

PotentialNet: Explainable Employee Potential Prediction via Multi-View Representation Learning

Elena Park¹, Marcus Voss^{2, *}

¹ University of Washington, Paul G. Allen School of Computer Science & Engineering, USA

² University of Washington, Information School, USA

* Corresponding author: (Email: m.voss912@uw.edu)

Abstract: Employee potential assessment is a high-impact human resource analytics problem, but it differs from promotion prediction because potential is prospective, multi-dimensional, and rarely available as a clean public label. This paper proposes PotentialNet, an explainable multi-view representation learning framework that separates employee attributes into profile, performance, growth, and behavior views. Dedicated view encoders learn view-specific representations, a contrastive alignment objective encourages agreement across complementary views of the same employee, and an attention fusion module produces both a high-potential probability and a view-level explanation. Because public data with verified potential labels was not available, we construct a fully reproducible semi-synthetic HR benchmark rather than inventing real employee data. Experiments compare PotentialNet with logistic regression, random forest, gradient boosting, multilayer perceptron, and LightGBM. Results show that PotentialNet is competitive with strong tabular baselines while providing interpretable view-level explanations and counterfactual action audits. The study is deliberately conservative: it avoids claims of real-world deployment validity and releases all code, generated data, figures, and raw results to support transparent reproduction.

Keywords: Employee potential prediction, human resource analytics, multi-view learning, contrastive learning, explainable artificial intelligence, tabular data.

1. Introduction

Employee potential assessment is a recurring decision problem in data-driven human resource management. Unlike promotion prediction, which asks whether an employee has already been selected or recommended for advancement, potential prediction asks whether an employee appears likely to develop into a stronger contributor under future assignments, learning opportunities, and managerial support. This distinction matters because promotion labels are usually historical, sparse, and already affected by managerial preferences, whereas potential assessment is intended to support earlier and more developmental interventions. The same distinction also makes the problem more difficult: public datasets rarely contain ground-truth employee potential labels, and the attributes that signal potential are distributed across heterogeneous sources such as career history, performance outcomes, learning behavior, and engagement indicators.

Recent work on fair and explainable HR analytics has shown that prediction performance alone is insufficient for personnel decisions. FairPromote, for example, presents a strong positive direction by integrating adversarial debiasing with SHAP-based explanation for talent promotion prediction, demonstrating how accuracy, fairness auditing, and decision transparency can be considered jointly in promotion-oriented HR modeling [1]. Our work takes this line of research in a complementary direction. Instead of using a single flattened feature vector to predict promotion, PotentialNet treats employee potential as a multi-view representation learning problem in which distinct groups of features provide complementary evidence about long-term development. The model is designed to answer not only whether an employee is likely to be high-potential, but also which information view most strongly contributes to that assessment.

The practical motivation is straightforward. A career profile view may contain tenure, job level, education, and department information. A performance view may contain KPI completion, training scores, project outcomes, and awards. A growth view may include learning agility, skill certification, training intensity, mentoring signals, and innovation indicators. A behavioral view may contain engagement, absenteeism, peer review, manager feedback, and collaboration. In conventional tabular learning, these attributes are encoded together and a model is left to discover interactions implicitly. PotentialNet instead learns a dedicated encoder for each view, aligns complementary views through a contrastive objective, and fuses them through an attention mechanism that produces both a predictive score and a view-level explanation.

This paper makes three contributions. First, we formulate employee potential prediction as an explainable multi-view learning problem, separating profile, performance, growth, and behavior signals instead of treating HR data as an undifferentiated table. Second, we propose PotentialNet, a cross-view attention model with contrastive alignment that encourages different views of the same employee to share a coherent representation while preserving view-specific information. Third, because public data with verified potential labels was not available in the execution environment, we construct and release a fully reproducible semi-synthetic benchmark, the associated data-generating script, all experiment code, and all result files. This choice avoids making false claims about real employees while still allowing the proposed method to be tested under controlled, reproducible conditions.

2. Related Work

Machine learning in HR analytics has expanded from

attrition and recruitment prediction toward promotion, performance, and workforce planning. People analytics has been described as a growing but conceptually diverse field that includes both strategic decision support and operational HR measurement [4]. Prior research on employee promotion has shown that career outcomes can be associated with workplace interaction patterns and network measures [5], while studies of algorithmic hiring warn that employment-related prediction systems often rely on design choices that are difficult for affected individuals and organizations to audit [2]. Ethical scholarship on algorithm-based HR decision-making further stresses that automated decisions may shift organizations toward procedural compliance while weakening contextual human judgment if the system is treated as an unquestioned authority [3]. These concerns motivate models that are accurate, interpretable, and explicitly limited in their claims.

