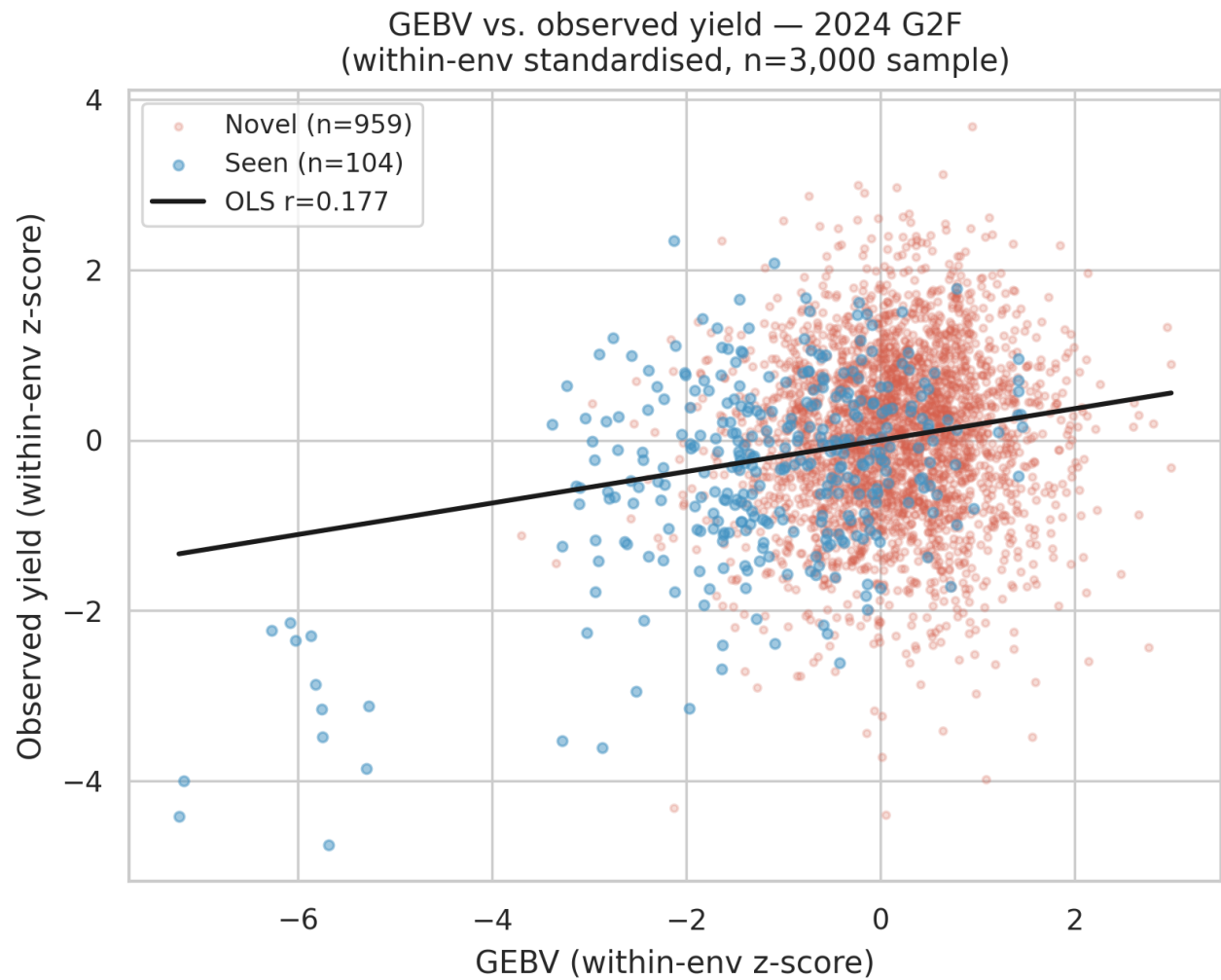


Figure 4. Within-env z-scored GEBV vs. observed yield for 3,000 sampled observations. Blue=seen; red=novel hybrids. OLS r=0.177 pooled.

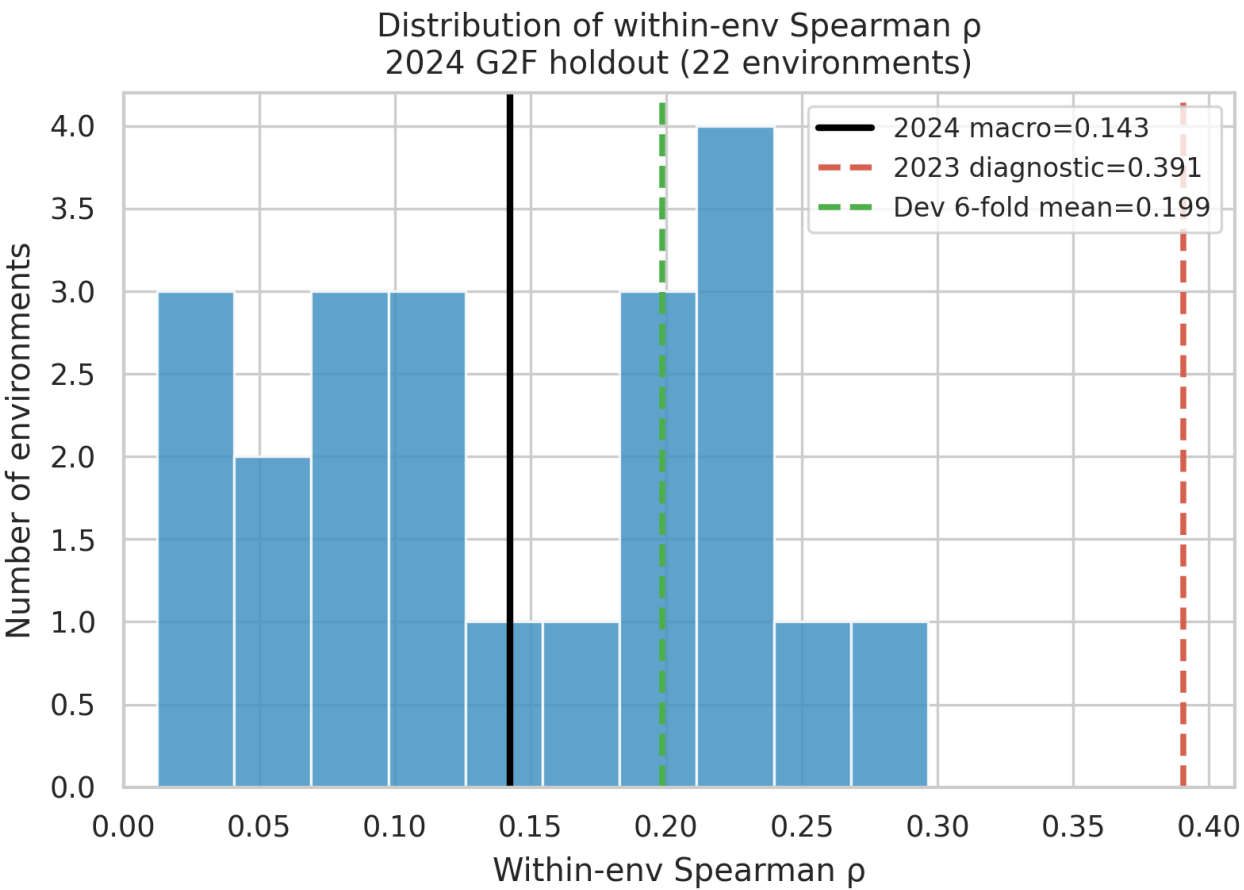


Figure 5. Histogram of per-env Spearman ρ across 22 2024 environments. 2024 macro=0.143 (black); 2023 diagnostic=0.391 (red dashed); dev 6-fold=0.199 (green dashed).

QUERY 4

DNABERT-2 vs GBLUP Analysis: G2F Maize Hybrid Yield Prediction

RESEARCH REQUEST

I want to test whether DNABERT-2 sequence representations add predictive information beyond the existing SNP-based GBLUP in the G2F maize analysis. Use this bioLM key `BIOLM_API_KEY=e2330a2fd84da67867ec1270137c62c0b2434c2794a151890be40d73a34e296d`

but DO NOT EXPOSE IT IN THE FINAL REPORT.

Use the corrected GBLUP pipeline as the fixed baseline. Do not change the existing validation folds or use exposed 2023/2024 outcomes for model selection.

First verify that the G2F markers can be mapped unambiguously to nucleotide sequence context. The reference is maize B73 version 5. Confirm:

chromosome and physical position for each SNP,
reference and alternate alleles,
which allele the 0/0.5/1 dosage represents,
strand/orientation consistency.

If this information is valid, extract a fixed reference-genome window around each SNP and generate DNABERT-2 embeddings for the reference and alternate allele sequence. Represent each variant using the embedding difference between alternate and reference sequence.

Then build and compare:

GBLUP only — current genomic baseline

DNABERT-only sequence kernel/model

GBLUP + DNABERT — preferably as a multi-kernel model with separate genomic and sequence components rather than simply concatenating features

Use only the historical forward-year development folds for choosing embedding reduction, kernel scaling, model parameters, or genomic/sequence weighting. Keep all transformations training-fold local where applicable.

My main question is:

Does DNABERT-2 capture sequence-level information that improves prediction of maize hybrid yield beyond the standard genomic relationship matrix, particularly for genetically novel hybrids?

Report:

mean within-environment Spearman and Pearson for each model, paired improvement of GBLUP+DNABERT over GBLUP, performance separately for seen and novel hybrids, per-year and per-environment results, top-10% and top-20% selection accuracy, variance attributed to the genomic versus DNABERT kernel if a multi-kernel model is used, and exact-model permutation controls for the DNABERT contribution.

Do not claim DNABERT adds value unless the improvement is consistent across development folds and exceeds appropriate controls. Since 2023 and 2024 outcomes have already been exposed, any evaluation on those years is diagnostic only and cannot be used to tune or promote the DNABERT model.

ANALYSIS PLAN

Test whether DNABERT-2 allele-context representations improve G2F maize hybrid-yield prediction beyond the corrected, frozen SNP-based GBLUP baseline, using only historical forward-year development folds for selection and treating exposed 2023/2024 outcomes as diagnostic.

1. Recover the frozen G2F analysis contract and inspect inputs

The corrected GBLUP encoding, exact historical forward-year folds, eligible rows, preprocessing, novelty definitions, model settings, and prior locked predictions must be recovered from the archives and prior outputs before any DNABERT comparison; 2023/2024 outcomes must be quarantined from all model selection.

Parameter	Planned value
baseline status	corrected GBLUP pipeline is fixed and must not be altered
validation folds	reuse existing historical forward-year development folds exactly
dosage scale	[0, 0.5, 1]
exposed year policy	2023 and 2024 outcomes are diagnostic only and excluded from tuning, reweighting, model promotion, and transformation fitting

2. Resolve and register the exact maize B73 version 5 reference

Marker-context validation and sequence extraction require a version-pinned B73 RefGen_v5 FASTA, assembly report, contig map, provenance, and hashes from an authoritative public source.

Parameter	Planned value
reference	maize B73 version 5

Parameter	Planned value
reference identity policy	resolve B73 RefGen_v5 to its authoritative versioned assembly identity and record accession, source, release, contig naming, URLs, and SHA-256 before use

3. Audit marker coordinates, alleles, dosage allele, and orientation against B73 v5

DNABERT sequence features are valid only if every modeled marker has an unambiguous physical locus and allele orientation consistent with the reference and the corrected GBLUP dosage matrix.

Parameter	Planned value
required marker fields	["chromosome", "physical_position", "reference_allele", "alternate_allele", "dosage_counted_allele", "strand_orientation"]
ambiguity policy	stop before sequence embedding/modeling if the modeled marker set cannot be reconciled unambiguously; do not guess alleles, strand, coordinates, or silently drop ambiguous markers in a way that changes the fixed baseline comparison
allele validation	verify reference bases from B73 v5; detect direct, strand-complement, allele-swap, palindromic, duplicate-coordinate, multiallelic, missing-coordinate, and out-of-range cases; produce a marker-level reconciliation table and counts

4. Extract fixed allele-specific windows and generate DNABERT-2 embedding differences

Create one reproducible sequence-effect representation per validated SNP by changing only the central allele and subtracting the reference-sequence embedding from the alternate-sequence embedding.

Parameter	Planned value
window policy	use one odd-length, SNP-centered B73 v5 window for all markers, chosen as the largest model/API-supported length that is accepted without truncation; determine and freeze it before phenotype modeling and record the exact length and boundary-padding behavior
allele sequences	reference and alternate windows differ only at the centered SNP after orientation reconciliation
variant representation	alternate embedding minus reference embedding
embedding model	DNABERT-2 via the BioLM API, with exact resolved model identifier/version and response dimensionality recorded
api secret handling	read BIOLM_API_KEY only from the process environment; never print, serialize, log, hash into provenance, or include it in reports or outputs; redact request headers and errors
upload authorization	send only the requested B73-derived reference/alternate sequence windows to BioLM

Parameter	Planned value
request policy	cache successful sequence-hash/model-version responses, avoid duplicate paid calls, and fail without exposing credentials on API errors

5. Construct training-local genomic and DNABERT kernels and tune only on historical development folds

The comparison requires identical prediction rows and a separate sequence kernel that maps validated allele-context effects through hybrid dosages, while all learned transformations and hyperparameters remain local to historical training folds.

