

LMIC Demonstration Report

Benchmarking Online Job Postings to the
Job Vacancy and Wage Survey to Improve
Vacancy Estimates

June 2025





The Labour Market Information Council (LMIC) is a pan-Canadian non-profit that produces accessible, evidence-based insights on Canada's labour market. Through research, collaboration, and data innovation, LMIC supports informed decision-making by governments, employers, workers, and educators. Our work helps bridge information gaps, improve labour market outcomes, and strengthen Canada's workforce development ecosystem.

To learn more about our research and initiatives, visit lmic-cimt.ca or contact us at info@lmic-cimt.ca.

Authors

Anne-Lore Fraikin

Ibrahim Abuallail

This report was authored by Anne-Lore Fraikin, formerly of LMIC, and Ibrahim Abuallail, an external consultant engaged through a competitive RFP process.

How to cite this report

Abuallail, I & Fraikin, A.-L. (2025). Benchmarking Online Job Postings to the Job Vacancy and Wage Survey to Improve Vacancy Estimates: LMIC Demonstration Report. Ottawa: Labour Market Information Council (LMIC).

Information contained in this publication or product may be reproduced, in whole or in part, and by any means, for personal or public non-commercial purposes without charge or further permission, unless otherwise specified. Commercial reproduction and distribution are prohibited except with written permission from the Labour Market Information Council (LMIC). To obtain permission to reproduce any content owned by LMIC for commercial purposes, please contact communications@lmic-cimt.ca.

For more information visit lmic-cimt.ca



Table of Contents

Introduction	1
Contextualizing the data sources used for this demonstration	3
Demonstration	5
Conclusion	14
References	16



Introduction

Labour market information (LMI) and data science are advancing, creating new opportunities for planning and policy development in strategic workforce and post-secondary enrolment. Outside traditional government surveys, online job postings (OJP) are one of the most significant sources of LMI. LMIC has observed that across Canada, jurisdictions and institutions are increasingly using data from OJP to complement traditional surveys. Because OJP are more commonly used as a complementary source of data, understanding how to use these properly—that is, accurately, robustly, and with careful thought toward applicability—is critical. This is particularly true in times of rapid and dynamic labour market change.

OJP data scraped from employer websites or online job boards are sometimes used to complement job vacancy data because these provide direct insight into employer-demanded skills. However, a job posting is not the same as a job vacancy.¹

Traditional approaches to understanding job vacancies, like Statistics Canada's Job Vacancy and Wage Survey (JVWS), provide robust structure as well as strict and validated parameters for their measurements of job vacancies. This matters because accurately measuring job vacancies is crucial to understanding labour market trends, guiding workforce planning, and supporting labour market policy decisions.

While OJP and JVWS data can offer complementary insights, this report explores whether the JVWS data can be used to strengthen OJP-based vacancy estimates—moving the coordinated analysis beyond side-by-side comparison.

In this demonstration, LMIC presents a novel method for estimating job vacancies using OJP data from Vicinity Jobs, benchmarked against the JVWS for accuracy. Through machine learning models, weighting functions, and robust regression, we show how OJP data can be refined to produce more accurate estimates of job vacancies.

¹ The JVWS defines a job as “vacant” if it meets several conditions: it must be vacant to the reference date (first of the month) or set to become vacant during the month; there must be tasks to be carried out during the month for the job in question; and the employer must be actively recruiting outside the organization to fill the job.

Specifically, this demonstration:

- ▶ introduces a weighting function benchmarked to the JVWS to improve the representativeness of OJP
- ▶ implements a winsorization algorithm to address the influence of outliers in the data
- ▶ introduces a robust regression using Huber weights to minimize survey prediction errors
- ▶ explores a combined approach, testing how multiple methods—like applying winsorization alongside a weights adjustment function—can improve accuracy
- ▶ tests and applies several machine learning models to predict job vacancies from OJP data
- ▶ combines improved forecast accuracy with a weighting function to ensure both precision and representativeness

By applying these methods, this report demonstrates that the combination of winsorization and robust regression with Huber weights outperforms other methods: this approach reduced prediction errors by an average of 15% compared to the winsorization algorithm alone.

Why this matters

OJP are a widely used source of LMI across Canada. They offer high-frequency, granular data that can complement existing administrative datasets and surveys to provide deeper insight into labour market dynamics.

Governments, post-secondary institutions, and workforce development organizations increasingly rely on OJP data to supplement traditional LMI. However, unlike survey-based sources, OJP data typically lack rigorous sample selection criteria and complex weighting methods. This introduces misrepresentation risks, especially when interpreting vacancy trends. As a result, there is a growing need for guidance on how to use OJP data responsibly and accurately (Rosenbaum & Feor, 2020).