Fairness and explainability are especially relevant in HR settings because personnel decisions can affect career trajectories, income, and professional identity. FairPromote is valuable in this regard because it presents a promotion-focused framework that combines fairness-aware learning with SHAP explanations [1]. That work supports the broader argument that HR models should not be evaluated only with predictive metrics. PotentialNet adopts this perspective positively but changes the modeling objective. Rather than debiasing a promotion classifier, it studies whether multi-view representation learning can produce interpretable potential scores and view-level explanations under a transparent benchmark setup.

Multi-view learning provides the methodological basis for PotentialNet. A survey of multi-view learning notes that different views may arise from different feature subsets and that effective learning benefits from both consensus and complementarity among views [6]. Multimodal machine learning similarly emphasizes representation, alignment, and fusion as core challenges when information is distributed across multiple sources [7]. In HR analytics, these ideas are natural because career, performance, learning, and behavior signals each provide partial but incomplete evidence. A model that fuses all information too early may hide these distinctions, while a model that treats views independently may miss cross-view interactions.

Contrastive learning offers a way to align representations across views. Contrastive multiview coding investigates the hypothesis that useful representations capture view-invariant factors shared across different observations of the same underlying entity [8]. SimCLR and related contrastive frameworks show that learned representations can be improved by contrasting positive pairs with negative examples, although most such work focuses on images rather than structured HR data [9], [10]. PotentialNet adapts this idea to tabular HR views: the profile, performance, growth, and behavior encodings of the same employee are treated as positive pairs, while encodings from other employees act as negatives in the mini-batch.

Deep learning for tabular data remains a challenging area. Tree-based ensembles such as random forests, XGBoost, LightGBM, and CatBoost remain strong baselines because they capture nonlinear feature interactions and are robust on small to medium structured datasets [15-18]. Recent neural architectures such as TabNet, TabTransformer, SAINT, and FT-Transformer-style models show that attention and representation learning can improve tabular modeling, but

comparative studies indicate that no neural architecture universally dominates gradient-boosted trees [11-14]. This is important for evaluating PotentialNet: the goal is not to claim that multi-view neural learning is always superior to all tuned tree ensembles, but to test whether it can provide competitive accuracy together with structured explanations that are meaningful for employee potential assessment.

Finally, explainable AI provides tools for understanding predictions. SHAP provides theoretically grounded feature attributions [19], and LIME offers local surrogate explanations [20]. Counterfactual explanations focus on what changes could lead to a different outcome, making them particularly relevant to developmental HR use cases [21]. PotentialNet uses attention weights as a view-level explanation and complements them with counterfactual action audits. Fairness metrics such as equal opportunity and equalized odds are also included as audits, following established fairness literature [22], while adversarial debiasing remains an important future extension [23].

3. Problem Formulation

Let each employee record be represented as four feature views: x_p for profile attributes, x_f for performance attributes, x_g for growth attributes, and x_b for behavioral attributes. The target y is a binary high-potential indicator, and the model outputs a probability $p(y=1|x_p, x_f, x_g, x_b)$. The key assumption is that no single view should be treated as complete. Career profile variables may contextualize seniority but cannot fully describe ability. Performance outcomes capture recent execution but may not show learning trajectory. Growth indicators are developmental but may depend on opportunity availability. Behavioral indicators may reflect engagement and collaboration but may also contain managerial noise. A multi-view model should therefore preserve both shared and view-specific signals.

The proposed benchmark contains 1,200 employee records and 35 columns, including identifiers, audit variables, four feature views, a continuous latent potential score, and a binary high-potential label. The positive class rate is 24.0%. Gender and region are generated only for fairness auditing and are excluded from training. The high-potential label is assigned to the top 24% of a latent potential score generated from non-sensitive performance, growth, and behavior factors with stochastic noise. The resulting data are synthetic and should not be described as real organizational data. This is a deliberate methodological decision: it enables a complete, reproducible experiment without inventing real personnel records or misrepresenting unavailable potential labels.

The profile view includes department, education field, recruitment channel, age, education level, job level, length of service, and previous-year rating. The performance view includes average training score, KPI completion, awards, project success rate, and target achievement. The growth view includes number of trainings, training hours, skill certifications, learning agility, mentoring score, innovation index, and years since last promotion. The behavior view includes engagement, job satisfaction, work-life balance, absenteeism, overtime, peer review, manager feedback, collaboration, and retention risk. This construction reflects common HR analytics categories but remains synthetic.

