

Nacelle modeling considerations for wind turbines using large-eddy simulations

Amaral, R.; Houtin-Mongrolle, F.; Von Terzi, D.; Viré, A.

DOI

[10.1088/1742-6596/2767/5/052056](https://doi.org/10.1088/1742-6596/2767/5/052056)

Publication date

2024

Document Version

Final published version

Published in

Journal of Physics: Conference Series

Citation (APA)

Amaral, R., Houtin-Mongrolle, F., Von Terzi, D., & Viré, A. (2024). Nacelle modeling considerations for wind turbines using large-eddy simulations. *Journal of Physics: Conference Series*, 2767(5), Article 052056. <https://doi.org/10.1088/1742-6596/2767/5/052056>

Important note

To cite this publication, please use the final published version (if applicable).
Please check the document version above.

Copyright

Other than for strictly personal use, it is not permitted to download, forward or distribute the text or part of it, without the consent of the author(s) and/or copyright holder(s), unless the work is under an open content license such as Creative Commons.

Takedown policy

Please contact us and provide details if you believe this document breaches copyrights.
We will remove access to the work immediately and investigate your claim.

PAPER • OPEN ACCESS

Nacelle modeling considerations for wind turbines using large-eddy simulations

To cite this article: R Amaral *et al* 2024 *J. Phys.: Conf. Ser.* **2767** 052056

View the [article online](#) for updates and enhancements.

You may also like

- [Power curve measurement of a floating offshore wind turbine with a nacelle-based lidar](#)
Umut Özinan, Dexing Liu, Raphaël Adam et al.
- [Framework for Design Validation of Control Algorithms used on Mechanical Hardware-in-the-Loop Nacelle Test Rigs](#)
Oliver Feindt, Muhammad Omer Siddiqui, Adam Zuga et al.
- [Investigation of the nacelle blockage effect for a downwind turbine](#)
Benjamin Anderson, Emmanuel Branlard, Ganesh Vijayakumar et al.



HONOLULU, HI
October 6-11, 2024

Joint International Meeting of
The Electrochemical Society of Japan (ECSJ)
The Korean Electrochemical Society (KECS)
The Electrochemical Society (ECS)



Early Registration Deadline:
September 3, 2024

MAKE YOUR PLANS NOW!



Nacelle modeling considerations for wind turbines using large-eddy simulations

R Amaral¹, F Houtin-Mongrolle², D von Terzi¹ and A Viré¹

¹ Technische Universiteit Delft (TU Delft), Kluyverweg 1, 2629 HS Delft, Netherlands

² Siemens Gamesa Renewable Energy (SGRE), Prinses Beatrixlaan 800, 2595 BN Den Haag, Netherlands

E-mail: r.pintoelisbaomartinsamaral@tudelft.nl

Abstract. Two setups are used to investigate differences between modeling a wind turbine nacelle by means of an actuator-line model (ALM) and a wall-model (WM) using large-eddy simulations. One advantage of the ALM is that it requires a lower mesh refinement, making it less computationally costly. In the first setup, the nacelle is in standalone configuration and the ALM results show a much lower turbulence intensity and a significantly slower wake recovery when compared to the WM cases. In the second setup, the nacelle is in a rotor-nacelle assembly configuration and many variations of the ALM are tested in order to match the results from the experiment addressed in the OC6 task phase III. Contrary to previous findings that the nacelle might affect the turbine loads, this study shows that the improved match with the experiment stems from the increased mesh refinement in the nacelle region rather than the actual presence of the nacelle. Nevertheless, the wake profiles in the near-wake show a very good agreement between the ALM and WM, regardless of the refinement in the nacelle region. These cases also show a higher wake deficit than not using any nacelle at all.

1. Introduction

A wind turbine nacelle produces an inner shear layer and vortex system that can change the wake and rotor dynamics. When using an actuator-line model (ALM) for the blades, the nacelle is frequently omitted or included as an ALM [1][2] because using a wall-model (WM) significantly increases the computational cost. Moreover, with floating offshore wind turbines (FOWT), the rotor experiences translation and rotation motions that are easier to replicate with ALMs because they don't require remeshing. Although faster, nacelle ALMs can lead to differences in the flow when compared to higher-fidelity WMs, ultimately impacting the loads and the wake of wind turbines. Churchfield *et al.* [3] tested both an actuator drag disk (ADD) and a power-coefficient-based (C_P -based) nacelle models. They showed that the ADD leads to a steady wake when only the nacelle is simulated, while the C_P -based nacelle is able to generate an unsteady wake. Both models used together with similar tower models lead to mean wake velocity profiles which are much closer to the experiment than not using any models. Further including a pair of equal and opposite side-forces on the tower leads to an unphysical force peak when the blade passes through the tower. Yang *et al.* [4] tested an actuator surface model (ASM) of the nacelle where the normal and tangential forces were calculated by satisfying the non-penetration condition and using a friction coefficient, respectively. Comparison with wall-resolved (WR) large-eddy simulations (LES) shows good agreement for the large scales of the flow field. Good agreement is also found in the wake deficit profiles except in the region from 3R-7R where the WR deficit is lower, i.e, the wake recovery is faster. In terms of turbulent kinetic energy (TKE), the WR



