

Projector augmented-wave method

P. E. Blöchl

IBM Research Division, Zurich Research Laboratory, CH-8803 Rüschlikon, Switzerland

(Received 13 June 1994; revised manuscript received 22 August 1994)

An approach for electronic structure calculations is described that generalizes both the pseudopotential method and the linear augmented-plane-wave (LAPW) method in a natural way. The method allows high-quality first-principles molecular-dynamics calculations to be performed using the original fictitious Lagrangian approach of Car and Parrinello. Like the LAPW method it can be used to treat first-row and transition-metal elements with affordable effort and provides access to the full wave function. The augmentation procedure is generalized in that partial-wave expansions are not determined by the value and the derivative of the envelope function at some muffin-tin radius, but rather by the overlap with localized projector functions. The pseudopotential approach based on generalized separable pseudopotentials can be regained by a simple approximation.

I. INTRODUCTION

In the past few decades, electronic structure calculations have made significant contributions to our understanding of solid-state properties. The majority of such calculations are based on the local-density approximation (LDA) of the density-functional theory.^{1,2} The density-functional theory maps the ground state of an interacting electron gas onto the ground state of noninteracting electrons, which experience an effective potential.

Numerous methods have been developed to solve the resulting one-particle Schrödinger equation of the LDA. The most widely used electronic structure methods can be divided into two classes: (i) the linear methods³ developed by Andersen from the augmented-plane-wave (APW) method^{4,5} and the Korringa-Kohn-Rostocker method^{6,7} and (ii) the pseudopotential method based on norm-conserving *ab initio* pseudopotentials invented by Hamann, Schlüter, and Chiang.⁸ A third class, primarily employed in chemistry, uses Gaussian basis sets to expand the full wave functions.

The linear methods can be subdivided into a variety of methods ranging from the most accurate linear augmented-plane-wave (LAPW) method to the linear muffin-tin orbital (LMTO) method, which, in a simplified version, even allows some electronic structure calculations to be performed with paper and pencil. The linear methods deal with the full wave functions and treat all elements in the Periodic Table, i.e., *s*-, *p*-, *d*-, and *f*-electron systems, on the same footing.

The pseudopotential method, when used in combination with a plane-wave basis set, on the other hand, has the advantage of formal simplicity. When applied to either first-row elements or systems with *d* or *f* electrons, even pseudopotentials become very "hard," so that in practice either very large or complicated basis sets instead of plane waves have to be used. Similarly, treating semicore states as valence states, which is often necessary for early transition-metal elements and alkali and alkaline earth metals, results in hard pseudopotentials and affects

the transferability of the pseudopotential. Vanderbilt's ultrasoft pseudopotentials^{9,10} have improved this situation significantly by relaxing the norm-conservation condition that is usually imposed on the pseudopotential approach. This method also allows first-row and transition-metal elements to be dealt with in an economical way.

Car and Parrinello have combined the density-functional theory with molecular-dynamics techniques.¹¹ Here both the electronic structure problem and the dynamics of the atoms are solved simultaneously by a set of Newton's equations. In this way not only has the structure determination become a straightforward technique, but the fully dynamic time evolution of the atomic structure has also become accessible.

The Car-Parrinello method was first applied in the context of the plane-wave pseudopotential method. There is considerable interest in applying the same technique to all-electron (AE) methods, which allow one to deal efficiently with first-row and transition-metal elements and which supply information about the wave function close to the nucleus probed by several experimental techniques, but not provided by the pseudopotential approach. These are, among many others, hyperfine parameters¹² and electric field gradients.^{13,14} Several features of the Car-Parrinello method have been implemented into existing AE methods such as the combined minimization of electronic and nuclear degrees of freedom.¹⁵⁻¹⁹ To my knowledge, however, no energy-conserving molecular-dynamics simulation has been performed to date that can compare in terms of quality with simulations using the pseudopotential approach.

This article describes an approach that combines the versatility of the LAPW method with the formal simplicity of the traditional plane-wave pseudopotential approach. The method extends the augmented-wave methods, such as the LAPW method, and the pseudopotential method in a natural way. As an AE method it provides the full wave functions that are not directly accessible with the pseudopotential approach, and the potential is determined properly from the full charge density.

ties. It will be demonstrated that the accuracy of the method described here compares well with the most accurate existing electronic structure methods based on the local-density approximation. The quality of first-principles molecular dynamics obtained with the present AE approach is in line with that of state-of-the-art Car-Parrinello calculations. Hence the first energy-conserving molecular-dynamics calculations based on the full wave functions were made possible. Finally, it can be implemented with relatively minor effort into existing pseudopotential codes.

The method has many similarities with both the existing linear methods and the pseudopotential approach. We can therefore expect that this method will close the gap between the two. The LAPW method is a special case of the present method, and the pseudopotential formalism is obtained by a well-defined approximation.

This article is organized as follows. Section II establishes the principles of the method. Section III describes which approximations are required in real calculations. Section IV derives the expressions for the Hamilton operator and forces. Section V describes the implementation in a first-principles molecular-dynamics scheme. Section VI describes the basic ingredients used in the method, such as partial waves and projector functions. Section VII contains a detailed analysis of the errors introduced in Section III. Section VIII is devoted to numerical test calculations. Section IX shows the relation between the new method and existing approaches.

II. FORMALISM

A. Projector augmented-wave functions

Wave functions of real materials have very different signatures in different regions of space: **in the bonding region the wave function is fairly smooth, whereas close to the nucleus the wave function oscillates rapidly owing to the large attractive potential of the nucleus.** This is the source of the difficulty of electronic structure methods to describe the bonding region to a high degree of accuracy while accounting for the large variations in the atom center. The strategy of the augmented-wave methods has been to divide the wave function into parts, namely, a partial-wave expansion within an atom-centered sphere and envelope functions outside the spheres. The envelope function is expanded into plane waves or some other convenient basis set. Envelope function and partial-wave expansions are then matched with value and derivative at the sphere radius.

Even though the present method has been inspired by the existing augmented-wave methods, I approach the problem in a somewhat different way. The relation of my approach to the commonly used one described above will be described in Sec. IX B. Concerning the following derivation it is emphasized that the present method is, in a certain sense, the most general augmentation scheme.

Let us consider the Hilbert space of all wave functions orthogonal to the core states. The physically relevant wave functions in this Hilbert space exhibit strong os-

cillations, which make a numerical treatment cumbersome. Therefore, we transform the wave functions of this Hilbert space into a new, so-called pseudo (PS) Hilbert space. Mapping the physical valence wave functions onto the fictitious PS wave functions thus defined shall be a linear transformation and it shall transform the physically relevant AE wave functions onto computationally convenient PS wave functions. The PS wave functions will be identified with the envelope functions of the linear methods or the wave functions of the pseudopotential approach. An AE wave function is a full one-electron Kohn-Sham wave function and is not to be confused with a many-electron wave function. All quantities related to the PS representation of the wave functions will henceforth be indicated by a tilde.

This transformation changes the representation of the wave functions in a way reminiscent of the change from a Schrödinger to a Heisenberg picture. Knowing the transformation $\mathcal{T}$ from the PS wave function to the AE wave functions, we can obtain physical quantities, represented as the expectation value $\langle A \rangle$ of some operator A , from the PS wave functions $|\tilde{\Psi}\rangle$ either directly as $\langle \Psi | A | \Psi \rangle$ after transformation to the true AE wave functions $|\Psi\rangle = \mathcal{T}|\tilde{\Psi}\rangle$ or as the expectation value $\langle A \rangle = \langle \tilde{\Psi} | \tilde{A} | \tilde{\Psi} \rangle$ of a PS operator $\tilde{A} = \mathcal{T}^\dagger A \mathcal{T}$ in the Hilbert space of the PS wave functions. Similarly we can evaluate the total energy directly as a functional of the PS wave functions. The ground-state PS wave functions can be obtained from

$$\frac{\partial E[\mathcal{T}|\tilde{\Psi}]}{\partial \langle \tilde{\Psi} |} = \epsilon \mathcal{T}^\dagger \mathcal{T} |\tilde{\Psi}\rangle. \quad (1)$$

Next, we choose a particular transformation. Since we will exploit the characteristics of particular atom types, we consider only transformations that differ from identity by a sum of local, atom-centered contributions $\hat{\mathcal{T}}_R$ such that

$$\mathcal{T} = 1 + \sum_R \hat{\mathcal{T}}_R. \quad (2)$$

Each local contribution $\hat{\mathcal{T}}_R$ acts only within some augmentation region Ω_R enclosing the atom. This implies that AE and PS wave functions coincide outside the augmentation regions. The equivalent of the augmentation region in the linear methods is the muffin-tin or atomic sphere. In the pseudopotential method the augmentation region corresponds to the so-called core region.

The local terms $\hat{\mathcal{T}}_R$ are defined for each augmentation region individually by specifying the target functions $|\phi_i\rangle$ of the transformation $\mathcal{T}$ for set of initial functions $|\tilde{\phi}_i\rangle$ that is orthogonal to the core states and otherwise complete in the augmentation region,²⁰ namely, by $|\phi_i\rangle = (1 + \hat{\mathcal{T}}_R)|\tilde{\phi}_i\rangle$ within Ω_R . I call the initial states $|\tilde{\phi}_i\rangle$ PS partial waves and the corresponding target functions $|\phi_i\rangle$ AE partial waves. A natural choice for these functions for the AE partial waves are solutions of the radial Schrödinger equation for the isolated atom, which are orthogonalized to the core states if necessary. Hence the index i refers to the atomic site R , the angular momentum quantum numbers $L = (\ell, m)$, and an additional

index n to label different partial waves for the same site and angular momentum. For each such AE partial wave let us choose a PS partial wave denoted by $|\phi_i\rangle$. The PS partial waves must be identical to the corresponding AE partial waves outside the augmentation region and should themselves form a complete set of functions within the augmentation region. The remaining degree of freedom in the choice of the PS partial waves will be exploited to map the physically relevant AE wave functions onto computationally convenient PS wave functions. In our case these are smooth functions.

This formal definition must be turned into a closed expression for the transformation operator. We make use of the fact that, within the augmentation region, every PS wave function can be expanded into PS partial waves:

$$|\tilde{\Psi}\rangle = \sum_i |\tilde{\phi}_i\rangle c_i \quad \text{within } \Omega_R. \quad (3)$$

Since $|\phi_i\rangle = \mathcal{T}|\tilde{\phi}_i\rangle$, the corresponding AE wave function is of the form

$$|\Psi\rangle = \mathcal{T}|\tilde{\Psi}\rangle = \sum_i |\phi_i\rangle c_i \quad \text{within } \Omega_R, \quad (4)$$

with identical coefficients c_i in both expansions. Hence we can express the AE wave function as

$$|\Psi\rangle = |\tilde{\Psi}\rangle - \sum_i |\tilde{\phi}_i\rangle c_i + \sum_i |\phi_i\rangle c_i, \quad (5)$$

where the expansion coefficients for the partial wave expansions remain to be determined.

Since we require the transformation $\mathcal{T}$ to be linear, the coefficients must be linear functionals of the PS wave functions. Hence the coefficients are scalar products

$$c_i = \langle \tilde{p}_i | \tilde{\Psi} \rangle \quad (6)$$

of the PS wave function with some fixed functions $\langle \tilde{p}_i |$, which I will call projector functions. There is exactly one projector function for each PS partial wave.

The projector functions must fulfill the condition $\sum_i |\tilde{\phi}_i\rangle \langle \tilde{p}_i| = 1$ within Ω_R , so that the one-center expansion $\sum_i |\tilde{\phi}_i\rangle \langle \tilde{p}_i | \tilde{\Psi} \rangle$ of a PS wave function is identical to the PS wave function $|\tilde{\Psi}\rangle$ itself. This implies that

$$\langle \tilde{p}_i | \tilde{\phi}_j \rangle = \delta_{ij}. \quad (7)$$

The projector functions are localized in the augmentation region, even though more extended projector functions could in principle also be chosen. The most general form for the projector functions is $\langle \tilde{p}_i | = \sum_j (\{ \langle f_k | \tilde{\phi}_i \rangle \})_{ij}^{-1} \langle f_j |$, where the $|f_j\rangle$ form an arbitrary, linearly independent set of functions. The projector functions are localized if the functions $|f_i\rangle$ are localized. The reader interested at this point in a practical procedure to determine partial waves and projector functions might wish to jump to Sec. VI.

In summary, a linear transformation

$$\mathcal{T} = 1 + \sum_i (|\phi_i\rangle - |\tilde{\phi}_i\rangle) \langle \tilde{p}_i| \quad (8)$$

between the valence wave functions and fictitious PS wave functions has been established. Using this transformation, the AE wave function can be obtained from the PS wave function by

$$|\Psi\rangle = |\tilde{\Psi}\rangle + \sum_i (|\phi_i\rangle - |\tilde{\phi}_i\rangle) \langle \tilde{p}_i | \tilde{\Psi} \rangle. \quad (9)$$

The three quantities that determine this transformation are (i) the AE partial waves $|\phi_i\rangle$ obtained by radially integrating the Schrödinger equation of the atomic energy for a set of energies ϵ_i^1 and orthogonalization to the core states; (ii) one PS partial wave $|\tilde{\phi}_i\rangle$, which coincides with the corresponding AE partial wave outside some augmentation region for each AE partial wave; and (iii) one projector function $|\tilde{p}_i\rangle$ for each PS partial wave localized within the augmentation region and which obeys the relation $\langle \tilde{p}_i | \tilde{\phi}_j \rangle = \delta_{ij}$.

The partial waves are functions on a radial grid, multiplied with spherical harmonics. In our case the PS wave functions are expanded into plane waves, but other choices are equally possible. The projectors are also calculated as a radial function times spherical harmonics, but are then transformed into the same representation as the PS wave functions, which, in our case, is a plane-wave representation. Since the projectors are tied to the atomic positions and since their shape is independent of the potential, their Fourier components are expressed as a product of a form factor and a structure factor.

The core states $|\Psi^c\rangle$ are decomposed in a way similar to the valence wave functions. They are decomposed into three contributions: a PS core wave function $|\tilde{\Psi}^c\rangle$, which is identical to the true core state outside the augmentation region and a smooth continuation inside; an “AE core partial wave” $|\phi^c\rangle$, which is identical to the AE core state $|\Psi^c\rangle$ and is expressed as a radial function times spherical harmonics; and finally a “PS core partial wave” $|\tilde{\phi}^c\rangle$, which is identical to the PS core state $|\tilde{\Psi}^c\rangle$, but represented as a radial function times spherical harmonics. The core state is therefore expressed as

$$|\Psi^c\rangle = |\tilde{\Psi}^c\rangle + |\phi^c\rangle - |\tilde{\phi}^c\rangle. \quad (10)$$

In contrast to the valence states, no projector functions need be defined for the core states, and the “coefficients” of the one-center terms are always unity. Furthermore, consistent with the frozen-core approximation, the core states are imported from an isolated atom. In practice, a soft core scheme with core states that adjust to the instantaneous potential is also conceivable (see Sec. VI D 2), but has not been implemented. In the following, the core states are implicitly included when summing over energy states. Note that the corresponding coefficients are not defined via the scalar product with a PS wave function, even though, for the sake of simplicity, I will still use the symbol for all states.

It should be noted that the frozen-core approximation allows certain nontrivial changes of the core wave function during the self-consistency or molecular-dynamics simulation. The frozen-core approximation only restricts the variational degree of freedom to a simple unitary transformation among the core states (and occupied va-

lence states). It does allow mixing among the core states due to changing potential. Therefore, to test the accuracy of the frozen-core approximation one should never compare the core states of the isolated atom on a one-to-one basis with those obtained from a relaxed-core calculation in a crystal or molecule.

At this point I will not discuss the components of the projector augmented-wave (PAW) method further. They are described in Sec. VI of this article. I will, however, continue to impose the condition that the AE and PS partial waves form complete sets of functions within the augmentation regions. In practical calculations the number of partial waves and projectors needs to be truncated. The way to truncate the series and the errors involved are also described in detail in later sections of the paper.

Here and in the following I will make extensive use of Dirac's bra and ket notation. A wave function in real space is written as $\langle r|\Psi\rangle = \Psi(r)$; its complex conjugate function is $\langle\Psi|r\rangle = \Psi^*(r)$. The Fourier components of the wave function are $\langle G|\Psi\rangle = \Psi(G)$ with a similar definition of its complex conjugate. A plane wave is of the form $\langle r|G\rangle = \exp(iGr)$. I have adopted the convention for the Fourier transform that the forward transform of a function f is $\langle r|f\rangle = \sum_G \langle r|G\rangle \langle G|f\rangle$ and the backward transform is of the form $\langle G|f\rangle = 1/V \int_V dr \langle G|r\rangle \langle r|f\rangle$, where V is the volume of the unit cell.

B. Operators

Since in the PAW method the PS wave functions instead of the AE functions play the role of the variational parameters, we need to be able to obtain observable quantities as the expectation values of the PS wave functions. As the representation of the wave functions has been changed, we also need to transform our operators into new, so-called PS operators.

Consider some operator A : Its expectation value $\langle A\rangle = \sum_n f_n \langle \Psi_n|A|\Psi_n\rangle$, where n is the band index and f_n is the occupation of the state, can be obtained alternatively as $\langle A\rangle = \sum_n f_n \langle \tilde{\Psi}_n|\tilde{A}|\tilde{\Psi}_n\rangle$. For quasiloc operators, such as the kinetic-energy operator $-\nabla^2/2$ and the real-space projection operator $|r\rangle\langle r|$, which are needed to evaluate total energy and charge density, the PS operator has the form

$$\begin{aligned} \tilde{A} &= \mathcal{T}^\dagger A \mathcal{T} \\ &= A + \sum_{i,j} |\tilde{p}_i\rangle (\langle \phi_i|A|\phi_j\rangle - \langle \tilde{\phi}_i|A|\tilde{\phi}_j\rangle) \langle \tilde{p}_j|. \end{aligned} \quad (11)$$

To arrive at Eq. (11) I expanded $\mathcal{T}$ using Eq. (8) and deleted terms that cancel because $\sum_i |\tilde{\phi}_i\rangle \langle \tilde{p}_i| = 1$ within the augmentation region Ω_R and $|\tilde{\phi}_i\rangle = |\phi_i\rangle$ outside the augmentation region. Note that only on-site terms contribute.

For truly nonlocal operators we need to add a term ΔA to the expression Eq. (11) given as

$$\begin{aligned} \Delta A &= \sum_i |\tilde{p}_i\rangle (\langle \phi_i| - \langle \tilde{\phi}_i|) A \left(1 - \sum_j |\tilde{\phi}_j\rangle \langle \tilde{p}_j| \right) \\ &\quad + (1 - |\tilde{p}_j\rangle \langle \tilde{\phi}_j|) A (|\phi_i\rangle - |\tilde{\phi}_i\rangle) \langle \tilde{p}_i|. \end{aligned} \quad (12)$$

The general form of an operator is strongly reminiscent of generalized separable pseudopotentials.²¹ The PS operator contains three parts: The first part is an operator that directly acts on the PS wave function and is evaluated either in real or reciprocal space. The remaining two parts contain the projectors and the expectation value of the operator either between the AE or the PS partial waves, which can be easily obtained on radial grids using spherical harmonics and Clebsch-Gordan coefficients. If the partial waves are unbound, the individual terms A^1 and $\tilde{A}^1$ are not defined. However, since the PS and AE partial waves are identical outside the augmentation region, these tails cancel exactly for each pair of partial waves. In practice, this problem is solved by truncating the AE and PS partial waves somewhere outside the augmentation region in a completely identical way.