Conventionally, job vacancies in an economy are measured through employer surveys. In 2014, for this reason, the JVWS was launched in Canada to assess job vacancies by region and occupation. To the benefit of researchers, the JVWS includes short job descriptions, allowing analysts to match vacancies to National Occupation Classification (NOC) codes. However, while robust and reliable, the JVWS alone does not capture the full scope or speed of changes in the labour market.

Integrating OJP and JVWS data could produce more timely and detailed vacancy insights. Together, these sources could offer a new and complementary source of data, enhancing the tools already available for understanding the labour market.

It is important to note that OJP data, like all information sources, have limitations. For example, Vu et al. (2019) explored how sectoral overrepresentation and educational bias in OJP data can distort vacancy estimates. To deal with outliers and extreme values in OJP data, Evans et al. (2023) used a winsorization² algorithm to predict job vacancies in Australia. When working with UK data, Turrell et al. (2022) used a weighting function, computed from the weights of jobs within their sectors in the survey, to improve the representativeness of OJP.

By benchmarking Vicinity Jobs data against the JVWS, we tested whether advanced statistical methods—including machine learning and robust regression—can address the known challenges with OJP data.

We want to emphasize that this is a demonstration project, not a conclusive solution. Still, our work provides a foundation for more accurate, representative, and timely use of OJP data in vacancy measurement.

Contextualizing the data sources used for this demonstration

The work outlined in this report relies on two primary data sources: the JVWS and Vicinity Jobs. The JVWS provides structured, validated, employer-reported vacancy data, while the OJP data from Vicinity Jobs offers a more granular perspective on job postings. Together, these sources allow for a comparative assessment of job vacancy trends in Canada.

The Canadian Job Vacancy and Wage Survey

The JVWS was launched to enable researchers to identify vacancies by region and detailed occupation information. The survey targets businesses with at least one employee, identified from Statistics Canada's Business Register (BR)—a central repository that contains basic information about businesses in Canada. As of December 2022, the BR contained approximately 1.22 million businesses.

² A winsorization technique is a method that deals with outliers by eliminating extreme values in a dataset and replacing them with certain minimum or maximum values.

The JVWS employs a stratified random sampling approach, selecting 100,000 business locations for participation each quarter. After each reference period, Statistics Canada publishes detailed occupation and region-specific JVWS data (with a several-month lag).

While the JVWS offers high-quality, employer-reported data, it has limitations, including reporting lags and sampling constraints. To complement this dataset, this demonstration analyzes OJP.

Online job postings

OJP data offer rich information about job demand. These data are scraped from employer websites and online job boards. Unlike the JVWS, OJP data are not gathered through sampling. They reflect immediate employer demand. While this offers higher-frequency labour market insights, it comes with caveats around data interpretation and representativeness.

Vicinity Jobs data description

After cleanup and deduplication, the Vicinity Jobs dataset for this demonstration contained 16,781,882 job postings for the period between the start of 2018 and the end of 2023. These data were collected using a web-scraping tool that collects OJP from multiple online sources. Approximately half of the data are from employer websites, while the rest are from third-party, non-employer job posting sites, such as [Indeed.ca](https://www.indeed.ca) and [Monster.com](https://www.monster.com). All third-party sources are subject to verification processes to confirm employer legitimacy and job authenticity.

Most third-party data sources included require employers to pay to post, which may introduce a barrier for smaller or lower-wage employers. Additionally, government job boards like Job Bank impose strict submission requirements, including multi-step forms and employer verification steps.

To assess the representativeness of the data, we reviewed the distribution of job postings by NOC category. Across Canada during the 2018 to 2023 period, the first four NOC categories—which include occupations in business, finance, the natural and applied sciences, education, and law—accounted for 46.5% of all job postings based on OJP data. By comparison, the JVWS shows that these same categories represented only 35.1% of total vacancies. The difference highlights a higher representation of digitally intensive occupations in OJP data (Abuallail & Vu, 2022).

Conversely, the remaining NOC groups—including trades, transport and equipment operators and related occupations, manufacturing and utilities occupations, and occupations related to natural resources, agriculture, and production—are under-represented in OJP compared to their JVWS share of 64.9%.

Limitations and potential biases

As with any dataset, OJP have some caveats and potential biases. One caveat is the potential for duplication: employers may post the same job opening on multiple websites to reach a wider audience. While web scrapers use techniques to reduce duplication (like identifying identical text even with formatting variations), complete elimination is not guaranteed. There are multiple ways to detect duplicates, and the chosen approach may vary by OJP provider. Many decide that a combination of methods achieves optimal results.