case shows higher values of this quantity up to 5R and they decay faster than the ASM case. In a hydrokinetic turbine setup, using the nacelle ASM yields a much better match in the mean axial velocity profiles up to 3D downstream, when compared with the no-nacelle case. The match is similar and good for the remainder positions up to 10D. Regarding the TKE, a better match is observed up to 10D downstream, except at 3D. Anderson *et al.* [5] investigated the effect of the nacelle in the rotor performance of a downwind turbine using multiple fidelity techniques. They report that the maximum change in C_P and thrust coefficient (C_T) for many nacelle shapes, aspect ratios, Reynolds number (Re) and turbulence intensity (TI) relative to a no-nacelle case was 0.54%. They also reported a much faster wake recovery with the standalone geometry-resolved nacelle when compared to the standalone body-force model. De Cillis *et al.* [6] compared the wake of a wind tunnel turbine setup with and without immersed-boundary method nacelle and tower by means of LES with ALM representations of the blades. They reached the conclusion that the nacelle and tower increase the wake deficit in the near-wake but promote a faster recovery. Proper orthogonal decomposition analysis of the wake modes suggests that the nacelle slows the root vortices' convection and significantly alters the modal energy distribution, disavouring modes associated with the tip and root vortices, both in the near and far wake. The TKE balance shows that this promotes the wake recovery across the whole wake. Gao *et al.* [7] used LES with blade and nacelle ALMs to validate an isotropic and anisotropic nacelle kernel against experimental data. Comparison with wind tunnel experiments suggests that the anisotropic kernel results are within 9% in terms of wake deficit at the centerline relative to the experiment. This contrasts with deviations above 25% for the isotropic kernel and above 42% for the no-nacelle case. Comparison with field experiments showed that the mean velocity went from a 30% to 8% deviation from the experiment and that the TI went from an mean deviation of 14.44% to 6.6% at two points in the wake when comparing the no-nacelle case to the anisotropic kernel. In conclusion, this literature review supports the fact that the inclusion of the nacelle will have little impact on the rotor loads but a significant impact on the wake evolution. However, some of our preliminary simulations suggested that the loads too would be impacted between 3% to 9% depending on the nacelle modeling approach which was something also reported by another participant in the OC6 task phase III [8]. This led us to further investigate the nacelle modeling and its impact on the rotor performance.

2. Objectives

This paper focuses on identifying caveats and recommendations of using the actuator-line model (ALM) versus using wall models (WMs) for nacelle modeling.

3. Methodology

YALES2 [9], a low-Mach number LES simulation tool that is 4th-order-accurate in space and time, was used in this paper. All of the meshes were unstructured and tetrahedral. First, the nacelle was studied in a standalone setup (SS) by comparing ALMs to WMs, in order to compare the wind fields. The nacelle WM geometry used to reproduce the experiment was in fact made up of two parts but only the hub was considered here to simplify the flow field. The hub was a half-sphere on top of a cylinder with a radius $R = 0.183\text{ m}$ and a height of $h = 0.064\text{ m}$. Afterwards, the ALM and the WM were studied in a rotor-nacelle assembly setup (RNAS). This study was conducted in two parts. In the first part, we investigated the nacelle modeling by testing many variations of the ALM, in order to improve the rotor loads and near-wake when compared to the WM and the experiment [8]. In the second part, we investigated the used meshes with the same goal. In this setup, the domain replicated the POLIMI wind tunnel [8] which has dimensions: 13.84 m wide x 3.84 m high x 35 m long. The blades were modeled as actuator-lines. In all cases of both setups, the inflow was uniform and steady and the inlet wind speed was $U_0 = 4\text{ m/s}$ which is the rated wind speed of the scaled-down DTU 10 MW turbine studied in the RNAS [8]. The turbine diameter was $D = 2.381\text{ m}$. Simulations in the SS were performed for 12 s with

load and wake convergence achieved after 3 s and 8 s, respectively. Simulations in the RNAS were performed for 17.86 s with load and wake convergence achieved after 8.93 s and 10.36 s, respectively. All statistics were taken after the convergence was achieved.

In SS, three ALMs with different kernels were tested. The first kernel, K_x , denotes the Gaussian kernel with steady thrust, i.e, with mollification only in the streamwise direction. The second kernel K_R was obtained by summing the remaining steady and unsteady force components to K_x . The last kernel, $K_x + K_{rc}$, was the sum of K_x with an outward steady radial force component relative to the hub axis which is denoted by K_{rc} , where r and c stand for radial and cylindrical, resp. Unlike K_x and K_R , the force field direction of $K_x + K_{rc}$ is not uniform. The slices of K_x and $K_x + K_{rc}$ force fields on the x-z plane are shown in fig. 1, where x and z are the streamwise and height directions, resp. K_R is not shown due to the similarity with K_x . The force projection for K_x and $K_x + K_{rc}$ followed equations 1 & 3 and 2 & 3, resp., where the bold font denotes a vector; $\mathbf{f}$ is the mollified force field; F_x and F_{rc} are appropriate scalars; $\eta(d)$ is the Gaussian mollification kernel; d is the distance to either the kernel center for K_x or axis for K_{rc} ; $\epsilon = 2\Delta x$ is the smearing length scale, where Δx is the maximum cell size in the vicinity of the ALM; $\mathbf{e}_x$ is a unitary vector directed in the streamwise direction and $\mathbf{e}_{rc}$ is a unitary vector perpendicular to $\mathbf{e}_x$. A fair comparison between the ALMs and the WM was attained by ensuring that the mollified forces were the same as in the WM. Therefore, the scalars F_x and F_{rc} represent the WM thrust and radial forces. These were calculated by dividing the WM into forty sections and by integrating the resultant force at each of these sections projected on the corresponding director vectors ($\mathbf{e}_x$ and $\mathbf{e}_{rc}$). The steady part of the forces was obtained by taking the time-average of the integrated WM forces after convergence. The unsteady part was obtained by summing Fourier modes such that 99% of the force spectral energy was included.