There is an additional freedom to add a term of the form

$$B - \sum_{i,j} |\tilde{p}_i\rangle \langle \tilde{\phi}_i| B |\tilde{\phi}_j\rangle \langle \tilde{p}_j| \quad (13)$$

to the right-hand side of Eq. (11), where B is an arbitrary operator that is localized within the augmentation regions. It is easily shown that the expectation value of this term is zero for any PS wave function, since $|\tilde{\Psi}\rangle = \sum_i |\tilde{\phi}_i\rangle \langle \tilde{p}_i|\tilde{\Psi}\rangle$ within the augmentation region. This freedom can be exploited when the operator A cannot easily be evaluated in a plane-wave expansion. An example is the Coulomb potential of the nucleus, which is problematic due to its singularity at the nuclear site. In this case we may construct a new potential that is identical to the true potential outside the augmentation region and a smooth continuation inside. The difference between the two potentials is localized within the augmentation region and therefore can act as the operator B in Eq. (13). By adding a term of the kind of Eq. (13) to the PS electrostatic potential obtained from Eq. (11), we can cancel the Coulomb singularity in the plane-wave part and obtain an expression that is less sensitive to a truncation of number of plane waves.

We now obtain the charge density following the description given in Eq. (11). The charge density at a point r in space is the expectation value of the real-space projection operator $|r\rangle\langle r|$. Hence the charge density is given by

$$n(r) = \tilde{n}(r) + n^1(r) - \tilde{n}^1(r), \quad (14)$$

where

$$\tilde{n}(r) = \sum_n f_n \langle \tilde{\Psi}_n|r\rangle \langle r|\tilde{\Psi}_n\rangle, \quad (15)$$

$$n^1(r) = \sum_{n,(i,j)} f_n \langle \tilde{\Psi}_n|\tilde{p}_i\rangle \langle \phi_i|r\rangle \langle r|\phi_j\rangle \langle \tilde{p}_j|\tilde{\Psi}_n\rangle, \quad (16)$$

and

$$\tilde{n}^1(r) = \sum_{n,(i,j)} f_n \langle \tilde{\Psi}_n|\tilde{p}_i\rangle \langle \tilde{\phi}_i|r\rangle \langle r|\tilde{\phi}_j\rangle \langle \tilde{p}_j|\tilde{\Psi}_n\rangle. \quad (17)$$

Note also that n^1 contains the contribution of the core states $\sum_n \langle \phi_n^c | r \rangle \langle r | \phi_n^c \rangle$ and that $\tilde{n}^1$ as well as $\tilde{n}$ contain the contribution of the PS core states $\sum_n \langle \tilde{\phi}_n^c | r \rangle \langle r | \tilde{\phi}_n^c \rangle$ and $\sum_n \langle \tilde{\Psi}_n^c | r \rangle \langle r | \tilde{\Psi}_n^c \rangle$, respectively.

In practice, we do not construct a PS core state for each core state individually unless we are interested in the physics related directly to the core states. Instead, we construct only a PS core density.

C. Total energy

Similar to the expectation values, the expression for the total energy functional

$$E = \sum_n f_n \langle \Psi_n | -\frac{1}{2} \nabla^2 | \Psi_n \rangle + \frac{1}{2} \int dr \int dr' \frac{(n + n^Z)(n + n^Z)}{|r - r'|} + \int dr n \epsilon_{xc}(n) \quad (18)$$

can also be divided as $E = \tilde{E} + E^1 - \tilde{E}^1$, into a smooth part $\tilde{E}$, which is evaluated on regular grids in Fourier or real space, and two one-center contributions E^1 and $\tilde{E}^1$, which are evaluated on radial grids in an angular momentum representation. Let us denote the point charge density of the nucleus by n^Z and the energy per electron from exchange and correlation as ϵ_{xc} . Here and in the following I use hartree atomic units ($\hbar = e = m_e = 1$). The three contributions to E are

$$\begin{aligned} \tilde{E} = & \sum_n f_n \langle \tilde{\Psi}_n | -\frac{1}{2} \nabla^2 | \tilde{\Psi}_n \rangle \\ & + \frac{1}{2} \int dr \int dr' \frac{(\tilde{n} + \hat{n})(\tilde{n} + \hat{n})}{|r - r'|} + \int dr \tilde{n} \bar{v} \\ & + \int dr \tilde{n} \epsilon_{xc}(\tilde{n}), \end{aligned} \quad (19)$$

$$\begin{aligned} E^1 = & \sum_{n,(i,j)} f_n \langle \tilde{\Psi}_n | \tilde{p}_i \rangle \langle \phi_i | -\frac{1}{2} \nabla^2 | \phi_j \rangle \langle \tilde{p}_j | \tilde{\Psi}_n \rangle \\ & + \frac{1}{2} \int dr \int dr' \frac{(n^1 + n^Z)(n^1 + n^Z)}{|r - r'|} \\ & + \int dr n^1 \epsilon_{xc}(n^1), \end{aligned} \quad (20)$$

$$\begin{aligned} \tilde{E}^1 = & \sum_{n,(i,j)} f_n \langle \tilde{\Psi}_n | \tilde{p}_i \rangle \langle \tilde{\phi}_i | -\frac{1}{2} \nabla^2 | \tilde{\phi}_j \rangle \langle \tilde{p}_j | \tilde{\Psi}_n \rangle \\ & + \frac{1}{2} \int dr \int dr' \frac{(\tilde{n}^1 + \hat{n})(\tilde{n}^1 + \hat{n})}{|r - r'|} + \int dr \tilde{n}^1 \bar{v} \\ & + \int dr \tilde{n}^1 \epsilon_{xc}(\tilde{n}^1). \end{aligned} \quad (21)$$

The potential $\bar{v}$ is an arbitrary potential localized in the augmentation regions. Its contribution to the total energy vanishes exactly because $\tilde{n} = \tilde{n}^1$ within the augmentation region. Since the potential $\bar{v}$ contributes only if the partial wave expansion is not complete, it is used to minimize truncation errors.

In addition, let us introduce the so-called compensation charge density $\hat{n}$. After adding an appropriate compensation charge density to the PS charge density and its one-center expansion, the difference of the AE and the PS one-center contributions $(n^1 + n^Z) - (\tilde{n}^1 + \hat{n})$ to the charge density has vanishing electrostatic multipole moments and hence no longer interacts with charges outside the augmentation region: This energy has been transferred to $\tilde{E}$. Here I made use of the fact that a localized charge distribution produces a potential that, outside the region of localization, depends only on the electrostatic multipole moments, but not on the shape of the charge distribution.

The identity $E = \tilde{E} + E^1 - \tilde{E}^1$ for a complete set of partial waves can be seen as follows. (Those not interested to follow through this detail may proceed to the next paragraph.) One divides space into augmentation regions and an interstitial region. Now we use the identities $n = n^1$ and $\tilde{n}^1 = \tilde{n}$ inside the augmentation regions and the identities $n = \tilde{n}$ and $n^1 = \tilde{n}^1$ in the interstitial region. One can convince oneself easily that the decomposition is true for the kinetic energy (see Sec. IIB), for the exchange and correlation energy, and for the term proportional to $\bar{v}$. The decomposition for the electrostatic energy is more complex to show: Let us add a charge density $n^1 + n^Z - \tilde{n}^1 - \hat{n}$ to $\tilde{n} + \hat{n}$ in Eq. (19) and to $\tilde{n}^1 + \hat{n}$ in Eq. (21). The effect of this addition vanishes: First, the term quadratic in $n^1 + n^Z - \tilde{n}^1 - \hat{n}$ cancels exactly because $\tilde{E}$ and $\tilde{E}^1$ are added with opposite sign. Second, the terms linear in $n^1 + n^Z - \tilde{n}^1 - \hat{n}$ are proportional to $\tilde{n} - \tilde{n}^1$, which is zero within the augmentation regions, and to the electrostatic potential of $n^1 + n^Z - \tilde{n}^1 - \hat{n}$, which is zero in the interstitial region, because the density itself is localized within the augmentation region and has zero electrostatic multipole moments. Once this term has been added, the electrostatic contributions of the last two terms Eqs. (20) and (21) are identical and cancel, while the first term is the true electrostatic interaction of the full charge density $n + n^Z = \tilde{n} + n^1 - \tilde{n}^1 + n^Z$. This special form of the total energy has been chosen in order to obtain a strict separation into partial-wave and plane-wave expansions and to achieve rapid convergence for both expansions.

The compensation charge density $\hat{n} = \sum_R \hat{n}_R$ with

$$\hat{n}_R(r) = \sum_L g_{RL}(r) Q_{RL} \quad (22)$$

is expressed as a sum of generalized Gaussians

$$g_{RL}(r) = C_\ell |r - R|^\ell Y_L(r - R) e^{-(|r - R|/r_c)^\alpha} \quad (23)$$

with the normalization constant C_ℓ determined such that its multipole moment $\int dr r^\ell Y_L(r) g_L(r)$ is unity. Y_L is a spherical harmonic function or its real counterpart. The decay length r_c is sufficiently small so that the compensation charge density is localized within the augmentation regions. The value of r_c depends on the particular atom type; R stands for a particular nuclear site and $L = (\ell, m)$ represents the angular momenta in the spherical harmonics expansion. The multipole moments Q_{RL} are given by

$$Q_{RL} = \int dr |r - R|^\ell [n_R^1(r) + n_R^Z(r) - \tilde{n}_R^1(r)] Y_L^*(r - R). \quad (24)$$

Since the Gaussians are required to decay within the augmentation regions, they often have high Fourier components. This would require a large plane-wave cutoff in the PS charge density. The problem is solved by a well-known trick already used in the pseudopotential approach:²² We introduce a second, primed compensation charge density $\hat{n}'$, which has the same multipole moments as $\hat{n}$, but uses generalized Gaussians $g'_{RL}(r)$ with a larger decay constant r'_c than the unprimed compensation charge density. It may extend over several atomic sites, but should not contribute higher Fourier components than the PS charge density itself does.

Now we rewrite the electrostatic energy in $\tilde{E}$:

$$\begin{aligned} \frac{1}{2} \int dr \int dr' \frac{(\tilde{n} + \hat{n})(\tilde{n} + \hat{n}')}{|r - r'|} \\ = \frac{1}{2} \int dr \int dr' \frac{(\tilde{n} + \hat{n}')(\tilde{n} + \hat{n}')}{|r - r'|} \\ + \int dr \tilde{n}(r) \hat{v}(r) + \sum_{R,R'} U_{R,R'}. \end{aligned} \quad (25)$$

The first term in the new expression (25) involves only smooth functions and can be evaluated in Fourier space as

$$2\pi V \sum_G \frac{|\tilde{n}(G) + \hat{n}'(G)|^2}{G^2}. \quad (26)$$

The second part of Eq. (25) introduces a potential

$$\hat{v}(r) = \int dr' \frac{\hat{n}(r') - \hat{n}'(r')}{|r - r'|}, \quad (27)$$

which has high Fourier components just as the original compensation charge density does. However, they do not contribute to the total energy because they are multiplied with the high Fourier components of the PS charge density, which are exactly zero if a plane-wave cutoff is imposed. Hence this term can also be exactly evaluated in Fourier space. The spacial extent of this potential in real space is identical to that of the smooth compensation charge density $\hat{n}'_R$.

The last term in Eq. (25) is a short-ranged pair potential between the atoms

$$U_{R,R'} = \frac{1}{2} \int dr \int dr' \frac{\hat{n}_R(r) \hat{n}_{R'}(r') - \hat{n}'_R(r) \hat{n}'_{R'}(r')}{|r - r'|}, \quad (28)$$

which can be evaluated analytically.^{23–25} The range of this pair potential is twice that of the smooth compensation charge densities $\hat{n}'$. It depends explicitly on the charge distributions via the multipole moments Q_{RL} . Note that the potential $\hat{v}$ and the pair potential $U_{R,R'}$ contain nonspherical terms and adjust to the actual charge density.

Finally, we need to evaluate the energy of exchange and correlation for a one-center expansion. We adopt a procedure from previous full-potential LMTO calculations^{26,27} and expand the corresponding energy density in the deviation of the one-center charge density from its spherical part $n_{R,\ell=0}^1$:

$$\begin{aligned} \int dr n_R^1 \epsilon_{xc}(n_R^1) &= \int dr n_{R,\ell=0}^1 \epsilon_{xc}(n_{R,\ell=0}^1) \\ &+ \frac{1}{2} \sum_{L,L' \neq 0} \int dr \frac{\partial \mu_{xc}(n_{R,\ell=0}^1)}{\partial n_R^1} (n_{R,L}^1)^2 \\ &+ O((n_{R,L}^1)^3), \end{aligned} \quad (29)$$

where $\mu_{xc}(n) = d[n\epsilon_{xc}(n)]/dn$. The angular momentum components of the one-center charge density are denoted by $n_{R,L}^1$. In practice a Taylor expansion up to the second order has shown to be sufficiently accurate.²⁶ The one-center contribution of the PS charge density is treated identically.

III. FROM AN EXACT FORMALISM TO A PRACTICAL SCHEME

Up to this point the PAW method is an exact implementation of the density-functional theory within the frozen-core approximation. However, we have required certain completeness conditions for the plane-wave basis set for the PS wave functions and the AE and PS partial waves. In order to arrive at a practical scheme, let me now introduce two approximations.

(i) Plane waves are included only up to a given plane-wave cutoff E_{PW} defined as the maximum of $G^2/2$.

(ii) The number of AE partial waves, PS partial waves, and projectors is finite. However, the truncation of AE and PS partial waves and projector functions are done in exactly the same way. That is, for each AE partial wave there is a corresponding PS partial wave and its projector function.

Both approximations can be controlled in a straightforward way, by increasing either the plane-wave cutoff and/or the number of partial waves. The convergence for both is rapid if a suitable set of partial waves and projectors has been selected. Typically good convergence is obtained for plane-wave cutoffs of 30 Ry and one or two partial waves per site and angular momentum, with a maximum angular momentum of typically $\ell = 1$ or $\ell = 2$. The partial-wave truncation will be discussed in detail in Sec. VII.

The two approximations define a new total energy functional, and we have to establish that this new functional is sufficiently close to the correct functional for the physically relevant states. Once this new functional is defined, no further approximations are allowed because they would destroy the energy conservation in a molecular-dynamics simulation. Energy conservation is the most important test of the quality of any molecular-dynamics simulation. Many previous electronic structure methods have concentrated on providing a satisfac-

tory description of the potential. For molecular-dynamics simulations the primary quantity is the total energy functional because small inconsistencies between forces and total energy can create substantial difficulties in a simulation. An accurate description of the potential follows, of course, from an accurate description of the total energy functional.

There are two further approximations that are not necessary, but are employed to accelerate the calculations: One can introduce a plane-wave cutoff in the representation of the PS charge density and an angular momentum truncation in the one-center PS and AE densities. Without these cutoffs, the PS charge density has plane-wave components corresponding to four times the plane-wave cutoff for the wave function and the one-center expansions have angular momentum components of up to twice the maximum angular momentum of the partial waves. However, a number of these terms contribute little to the total energy, so that these approximations are convenient ways to save computation time. One can truncate the angular momentum expansion safely at $\ell = 2$, and the plane-wave cutoff for the density can be chosen in many cases to be only twice the value of the wave function.

IV. FORCES, HAMILTON OPERATORS, AND OVERLAP MATRICES

In order to find the ground state of the density functional or to propagate wave functions and atoms in a molecular-dynamics simulation, one needs to calculate the gradients of the total energy functional with respect to all the variational parameters, namely, the PS wave functions and the atomic positions. In the following subsections I shall derive explicit expressions for forces, Hamilton operators, and overlap matrices.

A. Overlap operator

The overlap matrix in the AE representation is given simply by the matrix elements of the unity operator. Consequently plane waves form an orthogonal basis set in the AE representation. The PS version of the unity operator obtained via Eq. (11), however, is a nonlocal operator of the form

$$\tilde{O} = 1 + \sum_{i,j} |\tilde{p}_i\rangle [\langle \phi_i | \phi_j \rangle - \langle \tilde{\phi}_i | \tilde{\phi}_j \rangle] \langle \tilde{p}_j |. \quad (30)$$

Hence, in the PS representation, plane waves are no longer orthogonal, that is, $\langle G | \tilde{O} | G' \rangle \neq \delta_{G,G'}$, if the PS overlap operator $\tilde{O}$ differs from unity. This is a direct consequence of relaxing the condition of norm conservation: The PS overlap operator obviously reduces to the unity operator if the norm-conservation condition $\langle \tilde{\phi}_i | \tilde{\phi}_j \rangle = \langle \phi_i | \phi_j \rangle$ is imposed.

B. Hamilton operator

The Hamilton operator is the first derivative of the total energy functional with respect to the density operator

$\rho = \sum_n |\Psi_n\rangle f_n \langle \Psi_n|$, where f_n denotes the occupations and $|\Psi_n\rangle$ the orthogonal eigenfunctions of the density operator. This can be explained as follows: Since the expectation value of any one-particle operator A is the trace $\langle A \rangle = \text{Tr}[\rho A]$ of the product between the one-particle operator and the density operator, the first derivative of the total energy with respect to the density operator can be written as

$$\begin{aligned} \frac{\partial E}{\partial \rho} &= \frac{\partial \text{Tr}[-\frac{1}{2} \nabla^2 \rho]}{\partial \rho} + \int dr \frac{\partial E}{\partial n(r)} \frac{\partial \text{Tr}[|r\rangle \langle r| \rho]}{\partial \rho} \\ &= -\frac{1}{2} \nabla^2 + v, \end{aligned} \quad (31)$$

where the potential is $v(r) = |r\rangle \frac{\partial E}{\partial n(r)} \langle r|$, which is the well-known form of the Hamilton operator.

As the variational parameters of the PAW method are the PS wave functions, we construct a PS Hamilton operator $\tilde{H}$ defined as the derivative of the total energy, given by Eqs. (19)-(21), with respect to the PS density operator $\tilde{\rho} = \sum_i |\tilde{\Psi}_i\rangle f_i \langle \tilde{\Psi}_i|$ with wave functions that obey the orthogonality condition $\langle \tilde{\Psi}_n | \tilde{O} | \tilde{\Psi}_m \rangle = \delta_{nm}$. In this section we shall derive the explicit expressions for the PS Hamilton operator that will be needed to set up the Kohn-Sham equations or the equations of motion for first-principles molecular dynamics.

Let us treat the potential energy as a functional of four arguments $\tilde{n}$, n^1 , $\tilde{n}^1$, and the multipole moments Q_{RL} . The multipole moments, which determine the compensation charge densities, are themselves unique functions of the one-center densities. This choice—which is not an approximation—will simplify the bookkeeping in the following derivation. The derivative with respect to the PS density operator is then obtained as

$$\begin{aligned} \frac{\partial E}{\partial \tilde{\rho}} &= \frac{\partial \text{Tr}[\tilde{\rho} \tilde{T}]}{\partial \tilde{\rho}} + \int dr \frac{\partial E}{\partial \tilde{n}} \frac{\partial \tilde{n}}{\partial \tilde{\rho}} \\ &+ \int dr \left(\frac{\partial E}{\partial n^1} + \sum_{R,L} \frac{\partial E}{\partial Q_{RL}} \frac{\partial Q_{RL}}{\partial n^1} \right) \frac{\partial n^1}{\partial \tilde{\rho}} \\ &+ \int dr \left(\frac{\partial E}{\partial \tilde{n}^1} + \sum_{R,L} \frac{\partial E}{\partial Q_{RL}} \frac{\partial Q_{RL}}{\partial \tilde{n}^1} \right) \frac{\partial \tilde{n}^1}{\partial \tilde{\rho}}, \end{aligned} \quad (32)$$

where

$$\begin{aligned} \tilde{T} &= -\frac{1}{2} \nabla^2 \\ &+ \sum_{i,j} |\tilde{p}_i\rangle \left(\langle \phi_i | -\frac{1}{2} \nabla^2 | \phi_j \rangle - \langle \tilde{\phi}_i | -\frac{1}{2} \nabla^2 | \tilde{\phi}_j \rangle \right) \langle \tilde{p}_j | \end{aligned} \quad (33)$$

is the PS version of the kinetic-energy operator $T = -\nabla^2/2$. Note that the three densities $\tilde{n}$, n^1 , and $\tilde{n}^1$ are linear functions of the PS density operator and their derivatives with respect to the PS density operator are obtained easily from Eqs. (15)–(17).