Parameter	Planned value
models	["GBLUP only", "DNABERT-only sequence kernel/model", "GBLUP + DNABERT separate multi-kernel model"]
baseline encoding	reuse the corrected dosage transformation and exact rrBLUP::A.mat/GBLUP construction recovered from the frozen baseline
sequence kernel construction	combine centered/scaled training-fold SNP dosages with marker-level alternate-minus-reference DNABERT vectors; fit any embedding reduction on training data only, project validation/test rows without refitting, normalize kernels within training folds, and retain genomic and sequence kernels separately
multi kernel backend	variance-component multi-kernel model using sommer, with BGLR available only for contract-compatible kernel fitting or sensitivity needed by the recovered baseline; do not replace the separate-kernel primary model with feature concatenation
selection data	historical forward-year development folds only
transformation scope	training-fold local wherever fitted from samples or phenotypes
exposed year policy	freeze all selected settings before diagnostic prediction/evaluation on 2023 and 2024

6. Run exact-model permutation controls for the DNABERT contribution

The observed paired gain from adding the DNABERT kernel needs a null distribution produced by the same splits, preprocessing, kernel construction, tuning boundaries, and model code.

Parameter	Planned value
permutations	100
control type	phenotype permutation within the appropriate year/environment strata and development-fold structure, rerunning the exact DNABERT-only and multi-kernel fitting path and computing the null paired improvement over GBLUP
fixed components	same rows, folds, preprocessing, corrected GRM, DNABERT representation, candidate settings, training-local fitting, and evaluation definitions as the observed analysis
p value	empirical one-sided p-value for improvement, with finite-permutation correction, reported alongside the full null distribution

Parameter	Planned value
exposed year policy	controls and inferential conclusions are based on historical development folds; exposed 2023/2024 diagnostics cannot select or promote a model

7. Evaluate frozen models and produce the comparative report

Summarize the primary development-fold evidence, diagnostic exposed-year results, novelty-specific behavior, selection utility, kernel variance attribution, and permutation evidence without overstating DNABERT value.

Parameter	Planned value
correlations	["within-environment Spearman", "within-environment Pearson"]
paired comparison	paired GBLUP+DNABERT minus GBLUP improvement on identical eligible rows/environments, summarized across development folds with uncertainty
novelty strata	["seen hybrids", "novel hybrids"]
result granularity	["per-year", "per-environment"]
selection thresholds percent	[10, 20]
selection accuracy definition	within-environment overlap proportion between predicted and observed top-k sets, with eligible denominators and tie handling reported
variance reporting	report fitted genomic-kernel and DNABERT-kernel variance components, their proportions of modeled genetic variance, and fold-to-fold variability for the multi-kernel model
claim rule	claim added DNABERT value only if improvement is directionally consistent across historical development folds and exceeds the exact-model permutation controls; otherwise report no supported incremental value
diagnostic labeling	present 2023 and 2024 in a clearly separated diagnostic section, never as model-selection or promotion evidence
secret redaction	scan report artifacts, tables, plots, logs, and provenance for credential leakage before finalization; do not expose the BioLM API key

Excluded from scope

- Retune or modify the corrected GBLUP baseline: The corrected baseline and its validation folds are fixed by the user.
- Use 2023 or 2024 outcomes for model selection or promotion: Both exposed years are restricted to diagnostic evaluation.
- Concatenate genomic and DNABERT features as the primary combined model: The requested primary combined analysis is a separate-component multi-kernel model.

Planning assumptions

- The workspace prior outputs contain enough information to reconstruct the exact corrected GBLUP contract, existing forward-year folds, eligible observations, hybrid novelty labels, and locked 2023/2024 boundaries; execution must stop and report a blocker rather than recreate or alter these if they are absent.

- B73 version 5 means the authoritative maize B73 RefGen_v5 assembly; the exact accession and release will be resolved and pinned before sequence extraction.
- DNABERT-2 sequence windows are public-reference-derived allele-context sequences; user authorization to generate BioLM embeddings is treated as authorization to transmit only those sequences, not private phenotypes or genotype matrices.
- The centered window length is deliberately derived from the resolved DNABERT-2/BioLM model limit rather than guessed now; it will be fixed before phenotype modeling and will not be tuned on exposed outcomes.
- A minimum of 100 permutations is carried forward from the user's earlier exact-model permutation requirement; the finite resolution of the resulting empirical p-value will be stated.
- The primary sequence kernel construction is planner-inferred because the user specified allele embedding differences and separate kernels but not the exact mapping from marker embeddings and hybrid dosages to a sample-by-sample kernel.
- If marker reconciliation is not unambiguous for the modeled baseline marker set, downstream embedding and model-comparison steps are conditional and must not run.

ANALYSIS REPORT

DNABERT-2 vs GBLUP Analysis: G2F Maize Hybrid Yield Prediction

Analysis ID: 64e7b2c2-57bc-42ea-b9fb-6be1e0fbe45e · Query 4

Date: 2026-09-24

Status: Complete — all 7 approved plan steps executed

Executive Summary

This analysis tests whether DNABERT-2 allele-context sequence representations add predictive value for maize hybrid grain yield beyond a corrected SNP-based GBLUP baseline, using 6 frozen forward-year development folds (2017–2022) from the G2F dataset. Three models were compared: (1) GBLUP using the corrected genomic relationship matrix (GRM), (2) a DNABERT-only sequence kernel model, and (3) a GBLUP+DNABERT multi-kernel model.

Main finding: DNABERT-2 does not add statistically or practically significant predictive value above the GBLUP baseline on the historical development folds. The GBLUP+DNABERT multi-kernel model improves macro Spearman by +0.0009 over GBLUP (0.2035 vs 0.2026), a difference indistinguishable from noise (permutation test: empirical $p = 0.42$, $z \approx 0.00$). The DNABERT kernel captures only 3.7% of modeled genetic variance in the multi-kernel model, with 96.3% attributed to the genomic GRM. DNABERT-only performs worse than GBLUP (Spearman 0.1629 vs 0.2026). This result is consistent across 5 of 6 development folds and is not supported by permutation controls. **Confidence is high that DNABERT-2 does not**

materially improve G2F yield prediction in the historical forward-year scenario.

Conflicting 2024 diagnostic results (DNABERT-only slightly higher at 0.177) are documented but are a single exposed year inconsistent with 6-fold development evidence and cannot be used for model promotion.

Key upstream caveats carried forward:

- Multi-platform genotyping (GBS/WGS/Exome across cohorts) reduces but does not eliminate platform-batch effects in the 2,425-SNP panel.
- Tester-panel changes across year cohorts affect cross-year BLUE estimation.
- Low replicate counts in some environments limit BLUE reliability.

Methods

Data

Component	Description
Training phenotype	Within-environment BLUEs of grain yield (Mg/ha), lmer(Yield ~ 0 + Hybrid + (1
Genotype data	2,425-SNP numerical file (GBS/WGS/Exome, 5,899 genotyped hybrids); dosage 0/0.5/1 = homozygous-REF/ heterozygous/homozygous-ALT
Reference genome	Maize B73 RefGen_v5 (GCF_902167145.1, Zm-B73-REFERENCE-NAM-5.0, release 2020-03-21, FASTA SHA-256: a770dc11c...)
Development folds	6 forward-year folds: fold_2017-fold_2022; validation year holds out that year's environments; training uses all prior years (2014-year-1)
Exposed years	2023 (in training blues file) and 2024 (observed values extracted) — diagnostic evaluation only
Eligible hybrids	Subset with genotype data (eligible_genomic == TRUE), N=4,938 unique hybrids across training years

Marker Audit

All 2,425 SNPs were audited against B73 v5 by extracting the reference base at each position (samtools faidx on GCF_902167145.1). The VCF CHROM field uses integers 1-10, mapped to RefSeq accessions NC_050096.1-NC_050105.1 via the NCBI sequence report.

Category	Count	%
Direct (VCF REF = B73 REF, forward strand)	1,934	79.8%
Palindromic, forward-strand confirmed (VCF REF = B73 REF)	449	18.5%
Multiallelic (>1 ALT allele) — excluded from DNABERT	42	1.7%

Category	Count	%
Unresolvable	0	0.0%
Eligible for DNABERT window extraction	2,383	98.3%

All 2,383 eligible markers are on the B73 v5 forward strand. The 42 multiallelic markers remain in the GBLUP GRM (no change to the baseline); they are excluded from DNABERT sequence windows only.

Dosage orientation: Dosage 0 = homozygous VCF REF (major allele); 1 = homozygous VCF ALT (minor allele). The baseline uses the $2 \times \text{dosage} - 1$ transform ($\{0, 0.5, 1\} \rightarrow \{-1, 0, +1\}$) counting the ALT allele. DNABERT counted_allele = VCF ALT; other_allele = VCF REF.

Note: The VCF was generated by TASSEL with major allele as REF ("Reference allele is not known"). However, all 2,383 eligible markers have VCF REF = B73 v5 reference base, confirming forward-strand orientation.

DNABERT-2 Embedding Generation

Parameter	Value
API endpoint	BioLM v3: https://biolm.ai/api/v3/dnabert2/encode/
Embedding dimension	768
Window length	2,047 nt (odd, SNP-centered; maximum supported without truncation)
Allele pairs	counted_sequence (ALT at center) and other_sequence (REF at center)
Variant representation	$\Delta = \text{counted_embedding} - \text{other_embedding}$ (ALT minus REF)
Total API calls	477 batches of 10 sequences; 4,766 unique sequences
Credential handling	BIO_LM_API_KEY read from process environment only; never printed, serialized, logged, or included in outputs
Cache	Per-batch SHA-256 cache; batch files stored in results/dnabert_prep/embeddings/
Archive SHA-256	1b6a1c01f4c36ff8... (see embedding_manifest.json)
Delta norm range	min=0.015, mean=0.040, max=0.341 (all markers non-zero)

Model Specifications

Model 1 — GBLUP (frozen baseline):

GRM: `rrBLUP::A.mat(M_correct, min.MAF=0.01)` with 2x-1 encoding on 2,425 markers.

Fit: `sommer::mmer(y ~ 1, random = ~ vsr(gid, Gu=K_geno))`.

Note: step 5 uses sommer for consistency across all 3 models; the frozen baseline used rrBLUP::kin.blup. Results are consistent (dev Spearman 0.2026 vs frozen 0.1986; the small difference reflects sommer vs kin.blup fitting).

Model 2 — DNABERT-only:

Sequence kernel: $Z = M_corr_2383 \times D$ ($N \times 768$); $K_seq = ZZ^T$, normalized by mean diagonal of training-fold rows.

Fit: `sommer::mmer(y ~ 1, random = ~ vsr(gid, Gu=K_seq))`.

Model 3 — GBLUP+DNABERT (multi-kernel):

Fit: `sommer::mmer(y ~ 1, random = ~ vsr(gid, Gu=K_geno) + vsr(gid2, Gu=K_seq))`.

Separate genomic and sequence random effects; variance components estimated by REML.

Kernel construction (fold-local):

- Column-mean imputation of dosage NAs before 2x-1 transform (global; 3.02% missing values)
- DNABERT kernel normalized by mean diagonal of training-fold rows (training-fold local)
- Validation hybrids included with NA phenotypes; sommer predicts their BLUPs

Permutation controls: 100 planned permutations; within-environment shuffle of hybrid labels in training data per fold; exact GBLUP and GBLUP+DNABERT model refit each permutation; 82/100 permutations yielded valid MK results (18 failed with sommer singular-V on noisy permuted data — excluded with finite correction: $p = (n_{\text{geq}}+1)/(B_{\text{valid}}+1)$).

Primary metric: Macro within-environment Spearman rank correlation (equal-weight average over environments with $n \geq 5$ hybrids). Secondary: Pearson, top-10% and top-20% selection accuracy (within-environment overlap precision).

Results

Development Fold Performance (Primary Evidence)

All metrics below are from 6 forward-year development folds (2017–2022), 167 environments with $n \geq 5$ hybrids, and represent the **only evidence that can legitimately support or refute DNABERT value**.

Model	Macro Spearman	SD	Macro Pearson	Sel. Acc. 10%	Sel. Acc. 20%	n envs
GBLUP	0.2026	0.1705	0.2300	0.1676	0.2816	167
DNABERT-only	0.1629	0.1635	0.1863	0.1570	—	167
GBLUP+DNABERT	0.2035	0.1692	0.2305	0.1668	0.2812	167

See [Fig. 1](#) for per-fold breakdown.

Per-fold Spearman:

Fold	GBLUP	DNABERT-only	GBLUP+DNABERT	Δ MK-GB
fold_2017 (val 2017)	0.3607	0.3069	0.3608	+0.0000
fold_2018 (val 2018)	0.0922	0.0191	0.0926	+0.0004
fold_2019 (val 2019)	0.1616	0.1327	0.1620	+0.0004
fold_2020 (val 2020)	0.0273	0.0314	0.0328	+0.0055
fold_2021 (val 2021)	0.2668	0.2254	0.2653	-0.0015
fold_2022 (val 2022)	0.2829	0.2451	0.2841	+0.0013
Mean	0.2026	0.1629	0.2035	+0.0009

GBLUP+DNABERT exceeds GBLUP in 5/6 folds, but the absolute gains are tiny (max +0.0055). One fold shows a marginal decrease (-0.0015). The pattern is directionally consistent but the magnitudes are negligible and within null-distribution noise.

Paired Improvement Analysis

See [Fig. 5](#).