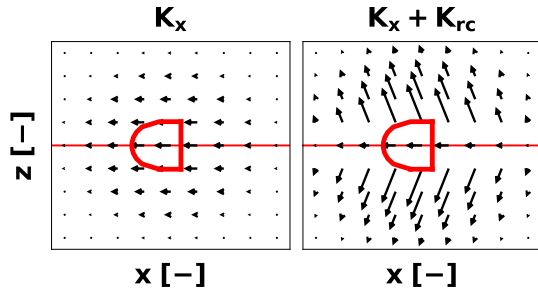


Figure 1: Mollification kernels. The vectors show the direction and norm of the mollified force. The red axis coincides with the hub rotation axis. The hub shape is shown in red as a reference. The figure is not at scale.

$$\mathbf{f}(x, y, z) = F_x \eta(d) \mathbf{e}_x \quad (1) \quad \mathbf{f}(x, y, z) = F_x \eta(d) \mathbf{e}_x + F_{rc} \eta(d) \mathbf{e}_{rc} \quad (2) \quad \eta(d) = \frac{1}{\epsilon^3 \pi^{3/2}} \exp \left[- \left(\frac{d}{\epsilon} \right)^2 \right] \quad (3)$$

Two meshes were used in the SS, MS1 for the ALM simulations and MS2 for the WM simulations (see fig. 2). Both had the same boundaries, external boundary conditions, growth rate ($GR = 1.07$) and three common levels of refinement that were approximately $\Delta x_T = 5.2$ cm for turbulent inflow (although no simulations with turbulence were included in the end), $\Delta x_B = 20$ cm for the background and $\Delta x_W = 2.1$ cm for the wake which was approximately the same refinement used by us in [10]. However, while in MS2, the hub geometry is present in the mesh as a no-slip boundary condition with a cell size $\Delta x_N = 0.16$ cm at the boundary, in MS1 there is no hub boundary and the hub volume is filled with cells of size $\Delta x_N = \Delta x_W$. This because the goal was to model the hub without increasing the computational time. Two WMs were tested, the viscous sub-layer (VS) and the log-law (LL). Fig. 3 shows the average y^+ (denoted by $\bar{y}^+$) contour for the VS case in the interval $t = [16; 23.5]$ s to give an idea of how it was distributed across the surface. The standard deviation $std(\bar{y}^+)$ had a maximum value of 2.51. The exterior dark blue contour does not represent $\bar{y}^+$. Despite the fact that $max(\bar{y}^+) = 23$ and that the VS should only be used when $y^+ \leq 1$, the ultimate goal in future work is to analyze

the nacelle in a FOWT where, due to the floater motion, there is likely to be separation. This is a situation where the LL model breaks down and hence, we decided to leave the door open to performing full rotor simulations under prescribed motion with the VS too. Ideally, we would have liked to have gone down to $y^+ \leq 1$ but this was too costly at the moment. In any case, we verify in the next section that both WMs perform very similarly in the standalone setup.

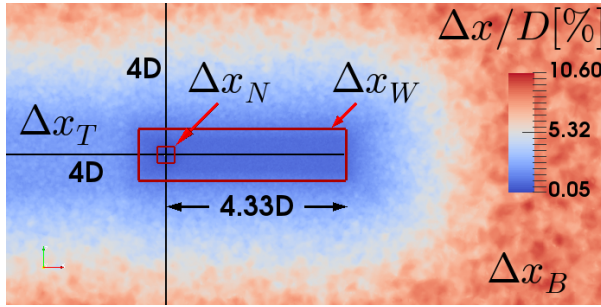


Figure 2: Cell size in the MS2 mesh. $\Delta x/D[\%]$ is the cell size as a percentage of the turbine diameter. The black lines indicate distances. The hub is contained in the red square. The red parallelepiped delimites the wake refinement area.

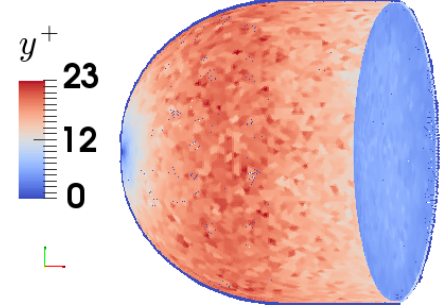


Figure 3: $\overline{y^+}$ in the interval $t = [16; 23.5] s$ for the viscous sub-layer (VS) case.