OJP can effectively attract talent for employers, but use of this tactic is not necessarily uniform across occupations. Some employers still use traditional forms of job advertising, such as word of mouth or physically posted “Help Wanted” signs. OJP data cannot capture these. For example, LMIC found that, when comparing JVWS to Vicinity Jobs data, OJP are more likely to represent positions requiring university degrees and less likely to represent positions that emphasize on-the-job training (Rosenbaum & Feor, 2020).

Additionally, OJP data can contain geographical differences, with certain locations showing higher proportions of jobs than others. Although these figures can fluctuate over time and across regions, relative comparisons remain useful for providing context and supporting interpretation.

Demonstration

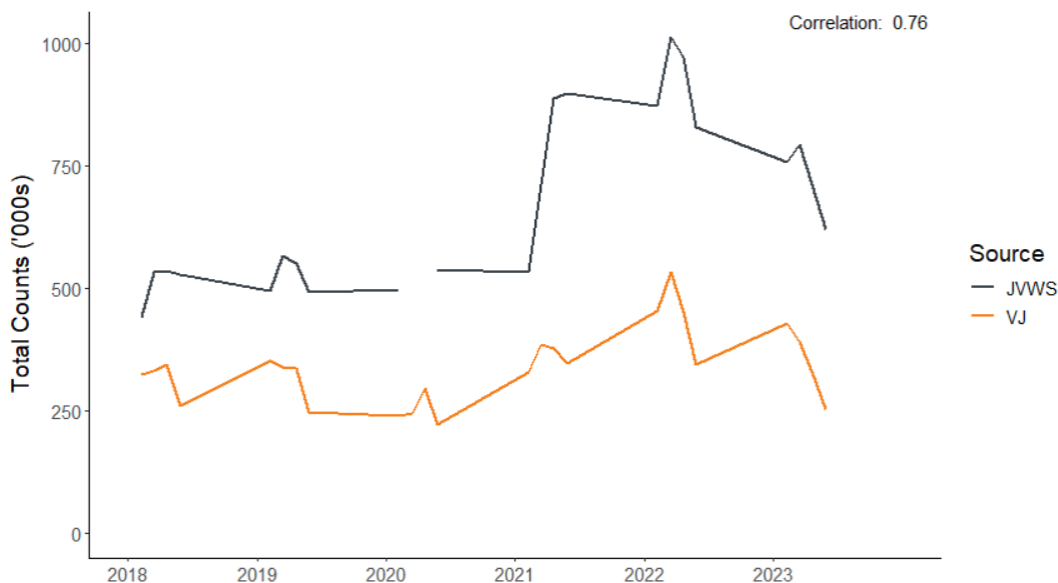
This section outlines how we benchmark and improve vacancy estimates from OJP data. We begin by comparing OJP and JVWS data, then test a series of models. The goal is to evaluate each model’s effectiveness in reducing prediction errors and improving representativeness.

Comparing online job postings with survey data

Our first step is to compare the total stock of OJP data (gathered by Vicinity Jobs) with the JVWS vacancy data. Each year, the OJP count is consistently lower than the JVWS total. Several factors could cause this, including employers opting not to post all job openings online.

Despite count differences, there is a positive correlation of 0.76 between the Vicinity Jobs and the JVWS data.

Figure 1: Total quarterly stock of JVWS job vacancies and OJP job postings



Note: The JVWS data are missing two values. Due to the COVID-19 pandemic, no data were collected during the second and third quarters of 2020.

After comparing the counts, we apply various machine learning models to assess the differences between Vicinity Jobs and JVWS data and use the results as a benchmark for comparing the performance of our improved algorithms (refer to Figure 2). Performance is assessed using root mean squared errors (RMSE) on a test sample.

We tested three machine learning models:

► **A tuned random forest model**

This model is well-suited to handling large datasets. It builds multiple decision trees and averages their predictions, which improves accuracy and reduces overfitting.

► **A gradient boosting model (GBM)**

This model sequentially combines multiple simple models, corrects prior errors, and captures important and subtle patterns in data.