The individual terms in Eq. (32) are evaluated as follows.

(i) The derivative with respect to the PS charge density is obtained from Eqs. (19) and (25) as

$$\tilde{v}(r) = \frac{\partial E}{\partial \tilde{n}(r)} = \int dr' \frac{\tilde{n}(r') + \hat{n}'(r')}{|r - r'|} + \tilde{v}(r) + \bar{v}(r) + \mu_{xc}[\tilde{n}(r)]. \quad (34)$$

(ii) Since the multipole moments enter the total energy expression only via the compensation charge densities $\tilde{n}$ and $\hat{n}'$, the corresponding derivative of the total energy is

$$\frac{\partial E}{\partial Q_{RL}} = \int dr \frac{\partial E}{\partial \tilde{n}(r)} \frac{\partial \tilde{n}(r)}{\partial Q_{RL}} + \int dr \frac{\partial E}{\partial \hat{n}'(r)} \frac{\partial \hat{n}'(r)}{\partial Q_{RL}}. \quad (35)$$

In order to obtain energies and forces in a fully consistent manner—a requirement for exact energy conservation—we must not evaluate the derivative of a term in one representation if the total energy is evaluated in another. Therefore we divide this expression further, by following exactly the way the corresponding total energy terms are evaluated. These representations are the Fourier mesh denoted by M and the radial grid RG . A stands for analytical evaluation as used for the pair potential defined by Eq. (28). We divide the energy derivative with respect to the multipole moments further into

$$\begin{aligned} \frac{\partial E}{\partial Q_{RL}} = & \int_M dr \int_M dr' \frac{g_{RL}(r)\tilde{n}(r') + g'_{RL}(r)\hat{n}'(r')}{|r - r'|} \\ & + \int_A dr \int_A dr' \frac{g_{RL}(r)\tilde{n}(r') - g'_{RL}(r)\hat{n}'(r')}{|r - r'|} \\ & - \int_{RG} dr \int_{RG} dr' \frac{g_{RL}(r)[\tilde{n}^1(r') + \hat{n}'(r')]}{|r - r'|}, \end{aligned} \quad (36)$$

using Eqs. (19), (21), (25), (27), and (28).

(iii) Using Eq. (24) we can resolve $\partial Q_{RL}/\partial n^1(r)$ and $\partial Q_{RL}/\partial \tilde{n}^1(r)$. We define the potential

$$\begin{aligned} v_R^0(r) = & \sum_L \frac{\partial E}{\partial Q_{RL}} \frac{\partial Q_{RL}}{\partial n^1(r)} = - \sum_L \frac{\partial E}{\partial Q_{RL}} \frac{\partial Q_{RL}}{\partial \tilde{n}^1(r)} \\ = & \sum_L (r - R)^L Y_L^*(|r - R|) \frac{\partial E}{\partial Q_{RL}}. \end{aligned} \quad (37)$$

(iv) With the help of Eqs. (20), (21), and (37), we evaluate the potentials

$$\begin{aligned} v_R^1(r) = & \frac{\partial E}{\partial n^1(r)} + \sum_L \frac{\partial E}{\partial Q_{RL}} \frac{\partial Q_{RL}}{\partial n^1(r)} \\ = & \int_R dr' \frac{n_R^1(r') + n_R^2(r')}{|r - r'|} + \mu_{xc}[n_R^1(r)] + v_R^0(r) \end{aligned} \quad (38)$$

and

$$\begin{aligned} \tilde{v}_R^1(r) = & - \left(\frac{\partial E}{\partial \tilde{n}^1(r)} + \frac{\partial E}{\partial Q_{RL}} \frac{\partial Q_{RL}}{\partial \tilde{n}^1(r)} \right) \\ = & \int_R dr' \frac{\tilde{n}_R^1(r') + \hat{n}_R(r')}{|r - r'|} + \mu_{xc}[\tilde{n}_R^1(r)] + v_R^0(r). \end{aligned} \quad (39)$$

The potential $v_R^1(r)$ is the one-center expansion of the AE potential and $\tilde{v}_R^1(r)$ is the one-center expansion of the PS potential $\tilde{v}(r)$ at site R .

(v) The expressions (33), (34), (38), and (39) can now be combined using Eq. (32) to form the PS Hamiltonian

$$\begin{aligned} \tilde{H} = & -\frac{1}{2}\nabla^2 + \tilde{v} + \sum_{i,j} |\tilde{p}_i\rangle (\langle \phi_i| - \frac{1}{2}\nabla^2 + v^1|\phi_j\rangle \\ & - \langle \tilde{\phi}_i| - \frac{1}{2}\nabla^2 + \tilde{v}^1|\tilde{\phi}_j\rangle) \langle \tilde{p}_j|. \end{aligned} \quad (40)$$

The full potential

$$v(r) = \tilde{v}(r) + v^1(r) - \tilde{v}^1(r) \quad (41)$$

is, like the full charge density, a superposition of a smooth plane-wave part and two one-center expansions per site. The smooth part is a plane-wave sum and the one-center expansions are radial functions times spherical harmonics. The gradient of the total energy functional with respect to the PS wave functions is then obtained as

$$\left. \frac{\partial E[\tilde{\Psi}, R]}{\partial \langle \tilde{\Psi}_n |} \right|_R = \tilde{H}|\tilde{\Psi}_n\rangle f_n. \quad (42)$$

C. Forces

1. Force theorems

Several force theorems have been discussed in the literature. Most of them exploit the variational principle for the electronic wave functions, which says that the total energy is insensitive to the first order to a change in the charge density. In other words, the forces acting on the electrons vanish in the electronic ground state. In the Hellmann-Feynman theorem^{28–30} the force on the atoms follows from an infinitesimal distortion of the atomic positions alone, while changes of the electronic wave functions do not contribute if the latter are determined variationally. In the so-called “force theorem,”^{31–34} not only the nucleus, but also the electronic charge density within some arbitrary volume enclosing the nucleus is rigidly displaced. Both force theorems produce the same result in the electronic ground state, where no net forces act on the electrons. Otherwise, their results differ, if the Kohn-Sham equations have not been obtained in a fully self-consistent manner.

However, if we use a Lagrangian formalism to calculate the trajectories, with the electronic wave functions and the atomic positions as independent parameters, there is only one choice for calculating the forces. The force must be identical to the partial derivative of the total energy with respect to the atomic positions while keeping the variational parameters of the wave functions fixed, $F_R = -\frac{\partial E}{\partial R}|\tilde{\Psi}\rangle$. Note that the variational parameters refer to the Hilbert space spanned by the occupied one-particle states. If an infinitesimal displacement of the atoms affects the orthonormality of the one-particle states, the latter must be rescaled as described below. This expression for the forces is uniquely defined even

arbitrarily far from the electronic ground state. The reason for using this expression is not that this force theorem is particularly insensitive to deviations from the Born-Oppenheimer surface, but rather that only the direct partial derivative avoids the double counting of forces acting on the wave functions and those acting on the atoms.

To calculate the derivative with respect to the atomic positions, one first has to calculate the forces on the nucleus and second to include the change of the AE wave functions for fixed PS wave functions, but changing atomic positions. This term appears because the augmentation depends on the atomic position. The force on the nucleus is obtained as $F_R^{HF} = -\frac{\partial E}{\partial R}|\Psi\rangle$. It is the product of the electric field at the nucleus and the nuclear charge. Its value is derived from an infinitesimal displacement of the nuclear charge density n^Z . The forces resulting from an infinitesimal change in the wave functions due to the atomic displacement can be written as $F_R^P = -\frac{\partial E}{\partial R}|\tilde{\Psi}\rangle + \frac{\partial E}{\partial R}|\Psi\rangle$ and are called Pulay forces.³⁵ To use the language of the force theorem, the Pulay forces describe forces on the electrons that are dragged along with the nucleus due to the position-dependent basis set. With the PAW method we must consider Pulay forces from the frozen-core electrons,³⁶ which shift rigidly with the nucleus, and the contributions from the augmentation.

When calculating the Pulay forces, we must also consider the change in the overlap between the wave functions. An infinitesimal change of the atomic positions must be accompanied by a change in the wave functions that restores the orthogonality. The new occupied wave functions must span the same portion of the Hilbert space that was occupied before displacement of the atoms. Hence the new wave function can be expressed as a linear combination

$$|\tilde{\Psi}_n(R + dR)\rangle = |\tilde{\Psi}_n(R)\rangle + \sum_m |\tilde{\Psi}_m(R)\rangle \Lambda_{mn} dR \quad (43)$$

of the PS wave functions with undistorted atomic positions. Note that these wave functions should not be confused with the *self-consistent* wave functions for displaced atomic positions. The new wave functions obey the orthogonality condition

$$\langle \tilde{\Psi}_n(R + dR) | \tilde{O}(R + dR) | \tilde{\Psi}_m(R + dR) \rangle = \delta_{nm} \quad (44)$$

to linear order in the displacement, which determines Λ_{nm} by

$$\Lambda_{nm} + \Lambda_{nm}^\dagger = -\langle \tilde{\Psi}_n | \nabla_R \tilde{O} | \tilde{\Psi}_m \rangle. \quad (45)$$

Here ∇_R corresponds to a derivative of a nuclear coordinate rather than to a derivative with respect to an electronic coordinate, which is denoted by ∇ . To arrive at this expression we used the orthonormality of the wave function $|\tilde{\Psi}_n(R)\rangle$ before the displacement. We retain the freedom of adding an arbitrary antisymmetric matrix to Λ , which is reminiscent of the invariance of the total energy with respect to a unitary transformation of the occupied wave functions. This can be seen as follows: A unitary matrix has the general form e^B , where B is an anti-Hermitian matrix (i.e., $B = -B^\dagger$). Hence

a unitary matrix can be written in the first order in B as $1 + B$. Since the total energy is invariant with respect to a unitary transformation between the occupied wave functions, the forces are invariant with respect to the antisymmetric part of Λ , given that Λ is block diagonal separating occupied and unoccupied states. Hence we obtain

$$\nabla_R |\tilde{\Psi}_n\rangle = -\frac{1}{2} \sum_m |\tilde{\Psi}_m\rangle [\langle \tilde{\Psi}_m | \nabla_R \tilde{O} | \tilde{\Psi}_n \rangle + B_{mn}]. \quad (46)$$

Using Eqs. (46) and (47) and

$$\left. \frac{\partial E[\tilde{\Psi}, R]}{\partial R} \right|_{|\tilde{\Psi}_n\rangle} = \sum_n f_n \langle \tilde{\Psi}_n | \nabla_R \tilde{H} | \tilde{\Psi}_n \rangle, \quad (47)$$

we can write the total force including the force on the nucleus and the Pulay force as

$$\begin{aligned} F_R = & - \sum_n f_n \langle \tilde{\Psi}_n | \nabla_R \tilde{H} | \tilde{\Psi}_n \rangle \\ & + \sum_{n,m} \frac{f_n + f_m}{2} \langle \tilde{\Psi}_m | \nabla_R \tilde{O} | \tilde{\Psi}_n \rangle \langle \tilde{\Psi}_n | \tilde{H} | \tilde{\Psi}_m \rangle \\ & + \sum_{n,m} \frac{f_n - f_m}{2} B_{mn} \langle \tilde{\Psi}_n | \tilde{H} | \tilde{\Psi}_m \rangle. \end{aligned} \quad (48)$$

The last term in Eq. (48) describes the effect of electronic excitations resulting from a unitary transformation between occupied and unoccupied states and cannot be specified further. However, its contribution to the force vanishes if the Hamilton matrix $\langle \tilde{\Psi}_n | \tilde{H} | \tilde{\Psi}_m \rangle$ commutes with the occupations (for example, if it is diagonal). This is the case if the wave functions are obtained by diagonalization of a Hamiltonian. In a Lagrangian formalism this term must be chosen in a well-defined way, as described in a later section.

2. Forces in the PAW method

The change of the AE wave functions is related to a displacement of the projector functions and the partial waves

$$\begin{aligned} \nabla_R |\Psi_n\rangle = & - \sum_i (|\nabla \phi_i\rangle - |\nabla \tilde{\phi}_i\rangle) \langle \tilde{p}_i | \tilde{\Psi}_n \rangle \\ & - \sum_i (|\phi_i\rangle - |\tilde{\phi}_i\rangle) \langle \nabla \tilde{p}_i | \tilde{\Psi}_n \rangle \\ & - \frac{1}{2} \sum_m |\Psi_m\rangle \langle \tilde{\Psi}_m | \nabla_R \tilde{O} | \tilde{\Psi}_n \rangle. \end{aligned} \quad (49)$$

The first summation corresponds to a rigid displacement of the partial waves, the second to change of shape of the one-center expansions, and the third is the force due to the change of the overlap. For simplicity, we derive the forces that result from the three contributions independently.

The force resulting from a rigid shift of the partial waves is treated together with those functions that do not

explicitly depend on the PS wave functions but are tied directly to the nuclear positions, such as the rigid shift of the nucleus n_R^Z , the potential $\bar{v}_R$, and the Gaussians g_{RL} and g'_{RL} used to expand the compensation charge densities $\hat{n}_R$ and $\hat{n}'_R$.

Let us first analyze the contribution of the smooth part $\tilde{E}$, given in Eq. (19), of the total energy

$$F_R^{(1)} = \int dr \nabla \hat{n}'_R \int dr' \frac{\tilde{n} + \hat{n}'}{|r - r'|} + \int dr \tilde{n} (\nabla \bar{v}_R + \nabla \bar{v}_R) - \int dr \tilde{n}^c \nabla \bar{v} - \sum_{R'} \nabla_R |Q U_{R,R'}|. \quad (50)$$

The first three terms are evaluated in G space. The fourth is related to the derivative of the pair potential for fixed multipole moments Q_{RL} and is evaluated analytically. Note that the multipole moments do not change under an infinitesimal rigid shift of the partial waves because the center of the Gaussians is also shifted analogously.

Second, let us consider the contribution of the rigid shift from the one-center terms E^1 and $\tilde{E}^1$. Their contribution is exactly zero, because *all* contributing densities and potentials are rigidly shifted, so that the change can be reduced to a change of coordinates, which does not affect the energy. Note that this term also contains the force on the nucleus.

Next we consider the change of shape of the one-center densities. Here we can use quantities that have already been calculated in the Hamiltonian. This force is proportional to the gradient

$$\nabla_R \Theta_{ij} = - \sum_n f_n (\langle \tilde{\Psi}_n | \nabla \tilde{p}_i \rangle \langle \tilde{p}_j | \tilde{\Psi}_n \rangle + \langle \tilde{\Psi}_n | \tilde{p}_i \rangle \langle \nabla \tilde{p}_j | \tilde{\Psi}_n \rangle) \quad (51)$$

of Θ_{ij} , which is defined as

$$\Theta_{ij} = \sum_n f_n \langle \tilde{\Psi}_n | \tilde{p}_i \rangle \langle \tilde{p}_j | \tilde{\Psi}_n \rangle, \quad (52)$$

where Θ_{ij} is a density matrix for the one-center expansions in terms of partial waves.

The force due to the change of shape has the form

$$F_R^{(2)} = - \sum_{i,j} \frac{\partial E}{\partial \Theta_{ij}} \nabla_R \Theta_{ij}. \quad (53)$$

The energy derivative with respect to Θ_{ij} is obtained similarly to the corresponding term in the Hamiltonian [cf. Eq. (32)]

$$\frac{\partial E}{\partial \Theta_{ij}} = \int dr \left(\frac{\partial E}{\partial n^1} + \sum_{R,L} \frac{\partial E}{\partial Q_{RL}} \frac{\partial Q_{RL}}{\partial n^1} \right) \frac{\partial n^1}{\partial \Theta_{ij}} + \int dr \left(\frac{\partial E}{\partial \tilde{n}^1} + \sum_{R,L} \frac{\partial E}{\partial Q_{RL}} \frac{\partial Q_{RL}}{\partial \tilde{n}^1} \right) \frac{\partial \tilde{n}^1}{\partial \Theta_{ij}}, \quad (54)$$

where

$$\frac{\partial n^1}{\partial \Theta_{ij}} = \langle \phi_i | r \rangle \langle r | \phi_j \rangle, \quad (55)$$

$$\frac{\partial \tilde{n}^1}{\partial \Theta_{ij}} = \langle \tilde{\phi}_i | r \rangle \langle r | \tilde{\phi}_j \rangle. \quad (56)$$

This force can be evaluated using the one-center expansions $v^1(r)$ and $\tilde{v}^1(r)$ of the AE and the PS potential

$$F_R^{(2)} = - \sum_{i,j} \nabla_R \Theta_{ij} \left(\langle \phi_i | -\frac{1}{2} \nabla^2 + v^1 | \phi_j \rangle - \langle \tilde{\phi}_i | -\frac{1}{2} \nabla^2 + \tilde{v}^1 | \tilde{\phi}_j \rangle \right). \quad (57)$$

Finally, we consider the forces resulting from the change of the orthogonality of the AE wave functions. The corresponding force is of the form

$$F^{(3)} = \sum_{n,m} \frac{f_n + f_m}{2} \langle \tilde{\Psi}_n | \tilde{H} | \tilde{\Psi}_m \rangle \langle \tilde{\Psi}_m | \nabla_R \tilde{O} | \tilde{\Psi}_n \rangle = - \sum_{i,j;n,m} \frac{f_n + f_m}{2} \langle \tilde{\Psi}_n | \tilde{H} | \tilde{\Psi}_m \rangle \times 2 \text{Re}(\langle \tilde{\Psi}_n | \nabla \tilde{p}_i \rangle \langle \tilde{p}_j | \tilde{\Psi}_m \rangle) (\langle \phi_i | \phi_j \rangle - \langle \tilde{\phi}_i | \tilde{\phi}_j \rangle). \quad (58)$$

The total force is given by the sum of the three terms

$$F_R = F_R^{(1)} + F_R^{(2)} + F_R^{(3)}. \quad (59)$$

V. FIRST-PRINCIPLES MOLECULAR DYNAMICS

A. Fictitious Lagrangian

The first-principles molecular dynamics is implemented in a straightforward manner once the exact expressions for the Hamiltonian and the forces have been obtained. After inclusion of the real kinetic energy of the nuclei and the fictitious kinetic energy of the PS wave functions, we obtain a Lagrangian

$$\mathcal{L} = \sum_n m_\Psi f_n \langle \dot{\tilde{\Psi}}_n | \dot{\tilde{\Psi}}_n \rangle + \sum_R \frac{1}{2} M_R \dot{R}^2 - E[|\tilde{\Psi}_n\rangle, R] + \sum_{n,m} (\langle \tilde{\Psi}_n | \tilde{O} | \tilde{\Psi}_m \rangle - \delta_{nm}) \Lambda_{mn}, \quad (60)$$

where the last term ensures orthogonality of the AE wave functions via the method of Lagrange parameters. The resulting Euler-Lagrange equations have the form

$$m_\Psi |\ddot{\tilde{\Psi}}_n\rangle = -\tilde{H} |\tilde{\Psi}_n\rangle + \sum_m \tilde{O} |\tilde{\Psi}_m\rangle \Lambda_{mn} \frac{1}{f_n} \quad (61)$$

and

$$M_R \ddot{R} = F_R \quad (62)$$

and are integrated with the Verlet algorithm.³⁷ The La-

grange parameters for a system at rest are related to the Hamiltonian via $\Lambda_{nm} = \langle \tilde{\Psi}_n | \tilde{H} | \tilde{\Psi}_m \rangle (f_n + f_m)/2$. Note that for a system at rest the Hamiltonian commutes with the occupations.