Moving to the RNAS, we first investigated the nacelle modeling. We considered the complete nacelle in the WM simulations, instead of just the hub, so that we would be closer to the experimental setup. Two meshes were used, MR1 for the ALM and no-nacelle simulations and MR2 for the WM simulations (see fig. 4). The similarities and differences between MR1 and MR2 are the same as between MS1 and MS2 in the SS. MR1 and MR2 had a $GR = 1.2$ and three common levels of refinement that were approximately $\Delta x_{BL} = 5.8 cm$ for the top and bottom boundary layers, $\Delta x_B = 40 cm$ for the background and $\Delta x_W = 1.9 cm$ for the wake. MR2 had $\Delta x_N = 0.16 cm$. MR1 was the mesh used by us in [10] where we obtained a reasonable validation for the wake dynamics. In MR1, $\Delta x_N = \Delta x_W$. Only the viscous sub-layer WM was tested in this setup and for the standard case (W1 in table 1), $\max(\overline{y^+}) = 22$ and $\max(std(y^+)) = 4$ at the nacelle. As for the top and bottom boundaries, $\max(\overline{y^+}) = 740$ and $\max(std(y^+)) = 71$. After the nacelle modeling, we investigated the mesh itself by introducing a third mesh, MR3, which is essentially the same as MR2 with one difference: while in MR2, the hub geometry is present in the mesh as a no-slip boundary condition with cell size Δx_N at the boundary, in MR3 there is no hub boundary and the hub volume is filled with cells of size Δx_N from MR2.

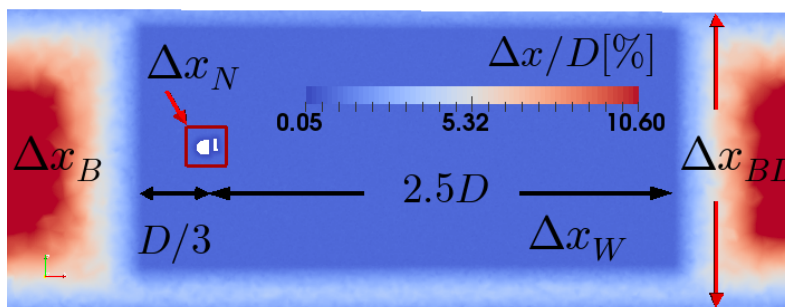


Figure 4: Cell size in the MR2 mesh. The nacelle is at a height of $0.88D$. The legend refers to $\Delta x/D[\%]$. The two-part nacelle is contained in the red box. The black two-sided arrows indicate horizontal distances.

The performed cases are summarized in table 1 where "Model" indicates the hub/nacelle modeling technique, "Geom." indicates the hub/nacelle geometry in the domain, "Mesh"

identifies the used mesh, " $\#$ Elem." is the number of elements in the mesh, "GR" is the growth rate in the wake refinement, " $\epsilon_N/\max(\Delta x_N)$ " is the ratio between the hub/nacelle kernel smearing length scale and the maximum cell size in its vicinity, " $\epsilon_B/\max(\Delta x_B)$ " is the ratio between the blade kernel smearing length scale and the maximum cell size in their vicinity, " T_N " is the integrated thrust value in WM cases and the prescribed thrust value in ALM cases and " CT " is the computation time. For all cases, except A9 and A10, each blade had 47 actuator-points. A9 and A10 are mesh refinement simulations of A1 and thus, their mesh (MR1/2) was obtained by uniformly splitting MR1's cell sides in two, leading to twice the number of actuator-points. The SS is covered by simulations K_x to LL. The first part of the RNAS is covered by cases A1 to W2. The second part of the RNAS is covered by cases W3 to A2*. W1 is the baseline case with the nacelle WM meant to be the closest to the experimental setup as possible and from which we obtain the nacelle loads to be prescribed in the ALM cases.

Table 1: Case definition.

Case Units	Model	Geom.	Mesh	# Elem. [$\times 10^6$]	GR [-]	$\frac{\epsilon_N}{\max(\Delta x_N)}$ [-]	$\frac{\epsilon_B}{\max(\Delta x_B)}$ [-]	Δx_N [cm]	T_N [N]	CT [CPUh]
K_x	ALM	K_x	MS1	186	1.07	2.0	-	2.1	0.047	NA
$K_x + K_{rc}$	ALM	$K_x + K_{rc}$	MS1	186	1.07	2.0	-	2.1	0.047	NA
K_R	ALM	K_R	MS1	186	1.07	2.0	-	2.1	0.047	NA
VS	VS	Hub	MS2	232	1.07	-	-	0.16	0.047	NA
LL	LL	Hub	MS2	232	1.07	-	-	0.16	0.053	NA
A1	ALM	K_x	MR1	97	1.2	2.0	2.0	1.9	0.192	2367
A2	-	-	MR1	97	1.2	-	2.0	1.9	0.192	1133
A3	ALM	K_x	MR1	97	1.2	2.0	2.0	1.9	0.384	3033
A4	ALM	K_x	MR1	97	1.2	1.0	2.0	1.9	0.192	2273
A5	ALM	K_x	MR1	97	1.2	4.0	2.0	1.9	0.192	2286
A6	ALM	K_R	MR1	97	1.2	2.0	2.0	1.9	0.192	2773
A7	ALM	15x1 K_x	MR1	97	1.2	2.0	2.0	1.9	0.192	2293
A8	ALM	3x5 K_x	MR1	97	1.2	2.0	2.0	1.9	0.192	2386
A9	ALM	K_x	MR1/2	772	1.2	2.0	2.0	0.95	0.192	46933
A10	ALM	K_x	MR1/2	772	1.2	2.0	4.0	0.95	0.192	50933
IB1	IBM	Hub	MR1	97	1.2	-	2.0	1.9	-	2620
W1	VS	Nacelle	MR2	99	1.2	-	2.0	0.16	0.192	35058
W2	VS	Hub	MR2	101	1.2	-	2.0	0.16	0.167	22960
W3	VS	Nacelle	MR2	115	1.07	-	2.0	0.16	0.197	35093
A1*	ALM	K_x	MR3	126	1.07	2.0	2.0	0.16	0.192	35093
A2*	-	-	MR3	126	1.07	-	2.0	0.16	0.192	27440