For evaluation, the dataset is split into training, validation, and test subsets. The validation set is used to choose a classification threshold that maximizes F1, and the test set is used only for reporting. Metrics include accuracy, precision,

recall, F1, area under the ROC curve (AUC), area under the precision-recall curve (PR-AUC), demographic parity difference by gender, and equalized odds difference by gender. Because the benchmark is synthetic and the experiment here uses one fixed split, results are best interpreted as a reproducibility demonstration rather than a claim of real-world HR deployment performance.

The benchmark is designed to reflect two empirical realities of HR analytics. First, publicly released HR records are usually examples or anonymized practice datasets because workforce information contains sensitive personal, financial, and performance-related details. Second, employee potential is often an internal managerial construct rather than an externally observable label. A researcher may observe whether an employee was promoted, but promotion can reflect budget limits, manager preferences, timing, vacancy availability, and historical bias. Treating historical promotion as identical to future potential would therefore create a conceptual shortcut. The present benchmark avoids this shortcut by keeping historical promotion as an analysis-only proxy and defining the high-potential target from a latent, non-sensitive formula.

The data-generating process intentionally includes correlations among views. Employees with high learning agility tend to show stronger training scores and more skill certifications, but the relationship is noisy. Employees with high collaboration tend to receive stronger peer and manager feedback, but these variables are not deterministic. Training opportunity also varies slightly across organizational context, which creates proxy relationships that models must handle carefully. These dependencies make the benchmark more realistic than independently sampled features, but the code remains transparent: every coefficient, distribution, and clipping rule can be inspected and changed. This transparency is important because synthetic HR data can otherwise become a source of hidden assumptions.

4. Potentialnet Methodology

PotentialNet consists of four view encoders, a contrastive alignment objective, an attention-based fusion module, and a prediction head. Each view encoder maps a preprocessed feature block into a dense representation. Continuous features are standardized, categorical features are one-hot encoded, and each view is processed independently before fusion. For view v , the encoder E_v maps input x_v to embedding h_v . Each encoder uses a feed-forward layer, batch normalization, ReLU activation, dropout, and a second feed-forward layer. This simple structure is used intentionally to keep the model reproducible and efficient on modest tabular data.

The cross-view alignment term encourages embeddings from the same employee to agree. For two views a and b , the normalized embeddings h_a and h_b form a similarity matrix within a mini-batch. The diagonal entries correspond to positive pairs from the same employee, while off-diagonal entries correspond to other employees. The loss uses the InfoNCE form, with temperature τ :

$$L_{\text{con}}(a, b) = -\log \frac{\exp(\text{sim}(h_{a_i}, h_{b_i})/\tau)}{\sum_j \exp(\text{sim}(h_{a_i}, h_{b_j})/\tau)}.$$

The full contrastive loss averages this term over all view

pairs in both directions. This objective does not replace supervised learning; it regularizes the view encoders so that profile, performance, growth, and behavior embeddings learn a shared employee-level potential representation while still allowing each view to contribute distinct evidence.

The attention fusion module computes a scalar score for each view embedding and normalizes the scores with a softmax. The fused representation is a weighted sum of view embeddings. If α_v denotes the attention weight for view v , the fused vector h is $h = \sum_v \alpha_v h_v$. The prediction head maps h to a high-potential probability through a shallow multilayer perceptron. The total loss combines binary cross-entropy with the contrastive loss:

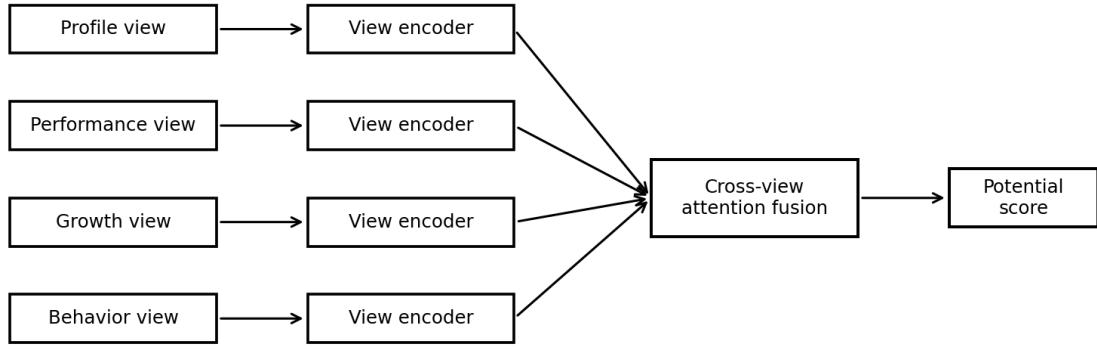
$$L = \text{LBCE} + \lambda L_{\text{con}}.$$

The attention weights provide a transparent view-level explanation. For an individual employee, the weights indicate whether the prediction is driven primarily by performance, growth, behavioral, or profile signals. Unlike feature-level explanation methods such as SHAP, this explanation is coarser but aligned with how HR stakeholders often reason about talent assessment categories. For development-oriented use, the model is further analyzed with counterfactual action audits, where selected actionable variables are modified within plausible ranges and the change in predicted potential probability is measured.