► **A support vector machine (SVM)**

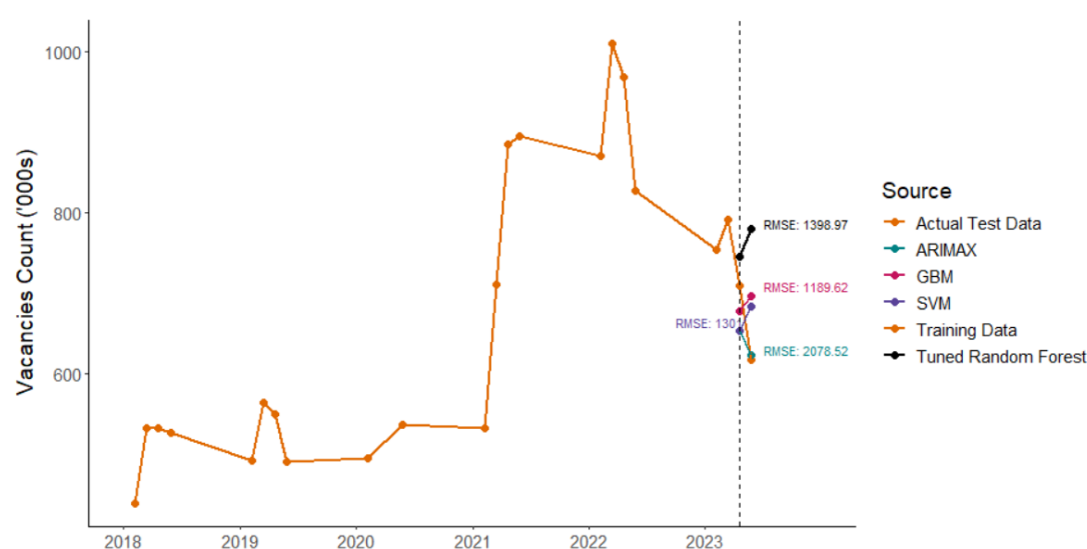
This model identifies the optimal hyperplane. In SVM, the optimal hyperplane separates and robustly classifies data. By doing so, the model can maximize the margin between data points of different classes, improving overall prediction accuracy.

We chose these models for their ability to handle diverse datasets and capture complex relationships. Each also balances accuracy, robustness, and computational efficiency with respect to the goal of improving job vacancy forecasts using both JVWS and OJP data.

For a comprehensive comparison of different machine learning models, see Hassan et al. (2018) and Osisanwo et al. (2017).

We will return to these machine learning models later in the report. First, we test and explore the performance of our algorithms.

Figure 2: Comparison of model predictions



Note: This figure shows a comparison of model predictions. The dashed vertical line indicates the first period of the two-period forecasts.

Algorithms and results

After testing and developing our benchmarks using machine learning models, we proceeded to develop our algorithm. Our approach incorporates elements of both the Turrell et al. (2022) and Evans et al. (2023) algorithms, aiming to modify parameters and enrich these approaches by introducing a robust regression with Huber weights and applying multiple approaches simultaneously.

Note that Vicinity Jobs vacancy data represent a **flow** variable (measured over a period of time), while JVWS vacancy data reflect a **stock** variable (measured at a specific point in time).

We need to measure the total number of job postings active during a specific month, not the number of job vacancies posted per month. As stated earlier, these are not definitively the same.

To find the total number of job postings active during a specific month, we convert the OJP flow variable into a stock variable. To do this, we first determine a representative figure for the average time it takes for a posted job offer to be filled. Using U.S. data, The Josh Bersin Company (2023) found that the average job posting is filled within 44 days. Given the absence of similar data for Canada, we assume the Canadian data would be similar, and aggregate the OJP data over the 44 days before the end of each quarter. The aggregation function is:

$$VJ_q = \sum_{d \in q} (VJ_d)$$

In this function, d represents a day and q is specified as a particular quarter of the JVWS survey represents the 44 days preceding the end of the quarter.

Weighting and scaling function

Next, we introduce a weighting function that adjusts vacancy distributions to align more closely with survey-based benchmarks from the JVWS. To achieve this, we consider the weight of each job vacancy from the JVWS in the total stock of vacancies for the quarter. Our weighting function, then, becomes:

$$W_{i,q} = \frac{(V_{i,q})^{jvws}}{V_q^{jvws}}$$

In this function, i stands for a NOC code, q stands for the quarter, $(V_{i,q})^{jvws}$ stands for the vacancies from the JVWS survey for a given NOC code in a given quarter, and V_q^{jvws} is the sum of vacancies for all nine NOC codes within the same quarter.

To incorporate the weights, we create a set of matched NOCs for all quarters between Vicinity Jobs and JVWS and rescale the distribution of the OJP in proportion to their relative weights. In other words, the scaling factor equals $(1 + w_{i,q})$. This ensures better alignment with the JVWS data while preventing overfitting. The adjustment also accounts for the under-representation of certain jobs and the overall gap between vacancy measurements from JVWS and OJP.

We also examined the effect of the matched and weighted approach on improving the distribution of the OJP data compared to the JVWS data. We found this had little

effect on the aggregate number of vacancies measured from OJP. We used scaled Vicinity Jobs data to calculate the error in absolute value (the absolute difference in mean percentage changes between consecutive quarters in Vicinity Jobs data and the same metric for JWVS). We then compared the error in absolute value with the error that would occur without applying the weight and scaling function. The reduction in average error was around 0.07%.