4. Standalone setup (SS) - Wake

Here, the wakes of the ALMs were compared to the wakes of the WMs. The integrated and time-averaged hub loads, F_x and F_{rc} , were 0.047 N and 0.215 N. The mean streamwise velocity fields around $K_x + K_{rc}$, VS and K_x are shown in fig. 5. It can be seen that $K_x + K_{rc}$ induces an unrealistic flow acceleration upstream of the ALM. The inspection of the pressure field reveals a suction zone where the acceleration happens, as expected. This is likely due to the force field driving the fluid away from the hub axis. Not only this, but the velocity field was very irregular because both the magnitude and direction of the mollified force varied for every point in the mesh, especially near the axis. Because of this, this kernel concept was discarded. Unlike $K_x + K_{rc}$, K_x leads exclusively to flow deceleration in the domain, although somewhat differently than VS.

In particular, the stagnation zone is less circular and less pronounced, there is no noticeable flow acceleration on the sides and, there is no recirculation in the wake. The mean streamwise velocity and TI_x profiles at a vertical line crossing the wake center can be seen in fig. 6. Not only do the WMs (VS and LL) produce a larger wake deficit than the ALMs (K_x and K_R) but also a larger value of TI_x at $x = 0.15 D$ which is consistent with the faster wake recovery seen afterwards. K_x does not produce any turbulence since the mollified force is steady. This significantly hampers the wake recovery when compared to the WMs due to virtually no mixing. K_R produces an improved recovery but only noticeable at $x = 4 D$, presumably because the force variations take time to develop into wake perturbations. Only when TI_x hits a certain level, do we have a meaningful recovery. It is reasonable to assume that the length scales generated by an ALM with unsteady force projection will differ significantly from the ones generated by a WM. In the ALM case, the wake may take more time to decay into turbulence since the mollified force direction is uniform and the magnitude changes smoothly from the core. Furthermore, we expect length scales of the order of the hub kernel smearing length scale for the ALMs, while for the WMs, we expect scales of the order of the hub diameter and of the boundary layer thickness, with the latter disturbing the flow right at the hub. These differences could explain the low turbulence generated in the vicinity of ALM and the slower wake recovery. Lastly, the LL case shows an increase in thrust of 13% compared to VS, with no material changes in the mean velocity profiles, suggesting a low boundary condition sensitivity at this Re . This suggests no objection about using the VS going forward. Overall, there are some differences between modeling the hub with an ALM and a WM. Still, this doesn't necessarily imply a large impact on the rotor loads and wake evolution in the RNAS.

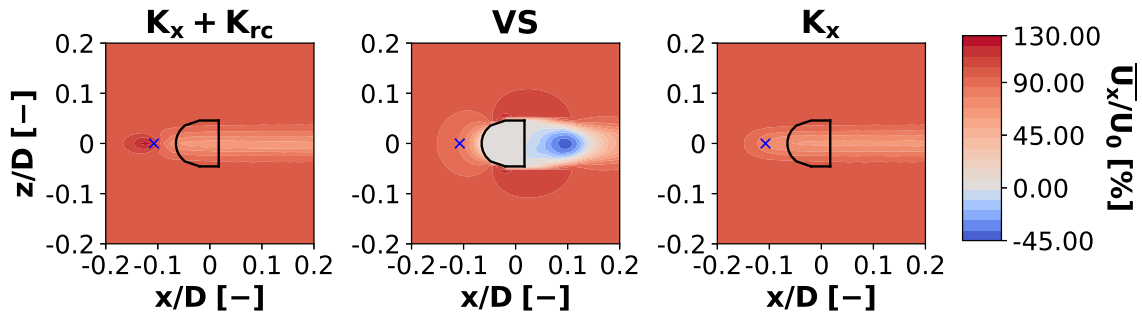


Figure 5: Mean streamwise velocity in the standalone setup. The blue cross indicates the kernel core position. The hub contour is shown in black. The subplot title is the case (see table 1).

5. Rotor-nacelle assembly setup (RNAS) - Loads

Differences identified before between the ALMs and the WMs may or may not impact the rotor loads so, in order to answer that question, we moved to a RNAS. We first tried to match the experimental results by working on the nacelle modeling. The load results of the cases defined in table 1 can be found in table 2, where " T_N " is the integrated nacelle thrust value in the WM cases and the prescribed nacelle thrust value in the ALM cases, " T_B " is the 3-blade thrust, " ΔT_{Exp} " is the relative difference between the total thrust $T = T_N + T_B$ and the experimental value of 36 N, and M denotes the torque, which follows the same nomenclature as the thrust, with an experimental value of 3.24 Nm. As expected, the nacelle wall-modeled case (W1) shows the best agreement with the experiment while using no nacelle at all (A2) worsens the results significantly. However, adding K_x to the rotor (A1) improves the loads very little which was unexpected. To better understand this, we tried many different variations of A1 to try and improve the results. Doubling the prescribed thrust (A3), halving or doubling ϵ_N (A4 and A5),