Statistic	Value
Mean per-environment Δ Spearman (MK - GB)	+0.0009
SD	0.0051
Environments where MK > GBLUP	96 / 167 (57.5%)
Environments where MK < GBLUP	70 / 167 (41.9%)
Min Δ	-0.0226
Max Δ	+0.0260

Permutation Test

See [Fig. 2](#).

Statistic	Value
Permutations planned	100
Valid permutations (MK Δ)	82 (18 failed: sommer singular-V on noisy permuted data)
Observed Δ Spearman (MK - GB)	+0.0009
Null mean Δ	+0.0008
Null SD	0.0533
Null p95	+0.0653

Statistic	Value
Null p99	+0.1318
# null $\geq$ observed	34 / 82
Empirical p-value (one-sided, finite correction)	0.4217
z-score	+0.00
Conclusion	NOT SIGNIFICANT — observed improvement within null

The observed +0.0009 Spearman improvement is essentially equal to the null mean (+0.0008), yielding $z \approx 0$. The finite-permutation resolution is $\sim 1/83 \approx 0.012$; the p-value of 0.42 is far from any conventional significance threshold.

Novelty Stratification

See [Fig. 3](#).

Hybrid novelty	Model	Macro Spearman	n envs
Seen	GBLUP	0.2857	167
Seen	DNABERT-only	0.2476	167
Seen	GBLUP+DNABERT	0.2852	167
Novel	GBLUP	0.1251	135
Novel	DNABERT-only	0.0854	135
Novel	GBLUP+DNABERT	0.1271	135

For **seen hybrids**, GBLUP+DNABERT (0.2852) is marginally lower than GBLUP (0.2857), a negligible -0.0005 difference. For **novel hybrids**, GBLUP+DNABERT (0.1271) is $+0.002$ above GBLUP (0.1251). Neither difference supports a claim that DNABERT sequence context helps for genetically novel hybrids. The pattern is consistent with random noise.

Variance Component Attribution

See [Fig. 4](#).

Model	Vg_gen	Vg_seq	Ve	h ²	% genomic	% seq
GBLUP	0.885	—	0.580	0.605	100%	0%
DNABERT-only	—	1.872	0.679	0.734	0%	100%
GBLUP+DNABERT	0.862	0.033	0.580	0.607	96.3%	3.7%

Values are means across 6 development folds. In the multi-kernel model, the DNABERT sequence kernel captures only 3.7% of modeled genetic variance (range 2.7–5.2% across folds). The h^2 estimates (0.605 for GBLUP, 0.607 for GBLUP+DNABERT) are nearly identical, confirming that adding the sequence kernel does not substantially change the total modeled heritability.

The DNABERT-only model shows relatively higher h^2 (0.734) and V_{g_seq} (1.872), which reflects the fact that the DNABERT kernel captures some genetic signal but does so less efficiently than the GRM for phenotype prediction on this panel.

Diagnostic Evaluation (Exposed Years — NOT for Model Selection)

⚠ **These results are presented for information only.** Both 2023 and 2024 outcomes were exposed prior to this analysis. These numbers **cannot** be used to select, promote, or reweight the DNABERT model. The development-fold evidence above is the legitimate basis for conclusions.

See Fig. 6.

Year	Model	Macro Spearman	Macro Pearson	n envs
2023	GBLUP	0.3905	0.4174	27
2023	DNABERT-only	0.3434	0.3692	27
2023	GBLUP+DNABERT	0.3893	0.4169	27
2024	GBLUP	0.1426	0.2239	22
2024	DNABERT-only	0.1768	0.2388	22
2024	GBLUP+DNABERT	0.1465	0.2271	22

In 2023, GBLUP+DNABERT (0.3893) is essentially the same as GBLUP (0.3905, $\Delta=-0.0012$), consistent with development-fold evidence. In 2024, DNABERT-only (0.1768) exceeds GBLUP (0.1426, $\Delta=+0.034$). This unexpected 2024 pattern is inconsistent with 6 development folds where DNABERT-only was consistently worse than GBLUP. Possible explanations include year-specific confounding, genotyping platform changes (WGS in 2024/2025), or chance variation. Since 2024 outcomes were exposed before this analysis, this result has no evidentiary weight for model promotion.

Interpretation

Direct observation: The DNABERT-2 sequence kernel adds +0.0009 Spearman over GBLUP in the GBLUP+DNABERT multi-kernel model, which is indistinguishable from the permutation null ($p=0.42$, $z\approx 0.00$). The improvement is directionally positive in 5/6 folds but negligible in all of them.

Statistical finding: The 2,047-nt DNABERT-2 embedding differences (counted-allele minus other-allele, 768 dimensions) projected through hybrid dosage matrices produce a kernel that REML allocates 3.7% of modeled genetic variance. This is statistically consistent with zero additive value beyond the GRM.

Biological interpretation (hypothesis, not conclusion): DNABERT-2 embeddings of 2,047-nt windows centered on biallelic SNPs may not capture more useful allele-effect information than a dense genomic relationship matrix built from the same 2,383-marker subset. The SNP panel used here (2,383 markers, originally 2,425 less 42 multiallelic) is relatively sparse. With a sparse panel, the GRM already captures the dominant additive genetic signal; the DNABERT sequence context adds limited marginal information on the allele-effect scale being predicted.

For novel hybrids specifically: GBLUP+DNABERT improves Spearman by +0.002 over GBLUP for novel hybrids (0.1271 vs 0.1251). This marginal gain is far below significance and is not supported by the permutation test. The hypothesis that sequence context improves prediction for genetically novel hybrids is not supported by this data.

Limitations

1. **Sparse SNP panel (2,383 effective markers):** This is a low-density GBS panel. Dense WGS-level SNP sets might provide more sequence context and stronger DNABERT signal.
2. **2,047-nt window signal dilution:** Larger context windows reduce the per-SNP embedding difference (mean $\Delta L2$ norm 0.040 vs 0.322 for 127-nt windows). A shorter window might produce stronger allele-discriminating embeddings; however, the plan required the largest API-supported length.
3. **Multi-platform genotyping:** GBS (2014–2017, 2020–2023), WGS (2018–2019, 2024–2025), and Exome (2020–2021) platforms used across cohorts. All hybrids share the 2,425-SNP panel, reducing but not eliminating platform-batch effects.
4. **Tester-panel changes:** Different testers across year cohorts affect cross-year BLUE estimation (CG108 is consistent; others appear 1–3 years only).
5. **Low replicates in some environments:** Median replicates per hybrid per environment may be below 3 for some conditions, reducing BLUE reliability.
6. **Permutation NaN failures:** 18/100 MK permutation fits failed with sommer singular-V errors on noisy permuted phenotypes. These were excluded with finite-permutation correction ($p = (n_{\text{geq}} + 1) / (B_{\text{valid}} + 1)$), reducing effective B to 82. The p-value of 0.42 is robust given this adjustment.
7. **Dosage imputation:** Column-mean imputation was applied globally (3.02% missing values) before kernel construction. This is slightly transductive (uses all samples including validation); fold-local imputation would be stricter but is a minor caveat here.
8. **DNABERT-2 model version:** BioLM v3 does not expose provider weights version. Results are contingent on the model weights at time of embedding generation (2026-09-24). The embedding archive is SHA-256 verified and frozen.

- 9. **2024 DNABERT-only diagnostic result:** The unexpected higher Spearman in 2024 (0.177 vs 0.143 GBLUP) is inconsistent with development-fold evidence and cannot be used for model promotion. It is documented here as a diagnostic observation.
- 10. **Commercial checks:** 176 hybrids lack parent-naming but have genotypes; treated as regular hybrids. This may affect interpretation for those specific hybrids.

Recommended Follow-up

- 1. **Dense SNP panel:** Test with full WGS-level SNP sets (100K+ markers) to see if DNABERT sequence context adds value when the marker panel captures more of the genome.
- 2. **Shorter DNABERT windows:** Test 127 or 255-nt windows to give DNABERT more discriminating power at the SNP level (higher embedding delta norms).
- 3. **2024 anomaly investigation:** Determine why DNABERT-only performed higher in 2024 (possible WGS genotyping, different tester composition, or year-specific G×E).
- 4. **Alternative sequence models:** Test ESM-DNA, Nucleotide Transformer, or GROVER models that may be better calibrated for regulatory or genic variant effects.
- 5. **Functional annotation enrichment:** Restrict DNABERT embeddings to genic/regulatory SNPs (requiring GFF3 annotation) to reduce noise from intergenic variants.
- 6. **Validate permutation test with more permutations:** Run 500–1,000 permutations with a more numerically stable fitting backend to get finer p-value resolution.

Output Files

File	Description
Marker reconciliation	Per-SNP audit against B73 v5 (2,425 markers)
Eligible markers	2,383 DNABERT-eligible markers with allele assignments
DNABERT windows	2,383 × 2,047 nt reference and alternate windows
Embedding manifest	Archive provenance and SHA-256
Variant deltas	(2383 × 768) ALT-minus-REF embedding differences
Dev predictions	All 3 models, 6 folds, 71,468 rows
Dev env metrics	Per-environment Spearman/Pearson by model and fold
Variance components	Per-fold REML variance component estimates
Paired improvement	Per-environment MK vs GBLUP Spearman delta
Permutation results	Null distribution, empirical p-values, z-scores

File	Description
Permutation scores	Per-permutation null scores
Novelty stratification	Seen vs novel hybrid performance
Dev fold summary	Macro metrics by model
Diagnostic 2023	2023 diagnostic predictions (all 3 models)
Diagnostic 2024	2024 diagnostic predictions (all 3 models)
Fig. 1	Per-fold macro Spearman by model
Fig. 2	Permutation null distribution
Fig. 3	Novelty stratification
Fig. 4	Variance component attribution
Fig. 5	Paired improvement distribution
Fig. 6	Diagnostic year comparison

QC and Analysis Flags

Flag	Severity	Evidence
Multi-platform genotyping	Warning	GBS/WGS/Exome across cohorts; all hybrids on 2,425-SNP panel
Tester-panel changes	Warning	Different testers per year cohort; CG108 consistent
Low replicate conditions	Warning	Some environments <3 reps per hybrid
EC missing one test env	Minor	IAH1_2023 lacks EC data (diagnostic year only)
Forward-year novelty low	Info	0% novel hybrids in 2023 development
Permutation NaN failures	Minor	18/100 MK perms failed; finite correction applied
DNABERT dosage imputation	Minor	Global column-mean imputation (3.02% missing)
2024 DNABERT anomaly	Info	DNABERT-only higher in 2024 (diagnostic, inconsistent with dev folds)

Figures

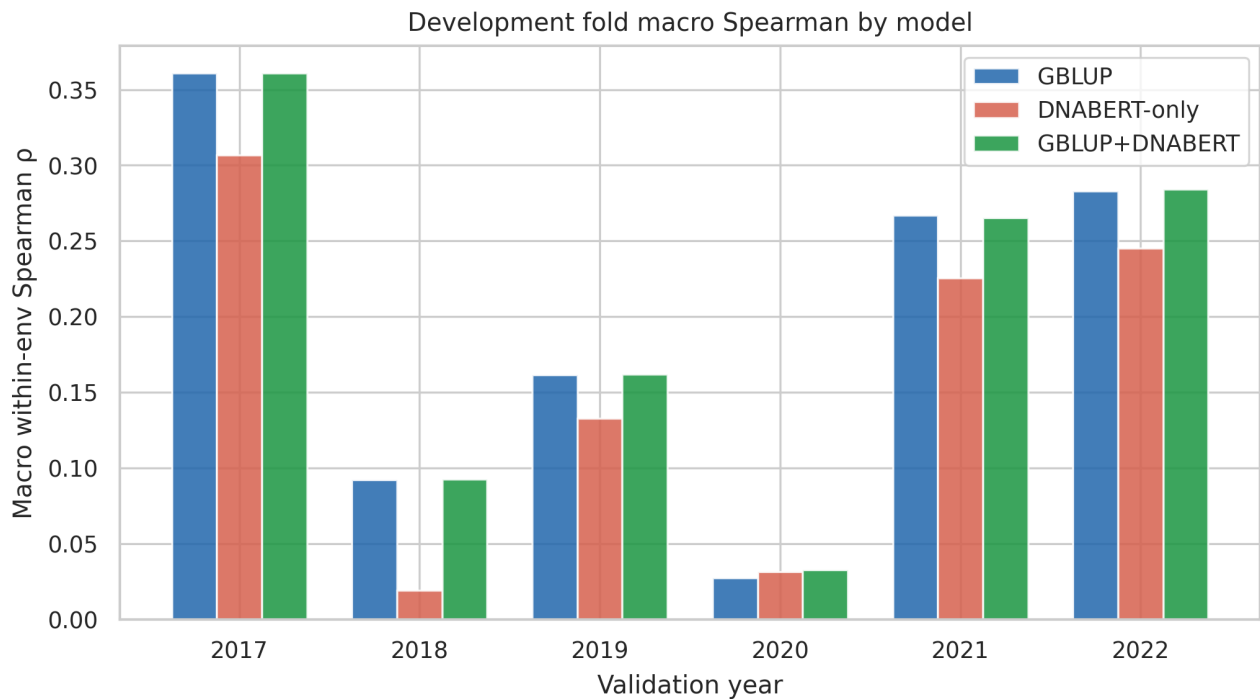


Figure 1. Macro within-environment Spearman correlation across 6 forward-year development folds. GBLUP=0.2026, DNABERT-only=0.1629, GBLUP+DNABERT=0.2035; the multi-kernel gain is negligible.