Winsorization algorithm

Our next step tests a winsorization algorithm inspired by Evans et al. (2023). Their algorithm uses quarter-to-quarter changes across all job posting sources to predict vacancies, weighing each source's signal according to the change in the source's count of postings.

The size of a job posting source can influence aggregate job vacancy estimates, creating several potential issues. Consider, when a single large source changes its number of postings, the overall vacancy measurement can shift, indicating a market change. However, this shift may simply represent a policy change at the job posting source—for example, changes in scraping policies or job posting strategies rather than genuine changes in the number of job vacancies.

To implement this algorithm, we consider the percentage changes in job postings from $q - 1$ to q . For each of these periods, the OJP data would be an aggregated set based on the aggregation function. In contrast, the JWVS would be the released vacancy level per NOC job at that specific quarter.

The first step is to split the dataset by job posting sources and identify sources $j=1,...,J$ with a certain cut-off, k , for postings per source. Then we compute the percentage change between $q - 1$ and q for the OJP per source:

$$\Delta p_{jq} = \frac{p_{jq} - p_{j,q-1}}{p_{j,q-1}}$$

We then apply the three-step winsorization described in Evans et al. (2023), testing different values of the cut-off $k\%$. The final step involves measuring the mean of the winsorized values.

To select the cut-off parameter (k), we apply a simulation that tests different possible values and finds the value that minimizes the mean absolute prediction error over the entire dataset. Based on observations from these tests, we chose a cut-off of 20%. This is higher than the optimal cut-off in Evans et al. (2023), which was around 10%. This parameter is highly data-dependent.

We compare the absolute value errors across all quarters when using a 5% winsorization algorithm versus a 20% algorithm. With the 20% winsorization cut-off level, the aggregate average error is reduced by around 27% across all quarters versus the 5% cut-off level.

Robust regression with Huber weights

The next step in developing our algorithm is to apply a robust regression with Huber weights to the winsorized data using the M-estimation method (Huber, 1967).

Robust regression methods are widely used to handle outliers in data. For instance, they may be applied for households, where outliers in consumption and earnings can have an influential impact (Gorodnichenko & Peter, 2007). They are also used extensively to control for outliers and influential observations in surveys, such as firm expectations on macroeconomic variables (Coibion et al., 2018).

Unlike ordinary least squares, which minimize least-squared regression errors, the M-estimation method minimizes a generalized function of the regression errors. More precisely, if we define a general objective function of the errors $S = \sum H(\varepsilon_i)$ (which is equal to ε_i^2 for ordinary least squares), and take the derivative with respect to a parameter β and set it equal to 0, we calculate:

$$\sum_{i=1}^n \frac{\partial H}{\partial \varepsilon_i} x_{ki} = 0$$

We can then rewrite this equation as:

$$\sum_{i=1}^n w_i \varepsilon_i x_{ki} = 0$$

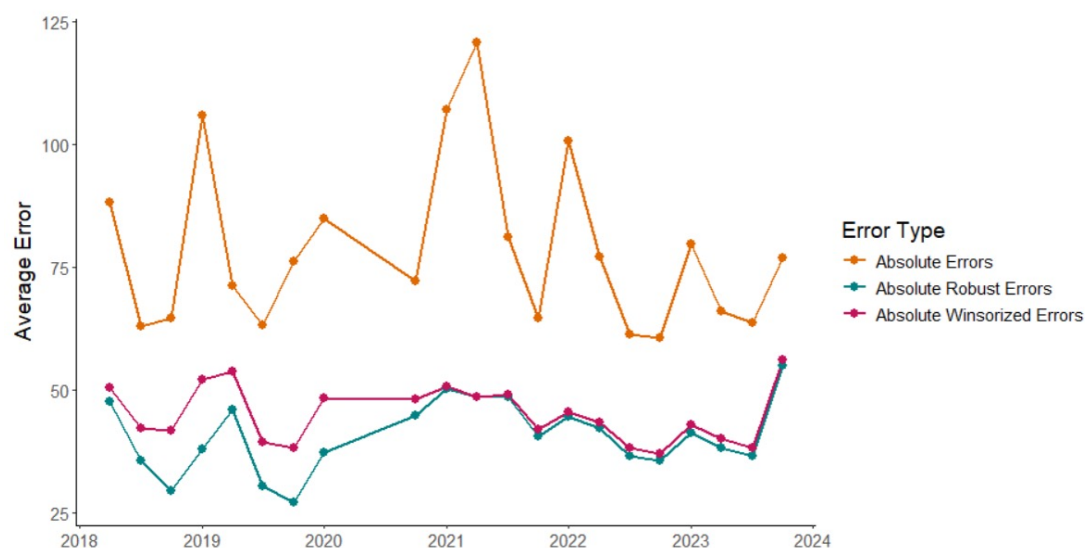
This is a weighted linear regression. Therefore, we can start by making a guess on w_i , fit the regression, and recalculate a new w_i . We repeat this process until we reach convergence. In a Huber M-estimation, the function $H(\varepsilon)$ takes the form:

$$H_{\varepsilon} = \begin{cases} \varepsilon^2/2 & \text{for } |\varepsilon| \leq \psi \\ \psi|\varepsilon| - \psi^2/2 & \text{for } |\varepsilon| > \psi \end{cases}$$