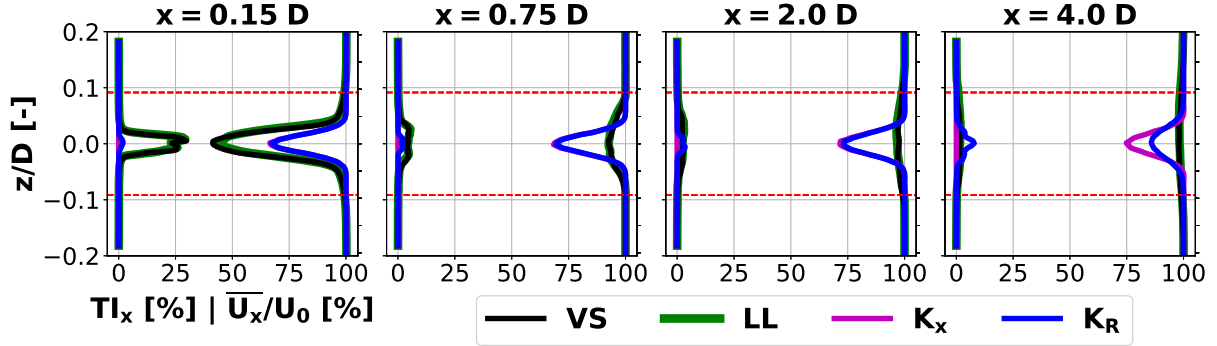


Figure 6: Mean streamwise velocity and turbulence intensity in the standalone setup. The plots come from vertical lines at the wake center. The legend indicates the cases (see table 1).

using K_R (A6) and changing the number of the ALM actuator points and their configuration in space (A7 and A8) all have a marginal impact. We also performed one mesh refinement step to assess the sensitivity to the mesh resolution. Doubling the mesh refinement while keeping $\epsilon_B/\max(\Delta x_B) = 2$ significantly decreases the loads away from the experiment (A9). However, if the refinement is performed while keeping the same kernel size by making $\epsilon_B/\max(\Delta x_B) = 4$ (A10), the loads come closer to the experimental value but still less than W1. Moving further, we tested another modeling technique, the immersed-boundary method (IBM) with a penalization function (IB1). In this case, we didn't calculate the nacelle loads so we used the blade loads as a proxy, since most of the load changes stem from the blades in the other cases. The loads did not differ significantly from A1 though. One explanation for the mismatch between the ALM/IBM and the WM could be the fact that the shape of the nacelle in W1 and the experiment had in fact two parts (see fig. 4) rather than a single hub. The back of the nacelle could be introducing flow perturbations that increase the deficit locally and lead to higher blockage when compared to a single-hub geometry. To test this hypothesis, we modified W1 to include just the hub (W2) and verified that it does have an impact but not sufficient to explain the mismatch.

The previous tests suggest that the rotor loads are insensitive to the nacelle ALM but such differences with regards to the WM were unexpected. This because both the nacelle ALM and WM are clearly disturbing the flow by introducing a deficit in the wake which should lead to a velocity excess elsewhere in the same plane (see A1 and W1 in fig. 8). Facing this, we shifted our focus to the mesh itself. The meshes used so far were designed with a growth rate (GR) of 1.2 in the wake region and this influences the cell size smoothness from the nacelle region to the wake region and from the wake region to the top and bottom boundaries of the domain where there is a boundary layer. It could be that a GR of 1.2 was too steep and would change the boundary layer thus influencing the blockage, so we remeshed MR2 to have a smoother GR of 1.07 (compare W1 to W3) but this had little impact. We later recalled that the cell size of MR2 in the nacelle region was much lower than in MR1 (compare W1 to A1), i.e., $\epsilon_N/\max(\Delta x_N) > 2$ and $\epsilon_B/\max(\Delta x_B) > 2$ in the vicinity of the nacelle for MR2. This difference could explain the mismatch and to test this hypothesis, MR2 was remeshed into MR3 (see "Methodology" for details) producing cases A1* and A2*, from A1 and A2, respectively. This way, the only difference between A1*/A2* and W3 was the nacelle model or its absence given that the refinement in the nacelle region was the same in all three cases. Curiously, A2* loads are as close to the experiment as W3, despite the fact that the nacelle is not present. Adding the nacelle ALM in the same mesh (A1*) has a marginal impact relative to A2*. We concluded therefore that the nacelle was not responsible for the differences but rather

Table 2: Case load results.

Case	T_N [N]	T_B [N]	ΔT_{Exp} [%] [N]	M_N [Nm]	M_B [Nm]	ΔM_{Exp} [%]
A1	0.192	34.19	-4.51	0	2.95	-8.93
A2	0	34.13	-5.19	0	2.93	-9.42
A3	0.384	34.25	-3.80	0	2.96	-8.59
A4	0.192	34.14	-4.63	0	2.94	-9.36
A5	0.192	34.15	-4.60	0	2.94	-9.21
A6	0.192	34.19	-4.50	0	2.95	-8.93
A7	0.192	34.17	-4.55	0	2.94	-9.14
A8	0.192	34.21	-4.44	0	2.95	-8.99
A9	0.192	31.90	-10.85	0	2.54	-21.36
A10	0.192	34.78	-2.87	0	3.05	-5.81
IB1	-	34.23	-	-	2.95	-
W1	0.192	35.39	-1.15	0.002	3.17	-2.18
W2	0.164	35.14	-1.94	0.001	3.11	-3.95
W3	0.197	35.35	-1.23	0.002	3.16	-2.43
A1*	0.192	35.31	-1.40	0	3.16	-2.34
A2*	0	35.45	-1.53	0	3.19	-1.41

the mesh. At least two explanations are possible: (1) The fact that $\epsilon_N/\max(\Delta x_N) > 2$ and $\epsilon_B/\max(\Delta x_B) > 2$ in the vicinity of the nacelle can directly affect the blade root and nacelle kernels, thus increasing their mollified force and (2) given the smaller cell size in the nacelle region, some root vortex length scales are no longer filtered by the mesh which leads to changes in the wake evolution. This can be clearly seen for A2* on the right plot of fig. 7, at the rotor centerline, where the root vortices are more pronounced and where there are more filaments of higher curl magnitude. Both these explanations can directly modify the thrust or indirectly by influencing the wake but further investigation is left as future work.