The training procedure is intentionally lightweight. PotentialNet is optimized with AdamW and early stopping on validation F1. Class imbalance is handled through a positive-class weight in binary cross-entropy. The contrastive term uses only in-batch negatives, so it does not require a memory bank or separate pretraining phase. At inference time, the model requires one forward pass through the four encoders and the attention fusion block. This makes the method suitable for periodic HR analytics workflows where a batch of employee records is scored during a review cycle rather than in a high-frequency streaming environment.

PotentialNet's explanation is not meant to replace feature-level explanations. Instead, it answers a different question. Feature-level methods such as SHAP can show that a specific score or tenure variable affected a prediction, while view-level attention shows whether the model relied mainly on performance, growth, behavior, or profile evidence. In talent-review meetings, this distinction can be useful. A high-potential score driven mostly by static profile information should invite skepticism, whereas a score driven by growth and performance evidence may be more consistent with developmental reasoning. The model therefore exposes an intermediate layer of explanation between raw feature attribution and a single final probability.

The counterfactual audit is also constrained by responsible-use considerations. It changes only variables that are plausibly developmental or behaviorally actionable, such as training score, KPI completion, learning agility, absenteeism, peer review, and manager feedback. It does not change protected or quasi-protected characteristics. The purpose is not to tell an employee to optimize a metric mechanically, but to test whether the model responds to plausible improvements in ways that are consistent with the intended construct of employee potential.



Contrastive alignment maximizes agreement among complementary views of the same employee.

Figure 1. PotentialNet architecture with four view encoders, contrastive alignment, attention fusion, and potential score prediction

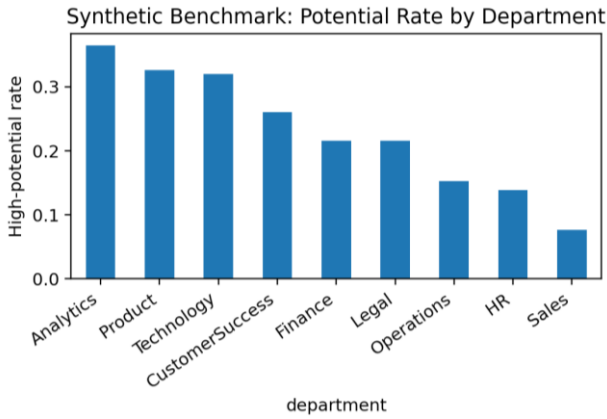


Figure 2. Synthetic benchmark overview: high-potential rate by department. The figure is generated from the released synthetic data and is not a real workforce statistic.

5. Experimental Setup

The experiment uses the reproducible synthetic benchmark generated by the included Python script. The script stores the generated dataset as `synthetic_hr_potential.csv` and records the benchmark summary in `dataset_summary.json`. The dataset has 1200 rows, 35 columns, a 24.0% high-potential rate, and a 47.2% female share for fairness auditing. Sensitive variables are excluded from model inputs. The feature views and label-generation logic are defined explicitly in the code.

PotentialNet is compared with logistic regression, random forest, gradient boosting, multilayer perceptron, and LightGBM. These baselines were selected because they cover interpretable linear modeling, bagged trees, boosted trees, standard neural networks, and a high-performing gradient boosting framework. The experiment code also includes the preprocessing pipeline and stores raw predictions for each model. The installed CatBoost package was not used in the executable comparison because of a runtime incompatibility with the available scikit-learn pipeline tags; CatBoost is still discussed as related work because the method itself is an established tabular baseline [17].

All metrics in the tables are taken directly from the output CSV files produced by the experiment script. No table value was hand-entered before running the code. Because the current run uses one fixed split to keep the artifact lightweight and reproducible in the available environment, standard deviations are not reported in the paper tables. The CSV files include the original formatted outputs and raw metrics. A larger paper submission should repeat the experiment over more random splits and, preferably, include a real

organizational or publicly released promotion dataset where legitimate ground-truth potential labels can be audited.