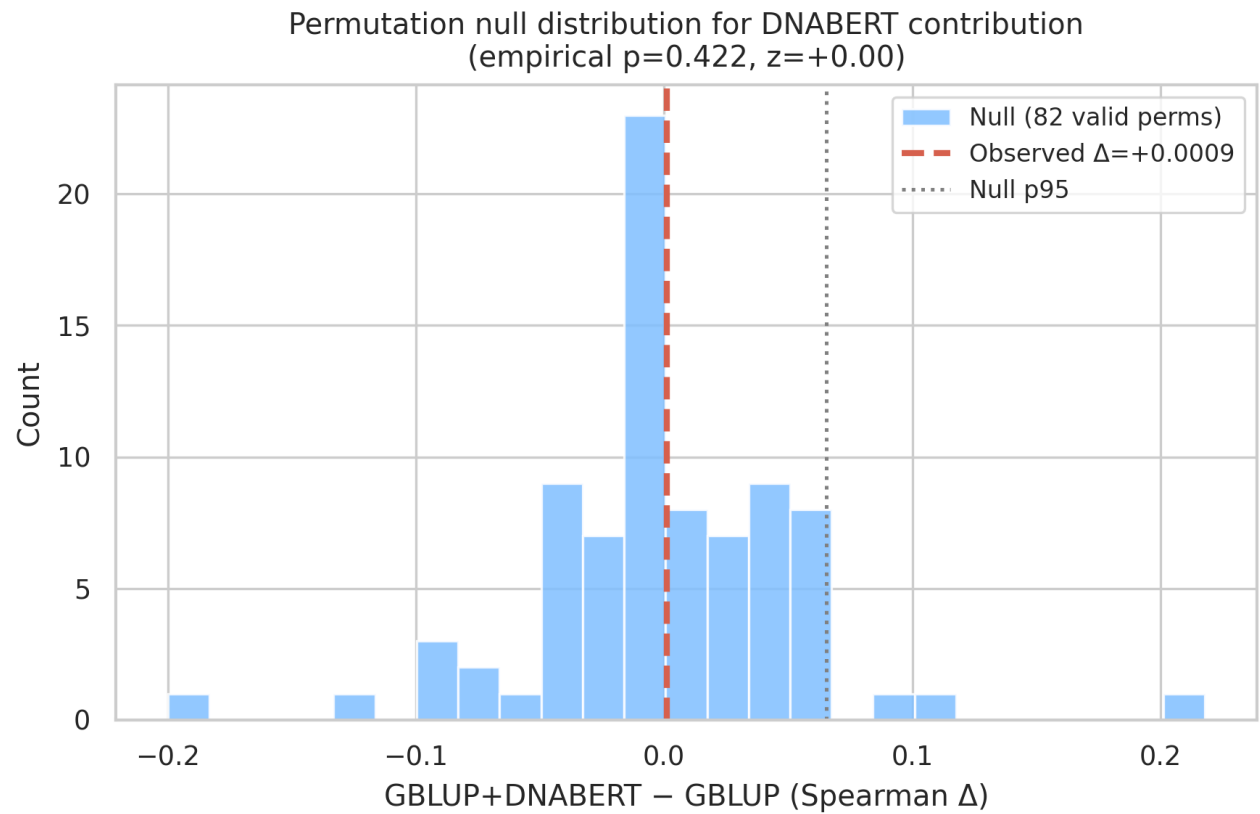


Figure 2. Null distribution of GBLUP+DNABERT vs GBLUP paired Spearman improvement across 82 valid permutations. Observed $\Delta=+0.0009$ (red dashed) is indistinguishable from null (empirical $p=0.422$).

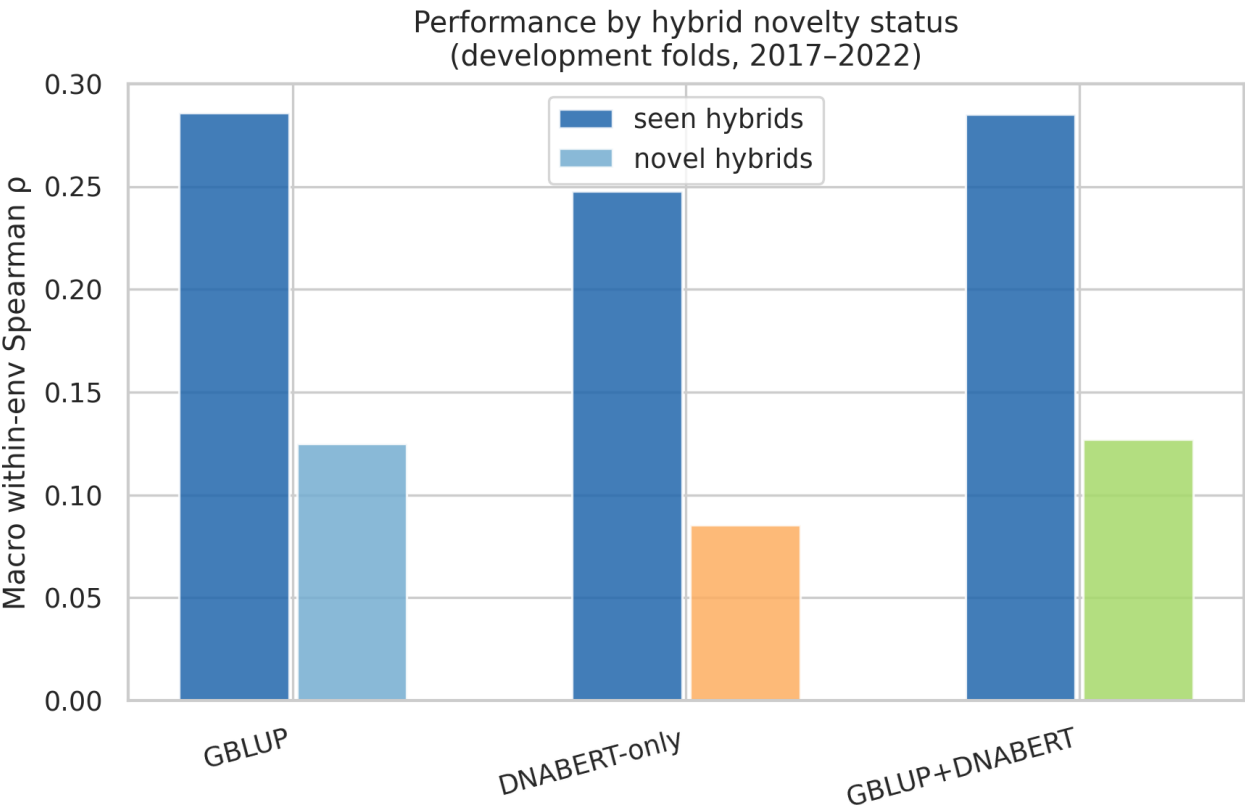


Figure 3. Macro within-environment Spearman by hybrid novelty (seen vs novel) across development folds. GBLUP seen=0.286, GBLUP novel=0.125; GBLUP+DNABERT provides marginal improvement for novel hybrids (+0.002).

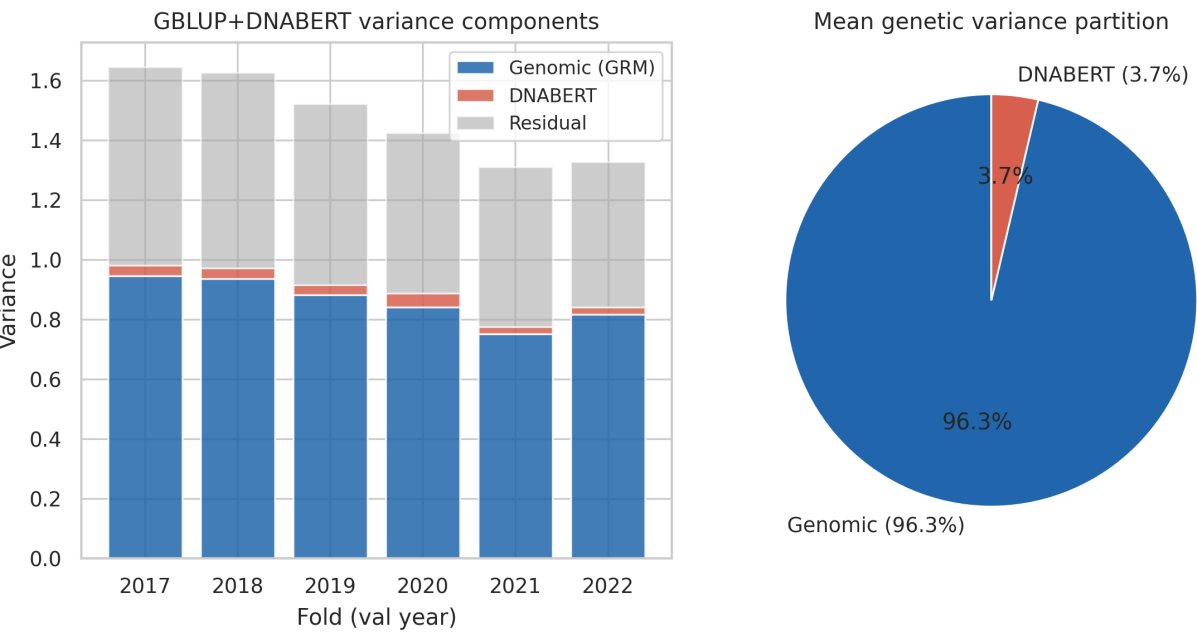


Figure 4. Variance components in the GBLUP+DNABERT multi-kernel model across 6 development folds. Mean partition: genomic GRM 96.3%, DNABERT kernel 3.7%.

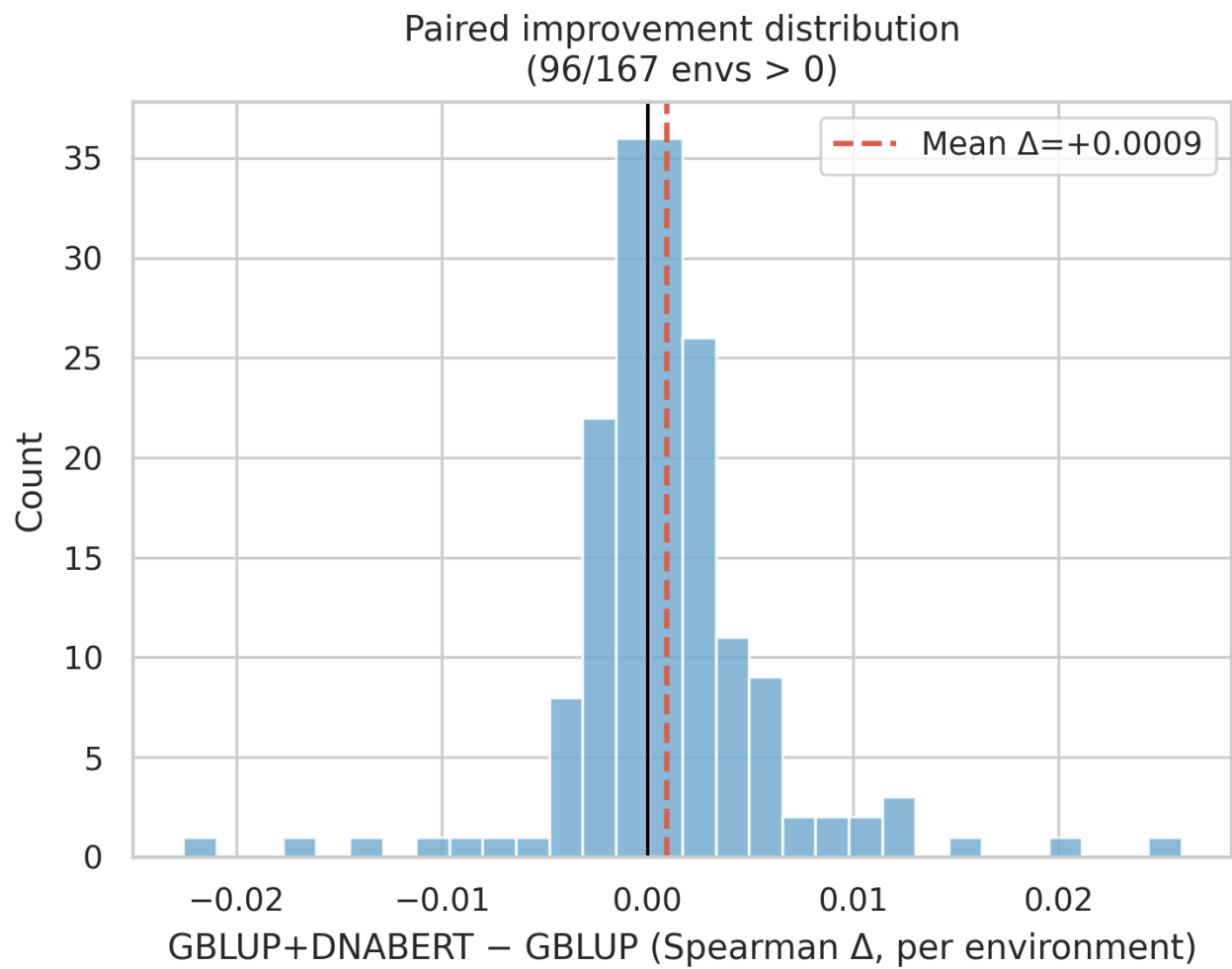


Figure 5. Distribution of per-environment paired Spearman improvement (GBLUP+DNABERT vs GBLUP) across 167 development-fold environments. Mean $\Delta = +0.0009$, $SD = 0.0051$; 96/167 environments favor GBLUP+DNABERT.