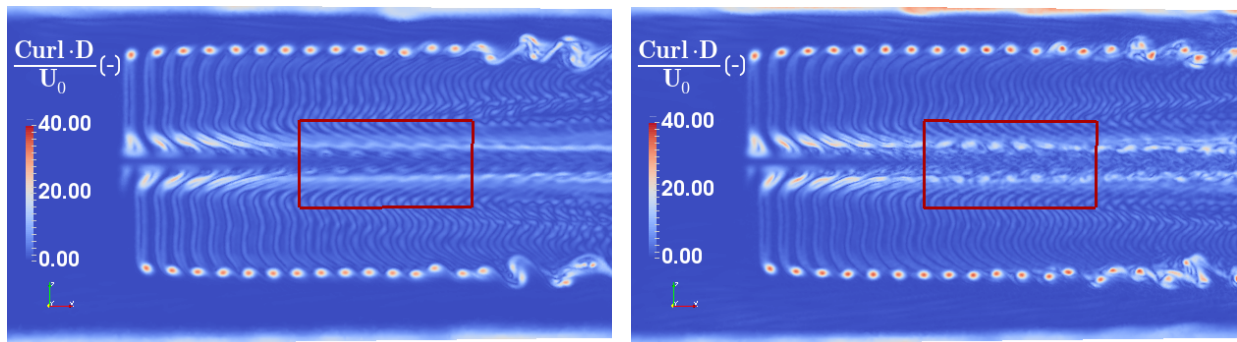


Figure 7: Velocity curl magnitude. Vertical slices of the domain passing by the rotor axis are shown for A2 on the left and for A2* on the right. The slices were taken at approximately $t = 17.86$ s. The red square highlights the differences seen at the wake core.

6. Rotor-nacelle assembly setup (RNAS) - Wake

In this section, the impact of the nacelle in the near-wake evolution is analyzed by looking at a smaller set of cases (A1, A2, A1*, A2* and W1). In terms of mean streamwise velocity profile

in the nacelle region, the simulations with nacelle modeling (A1, A1* and W1) agree to a very good extent. The same can be said for the simulations without nacelle (A2 and A2*). As for the blade region, the profiles agree to a very good degree for the simulations with a higher refinement in the nacelle region (A1*, A2* and W1). This is in accordance with the fact that the thrust is higher for these cases, leading to higher momentum extraction and velocity deficit. Vice-versa, the profiles of the simulations with a lower refinement in the nacelle region (A1 and A2) also agree. Regarding the TI, all the simulations have decent agreement outside of the nacelle region for all the positions, except W1 at top tip-vortex position at $x = 2.5 D$. In the nacelle region, simulations without the nacelle show relatively close TI_x profiles but with an increasing difference moving downstream. This is in accordance with the fact that the mean velocity profile in A2* is increasingly lagging that of A2 while the root vortices are more pronounced (see fig. 7). Simulations with the nacelle ALM (A1, A1*) lead to higher TI_x values than the simulation with the nacelle WM (W1) after $x = 0.5 D$. A possible explanation for this could be that the nacelle WM generates part of the turbulence with smaller length scales (the ones stemming from the boundary layer) than those generated by the ALM. Since in the WM case, the cell size increases from the nacelle surface up to the wake refinement, we can expect the scales below the increasing cell size to be filtered out, leading to lower turbulence intensity downstream. The overall conclusion in the RNAS with both meshes is that the nacelle ALM leads to near-wake profiles that are much closer to those of the WM, when compared with the no-nacelle situation and the standalone setup (SS). We reckon this is caused by the cyclical fluctuations and turbulence generated by the root vortices, which induce a faster wake convergence. We support this with the fact that the highest spectral peak, by far, in the nacelle thrust force spectrum of W1 happens at a frequency of 12 Hz which is the 3P frequency. Moreover, the hub thrust force increases from 0.047 N to 0.164 N (a factor of 3.5) when moving from the SS to the RNAS (see cases VS and W2). Although the blockage effect is much lower in the SS, it is unlikely that the blockage alone is responsible for such an increase. Instead, this is probably due to the combination of lower pressure in the back of the nacelle induced by the cyclical blade loading and the flow tunneling in between the blades.