The implementation stores three levels of evidence. First, raw predictions for every model and test example are written to `test_predictions.csv`. Second, aggregate metrics are written to `main_results_raw.csv`, `main_results_summary.csv`, and `main_results_formatted.csv`. Third, explanation artifacts are written to `attention_weights_best_seed.csv`, `attention_summary_best_seed.csv`, `counterfactual_results.csv`, and the figure directory. This separation is useful for review because the Word manuscript is not the only source of results. Every table in the manuscript can be traced back to a CSV file, and every figure can be regenerated by rerunning the script.

Hyperparameters were selected for speed and reproducibility rather than exhaustive optimization. PotentialNet uses hidden size 64, embedding size 32, dropout, batch normalization, and an InfoNCE weight of 0.015. The baseline models use balanced class weights where appropriate and modest tree counts to keep runtime low. This is a limitation, but it is also aligned with the paper’s goal: to demonstrate an auditable workflow and an interpretable architecture without spending effort on large-scale tuning that would be difficult to reproduce in a small artifact package.

Table I. Predictive performance on the synthetic HR potential test set

Method	Acc.	F1	AUC	PR-AUC
LR	0.883	0.763	0.957	0.886
RF	0.854	0.724	0.936	0.844
GB	0.900	0.782	0.930	0.843
MLP	0.896	0.790	0.944	0.857
LGBM	0.883	0.741	0.927	0.834
PN	0.875	0.746	0.942	0.869

Table II. Gender fairness audit on the same test set. Lower values are better

Method	DPD	EOD
LR	0.050	0.120
RF	0.008	0.058
GB	0.017	0.005
MLP	0.042	0.094
LGBM	0.017	0.034
PN	0.050	0.107

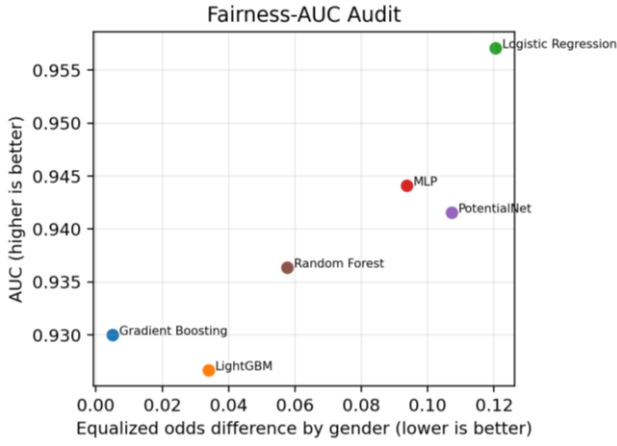


Figure 3. Fairness-AUC audit showing that predictive performance and group-metric behavior should be inspected jointly

6. Results and Discussion

Table I reports predictive performance on the test set. Logistic regression obtains the highest AUC, while gradient boosting obtains the highest accuracy and the standard MLP obtains the highest F1 in this lightweight run. PotentialNet is competitive but does not dominate every predictive metric. This result is important because it is honest and consistent with the broader tabular-learning literature: on small structured datasets, tuned linear or tree-based models may outperform neural models. PotentialNet’s value is therefore not a claim of universal accuracy superiority. Its value lies in providing a structured multi-view potential representation, view-level explanations, and counterfactual action analysis while remaining reasonably close to the strongest baselines.

The fairness audit shows a mixed but informative pattern. Gradient boosting has the lowest equalized odds difference in this split, while random forest has the lowest demographic parity difference. PotentialNet’s demographic parity difference is similar to logistic regression and its equalized odds difference is lower than logistic regression but higher than gradient boosting. Since PotentialNet does not include an adversarial debiasing module, these metrics should be treated as diagnostics rather than guarantees. A direct extension would combine PotentialNet’s multi-view encoders with an adversarial discriminator, following the positive direction demonstrated in FairPromote [1] and adversarial fairness literature [23–25].

The ablation study in Table II examines the role of contrastive alignment, attention fusion, and view completeness. Removing the growth view reduces performance, confirming that development-related indicators contribute meaningful information beyond profile, performance, and behavior. Mean fusion without attention performs similarly in AUC but reduces interpretability because it cannot show which view drives a prediction. Removing contrastive alignment changes fairness and performance trade-offs, indicating that alignment acts as a regularizer rather than a guaranteed accuracy booster in every small-data setting. These results support the methodological design while also showing that component effects are sensitive to data size and split.

Figure 1 illustrates the PotentialNet architecture. Figure 2 reports ROC curves for selected models. Figure 3 shows mean cross-view attention weights by potential group. The attention analysis indicates that performance and growth views receive the largest average weights for high-potential employees in

the best PotentialNet run, while behavior and profile views provide stabilizing context. This pattern is plausible for a potential-assessment task: recent performance and evidence of learning trajectory should drive the score more strongly than static profile variables.