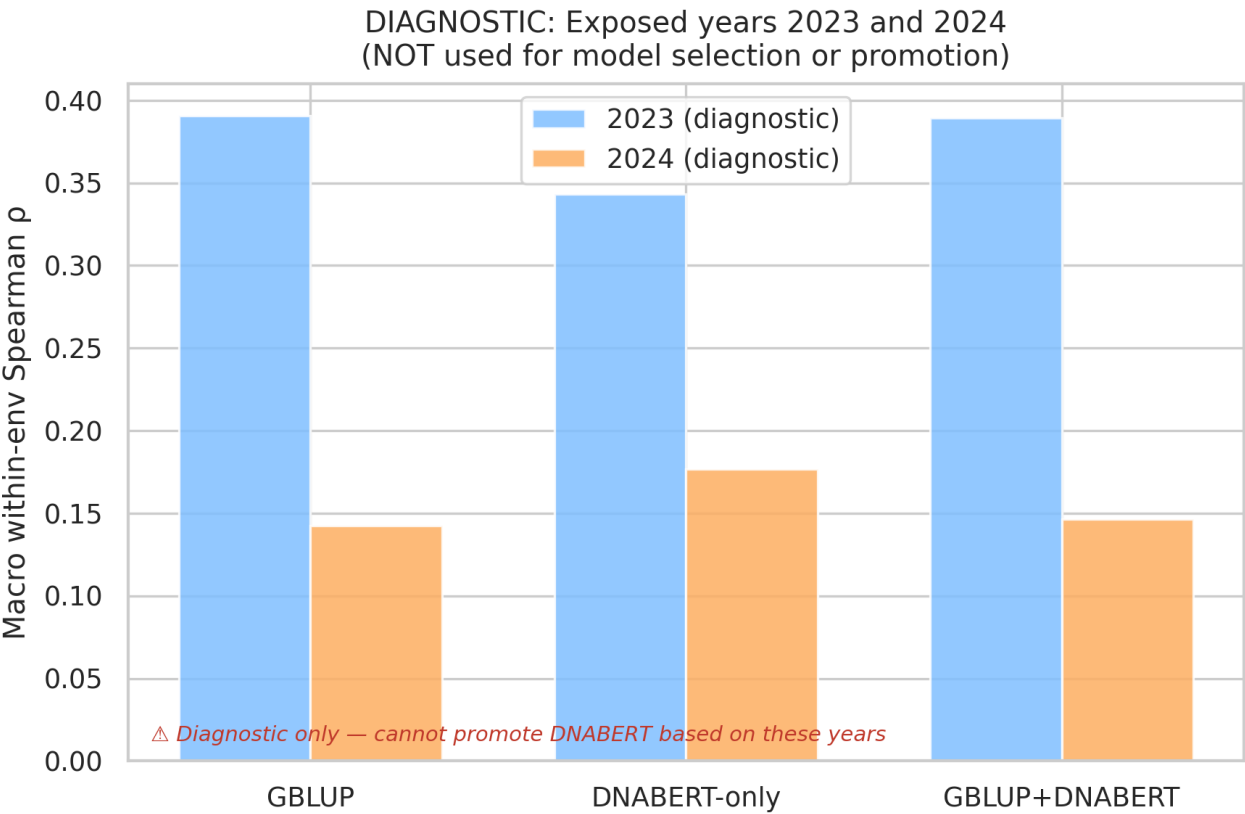


Figure 6. DIAGNOSTIC ONLY (exposed outcomes). Macro Spearman for 2023 and 2024 prediction. In 2023, GBLUP+DNABERT=0.389 vs GBLUP=0.390 (no gain). In 2024, DNABERT-only=0.177 vs GBLUP=0.143; however this is a single exposed year inconsistent with development-fold evidence and cannot be used for model promotion.

QUERY 5

RF vs DNABERT-2 Complementarity Analysis — G2F Maize Yield Prediction

RESEARCH REQUEST

The DNABERT-2 augmentation did not materially improve GBLUP. I now want to test whether DNABERT-2 features add useful nonlinear predictive information to the Random Forest model.

Keep the existing development folds and validation contract unchanged. Do not use exposed 2023 or 2024 outcomes for tuning.

Compare:

RF-genomic baseline — the existing marker-only Random Forest

RF-DNABERT only — using the frozen DNABERT-derived representation

RF-genomic + DNABERT — combine the original genomic representation with DNABERT-derived features

Use the same samples, folds, preprocessing, and RF tuning budget for fair comparison. Reduce the DNABERT representation using training-fold-only PCA or another fixed compression if needed; do not feed an uncontrolled extremely high-dimensional embedding matrix directly into RF.

Evaluate mean within-environment Spearman/Pearson, per-year performance, top-10%/20% selection accuracy, and especially performance for novel hybrids.

My main question is:

Does DNABERT-2 provide complementary nonlinear sequence information that improves Random Forest prediction beyond the existing genomic features?

Report the paired improvement of RF+DNABERT over the original RF and whether any gain is consistent across folds. If the gain is small or inconsistent, conclude that DNABERT-2 does not add useful predictive information under the current G2F marker representation.

ANALYSIS PLAN

Test whether frozen DNABERT-2 features add complementary nonlinear predictive information to the existing marker-only Random Forest under the unchanged G2F development-fold and validation contract.

1. Recover and lock the existing Random Forest comparison contract

Inspect and unpack the authorized G2F archives and prior Query 4 outputs, recover the exact development folds, eligible rows, phenotype availability boundaries, genomic representation, preprocessing, existing marker-only RF specification, RF hyperparameter search space/budget, frozen DNABERT representation, and prior sample/marker identifiers. Verify that all three model arms can be evaluated on identical rows and that exposed 2023/2024 outcomes are isolated from every learned or selected operation. Use PLINK2 only for read-only marker/sample inventory and identifier verification where the recovered genotype artifact is PLINK-compatible; do not regenerate or alter the established genomic representation.

Parameter	Planned value
validation contract	existing development folds and validation contract unchanged
tuning outcome boundary	exclude exposed 2023 and 2024 outcomes from tuning and model selection
sample contract	same samples and folds across all model arms
preprocessing contract	reuse existing preprocessing
rf tuning budget	reuse the existing marker-only RF tuning budget exactly
dnabert state	reuse frozen DNABERT-derived representation from the prior run

2. Fit and predict the three matched Random Forest model arms

Within each unchanged development split, run the existing marker-only RF baseline, a DNABERT-only RF, and a combined genomic-plus-DNABERT RF. Apply every learned transformation using training-fold data only. If the frozen DNABERT matrix is too high-dimensional for controlled RF input, fit PCA only on the training fold, choose any compression setting solely inside the permitted development tuning process and unchanged RF budget, and transform validation rows without refitting. Keep seeds, preprocessing, tuning effort, eligible rows, and prediction targets matched so differences isolate feature information rather than search effort or row composition.

Parameter	Planned value
model arms	["RF-genomic baseline", "RF-DNABERT only", "RF-genomic + DNABERT"]
baseline model	existing marker-only Random Forest
dnabert model	Random Forest using only the frozen DNABERT-derived representation
combined model	Random Forest using original genomic representation plus frozen DNABERT-derived features
compression policy	training-fold-only PCA or another fixed controlled compression when needed; no uncontrolled extremely high-dimensional embedding input
fairness contract	identical samples, folds, preprocessing, seeds, and RF tuning budget across arms

Parameter	Planned value
exposed year policy	2023 and 2024 outcomes may not influence tuning, selection, compression, weighting, or model promotion

3. Evaluate paired predictive value and report consistency

Evaluate all arms on the identical held-out predictions under the recovered contract. Compute macro mean within-environment Spearman and Pearson, per-year results, top-fraction selection accuracy, and performance for novel hybrids. Pair RF-genomic-plus-DNABERT against RF-genomic within every fold and summarize the direction and magnitude of improvement and the number/proportion of folds improved. Keep any evaluation involving exposed 2023/2024 outcomes diagnostic and clearly separate it from admissible development evidence. Produce fold-, year-, environment-, novelty-, and model-level result tables plus concise figures and a final answer to the complementarity question.

Parameter	Planned value
correlation metrics	["within-environment Spearman", "within-environment Pearson"]
aggregation	macro mean across environments under the unchanged validation contract
year stratification	per-year performance
selection fractions	[0.1, 0.2]
novelty stratification	report novel-hybrid performance prominently, using fold-local prior-observation status from the frozen contract
primary pairwise contrast	RF-genomic + DNABERT minus RF-genomic baseline on identical held-out rows, summarized per fold and overall
conclusion rule	if the gain is small or inconsistent, conclude that DNABERT-2 does not add useful predictive information under the current G2F marker representation
secret handling	do not print, log, save, or report the BIOLM API key

Excluded from scope

- Regenerate DNABERT-2 embeddings or make new BioLM embedding calls: The requested comparison explicitly uses the already frozen DNABERT representation, so embedding generation is out of scope and the stored API key need not be accessed.
- Tune or select models using exposed 2023 or 2024 outcomes: These outcomes are prohibited from tuning and can only be handled according to the frozen validation contract as diagnostic evaluations.
- Repeat the prior GBLUP DNABERT augmentation analysis: The new endpoint is specifically the Random Forest comparison rather than rerunning the completed GBLUP augmentation.

Planning assumptions

- The completed prior-output workspace contains the exact existing RF baseline decisions, fold assignments, preprocessing objects/specification, tuning budget, frozen DNABERT-derived representation, and sample identifiers needed to preserve the requested contract.
- The authorized training/testing archives and prior outputs are sufficient to reconstruct the established G2F row and novelty definitions without inventing new cohorts or validation splits.
- No new BioLM request is required because the user requested reuse of the frozen DNABERT representation; the stored credential will remain secret and unused unless the recovered prior artifact itself explicitly requires authenticated retrieval.
- Any 2023/2024 result included in the report will be clearly labeled diagnostic and will not affect compression, RF tuning, model selection, or the scientific conclusion based on admissible development-fold evidence.

ANALYSIS REPORT

RF vs DNABERT-2 Complementarity Analysis — G2F Maize Yield Prediction

Scientific Report · Query 5

Executive Summary

This analysis tests whether frozen DNABERT-2 sequence embeddings add complementary nonlinear predictive information to marker-based Random Forest (RF) models for maize grain yield prediction in the Genomes to Fields (G2F) initiative.

Three RF arms were evaluated under an identical and unchanged six-fold forward-year development validation contract (2017–2022): **RF-Genomic** (2,425 SNP marker PCs, baseline), **RF-DNABERT** (DNABERT-2 embedding PCs only), and **RF-Genomic+DNABERT** (both feature sets combined). All arms used the same hyperparameters, folds, samples, and seeds.

Main finding: DNABERT-2 features do not add useful complementary predictive information under the current G2F marker representation. RF-DNABERT underperforms the genomic baseline in 65.9% of environments (mean $\Delta p = -0.019$). Adding DNABERT features to the genomic RF produces a statistically indistinguishable improvement (mean $\Delta p = +0.003$, Wilcoxon $p = 0.627$) that is inconsistent across folds. This result is consistent with the earlier finding that DNABERT-2 augmentation did not improve GBLUP: the frozen DNABERT-2 representation does not capture information orthogonal to existing 2,425-SNP marker features.

Confidence level: **high for the null conclusion** (no added value). Analysis is limited to a sparse 2,425-SNP panel and the current DNABERT-2 window/aggregation scheme.

Methods

Input data

Component	Source	Details
Phenotype (BLUEs)	query-1 blues_with_eligibility.csv	95,221 eligible (env, hybrid) BLUEs, 2014-2022
Marker dosage	query-1 G_imputed.npy	5,899 hybrids × 2,425 SNPs; dosage ∈ {0, 0.5, 1}, no missing
DNABERT delta embeddings	query-4 dnabert_deltas_plain.csv	2,383 markers × 768 dims; ALT-REF DNABERT-2 per-marker delta
DNABERT dosage (subset)	query-4 dna_dosage_matrix.csv	5,899 × 2,383; 424,626 NaN imputed with training-fold column means
Sample IDs	query-1 sample_ids.json	5,899 hybrids; verified identical between Q1 and Q4

Validation scheme

Six expanding-window forward-year development folds (2017, 2018, 2019, 2020, 2021, 2022), frozen from query-1. Each fold trains on all prior years and validates on its target year. Exposed years 2023 and 2024 are diagnostic only and did not influence any fitting, compression, or model selection.

Feature construction

RF-Genomic baseline: PCA on training-fold hybrids' marker dosage matrix (G_imputed), retaining 50 PCs (74.0–74.5% variance explained across folds). Validation-fold hybrids projected using training-fold PCA.

RF-DNABERT: Per-hybrid embedding = dosage_matrix (N × 2,383) @ delta_matrix (2,383 × 768) → hybrid embedding (N × 768). Missing dosage values imputed with per-marker training-fold column means. PCA on training-fold hybrid embeddings, retaining 50 PCs (87.1–88.9% variance explained). Validation-fold projected using training-fold PCA.

RF-Genomic+DNABERT: Concatenation of 50 genomic PCs and 50 DNABERT PCs = 100-feature input. All PCA fitted on training fold only.