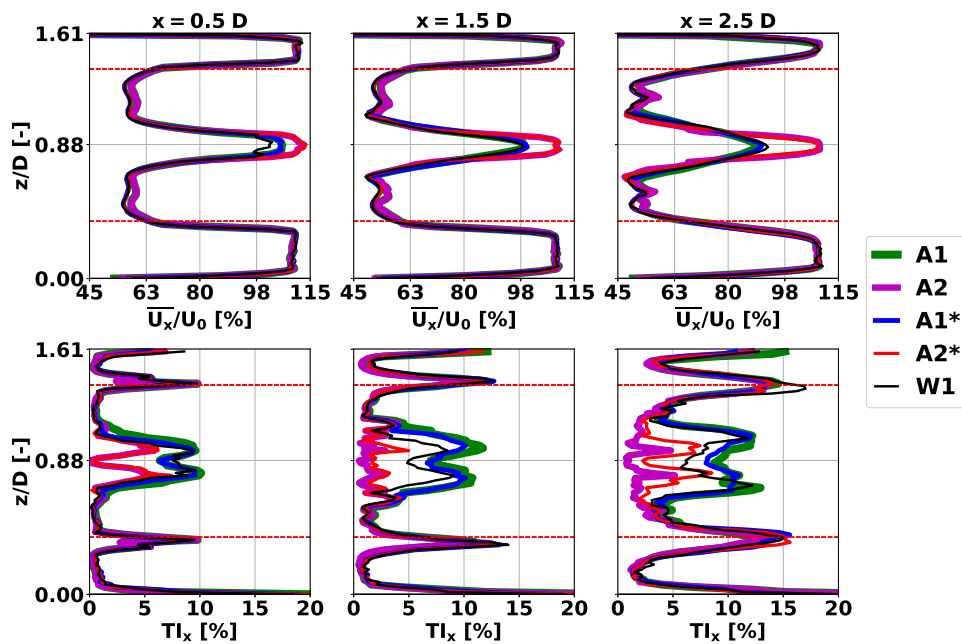


Figure 8: Mean stream-wise velocity (top) and TI_x (bottom) The dashed red lines delimit the rotor area.

7. Computation cost comparison

Lastly, we compared the cost between both techniques using the computation time $CT = T \times N_{cores}$ as a metric, where T is the elapsed simulation time and N_{cores} is the number of active CPU cores. The results are compiled in table 1. ALM cases on mesh MR1 required on average about 92% less CT than using the WM with mesh MR2. Although this is an estimate based on the cases we performed, the reported reduction is plausible taking into account that the meshes have about the same number of elements and that the minimum cell size (which limits the time step) is around 12 times larger, which theoretically leads to a reduction of around 92% in the time step. Comparing the ALM cases with the analogous no-nacelle cases, it can be seen that the computation time increases between 28% and 120% for the mesh MR1 and MR3, respectively. It is not possible to generalize this based on such a low number of cases but using the nacelle ALM model in the current implementation leads to a noticeable increase in CT .

8. Conclusions

This work focused on comparing two nacelle modeling approaches, the actuator-line model (ALM) and the wall model (WM), by assessing their impact on the rotor loads and near-wake profiles. This was performed in a standalone setup and in a rotor-nacelle assembly setup. The simulations in the standalone setup showcase key differences between the modeling approaches. It shows that the wake of the steady-thrust ALM recovers much slower than that of the WM and that the differences can be reduced to some degree in the far-wake of the nacelle, by including the unsteady part of the thrust in the ALM. However, this unsteady part appears in the wake at a location farther downstream, a fact we attribute to the difference in length scales generated by each of the models. Most of these differences vanish in the rotor-nacelle assembly setup, likely due to the dominance of the blade root-vortices action on the nacelle. In terms of mean streamwise wind velocity and TI_x profiles in the near-wake, the ALM is able to match the WM to a high degree. In terms of loads, it was concluded that the differences between the ALM and WM were caused by the mesh refinement in the nacelle region rather than the nacelle itself. All in all, balancing the improved wake characteristics and the computation time, using a single-point ALM seems to be a good idea for the studied conditions.

Acknowledgments

This work results from the STEP4WIND project, a European Doctorate programme granted under the H2020 Marie-Curie Innovative Training Network (H2020-MSCA-ITN-2019, grant 860737.) We thank SURF (www.surf.nl) for the use of the National Supercomputer Snellius. This work is part of W2ITASEC (Wind turbine Wake Interactions through Aero-Servo-Elastic Coupling) with computer resources provided by GENCI at TGCC (grant 2023-S142aspe00038 on the supercomputer Joliot Curie).

References

- [1] Troldborg N, Zahle F, Réthoré P-E and Sørensen N N 2015 *Wind Energ.* **8** p 1239–1250
- [2] Houtin-Mongrolle F 2022 Investigations of yawed offshore wind turbine interactions through aero-servo-elastic Large Eddy Simulations PhD Thesis (Rouen: INSA)
- [3] Churchfield M J, Wang Z, Schmitz S 2015 *33rd Wind Energy Symp. (Kissimmee, Florida)* (Reston: Virginia/American Institute of Aeronautics and Astronautics) p 0214
- [4] Yang X and Sotiropoulos F 2017 *Wind Energ.* **00** 1-19
- [5] Anderson B, Branlard E, Vijayakumar G and Johnson N 2020 *J. Phys.: Conf. Ser.* **1618** 062062
- [6] De Cillis G, Cherubini S, Semeraro O, Leonardi S and De Palma P 2020 *Wind Energ.* **24** p 609-633
- [7] Gao Z, Li Y, Wang T, Shen W, Zheng X, Pröbsting S, Li D and Li R 2021 *Ren. Ener.* **172** p 263-275
- [8] Bergua R *et al.* 2023 *Wind Energ. Sci.* **8** p 465–485
- [9] Moureau V, Domingo P and Vervisch L 2011 Design of a massively parallel cfd code for complex geometries *Comptes Rendus Mécanique* **339** 1239-1250
- [10] Cioni S *et al.* 2023 *Wind Energ. Sci.* **8** p 1659–1691