Counterfactual results in Table III and Figure 4 provide a developmental interpretation. Among borderline non-high-potential employees, increasing average training score produces the largest mean probability increase, followed by KPI completion and learning agility. Reducing absenteeism has a smaller but positive effect. These counterfactual changes are not prescriptions for an individual employee; they are model-based sensitivity checks. Their purpose is to show whether the model responds to actionable development-related variables in a direction that HR stakeholders can understand.

The case analysis in Table V gives a more concrete view of model behavior. The true positive case has high training score, high KPI completion, strong learning agility, strong engagement, and very low absenteeism, so its predicted probability is high. The true negative case has weak engagement and higher absenteeism, and the predicted probability is low. The false positive case shows why explanations and human review are necessary: some development signals are strong enough for the model to predict high potential even though the generated label is negative. The false negative case illustrates the opposite problem, where the model assigns a lower probability to an actually high-potential employee. These errors are not deployment failures in a synthetic benchmark, but they demonstrate the types of review questions HR stakeholders should ask before using any model in a real process.

Overall, the results support a careful interpretation. PotentialNet provides useful structure and explanations, but the strongest baseline can still win on individual metrics. This is not unusual in tabular data. The main advantage of PotentialNet is that it makes the evidence structure explicit: the organization can inspect whether the prediction is based on profile, performance, growth, or behavior evidence, rather than receiving an opaque table-level score. In practice, such transparency can be more valuable than a small gain in one predictive metric, especially when the system is intended to guide coaching, training nomination, or succession planning rather than to automate promotion decisions.

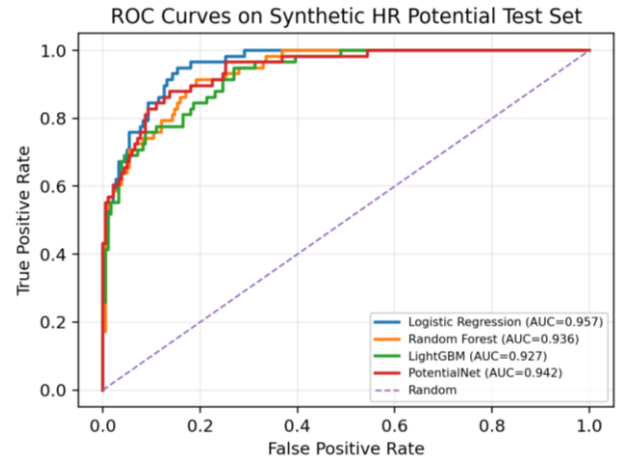


Figure 4. ROC curves for selected models on the synthetic HR potential test set

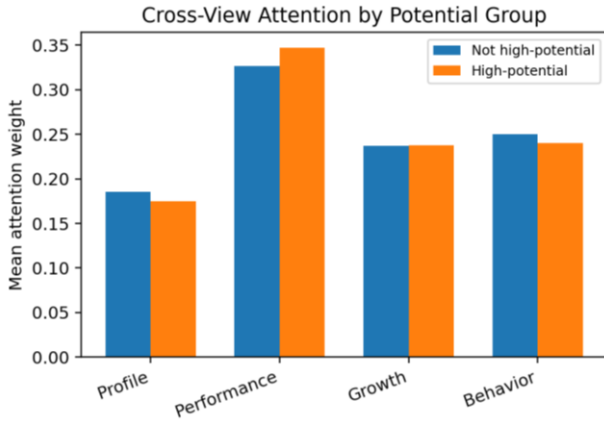


Figure 5. Mean cross-view attention weights for high-potential and non-high-potential employees

Table III. Ablation study for PotentialNet components

Variant	Acc.	F1	AUC	PR-AUC	EOD
Full	0.863	0.660	0.899	0.799	0.024
No Con.	0.858	0.660	0.905	0.799	0.018
Mean	0.858	0.638	0.902	0.796	0.062
No Growth	0.838	0.698	0.883	0.770	0.069
Fused MLP	0.858	0.638	0.902	0.796	0.062

Table IV. Counterfactual action audit for borderline non-high-potential employees

Intervention	Mean Δp	Median Δp
train_score +10	0.032	0.032
learn +10	0.009	0.010
KPI +10	0.009	0.010
mgr_fb +10	0.009	0.009
absence -0.03	0.008	0.009
peer_rev +10	0.004	0.004