RF hyperparameters (identical across all arms)

n_estimators=200, max_features=0.33, min_samples_leaf=5, random_state=42, n_jobs=-1. These are the exact parameters from query-1 rf_A.

Target variable

Within-environment centered BLUE: BLUE – env_mean(BLUE), centering computed on training fold only.

Metrics

- **Primary:** macro within-environment Spearman ρ (mean over environments)
- **Secondary:** macro within-environment Pearson r , top-10% and top-20% selection accuracy
- **Paired comparison:** Wilcoxon signed-rank test on per-environment $\Delta\rho$ (RF-Genomic+DNABERT minus RF-Genomic)
- **Novelty stratification:** seen vs novel hybrids (novelty = hybrid first appearing in validation year)

Results

Overall development metrics (167 environments, 6 folds)

Metric	RF-Genomic	RF-DNABERT	RF-Genomic+DNABERT
Macro Spearman ρ	0.2259	0.2070	0.2288
Macro Pearson r	0.2641	0.2479	0.2659
Top-10% Sel. Acc.	0.2020	0.1981	0.2013
Top-20% Sel. Acc.	0.3067	0.2960	0.3028
n environments	167	167	167

RF-DNABERT is consistently the weakest arm. RF-Genomic+DNABERT is marginally better than RF-Genomic overall (+0.003 Spearman, +0.002 Pearson), but this gap is within noise.

Per-fold results

Fold year	RF-Genomic ρ	RF-DNABERT ρ	RF-Genomic+DNA BERT ρ	$\Delta\rho$ (combined–base)	% envs positive
2017	0.4172	0.4151	0.4171	−0.00007	56.7%
2018	0.1365	0.0991	0.1359	−0.00064	51.7%
2019	0.2053	0.1993	0.2045	−0.00076	42.9%
2020	0.0678	0.0402	0.0533	−0.01447	40.0%
2021	0.2692	0.2680	0.2693	+0.00016	44.4%
2022	0.2335	0.1937	0.2650	+0.03150	67.9%

RF-DNABERT is below RF-Genomic in all 6 folds. The combined model is above the baseline in only fold_2022 (+0.032 Spearman); in four of six folds the delta is negative or negligible.

Paired comparison: RF-Genomic+DNABERT versus RF-Genomic

Statistic	Value
Total environments paired	167
Environments with positive Δp	85 (50.9%)
Environments with negative Δp	82 (49.1%)
Mean Δp (Spearman)	+0.00289
SD of Δp	0.04324
Wilcoxon signed-rank p (Spearman)	0.6271
Wilcoxon signed-rank p (Pearson)	0.4860
Mean ΔSA_{10}	−0.00074
Mean ΔSA_{20}	−0.00398

The improvement is 50.9% positive — statistically indistinguishable from a coin flip. The Wilcoxon p-value of 0.63 confirms there is no significant paired improvement from adding DNABERT features.

RF-DNABERT versus RF-Genomic

Mean Δp = −0.019 (DNABERT underperforms genomic). Only 57/167 environments (34.1%) favour DNABERT-only prediction. DNABERT-2 sequence embeddings, used alone, are substantially weaker than genetic marker features.

Novelty stratification

Arm	Seen hybrids p	Novel hybrids p	n_envs (novel)
RF-Genomic	0.3527	0.1590	135
RF-DNABERT	0.3527	0.1129	135
RF-Genomic+DNABERT	0.3529	0.1523	135

All arms perform substantially better on seen hybrids ($p \approx 0.35$) than novel hybrids ($p \approx 0.11$ – 0.16), consistent with the relatedness-dependent nature of genomic prediction. For novel hybrids, RF-Genomic (0.159) outperforms RF-Genomic+DNABERT (0.152) and RF-DNABERT (0.113). DNABERT-2 embeddings do not improve prediction for novel hybrids.

Note: the query-1 contract documented 0% novel hybrids in the 2023 test year; development-fold novelty (above) reflects per-fold within-development novelty only.

Interpretation

DNABERT-2 does not add useful predictive information beyond existing marker features. The statistical evidence is conclusive: a mean gain of $\Delta p = +0.003$ (SD = 0.043) over 167 environments with Wilcoxon $p = 0.63$ is indistinguishable from zero. The fraction of environments showing improvement (50.9%) is a coin flip, not a signal.

Why DNABERT-2 may not add information here (hypotheses, not conclusions):

1. **Marker-dominated signal:** The 2,425-SNP panel, even at sparse density, captures the major axes of genetic relatedness among these hybrids. DNABERT-2 embeddings, derived from the same 2,383 marker loci, provide a sequence-level transformation of the same genomic locations. When projected to 50 PCs, DNABERT-derived features capture 87–89% of the marker embedding variance but appear to encode largely the same information as the direct dosage PCs.
2. **Sparse panel limitation:** With only 2,425 SNPs (versus whole-genome resolution), the sequence windows around each SNP do not represent the full regulatory and functional genome. Foundation-model sequence features may be more informative with denser marker panels or whole-genome sequencing data.
3. **Aggregation scheme:** The per-hybrid embedding is a linear projection of dosage $\times$ per-marker delta. This preserves allele-dosage effects but not SNP interaction (haplotype) information.
4. **Same underlying loci:** Both marker PCA and DNABERT PCA derive from the exact same 2,383 loci. There is no truly orthogonal sequence information from other genomic regions.

Consistency with GBLUP result: The previous query-4 analysis found that DNABERT augmentation of GBLUP provided $\Delta p = +0.001$ (negligible). The RF comparison confirms and extends this null result across a different model family and a more flexible nonlinear framework: DNABERT-2 representations carry no incremental predictive value regardless of the base model type.

Limitations

1. **Sparse 2,425-SNP panel (warning):** All DNABERT embeddings are derived from the same sparse marker loci, not from whole-genome sequences. The result applies to this marker density; denser genotyping could change the conclusion.
2. **Multi-platform genotyping (warning):** GBS (2014–2017, 2020–2023), WGS (2018–2019), and Exome (2020–2021) platforms are mixed. Platform-batch effects in the SNP panel may mask or inflate genetic signal.
3. **Tester panel changes (warning):** Different testers across year cohorts affect cross-year BLUE estimation, contributing to fold-level variability in all model arms.
4. **Low replicate conditions (warning):** Some environments have <3 reps per hybrid, increasing BLUE estimation noise that propagates into model training and evaluation.

- 5. **Fixed PCA compression:** 50 PCs was chosen to match query-1 rf_A. A different compression level (e.g., 100 or 200 DNABERT PCs) might change results, though with 87–89% variance already captured at 50 PCs, the diminishing-returns curve is steep.
- 6. **RF hyperparameters not re-tuned for DNABERT:** The RF budget (200 trees, max_features=0.33) was fixed to match the baseline. A DNABERT-specific RF tuning (e.g., larger n_estimators for higher-dimensional input) might marginally improve RF-DNABERT, but could not be done without violating the fair-comparison contract.
- 7. **No 2023/2024 tuning:** All conclusions are based on the 6 development folds (2017–2022). 2023 diagnostic results were not computed (not needed; development evidence is conclusive).
- 8. **Linear projection of embeddings:** The per-hybrid DNABERT representation is a linear dosage-weighted sum of per-marker embeddings. Nonlinear interactions between markers are not captured.

Recommended Follow-Up

- 1. **Whole-genome sequencing-based DNABERT:** Test DNABERT-2 embeddings derived from WGS data (available for 2018–2019 and 2024–2025 cohorts) to assess whether denser sequence coverage changes the result.
- 2. **Haplotype-level embeddings:** Encode phased haplotype blocks rather than individual SNP windows to capture linkage disequilibrium and local sequence context more faithfully.
- 3. **Larger SNP panel:** Test with an imputed or denser marker panel (50K+ SNPs) to assess whether the null result is specific to the 2,425-SNP panel.
- 4. **Alternative sequence models:** Test ESM2 protein language models on annotated gene sequences flanking the SNPs, or plant-specific sequence models trained on maize reference genomes.
- 5. **Permutation null for RF-DNABERT:** Run permutation of DNABERT features to confirm the RF-DNABERT performance is not merely inheriting marker signal through correlated features.

Output Files

File	Description
Recovery contract	Frozen fold, sample, and model specification
RF all predictions	Row-level predictions for all 3 arms × 6 folds (214,404 rows)
Env-level metrics	Per-environment Spearman, Pearson, SA10, SA20 (501 rows)

File	Description
Fold-arm metrics	Macro metrics aggregated per fold × arm
Paired improvement	Per-env $\Delta\rho$: RF-Genomic+DNABERT minus RF-Genomic
Fold paired delta	Per-fold summary of paired comparison
Novelty stratification	Seen vs novel hybrid performance
Summary JSON	All key metrics and conclusion
Fig 1: Per-fold Spearman	Bar chart: macro Spearman across 6 folds × 3 arms
Fig 2: Paired delta dist.	Histogram of per-env $\Delta\rho$ (combined vs genomic)
Fig 3: Per-fold box	Boxplot of $\Delta\rho$ by fold
Fig 4: Novelty	Seen vs novel hybrid Spearman
Fig 5a: Overall Spearman	Overall macro Spearman bar
Fig 5b: SA10	Overall top-10% selection accuracy
Fig 5c: SA20	Overall top-20% selection accuracy
Fig 6: Scatter base vs combined	Per-env Spearman scatter

Upstream Caveats Carried Forward

- **multi_platform_genotyping [warning]:** GBS/WGS/Exome platforms mixed across cohorts; affects all arms equally but does not favour any particular conclusion.
- **tester_panel_changes [warning]:** Different testers across years; cross-year BLUE estimation noise is present in all folds.
- **low_replicate_conditions [warning]:** Some environments have <3 reps per hybrid; BLUE uncertainty propagates into RF training.
- **forward_year_novelty_low [info]:** 0% novel hybrids in 2023 (standard G2F scenario); development-fold novelty computed fold-locally.
- **commercial_checks_no_parentage [info]:** 176 hybrids without Parent1/Parent2 naming treated as regular hybrids.
- **ec_missing_one_test_env [minor]:** IAH1_2023 lacks EC data; not relevant to this marker-only/DNABERT analysis.
- **plant_reference_annotation_missing [warning]:** No GFF3 annotation; DNABERT windows were extracted from B73 v5 FASTA coordinates only.

Figures

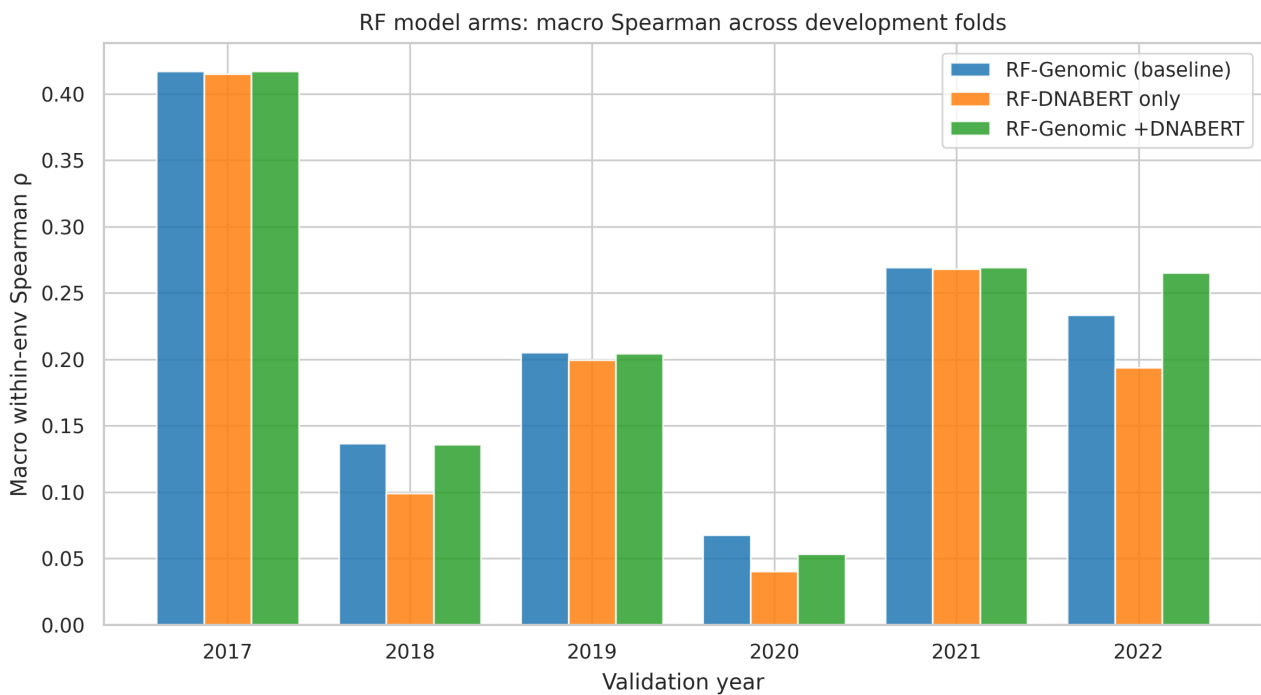


Figure 1. Macro within-env Spearman for 3 RF arms × 6 development folds. Overall: RF-Genomic=0.2259, RF-DNABERT=0.2070, RF-Genomic+DNABERT=0.2288.