We run the robust regression on a constant, with the dependent variable being the percentage change in the number of job postings from one quarter to the next after winsorization. This allows us to get a robust measure of central tendency to the percentage movements of jobs across all sources from one quarter to another. Figure 3 shows how the robust regression approach produces fewer errors than the winsorization-only algorithm due to the adjustment to all outliers identified through the Huber function.

We retain all calculations and errors at the NOC level, allowing for further analyses at the level of a certain job or broad occupational category. By implementing the robust regression on the winsorized changes, the average absolute value error of prediction across all quarters is reduced by an additional 6% compared to the average errors calculated using the winsorization-only approach.

Figure 3: Average error by quarter (winsorized and robust)



Note: The average error for each quarter is calculated using the absolute value distance between the measure of vacancies from OJP and the same measure from the JVWS.

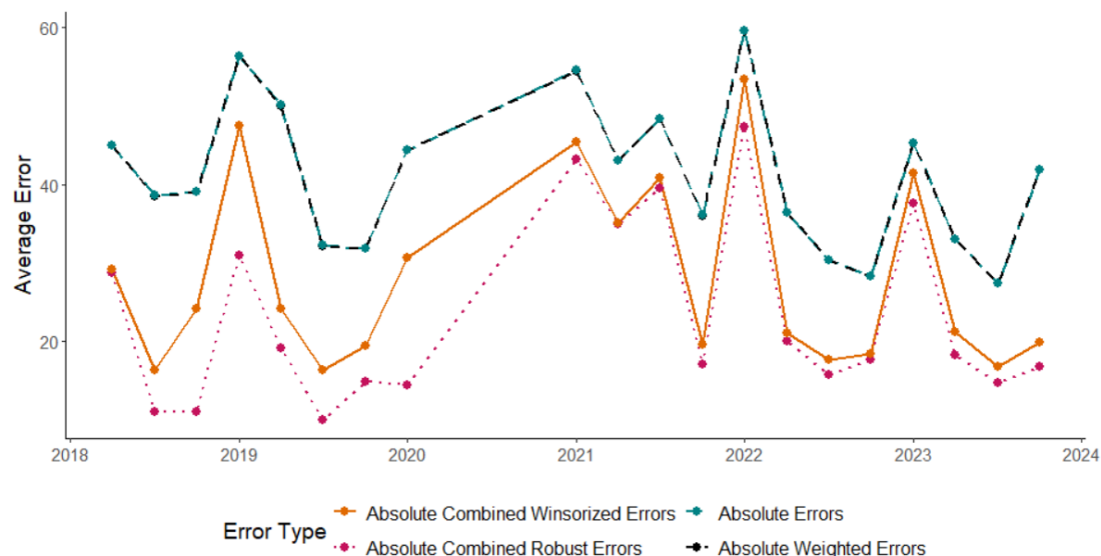
Combined approaches

The next step in developing our algorithm combines multiple adjustment methods in an effort to improve both the representativeness and predictive accuracy of OJP data. We start by combining our weighting function with the winsorization approach. The goal is to ensure that, while we tackle the issues of potential duplication and outsized effects of larger job posting sources, we also account for biases toward certain jobs.

To create the combined approach, we aggregate the data to find the number of job postings that correspond to the winsorized percentage changes. We apply the weighting function to get the absolute value errors of the winsorized-weighted algorithm. This combined approach leads to further, large improvements in the prediction errors: a reduction of around 32% compared to the raw data.

To create a winsorized-weighted-robust algorithm, we apply the robust regression method to the winsorized percentage changes. This results in an additional 15% reduction in absolute value errors over all quarters. Figure 4 compares the two combined algorithms with the weighting function approach as well as the original errors. To test whether this improves estimates, we use a Diebold-Mariano test. This test shows a statistically significant improvement in errors in the combined algorithms approach compared to raw data, with a P value close to 0 (see Table 6 in the appendix).