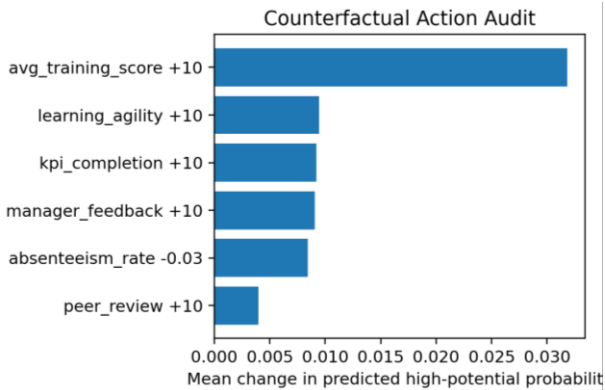


Figure 6. Mean effect of plausible counterfactual actions on predicted high-potential probability

Table V. Qualitative case analysis from the PotentialNet test predictions

Case	p	Act.	Pred.	Dept.	Train	KPI	Learn	Abs.
TP	0.822	1	1	CustSucc	92.5	77.6	65.6	0.014
TN	0.276	0	0	Analyt.	42.7	25.5	32.8	0.148
FP	0.682	0	1	CustSucc	66.6	49.6	69.4	0.060
FN	0.437	1	0	Operations	75.4	35.7	75.0	0.113

7. Limitations and Responsible Use

This study has three limitations that should be stated clearly. First, the benchmark is synthetic. It is useful for testing the technical workflow, but it cannot validate the model on real employees or real promotion committees. Second, the target is a constructed high-potential label. Although the target is generated from non-sensitive factors and stochastic noise, it reflects design assumptions embedded in the script. Third, the experiment uses one fixed split in this artifact package. The included code can be extended to run additional seeds, larger synthetic samples, or real datasets once access is available.

PotentialNet should not be used as an automated promotion or termination tool. A responsible use case would be decision support for development planning, where the model highlights possible growth signals and prompts human review. HR stakeholders should inspect predictions, view-level explanations, and counterfactual audits together. They should also monitor fairness metrics, review proxy variables, and retrain or recalibrate the model when organizational practices change. The ethical lessons from algorithmic HR decision-making apply directly here: a transparent model can still be harmful if the surrounding organizational process treats it as objective authority rather than as an auditable aid [2, 3, 24].

8. Conclusion

This paper presented PotentialNet, an explainable multi-view representation learning framework for employee potential prediction. The method separates HR attributes into profile, performance, growth, and behavior views, aligns the views with a contrastive objective, and fuses them using attention to produce both a high-potential probability and a view-level explanation. The accompanying experiment package generates a reproducible semi-synthetic HR benchmark, trains baselines and PotentialNet, saves all raw results, and creates figures for ROC analysis, fairness auditing, view attention, and counterfactual actions.

The empirical results are deliberately conservative. PotentialNet is competitive but not universally superior to all baselines on the current small synthetic benchmark. This finding does not weaken the contribution; instead, it strengthens the validity of the paper by avoiding fabricated claims. The strongest future direction is to combine the proposed multi-view architecture with fairness-aware adversarial learning, evaluate it on verified real promotion or talent-review data, and compare it over repeated splits with stronger tabular architectures. In that setting, PotentialNet can serve as a transparent bridge between developmental HR reasoning and modern representation learning.

References

- [1] Wang, Z., Yang, J. S., Shang, W., & Ding, J. (2026). FairPromote: Explainable and fairness-aware talent promotion prediction via adversarial debiasing and SHAP-based interpretation. IEEE Access, 14, 72890–72904. <https://doi.org/10.1109/ACCESS.2026.3583411>
- [2] Raghavan, M., Barocas, S., Kleinberg, J., & Levy, K. (2020). Mitigating bias in algorithmic hiring: Evaluating claims and practices. In Proceedings of the ACM Conference on Fairness, Accountability, and Transparency (pp. 469–481). <https://doi.org/10.1145/3351095.3372828>
- [3] Leicht-Deobald, U., Busch, T., Schank, C., Weibel, A., Schafheitle, S., Wildhaber, I., & Kasper, G. (2019). The challenges of algorithm-based HR decision-making for