B. Propagation of the wave functions

The equations of motion for the electrons are integrated using the Verlet algorithm³⁷

$$|\tilde{\Psi}_n(+)\rangle = 2|\tilde{\Psi}_n(0)\rangle - |\tilde{\Psi}_n(-)\rangle - \tilde{H}|\tilde{\Psi}(0)\rangle + \sum_m \tilde{O}|\tilde{\Psi}_m(0)\rangle \Lambda_{mn} \frac{1}{f_n}. \quad (63)$$

The following notation is used here for the time steps. The time step is denoted by Δ . The wave functions have an integer argument denoting the number of the particular time step relative to the actual configuration. Hence the series of coordinates is $\dots, |\tilde{\Psi}_n(-2)\rangle, |\tilde{\Psi}_n(-)\rangle, |\tilde{\Psi}_n(0)\rangle, |\tilde{\Psi}_n(+)\rangle, |\tilde{\Psi}_n(+2)\rangle, \dots$. For convenience let us abbreviate the arguments for the previous and the subsequent time step. Similar notation is used for the atomic positions.

As a rule of thumb, the equations of motion for the wave functions are integrated properly if the time step Δ is related to the fictitious mass of the wave functions m_Ψ so that Δ^2/m_Ψ lies in the range 0.1–0.15 a.u.³⁸ For most systems the mass m_Ψ can be chosen to be 1000 a.u., resulting in a time step of about 10 a.u. or 0.25 fsec.

During the dynamical simulation, we have to ensure the orthogonality of the wave functions in an energy-conserving manner consistent with the accuracy of the Verlet algorithm.³⁷ The methods were originally invented for molecular-dynamics simulations of polymers³⁹ that obey bond-length constraints. Car and Parrinello^{11,40} adopted this algorithm in their formulation of first-principles molecular dynamics. Later Laasonen *et al.*⁴¹ extended this method to the case of a position-dependent overlap matrix, as used for Vanderbilt's ultrasoft pseudopotentials. I have adopted their strategy and extended it to include the possibility of different occupations for different states. This will be described in the following.

The wave functions are first propagated without constraints, which yields $|\tilde{\Psi}\rangle$

$$|\tilde{\Psi}_n\rangle = 2|\tilde{\Psi}_n(0)\rangle - |\tilde{\Psi}_n(-)\rangle - \tilde{H}|\tilde{\Psi}_n(0)\rangle \frac{\Delta^2}{m_\Psi}. \quad (64)$$

The forces of constraint are added in a second step

$$|\tilde{\Psi}_n(+)\rangle = |\tilde{\Psi}_n\rangle - \sum_m \tilde{O}(0)|\tilde{\Psi}_m(0)\rangle \Lambda_{mn} \frac{\Delta^2}{m_\Psi f_n}, \quad (65)$$

and the Lagrange parameters Λ_{nm} are determined iteratively so that the constraints for the next time step

$$\langle \tilde{\Psi}_n(+) | \tilde{O}(+) | \tilde{\Psi}_m(+) \rangle = \delta_{nm} \quad (66)$$

are exactly fulfilled within a given tolerance for the overlap matrix. $\tilde{O}(+)$ is the PS overlap operator for the next

time step. It differs from $\tilde{O}(0)$ only if the atoms are moving. The trajectories determined with this procedure are exactly symmetric under time reversal, which is crucial to obtaining energy conservation and predicts the wave functions accurately to the order Δ^3 , consistent with the overall accuracy of the Verlet algorithm. Furthermore, the constraints cannot deteriorate if a finite time step is used.

In practice we first evaluate the forces of constraints as

$$|\chi_n\rangle = \tilde{O}(0)|\tilde{\Psi}_n(0)\rangle. \quad (67)$$

Using Eq. (65) I rewrite the constraint equation Eq. (66) as

$$A^{(0)} + X^\dagger B + B^\dagger X + X^\dagger C X = 1 \quad (68)$$

with the definitions

$$A_{nm}^{(0)} = \langle \tilde{\Psi}_n | \tilde{O}(+) | \tilde{\Psi}_m \rangle, \quad (69)$$

$$B_{nm} = \langle \chi_n | \tilde{O}(+) | \tilde{\Psi}_m \rangle, \quad (70)$$

$$C_{nm} = \langle \chi_n | \tilde{O}(+) | \chi_m \rangle, \quad (71)$$

and

$$X_{mn} = \Lambda_{mn} \frac{\Delta^2}{m_\Psi f_n}. \quad (72)$$

Equation (68) cannot be solved directly for X . Therefore we obtain X iteratively. Let us first analyze Eq. (68) in orders of the time step Δ .

We will see that the leading order is proportional to Δ^2 : $A_{nm}^{(0)} = \delta_{nm} + O(\Delta^2)$, because the forces of constraint contribute only in the second order and therefore $|\Psi(+)\rangle = |\tilde{\Psi}\rangle + O(\Delta^2)$. As the leading order of X_{nm} is proportional to Δ^2 , the term $X^\dagger C X$ vanishes in leading order and only the zeroth order of B contributes. The zeroth order of B is equal to $\langle \tilde{\Psi}_n(0) | \tilde{O}(0) | \tilde{\Psi}_m(0) \rangle$ and therefore Hermitian. Hence the lowest order of Eq. (68) in Δ is

$$\sum_i B_{ni}^* X_{im}^{(i)} + X_{ni}^{(i)\dagger} B_{im}^* = -(A_{nm}^{(i)} - \delta_{nm}), \quad (73)$$

where $i = 0$ and $B^* = (B + B^\dagger)/2$ is the Hermitian part of B . Equation (73) determines $X = X^{(0)} + O(\Delta^3)$ accurately in leading order.

Analogously to the previous discussion we find that also the higher orders of $X = \sum_i X^{(i)}$ are obtained from Eq. (73), with $A^{(i)}$ given by

$$A^{(i)} = A^{(i-1)} + X^{(i-1)\dagger} B + B X^{(i-1)} + X^{(i-1)\dagger} C X^{(i-1)}. \quad (74)$$

However, there are many solutions to Eq. (73). They differ by a matrix $\delta X = (B^*)^{-1} D$, where D is an arbitrary antisymmetric matrix. Only a solution that fulfills

$$X_{ij} f_j = f_i X_{ji} \quad (75)$$

will conserve the energy and is of interest.

To obtain X , we diagonalize $B_{nm} = \sum_l U_{nl}^\dagger b_l U_{lm}$ and obtain its eigenvalues b_i and the unitary matrix U_{nm} formed by its eigenvectors. From that we determine

$$X_{nm}^{(i+1)} = \sum_{p,j,k,l} \frac{U_{np}^\dagger U_{pj} (A_{jk}^{(i)} - \delta_{jk}) U_{kl}^\dagger U_{lm}}{b_p + b_l} \frac{2f_n}{f_n + f_m}. \quad (76)$$

The iterative procedure Eq. (76) for $X = \sum_i X^i$ has a fixed point at the correct solution for Eqs. (75) and (68). In each iteration Eq. (75) is exactly fulfilled, which ensures energy conservation in each step.

This iterative scheme for X is a Taylor expansion in Δ if either all occupations are identical or if B^s is unity times a constant factor. As these requirements often are not fulfilled, each order in Δ requires an additional iteration, which is similar to that described above, but without the term quadratic in $X^{(i)}$ and the non-Hermitian part of B in Eq. (74). Owing to the close relationship between the two nested iteration schemes, in practice I perform only the outer iteration, which now also plays the role of the inner iteration.

Evaluation of the overlap matrix $A^{(i)}$ does not require that the scalar products of the wave functions be reevaluated. Instead, the matrix $A^{(i)}$ is calculated from Eq. (68) using the matrices $X^{(i)}$ from the previous iteration steps. Note also that B^s needs to be diagonalized only once for every time step, i.e., once for the entire iterative scheme described in this subsection. Convergence is reached if every element of the right-hand side in Eq. (73) is smaller than a certain given tolerance. Finally we can predict the new PS wave functions according to Eq. (65).

C. Propagating the atoms

The equations of motion for the ions are integrated as

$$R_i(+) = 2R_i(0) - R_i(-) + F_{R_i} \frac{\Delta^2}{M_{R_i}}. \quad (77)$$

1. Forces consistent with the Lagrange multipliers

As mentioned before, the force component due to the changing overlap operator described in Eq. (58) must be modified in the molecular-dynamics formalism. Instead of $\langle \tilde{\Psi}_n | \hat{H} | \tilde{\Psi}_m \rangle (f_n + f_m)/2$, one has to use the Lagrange multipliers. This results from the condition of energy conservation. The change of the total energy is

$$\begin{aligned} \dot{E} = M\dot{R}\ddot{R} + \dot{R} \frac{dE}{dR} + \sum_n f_n \left[m_\Psi (\langle \ddot{\tilde{\Psi}}_n | \dot{\tilde{\Psi}}_n \rangle + \langle \dot{\tilde{\Psi}}_n | \ddot{\tilde{\Psi}}_n \rangle) \right. \\ \left. + \langle \dot{\tilde{\Psi}}_n | \frac{dE}{d\langle \tilde{\Psi}_n |} + \frac{dE}{d|\tilde{\Psi}_n \rangle} | \dot{\tilde{\Psi}}_n \rangle \right]. \end{aligned} \quad (78)$$

We insert the equations of motion and resolve the energy derivatives using Eqs. (42) and (47) to obtain

$$\begin{aligned} \dot{E} = \dot{R}F_R + \dot{R} \sum_n f_n \langle \tilde{\Psi}_n | \nabla_R \tilde{H} | \tilde{\Psi}_n \rangle \\ + \sum_{n,m} [\Lambda_{nm}^\dagger \langle \tilde{\Psi}_m | \tilde{O} | \dot{\tilde{\Psi}}_n \rangle + \langle \dot{\tilde{\Psi}}_m | \tilde{O} | \tilde{\Psi}_n \rangle \Lambda_{nm}]. \end{aligned} \quad (79)$$

Now we use the first energy derivative of the constraint equation $\langle \tilde{\Psi}_n | \tilde{O} | \tilde{\Psi}_m \rangle$ and the requirement that Λ be Hermitian to obtain

$$\begin{aligned} \dot{E} = \dot{R} \left[F_R + \sum_n f_n \langle \tilde{\Psi}_n | \nabla_R \tilde{H} | \tilde{\Psi}_n \rangle \right. \\ \left. - \sum_{n,m} \langle \tilde{\Psi}_n | \nabla_R \tilde{O} | \tilde{\Psi}_m \rangle \Lambda_{mn} \right]. \end{aligned} \quad (80)$$

The requirement $\dot{E} = 0$ results in an expression for the force

$$F_R = - \sum_n f_n \langle \tilde{\Psi}_n | \nabla_R \tilde{H} | \tilde{\Psi}_n \rangle + \sum_{n,m} \langle \tilde{\Psi}_n | \nabla_R \tilde{O} | \tilde{\Psi}_m \rangle \Lambda_{mn}, \quad (81)$$

which in the stationary case is identical to Eq. (48) above.

The propagation of the ionic positions is straightforward once the Lagrange multipliers are known. As seen in Sec. VB, those will be calculated only after the new positions are determined because the PS overlap operator and hence the orthogonalization that yields the Lagrange multipliers depend explicitly on the atomic positions. However, the Verlet algorithm has only limited accuracy in the time step Δ . Hence it is sufficient if we can predict Λ_{nm} up to the order Δ^1 . This implies that a linear extrapolation from the last two time steps

$$\Lambda_{nm}(0) = 2\Lambda_{nm}(-) - \Lambda_{nm}(-2) + O(\Delta^2) \quad (82)$$

is sufficient. It is therefore not necessary to achieve self-consistency of atomic positions and Lagrange multipliers iteratively, which would be computationally extremely demanding. This is in contrast to the approach used by Laasonen *et al.*,⁴¹ who suggest that Λ and forces be determined self-consistently.

2. Renormalization of the atomic masses

The fictitious dynamics of the electronic wave functions has two main effects on the atomic trajectories. First, energy is transferred constantly to the wave functions, which have a temperature that is very low compared to that of the ionic subsystem. The rate of heat transfer is roughly proportional to the magnitude of the forces acting on the ions and to the band gap between occupied and unoccupied states.³⁸ This effect can be controlled for long simulations using two Nosé thermostats,^{42,43} one to keep the ions at their physical temperature and one to keep the wave functions close to the Born-Oppenheimer surface.

The second effect is that the ions propagate as quasi-particles dressed by the wave functions,⁴⁴ which increases the effective mass. This effect can be compensated by renormalizing the masses of the nuclei. The magnitude of the effect can be estimated from an isolated atom that experiences external forces, described by a

potential $V_{\text{ext}}(R)$ acting on the nuclei. Assuming that the wave functions of the atom reside exactly on the Born-Oppenheimer surface, the wave functions do not change, except that they undergo a rigid displacement and $|\tilde{\Psi}\rangle_{\text{at}} = -|\nabla\tilde{\Psi}\rangle_{\text{at}}\dot{R}$. Hence the Lagrangian can be simplified to

$$\mathcal{L}' = \frac{1}{2} \sum_{i,j} \dot{R}_i \left(2m_{\Psi} \sum_n f_{\text{at},n} \langle \nabla_i \tilde{\Psi}_n | \nabla_j \tilde{\Psi}_n \rangle_{\text{at}} + M_R \delta_{ij} \right) \dot{R}_j - E_0[\Psi] - V_{\text{ext}}. \quad (83)$$

Here E_0 is the total energy of the isolated atom, which is constant during the simulation. The constraints of orthonormal wave functions are automatically fulfilled because here $\tilde{\Psi}$ denotes rigidly displaced PS wave functions of the isolated atom. The effective mass tensor $2m_{\Psi} \sum_n f_n \langle \nabla_i \tilde{\Psi}_n | \nabla_j \tilde{\Psi}_n \rangle + M_R \delta_{ij}$ is diagonal because the isolated atom is spherically symmetric. The effective mass

$$\tilde{M}_R = \frac{4}{3} m_{\Psi} \sum_n f_{\text{at},n} \langle \tilde{\Psi} | -\frac{1}{2} \nabla^2 | \tilde{\Psi} \rangle_{\text{at}} + M_R \quad (84)$$

is therefore one-third of the trace of the mass tensor, which has been modified after applying Gauss's theorem. The expectation value $\sum_n f_n \langle \tilde{\Psi} | -\frac{1}{2} \nabla^2 | \tilde{\Psi} \rangle_{\text{at}}$ is nothing other than the plane-wave part of the true electronic kinetic energy. Hence the bare mass of the ions used in the Lagrangian should be reduced by

$$\delta M = \frac{4}{3} m_{\Psi} \sum_n f_{\text{at},n} \langle \tilde{\Psi}_n | -\frac{1}{2} \nabla^2 | \tilde{\Psi}_n \rangle_{\text{at}}. \quad (85)$$

This correction has been included in all our simulations described here. The quality of this correction can be estimated by comparing the kinetic energy related to the PS wave functions of the system of interest to that of the isolated atoms.

VI. CONSTRUCTION OF PARTIAL WAVES AND PROJECTORS

The basic ingredients of the PAW method are partial waves and projectors. There is an infinite number of ways to construct them. I will describe here in detail the particular choice I made for this application. Even though the solution of the problem is quite satisfactory, there may be better choices than the ones described here. In particular the construction of PS partial waves is completely analogous to the construction of pseudopotentials with the pseudopotential method. The expertise acquired with the pseudopotential method⁴⁵⁻⁴⁸ is likely to create choices that permit a further reduction of the number of plane waves. The partial waves and projector functions obtained with the procedure described below are shown in Fig. 1.

A. All-electron partial waves

The AE partial waves are obtained by radially integrating the Schrödinger equation

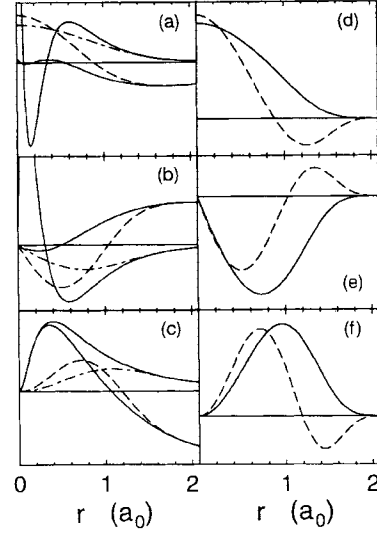


FIG. 1. Partial waves and projectors for Mn. Left panel: AE partial waves (solid lines) and PS partial waves (dashed and dash-dotted lines). The “first” PS partial wave is a dash-dotted line. Right panel: first (solid line) and second (dashed line) projector functions. (a) and (d) show the results for the first and the second partial wave of the s angular momentum channel, respectively, (b) and (e) for the p channel, and (c) and (f) for the d channel. $3s$ and $3p$ functions are treated as valence states. Functions are scaled individually.

$$\left(-\frac{1}{2} \nabla^2 + v_{\text{at}} - \epsilon_i^1 \right) |\phi_i\rangle = 0 \quad (86)$$

outward for the self-consistent atomic AE potential v_{at} and a set of energies ϵ_i^1 . In practice we use the scalar relativistic version of Koelling and Harmon.⁴⁹

The AE partial waves are chosen to describe the physically relevant states, i.e., those from the valence band region, reasonably well. The energy ϵ^1 of the first partial wave per angular momentum and site is usually chosen as the energy of the lowest bound valence state of the atom. The energy of the second partial wave is chosen after inspecting the scattering properties of the PAW method for the isolated atom with only one partial wave. It is placed at the energy where the scattering properties begin to deteriorate. There is no need for the partial wave to be bound states because the exponentially increasing tail will be canceled exactly by the identical behavior of the PS partial waves. The number of partial waves is then further increased until the scattering properties of the valence band region are described satisfactorily.

An equally justified approach, more similar in spirit to the linear methods, is to use increasingly higher energy derivatives of the energy-dependent partial waves obtained at one fixed energy. In principle, one can also add partial waves from atoms with various occupations.

If core states extend beyond the augmentation region, we subsequently orthogonalize the AE partial waves *within the augmentation region* to core states centered on the same site. We find that one AE partial wave per angular momentum and site is often sufficient and that

two terms yield a satisfactory description even for difficult cases such as transition metals and systems in which semicore states are treated as valence states. However, if one is interested in states that lie far above the valence band region, the number of AE partial waves can be increased further until the desired accuracy is achieved.

The entire construction is fairly insensitive to the choice of the energies of the AE partial waves. Since the valence band region can be described fairly well with two partial waves, as shown in the linear methods, any two partial waves from this region will span a very similar portion of the Hilbert space. Even though the individual partial waves and projectors will differ, they represent an almost equivalent choice.

B. Pseudopartial waves

To construct PS partial waves, I proceed in loose analogy to the pseudopotential approach described in the work of Hamann *et al.*,^{8,50,51} which I have extended to include several terms into the separable form.²¹ However, in general we do not perform the norm-conservation step suggested there.

I first select a PS potential $\tilde{v}_{\text{at}}$. This is done in two different ways, depending on the element.

(i) For transition metals, a polynomial of fourth order is matched differentially to the AE potential. Outside the matching radius the two potentials coincide. The remaining free parameter, the value of $\tilde{v}_{\text{at}}$ at the nuclear site, is adjusted by hand.