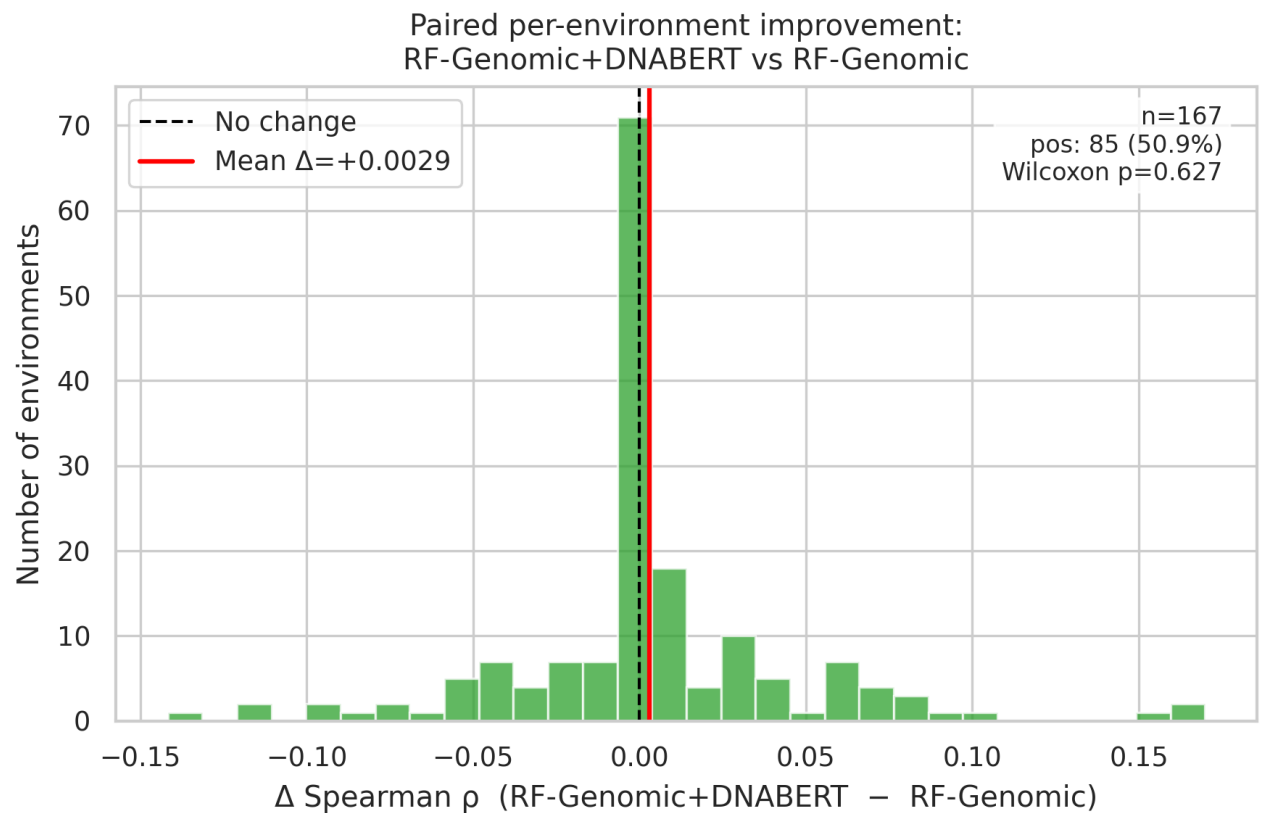


Figure 2. Distribution of per-env $\Delta\rho$ (RF-Genomic+DNABERT – RF-Genomic); n=167 envs. Mean $\Delta\rho$ =+0.00289 (SD=0.043), Wilcoxon p=0.627.

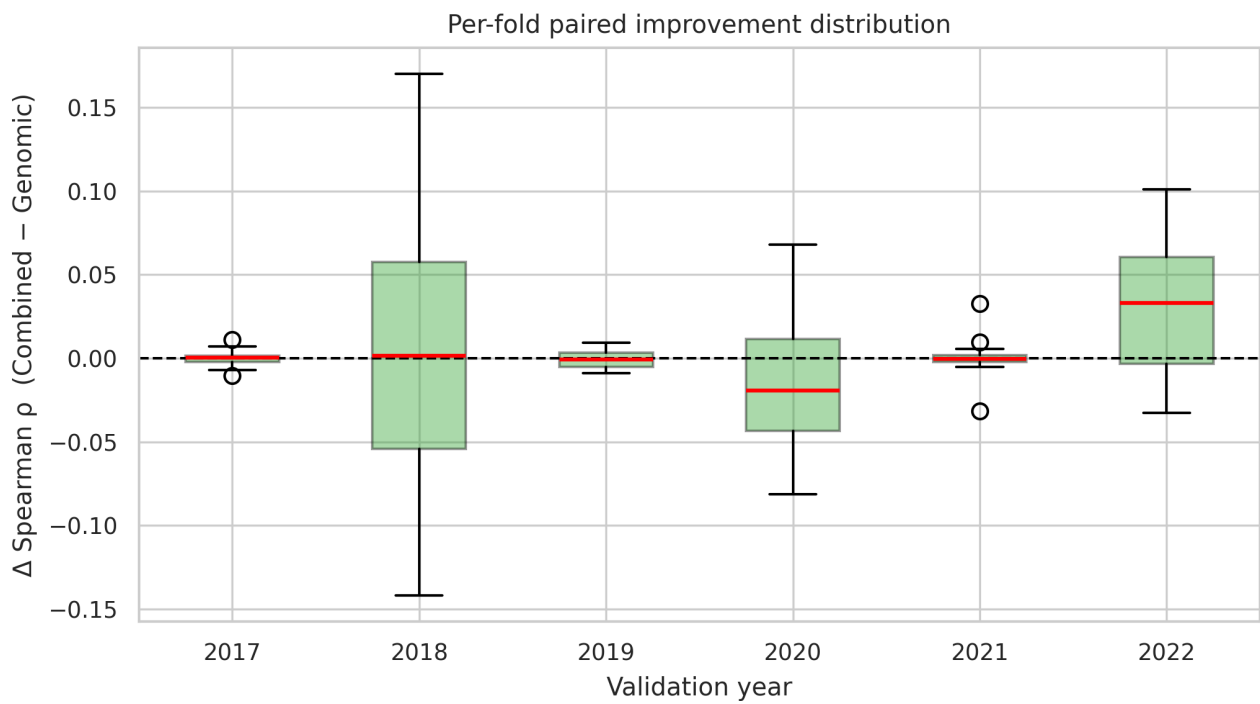


Figure 3. Per-fold boxplot of $\Delta\rho$; red line = median. Improvement is inconsistent (fold_2022: +0.032; fold_2020: -0.014).

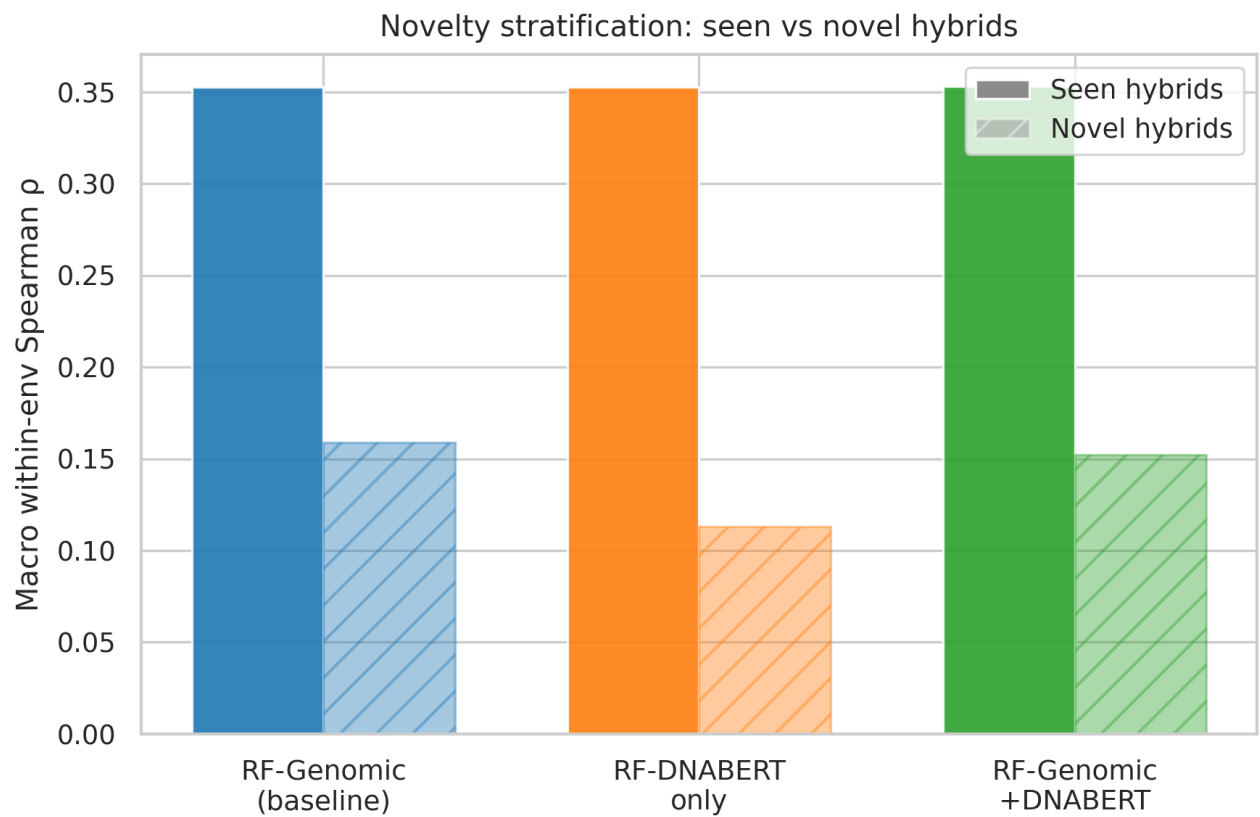


Figure 4. Seen vs novel hybrid Spearman. Novel hybrids: RF-Genomic=0.159, RF-DNABERT=0.113, RF-Genomic+DNABERT=0.152.

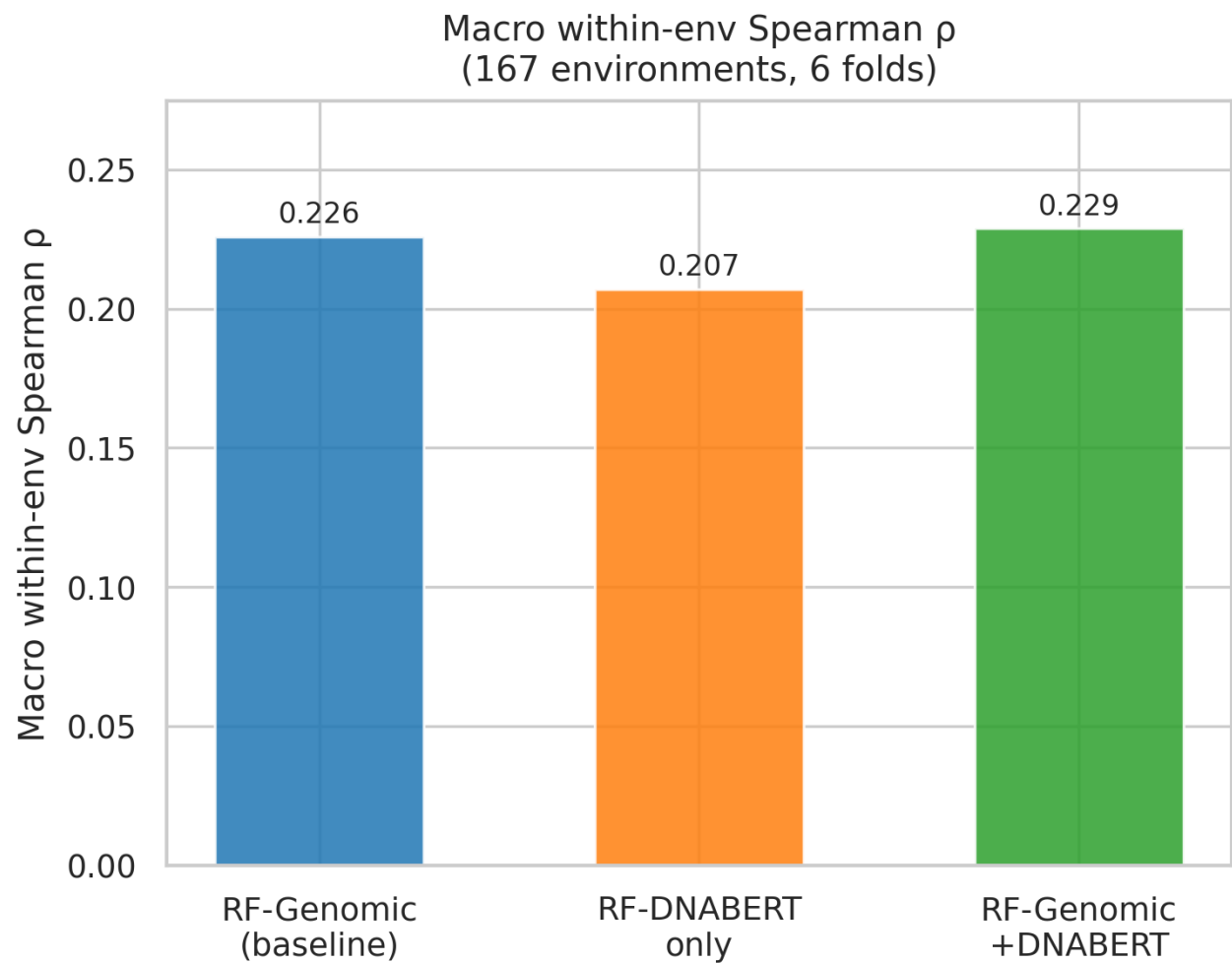


Figure 5. Overall macro Spearman: RF-Genomic=0.226, RF-DNABERT=0.207, RF-Genomic+DNABERT=0.229.