Figure 4: Comparison of combined algorithm errors with weighting function errors and original errors



Machine learning forecasts

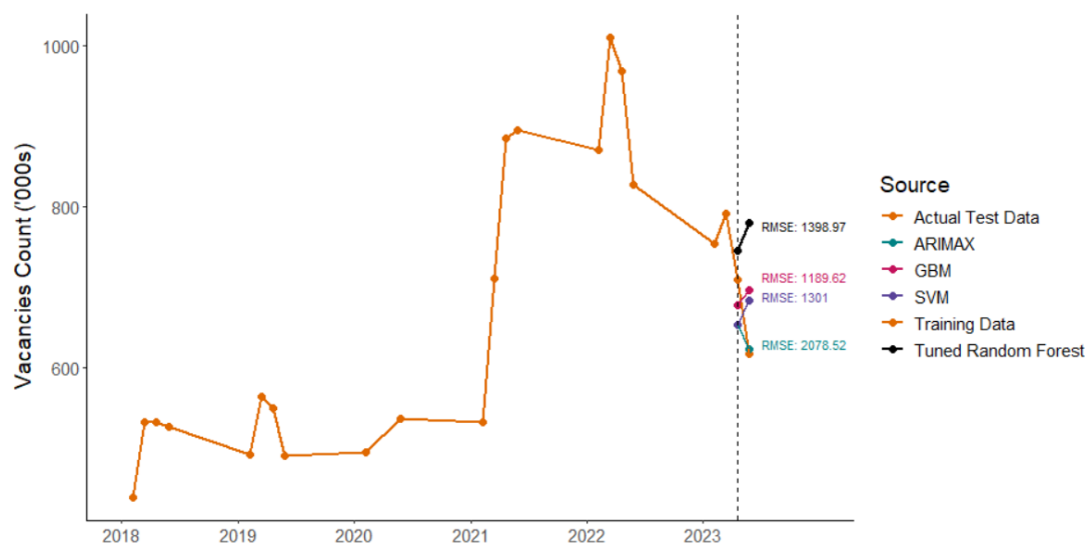
Our final step is to evaluate the predictive accuracy of our combined approach to adjusting the dataset. For this assessment, we reapply the machine learning models.

Figure 5 shows an improvement in machine learning predictive capabilities for JVWS data points within the test dataset, as measured by the RMSE. Of all models tested, the gradient boosting model (GBM) and support vector machine (SVM) yield the most accurate forecasts, reducing RMSE by approximately 50%.

The success of this combined approach demonstrates the value of a logical, step-by-step approach to building and testing different models to improve the accuracy of estimating job vacancies using OJP. By taking a step-by-step approach to testing adjustments and by evaluating outcomes at each stage, we improve the accuracy of vacancy estimates based on OJP data.

Due to the scope of this report, this is the final step of our demonstration. However, there are opportunities to expand on this work. Future research could explore improving estimates and deepening the understanding of our findings (such as the accuracy of GBM and SVM). We touch upon this in our conclusion.

Figure 5: Performance of machine learning models on combined algorithm data



Conclusion

OJP have emerged as a valuable, complementary data source for economists and policy analysts seeking to understand labour market dynamics. In particular, OJP data complement data from the JVWS by providing high-frequency, in-depth insights into the skills employers seek—an increasingly important advantage in the context of a rapidly changing labour market.

It is important to emphasize that OJP data are not the same as job vacancy data. However, when used alongside traditional survey data (such as those from the JVWS), they offer a more complete picture of labour market dynamics and business cycles, informing policy decisions and responses to economic fluctuations.

This demonstration explored how the accuracy and representativeness of OJP data can be improved. Using a comprehensive, step-by-step approach, we tested a series of adjustments: a weighting function, a winsorization algorithm (to align OJP data more closely with official JVWS data), and a robust regression with Huber weights. We then compared their performance in predicting JVWS vacancy counts. We found that applying these adjustments reduced prediction errors and increased representativeness with respect to JVWS benchmarks. Notably, combining winsorization and weighting reduced prediction errors in the combined dataset by 15%. These findings highlight the potential of OJP data when refined using targeted methods.

Our use of machine learning models also revealed interesting results, suggesting that SVM and GBM approaches are most accurate. In our work, they reduced prediction error by nearly 50%.

Our outcomes highlight the value of combining methodological rigour with data science. Deeper research to advance understanding and refine these approaches was beyond the scope of this report; however, this work indicates a promising direction for future studies that take a deeper dive into the data science. Ultimately, our goal was to demonstrate what is possible with emergent sources of LMI and advancements in data science.

Refining OJP data can improve its predictive power, which could have major implications for monitoring occupational shortages, responding to regional shifts, and filling known gaps in JVWS data. For instance, future work could explore weightings that are rooted in other (or more detailed) parameters for accuracy. This could include exploring benchmarking or other algorithm development to improve accuracy, or investigating accuracy at sub-NOC levels or in occupational sectors where JVWS data are sparse.