(ii) For elements without d electrons in the valence shell, $\tilde{v}_{\text{at}}$ is obtained from the AE potential as $\tilde{v}_{\text{at}}(r) = \tilde{v}_{\text{at}}(0)k(r) + [1 - k(r)]v_{\text{at}}(r)$, using a cutoff function of the form

$$k(r) = \exp[-(r/r_k)^\lambda]. \quad (87)$$

In order to obtain the PS partial wave, I define for each AE partial wave a PS potential of the form

$$w_i(r) = \tilde{v}_{\text{at}}(r) + c_i k(r). \quad (88)$$

The values of the cutoff radius r_k and the exponent λ are chosen such that this potential is virtually identical to the AE atomic potential outside the augmentation region. Often we choose $\lambda = 6$ and r_k as three-quarters of the covalent radius. The values used in our calculations are listed in Table I. The PS partial wave is obtained as a solution of the nonrelativistic Schrödinger equation

$$\left(-\frac{1}{2}\nabla^2 + w_i(r) - \epsilon_i^1\right)|\tilde{\phi}_i\rangle = 0 \quad (89)$$

for the energy of the corresponding AE partial wave and the potential $w_i(r)$. The free coefficient c_i in Eq. (88) is then determined such that the PS partial wave coincides with the corresponding AE partial wave outside the augmentation region.

C. Projector functions

Next we calculate preliminary projector functions according to

TABLE I. Parameters for the construction of PS partial waves. The cutoff parameter has been chosen as $\lambda = 6$ for all atoms and angular momentum channels.

Symbol	$\tilde{v}(0)[H]$	$\epsilon_s[H]$	r_{ks}	ϵ_p	r_{kp}	ϵ_d	r_{kd}
H	-3.43	-0.234	0.45				
Li	-1.43	-0.106	1.20				
Be	-1.40	-0.207	1.20	-0.079	1.20		
B	-2.20	-0.346	1.00	-0.137	1.00		
C	-2.47	-0.501	1.00	-0.199	1.00		
N	-2.58	-0.677	1.00	-0.266	1.00		
O _a	-3.19	-0.873	0.90	-0.338	0.90		
O _b	-2.60	-0.873	1.00	-0.338	1.00		
F	-2.55	-1.090	1.00	-0.415	1.00		
Fe	1.88	-0.020	1.50	-0.058	1.50	-0.287	1.50
		0.000	1.50			0.000	1.50
Mn _a	0.0	-0.194	1.50	-0.054	1.50	-0.257	1.50
			1.50		1.50	0.500	1.50
Mn _b	-3.20	-3.138	1.50	-2.002	1.50	-0.257	1.50
		-0.194	1.50	-0.054	1.50	0.500	1.50

$$|\tilde{p}_i\rangle = (-\frac{1}{2}\nabla^2 + \tilde{v}_{\text{at}} - \epsilon_i^1)|\tilde{\phi}_i\rangle. \quad (90)$$

If the result is zero, we set the projector function equal to the cutoff function $k(r)$. These projector functions must be modified such that they fulfill the condition $\langle\tilde{p}_i|\tilde{\phi}_j\rangle = \delta_{ij}$. This condition is now imposed iteratively beginning with the lowest partial wave. The following equations [(91)–(96)] should not be read as mathematical identities but rather like a computer program: The left-hand side is the product of the right-hand side, with the new value overwriting previous values for the same symbol. I have done this to avoid multiple symbols for the same quantity.

For a given partial wave denoted by subscript i , and assuming that $\langle\tilde{p}_k|\tilde{\phi}_j\rangle = \delta_{kj}$ for $k, j < i$, the projector functions are first orthogonalized to all lower PS partial waves by

$$|\tilde{p}_i\rangle = |\tilde{p}_i\rangle - \sum_{j=1}^{i-1} |\tilde{p}_j\rangle \langle\tilde{\phi}_j|\tilde{p}_i\rangle. \quad (91)$$

Then the AE and PS partial waves are modified in order to ensure the orthogonality of the PS partial wave with the lower projector functions

$$|\phi_i\rangle = |\phi_i\rangle - \sum_{j=1}^{i-1} |\phi_j\rangle \langle\tilde{p}_j|\tilde{\phi}_i\rangle, \quad (92)$$

$$|\tilde{\phi}_i\rangle = |\tilde{\phi}_i\rangle - \sum_{j=1}^{i-1} |\tilde{\phi}_j\rangle \langle\tilde{p}_j|\tilde{\phi}_i\rangle. \quad (93)$$

Finally the projector function and partial waves are scaled so that $\langle\tilde{p}_i|\tilde{\phi}_i\rangle$ is unity

$$|\tilde{p}_i\rangle = |\tilde{p}_i\rangle / \langle\tilde{p}_i|\tilde{\phi}_i\rangle \times c, \quad (94)$$

$$|\tilde{\phi}_i\rangle = |\tilde{\phi}_i\rangle / c, \quad (95)$$

$$|\phi_i\rangle = |\phi_i\rangle / c. \quad (96)$$

The free constant c is used to avoid very large projector functions, while the partial waves are very small and vice versa. It has no influence other than to prevent very small and very large numbers, which may create problems on the computer. Once the set of projectors and partial waves with index i are modified to obey the orthogonality condition, one proceeds to the next set of projectors and partial waves $|\tilde{p}_{i+1}\rangle, |\phi_{i+1}\rangle, |\tilde{\phi}_{i+1}\rangle$.

D. The potential $\bar{v}$

The potential $\bar{v}$ is now obtained by subtracting the potential of the self-consistent atomic PS density from the PS potential used to define the PS waves:

$$\bar{v}(r) = \bar{v}_{\text{at}}(r) - \int d\mathbf{r}' \frac{\tilde{n}(\mathbf{r}') + \hat{n}(\mathbf{r}')}{|\mathbf{r} - \mathbf{r}'|} - \mu_{\text{xc}}[\tilde{n}(r)]. \quad (97)$$

This step is the analog to the unscreening of a pseudopotential performed in the pseudopotential approach.

Since the PS partial waves do not necessarily correspond to the atomic bound states, which are needed in Eq. (97), the latter are obtained from the radial, separable Schrödinger equation

$$\left(-\frac{1}{2}\nabla^2 + \bar{v}_{\text{at}} - \epsilon + \sum_{i,j} |\tilde{p}_i\rangle (dH_{ij} - \epsilon dO_{ij}) \langle \tilde{p}_j| \right) |\tilde{\Psi}_j\rangle = 0 \quad (98)$$

with dH_{ij} and dO_{ij} given as

$$dH_{ij} = \langle \phi_i | -\frac{1}{2}\nabla^2 + v_{\text{at}} | \phi_j \rangle - \langle \tilde{\phi}_i | -\frac{1}{2}\nabla^2 + \bar{v}_{\text{at}} | \tilde{\phi}_j \rangle, \quad (99)$$

$$dO_{ij} = \langle \phi_i | \phi_j \rangle - \langle \tilde{\phi}_i | \tilde{\phi}_j \rangle. \quad (100)$$

A way to solve this equation is sketched in the Appendix.

To obtain the PS density we still need to define its core contribution. The PS core density $\tilde{n}^c$ is obtained by matching a parabola differentially to the AE core density. Outside the matching radius PS and AE core densities are identical. Using the wave functions and occupations of the atom one constructs a PS charge density $\tilde{n}$ and from that $\bar{v}$ according to Eq. (97).

E. Outlook

The procedure described above is by no means the only way to create PS partial waves. There are a number of different ways to construct first-principles pseudopotentials.^{45,46,48} These methods can easily be adapted to relax the norm-conservation condition, to allow a larger augmentation region, and to include unbound states. Each of them can be used to construct PS partial waves corresponding to given AE partial waves and, with the procedure outlined above, to construct projector functions.

Once the PS partial waves are defined, it is recommended that the procedure described above be followed

in detail. In particular, this approach ensures that the states used to construct the PS partial waves are reproduced correctly as atomic PS wave functions of the PAW method, irrespective of the quality of the partial-wave expansion.

Let me summarize which quantities we import from the atomic calculation into the *ab initio* molecular dynamics simulation: (i) the AE and PS core densities, (ii) AE and PS partial waves $|\phi_i\rangle$ and $|\tilde{\phi}_i\rangle$ and PS projector functions $\langle \tilde{p}_i|$, (iii) the matrices $\langle \phi_i | -\frac{1}{2}\nabla^2 | \phi_j \rangle - \langle \tilde{\phi}_i | -\frac{1}{2}\nabla^2 | \tilde{\phi}_j \rangle$ and $\langle \phi_i | \phi_j \rangle - \langle \tilde{\phi}_i | \tilde{\phi}_j \rangle$ for the calculation of the one-center contributions of kinetic energy and overlap matrix (note that the Laplacian for the AE partial waves is replaced by its scalar relativistic counterpart), and (iv) the cutoff r_c that determines the range of the short-ranged compensation densities.

VII. ANALYSIS OF THE PARTIAL-WAVE TRUNCATION ERROR AND EXTENSIONS OF THE PAW METHOD

In the previous sections we have taken the point of view that the PAW method is an exact formulation of the Kohn-Sham equations, from which a practical scheme is obtained by truncating two rapidly converging series expansions. Here I will analyze the truncation errors of the partial-wave expansion in detail and thus justify the choices I made for wave functions and total energy expressions. This section can be skipped by the practitioner. I recommend this section to those who are interested in the underlying principles and possible extensions of the present implementation of the PAW method.

A. Truncation error in the wave function

1. Orthogonality to the core states

The main effect of the truncation of the partial-wave expansions for the wave function is to redefine the transformation from the valence wave functions to the PS wave functions. This in itself does not introduce errors, but it affects the orthogonality to the core states. Whereas the AE partial waves are constructed to be orthogonal to the core states, a nonzero remainder of $|\tilde{\Psi}\rangle - \sum_i |\tilde{\phi}_i\rangle \langle \tilde{p}_i | \tilde{\Psi}\rangle$ can create a nonzero overlap with the core states.

Therefore, I introduce a new definition of the transformation $\mathcal{T}$ that explicitly ensures that any PS wave function is transformed onto an AE wave function that is exactly orthogonal to the core states $|\phi^c\rangle$:

$$|\Psi\rangle = |\tilde{\Psi}\rangle + \sum_i (|\phi_i\rangle - |\tilde{\phi}_i\rangle) \langle \tilde{p}_i | \tilde{\Psi}\rangle - \sum_i |\phi^c\rangle \langle \phi^c | \left(1 - \sum_j |\tilde{\phi}_j\rangle \langle \tilde{p}_j|\right) |\tilde{\Psi}\rangle. \quad (101)$$

In the analysis of truncation errors it should be kept in mind that Eq. (101) rather than Eq. (9) is the true definition of the transformation between PS and AE wave

functions. Of course, if the partial waves form a complete set, the two expressions are identical.

2. Additive augmentation

When truncating the partial-wave expansions it is important that the partial-wave expansions of the AE and the PS wave functions are truncated in a completely analogous way. This principle is called additive augmentation and has important advantages.

First, the wave functions of the PAW method are differentiable to an arbitrary order if the PS partial waves have been constructed to be differentiable to an infinite order. (In many implementations of the LAPW method the wave function is even discontinuous.)

Second, higher partial waves not explicitly included in the partial-wave expansions are represented by the tails of the plane-wave part that extend into the augmentation region.

Finally, the PAW basis set is complete whenever the plane waves form a complete basis set, irrespective of the partial-wave truncation. This justifies the use of partial waves imported from the isolated atom without adjusting them to the actual potential, as done in the linear methods. The use of frozen partial waves has substantial advantages in combination with the first-principles molecular-dynamics approach because it eliminates a large number of parameters that otherwise have to be treated as dynamical variables or determined variationally in each time step to a very high degree of accuracy.

The principle of additive augmentation itself is not new and has been exploited to some extent in the LMTO method and in the APW method of Soler and Williams.^{15–17} There the angular momentum expansions of the wave function and charge density were truncated in the same way, resulting in a very rapid ℓ convergence. As a result of the projector augmentation, however, we can exploit this principle even on the level of individual partial waves.

Here I will show that truncation of the partial-wave expansions does not affect the completeness of the basis set: If a set of PS wave functions forms a complete basis, the corresponding basis of projector augmented wave functions is complete in the Hilbert space orthogonal to the core states. For this to be true two weak conditions must be fulfilled: There is a matrix a_{ij} such that $\sum_k a_{ik} \langle \tilde{p}_k | \phi_j \rangle = \delta_{ij}$, which has no zero right-hand eigenvalues,⁵² and the differences between AE and PS partial waves $|\phi_i\rangle - |\tilde{\phi}_i\rangle$ are not linearly dependent.

To prove this statement, it must be shown that for every AE wave function orthogonal to the core states there exists one and only one well-defined PS wave function. For linear transformations such as the ones considered here, this implies first that we can define an inverse transformation $\mathcal{T}^{-1}$ from the AE wave function to the PS wave function. This is indeed possible and the expression is formally very similar to the forward transformation:

$$|\tilde{\Psi}\rangle = |\Psi\rangle + \sum_i (|\tilde{\phi}_i\rangle - |\phi_i\rangle) \langle p_i | \Psi \rangle, \quad (102)$$

with the “AE projector functions” defined as

$$\langle p_i | = \sum_j (\langle \tilde{p} | \phi \rangle^{-1})_{ji} \langle \tilde{p}_j |. \quad (103)$$

Note the difference between the AE and the PS projector functions.

To show that the transformation is unique, we test whether any nonzero function orthogonal to the core states is mapped onto a zero PS wave function. This would be the case only if

$$|\Psi\rangle + \sum_i (|\tilde{\phi}_i\rangle - |\phi_i\rangle) \langle p_i | \Psi \rangle = 0 \quad (104)$$

for any function $|\Psi\rangle$ orthogonal to the core states. Hence such a function $|\Psi\rangle$ must be a superposition $\sum_i (|\phi_i\rangle - |\tilde{\phi}_i\rangle) c_i$. If we insert this ansatz into Eq. (104) and exploit $\langle p_i | \phi_j \rangle = \delta_{ij}$, it is clear that the coefficients must fulfill

$$\sum_{i,j} (|\phi_i\rangle - |\tilde{\phi}_i\rangle) \langle p_i | \tilde{\phi}_j \rangle c_j = 0. \quad (105)$$

The matrix $\langle p_i | \tilde{\phi}_j \rangle$ is none other than the matrix a_{ij} defined above, as it fulfills the relation $\sum_k \langle p_i | \tilde{\phi}_k \rangle \langle \tilde{p}_k | \phi_j \rangle = \delta_{ij}$. Equation (105) can only be fulfilled if either a_{ij} has a zero right-hand eigenvalue or the functions on the left-hand side can add up to zero; these are the exceptions given above. This concludes the proof of the completeness of the PAW basis set.

B. Truncation error in the expectation values

While evaluating expectation values in Sec. VII A, it was reasonable to neglect cross terms between the three contributions of the wave function because $|\tilde{\Psi}\rangle - \sum_i |\tilde{\phi}_i\rangle \langle \tilde{p}_i | \tilde{\Psi}\rangle = 0$ in the augmentation regions, if the partial waves form a complete set of functions in the augmentation region. However, if the partial-wave expansions are truncated, this condition is no longer exactly fulfilled.

In the following I will use the symbols $|\Psi^1\rangle = \sum_i |\phi_i\rangle \langle \tilde{p}_i | \Psi\rangle$ for the one-center expansion of the AE wave function and $|\tilde{\Psi}^1\rangle = \sum_i |\tilde{\phi}_i\rangle \langle \tilde{p}_i | \tilde{\Psi}\rangle$ for the one-center expansion of the PS wave function. The difference $\Delta A_{nm} = \Delta A_{nm}^{(1)} + \Delta A_{nm}^{(2)}$ between the matrix elements of an operator A calculated directly from Eq. (101) and $\langle \tilde{\Psi}_n | \tilde{A} | \tilde{\Psi}_m \rangle$, with the PS operator from Eq. (11), is given by

$$\begin{aligned} \Delta A_{nm}^{(1)} = & [\langle \Psi_n^1 | A (1 - P^c) - \langle \tilde{\Psi}_n^1 | A | \tilde{\Psi}_m - \tilde{\Psi}_m^1 \rangle \\ & + \langle \tilde{\Psi}_n - \tilde{\Psi}_n^1 | [(1 - P^c) A | \Psi_m^1 \rangle - A | \tilde{\Psi}_m^1 \rangle], \end{aligned} \quad (106)$$

$$\Delta A_{nm}^{(2)} = \langle \tilde{\Psi}_n - \tilde{\Psi}_n^1 | (1 - P^c) A (1 - P^c) - A | \tilde{\Psi}_m - \tilde{\Psi}_m^1 \rangle, \quad (107)$$

where $P^c = \sum_i |\phi_i^c\rangle \langle \phi_i^c|$ is the projection operator on the core states.

The first term $\Delta A_{nm}^{(1)}$ is proportional to the difference $|\tilde{\Psi}\rangle - |\tilde{\Psi}^1\rangle$ between a PS wave function and its one-center expansion, whereas $\Delta A_{nm}^{(2)}$ depends quadratically on it. Consequently both terms converge to zero as the partial-wave expansion is made complete.

The first term is a matrix element between the difference between the PS wave function and its one-center expansion, which is largest at the surface of the augmentation region, and a function that is localized in the center of the atom, namely, the difference between the one-center expansion of the AE wave function and the one for the PS wave functions. For quasilocal operators such as the kinetic energy or the real-space projection operator needed for the total energy, one profits from the fact that the two functions are largest in opposite regions of space, resulting in small errors.

The fact that the dominant truncation error $\Delta A_{nm}^{(1)}$ is proportional to the difference between AE and PS partial waves is the reason for truncating both partial-wave expansions in exactly the same way. Partial waves for higher energies become increasingly insensitive to

the shape of the potential: High-energy electrons pass through the atom too fast to be seriously affected by the potential. Hence the difference between the AE and PS partial waves vanishes for high energies, even though each partial wave itself is still sizable. Since they appear in pairs of opposite sign in the truncation error, the consistent truncation of both expansions is highly favorable. I shall return to this point when comparing the PAW with the LAPW method.

C. Truncation error in the total energy

The error in the total energy can be divided into two parts. The first is due to the difference between expressions (19)–(21) and the total energy calculated directly using the expectation values for charge density, kinetic energy, and overlap obtained via Eq. (11). The second source of the error is due to the approximations described in Sec. VII B.

The first error term is of the form

$$\begin{aligned} \Delta E^{(1)} = & \sum_R \int_{\Omega_R} dr [\tilde{n}(r) - \tilde{n}_R^1(r)] \int dr' \left(\frac{n_R^1(r') + n_R^Z(r')}{|r - r'|} - \frac{\tilde{n}_R^1(r') + \tilde{n}_R(r')}{|r - r'|} - \bar{v}(r) \right) \\ & + \int_{\Omega_R} dr \{ [\tilde{n}(r) + n^1(r) - \tilde{n}^1(r)] \epsilon_{xc}(\tilde{n} + n^1 - \tilde{n}^1) - \tilde{n}(r) \epsilon_{xc}(\tilde{n}) - n^1(r) \epsilon_{xc}(n^1) - \tilde{n}^1(r) \epsilon_{xc}(\tilde{n}^1) \}. \end{aligned} \quad (108)$$

It is easily seen that both terms vanish as $\tilde{n} - \tilde{n}^1$, i.e., if the partial-wave expansions are complete. Furthermore the integrands go smoothly to zero at the boundary of the augmentation region.