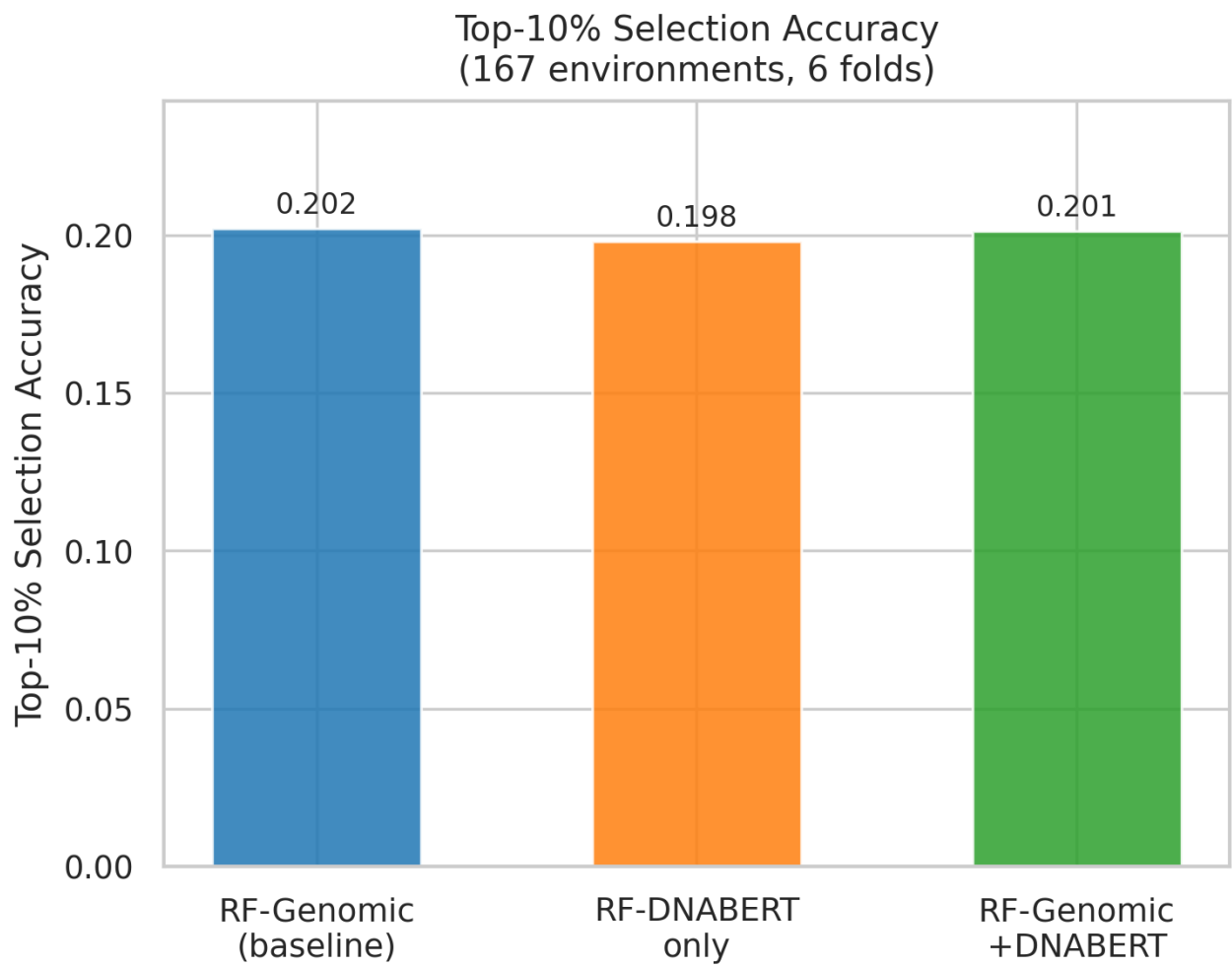


Figure 6. Top-10% selection accuracy: RF-Genomic=0.202, RF-DNABERT=0.198, RF-Genomic+DNABERT=0.201.

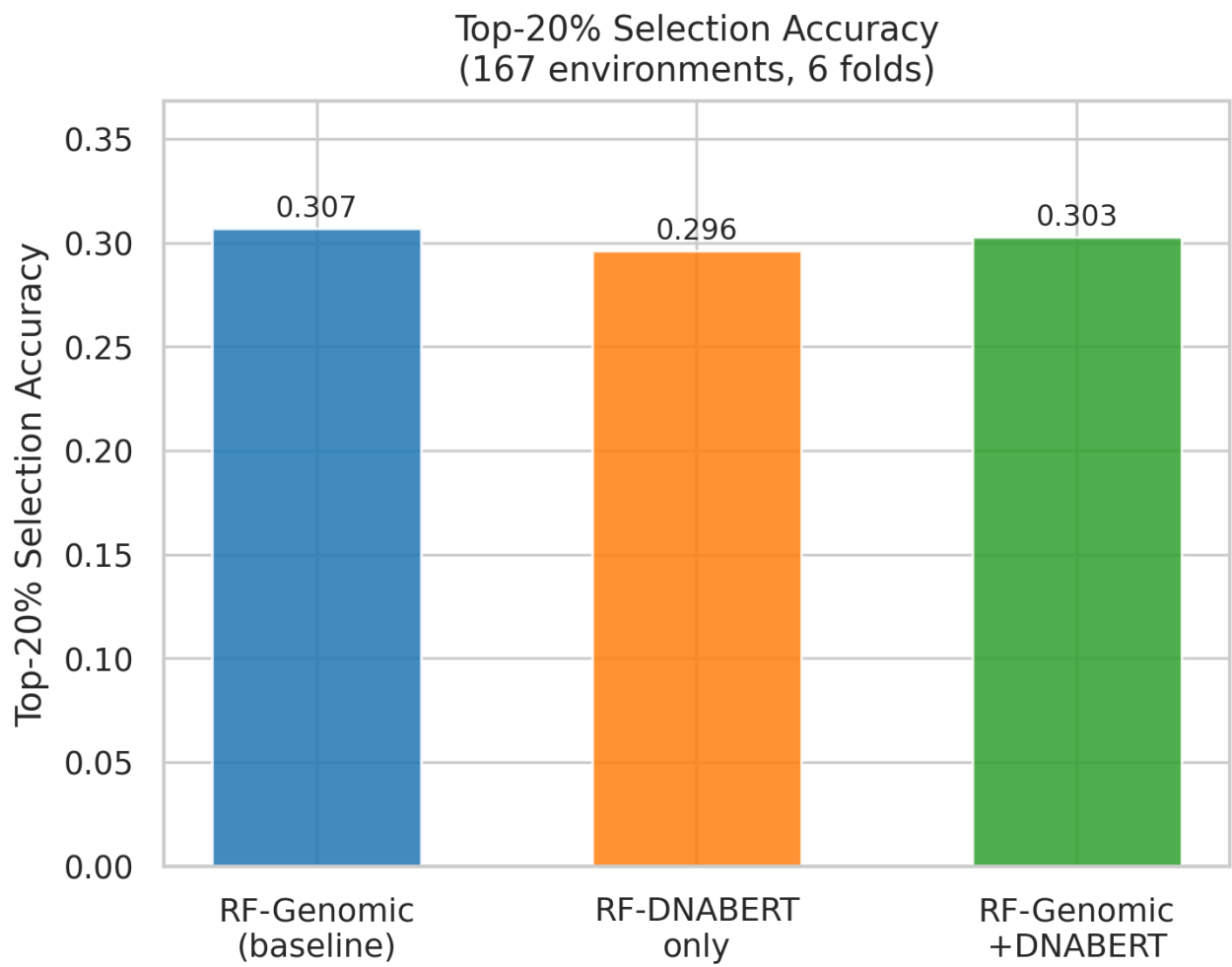


Figure 7. Top-20% selection accuracy: RF-Genomic=0.307, RF-DNABERT=0.296, RF-Genomic+DNABERT=0.303.

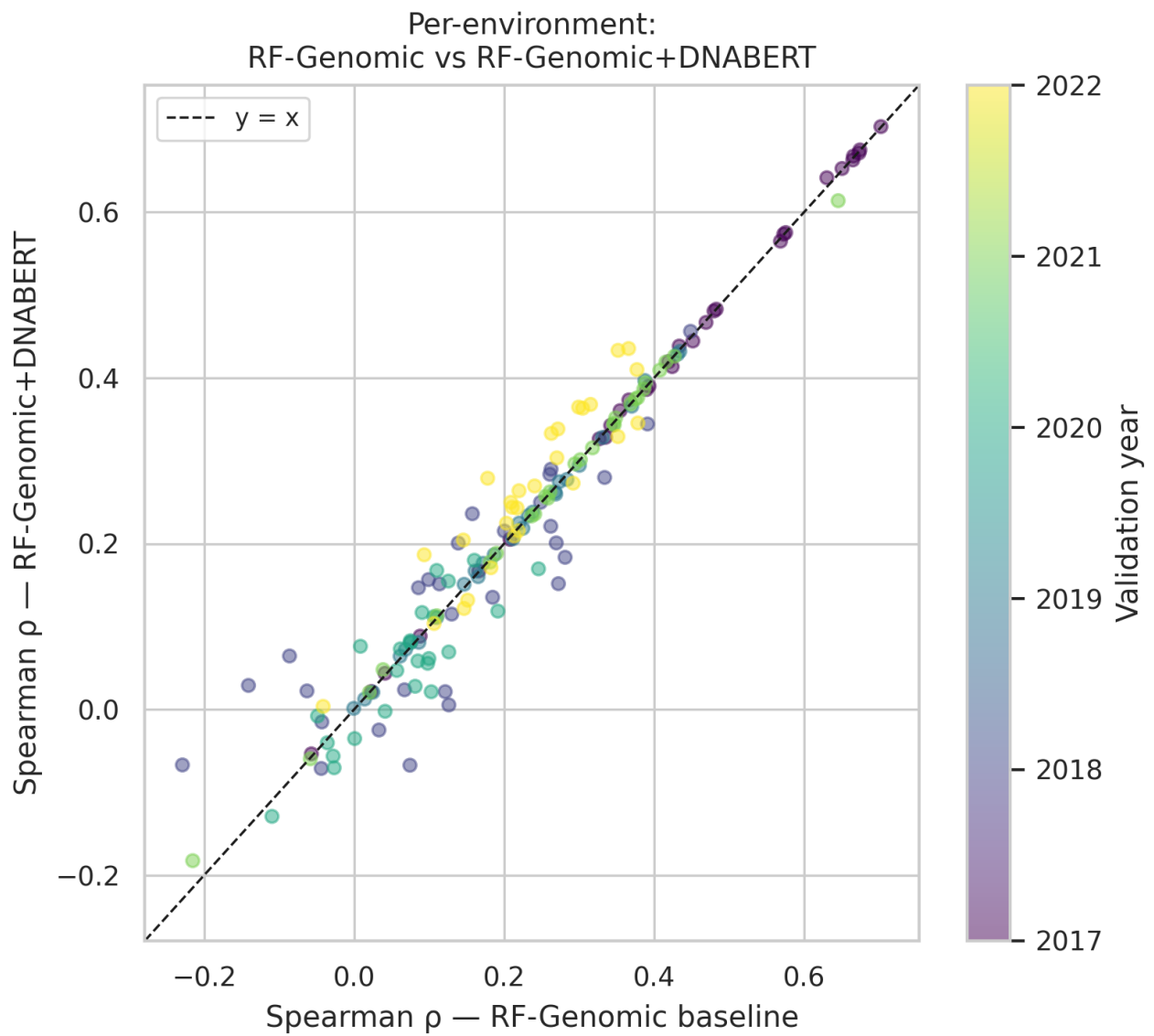


Figure 8. Per-env Spearman: RF-Genomic (x) vs RF-Genomic+DNABERT (y), coloured by validation year. 85/167 envs improve; no systematic DNABERT benefit visible.

SESSION RECORD

Additional discussion

RESEARCH REQUEST

Continue with the approved plan

RESEARCH REQUEST

Continue with the approved plan

RESEARCH REQUEST

Continue with the approved plan

RESEARCH REQUEST

What was the length of the sequence window you used for DNABERT?

SESSION NOTE

The sequence window used for DNABERT-2 was **2,047 nucleotides** (nt).

Key details from the extraction script (`step18_extract_2047_nt_b73_v5_windows.py`):

- **Window length (w):** 2,047 nt — chosen as an odd number and the maximum length supported by the bioLM API
- **Half-width (HALF):** 1,023 nt on each side of the focal SNP
- **Center:** The SNP position is placed exactly at index 1,023 (0-based) within the window
- **Two sequences per SNP:** A reference (`other_allele`) sequence and an alternate (`counted_allele`) sequence, differing only at the center base
- **DNABERT feature:** The embedding difference (alternate – reference) was then used as the per-SNP representation

Windows were excluded if they fell within 1,023 nt of a chromosome boundary, had a center-base mismatch with the B73 v5 reference, or contained non-ACGT characters.