This study also emphasizes a broader opportunity to bring data science into the LMI ecosystem. Our aim was not to develop a definitive model, but to demonstrate what is

possible by using a transparent, step-by-step approach to testing and refining OJP-based job vacancy estimates. We encourage others to build on this work by applying and improving upon these methods in other contexts.

The JVWS remains an invaluable resource in Canada's labour market as a core dataset. Its quality, depth, and longer data time series are unmatched. Still, the inclusion of OJP data offers a complementary perspective that could improve agility and lead to frequently producible, accessible, and expansive datasets. This complementary approach offers researchers and policymakers a robust and responsive toolkit for understanding and responding to labour market dynamics.

This is only a first step—ongoing research and collaboration among OJP providers, researchers, and policymakers are necessary to develop more standardized approaches for improving OJP-based data.



References

- Abuallail, I., & Vu, V. (2022). Race alongside the machines: Occupational digitalization trends in Canada, 2006–2021. Brookfield Institute for Innovation + Entrepreneurship. <https://dais.ca/wp-content/uploads/2023/10/Race-Alongside-the-Machines-report-FINAL.pdf>
- Coibion, O., Gorodnichenko, Y., & Kumar, S. (2018). How do firms form their expectations? New survey evidence. *American Economic Review*, 108(9), 2671–2713. <https://doi.org/10.1257/aer.20151299>
- Evans, D., Mason, C., Chen, H., & Reeson, A. (2023). An algorithm for predicting job vacancies using online job postings in Australia. *Humanities & Social Sciences Communications*, 10(1), 102. <https://doi.org/10.1057/s41599-023-01562-9>
- Gorodnichenko, Y., & Peter, K. S. (2007). Public sector pay and corruption: Measuring bribery from micro data. *Journal of Public Economics*, 91(5–6), 963–991. <https://doi.org/10.1016/j.jpubeco.2006.12.003>
- Hassan, C. A. U., Khan, M. S., & Shah, M. A. (2018). Comparison of machine learning algorithms in data classification. In X. Ma (Ed.), 2018 24th International Conference on Automation and Computing: Improving productivity through automation and computing (pp. 1–6). Institute of Electrical and Electronics Engineers. <https://doi.org/10.23919/IConAC.2018.8748995>
- Huber, P. J. (1967). The behavior of maximum likelihood estimates under non-standard conditions. In L. M. Le Cam, & J. Neyman (Eds.), *Proceedings of the fifth Berkeley symposium on mathematical statistics and probability*, Volume 1: Statistics (pp. 221–233). University of California Press. <https://projecteuclid.org/proceedings/berkeley-symposium-on-mathematical-statistics-and-probability/Proceedings-of-the-Fifth-Berkeley-Symposium-on-Mathematical-Statistics-and/Chapter/The-behavior-of-maximum-likelihood-estimates-under-nonstandard-condition>
- Osisanwo, F. Y., Akinsola, J. E. T., Awodele, O., Hinmikaiye, J. O., Olakanmi, O., & Akinjobi, J. (2017). Supervised machine learning algorithms: Classification and comparison. *International Journal of Computer Trends and Technology*, 48(3), 128–138. <https://doi.org/10.14445/22312803/IJCTT-V48P126>
- Rosenbaum, Z., & Feor, B. (2020). How representative are online job postings? (LMI Insight Report no. 36). LMIC. <https://lmic-cimt.ca/publications-all/lmi-insight-report-no-36/>

The Josh Bersin Company. (2023, June 1). New research shows that hiring is harder than ever: Time to hire increasing significantly for almost all roles. PR Newswire. <https://www.prnewswire.com/news-releases/new-research-shows-that-hiring-is-harder-than-ever-time-to-hire-increasing-significantly-for-almost-all-roles-301839785.html>

Turrell, A., Speigner, B., Djumalieva, J., Copple, D., & Thurgood, J. (2022). Transforming naturally occurring text data into economic statistics. In K. G. Abraham, R. S. Jarmin, B. C. Moyer, & M. D. Shapiro (Eds.), *Big data for twenty-first-century economic statistics* (pp. 173–208). University of Chicago Press. <https://www.degruyterbrill.com/document/doi/10.7208/chicago/9780226801391-008/html?srsId=AfmBOoomHsBanSFIROhI7bncw1DStKVM1B4DA0zrMGUQ7FjKprDTjjSw>

Vu, V., Lamb, C., & Willoughby, R. (2019). *I, human: Digital and soft skills in a new economy*. Brookfield Institute for Innovation + Entrepreneurship. <https://dais.ca/wp-content/uploads/2023/10/I-Human-ONLINE-FA-1.pdf>