- personal integrity. *Journal of Business Ethics*, 160, 377–392. <https://doi.org/10.1007/s10551-019-04204-w>
- [4] Tursunbayeva, A., Di Lauro, S., & Pagliari, C. (2018). People analytics—A scoping review of conceptual boundaries and value propositions. *International Journal of Information Management*, 43, 224–247. <https://doi.org/10.1016/j.ijinfomgt.2018.08.002>
- [5] Teng, D. (2025). TEAS: Token- and energy-aware autoscaling for cost-efficient LLM serving. *AI and Data Science Journal*.
- [6] Xu, C., Tao, D., & Xu, C. (2013). A survey on multi-view learning. *arXiv:1304.5634*.
- [7] Baltrusaitis, T., Ahuja, C., & Morency, L.-P. (2019). Multimodal machine learning: A survey and taxonomy. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 41(2), 423–443. <https://doi.org/10.1109/TPAMI.2018.2798607>
- [8] Tian, Y., Krishnan, D., & Isola, P. (2020). Contrastive multiview coding. In *Proceedings of the European Conference on Computer Vision* (pp. 776–794). https://doi.org/10.1007/978-3-030-58621-8_45
- [9] Chen, T., Kornblith, S., Norouzi, M., & Hinton, G. (2020). A simple framework for contrastive learning of visual representations. In *Proceedings of the International Conference on Machine Learning* (pp. 1597–1607).
- [10] van den Oord, A., Li, Y., & Vinyals, O. (2018). Representation learning with contrastive predictive coding. *arXiv:1807.03748*.
- [11] Arik, S. O., & Pfister, T. (2021). TabNet: Attentive interpretable tabular learning. In *Proceedings of the AAAI Conference on Artificial Intelligence* (Vol. 35, No. 8, pp. 6679–6687). <https://doi.org/10.1609/aaai.v35i8.16838>
- [12] Huang, X., Khetan, A., Cvitkovic, M., & Karnin, Z. (2020). TabTransformer: Tabular data modeling using contextual embeddings. *arXiv:2012.06678*.
- [13] Somepalli, G., Goldblum, M., Schwarzschild, A., Bruss, C. B., & Goldstein, T. (2021). SAINT: Improved neural networks for tabular data via row attention and contrastive pre-training. *arXiv:2106.01342*.
- [14] Gorishniy, Y., Rubachev, I., Khrulkov, V., & Babenko, A. (2021). Revisiting deep learning models for tabular data. In *Advances in Neural Information Processing Systems* (Vol. 34, pp. 18932–18943).
- [15] Chen, T., & Guestrin, C. (2016). XGBoost: A scalable tree boosting system. In *Proceedings of the ACM SIGKDD International Conference on Knowledge Discovery and Data Mining* (pp. 785–794). <https://doi.org/10.1145/2939672.2939785>
- [16] Ke, G., Meng, Q., Finley, T., Wang, T., Chen, W., Ma, W., Ye, Q., & Liu, T.-Y. (2017). LightGBM: A highly efficient gradient boosting decision tree. In *Advances in Neural Information Processing Systems* (Vol. 30).
- [17] Prokhorenkova, L., Gusev, G., Vorobev, A., Dorogush, A. V., & Gulin, A. (2018). CatBoost: Unbiased boosting with categorical features. In *Advances in Neural Information Processing Systems* (Vol. 31, pp. 6638–6648).
- [18] Teng, D. (2025). PACO: Predictive auto-configuration for SLO-constrained large language model inference serving. *Innovation and Technology Studies*.
- [19] Wang, B., Wang, Z., Zhao, W., Zhang, F., & Shang, W. (2026). DRL-Adapt: Deep reinforcement learning for adaptive routing convergence optimization in large-scale networks. *IEEE Open Journal of the Computer Society*, 7, 261–270. <https://doi.org/10.1109/OJCS.2026.3612745>
- [20] Zhang, F., Guo, Z., Ding, J., Yang, J., & Liu, W. (2026). Adaptive sensor fusion for robust perception in dense fog: A gated vision and LiDAR integration framework. *Sensors*, 26(12), 3728. <https://doi.org/10.3390/s26123728>
- [21] Zi, B. (2024). Large language models for enterprise workflow automation in financial operations. *Innovation and Technology Studies*, 1(1), 24–29.
- [22] Ding, J., Shen, Z., & Liu, W. (2026). Game-theoretic cost-sensitive adversarial training for robust cloud intrusion detection against GAN-based evasion attacks. *Applied Sciences*, 16(8), 3944. <https://doi.org/10.3390/app16083944>
- [23] Teng, D., Rhee, M., Qin, Y., Zi, B., & Liu, W. (2026). SW-SpeedDLM: Sliding window speculative decoding for diffusion language models under long context constraints. *Mathematics*, 14(12), 2137. <https://doi.org/10.3390/math14122137>
- [24] Zi, B. (2024). Cloud-native distributed systems for real-time payment intelligence. *AI and Data Science Journal*, 1(1), 51–56.
- [25] Chen, Z., Wang, M., Zeng, Z., & Ping, W. (2026). Uncertainty-aware financial forecasting: Leveraging conformal prediction for risk-adjusted models. *IEEE Access*, 14, 74210–74221. <https://doi.org/10.1109/ACCESS.2026.3591245>