To get a better idea of these terms, we can expand them in orders of $\tilde{n} - \tilde{n}^1$ and $|\tilde{\Psi}\rangle - |\tilde{\Psi}^1\rangle$ and consider only terms up to the first order:

$$\begin{aligned} \Delta E^{(1)} = & \sum_R \int dr [\tilde{n}(r) - \tilde{n}_R^1(r)] [v_R^1(r) - \bar{v}_R^1(r)] \\ & + O(\tilde{n} - \tilde{n}^1)^2 \\ = & 2\text{Re} \sum_{Rn} f_n \langle \tilde{\Psi}_n - \tilde{\Psi}_n^1 | (v_R^1 - \bar{v}_R^1) | \tilde{\Psi}_n^1 \rangle \\ & + O(|\tilde{\Psi}\rangle - |\tilde{\Psi}^1\rangle)^2. \end{aligned} \quad (109)$$

Before returning to $\Delta E^{(1)}$, I consider the errors that propagate from the approximation of the expectation values. The error can be obtained via Eqs. (106) and (107),

where the operator A is $-\frac{1}{2}\nabla^2 + v - \epsilon_n$, where v is the exact potential and ϵ_n the exact energy eigenvalues. The error is the sum of the diagonal matrix elements with the eigenfunctions. Let us consider again only the lowest order in $|\tilde{\Psi}\rangle - |\tilde{\Psi}^1\rangle$:

$$\begin{aligned} \Delta E^{(2)} = & 2\text{Re} \sum_{R,n,i} f_n \langle \tilde{\Psi}_n - \tilde{\Psi}_n^1 | \\ & \times \{ (1 - P^c) [(v - v_{\text{at}}) - (\epsilon_n - \epsilon_i^1)] | \phi_i \rangle \\ & - [(v - w_i) - (\epsilon_n - \epsilon_i^1)] | \tilde{\phi}_i \rangle \} \langle \tilde{p}_i | \tilde{\Psi}_n \rangle \\ & + O(|\tilde{\Psi}\rangle - |\tilde{\Psi}^1\rangle)^2, \end{aligned} \quad (110)$$

using Eqs. (86) and (89), which define the partial waves, and (106). (Note that we need to apply a linear transformation to partial waves and projectors to undo the scrambling of partial waves described in Sec. VI C.)

Combining the two sources of error, $\Delta E^{(1)}$ from Eq. (109) and $\Delta E^{(2)}$ from Eq. (110), we find

$$\begin{aligned} \Delta E^{(1)} + \Delta E^{(2)} = & 2\text{Re} \sum_{R,n,i} f_n \langle \tilde{\Psi}_n - \tilde{\Psi}_n^1 | \{ (1 - P^c) [(v - v_{\text{at}}) - (\epsilon_n - \epsilon_i^1)] | \phi_i \rangle \\ & - [(v - v^1) - (w_i - \bar{v}^1) - (\epsilon_n - \epsilon_i^1)] | \tilde{\phi}_i \rangle \} \langle \tilde{p}_i | \tilde{\Psi}_n \rangle. \end{aligned} \quad (111)$$

Let us simplify Eq. (111) by replacing v with v^1 , which is justified since the difference between them also vanishes if the partial-wave expansions are complete and therefore does not contribute to the lowest order of the truncation error in $|\tilde{\Psi}\rangle - |\tilde{\Psi}^1\rangle$. Furthermore, we exploit the fact that $\langle \tilde{\Psi} - \tilde{\Psi}^1 | (w_i - \bar{v}_{\text{at}}) | \phi_i \rangle$ vanishes because $(w_i - \bar{v}_{\text{at}}) | \phi_i \rangle$ can be expressed as a superposition of projector functions [Eqs. (89) and (90)] and $\langle \tilde{\Psi} - \tilde{\Psi}^1 | \tilde{p}_i \rangle = \langle \tilde{\Psi} | (1 - \sum_j |\tilde{p}_j\rangle \langle \phi_j|) | \tilde{p}_i \rangle = 0$:

$$\Delta E^{(1)} + \Delta E^{(2)} = 2\text{Re} \sum_{R,n} f_n \langle \tilde{\Psi}_n - \tilde{\Psi}_n^1 | \times \{ (1 - P^c)[(v^1 - v_{\text{at}}) - (\epsilon_n - \epsilon_i^1)]|\phi_i\rangle - [(\tilde{v}^1 - \tilde{v}_{\text{at}}) - (\epsilon_n - \epsilon_i^1)]|\tilde{\phi}_i\rangle \} \langle \tilde{p}_i | \tilde{\Psi}_n \rangle. \quad (112)$$

Equation (112) is an important result. It tells us that the strongly varying potential of the core and the nucleus does *not* contribute to the error. This is a result of an efficient cancellation of two errors, namely, $\Delta E^{(1)}$ and $\Delta E^{(2)}$. In a hand-waving way, the term proportional to $(1 - P^c)(v^1 - v_{\text{at}})|\phi_i\rangle - (\tilde{v}^1 - \tilde{v}_{\text{at}})|\tilde{\phi}_i\rangle$ describes the charge density transferability error because it depends on the deviation of the potential from the atomic one, whereas the term proportional to $(\epsilon_n - \epsilon_i^1)|\tilde{\phi}_i - \phi_i\rangle$ describes the error in the scattering properties or the energy transferability. Note that the constant term in v and $\tilde{v}$ almost cancels a similar term in $\epsilon_n - \epsilon_i^1$. It is also worthwhile to note that the right-hand side of the product vanishes differentially at the boundary of the augmentation region, whereas the term $\langle \tilde{\Psi} | - \langle \tilde{\Psi}^1 |$ is a small quantity, which is expected to be largest far from the center of the atom. The resulting expression could actually be used to estimate the partial-wave truncation error in practice.

D. Extensions of the PAW method

The PAW method lends itself to a number of extensions which, though not yet implemented, may be interesting to keep in mind. These extensions concern the use of partial waves that adjust to the actual potential and the relaxation of the frozen-core approximation. I will show in Sec. VIII that the PAW method is highly accurate even without these features owing to the rigorous exploitation of the principle of additive augmentation. However, I want to demonstrate that the PAW method is sufficiently flexible to accommodate them, if desired. This will turn the linear transformation between AE and PS wave functions into a nonlinear one. As in the linear methods with adjusting partial waves, the nonlinear degrees of freedom can be relaxed during a self-consistent procedure. However, some caution is required if they are to be used in combination with a fictitious Lagrangian formalism because all nonlinearities must be treated consistently. In contrast to the linear methods in their present implementation, the partial waves will be adjusted here to both spherical and nonspherical parts of the on-center potentials.

1. Optimization of partial waves to the actual potential

Here I describe how the partial waves can be adjusted to the actual potential within the frozen-core approximation. In Sec. VIID 2, I will describe how to relax also the frozen-core approximation.

We start out with a large set of partial waves, one that is sufficiently complete to describe the wave functions accurately for all possible potentials that may occur in a

molecular or crystalline environment. This set of partial waves will be divided into a subset of “lower” partial waves that will be used as partial waves in the way outlined in the previous sections and a subset of “higher” partial waves. The lower partial waves are adjusted to the actual potential by mixing with the higher partial waves. This approach avoids the inclusion of additional projector functions and the corresponding increase of the computational effort by determining the coefficients of the higher partial waves self-consistently from the potential within the sphere. The procedure can be termed downfolded augmentation.

Let us denote the higher partial waves by the subscripts h, h' and the lower partial waves that adjust to the potential by the subscripts l, l' . The rigid partial waves that have been constructed from an atom and that will be used as a reference are distinguished from the adjusting partial waves by the superscript 0.

I make the following ansatz for the lower partial waves and projectors

$$\begin{aligned} |\phi_l\rangle &= |\phi_l^0\rangle + \sum_h |\phi_h^0\rangle a_{hl}, \\ |\tilde{\phi}_l\rangle &= |\tilde{\phi}_l^0\rangle + \sum_h |\tilde{\phi}_h^0\rangle a_{hl}, \\ |\tilde{p}_l\rangle &= \sum_{l'} |\tilde{p}_{l'}^0\rangle b_{l'l} + \sum_h |\tilde{p}_h^0\rangle c_{hl}. \end{aligned} \quad (113)$$

The coefficients will be determined such that (i) the orthogonality relation between projector functions and partial waves is fulfilled and (ii) the scattering into the higher partial waves vanishes:

$$\begin{aligned} \langle \tilde{p}_{l'} | \tilde{\phi}_l \rangle &= \delta_{l,l'}, \\ \langle \phi_h^0 | \theta_\Omega(-\tfrac{1}{2}\nabla^2 + v^1 - \epsilon_l) | \phi_l \rangle &= 0, \end{aligned} \quad (114)$$

$$\begin{aligned} \langle \tilde{\phi}_h^0 | \theta_\Omega\left(-\tfrac{1}{2}\nabla^2 + \tilde{v}^1 - \epsilon_l + \sum_{l',l''} |\tilde{p}_{l'}\rangle \right. \\ \left. \times (dH_{l',l''} - \epsilon_l dO_{l',l''}) \langle \tilde{p}_{l''} | \right) | \tilde{\phi}_l \rangle &= 0. \end{aligned}$$

Inserting the ansatz, we obtain expressions for the matrices a, b, c

$$a_{h,l} = - \sum_{h'} \Xi_{hh'}(\epsilon_l) \langle \phi_{h'}^0 | \theta_\Omega(-\tfrac{1}{2}\nabla^2 + v^1 - \epsilon_l) | \phi_l^0 \rangle, \quad (115)$$

$$c_{h,l} = - \sum_{l'} \langle \tilde{\phi}_h^0 | \theta_\Omega(-\tfrac{1}{2}\nabla^2 + \tilde{v}^1 - \epsilon_l) | \tilde{\phi}_{l'}^0 \rangle [dH - \epsilon_l dO]_{l'l}^{-1}, \quad (116)$$

$$b_{l,l'} = \delta_{l,l'} - \sum_h c_{lh} a_{hl'}, \quad (117)$$

where $\Xi_{l,hh'}(\epsilon)$ is defined by

$$\sum_{h''} \Xi_{l,hh''}(\epsilon) \langle \phi_{h''} | \theta_{\Omega} (-\frac{1}{2} \nabla^2 + v^1 - \epsilon) | \phi_{h'} \rangle = \delta_{h,h'}. \quad (118)$$

The matrices $dH_{l,l'}$ and $dO_{l,l'}$ are defined as in Eqs. (99) and (100) using only the lower, optimized partial waves and the actual one-center potentials instead of the atomic potential. If the partial waves are made self-consistent in each time step, also the projections $\langle \tilde{p}_l | \tilde{\Psi}_n \rangle$ must be transformed according to the change in the projector functions. When adjusting the partial waves and projectors, the energies ϵ_l are chosen relative to a potential reference, such as the average potential in the augmentation sphere, in order to avoid a dependence on an arbitrary overall shift of the potential.

2. Beyond the frozen-core approximation

Even though the present method has been implemented in the frozen-core approximation, it is not limited to it. The core states can of course always be relaxed within an internal self-consistency loop for core wave functions and one-center potentials. The only difficulty is that the orthogonality between AE partial waves and core states must be restored. One could simply imagine using a Gram-Schmitt orthogonalization procedure. However, this would produce partial waves that are no longer close to the solution of the Schrödinger equation, resulting in large partial-wave truncation errors.

Let us therefore mix the partial waves with both the frozen- and the relaxed-core states to impose orthogonality and to minimize the total energy. We can make an ansatz for the new AE partial waves

$$|\phi_i\rangle = |\phi_i^0\rangle + \sum_j |\phi_j^{c0}\rangle a_{ji} + \sum_j |\phi_i^c\rangle b_{ji}, \quad (119)$$

where $|\phi_i\rangle$ are the AE partial waves orthogonalized to the relaxed-core states, $|\phi_i^0\rangle$ are the partial waves orthogonalized to the frozen-core states, $|\phi_i^c\rangle$ are the relaxed-core states, and $|\phi_i^{c0}\rangle$ are the frozen-core states. The coefficients a_{ji} and b_{ji} are determined such that the total energy is minimized

$$\langle \phi_i^{c0} | -\frac{1}{2} \nabla^2 + v^1 - \epsilon | \phi_j \rangle = 0, \quad (120)$$

$$\langle \phi_i^c | -\frac{1}{2} \nabla^2 + v^1 - \epsilon | \phi_j \rangle = 0. \quad (121)$$

The relaxed-core states $|\phi_i^c\rangle$ depend explicitly on the coefficients a_{ij} and b_{ij} . Using the fact that the core states are solutions to some Schrödinger equations, we see that the second equation reduces directly to the orthogonality condition between relaxed core states and the new partial waves. The resulting partial waves for the relaxed core are

$$|\phi_i\rangle = (1 - P^c) \left[|\phi_i^0\rangle - \sum_{j,k} |\phi_j^{c0}\rangle \Xi_{jk}^c(\epsilon_i) \times \langle \phi_k^{c0} | (1 - P^c) (-\frac{1}{2} \nabla^2 + v^1 - \epsilon_i) \times (1 - P^c) | \phi_i^0 \rangle \right], \quad (122)$$

where $\Xi_{jk}^c(\epsilon)$ is defined by

$$\Xi_{jk}^c(\epsilon) \langle \phi_k^{c0} | (1 - P^c) (-\frac{1}{2} \nabla^2 + v^1 - \epsilon) \times (1 - P^c) | \phi_i^{c0} \rangle = \delta_{jl}. \quad (123)$$

$P^c = \sum_i |\phi_i^c\rangle \langle \phi_i^c|$ denotes the projection onto the relaxed-core states. The validity of the result can be verified by inserting it into the defining equations. The result can be further simplified using the fact that the core states are solutions to some Schrödinger equations and the known orthogonality relations. It should, however, be noted that $(1 - P^c) |\phi_i^{c0}\rangle$ vanishes as the frozen-core states and relaxed-core states become similar. Hence these functions should be normalized before inserting them into the above equation in order to avoid a division by zero. If semicore states are present, the same transformation should be performed between the PS partial waves and PS core states in order to guarantee that PS and AE partial waves match at the boundary of the augmentation region. This linear system of equations can be solved and iterated until self-consistency is achieved among core states, AE valence partial waves, and the potential.

VIII. NUMERICAL TESTS

A. Scattering properties

It can easily be shown that the scattering properties of the atom are reproduced correctly in the neighborhood of the energies for which partial waves have been included. The logarithmic derivative of the PS and AE wave functions and their first derivatives agree beyond the core region. The proof is analogous to that for local potentials, which can be found in Skriver's book.⁵³

Thus the scattering properties of the atom can be improved systematically for an arbitrarily large energy region by increasing the number of partial waves. This principle is illustrated in Fig. 2 for the example of manganese. The semicore states have been treated as valence states, which is not necessary in most applications. The logarithmic derivatives obtained with only one partial wave per angular momentum—not to be confused with the setup Mn₆ in Table I—result in a poor description of the valence region. If the number of partial waves is doubled, the scattering properties are accurate up to approximately 1.5 Ry above the occupied states as shown in Fig. 2, which is more than sufficient for most calculations. This choice corresponds to the number of partial waves used in the linear methods. Our experience is that first-row elements can be well described with only one partial wave per angular momentum and that two projectors per angular momentum are sufficient for the narrow *d* states and if semicore states are treated as valence states.

B. Accuracy of the AE wave functions

In order to analyze the accuracy of the AE wave functions obtained with the PAW method, let us calculate the

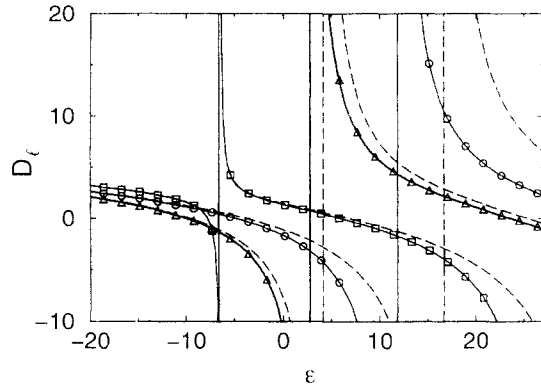


FIG. 2. Scattering properties of the Mn atom. Logarithmic derivative function $D_l(\epsilon) = r \partial_r \phi_l(r, \epsilon) / \phi_l(r, \epsilon)$ with $r = 3a_0$ for s , p , and d angular momenta versus energy. Triangles, circles, and squares indicate the exact result for s , p , and d angular momenta, respectively. Solid lines are obtained with the PAW method using the setup denoted as Mn_b in Table I. Dashed lines have been obtained with the same setup, but without the second partial wave per ℓ .

scattering states of an isolated manganese atom using the PAW method and compare the resulting AE wave functions with the direct integration of the radial Schrödinger equation. Figures 3(a)–3(c) shows the wave functions for an energy of -8.16 eV, which lies in the range of the occupied states. The deviation between the PS and the AE

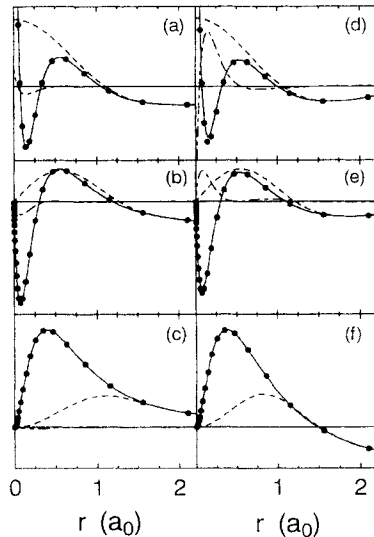


FIG. 3. Comparison of atomic wave functions of Mn using the PAW method with the exact result. Each graph shows the wave function obtained from the PAW method (solid line), the exact AE wave function (bullets), their difference magnified by a factor of 10 (dash-dotted line), and the PS wave function (dashed line) for a given energy and angular momentum. (a)–(c) are obtained at an energy of -8.16 eV, which lies in the valence band region; (d)–(f) are obtained at an energy of $+13.61$ eV, which is far above the valence band region. (a) and (d) are s -type wave functions, (b) and (e) are p -type wave functions, and (c) and (f) are d -type wave functions.

wave function is smaller than 1%.

In order to see the deviations let us compare the wave functions for an energy of $+13.6$ eV; see Figs. 3(d)–3(f). This is about 20 eV above the valence states. Whereas the wave functions of the p and d angular momenta are still quite accurate, we see deviations of 15% in the s -type wave function. This type of deviation, namely, an underestimation of the maxima of the wave function close to the nucleus, is a typical signature of partial-wave truncation errors in the wave function. We can conclude that the PAW method predicts the wave function with high accuracy over a wide energy range.

C. Structural properties

I have performed a number of test calculations on simple molecules to establish the accuracy of the PAW method. I have chosen small molecules because they exhibit the strongest nonspherical potentials and the smallest bond lengths and therefore should provide stringent test systems for any electronic structure method. The results are summarized in Table II together with the results of other recent accurate all-electron LDA

TABLE II. Comparison of binding energies, structural properties, and vibrational properties for dimers obtained with the PAW method at a plane-wave cutoff of 30 Ry with those of other all-electron LDA calculations. Note that the B₂ and the O₂ dimer of Ref. 54 are non-spin-polarized.

Molecule	Quantity	PAW	Other LDA
H ₂	E_B (eV)	4.62	4.65, ^a 4.6 ^b
	d (a ₀)	1.46	1.45, ^c 1.45 ^b
	ω (cm ⁻¹)	4040	4160 ^b
Li ₂	E_B (eV)	1.04	1.00, ^a 1.02 ^b
	d (a ₀)	5.13	5.12, ^c 5.20 ^b
	ω (cm ⁻¹)	335	322 ^b
Be ₂	E_B (eV)	0.53	
	d (a ₀)	4.51	4.52 ^c
	ω (cm ⁻¹)	367	
B ₂	E_B (eV)	3.78	3.5 ^b
	d (a ₀)	3.03	3.05 ^b
	ω (cm ⁻¹)	1060	1030 ^b
N ₂	E_B (eV)	11.38	11.47, ^a 11.3 ^b
	d (a ₀)	2.09	2.07, ^c 2.07 ^b
	ω (cm ⁻¹)	2417	2380 ^b
O ₂	E_B (eV)	7.33	7.48, ^a 6.2 ^b
	d (a ₀)	2.32	2.29 ^b
	ω (cm ⁻¹)	1660	1620 ^b
F ₂	E_B (eV)	3.11	3.33, ^a 3.1 ^b
	d (a ₀)	2.67	2.62, ^c 2.63 ^b
	ω (cm ⁻¹)	1148	1060 ^b
Fe ₂	E_B (eV)	3.99	4.05, ^d 2.89 ^e
	d (a ₀)	3.68	3.74, ^d 3.70 ^e
	ω (cm ⁻¹)	441	418, ^d 412 ^e

^aReference 54.

^bReference 55.

^cReference 56.

^dReference 58.

^eReference 57.

calculations.^{54–58}

I used fcc supercell with a lattice constant of 30 a.u. The PS wave functions have been expanded up to a plane wave cutoff of 30 Ry. The charge density has been expanded to 60 Ry. The augmentation charge density has been expanded up to $\ell = 2$. I used the Perdew-Zunger⁵⁹ parameterization of Ceperley and Alder's Monte Carlo calculations of the free-electron gas.⁶⁰ The parameters used to define the PS partial waves are listed in Table I. For the O₂ dimer calculation we used the setup denoted by O_a in Table I. The setup O_b resulted in a frequency about 5% too low due to the overlap of the augmentation regions in the O₂ molecule. This setup (O_b) has been used in the calculations of MnFO₃ described below.

The atoms have been calculated with integer occupations, which makes them nonspherical. An exception is F, which prefers to spread the two $2p_{\uparrow}$ electrons over three orbitals. In the case of Fe₂ we use a numerical spherical, spin-polarized calculation on radial grids in the experimentally observed valence configuration $d^5d^1s^2$. This is not the ground-state configuration of the LSDA, but it has been used to allow comparisons with previous calculations.

Dioxygen and the boron dimer have been calculated in the triplet configuration, and the iron dimer has been treated in the septet configuration. All other dimers have been calculated in a non-spin-polarized fashion. Binding energies have been reduced by the zero-point vibration energies.

Vibrational frequencies have been obtained by dynamical simulation. Hereby I expanded the bond length by approximately 5% and let the system evolve unperturbed according to equations of motion. No thermostatting has been used for electrons or ions. I used a time step of 10 a.u. for all molecules, except for the hydrogen molecule, for which I used a time step of 1 a.u. The H₂ vibrational frequency obtained with a time step of 10 a.u. is about 10% lower and that obtained with 5 a.u. is 1.2% lower. For N₂, the reduction of the time step to 5 a.u. lowers the frequency by less than 1%.

The results agree very well with other AE calculations. The bond length deviates typically by less than 1% and vibrational frequencies deviate typically by 4%. Binding energies have deviations of 0.1–0.2 eV, which is within the accuracy of previous calculations. Note that the largest discrepancies can be reduced further by increasing the plane-wave cutoff.

In order to study the other worst case for the PAW, i.e., that in which the density deviates strongly from the atomic density, I studied MnFO₃. In this compound, the manganese occurs in a formal oxidation state of seven. The structure is that of a slightly distorted tetrahedron formed by the oxygen and fluorine atoms, with the manganese in its center. In this system I also compared a calculation with frozen-semicore $3s$ and $3p$ states with one that included these electrons explicitly as valence electrons. Treating the semicore states as valence states results in bond distances of $d_{\text{Mn-F}} = 3.205a_0$ (a_0 is the Bohr radius) and $d_{\text{Mn-O}} = 2.973a_0$, which agrees very well with the results when keeping the semicore state frozen, namely, $d_{\text{Mn-F}} = 3.189a_0$ and $d_{\text{Mn-O}} = 2.976a_0$.

The results agree very well with the AE calculations of Chen *et al.*,⁶¹ who predict $d_{\text{Mn-F}} = 3.20a_0$ and $d_{\text{Mn-O}} = 2.96a_0$. As a comparison, the *pseudopotential* calculations of Chen *et al.* overestimate these bond distances by about 2.5% relative to the AE results.

The calculations of MnFO₃ demonstrate that the PAW method faces no problems in dealing with high-oxidation states. Furthermore they show that also calculations with unfrozen-semicore states are feasible with moderate computational effort.

I conclude that the PAW method matches the accuracy of the best existing schemes within the LDA. Even though we have examined simple molecules here, studies of larger systems have also yielded a similar accuracy of structural parameters for which data existed for comparison.^{62–67}

D. First-principles molecular dynamics

To illustrate the quality of dynamical simulations I show in Fig. 4 the various energy contributions to the dynamics. The potential energy and the fictitious kinetic energy undergo a regular oscillation with a period that corresponds to 441 cm⁻¹, which agrees well with the previously calculated vibration frequency of 412–418 cm⁻¹.^{57,58} The solid line represents the conserved energy which should be constant in a high-quality molecular-dynamics simulation. Here the conserved energy deviates less than 0.8 meV from the initial value. No drift has been observed within 10⁻⁵ H. Hence the quality of energy conservation is as good as that obtained with traditional Car-Parrinello simulations using pseudopotentials.

The oscillations of the fictitious kinetic energy should not be misinterpreted as deviations from the Born-Oppenheimer surface. The latter are irreversible and their signature is a monotonous drift of the fictitious kinetic energy to higher values, accompanied by a simultaneous drift of the nuclear kinetic energy towards lower values. The oscillations of the fictitious kinetic energy seen in Fig. 4 represent the motion of the wave function

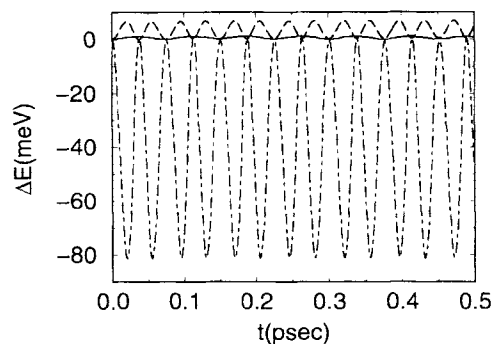


FIG. 4. Energy contributions during a first-principles molecular-dynamics simulation of an iron dimer. The dash-dotted line is the LDA total energy, the dashed line the fictitious kinetic energy of the wave functions, and the solid line the conserved energy. All energies are plotted relative to their initial values. See text for discussion.

on the Born-Oppenheimer surface. Its influence on the motion has been compensated by the renormalization of the nuclear masses by 8% as described in Sec. VC 2.

In conclusion it has been demonstrated that the PAW method allows high-quality energy-conserving molecular dynamics. Molecular-dynamics simulations have been performed on larger systems and the results are published elsewhere.^{62,63}

E. Plane-wave convergence

Figures 5–7 show the plane-wave convergence of total energy, binding energies, and bond distances obtained with the PAW method for all first-row elements (except carbon) and iron as an example of a transition metal.⁶⁸ Convergence to 0.1 eV is achieved at about 30–40 Ry even for difficult cases, such as oxygen and fluorine, and substantially earlier for other elements. Structural properties such as the bond distances are accurate to $0.02a_0$ at a cutoff of 30 Ry, which is less than 1% of the bond length. Binding energies converge faster than absolute total energies, and the error at 30 Ry is as small as 0.1 eV.

For oxygen, we can compare this plane-wave convergence with that of Vanderbilt's ultrasoft pseudopotential.⁴¹ The accuracy and the plane-wave convergence of our calculations are apparently comparable to the "hardest" pseudopotential used. It is not unexpected that a similar plane-wave convergence is obtained for the PAW method and Vanderbilt's ultrasoft pseudopotentials if a similar choice of PS wave functions is used in our calculation and in the construction of the ultrasoft pseudopotentials.

In all calculations shown here we expand the charge density in plane waves up to a plane-wave cutoff that is only twice that used to expand the wave functions. Hence many operations, such as Fourier transforms evaluation of the potential-energy part of $\tilde{E}$, are performed as if the plane-wave cutoff were only one-half of that actually used. This is one advantage over Vanderbilt's ultrasoft pseudopotentials, which either require a plane-wave

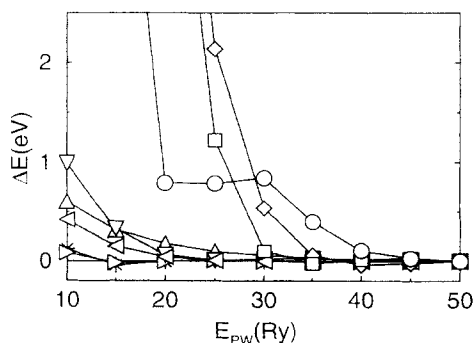


FIG. 5. Plane-wave convergence of the atomic total energy for first-row elements and iron. E_{PW} is the plane-wave cutoff for the wave function. ΔE is the total energy relative to the result obtained with $E_{PW} = 50$ Ry. The following symbols are used: H (Δ), Li ($*$), Be ($\triangleright$), B ($\triangledown$), N ($\triangleleft$), O ($\diamond$), F ($\square$), and Fe ($\circ$). For details see text.

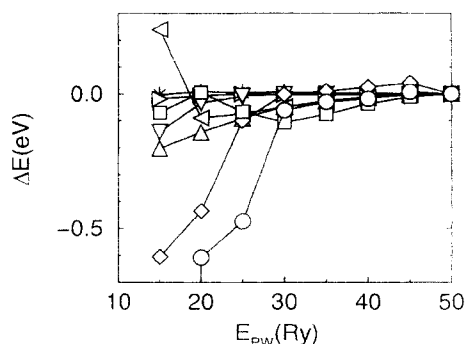


FIG. 6. Plane-wave convergence of the binding energy for dimers. ΔE is the binding energy relative to the result obtained with $E_{PW} = 50$ Ry. The symbols are the same as in Fig. 5. See text for details.

cutoff that is substantially larger than twice the wave-function plane-wave cutoff or where one has to resort to multigrid techniques such as described by Laasonen *et al.*⁴¹

The plane-wave convergence of the PAW method is already close to that of the LAPW method, which predicts mRy convergence of the total energy at a plane-wave cutoff given by $E_{PW} = (5 + \ell_{\max})^2 / r_{MT}^2 a_0^2$ Ry, where $\ell_{\max}$ is the highest angular momentum of the wave functions and r_{MT} is the muffin-tin radius. For a system such as oxygen with a muffin-tin radius of $1.1a_0$ as is necessary for molecular bonds, the plane-wave cutoff should be about 30 Ry. This is a better convergence than that produced by the present implementation of the PAW method since the former corresponds to mRy convergence.

IX. COMPARISON WITH EXISTING METHODS

One can observe that AE and the pseudopotential methods introduced in the past few years seem to converge. Vanderbilt's ultrasoft pseudopotentials⁹ opened the way to the efficient study of first-row and transition metals using a plane-wave-based pseudopotential approach. Conversely Goedecker and Maschke⁶⁹ have

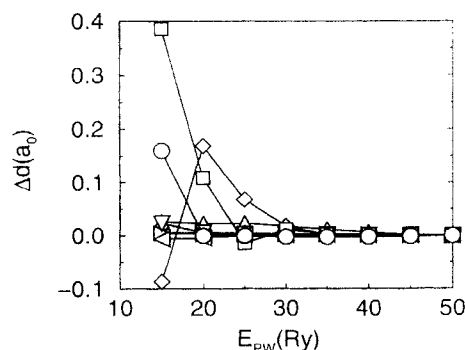


FIG. 7. Plane-wave convergence of dimer bond lengths. Δd is the bond length relative to the result obtained with $E_{PW} = 50$ Ry. The symbols are the same as in Fig. 5. See text for details.

shown that the LAPW method can be simplified to yield pseudopotentials and introduced techniques very similar to the pseudopotential approach into the LAPW method. I believe that the PAW method actually bridges the gap between these two approaches, namely, the pseudopotential approach and the augmented-wave methods. On the one hand, it is an augmented-wave method that yields the full wave functions. It can be regarded, in a sense, as a generalization of the LAPW method. On the other hand, most of the operations done on the plane-wave part are in fact identical to those performed in the plane-wave-based pseudopotential approach. I view this as a promising development, as it allows the virtues of two separate techniques to be combined. In this section I will therefore discuss the relationship between the PAW method and other approaches.

A. Pseudopotential method

1. Norm-conserving pseudopotentials

The pseudopotential approach, based on generalized separable potentials,^{21,70} can actually be derived from the PAW method by making one well-defined approximation. This way of approaching the pseudopotential theory sheds light on the underlying principles of the pseudopotential approach. Furthermore, the comparison will provide a one-to-one correspondence between quantities in the PAW method and those in the pseudopotential method. In order to alert the reader to this fact, I have called the plane-wave part of the wave function the “PS wave function.”

Before we start comparing PAW and the pseudopotential approach, let us recall what is termed the “pseudopotential approach.” A more extensive description can be found in Refs. 8 and 50.

A valid first-principles pseudopotential obeys the following conditions.

- (i) For the atom, the PS wave functions agree with the AE wave functions beyond the core region.
- (ii) Atomic eigenvalues and the first energy derivative of the logarithmic derivative of the energy-dependent partial waves agree with those of the AE calculation.
- (iii) The norm of the atomic PS wave function is identical to that of the true wave function.

Owing to the construction described in Sec. VI, the PAW method fulfills the first two conditions automatically. Even though the third condition is usually not fulfilled, it can be imposed.

In order to reduce the PAW method to the traditional pseudopotential approach, we must make a Taylor expansion of the total energy in terms of the deviation of the one-center densities n^1 and $\tilde{n}^1$ from their isolated atom values. This yields, to first order,

$$E \approx E_{\text{self}} + \tilde{E} + \sum_n \langle \tilde{\Psi}_n | v_{nl} | \tilde{\Psi}_n \rangle. \quad (124)$$

The term linear in the density operator is a nonlocal pseudopotential

$$v_{nl}(r, r') = \sum_{i,j} \langle r | \tilde{p}_i \rangle \langle \phi_i | -\frac{1}{2} \nabla^2 + v_{\text{at}}^1 | \phi_j \rangle - \langle \tilde{\phi}_i | -\frac{1}{2} \nabla^2 + \tilde{v}_{\text{at}}^1 | \tilde{\phi}_j \rangle \langle \tilde{p}_j | r' \rangle \quad (125)$$

and the constant is the so-called self-energy

$$E_{\text{self}} = E_{\text{at}}^1 - \tilde{E}_{\text{at}}^1 - \sum_n \langle \tilde{\Psi}_{\text{at},n} | v_{nl} | \tilde{\Psi}_{\text{at},n} \rangle. \quad (126)$$

The potentials v_{at}^1 , $\tilde{v}_{\text{at}}^1$, and $\tilde{v}$ are obtained from the charge densities n and $\tilde{n}$ of the atoms. Also compensation charge densities in $\tilde{E}$ are kept frozen and are imported from the isolated atoms. The first-order change of the total energy with respect to a change in the compensation densities has been absorbed by the nonlocal PS potential. The approximations are well justified as long as the PS charge density is sufficiently close to the AE charge density. If the condition of norm conservation is imposed, the overlap operator is identical to unity.

We thus obtain the form for the total-energy functional that is well known from the pseudopotential method, which is none other than that actually implemented in the original Car-Parrinello codes²² using generalized separable pseudopotentials.²¹ Achieving a close similarity of the computationally demanding plane-wave expressions to existing Car-Parrinello codes was one of my goals while deriving the PAW method. This allows us to make direct use of the technology of existing methods, while incorporating the full-wave functions.

The major difference in the computational effort required for the PAW method and the traditional pseudopotential approach is the explicit inclusion of the one-center terms in the PAW method. The related effort, however, scales linearly with the number of atoms and is in practice negligibly small. Furthermore, the PAW method provides new flexibility to increase computational efficiency and accuracy. We can use larger core radii, and by relaxing the norm-conservation condition we can use smoother PS partial waves and reduce the basis set just as in Vanderbilt's ultrasoft pseudopotentials. The trade-off of this, however, is that the PS wave functions have to be orthogonalized in the presence of a nontrivial overlap operator. Hence the decision in favor of or against the norm-conservation condition depends on the system under study.

2. Vanderbilt's ultrasoft pseudopotentials

Vanderbilt recently introduced ultrasoft pseudopotentials that relax the norm-conservation condition.^{9,10} This approach has two basic ingredients: First, the wave functions are not norm conserving and second, generalized separable pseudopotentials^{9,21}—an extension of the Kleinman-Bylander potentials⁷⁰—are used.

These two ingredients of the ultrasoft pseudopotentials and the PAW method are similar. In addition, the projector augmentation uses concepts from the generalized separable pseudopotentials and, like other augmented wave schemes, its PS wave function or envelope function has a different norm than that of the AE wave function. The

difference, however, is that the PAW method is an all-electron method and Vanderbilt's approach a pseudopotential method.

(i) The PAW method avoids "pseudization" steps to which one resorts when using pseudopotential approaches, be they norm-conserving pseudopotential approaches or not. The PAW method, on the other hand, works directly with the full-wave functions and potentials and includes the core states. This is a nontrivial problem; the notion that the full-wave functions cannot be treated in a reasonable way on a regular grid was the reason to introduce augmented-wave and pseudopotential methods. The PAW method follows here the tradition of the augmented-wave methods, where the full-wave function is decomposed into various parts, each of which can be handled conveniently in its own representation.

(ii) The PAW method provides a prescription to go back and forth between the PS wave functions and the physical AE wave functions. The analogy between the PAW method and the pseudopotential approach has been exploited successfully by Van de Walle and myself¹² to reconstruct an approximate full-wave function from a pseudopotential calculation and obtain quantities that are not directly accessible by the pseudopotential approach. The PAW method is more rigorous than a mere reconstruction of the wave functions because the full wave function takes part in the screening process.

(iii) In the ultrasoft pseudopotentials the overlap operator and the local charges have been introduced to restore the scattering properties of the pseudopotential when the norm-conservation condition is relaxed in order to obtain smoother PS wave functions. In the PAW method, the non-norm-conserving PS wave functions enter naturally as in all other augmented-wave methods.

(iv) From the point of view of computational effort, the PAW method and the ultrasoft pseudopotentials are equivalent with respect to plane-wave convergence if a similar construction of non-norm-conserving PS wave functions is used. However, the PAW method is more efficient because it treats the one-center expansions on radial grids, which effectively eliminates the related computational cost, rather than in a plane-wave representation. Furthermore, the plane-wave cutoff for the charge density can be chosen substantially lower in the PAW method because the augmentation density is not directly added to the density grid. This cuts the computational effort for the Fourier transforms substantially.

B. Linear augmented-plane-wave method

The present method is in the tradition of existing augmented-wave methods. Augmented-wave functions were originally invented by Slater.⁴ There, the Lippmann-Schwinger equation for a muffin-tin type potential is solved by matching the energy-dependent partial solution of the Schrödinger equation inside and outside the muffin-tin sphere. As this matching procedure is computationally extremely demanding, Andersen introduced the linear methods,³ where the partial solutions from within the muffin-tin sphere are linearized with

respect to the one-particle energy, resulting in energy-independent basis functions of a type previously suggested by Marcus⁵ in the context of the APW method. This allowed the wave functions to be obtained by matrix diagonalization. It also led to the definition of basis functions such as linear augmented-plane waves, linear augmented muffin-tin orbitals, and many more.

All linear methods have in common that two partial waves of the spherical part of the potential within the muffin-tin sphere are matched at the sphere radius to a so-called envelope function, which corresponds to the PS wave function of the PAW method. The PAW method modifies precisely that principle. Instead of determining the coefficients of the partial-wave expansion from the value and derivative of the PS wave function at some sphere radius, it uses the more general principle of projector augmentation. We have seen in Sec. II that the scalar product with some projector function is the *most general* way to determine these coefficients linearly from the PS wave function. There are other differences to the linear methods made possible by this more general approach which provide the important practical advantages. Those will be discussed later.

Here we show that the LAPW method is, in some sense, a special case of the PAW method, namely, that it is possible to formulate the augmentation by matching as in the linear methods by projector functions. That the matching of the LAPW method can be expressed by projector functions has been observed earlier by Goedecker and Maschke.⁶⁹

In the LAPW method the wave functions are expressed by partial waves $|\phi_\nu\rangle$ and their energy derivatives $|\dot{\phi}_\nu\rangle$ at some energy ϵ_ν within atom-centered muffin-tin spheres Ω_R and by plane waves in the interstitial region Ω_i :

$$|\Psi\rangle = (1 - \theta_{\Omega_R})|\tilde{\Psi}\rangle + \theta_{\Omega_R}(|\phi_\nu\rangle a - |\dot{\phi}_\nu\rangle b), \quad (127)$$

where θ_{Ω_R} is a step function that is unity within the augmentation sphere Ω_R and zero outside. The coefficients a and b ,

$$\begin{bmatrix} a \\ b \end{bmatrix} = \frac{1}{(\phi_\nu \partial_r \dot{\phi}_\nu - \dot{\phi}_\nu \partial_r \phi_\nu)} \begin{bmatrix} \tilde{\Psi} \partial_r \dot{\phi}_\nu - \dot{\phi}_\nu \partial_r \tilde{\Psi} \\ \phi_\nu \partial_r \tilde{\Psi} - \tilde{\Psi} \partial_r \phi_\nu \end{bmatrix}, \quad (128)$$

are determined such that the wave function is differentiable at the sphere radius. For reasons of simplicity we can write these equations for only one angular momentum component of the wave function.

The same result can also be reproduced using projector functions: First we construct two AE partial waves per angular momentum and site as the partial wave $|\phi(\epsilon)\rangle$ for a given energy ϵ_ν , yielding $|\phi_\nu\rangle$, and its energy derivative $|\dot{\phi}\rangle$, where the overdot stands for energy derivative. Then we construct PS partial waves analogously to the procedure described in Sec. VI. When constructing the projector functions, the augmented-wave methods deviate from the recipe given in Sec. VI and build them up as superpositions $|\tilde{p}_i\rangle = \sum_j (\nabla^2 \theta_{\Omega_R}) |\phi_j\rangle c_{ji}$. The coefficients c_{ij} are determined such that the orthogonality condition $\langle \tilde{p}_i | \tilde{p}_j \rangle = \delta_{ij}$ is fulfilled.⁷¹ The projector functions corresponding to the linear methods are localized on the

two-dimensional muffin-tin surface, which excludes real space summation to obtain the scalar products with the PS wave function. However, the latter can be evaluated in G space for any other representation of the PS wave function in terms of analytical functions. There are extensions of the LAPW method that use more than two partial waves,⁷² but they will not be discussed here, as I only wish to illustrate the principle. The computational effort for this part of the augmentation is similar in the LAPW and the PAW methods if such a projection-type approach is used in the LAPW method.

In contrast to the PAW method, the LAPW method requires the one-center expansion of the PS wave function to be exactly identical to the PS wave function itself, which is computationally demanding. In the PAW method there is no need to make the one-center expansion of the PS wave function more accurate than the one-center of the AE wave function and both are therefore obtained in a completely analogous way. In fact, this procedure results in an extremely rapid convergence of the partial-wave expansion, which is due to the error cancellations discussed earlier. The one-center expansions of the PS wave function are obtained *without extra effort* because their coefficients are identical to those obtained previously with the AE wave function.

Another difference between the LAPW and the PAW methods is the use of frozen partial waves imported from an isolated atom as opposed to partial waves that adjust to the actual potential. At first glance this appears to be a disadvantage of the PAW method. However, if we count the number of variational degrees of freedom, we find that, given the same number of plane waves, the PAW method has a more flexible basis set, even though the partial-wave expansions often have fewer terms than the LAPW method. Whereas the LAPW method has complete flexibility to adjust to the spherical part of the potential inside the muffin-tin spheres, the PAW method lets the plane waves of the unaugmented part of the PS wave function extend into the spheres and thus allows the wave functions to adjust to both spherical and nonspherical potentials. This is the result of what is often called “additive augmentation” and has been discussed in the Sec. VII. In combination with a fictitious Lagrangian formalism, the use of fixed partial waves has the important advantage that the total energy is a unique function of only the PS wave functions and the atomic positions.

1. The APW method of Soler and Williams

The APW method of Soler and Williams^{15–17} differs from other implementations of the LAPW method in several ways. I will compare the main differences between their APW method and the PAW method. Most of what has already been said about the difference between the PAW method and the LAPW method also applies here.

First, Soler’s APW method employs an iterative minimization of the total energy, which is similar in spirit to the Car-Parrinello method. To my knowledge no molecular-dynamics features have been implemented. The reason has been attributed⁷³ to the problem of “hid-

den” variables in the APW method, e.g., the shape of the partial waves, that adjust to the actual spherical potential. As the PAW method uses frozen partial waves, this problem does not arise.

Second, in the Soler-Williams APW the linear approximation has been abandoned in favor of so-called infinitesimal paneling. This means that every state has its own ϕ, ϕ functions obtained from the spherical potential at the energy of that state. As a result the scattering properties are correct for all one-particle energies given a spherical potential. The PAW method reproduces the scattering properties correctly over the entire energy range only in the ideal case of an infinite number of partial waves. In this case the PAW wave functions are exact solutions for the full *nonspherical* potential. In practice the wave functions of the PAW are accurate in an arbitrarily large, but finite energy region.

Like previous implementations of the LMTO method, the Soler-Williams APW employs the principle of additive augmentation. In the PAW method the principles of additive augmentation are applied even more rigorously. Not only the ℓ convergence, but also the partial-wave convergence for each angular momentum channel is accelerated using this idea.

2. Singh’s projector-basis technique

Recently, Singh⁷⁴ introduced an implementation of the LAPW method (and the mixed-basis pseudopotential) method that is substantially more efficient than previous implementations. Instead of matching the partial waves directly to the plane-wave part of the wave function, a set of analytical functions is first fitted to the PS wave function on the real-space grid points in a localized region around each atom. Once an analytical expression for the plane-wave part is given, the matching of partial waves and the integrations over the muffin-tin sphere can be evaluated quickly. The advantage of this approach is that computational effort for the augmentation scales quadratically [or $N^2 \ln(N)$] with the number of atoms N , compared to an N^3 scaling in a pure G space formulation. Another solution of the same problem has been developed in the pseudopotential approach.⁷⁵

There has been some confusion about the terms “projector basis function” in Singh’s paper and the “projector functions” of the PAW method. The two refer to different objects. Since the projector functions are used to represent the PS wave function, the projector basis functions of Singh’s approach relate more closely to the PS partial waves of the PAW method rather than to its projector functions. However, also here are differences: The PS partial waves of the PAW method are per construction adapted to the PS potential, and the coefficients of the partial wave expansion are obtained by a scalar product of my projector functions with the PS wave function, whereas in Singh’s approach these coefficients are obtained by a least-squares-type fit to the plane-wave part at the real-space grid points. The difference between the two approaches is most apparent from the number of projector functions and PS waves, which in PAW vary

typically between four and fourteen per site, whereas in Singh's approach several hundred functions are used.

Singh's approach can be applied equally to the pseudopotential formalism and to all augmented-wave schemes including the PAW method. When convergence with the number of Singh's projector functions and plane-wave cutoff is obtained, it reproduces exactly the electronic structure method to which it has been applied.

X. CONCLUSIONS

An electronic structure method has been described that works directly on the full valence and core wave functions and allows highly accurate first-principles molecular-dynamics simulation to be performed. It has been demonstrated that the accuracy of the PAW method matches that of other state-of-the-art electronic structure methods based on the LDA and that high-quality first-principles molecular-dynamics simulations are possible using this approach. To the best of my knowledge, this method was used in the first molecular-dynamics simulation using an all-electron method.⁶² The method bridges the gap between the existing augmented-wave methods and the pseudopotential methods and underscores the relationship between these two approaches. Compared to the existing approaches, the present method is expected to be more efficient, given a similar level of optimization. The method can be incorporated into existing pseudopotential codes with relatively minor additional effort.

ACKNOWLEDGMENTS

I would like to express my gratitude to O. K. Anderson, who taught me the basics that allowed me to develop PAW. Furthermore, I am indebted to S. T. Pantelides

for his continuous support and encouragement. I would also like to thank K. Schwarz, D. Vanderbilt, P. Margl, J. Hutter, and M. Parrinello for useful comments after reading the manuscript.

APPENDIX A: RADIAL SCHRÖDINGER EQUATION WITH A SEPARABLE POTENTIAL

In order to test the scattering properties it is necessary to solve the radial Schrödinger equation using a separable pseudopotential. An approach is sketched here briefly:

$$\left(-\frac{1}{2}\nabla^2 + \tilde{v} - \epsilon\right)|\tilde{\Psi}\rangle + \sum_{i,j} |\tilde{p}_i\rangle \langle dH_{ij} + \epsilon dO_{ij} | \langle \tilde{p}_j | \tilde{\Psi} \rangle = 0. \quad (\text{A1})$$

An ansatz for the solution is

$$|\tilde{\Psi}\rangle = |u\rangle + \sum_i |w_i\rangle c_i \quad (\text{A2})$$

with $|u\rangle$ and $|w_i\rangle$ defined as

$$\left(-\frac{1}{2}\nabla^2 + \tilde{v} - \epsilon\right)|u\rangle = 0 \quad (\text{A3})$$

and

$$\left(-\frac{1}{2}\nabla^2 + \tilde{v} - \epsilon\right)|w_i\rangle = |\tilde{p}_i\rangle. \quad (\text{A4})$$

After inserting this ansatz into Eq. (A1) we obtain an equation that determines the coefficients c_i as

$$c_i = - \sum_{j,l} \left[\delta_{ij} + \sum_k (dH_{ik} - \epsilon dO_{ik}) \langle \tilde{p}_k | w_j \rangle \right]^{-1} \times (dH_{jl} - \epsilon dO_{jl}) \langle \tilde{p}_l | u \rangle. \quad (\text{A5})$$

- ¹ P. Hohenberg and W. Kohn, Phys. Rev. **136**, B664 (1964).
- ² W. Kohn and L.J. Sham, Phys. Rev. **140**, B1133 (1965).
- ³ O.K. Andersen, Phys. Rev. B **12**, 3060 (1975).
- ⁴ J.C. Slater, Phys. Rev. **51**, 846 (1937).
- ⁵ P.M. Marcus, Int. J. Quantum. Chem. **1S**, 567 (1967).
- ⁶ J. Korringa, Physica **13**, 392 (1947).
- ⁷ W. Kohn and N. Rostocker, Phys. Rev. **94**, 111 (1954).
- ⁸ D.R. Hamann, M. Schlüter, and C. Chiang, Phys. Rev. Lett. **43**, 1494 (1979).
- ⁹ D. Vanderbilt, Phys. Rev. B **41**, 7892 (1985).
- ¹⁰ K. Laasonen, R. Car, C. Lee, and D. Vanderbilt, Phys. Rev. B **43**, 6796 (1991).
- ¹¹ R. Car and M. Parrinello, Phys. Rev. Lett. **55**, 2471 (1985).
- ¹² C.G. Van de Walle and P.E. Blöchl, Phys. Rev. B **47**, 4244 (1993).
- ¹³ P. Blaha, K. Schwarz, and P.H. Dederichs, Phys. Rev. B **37**, 2792 (1988).
- ¹⁴ K. Schwarz and P. Blaha, Z. Naturforsch. A **47**, 197 (1992).
- ¹⁵ J.M. Soler and A.R. Williams, Phys. Rev. B **40**, 1560 (1989).

- ¹⁶ J.M. Soler and A.R. Williams, Phys. Rev. B **42**, 9728 (1990).
- ¹⁷ J.M. Soler and A.R. Williams, Phys. Rev. B **47**, 6784 (1993).
- ¹⁸ D. Singh, Phys. Rev. B **40**, 5428 (1989).
- ¹⁹ T. Oguchi, in *Interatomic Potential and Structural Stability*, edited by K. Terakura and H. Akai (Springer, Berlin, 1993), p. 33.
- ²⁰ A set of functions that is complete in a region can be used to expand any function within that region.
- ²¹ P.E. Blöchl, Phys. Rev. B **41**, 5414 (1990).
- ²² R. Car (private communication).
- ²³ S.F. Boys, Proc. R. Soc. London Ser. A **200**, 542 (1950).
- ²⁴ S. Obara and A. Saika, J. Chem. Phys. **84**, 3963 (1986).
- ²⁵ V.R. Saunders, in *Methods in Computational Molecular Physics*, edited by G.H.F. Diercksen and S. Wilson (Reidel, Dordrecht, 1983), pp. 1-36.
- ²⁶ K.H. Weyrich, Phys. Rev. B **37**, 10 269 (1988).
- ²⁷ P.E. Blöchl, Ph.D. thesis, University of Stuttgart, Germany, 1989.

- ²⁸ P. Ehrenfest, Z. Phys. **45**, 455 (1927).
- ²⁹ H. Hellmann, *Einführung in die Quantenchemie* (Deuticke, Leipzig, 1937).
- ³⁰ R.P. Feynman, Phys. Rev. **56**, 340 (1939).
- ³¹ D.G. Pettifor, Commun. Phys. **1**, 141 (1979).
- ³² A.R. Mackintosh and O.K. Andersen, in *Electrons at the Fermi Surface*, edited by M. Springford (Cambridge University Press, Cambridge, England, 1980), p. 187.
- ³³ V. Heine, Solid State Phys. **35**, 1 (1980).
- ³⁴ O.H. Nielsen and R.M. Martin, Phys. Rev. B **32**, 1780 (1985).
- ³⁵ P. Pulay, Mol. Phys. **17**, 197 (1969).
- ³⁶ J. Harris, R.O. Jones, and J. E. Müller, J. Chem. Phys. **75**, 3904 (1981).
- ³⁷ L. Verlet, Phys. Rev. **159**, 98 (1967).
- ³⁸ G. Pastore, E. Smargiassi, and F. Buda, Phys. Rev. A **44**, 6334 (1991).
- ³⁹ J.-P. Ryckaert, G. Ciccotti, and H.J.C. Berendsen, J. Comput. Phys. **23**, 327 (1977).
- ⁴⁰ R. Car and M. Parrinello, in *Simple Molecular Systems at Very High Density*, edited by A. Polian *et al.* (Plenum, New York, 1989), p. 455.
- ⁴¹ K. Laasonen, A. Pasquarello, R. Car, C. Lee, and D. Vanderbilt, Phys. Rev. B **47**, 10 142 (1993).
- ⁴² S. Nosé, J. Chem. Phys. **81**, 511 (1984).
- ⁴³ P.E. Blöchl and M. Parrinello, Phys. Rev. B **45**, 9413 (1992).
- ⁴⁴ M. Parrinello (private communication).
- ⁴⁵ G. P. Kerker, J. Phys. C **13**, 1189 (1980).
- ⁴⁶ A.M. Rappe, K.M. Rabe, E. Kaxiras, and J.D. Joannopoulos, Phys. Rev. B **41**, 1227 (1990).
- ⁴⁷ J.S. Lin, A. Qteish, M.C. Payne, and V. Heine, Phys. Rev. B **47**, 4174 (1993).
- ⁴⁸ N. Troullier and J.L. Martins, Phys. Rev. B **43**, 1993 (1991).
- ⁴⁹ D.D. Koelling and B.N. Harmon, J. Phys. C **10**, 3107 (1977).
- ⁵⁰ G.B. Bachelet, D.R. Hamann, and M. Schlüter, Phys. Rev. B **26**, 4199 (1982).
- ⁵¹ D.R. Hamann, Phys. Rev. B **40**, 2980 (1989).
- ⁵² A right-hand eigenvalue is a value λ that has a nontrivial solution c_i to $\sum_j a_{ij}c_j = c_i\lambda$.
- ⁵³ H.L. Skriver, in *The LMTO Method*, edited by M. Cardona, P. Fulde, and H. J. Queisser, Springer Series in Solid State Sciences Vol. 41 (Springer, Berlin, 1984).
- ⁵⁴ A.D. Becke, J. Chem. Phys. **97**, 9173 (1992).
- ⁵⁵ P.A. Serena, A. Baratoft, and J.M. Soler, Phys. Rev. B **48**, 2046 (1993).
- ⁵⁶ R.M. Dickson and A. Becke, J. Chem. Phys. **99**, 3898 (1993).
- ⁵⁷ S. Dhar and N.R. Kestner, Phys. Rev. A **38**, 1111 (1988).
- ⁵⁸ J.L. Chen, C.S. Wang, K.A. Jackson, and M.R. Pederson, Phys. Rev. B **44**, 6558 (1991).
- ⁵⁹ J.P. Perdew and A. Zunger, Phys. Rev. B **23**, 5048 (1981).
- ⁶⁰ D.M. Ceperley and B.J. Alder, Phys. Rev. Lett. **45**, 566 (1980).
- ⁶¹ H. Chen, M. Krasowski, and G. Fitzgerald, J. Chem. Phys. **98**, 8710 (1993).
- ⁶² P. Margl, P.E. Blöchl, and K. Schwarz, in *Computations for the Nano Scale*, edited by P.E. Blöchl *et al.* (Kluwer, Dordrecht, 1993), p. 153.
- ⁶³ P. Margl, P.E. Blöchl, and K. Schwarz, J. Chem. Phys. **100**, 8194 (1994).
- ⁶⁴ A.J. Fisher and P.E. Blöchl, Phys. Rev. Lett. **70**, 3263 (1993).
- ⁶⁵ A.J. Fisher and P.E. Blöchl, in *Computations for the Nano Scale* (Ref. 62), p. 185.
- ⁶⁶ R. Nesper, K. Vogel, and P.E. Blöchl, Angew. Chem. Int. Ed. Engl. **32**, 701 (1993).
- ⁶⁷ P. Carloni, P.E. Blöchl, and M. Parrinello (unpublished).
- ⁶⁸ The convergence of the binding energy of Fe₂ has been studied by comparing a plane-wave PAW calculation of Fe²⁺.
- ⁶⁹ S. Goedecker and K. Maschke, Phys. Rev. B **42**, 8858 (1990).
- ⁷⁰ L. Kleinman and D.M. Bylander, Phys. Rev. Lett. **48**, 1425 (1982).
- ⁷¹ Matrix elements of $\nabla^2\theta_n$ can be resolved by performing partial integration two times and by the application of Gauss's theorem as $\langle f_L | (\nabla^2\theta_n) | g_L \rangle = \delta_{L,L'} r_n^2 (f_L \partial_r g_L + g_L \partial_r f_L)$, where r_n is the radius of the muffin-tin sphere. f_L and $\partial_r f_L$ are the value and the derivative of the radial part of $|f_L\rangle$ at the muffin-tin radius.
- ⁷² D.J. Singh, Phys. Rev. B **43**, 6388 (1991).
- ⁷³ J.M. Soler (private communication).
- ⁷⁴ D. Singh, Phys. Rev. B **46**, 13 065 (1992).
- ⁷⁵ R.D. King-Smith, M.C. Payne, and J.S. Lin, Phys. Rev. B **44**, 13 063 (1991).